
– Mathématiques –
MP - 2026/2027

Table des matières

1	Séries numériques	7
1	Généralités	8
2	Séries de référence	13
3	Séries à termes positifs	17
4	Séries absolument convergentes	24
5	Comparaison séries - intégrales	28
6	Séries alternées	31
7	Sommation des relations de comparaison pour les séries à termes positifs	33
2	Algèbre linéaire et éléments propres	40
1	Rappels	40
2	Rappels sur les matrices	46
3	Formes linéaires et hyperplans	48
4	Compléments en algèbre linéaire	52
5	Éléments propres d'un endomorphisme	61
6	Éléments propres d'une matrice	66
3	Intégration	72
1	Intégrale généralisée sur un intervalle de la forme $[a, +\infty[$	74
2	Intégration sur un intervalle quelconque	85
3	Les théorèmes de Lebesgue	100
4	Polynôme caractéristique et réduction	109
1	Polynôme caractéristique	109
2	Matrices et endomorphismes trigonalisables	116
3	Matrices et endomorphismes nilpotents	122

5	Suites et séries de fonctions	125
1	Suites de fonctions	125
2	Continuité et double limite	134
3	Intégration et dérivation	138
4	Séries de fonctions	143
6	Dénombrabilité et familles sommables	154
1	Ensembles dénombrables	154
2	Familles sommables	159
7	Probabilités	175
1	Espaces probabilisés	175
2	Propriétés élémentaires des probabilités	179
3	Indépendance et probabilités conditionnelles	183
4	Variables aléatoires discrètes	188
5	Lois usuelles	198
8	Polynôme minimal et réduction	203
1	Algèbres	204
2	Polynômes d'endomorphismes et de matrices	206
3	Compléments d'algèbre	210
4	Polynômes annulateurs	219
5	Critère de diagonalisabilité	226
6	Compléments et applications	234
9	Espaces vectoriels normés I	239
1	Normes	240
2	Exemple de normes équivalentes	253
3	Suites à valeurs dans un espace vectoriel normé	256
4	Séries à valeurs dans un espace vectoriel de dimension finie	264
5	Applications linéaires lipschitziennes	266
10	Séries entières	272
1	Généralités et rayon de convergence	272
2	Étude de la somme d'une série entière	281
3	Fonctions développables en séries entières	286
4	Utilisation des séries entières dans l'étude des équations différentielles	288
11	Espaces préhilbertiens réels I	297
1	Rappels	297
2	Adjoint d'un endomorphisme	304

3	Matrices orthogonales	307
4	Isométries vectorielles d'un espace euclidien	312
12	Probabilités II	325
1	Espérance d'une variable aléatoire discrète	325
2	Variance et écart type d'une variable aléatoire réelle	334
3	Inégalités probabilistes et loi des grands nombres	341
4	Fonctions génératrices	345
13	Espaces vectoriels normés II	349
1	Topologie d'un espace vectoriel normé	350
2	Limite d'une application	364
3	Continuité	369
4	Parties compactes d'un espace vectoriel normé	379
5	Espaces vectoriels de dimension finie	383
6	Séries à valeurs dans un espace vectoriel de dimension finie	387
7	Parties connexes par arcs	392
14	Espaces préhilbertiens réels 2	396
1	Endomorphismes autoadjoints d'un espace euclidien	396
2	Endomorphismes autoadjoints positifs, définis positifs	401
15	Fonctions à valeurs vectorielles	405
1	Dérivabilité	405
2	Intégration sur un segment	412
3	Intégrale fonction de sa borne supérieure	416
4	Formules de Taylor	417
5	Suites et séries de fonctions à valeurs vectorielles	419
16	Équations différentielles linéaires	426
1	Généralités	426
2	Théorème de Cauchy linéaire	432
3	Équations différentielles à coefficients constants	437
4	Equations différentielles scalaires du second ordre	443
17	Intégrales à paramètres	450
1	Théorèmes généraux	450
2	Exemples	459
18	Calcul différentiel	460
1	Différentielle et dérivées partielles	461

2	Opérations sur les applications différentiables et les applications de classe $\mathcal{C}^1$	474
3	Dérivée le long d'un arc	482
4	Fonctions de classe $\mathcal{C}^k$	488
5	Exemple d'équations aux dérivées partielles	493
6	Optimisation	495
19	Algèbre générale	504
1	Groupes	504
2	Groupes monogènes	512
3	Ordre d'un élément	517
4	Anneaux	522
5	L'anneau $(\mathbb{Z}/n\mathbb{Z}, +, \times)$	527
6	Factorisation des polynômes de $\mathbb{K}[X]$	539

Présentation

Voici les notes de cours de MP (programme de 2022). Ce sont des notes **préliminaires**. Il y a probablement de nombreux oublis et coquilles. Je remercie Felix Faisant (MP 2016-2017) pour m'avoir signalé quelques erreurs ainsi que @YetAnother_MT pour avoir partagé le code TikZ de ses dessins sur les fonctions convexes.

Séries numériques

1	Généralités	8
1.1	Définitions	8
1.2	Propriétés	10
1.3	Lien suite-série	10
1.4	Divergence grossière	12
2	Séries de référence	13
2.1	Séries géométriques	13
2.2	Séries de Riemann	13
2.3	Série exponentielle	14
2.4	Exemples de calculs de sommes	15
3	Séries à termes positifs	17
3.1	Généralités	17
3.2	Théorème de comparaison par inégalités pour les séries à termes positifs	18
3.3	Rappels sur les relations de comparaison	20
3.4	Théorème de comparaison par équivalents pour les séries à termes positifs	22
3.5	Règle de d'Alembert	23
4	Séries absolument convergentes	24
4.1	Définition	25
4.2	Applications aux séries dont le terme général n'a pas nécessairement un signe constant	25
5	Comparaison séries - intégrales	28
5.1	Généralités	28
5.2	Exemples et applications	29
6	Séries alternées	31
6.1	Généralités	31
6.2	Etude du reste et de la somme	32
7	Sommation des relations de comparaison pour les séries à termes positifs	33
7.1	Sommation des relations de comparaison pour les séries à termes positifs	33
7.2	Applications classiques	36

Dans ce chapitre, $\mathbf{K}$ désigne $\mathbf{R}$ ou $\mathbf{C}$.

1 Généralités

1.1 Définitions

Définition 1.1

Soit $(u_n)_{n \in \mathbf{N}}$ une suite à valeurs dans $\mathbf{K}$. On note $\sum u_n$ la série de terme général u_n . On appelle alors suite des sommes partielles la suite $(S_n)_{n \in \mathbf{N}}$ définie par

$$\forall n \in \mathbf{N}, S_n = \sum_{k=0}^n u_k.$$

Remarque : On peut considérer que le terme général n'est défini par qu'à partir d'un rang n_0 (la plupart de temps 1). On note alors $\sum_{n \geq n_0} u_n$ la série dont les sommes partielles sont, pour $n \geq n_0$,

$$\sum_{k=n_0}^n u_k.$$

Dans la suite, on traitera principalement des séries qui commencent à 0 ou à 1.

Exemples :

1. Pour la série $\sum 1$. On a donc pour $n \geq 0$, $S_n = \sum_{k=0}^n 1 = n + 1$.
2. Pour la série $\sum \frac{1}{2^n}$. On a donc pour $n \geq 0$, $S_n = \sum_{k=0}^n \frac{1}{2^k} = \frac{1 - (1/2)^{n+1}}{1 - (1/2)} = 2 - \left(\frac{1}{2}\right)^n$.

Définition 1.2

Soit $\sum u_n$ une série.

1. Si la suite des sommes partielles converge, on dit que la série est **convergente**.

On note alors $S = \sum_{k=0}^{+\infty} u_k$ la limite des sommes partielles. On l'appelle la somme de la série.

2. Si la suite des sommes partielles diverge, on dit que la série est **divergente**.

Remarque : Le fait d'être convergente ou divergente est la nature de la série.

Exemples :

1. La série $\sum 1$ diverge car pour tout entier n , $S_n = \sum_{k=0}^n 1 = n + 1$ et $(S_n) \rightarrow +\infty$.
2. La série $\sum \frac{1}{k}$ diverge aussi. Supposons par l'absurde qu'elle converge. On note alors $S = \sum_{k=0}^{+\infty} \frac{1}{k}$. La suite (S_n) tend vers S de ce fait la suite extraite (S_{2n}) tend aussi vers S . On en déduit que la suite $(S_{2n} - S_n)_{n \geq 0}$ est une suite convergente de limite nulle. Cependant, pour tout entier

$n \geq 1$,

$$S_{2n} - S_n = \sum_{k=1}^{2n} \frac{1}{k} - \sum_{k=1}^n \frac{1}{k} = \sum_{k=n+1}^{2n} \frac{1}{k} \geq \sum_{k=n+1}^{2n} \frac{1}{2n} \geq \frac{1}{2}.$$

Finalement la série $\sum \frac{1}{k}$ diverge.

Cette série s'appelle la *série harmonique*.

3. La série $\sum \frac{1}{2^n}$ converge. Sa somme vaut

$$\sum_{n=0}^{\infty} \frac{1}{2^n} = 2.$$

ATTENTION

Il faut faire attention à la terminologie. Ne pas confondre

- $\sum u_n$ qui est une série
- $\sum_{k=0}^n u_k = S_n$ qui est un élément de $\mathbb{K}$. C'est la somme partielle.
- $\sum_{k=0}^{+\infty} u_k$ la somme qui est un élément de $\mathbb{K}$.

En particulier, cela n'a pas de sens d'écrire « $\sum_{k=0}^{+\infty} u_k$ converge » puisque cela représente un nombre réel ou complexe.

Définition 1.3

Soit $\sum u_n$ une série **convergente**. Pour tout $p \geq 0$ on appelle *reste d'ordre p* ,

$$R_p = \sum_{k=p+1}^{\infty} u_k = \sum_{k=0}^{\infty} u_k - \sum_{k=0}^p u_k = S - S_p.$$

ATTENTION

On ne peut pas parler de reste d'une série divergente.

Exemple : Si on reprend encore la série $\sum \frac{1}{2^n}$ on a

$$R_p = \sum_{k=p+1}^{\infty} \frac{1}{2^k} = \lim_{n \rightarrow \infty} \sum_{k=p+1}^n \frac{1}{2^k} = \lim_{n \rightarrow \infty} \frac{1}{2^{p+1}} \frac{1 - (1/2)^{n-p}}{1 - (1/2)} = \frac{1}{2^p}$$

1.2 Propriétés

Proposition 1.4 (Linéarité de la somme)

Soit $\sum u_n$ et $\sum v_n$ deux séries et λ, μ deux scalaires.

1. Si $\sum u_n$ et $\sum v_n$ convergent alors $\sum(\lambda u_n + \mu v_n)$ converge.
2. Dans ce cas,

$$\sum_{k=0}^{\infty} (\lambda u_k + \mu v_k) = \lambda \sum_{k=0}^{\infty} u_k + \mu \sum_{k=0}^{\infty} v_k$$

Remarque : Cela signifie que l'ensemble des séries convergentes est un sous-espace vectoriel de l'ensemble des séries. De plus l'application définie sur ce sous-espace vectoriel qui associe à une série sa somme est linéaire.

Démonstration : Soit $n \in \mathbb{N}$. Par linéarité :

$$\sum_{k=0}^n (\lambda u_k + \mu v_k) = \lambda \sum_{k=0}^n u_k + \mu \sum_{k=0}^n v_k.$$

Maintenant, par hypothèses les suites $\left(\sum_{k=0}^n u_k\right)_n$ et $\left(\sum_{k=0}^n v_k\right)_n$ convergent vers $\sum_{k=0}^{\infty} u_k$ et $\sum_{k=0}^{\infty} v_k$ respectivement. La conclusion en découle d'après le théorème d'opération sur les limites de suites. $\square$

ATTENTION

Par contre la somme de deux séries divergentes n'est pas toujours divergente. Par exemple $\sum 1 + \frac{1}{2^n}$ et $\sum(-1)$.

1.3 Lien suite-série

Nous allons voir par la suite de nombreuses méthodes pour étudier les séries. Il peut-être tentant dès lors de vouloir voir une **suite** numérique comme étant la suite des sommes partielles d'une série. L'outil de base pour faire cela est le calcul de somme par télescopage.

Lemme 1.5 (Télescopage)

Soit $(u_k)_{k \in \mathbb{N}}$ une suite numérique. Pour tout entier n ,

$$u_n = u_0 + \sum_{k=0}^{n-1} (u_{k+1} - u_k)$$

Démonstration : Cela peut se démontrer de nombreuses manières :

- Avec des points de suspensions

$$\begin{aligned} \sum_{k=0}^{n-1} (u_{k+1} - u_k) &= (u_1 - u_0) + (u_2 - u_1) + (u_3 - u_2) + \cdots + (u_{n-1} - u_{n-2}) + (u_n - u_{n-1}) \\ &= u_n - u_0 \end{aligned}$$

– Avec un décalage d'indice dans les sommes

$$\begin{aligned}
 \sum_{k=0}^{n-1} (u_{k+1} - u_k) &= \sum_{k=0}^{n-1} u_{k+1} - \sum_{k=0}^{n-1} u_k \\
 &= \sum_{k=1}^n u_k - \sum_{k=0}^{n-1} u_k \\
 &= u_n + \sum_{k=1}^{n-1} u_k - u_0 - \sum_{k=1}^{n-1} u_k \\
 &= u_n - u_0
 \end{aligned}$$

– Par récurrence.

On note pour tout entier naturel n , $\mathcal{P}_n$ le prédicat :

$$u_n = u_0 + \sum_{k=0}^{n-1} (u_{k+1} - u_k)$$

– **I** Pour $n = 0$, $u_0 = u_0 + 0$ car la somme pour $k = 0$ à $k = -1$ est une somme vide qui vaut donc 0.

On peut aussi initialiser pour $n = 1$.

$$u_1 = u_0 + (u_1 - u_0)$$

– **H** Soit n un entier naturel fixé. On suppose $\mathcal{P}_n$ et on veut montrer $\mathcal{P}_{n+1}$.

$$\begin{aligned}
 \sum_{k=0}^{(n+1)-1} (u_{k+1} - u_k) &= u_{n+1} - u_n + \sum_{k=0}^{n-1} (u_{k+1} - u_k) \\
 &\stackrel{(HR)}{=} u_{n+1} - u_n + u_n - u_0 \\
 &= u_{n+1} - u_0
 \end{aligned}$$

– **C** Pour tout entier naturel $n \in \mathbb{N}$, $\mathcal{P}_n$.

□

Corollaire 1.6

Soit (u_n) une suite. Les propositions suivantes sont équivalentes

- i) La **suite** (u_n) est convergente.
- ii) La **série** de terme général $u_{k+1} - u_k$ est convergente.

De plus, si elles sont satisfaites on a

$$\lim (u_n) = u_0 + \sum_{k=0}^{\infty} (u_{k+1} - u_k).$$

Démonstration : Il suffit d'utiliser le télescopage :

$$\forall n \in \mathbb{N}, \sum_{k=0}^n (u_{k+1} - u_k) = u_{n+1} - u_0$$

puis de passer à la limite quand n tend vers $+\infty$. $\square$

★ **Méthode** : Pour étudier la nature d'une suite (u_n) , il arrive que l'on étudie la nature de la série $\sum (u_{k+1} - u_k)$. Nous verrons à la fin de ce chapitre que cela permet aussi de déterminer des équivalents.

Exemple : Soit (u_n) définie par $u_0 \in \mathbb{R}_+^*$ et pour tout $n \geq 0$,

$$u_{n+1} = u_n + \frac{1}{n^2 u_n}$$

On veut étudier la nature de (u_n) . Il est clair (par récurrence) que (u_n) est à termes positifs et donc croissante. De ce fait, pour tout entier n ,

$$0 \leq u_{n+1} - u_n \leq \frac{1/u_0}{n^2}$$

En utilisant les théorèmes de comparaisons (que nous allons revoir) on obtient que $\sum (u_{n+1} - u_n)$ converge et donc (u_n) converge.

Remarque : De même que l'on peut retrouver le terme général à partir de la suite des sommes partielles, on peut retrouver le terme général à partir de la suite des restes (dans le cas où la série est convergente). Pour tout entier naturel non nul n :

$$u_n = R_{n-1} - R_n$$

La formule reste vraie pour $n = 0$ en prenant la convention $R_{-1} = \sum_{k=0}^{+\infty} u_k$.

1.4 Divergence grossière

Proposition 1.7

Soit $\sum u_n$ une série. Si elle converge alors le terme général u_n tend vers 0.

ATTENTION

Ce n'est pas une équivalence. On a vu par exemple que $\sum \frac{1}{n}$ diverge.

Démonstration : Soit $\sum u_n$ une série convergente. Pour tout entier n non nul on a

$$u_n = \sum_{k=0}^n u_k - \sum_{k=0}^{n-1} u_k.$$

Maintenant si la série converge, $\sum_{k=0}^n u_k$ et $\sum_{k=0}^{n-1} u_k$ tendent toutes les deux vers la somme de la série. La différence tend donc vers 0. $\square$

Remarque : On peut se servir de la contraposée pour montrer la divergence d'une série. Si le terme général ne tend pas vers 0 alors la série diverge. On dit qu'elle diverge grossièrement.

2 Séries de référence

2.1 Séries géométriques

Les séries géométriques sont les plus simples.

Théorème 1.8 (Séries géométriques)

Soit $z \in \mathbb{C}$. La série de terme général z^n est convergente si et seulement si $|z| < 1$. Dans ce cas la somme de la série est

$$\sum_{k=0}^{\infty} z^k = \frac{1}{1-z}.$$

Démonstration :

- $\Rightarrow$ On remarque que si $|z| \geq 1$ alors pour tout entier k , $|z|^k \geq 1$. De ce fait la série diverge grossièrement. On en déduit le sens voulu par contraposée.
- $\Leftarrow$ Si $|z| < 1$, il suffit de calculer les sommes partielles. Soit $n \in \mathbb{N}$,

$$\sum_{k=0}^n z^k = \frac{1 - z^{n+1}}{1 - z} = \frac{1}{1 - z} - \frac{z^{n+1}}{1 - z}.$$

On a $\left| \frac{z^{n+1}}{1 - z} \right| = |z|^{n+1} \times \left| \frac{1}{1 - z} \right|$ qui tend donc vers 0 quand $n \rightarrow +\infty$. La série converge bien vers $\frac{1}{1 - z}$.

□

2.2 Séries de Riemann

Définition 1.9 (Séries de Riemann)

On appelle *séries de Riemann* les séries de la forme $\sum_{n \geq 1} n^{-s} = \sum_{n \geq 1} \frac{1}{n^s}$ où $s \in \mathbb{R}$.

On considère ici des séries qui commencent à $n = 1$.

Théorème 1.10 (TRES IMPORTANT)

La série de Riemann $\sum_{n \geq 1} \frac{1}{n^s}$ est convergente si et seulement si $s > 1$.

Démonstration : C'est un cas particulier de la comparaison série-intégrale qui sera revue un peu plus loin.

□

Remarque culturelle : La limite des sommes de Riemann est difficile à calculer dans la majorité des cas. On pose

$$\forall s \in]1, +\infty[, \quad \zeta(s) = \sum_{n=1}^{+\infty} \frac{1}{n^s}.$$

On sait calculer la valeur de la fonction ζ aux entiers pairs :

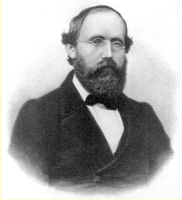
$$\zeta(2) = \frac{\pi^2}{6}, \quad \zeta(4) = \frac{\pi^4}{90}.$$

De manière générale, $\zeta(2k) = C_k \pi^{2k}$ où $C_k \in \mathbf{Q}$. On sait peu de choses que les $\zeta(2k+1)$. On conjecture que ce sont tous des irrationnels comme $\zeta(3)$ (Apéry : 1979).

L'étude de la fonction ζ est extrêmement importante en mathématiques et en particulier en arithmétique.

On peut montrer que l'on peut définir $\zeta(s)$ pour tout nombre complexe s tel que $\Re(s) > 1$. Ensuite, on peut prolonger cette fonction de manière unique à $\mathbf{C} \setminus \{1\}$ en une fonction holomorphe (une généralisation des fonctions dérivables pour les fonctions de $\mathbf{C}$ dans $\mathbf{C}$). On montre alors simplement que $\zeta(-1) = -\frac{1}{12}$ et que pour tout entier naturel $p > 0$, $\zeta(-2p) = 0$. L'hypothèse de Riemann affirme alors que pour tout s qui n'est pas de la forme $-2p$ avec $p \in \mathbf{N}^*$, $\zeta(s) = 0$ implique que $\Re(s) = \frac{1}{2}$.

Matheux (Bernhard Riemann : 1826 - 1866)



Georg Friedrich Bernhard Riemann, né le 17 septembre 1826 à Breselenz, mort le 20 juillet 1866 à Selasca en Italie, est un mathématicien allemand. Influent sur le plan théorique, il a apporté de nombreuses contributions importantes à la topologie, l'analyse, la géométrie différentielle et le calcul, certaines d'entre elles ayant permis par la suite le développement de la relativité générale.

2.3 Série exponentielle

Donnons le résultat sur la série exponentielle

Théorème 1.11 (Séries géométriques)

Soit $z \in \mathbf{C}$. La série de terme général $\frac{z^n}{n!}$ est convergente. La somme de la série est

$$\sum_{k=0}^{\infty} \frac{z^k}{k!} = \exp(z)$$

Démonstration : Soit $z \in \mathbf{C}$

On pose $f : z \mapsto \exp(tz)$. C'est une fonction de classe $\mathcal{C}^\infty$.

On peut donc utiliser les formules de Taylor-Lagrange avec reste intégral entre 0 et 1.

Rappelons que si $f \in \mathcal{C}^\infty([a, b], \mathbf{R})$ alors pour tout entier n ,

$$f(b) = \sum_{k=0}^n \frac{f^{(k)}(a)}{k!} (b-a)^k + \int_a^b \frac{(b-t)^n}{n!} f^{(n+1)}(t) dt$$

Comme pour tout entier k , $f^{(k)} = z^k f$, on a $f^{(k)}(0) = z^k$. Précisément, pour tout entier n ,

$$f(1) = \sum_{k=0}^n \frac{z^k}{k!} + \int_0^1 \frac{(1-t)^n}{n!} z^{n+1} e^{tz} dt$$

On va alors montrer que

$$\lim_{n \rightarrow \infty} \int_0^1 \frac{(1-t)^n}{n!} z^{n+1} e^{tz} dt = 0$$

Pour cela on va considérer son module.

$$\left| \int_0^1 \frac{(1-t)^n}{n!} z^{n+1} e^{tz} dt \right| \leq |z|^{n+1} e^{|z|} \int_0^1 \frac{(1-t)^n}{n!} dt = \frac{|z|^{n+1} e^{|z|}}{(n+1)!}$$

On obtient par croissance comparée que $\lim_{n \rightarrow \infty} \int_0^1 \frac{(1-t)^n}{n!} z^{n+1} e^{tz} dt = 0$ ce qui prouve que la série $\sum \frac{z^k}{k!}$ est convergente et que

$$\sum_{k=0}^{+\infty} \frac{z^k}{k!} = \exp(z)$$

□

Remarque : Nous verrons un peu plus loin une autre preuve de la convergence de la série exponentielle. On pourra alors définir, pour $z \in \mathbb{C}$, $\exp(z)$ comme la somme de cette série.

2.4 Exemples de calculs de sommes

Nous allons donner quelques techniques usuelles de calcul de sommes. Nous nous contenterons de les illustrer sur des exemples.

Séries géométriques dérivées :

Soit $x \in \mathbb{R}$ tel que $|x| < 1$. On a vu que les sommes partielles de la série géométrique de raison x étaient données par

$$\forall n \in \mathbb{N}, S_n = \sum_{k=0}^n x^k = \frac{1 - x^{n+1}}{1 - x}.$$

En considérant les deux termes comme des fonctions définies sur $] -1, 1[$, on peut les dériver (par rapport à x) pour obtenir

$$\forall n \in \mathbb{N}, \sum_{k=0}^n kx^{k-1} = \frac{-(n+1)x^n(1-x) + (1-x^{n+1})}{(1-x)^2}$$

Comme $|x| < 1$, le terme de droite tend vers $\frac{1}{(1-x)^2}$. Donc la série $\sum kx^{k-1}$ converge et $\sum_{k=0}^{+\infty} kx^{k-1} = \frac{1}{(1-x)^2}$.

Remarques :

1. On peut aussi obtenir ce résultat avec une somme double

$$\sum_{k=1}^n kx^{k-1} = \sum_{k=1}^n \left(\sum_{i=1}^k 1 \right) x^{k-1} = \sum_{i=1}^n \left(\sum_{k=i}^n x^{k-1} \right) = \frac{1}{1-x} \sum_{i=1}^n (x^{i-1} - x^n) = \frac{1-x^n}{(1-x)^2} - \frac{nx^n}{1-x}.$$

Il ne reste plus qu'à faire tendre n vers $+\infty$ pour retrouver que la série est convergente et que

$$\sum_{k=0}^{+\infty} kx^{k-1} = \frac{1}{(1-x)^2}.$$

2. Ce résultat reste vrai pour $x \in \mathbb{C}$, pour le vérifier il suffit de considérer la fonction

$$\varphi : t \mapsto \sum_{k=0}^n (tx)^k$$

Utilisation d'un télescopage

On veut montrer que la série $\sum_{k \geq 1} \frac{1}{k(k+1)}$ converge et calculer la somme.

On commence en remarquant que pour tout $k \geq 1$, $\frac{1}{k(k+1)} = \frac{1}{k} - \frac{1}{k+1}$. De ce fait, pour tout entier $n \geq 1$,

$$\sum_{k=1}^n \frac{1}{k(k+1)} = \sum_{k=1}^n \frac{1}{k} - \frac{1}{k+1} = 1 - \frac{1}{n+1}$$

En faisant alors tendre n vers $+\infty$ on voit que la série converge et que

$$\sum_{k=1}^{+\infty} \frac{1}{k(k+1)} = 1$$

Exercice : Montrer que la série $\sum_{k \geq 2} \frac{5k+3}{k(k^2-1)}$ converge et calculer la somme.

On utilise une décomposition en éléments simples de la fraction rationnelle $\frac{5X+3}{X(X^2-1)}$. On trouve

$$\frac{5X+3}{X(X^2-1)} = -\frac{3}{X} - \frac{1}{X+1} + \frac{4}{X-1}.$$

De ce fait pour $n \geq 2$,

$$\sum_{k=2}^n \frac{5k+3}{k(k^2-1)} = -\sum_{k=2}^n \frac{3}{k} - \sum_{k=2}^n \frac{1}{k+1} + \sum_{k=2}^n \frac{4}{k-1} = -\sum_{k=2}^n \frac{3}{k} - \sum_{k=3}^{n+1} \frac{1}{k} + \sum_{k=1}^{n-1} \frac{4}{k}.$$

En simplifiant on obtient :

$$\sum_{k=2}^n \frac{5k+3}{k(k^2-1)} = -\frac{3}{2} - \frac{3}{n} - \frac{1}{n+1} - \frac{1}{n} + 4 + \frac{4}{2} \xrightarrow{n \rightarrow \infty} \frac{9}{2}.$$

On en déduit que la série est convergente et que $\sum_{k=2}^{+\infty} \frac{5k+3}{k(k^2-1)} = \frac{9}{2}$.

Utilisation des formules de Taylor

On peut mettre en place la méthode utilisée pour calculer la somme de la série exponentielle dans d'autres cas. De manière générale si f est une fonction de classe $\mathcal{C}^\infty$ sur un intervalle I , d'après la formule de Taylor avec reste intégrale pour a, b dans I et tout $n \in \mathbb{N}$,

$$f(b) - f(a) = \sum_{k=0}^n \frac{f^{(k)}(a)}{k!} (b-a)^k + \int_a^b \frac{(b-t)^n}{n!} f^{(n+1)}(t) dt.$$

En particulier,

$$\sum_{k=0}^n \frac{f^{(k)}(a)}{k!} (b-a)^k = f(b) - f(a) - \int_a^b \frac{(b-t)^n}{n!} f^{(n+1)}(t) dt.$$

Si on arrive à montrer que le terme $\int_a^b \frac{(b-t)^n}{n!} f^{(n+1)}(t) dt$ tend vers 0 quand n tend vers $+\infty$, on obtient que la série $\sum_{k \geq 0} \frac{f^{(k)}(a)}{k!} (b-a)^k$ converge et que

$$\sum_{k=0}^{+\infty} \frac{f^{(k)}(a)}{k!} (b-a)^k = f(b) - f(a).$$

Exercice : En utilisant cette méthode montrer que la série $\sum_{n \geq 1} \frac{(-1)^{n-1}}{n}$ converge et que

$$\sum_{n=1}^{+\infty} \frac{(-1)^{n-1}}{n} = \ln 2.$$

On considère la fonction $f : x \mapsto \ln(1+x)$ sur l'intervalle $[0, 1]$. La fonction est $\mathcal{C}^\infty$. De plus, pour tout $k \in \mathbb{N}^*$,

$$f^{(k)} : x \mapsto (-1)^{(k-1)} (k-1)! (1+x)^{-k}.$$

En particulier, $f^{(k)}(0) = (-1)^{k-1} (k-1)!$ et donc pour $a = 0$ et $b = 1$:

$$\sum_{k=0}^n \frac{f^{(k)}(0)}{k!} = 0 + \sum_{k=1}^n \frac{(-1)^{k-1}}{k}.$$

De plus, le reste intégral vaut

$$r_n = \int_0^1 \frac{(1-t)^n}{n!} \frac{(-1)^n n!}{(1+t)^{n+1}} dt.$$

De ce fait,

$$|r_n| \leq \int_0^1 (1-t)^n dt = \frac{1}{n+1}.$$

Finalement $(r_n) \rightarrow 0$ et donc la série $\sum_{n \geq 1} \frac{(-1)^{n-1}}{n}$ converge et

$$\sum_{n=1}^{+\infty} \frac{(-1)^{n-1}}{n} = \ln 2.$$

3 Séries à termes positifs

Dans tout ce paragraphe nous supposons que les termes de la série (les u_k) sont des réels positifs. Tous les résultats se généralisent aux séries à termes négatifs en multipliant par -1 .

3.1 Généralités

Proposition 1.12

Soit $\sum u_n$ une série à termes positifs. La suite des sommes partielles est croissante.

Démonstration : Évident. □

On peut appliquer le théorème de convergence monotone.

Proposition 1.13

Soit $\sum u_n$ une série à termes positifs. Elle est convergente si et seulement si la suite des sommes partielles est majorée.

Dans ce cas

$$\sum_{k=0}^{+\infty} u_k = \sup \left\{ \sum_{k=0}^n u_k \mid n \in \mathbb{N} \right\}.$$

ATTENTION

Dans le cas d'une série à termes positifs, la suite $(S_n)_{n \in \mathbb{N}}$ des sommes partielles n'a que deux comportements asymptotiques possibles. Elle peut tendre vers une limite finie (série convergente) ou tendre vers $+\infty$ (série divergente).

Dans le cas où les termes sont positifs, on peut dans tous les cas calculer la somme $\sum_{n=0}^{\infty} u_n$ en convenant que si le résultat est un réel la série converge et que, dans le cas contraire, la série diverge.

On note alors $\sum_{n=0}^{\infty} u_n = +\infty$.

On précisera explicitement que les termes que l'on somme sont positifs et que les calculs sont réalisés dans $\mathbb{R}_+^*$.

3.2 Théorème de comparaison par inégalités pour les séries à termes positifs

En pratique on majore (ou minore) souvent par une autre série.

Théorème 1.14

Soit $\sum u_n$ et $\sum v_n$ deux séries à termes positifs. On suppose que pour tout entier n , $u_n \leq v_n$.

1. Si $\sum v_n$ converge alors $\sum u_n$ aussi et $\sum_{n=0}^{+\infty} u_n \leq \sum_{n=0}^{+\infty} v_n$.
2. Si $\sum u_n$ diverge alors $\sum v_n$ aussi.

ATTENTION

Par contre, avec les hypothèses ci-dessus, le fait que $\sum v_n$ diverge ne donne aucune information sur la nature de la série $\sum u_n$. De même le fait que $\sum u_n$ converge ne donne pas d'information sur la nature $\sum v_n$.

Remarque : Comme dans le cas classique pour les suites, on peut supposer que l'inégalité $u_n \leq v_n$ ne soit vraie qu'à partir d'un certain rang. Cependant, on n'a plus l'information sur les sommes dans ce cas.

Démonstration : Si on note U_n et V_n les sommes partielles. On a donc

$$\forall n \in \mathbb{N}, U_n \leq V_n.$$

1. Si on suppose que $\sum v_n$ converge et si on note $V = \lim_{n \rightarrow +\infty} V_n$ on en déduit que pour tout entier n ,

$$U_n \leq V_n \leq V.$$

De ce fait (U_n) est majorée par V (et pas par V_n ce qui ne voudrait rien dire...). Comme elle est croissante elle converge et

$$\lim_{n \rightarrow \infty} U_n \leq V.$$

2. C'est juste la contraposée du point 1.

□

Exercice : Écrire la démonstration de la première partie dans le cas où l'on suppose juste :

$$\exists N \in \mathbb{N}, \forall n \in \mathbb{N}, n \geq N \Rightarrow u_n \leq v_n$$

Exemples :

1. Cherchons la nature $\sum_{n \geq 1} \frac{1}{n^2}$. On remarque que

$$\forall n \geq 2, \quad 0 \leq \frac{1}{n^2} \leq \frac{1}{n(n-1)}$$

On a déjà vu que l'on peut montrer que la série $\sum_{n \geq 2} \frac{1}{n(n-1)}$ est convergente par un télescopage.

Le théorème ci-dessus affirme donc que $\sum_{n \geq 2} \frac{1}{n^2}$ est une série de convergente.

De plus :

$$\zeta(2) = \sum_{n=1}^{+\infty} \frac{1}{n^2} \leq 1 + \sum_{n=2}^{+\infty} \frac{1}{n^2} \leq 1 + \sum_{n=2}^{+\infty} \frac{1}{n(n-1)} = 2$$

car

$$\sum_{n=2}^N \frac{1}{n(n-1)} = \sum_{n=2}^N \frac{1}{n-1} - \frac{1}{n} = 1 - \frac{1}{N} \xrightarrow{N \rightarrow +\infty} 1$$

2. Cherchons la nature $\sum_{n \geq 2} \frac{1}{\ln n}$. On a pour tout entier $n \geq 2$,

$$0 \leq \frac{1}{n} \leq \frac{1}{\ln n}.$$

On sait que la série harmonique diverge donc $\sum_{n \geq 2} \frac{1}{\ln n}$ diverge aussi.

ATTENTION

Cela n'est bien évidemment plus vrai si la série n'est pas à termes positifs.

1. On sait que c'est une série convergente car $2 > 1$ mais on peut faire une preuve élémentaire

3.3 Rappels sur les relations de comparaison

Dans le cas précédent on a vu que l'on pouvait déduire la nature de $\sum u_n$ si $\sum v_n$ convergait mais que l'on avait aucune information si cette dernière divergeait. On va résoudre ce problème grâce à des équivalents.

Définition 1.15

Soit (u_n) et (v_n) deux suites. On suppose qu'elles ne s'annulent pas à partir d'un certain rang.

1. On dit que (u_n) et (v_n) sont équivalentes et on note $(u_n) \sim (v_n)$ si $\frac{u_n}{v_n} \rightarrow 1$.
2. On dit que (u_n) est négligeable devant (v_n) et on note $(u_n) = o(v_n)$ si $\frac{u_n}{v_n} \rightarrow 0$.
3. On dit que (u_n) est dominée par (v_n) et on note $(u_n) = O(v_n)$ si $\left(\frac{u_n}{v_n}\right)$ est bornée.

Remarques :

1. Les notations ci-dessus s'appliquent à des suites, cependant, pour alléger les notations on notera souvent $u_n \sim v_n$ à la place de $(u_n) \sim (v_n)$.
2. L'hypothèse de non annulation à partir d'un certain rang et juste une hypothèse technique qui sera vérifiée dans la très grande majorité des cas étudiés. Elle sert à définir la suite de terme général $\frac{u_n}{v_n}$ à partir d'un certain rang. On peut cependant traiter le cas général en disant que $(u_n) \sim (v_n)$ (resp. $u_n = o(v_n)$; $u_n = O(v_n)$) s'il existe une suite (ε_n) tendant vers 1 (resp. tendant vers 0 ; qui est bornée) telle que pour tout entier $u_n = v_n \times \varepsilon_n$. Avec cette notation la suite (ε_n) « joue le rôle » de la suite de terme général $\frac{u_n}{v_n}$.
3. La notion d'équivalents permet de distinguer entre deux suites qui tendent vers 0 ou qui tendent vers $+\infty$. Par contre cela ne sert à rien pour les suites qui tendent vers $\ell \notin \{0, \pm\infty\}$. En effet si (u_n) et (v_n) tendent vers $\ell \in \mathbb{R}^*$ alors $u_n \sim v_n$.

Exemple : On peut vérifier que $\frac{2n^2 + 3}{3n - 1} \sim \frac{2n}{3}$. En effet

$$\frac{\frac{2n^2 + 3}{3n - 1}}{\frac{2n}{3}} = \frac{6n^2 + 9}{6n^2 - 2n} \longrightarrow 1$$

Proposition 1.16 (Propriétés de la relation d'équivalence)

1. La relation d'équivalence est une relation d'équivalence : symétrie, réflexivité, transitivité
2. Si $u_n \sim v_n$ alors les suites (u_n) et (v_n) sont de même signe à partir d'un certain rang.

Démonstration :

1. Montrons que la relation d'équivalence vérifie les axiomes des relations d'équivalences²
 - Réflexivité : Soit (u_n) une suite qui ne s'annule pas à partir d'un certain rang. Il est clair que $\frac{u_n}{u_n} \longrightarrow 1$ donc $u_n \sim u_n$.

² et passons sous silence sur quel ensemble nous travaillons pour pas avoir de problème avec l'hypothèse technique que nous avons fait.

- Symétrie : Soit (u_n) et (v_n) deux suites qui ne s'annulent pas à partir d'un certain rang. On suppose que $u_n \sim v_n$ ce qui signifie que $\frac{u_n}{v_n} \longrightarrow 1$. En remarquant que pour tout entier n assez grand,

$$\frac{v_n}{u_n} = \frac{1}{\frac{u_n}{v_n}} \longrightarrow \frac{1}{1} = 1$$

on obtient que $v_n \sim u_n$.

- Transitivité : Soit (u_n) , (v_n) et (w_n) trois suites qui ne s'annulent pas à partir d'un certain rang. On suppose que $u_n \sim v_n$ et que $v_n \sim w_n$. Pour n assez grand,

$$\frac{u_n}{w_n} = \frac{u_n}{v_n} \times \frac{v_n}{w_n} \longrightarrow 1 \times 1 = 1$$

Cela signifie que $u_n \sim w_n$.

2. Laissée au lecteur.

□

ATTENTION

- On ne peut pas ajouter des équivalents sans réfléchir ! Par exemple $n^2 + n \sim n^2$ et $-n^2 + 1 \sim -n^2$ mais $n + 1 = (n^2 + n) + (-n^2 + 1) \not\sim 0$
- On ne peut pas composer (à droite) des équivalents. Par exemple $n^2 \sim n^2 + n$ mais $f(n^2 + n)$ n'est pas nécessairement équivalent à $f(n^2)$ par exemple si $f : x \mapsto e^x$.

Proposition 1.17 (Opérations sur les équivalents)

1. Si $u_n \sim v_n$ et $u'_n \sim v'_n$ alors $u_n u'_n \sim v_n v'_n$ et $\frac{u_n}{u'_n} \sim \frac{v_n}{v'_n}$
2. Si $u_n \sim v_n$ et α est un réel fixé alors $u_n^\alpha \sim v_n^\alpha$
3. Si $v_n = o(u_n)$ alors $u_n + v_n \sim u_n$

Démonstration :

À écrire

□

ATTENTION

Il est important de voir que l'élévation à la puissance α ne donne un résultat correct (en généralité) que si α ne dépend pas de n . Par exemple on a $1 + \frac{1}{n} \sim 1$ mais $(1 + \frac{1}{n})^n \not\sim 1^n$. Par contre on a bien $n^2 - n \sim n^2$ et $\sqrt{n^2 - n} \sim \sqrt{n^2} = n$.

★ **Méthode** : Équivalents et somme : On peut gérer les sommes d'équivalents en faisant attention

- Un terme est négligeable devant l'autre : On cherche un équivalent à $\ln(1 + \frac{1}{n}) + \frac{1}{n^2}$. On sait que $\ln(1 + \frac{1}{n}) \sim \frac{1}{n}$ donc $\frac{1}{n^2}$ est négligeable devant $\ln(1 + \frac{1}{n})$. On en déduit que

$$\ln\left(1 + \frac{1}{n}\right) + \frac{1}{n^2} \sim \frac{1}{n}$$

- Les deux termes sont du « même ordre » mais ne se compensent pas : Si on regarde $\ln(1 + \frac{1}{n}) + \frac{1}{n}$.

Cette fois $\ln(1 + \frac{1}{n}) \sim \frac{1}{n}$. On peut se douter que $\ln(1 + \frac{1}{n}) + \frac{1}{n} \sim \frac{2}{n}$. On a effet, par calcul sur les limites

$$\frac{\ln(1 + \frac{1}{n}) + \frac{1}{n}}{\frac{2}{n}} = \frac{1}{2} \left(\frac{\ln(1 + \frac{1}{n})}{\frac{1}{n}} + 1 \right) \xrightarrow{n \rightarrow \infty} 1$$

On peut aussi utiliser des développements limités.

- Les deux termes sont du « même ordre » et se compensent : si on regarde $\ln(1 + \frac{1}{n}) - \frac{1}{n}$, on peut essayer de faire de même. Mais là, les deux sont équivalents à des termes en $\frac{1}{n^2}$ qui se compensent. On doit faire un développement limité afin d'obtenir le terme en $\frac{1}{n^2}$ ou $\frac{1}{n^3}$.

Précisément,

$$\ln\left(1 + \frac{1}{n}\right) = \frac{1}{n} - \frac{1}{2n^2} + o\left(\frac{1}{n^2}\right)$$

On a donc

$$\ln\left(1 + \frac{1}{n}\right) - \frac{1}{n} \sim -\frac{1}{2n^2}$$

3.4 Théorème de comparaison par équivalents pour les séries à termes positifs

Théorème 1.18 (Théorème de comparaison pour les séries à termes positifs)

Soit $\sum u_n$ et $\sum v_n$ deux séries à termes positifs. On suppose que $u_n \sim v_n$ alors $\sum u_n$ et $\sum v_n$ sont de même nature.

Remarques :

1. Ce théorème est TRES utile. Dans la très grande majorité des cas pour étudier la **nature** d'une série à termes **positifs** on cherche un équivalent simple (via un DL par exemple) et on en déduit ainsi sa nature.
2. On utilisera dans ce cours l'abréviation TCPSP.
3. Il suffit que les deux séries soient positives à partir d'un certain rang.

ATTENTION

Cela n'est plus vrai si la série n'est pas à termes positifs (ou à termes de signe constant). Par exemple nous allons voir que la série de terme général $u_n = \frac{(-1)^n}{\sqrt{n}}$ converge mais la série de terme général $v_n = \frac{(-1)^n}{\sqrt{n}} + \frac{1}{n}$ diverge (somme d'une série convergente et d'une série divergente) par contre on a bien $u_n \sim v_n$.

Démonstration : Par définition $\frac{u_n}{v_n} \rightarrow 1$. De ce fait il existe $N \in \mathbb{N}$ tel que

$$\forall n \in \mathbb{N}, n \geq N \Rightarrow \frac{1}{2} \leq \frac{u_n}{v_n} \leq \frac{3}{2}$$

(définition de la limite en prenant $\varepsilon = 1/2$). C'est-à-dire que

$$\forall n \in \mathbb{N}, n \geq N \Rightarrow \frac{1}{2}v_n \leq u_n \leq \frac{3}{2}v_n$$

On en déduit (grâce à l'inégalité de droite) que si $\sum v_n$ converge alors $\sum u_n$ aussi. De plus, l'inégalité de gauche nous permet de dire que $\sum v_n$ diverge alors $\sum u_n$ aussi.

On en déduit que $\sum u_n$ a la même nature que $\sum v_n$.

□

Remarque : On voit dans la preuve que ce qui compte est qu'il existe $(m, M) \in \mathbb{R}_+^*$ tels que, à partir d'un certain rang,

$$\forall n \in \mathbb{N}, n \geq N \Rightarrow m \leq \frac{u_n}{v_n} \leq M$$

C'est en particulier vrai si $\frac{u_n}{v_n}$ a une limite finie non nulle (et pas nécessairement 1). Cela montre que si $u_n = O(v_n)$ et que $\sum v_n$ converge alors $\sum u_n$ aussi.

Exemples :

1. La série de terme général $\sum \frac{3n+2}{n^3+\ln n}$ converge car $\frac{3n+2}{n^3+\ln n} \sim \frac{3}{n^2}$ et que $\sum \frac{1}{n^2}$ est une série de Riemann convergente.
2. La série de terme général $\sum \frac{3+n \ln n}{n^2+n+1}$ diverge car $\frac{3+n \ln n}{n^2+n+1} \sim \frac{\ln n}{n}$ et que $\sum \frac{\ln n}{n}$ diverge car $\frac{\ln n}{n} \geq \frac{1}{n}$ et que cette dernière est une série de Riemann divergente.

3.5 Règle de d'Alembert

Le principe de la règle de d'Alembert va être de comparer le terme général d'une série à termes strictement positifs avec une série géométrique. On rappelle que pour $x \in \mathbb{R}_+$, $\sum x^k$ converge si et seulement si $x < 1$.

Proposition 1.19 (Comparaison exponentielle)

Soit $\sum u_n$ et $\sum v_n$ deux séries à termes strictement positifs. Si pour tout entier n , $\frac{u_{n+1}}{u_n} \leq \frac{v_{n+1}}{v_n}$ et que $\sum v_n$ converge alors $\sum u_n$ aussi.

Remarque : Comme toujours, on peut se contenter du fait que u_n et v_n soit strictement positifs APCR et de même pour l'inégalité.

Démonstration : Il suffit d'avoir recours à un produit télescopique. En effet pour tout entier $n \geq 0$,

$$\frac{u_n}{u_0} = \prod_{k=0}^{n-1} \frac{u_{k+1}}{u_k} \leq \prod_{k=0}^{n-1} \frac{v_{k+1}}{v_k} = \frac{v_n}{v_0}.$$

Finalement, $u_n \leq C.v_n$ où $C = \frac{u_0}{v_0}$. On peut alors conclure en utilisation le théorème de comparaison pour les séries à termes positifs. □

Remarque : Par contraposée, on obtient que si pour tout entier n , $\frac{u_{n+1}}{u_n} \leq \frac{v_{n+1}}{v_n}$ et que $\sum u_n$ diverge alors $\sum v_n$ aussi.

Corollaire 1.20 (Règle de d'Alembert)

Soit $\sum u_n$ une série à termes strictement positifs. On suppose que $\frac{u_{n+1}}{u_n} \rightarrow \ell \in \overline{\mathbb{R}}_+$.

- Si $\ell > 1$ alors $\sum u_n$ diverge.
- Si $\ell < 1$ alors $\sum u_n$ converge.
- Si $\ell = 1$. On ne peut rien dire....

Démonstration : Il suffit de faire une comparaison logarithmique avec une série géométrique. Traitons le premier cas.

On suppose que $\ell > 1$. On pose $x \in]1, \ell[$ et $v_n = x^n$.

On sait que $\frac{u_{n+1}}{u_n}$ tend vers ℓ , donc, à partir d'un certain rang, $\frac{u_{n+1}}{u_n} \geq x = \frac{v_{n+1}}{v_n}$. Or $\sum v_n$ diverge donc $\sum u_n$ diverge aussi par comparaison exponentielle. $\square$

Matheux (Jean Le Rond d'Alembert : 1717 - 1783)



Jean Le Rond d'Alembert, parfois écrit « Jean le Rond D'Alembert » ou « Dalember », voire « Dalambert », est un mathématicien, physicien, philosophe et encyclopédiste français, né le 16 novembre 1717 à Paris où il est mort le 29 octobre 1783. Il est célèbre pour avoir été l'inventeur d'un principe de l'équilibre que Condorcet explique dans son Éloge de d'Alembert. Il a ainsi fixé une liaison entre les lois du mouvement. Par son théorème maintenant nommé « théorème d'Alembert », il perçoit la présence de n racines dans toute équation algébrique de degré n . En 1744, il est l'inventeur de cette nouvelle branche des mathématiques, le calcul aux dérivées partielles, qui introduit des fonctions arbitraires. En 1749, à la suite de ses recherches en mathématiques sur les équations différentielles et les dérivées partielles, il est appelé pour diriger l'Encyclopédie avec Denis Diderot. Des écoles, des rues et des centres de recherche portent son nom.

Exemple : Soit x un réel positif. On pose $u_n = \frac{x^n}{n!}$. On cherche la nature de $\sum u_n$.

- Si $x = 0$ le résultat est évident.
- Si $x > 0$,

$$\frac{u_{n+1}}{u_n} = \frac{x^{n+1}}{(n+1)!} \cdot \frac{n!}{x^n} = \frac{x}{n+1} \rightarrow 0.$$

On en déduit que la série converge.

On retrouve ainsi que la série exponentielle converge (au moins pour x réel positif) sans supposer connaître au préalable l'existence de la fonction exponentielle.

Exercice : Étudier $\sum \frac{n!}{n^n}$.

Remarques :

1. Cette règle n'est « pas très fine ». En particulier, les séries de Riemann ne peuvent pas être traitées de cette manière car si $u_n = \frac{1}{n^s}$ alors $\frac{u_{n+1}}{u_n} \rightarrow 1$.
2. Cette règle sera très utile dans le chapitre sur les séries entières.

4 Séries absolument convergentes

On a vu les critères de convergence pour les séries à termes de signe constant. Voyons un peu ce que l'on peut faire pour des séries dont le terme général ne garde pas un signe constant (ou même des séries à termes complexes).

4.1 Définition

Définition 1.21

Soit $\sum u_n$ une série éventuellement complexe. On dit qu'elle est absolument convergente si la série (à termes réels positifs) $\sum |u_n|$ est convergente.

Remarque : L'intérêt est que pour savoir si une série est absolument convergente, on est ramené à étudier une série à termes (réels) positifs.

Exemples :

1. La série $\sum_{n \geq 1} \frac{(-1)^n}{n}$ n'est pas absolument convergente car $\left| \frac{(-1)^n}{n} \right| = \frac{1}{n}$ et que $\sum \frac{1}{n}$ n'est pas convergente.
2. La série $\sum \frac{\cos(n)}{n^2}$ est absolument convergente car $\sum \left| \frac{\cos(n)}{n^2} \right|$ est une série à termes positifs et que $\left| \frac{\cos(n)}{n^2} \right| \leq \frac{1}{n^2}$ cette dernière étant une série de Riemann convergente.

4.2 Applications aux séries dont le terme général n'a pas nécessairement un signe constant

Le résultat principal est qu'une série complexe (ou réelle) absolument convergente est convergente.

Théorème 1.22

Soit $\sum u_n$ une série dont le terme général est éventuellement complexe. Si elle est absolument convergente alors elle est convergente.

Démonstration :

- Commençons par étudier le cas où $\sum u_n$ une série à termes réels absolument convergente. Pour tout réel x , on note $x^+ = \max(x, 0)$ sa partie positive et $x^- = \max(-x, 0)$ sa partie négative de sorte que $|x| = x^+ - x^-$ et $|x| = x^+ + x^-$.

On considère les séries à termes positifs $\sum u_n^+$ et $\sum u_n^-$. Comme on sait que pour tout entier n , $0 \leq u_n^+ \leq |u_n|$ et $0 \leq u_n^- \leq |u_n|$ et que $\sum |u_n|$ est une série convergente, on en déduit que $\sum u_n^+$ et $\sum u_n^-$ sont des séries convergentes. Par linéarité on obtient alors que $\sum u_n$ est une série convergente.

- On suppose que $\sum u_n$ une série à termes complexes absolument convergente. On considère les séries à termes réels $\sum \Re(u_n)$ et $\sum \Im(u_n)$. Là encore on sait que $0 \leq \Re(u_n) \leq |u_n|$ et $0 \leq |\Im(u_n)| \leq |u_n|$. Les séries $\sum \Re(u_n)$ et $\sum \Im(u_n)$ sont donc des séries à termes réels absolument convergentes. D'après ce qui précède, elles convergent. On en déduit que $\sum u_n$ converge par linéarité.

□

ATTENTION

Ce n'est pas une équivalence. Par exemple $\sum \frac{(-1)^n}{n+1}$ n'est pas absolument convergente mais elle est convergente. On dit que c'est une série semi-convergente.

Exemples :

1. Soit $z \in \mathbb{C}$. On pose pour tout entier naturel n , $u_n = \frac{z^n}{n!}$. On a alors $|u_n| = \frac{|z|^n}{n!}$. On a vu ci-dessus via la règle de d'Alembert que la série $\sum \frac{|z|^n}{n!}$ convergeait (car $|z|$ est un réel positif). On en déduit que la série $\sum \frac{z^n}{n!}$ est absolument convergente donc convergente. C'est « la bonne » définition de l'exponentielle d'un nombre réel ou complexe. **Par définition,**

$$\forall z \in \mathbb{C}, \exp(z) = \sum_{n=0}^{\infty} \frac{z^n}{n!}$$

Il restera à vérifier que cela vérifie bien les propriétés attendues : la fonction $x \mapsto \exp(x)$ est dérivable sur $\mathbb{R}$ et pour tout $x \in \mathbb{R}$, $\exp'(x) = \exp(x)$; pour z, z' dans $\mathbb{C}$, $\exp(z + z') = \exp(z) \exp(z')$.

2. Il faut faire attention aux séries semi-convergentes.

Posons pour $k \geq 1$, $u_k = \frac{(-1)^{k-1}}{k}$. On a vu que $\sum_{k=1}^{\infty} u_k = 1 - \frac{1}{2} + \frac{1}{3} - \frac{1}{4} + \dots = \ln(2)$.

Soit σ la permutation de $\mathbb{N}^*$ qui consiste à prendre les deux premiers nombres impairs et puis le premier nombre pair et ainsi de suite :

k	1	2	3	4	5	6	7	8	9	10	11	12	13	14	...
$\sigma(k)$	1	3	2	5	7	4	9	11	6	13	15	8	17	19	...

C'est-à-dire que l'on pose

$$\sigma(k) = \begin{cases} 2p & \text{si } k = 3p \\ 4p + 1 & \text{si } k = 3p + 1 \\ 4p + 3 & \text{si } k = 3p + 2 \end{cases}$$

Pour $N \geq 1$, on calcule

$$\sum_{k=1}^{3N} u_{\sigma(k)} = -\frac{1}{2} \sum_{p=0}^{N-1} \frac{1}{p+1} + \sum_{p=0}^{N-1} \left(\frac{1}{4p+1} + \frac{1}{4p+3} \right) = -\frac{1}{2} \sum_{p=1}^N \frac{1}{p} + \sum_{\substack{p=1 \\ p \text{ impair}}}^N \frac{1}{p} = -\frac{1}{2} \sum_{p=1}^N \frac{1}{p} + \left(\sum_{p=1}^{4N} \frac{1}{p} - \frac{1}{2} \sum_{p=1}^{2N} \frac{1}{p} \right)$$

En utilisant que $\sum_{k=1}^n \frac{1}{k} = H_n = \ln(n) + \gamma + o(1)$ on obtient que

$$\lim_{N \rightarrow +\infty} \sum_{k=1}^{3N} u_{\sigma(k)} = \frac{3}{2} \ln(2)$$

Comme $u_{\sigma(3N+1)}$ et $u_{\sigma(3N+2)}$ tendent vers 0 on a montré que la série $\sum_{k \geq 1} u_{\sigma(k)}$ converge et que sa somme est

$$\sum_{k=1}^{\infty} u_{\sigma(k)} = 1 + \frac{1}{3} - \frac{1}{2} + \frac{1}{5} + \frac{1}{7} - \frac{1}{4} + \frac{1}{9} + \frac{1}{11} + \dots = \frac{3}{2} \ln(2)$$

On voit qu'en changeant l'ordre des termes de la série semi-convergente $\sum_{k \geq 1} u_k$ on modifie la valeur de la somme....

En prenant la permutation τ qui place dans l'ordre les p premiers entiers impairs puis les q premiers entiers pairs on obtient

$$\sum_{k=1}^{+\infty} u_{\tau(k)} = \ln(2) + \frac{1}{2} \ln(p/q)$$

le cas étudié est $p = 2$ et $q = 1$.

Corollaire 1.23

Soit $\sum u_n$ une série à termes (éventuellement) complexes et $\sum v_n$ une série à termes réels positifs. On suppose que

- on a $u_n = O(v_n)$
- la série $\sum v_n$ converge (absolument car c'est une série positive)

Alors la série $\sum u_n$ est absolument convergente (donc convergente).

Démonstration : On a $|u_n| = O(v_n)$. Dès lors il existe $K \in \mathbf{R}_+$ tel que $|u_n| \leq K v_n$. Donc $\sum u_n$ est absolument convergente donc convergente. $\square$

Remarque : Si on a $u_n = o(v_n)$ cela marche encore car cela implique que $u_n = O(v_n)$.

★ **Méthode :** On utilise souvent ce critère avec les séries de Riemann. En particulier si la série $\sum u_n$ vérifie qu'il existe α strictement supérieur à 1 tel que $n^\alpha u_n$ est bornée (ce qui est vrai si elle a une limite finie) alors $u_n = O\left(\frac{1}{n^\alpha}\right)$ et donc $\sum u_n$ converge.

ATTENTION

Il faut être précis dans la rédaction quand on utilise ce genre de raisonnements (et on va l'utiliser très fréquemment). Pour démontrer qu'une série $\sum u_n$ converge en la comparant (par exemple) à la série de Riemann $\sum \frac{1}{n^2}$ il faut écrire :

- On a $u_n = O\left(\frac{1}{n^2}\right)$ et $\sum \frac{1}{n^2}$ est **absolument** convergente donc $\sum u_n$ est absolument convergente donc convergente.

Dans le cas où $\sum u_n$ est une série à termes positifs, on peut aussi invoquer le TCPSP :

- On a $u_n = O\left(\frac{1}{n^2}\right)$ et $\sum \frac{1}{n^2}$ est une série convergente donc par comparaison pour les séries à termes positifs la série $\sum u_n$ est convergente.

Par contre, le simple fait que $\sum u_n$ converge n'implique pas que $\sum u_n$ converge quand $u_n \sim v_n$, $u_n = O(v_n)$ ou même $u_n = o(v_n)$. Par exemple en prenant $u_n = \frac{(-1)^n}{n}$ et $v_n = \frac{1}{n \ln n}$.

Exemples :

1. On veut étudier la nature $\sum \frac{\cos(n\theta)}{n(n-1)}$. On a

$$\left| \frac{\cos(n\theta)}{n(n-1)} \right| \leq \frac{1}{n(n-1)} \sim \frac{1}{n^2}.$$

Comme $\sum \frac{1}{n^2}$ est absolument convergente, $\sum \frac{1}{n(n-1)}$ est aussi absolument convergente et donc $\sum \frac{\cos(n\theta)}{n(n-1)}$ est absolument convergente donc convergente.

2. Discuter en fonction de $x > 0$ la nature de $\sum n(-1)^n x^n$.

On voit que si $x \geq 1$ la série diverge grossièrement. Maintenant si $x < 1$, on a $|n(-1)^n x^n| = nx^n$ et si on pose $\alpha \in]x, 1[$ alors $nx^n = o(\alpha^n)$. La série $\sum \alpha^n$ est absolument convergente donc $\sum (-1)^n nx^n$ est absolument convergente donc convergente.

Exercice : Nature de $\sum \frac{n \cos n}{(\ln n)^n}$.

5 Comparaison séries - intégrales

5.1 Généralités

Théorème 1.24

Soit f une fonction continue monotone sur $\mathbf{R}_+$ et $0 \leq p < n$.

1. Si f est croissante on a

$$\int_p^n f(t)dt \leq \sum_{k=p+1}^n f(k) \leq \int_{p+1}^{n+1} f(t)dt.$$

2. Si f est décroissante on a

$$\int_{p+1}^{n+1} f(t)dt \leq \sum_{k=p+1}^n f(k) \leq \int_p^n f(t)dt.$$

Démonstration :

On suppose que la fonction est croissante. que

De ce fait pour $k \geq 0$ et $t \in [k, k+1]$ on a $f(k) \leq f(t) \leq f(k+1)$. En intégrant entre k et $k+1$ on obtient que

$$\underbrace{\int_k^{k+1} f(k)dt}_{=f(k)} \leq \int_k^{k+1} f(t)dt \leq \underbrace{\int_k^{k+1} f(k+1)dt}_{=f(k+1)}.$$

En sommant les termes de l'inégalité de gauche pour k allant de $p+1$ à n on obtient que

$$\sum_{k=p+1}^n f(k) \leq \sum_{k=p+1}^n \int_k^{k+1} f(t)dt = \int_{p+1}^{n+1} f(t)dt.$$

De même, en sommant les termes de l'inégalité de droite pour k allant de p à $n-1$ on obtient

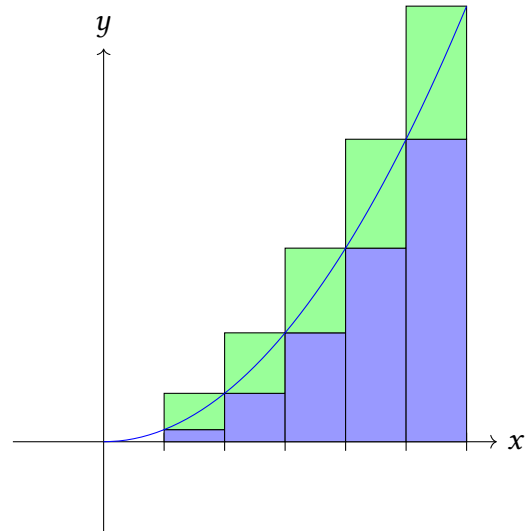
$$\int_p^n f(t)dt \leq \sum_{k=p+1}^n f(k) \leq \int_{p+1}^{n+1} f(t)dt.$$

Le cas où f est décroissante se déduit du précédent en l'appliquant à $-f$.

□

Remarques :

1. Si f est définie sur $\mathbf{R}$ en entier on peut même prendre $p = 0$.



2. La démonstration des séries de Riemann repose sur ces cas avec $f : t \mapsto \frac{1}{t^s}$ qui est décroissante.
3. Ces calculs sont juste les inégalités que l'on obtient en cherchant une valeur approchée d'une intégrale par la méthode des rectangles (pour un pas $h = 1$.)

5.2 Exemples et applications

Séries de Riemann

On étudie la série $\sum_{k \geq 1} \frac{1}{k^s}$.

Si $s \leq 0$, la série diverge grossièrement.

On suppose par la suite que $s > 0$. On sait que la fonction $t \mapsto t^s$ est décroissante sur $\mathbf{R}_+^*$. De ce fait pour tout $k \geq 1$,

$$\frac{1}{(k+1)^s} = \int_k^{k+1} \frac{dt}{(k+1)^s} \leq \int_k^{k+1} \frac{dt}{t^s} \leq \int_k^{k+1} \frac{dt}{k^s} = \frac{1}{k^s}.$$

On en déduit que, pour $k \geq 2$:

$$\int_k^{k+1} \frac{dt}{t^s} \leq \frac{1}{k^s} \leq \int_{k-1}^k \frac{dt}{t^s}.$$

On somme et on utilise la relation de Chasles :

$$\forall n \geq 2, \int_2^{n+1} \frac{dt}{t^s} = \sum_{k=2}^n \int_k^{k+1} \frac{dt}{t^s} \leq \sum_{k=2}^n \frac{1}{k^s} \leq \sum_{k=2}^n \int_{k-1}^k \frac{dt}{t^s} = \int_1^n \frac{dt}{t^s}.$$

— Si $s > 1$ on a donc

$$\sum_{k=2}^n \frac{1}{k^s} \leq \int_1^n \frac{dt}{t^s} = \frac{1}{1-s} [t^{1-s}]_1^n = \frac{1}{s-1} (1 - n^{1-s}) \leq \frac{1}{s-1}$$

On en déduit que la suite des sommes partielles est majorée. Comme elle est croissante, elle converge. La série est bien convergente.

— Si $s = 1$ on a cette fois,

$$\ln(n+1) - \ln(2) \leq \int_2^{n+1} \frac{dt}{t} \leq \sum_{k=2}^n \frac{1}{k}.$$

On en déduit que la suite des sommes partielles tend vers $+\infty$ et donc la série diverge.

— Si $s < 1$ on procède comme pour le cas $s = 1$.

Notons que on ne regarde les sommes partielles qu'à partir de $k = 2$ mais cela ne change rien à la convergence.

Note

Pour montrer la convergence, on majore les sommes partielles ; pour établir la divergence on minore les sommes partielles.

Un autre exemple

Si on cherche la nature de $\sum_{n \geq 2} \frac{1}{n \ln n}$. On procède de même. La suite des sommes partielles est croissante on cherche donc juste à majorer ou minorer. Comme $t \mapsto \frac{1}{t \ln t}$ est décroissante on est dans le deuxième cas :

$$\int_3^{n+1} \frac{dt}{t \ln t} \leq \sum_{k=3}^n \frac{1}{k \ln k} \leq \int_2^n \frac{dt}{t \ln t}.$$

Or $\int \frac{dt}{t \ln t} = \ln(\ln x)$. En prenant la minoration, on montre que la série diverge.

Exercice : Soit α et β . Étudier la nature des séries de Bertrand $\sum \frac{1}{n^\alpha (\ln n)^\beta}$.

Équivalent des sommes partielles (série divergente)

Si on reprend la série précédente, on a montré la divergence mais on peut même trouver un équivalent de la suite des sommes partielles. En effet

$$\ln(\ln(n)) \sim \ln(\ln(n+1)) - \ln(\ln(3)) = \int_3^{n+1} \frac{dt}{t \ln t} \leq \sum_{k=3}^n \frac{1}{k \ln k} \leq \int_2^n \frac{dt}{t \ln t} = \ln(\ln n) - \ln(\ln 2) \sim \ln(\ln n).$$

On en déduit que

$$\sum_{k=3}^n \frac{1}{k \ln k} \sim \ln(\ln n).$$

Équivalent du reste (série convergente)

Dans le cas d'une série convergente, ce genre de méthode permet aussi de trouver un équivalent du reste (qui tend vers 0). Si on reprend la série de Riemann pour $s > 1$. On a pour tout $n \geq p \geq 1$

$$\frac{1}{s-1} \left(\left(\frac{1}{p+1} \right)^{s-1} - \left(\frac{1}{n+1} \right)^{s-1} \right) = \int_{p+1}^{n+1} \frac{dt}{t^s} \leq \sum_{k=p+1}^n \frac{1}{k^s} \leq \int_p^{n+1} \frac{dt}{t^s} = \frac{1}{s-1} \left(\left(\frac{1}{p} \right)^{s-1} - \left(\frac{1}{n+1} \right)^{s-1} \right).$$

On sait que le terme central tend vers le reste R_p quand n tend vers $+\infty$. De plus, la limite

$$\frac{1}{s-1} \left(\frac{1}{p+1} \right)^{s-1} \leq R_p \leq \frac{1}{s-1} \left(\frac{1}{p} \right)^{s-1}.$$

On en déduit donc que

$$R_p \sim \frac{1}{s-1} \left(\frac{1}{p} \right)^{s-1}.$$

Développement asymptotique de la somme partielle de la série harmonique

Pour tout entier $n \geq 1$, on pose $H_n = \sum_{k=1}^n \frac{1}{k}$. En reprenant les calculs précédent, pour tout $k \geq 1$,

$$\frac{1}{k+1} \leq \int_k^{k+1} \frac{1}{t} dt \leq \frac{1}{k}$$

Donc

$$\sum_{k=1}^n \frac{1}{k+1} \leq \int_1^{n+1} \frac{1}{t} dt \leq \sum_{k=1}^n \frac{1}{k}$$

On obtient alors

$$\ln(n+1) \leq H_n \leq \ln(n+1) + 1 - \frac{1}{n+1}$$

Finalement $H_n \sim \ln(n+1) \sim \ln(n)$ cela peut s'écrire $H_n = \ln(n) + o(\ln n)$.

On peut être plus précis sur le $o(\ln(n))$. On considère donc

$$H_n - \ln(n+1) = \sum_{k=1}^n \frac{1}{k} - \sum_{k=1}^n \int_k^{k+1} \frac{dt}{t} = \sum_{k=1}^n \underbrace{\left(\frac{1}{k} - \int_k^{k+1} \frac{dt}{t} \right)}_{v_k}$$

Or l'inégalité ci-dessus, donne que

$$0 \leq v_k \leq \frac{1}{k} - \frac{1}{k+1} = \frac{1}{k(k+1)}$$

Cela montre que la série $\sum_{k \geq 1} v_k$ converge. Si on note γ sa somme, on obtient que $H_n - \ln(n+1) \xrightarrow{n \rightarrow +\infty} \gamma$.

En utilisant que $\ln(n+1) = \ln(n) + o(1)$ on obtient

$$H_n = \ln(n) + \gamma + o(1)$$

La constante γ s'appelle la constante d'Euler (Mascheroni). On a $\gamma \simeq 0,57721566$. Il est notable que l'on ne sait pas si γ est rationnel ou pas.

★ **Méthode** : Pour déterminer l'équivalent du reste d'une série convergente ou des sommes partielles d'une série convergente, on utilise souvent une comparaison série-intégrale.

Exercices :

1. Déterminer un équivalent de $\sum_{k=1}^n \frac{1}{k^s}$ pour $0 < s < 1$.
2. Montrer que la série de terme général $\sum \frac{1}{k \ln^2 k}$ converge et déterminer un équivalent du reste R_p .

6 Séries alternées

6.1 Généralités

La série $\sum \frac{(-1)^k}{k+1}$ (série harmonique alternée) n'est pas absolument convergente mais elle converge. Cela a été vu en 2.4

Définition 1.25

Soit $\sum u_n$ une série réelle. Elle est dite *alternée* si pour tout entier n , u_n et u_{n+1} sont de signes opposés c'est-à-dire $u_n u_{n+1} \leq 0$.

Théorème 1.26 (Critère spécial des séries alternées)

Soit $\sum u_n$ une série on suppose :

- Elle est alternée
- La suite $(|u_n|)$ décroît
- La suite (u_n) tend vers 0

La série $\sum u_n$ converge

Démonstration :

FAIRE UN DESSIN

On suppose que pour tout entier k , $u_{2k} \geq 0$ et $u_{2k+1} \leq 0$. Dans le cas contraire on remplace u_k par $-u_k$.

On considère $\alpha_n = S_{2n} = \sum_{k=0}^{2n} u_k$ et $\beta_n = S_{2n+1} = \sum_{k=0}^{2n+1} u_k$. Ce sont les suites extraites des termes de rang pairs et de rang impairs de la suite des sommes partielles

ATTENTION

On s'arrête à $2n$ ou à $2n+1$ mais dans les deux cas, k prend toutes les valeurs :

$$\alpha_2 = u_0 + u_1 + u_2 + u_3 + u_4 \text{ et } \beta_2 = u_0 + u_1 + u_2 + u_3 + u_4 + u_5$$

Nous allons montrer que les suites (α_n) et (β_n) sont adjacentes. En effet :

- On a $\alpha_{n+1} - \alpha_n = u_{2n+2} + u_{2n+1} = |u_{2n+2}| - |u_{2n+1}| \leq 0$ car $|u_n|$ décroît.
- On a $\beta_{n+1} - \beta_n = u_{2n+3} + u_{2n+2} = -|u_{2n+3}| + |u_{2n+2}| \geq 0$ car $|u_n|$ décroît.
- On a $\alpha_n - \beta_n = -u_{2n+1} \rightarrow 0$

Les deux suites (α_n) et (β_n) sont adjacentes. Elles convergent vers une même limite ℓ . On a donc que (S_{2n}) et (S_{2n+1}) convergent vers une même limite ℓ donc (S_n) converge vers ℓ et la série est convergente.

□

Exemples :

1. On retrouve simplement que la série harmonique alternée converge (mais sans la valeur de la somme).
2. Étudions $\sum u_n$ où $u_n = (-1)^n \frac{\ln n}{n}$.

Elle n'est pas absolument convergente car $|u_n| = \frac{\ln n}{n} \geq \frac{1}{n}$ est le terme général d'une série divergente. Par contre, elle est alternée et $(u_n) \rightarrow 0$ par croissances comparées. Pour finir, si on étudie $f : x \mapsto \frac{\ln(x)}{x}$ de dérivée $f' : x \mapsto \frac{1 - \ln x}{x^2}$ on voit que f décroît à partir de $x = e$ et donc $\left(\frac{\ln n}{n}\right)$ décroît à partir de $n = 3$.

6.2 Etude du reste et de la somme

Dans le cas où le critère spécial s'applique on obtient aussi des informations sur le reste.

Proposition 1.27

Soit $\sum u_n$ une série alternée qui vérifie les conditions du critère spécial,

- Le reste est du signe de son premier terme : R_n est du signe de u_{n+1}
- Le reste est en valeur absolue inférieur à son premier terme : $|R_n| \leq |u_{n+1}|$.

En particulier, la somme est du signe de u_0 et $\left| \sum_{n=0}^{+\infty} u_n \right| \leq |u_0|$.

Démonstration :

- Traitons pour commencer le cas $n = 0$.
 - Si on suppose que $u_0 \geq 0$: On a d'après la démonstration précédente que (S_{2n}) décroît et (S_{2n+1}) croît. On a donc

$$S_1 \leq S_3 \leq \sum_{k=0}^{+\infty} u_k = \lim_{n \rightarrow +\infty} S_n \leq S_2 \leq S_0.$$

Maintenant $S_1 = u_0 + u_1 \geq 0$ donc $S_1 \geq 0$ et de ce fait $\sum_{k=0}^{+\infty} u_k \geq 0$. Il est bien du signe de

u_0 . On a de plus $\sum_{k=0}^{+\infty} u_k \leq S_0 = u_0$.

- Pareil dans l'autre sens.
- Pour le cas général, il suffit de considérer la série $\sum_{k \geq 0} v_k$ où $v_k = u_{n+k+1}$. Il est clair que c'est encore une série alternée et que, par définition,

$$\sum_{k=0}^{+\infty} v_k = \sum_{k=n+1}^{+\infty} u_k = R_n.$$

De ce fait, R_n est du signe de $v_0 = u_{n+1}$ et $|R_n| \leq |u_{n+1}|$.

□

Exemple : Si on applique à la série harmonique alternée, on trouve que $R_{n-1} = \sum_{k=n}^{+\infty} \frac{(-1)^k}{k+1} = \ln 2 -$

$$\sum_{k=0}^{n-1} \frac{(-1)^k}{k+1} \leq |u_n| = \frac{1}{n+1}.$$

7 Sommation des relations de comparaison pour les séries à termes positifs

7.1 Sommation des relations de comparaison pour les séries à termes positifs

On a revu les trois relations de comparaisons : équivalence, négligeabilité et domination. On va voir que l'on peut « sommer » ces relations pour obtenir les mêmes relations sur les restes (séries convergentes) ou les sommes partielles (séries convergentes).

Théorème 1.28 (Cas des séries convergentes)

Soit $\sum v_n$ une série à termes positifs convergente. Soit $\sum u_n$ une autre série à coefficients éventuellement complexes. On note $R_n = \sum_{k=n+1}^{+\infty} u_k$ et $T_n = \sum_{k=n+1}^{+\infty} v_k$

1. Si $u_n = o(v_n)$ alors $R_n = o(T_n)$
2. Si $u_n = O(v_n)$ alors $R_n = O(T_n)$
3. Si $u_n \sim v_n$ alors $R_n \sim T_n$

Remarques :

1. Du fait que $\sum v_n$ soit, une série à termes positifs convergente, les hypothèses $u_n \sim v_n$; $u_n = o(v_n)$ ou $u_n = O(v_n)$ assure que la série $\sum |u_n|$ est aussi convergente et donc $\sum u_n$ aussi. On peut donc parler de la suite des restes de la série $\sum u_n$.
2. Le résultat portant sur les restes, il est clair que cela s'applique aussi si $\sum v_n$ est positive seulement à partir d'un certain rang.

Démonstration : On suppose que $\sum v_n$ converge. On remarque que dans les trois cas, les séries $\sum |u_n|$ et $\sum u_n$ sont convergentes. On peut donc considérer leurs restes.

1. On suppose que $u_n = o(v_n)$. Cela signifie que pour tout $\varepsilon > 0$, il existe un entier N tel que pour tout entier $n \geq N$, $|u_n| \leq \varepsilon v_n$.

En particulier, pour $n \geq N$,

$$|R_n| = \left| \sum_{k=n+1}^{+\infty} u_k \right| \leq \sum_{k=n+1}^{+\infty} |u_k| \leq \varepsilon \sum_{k=n+1}^{+\infty} v_k = \varepsilon T_n$$

Cela montre que $R_n = o(T_n)$.

2. On suppose que $u_n = O(v_n)$. Cela signifie qu'il existe $M \in \mathbf{R}_+$ tel que pour tout entier n , $|u_n| \leq M v_n$.

En particulier, pour tout entier n ,

$$|R_n| = \left| \sum_{k=n+1}^{+\infty} u_k \right| \leq \sum_{k=n+1}^{+\infty} |u_k| \leq M \sum_{k=n+1}^{+\infty} v_k = M T_n$$

Cela montre que $R_n = O(T_n)$.

3. On suppose que $u_n \sim v_n$. Cela implique que $u_n = v_n + o(v_n)$ ou que $u_n - v_n = o(v_n)$. On peut donc appliquer ce qui a été vu en 1. On obtient que $R_n - T_n = o(T_n)$ et donc $R_n \sim T_n$.

Exemples :

1. Cela permet de retrouver les équivalents des restes des séries de Riemann déjà vu par comparaison série-intégrale. En effet, si $s > 1$ (série de Riemann convergente), on a

$$\int_n^{n+1} \frac{dt}{t^s} = \left[\frac{1}{1-s} t^{-s+1} \right]_n^{n+1} = \frac{1}{1-s} \left(\frac{1}{(n+1)^{s-1}} - \frac{1}{n^{s-1}} \right) = \frac{1}{1-s} \frac{1}{n^{s-1}} ((1 + 1/n)^{1-s} - 1)$$

Et donc

$$\int_n^{n+1} \frac{dt}{t^s} \sim \frac{1}{1-s} \frac{1}{n^{s-1}} \frac{1-s}{n} = \frac{1}{n^s}.$$

On en déduit que le reste $R_n = \sum_{k=n+1}^{+\infty} \frac{1}{k^s}$ est équivalent au reste $R'_n = \sum_{k=n+1}^{+\infty} \int_n^{n+1} \frac{dt}{t^s}$. Or

$$R'_n = \lim_{n' \rightarrow +\infty} \int_n^{n'} \frac{dt}{t^s} = \frac{1}{s-1} n^{1-s}$$

2. On se donne $u_n = \sqrt[3]{n^3 + 1} - \sqrt[4]{n^4 + 2n}$. On étudie $\sum u_n$

On a

$$u_n = n \left((1 + 1/n^3)^{1/3} - (1 + 2/n^3)^{1/4} \right) = n \left(1 + \frac{1}{3n^3} - 1 - \frac{2}{4n^3} + o(1/n^3) \right) = -\frac{1}{6n^2} + o\left(\frac{1}{n^2}\right)$$

On a donc $u_n \sim -\frac{1}{6n^2}$. De ce fait on voit que la série est à terme négatif à partir d'un certain rang. Comme la série de terme général $\sum \frac{1}{n^2}$ converge (série de Riemann), la série $\sum u_n$ converge. De plus, si on note R_n le reste d'ordre n de la série. On a alors en utilisant les calculs faits ci-dessus :

$$R_n \sim -\sum_{k=n+1}^{+\infty} \frac{1}{6k^2} \sim -\frac{1}{6n}.$$

Théorème 1.29 (Cas des séries divergentes)

Soit $\sum v_n$ une série à terme positifs divergente. On note $U_n = \sum_{k=0}^n u_k$ et $V_n = \sum_{k=0}^n v_k$ les sommes partielles

1. Si $u_n = o(v_n)$ alors $U_n = o(V_n)$
2. Si $u_n = O(v_n)$ alors $U_n = O(V_n)$
3. Si $u_n \sim v_n$ alors $U_n \sim V_n$

Remarques :

1. Dans le cas où $u_n \sim v_n$, la série $\sum u_n$ est aussi divergente (et à termes positifs). Dans les deux autres cas si $\sum u_n$ est convergente le résultat est évident.
2. Là encore, on ne change rien si v_n est seulement positive à partir d'un certain rang. En effet, en faisant commencer les sommes partielles à partir d'un rang n_0 on les modifie d'une valeur constante, ce qui ne modifie rien car $\lim_{n \rightarrow +\infty} V_n = +\infty$.

Démonstration : Les démonstrations sont similaires à celles du théorème précédent. Il faut quand même faire un peu attention. En effet, dans le cas convergent, on s'intéressait aux restes. De ce fait, on pouvait faire la somme que pour les indices k où la propriété donnée par la relation de comparaison était vraie. Maintenant, comme on regarde des sommes partielles, il va falloir traiter aussi les « premiers » termes des séries là où on n'a pas d'informations données par la relation de comparaison.

On suppose que $\sum v_n$ diverge.

1. On suppose que $u_n = o(v_n)$. Cela signifie que pour tout $\varepsilon > 0$, il existe un entier N tel que pour tout entier $n \geq N$, $|u_n| \leq \varepsilon v_n$.

En particulier, pour $n \geq N$,

$$\begin{aligned}
 |U_n| = \left| \sum_{k=0}^n u_k \right| &\leq \left| \sum_{k=0}^N u_k \right| + \left| \sum_{k=N+1}^n u_k \right| \\
 &\leq A + \sum_{k=N+1}^n |u_k| \quad \text{où } A = \left| \sum_{k=0}^N u_k \right| \text{ ne dépend pas de } n \\
 &\leq A + \varepsilon \sum_{k=N+1}^n v_k \\
 &\leq A + \varepsilon V_n \quad \text{car } \sum v_k \text{ à termes positifs}
 \end{aligned}$$

Or, par hypothèse, $\lim_{n \rightarrow +\infty} V_n = +\infty$. Il existe donc N' (que l'on peut choisir supérieur à N) tel que pour $n \geq N'$, $V_n \geq \frac{A}{\varepsilon}$.

Pour $n \geq N' \geq N$ on a donc : $|U_n| \leq 2\varepsilon V_n$.

Cela montre que $U_n = o(V_n)$.

2. On suppose que $u_n = O(v_n)$. Cela signifie qu'il existe $M \in \mathbf{R}_+$ tel que pour tout entier n , $|u_n| \leq Mv_n$.

En particulier, pour tout entier n ,

$$|U_n| = \left| \sum_{k=0}^n u_k \right| \leq \sum_{k=0}^n |u_k| \leq M \sum_{k=0}^n v_k = MV_n$$

Cela montre que $U_n = O(V_n)$.

3. On suppose que $u_n \sim v_n$. Cela implique que $u_n = v_n + o(v_n)$ ou que $u_n - v_n = o(v_n)$. On peut donc appliquer ce qui a été vu en 1. On obtient que $U_n - V_n = o(V_n)$ et donc $U_n \sim V_n$.

□

Exemple : En reprenant les calculs similaires à ci-dessus pour la série harmonique. On a

$$u_n = \int_n^{n+1} \frac{dt}{t} = [\ln t]_n^{n+1} = \ln(1 + 1/n)$$

On a donc

$$\frac{1}{n} \sim u_n.$$

On en déduit que $H_n = \sum_{k=1}^n \frac{1}{k}$ est équivalente à $\sum_{k=1}^n \int_k^{k+1} \frac{dt}{t} = \int_1^{n+1} \frac{dt}{t} = \ln(n+1)$. Donc $H_n \sim \ln(n)$.

7.2 Applications classiques

Méthode de Césaro

Théorème 1.30 (Théorème de Césaro)

Soit (u_n) une suite (éventuellement complexe). Pour tout $n \in \mathbb{N}$, on pose

$$c_n = \frac{1}{n+1} \sum_{k=0}^n u_k$$

1. Si $(u_n)_{n \geq 0}$ tend vers $\ell \in \mathbb{R}$ alors $(c_n)_{n \geq 0}$ aussi.
2. Si (u_n) est à valeurs réelles et tend vers $\ell \in \{\pm\infty\}$ alors $(c_n)_{n \geq 0}$ aussi.

Démonstration :

1. On peut poser $u_n = \ell + \underbrace{u_n - \ell}_{v_n}$ où $v_n = o(1)$.

On a alors

$$c_n = \frac{1}{n+1} \sum_{k=0}^n u_k = \ell + \frac{1}{n+1} \sum_{k=0}^n v_k.$$

On utilise alors la sommation des relations de comparaison dans le cas divergent car la série $\sum_{n \geq 1} 1$ diverge. On en déduit que

$$\sum_{k=0}^n v_k = o\left(\sum_{k=0}^n 1\right) = o(n+1) \quad \text{et donc} \quad \frac{1}{n+1} \sum_{k=0}^n v_k = o(1)$$

Cela revient à dire $\frac{1}{n+1} \sum_{k=0}^n v_k \rightarrow 0$ et donc

$$\lim_{n \rightarrow +\infty} c_n = \ell$$

2. Quitte à remplacer u_n par $-u_n$ on peut supposer que $\ell = +\infty$. En particulier, la suite $(u_n)_{n \geq 0}$ est strictement positive à partir d'un certain rang, $\sum u_n$ diverge et $1 = o(u_n)$.

On en déduit par sommation des relations de comparaison que $n+1 = o\left(\sum_{k=0}^n u_k\right)$ et donc $1 = o(c_n)$ ce qui signifie que $(c_n)_{n \geq 0}$ tend vers $+\infty$ (car elle est positive à partir d'un certain rang).

□

Etude de suite récurrente

On veut étudier la série définie par

$$u_0 \in]0, 1[, \forall n \in \mathbb{N}, u_{n+1} = u_n - u_n^2 = f(u_n)$$

où $f : x \mapsto x - x^2 = x(1 - x)$. Rappelons rapidement l'étude de ce genre de suite.

- On voit que si $x \in [0, 1]$ alors $f(x) \in [0, 1]$ (l'intervalle $[0, 1]$ est stable par f). De ce fait, par une récurrence immédiate, pour tout entier n , $u_n \in [0, 1]$.

- Pour tout entier n , $u_{n+1} - u_n = -u_n^2 \leq 0$. On en déduit que la suite est décroissante.
- La suite (u_n) est décroissante et minorée (par 0) donc elle converge.
- On cherche les points fixes de f . Pour $x \in [0, 1]$, $f(x) = x \iff -x^2 = 0 \iff x = 0$.
- Soit ℓ la limite de (u_n) . On sait que f est continue sur $[0, 1]$ donc $(f(u_n))$ converge vers $\lim_{x \rightarrow \ell} f(x) = f(\ell)$. Or $(f(u_n)) = (u_{n+1})$ est une suite extraite de (u_n) donc elle tend vers ℓ . De ce fait, $\ell = f(\ell)$, ce qui implique que $\ell = 0$.

Exercice : Justifier que la suite définie par $u_0 \in [0, \pi]$ et $\forall n \in \mathbb{N}, u_{n+1} = \sin(u_n)$ tend vers 0.

On cherche maintenant un équivalent de u_n . Pour cela on va chercher α tel que

$$\frac{1}{u_{n+1}^\alpha} - \frac{1}{u_n^\alpha}$$

ait une limite finie non nulle. On en déduira ensuite un équivalent de $\frac{1}{u_n^\alpha}$.

On a ici,

$$\frac{1}{u_{n+1}^\alpha} - \frac{1}{u_n^\alpha} = \frac{1}{(u_n - u_n^2)^\alpha} - \frac{1}{u_n^\alpha} = \frac{1}{u_n^\alpha} ((1 - u_n)^{-\alpha} - 1) \sim \frac{\alpha u_n}{u_n^\alpha} = \alpha u_n^{1-\alpha}$$

On prend donc $\alpha = 1$.

Exercice : En faire de même avec la suite de l'exercice précédent.

On a donc $\frac{1}{u_{n+1}} - \frac{1}{u_n} \sim 1$ et donc par sommation des équivalents,

$$\frac{1}{u_n} - \frac{1}{u_0} = \sum_{k=0}^{n-1} \left(\frac{1}{u_{k+1}} - \frac{1}{u_k} \right) \sim \sum_{k=0}^{n-1} 1 = n$$

On en déduit que $\frac{1}{u_n} \sim n$ et donc $u_n \sim \frac{1}{n}$.

Exercice : En faire de même avec la suite de l'exercice précédent. On trouvera $u_n \sim \sqrt{3/n}$.

On peut alors continuer pour essayer d'obtenir un développement asymptotique.

En effet, si on reprend le calcul précédent,

$$\frac{1}{u_{n+1}} - \frac{1}{u_n} = \frac{1}{u_n} ((1 - u_n)^{-1} - 1) = 1 + u_n + o(u_n)$$

Si bien que

$$\frac{1}{u_n} - \frac{1}{u_0} = n + \sum_{k=0}^{n-1} (u_k + o(u_k))$$

Or $u_k + o(u_k) \sim u_k \sim \frac{1}{k}$. On en déduit que

$$\sum_{k=0}^{n-1} (u_k + o(u_k)) \sim \sum_{k=0}^{n-1} \frac{1}{k+1} \sim \ln(n)$$

Finalement,

$$\frac{1}{u_n} = n + \ln n + o(\ln n)$$

et donc

$$u_n = \frac{1}{n} - \frac{\ln n}{n^2} + o\left(\frac{\ln n}{n^2}\right).$$

Exercice : En faire de même avec la suite de l'exercice précédent. On trouvera $u_n = \sqrt{\frac{n}{3}} - \frac{6 \ln n}{5 n^{3/2}} + o\left(\frac{\ln n}{n^{3/2}}\right)$.

Algèbre linéaire et éléments propres

1	Rappels	40
1.1	Définitions	41
1.2	Familles de vecteurs	41
1.3	Rang d'une application linéaire	44
2	Rappels sur les matrices	46
2.1	Changements de bases	46
2.2	Matrices semblables	47
2.3	Trace	48
3	Formes linéaires et hyperplans	48
3.1	Changement de bases	49
3.2	Bases duales et formes coordonnées	50
3.3	Formes linéaires et hyperplans	51
4	Compléments en algèbre linéaire	52
4.1	Somme	52
4.2	Parties stables	55
4.3	Calcul par blocs	57
4.4	Calculs par blocs et déterminants	58
5	Éléments propres d'un endomorphisme	61
5.1	Définition	62
5.2	Sous-espaces propres	64
6	Éléments propres d'une matrice	66
6.1	Définition	66
6.2	Endomorphismes et matrices diagonalisables	69

Dans ce chapitre, K désigne $\mathbb{R}$ ou $\mathbb{C}$. En pratique tout restera vrai si K est un corps en général.

1 Rappels

1.1 Définitions

Définition 2.1 (Espace vectoriel)

On appelle *espace vectoriel sur $\mathbf{K}$* ou *$\mathbf{K}$ -espace vectoriel* un triplet $(E, +, \cdot)$ où E est un ensemble, $+$ une loi interne de E et $\cdot$ une loi externe :

$$\begin{aligned} + : E \times E &\rightarrow E & \text{et} & \quad \cdot : \mathbf{K} \times E \rightarrow E \\ (u, v) &\mapsto u + v & & (\lambda, u) \mapsto \lambda \cdot u \end{aligned}$$

Vérifiant

- $(E, +)$ est un groupe abélien
- La loi $\cdot$ vérifie
 - $\forall (\lambda, \mu) \in \mathbf{K}^2, \forall u \in E, (\lambda + \mu) \cdot u = \lambda \cdot u + \mu \cdot u$
 - $\forall (\lambda, \mu) \in \mathbf{K}^2, \forall u \in E, (\lambda \times \mu) \cdot u = \lambda \cdot (\mu \cdot u)$
 - $\forall u \in E, 1 \cdot u = u$

Remarque : Rappelons que le fait que $(E, +)$ soit un groupe abélien signifie que $+$ est associative, commutative, qu'il existe un élément neutre noté 0_E et que pour tout u de E il existe un élément de E noté $-u$ qui s'appelle l'opposé qui vérifie que $u + (-u) = 0_E$.

Définition 2.2 (Sous-espaces vectoriels)

Soit E un $\mathbf{K}$ -espace vectoriel, on appelle *sous-espace vectoriel* de E une partie $F \subset E$ telle que les lois $+$ et $\cdot$ induisent une structure d'espace vectoriel sur F .

Remarque : Soit $F \subset E$. C'est un sous-espace vectoriel si et seulement si $F \neq \emptyset$ et pour tout $(u, v) \in F^2$ et $(\lambda, \mu) \in \mathbf{K}^2$, $\lambda u + \mu v \in F$.

1.2 Familles de vecteurs

Dans tout ce paragraphe, E désigne un $\mathbf{K}$ -espace vectoriel.

Définition 2.3 (Combinaison linéaire)

Soit $(u_i)_{i \in I}$ une famille de vecteurs. On appelle *combinaison linéaire* des $(u_i)_{i \in I}$ tout vecteur de la forme

$$\sum_{i \in I} \lambda_i u_i$$

où $(\lambda_i)_{i \in I}$ est une famille presque nulle de scalaires ; c'est-à-dire qu'il existe un ensemble fini $J \subset I$ tel que si $i \in I \setminus J$, $\lambda_i = 0$.

Notation : On note $\mathbf{K}^{(I)}$ l'ensemble des familles de scalaires presque nulle.

Remarque : Dans le cas où I est un ensemble fini (par exemple $I = \llbracket 1 ; n \rrbracket$), la condition « presque nulle » est toujours vérifiée et on note juste

$$\sum_{i=1}^n \lambda_i u_i.$$

Notation : Soit $(u_i)_{i \in I}$ une famille de vecteurs.

On note $\text{Vect}((u_i)_{i \in I})$ l'ensemble des combinaisons linéaires. C'est le plus petit espace vectoriel qui contient tous les éléments de la famille. Il s'appelle l'espace vectoriel engendré par la famille.

Exemple : Dans l'espace vectoriel des suites $\mathbb{R}^{\mathbb{N}}$. On pose $u(i)$ la suite définie par

$$\forall n \in \mathbb{N}, u(i)_n = \delta_{i,n} = \begin{cases} 1 & \text{si } i = n \\ 0 & \text{sinon} \end{cases}$$

On a alors $\text{Vect}(u(i))_{i \in \mathbb{N}}$ qui est l'ensemble des suites à support fini c'est-à-dire stationnaires à 0. Cela correspond aux polynômes en notant X^i pour $u(i)$. On voit par exemple que la suite constante égale à 1 n'est pas dans $\text{Vect}(u(i))_{i \in \mathbb{N}}$.

Définition 2.4 (Famille libre)

Soit $(u_i)_{i \in I}$ une famille de vecteurs. Elle est libre si la famille nulle est la seule famille de scalaire presque nulle $(\lambda_i)_{i \in I}$ telle que $\sum_{i \in I} \lambda_i u_i = 0$.

C'est-à-dire :

$$((u_i)_{i \in I}) \text{ libre} \iff (\forall (\lambda_i)_i \in \mathbf{K}^{(I)}, \sum_{i \in I} \lambda_i u_i = 0 \Rightarrow \forall i \in I, \lambda_i = 0)$$

Remarque : Cela revient à demander que toute sous-famille finie est libre.

Exemples :

1. Pour tout $k \in \mathbb{N}$, on pose $u_k = (t \mapsto \sin(t^k))$. La famille $(u_k)_{k \in \mathbb{N}}$ est libre. En effet supposons par l'absurde qu'il existe une famille $(\lambda_k) \in \mathbb{R}^{(\mathbb{N})}$ qui ne soit pas nulle telle que

$$\sum_{k \in \mathbb{N}} \lambda_k u_k = 0.$$

Si on note $p = \min(k \in \mathbb{N} \mid \lambda_k \neq 0)$ le plus petit indice tel que $\lambda_k \neq 0$. On a

$$\sum_{k \in \mathbb{N}} \lambda_k u_k \underset{0}{\sim} \lambda_p t^p$$

Donc $\lambda_p = 0$ ce qui est absurde. La famille est bien libre

2. Soit $(P_i)_{i \in I}$ une famille de polynômes tels que $\forall (i, j) \in I^2, i \neq j \Rightarrow \deg(P_i) \neq \deg(P_j)$. Montrons que la famille est libre. Supposons par l'absurde qu'il existe une combinaison linéaire non triviale $\sum_{i \in I} \lambda_i P_i = 0$. Soit i_0 l'indice tel que P_{i_0} soit de degré maximal parmi les polynômes P_i où $\lambda_i \neq 0$ (les polynômes qui apparaissent réellement dans la combinaison linéaire). On a alors

$$\lambda_{i_0} P_{i_0} = - \sum_{i \in I \setminus \{i_0\}} \lambda_i P_i$$

Le polynôme de droite est de degré strictement inférieur à celui de gauche. C'est absurde.

La famille est donc libre.

Exercice : On pose $f_a : x \mapsto e^{ax}$. Montrer que $(f_a)_{a \in \mathbb{R}}$ est libre.

Définition 2.5 (Famille Génératrice)

Soit $(u_i)_{i \in I}$ une famille de vecteurs. On dit qu'elle engendre E si $E = \text{Vect}((u_i)_{i \in I})$ c'est-à-dire que tous les vecteurs de E sont des combinaisons linéaires de vecteurs de la famille.

ATTENTION

La notion de famille libre est intrinsèque à la famille, c'est-à-dire que si $(u_i)_{i \in I}$ est une famille d'un sous-espace vectoriel F , on peut travailler dans E ou dans F pour montrer que la famille est libre. Par contre, la notion de famille génératrice ne l'est pas. Une famille peut engendrer F sans engendrer E .

Exemple : La famille $(X^i)_{i \in \mathbb{N}}$ engendre $\mathbb{K}[X]$.

Définition 2.6

Une famille de vecteurs $(u_i)_{i \in I}$ est une base de E si c'est une famille libre et génératrice de E .

Théorème 2.7

Soit $(u_i)_{i \in I}$ est une base de E . Pour tout vecteur w il existe une unique famille presque nulle $(\lambda_i)_{i \in I} \in \mathbb{K}^{(I)}$ telle que

$$w = \sum_{i \in I} \lambda_i u_i$$

Remarque : Le fait que la famille est génératrice implique qu'il existe (au moins une) famille presque nulle $(\lambda_i)_{i \in I} \in \mathbb{K}^{(I)}$ telle que $w = \sum_{i \in I} \lambda_i u_i$. Le fait qu'elle soit libre implique qu'il existe au plus une famille presque nulle $(\lambda_i)_{i \in I} \in \mathbb{K}^{(I)}$ telle que $w = \sum_{i \in I} \lambda_i u_i$. En effet s'il y en a deux distinctes en faisant la différence on trouve une famille encore presque nulle $(\mu_i)_{i \in I}$ telle que $\sum_{i \in I} \mu_i u_i = 0$.

Définition 2.8

Soit $\mathcal{B} = (u_i)_{i \in I}$ une base de E . Pour tout vecteur w , l'unique famille presque nulle $(\lambda_i) \in \mathbb{K}^{(I)}$ telle que

$$\sum_{i \in I} \lambda_i u_i = w$$

s'appelle les coordonnées de w dans la base $\mathcal{B}$.

Remarque : Il a été vu en première année que tout espace vectoriel de dimension finie admettait une base. Il est possible de démontrer que tout espace vectoriel admet des bases (hors programme) cependant dans certains cas on ne peut pas en exhiber. Par exemple $\mathbb{R}^{\mathbb{R}}$ ou $\mathbb{R}^{\mathbb{N}}$.

Exercice : Déterminer une base de $\mathbb{C}(X)$. On pourra penser à la décomposition en éléments simples.

Définition 2.9 (Rang)

Soit $\mathcal{F} = (u_i)_{i \in I}$. On appelle rang de la famille $\mathcal{F}$ et on note $\text{rg}(\mathcal{F})$ la dimension de l'espace vectoriel engendré par la famille quand il est de dimension finie.

Proposition 2.10

Soit $\mathcal{F}$ une famille de vecteurs de E .

1. Si E est de dimension finie n . Alors $\text{rg}(\mathcal{F}) \leq n$.
2. Si $\mathcal{F}$ est de cardinal n . Alors $\text{rg}(\mathcal{F}) \leq n$.

ATTENTION

Ne pas mélanger les notions de rang, cardinal et dimension :

- On peut parler du rang ou du cardinal d'une famille mais pas de sa dimension.
- On peut parler de la dimension (et du cardinal) d'un sous-espace vectoriel mais pas de son rang.

Proposition 2.11

Soit $\mathcal{F}$ une famille de vecteurs de E

1. On a $(\text{rg}(\mathcal{F}) = \dim(E)) \iff (\text{La famille } \mathcal{F} \text{ engendre } E)$.
2. On a $(\text{rg}(\mathcal{F}) = \#\mathcal{F}) \iff (\text{La famille } \mathcal{F} \text{ est libre})$.
3. On a $(\text{rg}(\mathcal{F}) = \#\mathcal{F} = \dim E) \iff (\text{La famille } \mathcal{F} \text{ est une base de } E)$.

1.3 Rang d'une application linéaire

Définition 2.12

Soit u une application linéaire. On appelle rang de u et on note $\text{rg}(u)$ la dimension de $\text{Im}(u)$ quand elle est finie. On a donc

$$\text{rg}(u) = \dim(\text{Im}(u)).$$

Proposition 2.13

Soit u une application linéaire de E dans F et $\mathcal{F}$ une famille génératrice de E .

1. La famille image $u(\mathcal{F})$ engendre $\text{Im}(u)$.
2. On a donc $\text{rg}(u) = \text{rg}(u(\mathcal{F}))$ (dans le cas où ce sont des nombres finis).

Remarque : On utilisera ce résultat essentiellement avec $\mathcal{F}$ une base de E .

Théorème 2.14

Soit u une application linéaire de E dans F et S un supplémentaire de $\text{Ker}(u)$ dans E .

1. La restriction de u à S est injective : $\text{Ker}(u|_S) = \{0\}$
2. L'application u induit un isomorphisme de S sur $\text{Im}(u)$:

$$\begin{aligned}\tilde{u} : S &\rightarrow \text{Im}(u) \\ x &\mapsto u(x)\end{aligned}$$

Corollaire 2.15 (Théorème du rang)

Soit u une application linéaire de E dans F . On suppose que E est un espace vectoriel de dimension finie.

1. Le noyau $\text{Ker}(u)$ admet un supplémentaire S
2. On a

$$\dim S = \dim(\text{Im}(u)) \text{ et donc } \dim(\text{Ker}(u)) + \text{rg}(u) = \dim E$$

Proposition 2.16

Soit $u : E \rightarrow F$ une application linéaire entre deux espaces vectoriels de dimension finie.

1. On a $(\text{rg}(u) = \dim F) \iff (u \text{ est surjective})$
2. On a $(\text{rg}(u) = \dim E) \iff (u \text{ est injective})$
3. On a $(\text{rg}(u) = \dim E = \dim F) \iff (u \text{ est un isomorphisme}).$

Proposition 2.17

Soit $u : F \rightarrow G$ et $v : E \rightarrow F$ deux applications linéaires.

1. On a $\text{rg}(u \circ v) \leq \text{rg}(u)$ et $\text{rg}(u \circ v) \leq \text{rg}(v)$.
2. Si u est un isomorphisme alors $\text{rg}(u \circ v) = \text{rg}(v)$. Si v est un isomorphisme alors $\text{rg}(u \circ v) = \text{rg}(u)$.

Démonstration :

1. On sait que $(u \circ v)(E) \subset u(F)$ donc, en prenant les dimensions, $\text{rg}(u \circ v) \leq \text{rg}(u)$.

De même soit $(e_1, \dots, e_p)$ une base de $\text{Im} v$ (où $p = \text{rg} v$). Alors $(u(e_1), \dots, u(e_p))$ engendre $\text{Im}(u \circ v)$ donc

$$\text{rg}(u \circ v) = \dim(\text{Im}(u \circ v)) \leq p = \text{rg} v.$$

2. On suppose que u est un isomorphisme. On sait que $\text{rg}(u \circ v) \leq \text{rg}(v)$. De plus $v = u^{-1} \circ (u \circ v)$ donc, $\text{rg}(v) = \text{rg}(u^{-1} \circ (u \circ v)) \leq \text{rg}(u \circ v)$. En combinant les deux inégalités, on obtient $\text{rg}(u \circ v) = \text{rg} v$.

On suppose que v est un isomorphisme. On sait que $\text{rg}(u \circ v) \leq \text{rg}(u)$. De plus $u = (u \circ v) \circ v^{-1}$ donc, $\text{rg}(u) = \text{rg}((u \circ v) \circ v^{-1}) \leq \text{rg}(u \circ v)$. En combinant les deux inégalités, on obtient $\text{rg}(u \circ v) = \text{rg}(u)$.

□

2 Rappels sur les matrices

Faisons quelques rappels sur les matrices

2.1 Changements de bases

Soit E un espace vectoriel de dimension finie. Soit $\mathcal{B} = (e_1, \dots, e_n)$ et $\mathcal{B}' = (e'_1, \dots, e'_n)$ deux bases.

Définition 2.18 (Matrice de changement de bases)

On appelle matrice de changement de base de $\mathcal{B}$ à $\mathcal{B}'$ la matrice des coordonnées des vecteurs de $\mathcal{B}'$ dans la base $\mathcal{B}$.

On la note souvent $P_{\mathcal{B}}^{\mathcal{B}'}$. On a donc

$$P_{\mathcal{B}}^{\mathcal{B}'} = \text{Mat}_{\mathcal{B}}(e'_1, \dots, e'_n)$$

Théorème 2.19 (Changement de bases pour les vecteurs)

Soit w un vecteur de E . On note $X = \text{Mat}_{\mathcal{B}}(w)$ et $X' = \text{Mat}_{\mathcal{B}'}(w)$ ses coordonnées dans les bases $\mathcal{B}$ et $\mathcal{B}'$.

On a

$$X = P_{\mathcal{B}}^{\mathcal{B}'} X'$$

ATTENTION

La matrice de changement de base de $\mathcal{B}$ à $\mathcal{B}'$ permet de calculer simplement les coordonnées d'un vecteur dans la base $\mathcal{B}$ à partir de celles du vecteur dans la base $\mathcal{B}'$.

On se donne de plus un autre espace vectoriel F avec des bases $\mathcal{F}$ et $\mathcal{F}'$.

Théorème 2.20 (Changement de bases pour les applications linéaires)

Soit $u \in \mathcal{L}(E, F)$. On note $A = \text{Mat}_{\mathcal{B}, \mathcal{F}}(u)$ et $A' = \text{Mat}_{\mathcal{B}', \mathcal{F}'}(u)$. On a alors

$$A' = Q^{-1}AP$$

où P est la matrice du changement de bases de $\mathcal{B}$ à $\mathcal{B}'$ et Q celle de $\mathcal{F}$ à $\mathcal{F}'$.

Démonstration : On peut illustrer cette formule par le diagramme suivant.

On en déduit que $A'X' = Q^{-1}APX'$ pour toute matrice colonne X' . De ce fait, $A' = Q^{-1}AP$.

□

Dans le cas où on regarde des endomorphismes et non plus des applications linéaires générales, on utilise quasi-systématiquement la même base pour l'espace de départ et d'arrivée (afin que la composition des endomorphismes corresponde au produit des matrices). On note alors $\text{Mat}_{\mathcal{B}}(u)$ pour $\text{Mat}_{\mathcal{B}, \mathcal{B}}(u)$.

Corollaire 2.21 (Cas des endomorphismes)

Soit $u \in \mathcal{L}(E)$. On note $A = \text{Mat}_{\mathcal{B}}(u)$ et $A' = \text{Mat}_{\mathcal{B}'}(u)$. On a alors

$$A' = P^{-1}AP$$

où P est la matrice du changement de bases de $\mathcal{B}$ à $\mathcal{B}'$.

2.2 Matrices semblables

Définition 2.22 (Matrices semblables)

Soit A et A' deux matrices de $\mathcal{M}_n(\mathbf{K})$. On dit que A et A' sont semblables s'il existe $P \in \text{GL}_n(\mathbf{K})$ telle que

$$A' = P^{-1}AP.$$

ATTENTION

Ne pas confondre semblables et équivalentes. Deux matrices sont équivalentes si elles ont le même rang, ce qui revient au fait qu'il existe $(P, Q) \in \text{GL}_n(\mathbf{K})^2$ tel que

$$A' = PAQ$$

Proposition 2.23

La relation « est semblable » est une relation d'équivalence.

Démonstration : Il suffit de démontrer qu'elle est réflexive, symétrique et transitive. $\square$

Remarque : Soit A une matrice, l'ensemble des matrices semblables à A forment ce que l'on appelle sa classe de similitude.

Proposition 2.24

Soit E un espace vectoriel de dimension n et $\mathcal{B}$ une base de E . Soit $u \in \mathcal{L}(E)$, on note $A = \text{Mat}_{\mathcal{B}}(u)$.

Soit $A' \in \mathcal{M}_n(\mathbf{K})$, elle est semblable à A si et seulement s'il existe une base $\mathcal{B}'$ de E telle que $A' = \text{Mat}_{\mathcal{B}'}(u)$.

2.3 Trace

Définition 2.25 (Trace)

Soit $A = (a_{ij}) \in \mathcal{M}_n(\mathbf{K})$, on appelle trace de A et on note $\text{tr}(A)$ la somme des coefficients diagonaux.

$$\text{tr}(A) = \sum_{i=1}^n a_{ii}$$

Proposition 2.26 (Propriétés de la trace)

1. La trace est une forme linéaire.
2. Soit $A \in \mathcal{M}_n(\mathbf{K})$, $\text{tr}(A^\top) = \text{tr}(A)$
3. Soit $(A, B) \in \mathcal{M}_n(\mathbf{K})^2$, $\text{tr}(AB) = \text{tr}(BA)$

Exercice : Calculer pour $A \in \mathcal{M}_n(\mathbf{K})$, $\text{tr}(A^\top A)$.

ATTENTION

La trace d'un produit n'est pas (en général) le produit des traces.

Corollaire 2.27

Soit A et B deux matrices semblables, $\text{tr}(A) = \text{tr}(B)$.

Démonstration : Par définition si A et B sont semblables, il existe $P \in \text{GL}_n(\mathbf{K})$ telle que $B = P^{-1}AP$. On en tire que

$$\text{tr}(B) = \text{tr}(P^{-1}AP) = \text{tr}(APP^{-1}) = \text{tr}(A).$$

□

Exercices :

1. Trouver deux matrices A et B ayant la même trace mais n'étant pas semblables.
2. Deux matrices équivalentes ont-elles la même trace ?

Proposition-Définition 2.28

Soit u un endomorphisme de E (espace vectoriel de dimension finie). Pour toute base $\mathcal{B}$ la trace de $\text{Mat}_{\mathcal{B}}(u)$ est la même. On l'appelle la trace de u et on la note $\text{tr}(u)$.

3 Formes linéaires et hyperplans

On note $E^* = \mathcal{L}(E, \mathbf{K})$ l'ensemble des formes linéaires que l'on appelle aussi l'espace dual.

3.1 Changement de bases

Soit $\mathcal{B}$ une base de E , on a coutume de calculer les matrices d'une forme linéaire f en prenant la base canonique (c'est-à-dire 1_K) pour l'espace d'arrivée qui est K . On note alors $\text{Mat}_{\mathcal{B}}(f)$ pour $\text{Mat}_{\mathcal{B}, \text{can}}(f)$.

ATTENTION

Ne pas confondre avec l'abus de notations $\text{Mat}_{\mathcal{B}}(u) = \text{Mat}_{\mathcal{B}, \mathcal{B}}(u)$ dans le cas des endomorphismes.

Soit x un vecteur de E et $X = \text{Mat}_{\mathcal{B}}(x)$, si on note $L = \text{Mat}_{\mathcal{B}}(f)$ alors $LX = \text{Mat}_{\text{can}}(f(x))$, on obtient donc $f(x)$ en identifiant $\mathcal{M}_1(K)$ avec K .

Exemple : Soit

$$\begin{aligned} f: \mathbb{C}_3[X] &\rightarrow \mathbb{C} \\ P &\mapsto \int_0^1 P(t) dt \end{aligned}$$

Si on prend pour $\mathcal{B}$ la base canonique de $\mathbb{C}_3[X]$, On a alors $L = \begin{pmatrix} 1 & \frac{1}{2} & \frac{1}{3} & \frac{1}{4} \end{pmatrix}$.

Pour $P = 1 - 4X^2$ par exemple, on a $X = \begin{pmatrix} 1 \\ 0 \\ -4 \\ 0 \end{pmatrix}$ et on obtient bien

$$LX = \int_0^1 1 - 4t^2 dt = -\frac{1}{3}.$$

Proposition 2.29

Soit $\mathcal{B}$ et $\mathcal{B}'$ deux bases de E et $f \in \mathcal{L}(E, K)$. On note P la matrice de changement de bases de $\mathcal{B}$ à $\mathcal{B}'$, si on note $L = \text{Mat}_{\mathcal{B}}(f)$ et $L' = \text{Mat}_{\mathcal{B}'}(f)$ alors

$$L' = LP$$

ATTENTION

Cette forme est à l'envers par rapport au changement de bases des vecteurs. On a

$$X = PX'$$

et

$$L' = LP$$

Démonstration : Il suffit d'utiliser la formule du changement de bases en prenant garde au fait que l'on ne change pas la base à l'arrivée.

$$L' = I_n LP$$

□

3.2 Bases duales et formes coordonnées

Définition 2.30 (Formes coordonnées)

Soit E un espace vectoriel de dimension finie (notée n). Soit $(e_1, \dots, e_n)$ une base de E . On sait que tout vecteur u de E se décompose de manière unique :

$$u = \sum_{i=1}^n x_i e_i.$$

Pour $i \in \llbracket 1; n \rrbracket$, on appelle alors i -ème forme coordonnée et on note e_i^* l'application qui associe au vecteur v le scalaire x_i .

ATTENTION

Même si la notation ne fait apparaître que le vecteur e_i , la forme e_i^* dépend de toute la base. Par exemple pour $E = \mathbb{R}^2$, si on pose $e_1 = (1, 0)$, $e_2 = (0, 1)$ et $e'_2 = (1, 1)$. Si on considère la base (e_1, e_2) et $u = (7, 3)$ on a $u = 7e_1 + 3e_2$ et donc $e_1^*(u) = 7$. Par contre, si on considère la base (e_1, e'_2) et toujours $u = (7, 3)$, on a $u = 4e_1 + 3e'_2$ et donc $e_1^*(u) = 4$.

Remarque : Soit $(i, j) \in \llbracket 1; n \rrbracket$,

$$e_i^*(e_j) = \delta_{ij}.$$

Proposition 2.31

Avec les notations précédentes,

1. Les formes coordonnées sont des formes linéaires.
2. La famille $(e_1^*, \dots, e_n^*)$ est une base de E^* .

Démonstration :

1. Soit $u = \sum_{i=1}^n x_i e_i$ et $v = \sum_{i=1}^n y_i e_i$. Si on pose λ et μ dans $\mathbf{K}$ et

$$w = \lambda u + \mu v = \sum_{i=1}^n (\lambda x_i + \mu y_i) e_i.$$

On en déduit que

$$e_i^*(w) = \lambda x_i + \mu y_i = \lambda e_i^*(u) + \mu e_i^*(v).$$

Les formes coordonnées sont bien linéaires.

2. Montrons que $(e_1^*, \dots, e_n^*)$ est une base de E^* . Comme on sait que $\dim E^* = n$, il suffit de montrer qu'elle est libre. Soit $\lambda_1, \dots, \lambda_n$ des scalaires tels que $\lambda_1 e_1^* + \dots + \lambda_n e_n^* = 0$. En particulier, si on évalue en e_i , on obtient que $\lambda_i = 0$.

La famille est bien libre est c'est une base.

Exemple : Soit $E = \mathbb{R}_n[X]$ et $(1, X, \dots, X^n)$ la base canonique. On sait, d'après la formule de Taylor que pour tout polynôme P ,

$$P = \sum_{k=0}^n \frac{P^{(k)}(0)}{k!} X^k.$$

Cela signifie que si on note $e_i = X^i$ alors

$$e_i^* : P \mapsto \frac{P^{(i)}(0)}{i!}.$$

3.3 Formes linéaires et hyperplans

Définition 2.32

Un hyperplan de E est le noyau d'une forme linéaire non nulle.

Proposition 2.33

Si E est de dimension finie avec $\dim E = n$. Les hyperplans sont les espaces vectoriels de dimension $n - 1$.

Démonstration :

- Si H est un hyperplan et f une forme linéaire telle que $H = \ker f$, le théorème du rang nous dit que $\dim H = n - 1$ car $\text{rg}(f) = 1$.
- Soit H un sous-espace vectoriel de dimension $n - 1$. On considère une base $(u_1, \dots, u_{n-1})$ de H . On peut la compléter en une base $(u_1, \dots, u_{n-1}, u_n)$ de E . Il suffit de prendre $f = u_n^*$ la forme coordonnée.

□

Proposition 2.34

Soit H un hyperplan et D une droite qui n'est pas contenue dans H . On a $H \oplus D = E$.

Démonstration : Soit a un vecteur non nul de D . On a $D = \text{Vect}(a)$. Comme D n'est pas contenue dans H alors $H \cap D = \{0\}$. On en déduit que H et D sont en somme directe. Soit f maintenant une forme linéaire qui définit H . On a $f(a) \neq 0$ et, quitte à diviser par $f(a)$ on peut supposer que $f(a) = 1$. Dès lors, pour tout vecteur x de E

$$x = f(x)a + x - f(x)a$$

où $f(x)a \in D$ et $x - f(x)a \in H$.

□

Proposition 2.35

Soit f et g deux formes linéaires.

$$\text{Ker}(f) = \text{Ker}(g) \iff \exists \lambda \in \mathbb{K}^*, f = \lambda g.$$

Démonstration :

- $\Leftarrow$ Evident.
- $\Rightarrow$ Soit $D = \text{Vect}(u)$ un supplémentaire de H . On a $f(u)$ et $g(u)$ qui sont non nuls donc il existe λ tel que $f(u) = \lambda g(u)$. Maintenant tout vecteur w s'écrit $w_1 + \alpha u$ où $w_1 \in H$. Il suffit alors de faire le calcul.

□

Remarque : Cela revient à dire que deux équations d'un même hyperplan sont les mêmes à une constante multiplicative près.

4 Compléments en algèbre linéaire

4.1 Somme

Vous avez vu en première année la somme de deux sous-espace vectoriel. Nous allons généraliser cette notion à plus de deux.

Définition 2.36

Soit $F_1, F_2, \dots, F_n$ des sous-espaces vectoriels, on appelle somme de $F_1, F_2, \dots, F_n$ et on note $F_1 + F_2 + \dots + F_n$ ou $\sum_{i=1}^n F_i$ l'ensemble de tous les vecteurs w de E qui peuvent s'écrire sous la forme $w = u_1 + u_2 + \dots + u_n$ où, pour tout i compris entre 1 et n , $u_i \in F_i$. On a donc

$$F_1 + F_2 + \dots + F_n = \{u_1 + \dots + u_n \mid (u_1, \dots, u_n) \in F_1 \times \dots \times F_n\}$$

Proposition 2.37

Soit $F_1, F_2, \dots, F_n$ des sous-espaces vectoriels.

1. La somme $F_1 + \dots + F_n$ est un sous-espace vectoriel de E .
2. La somme $F_1 + \dots + F_n$ est le plus petit sous-espace vectoriel contenant chaque F_i , c'est-à-dire contenant l'ensemble $\bigcup_{i=1}^n F_i$

Démonstration :

1. On utilise la caractérisation usuelle. Pour tout $i \in \llbracket 1 ; n \rrbracket$, $0_E \in F_i$. On en déduit que $0_E = 0_E + \dots + 0_E \in F_1 + \dots + F_n$. De ce fait, $F_1 + \dots + F_n$ n'est pas vide.

Soit w, w' dans $F_1 + \dots + F_n$ et λ, λ' deux scalaires. Par définition il existe $(u_1, \dots, u_n) \in F_1 \times \dots \times F_n$ et $(u'_1, \dots, u'_n) \in F_1 \times \dots \times F_n$ tels que $w = u_1 + \dots + u_n$ et $w' = u'_1 + \dots + u'_n$. On en déduit

$$\lambda w + \lambda' w' = \lambda(u_1 + \dots + u_n) + \lambda'(u'_1 + \dots + u'_n) = (\lambda u_1 + \lambda' u'_1) + \dots + (\lambda u_n + \lambda' u'_n) \in F_1 + \dots + F_n$$

On a bien montré que $F_1 + \dots + F_n$ était non vide et stable par combinaison linéaire. C'est un sous-espace vectoriel de E .

2. Commençons par remarquer que pour tout $i \in \llbracket 1 ; n \rrbracket$, $F_i \subset F_1 + \dots + F_n$. En effet pour $u_i \in F_i$,

$$u_i = 0_E + \dots + 0_E + u_i + 0_E + \dots + 0_E \in F_1 + \dots + F_n$$

Réciproquement, soit G un sous-espace vectoriel de E contenant tous les F_i montrons que $F_1 + \dots + F_n \subset G$. Soit $w \in F_1 + \dots + F_n$. Il existe $(u_1, \dots, u_n) \in F_1 \times \dots \times F_n$ tels que $w = u_1 + \dots + u_n$. Or, pour tout $i \in \llbracket 1 ; n \rrbracket$, $u_i \in G$ donc, par combinaison linéaire, $w \in G$.

On a bien montré que $F_1 + \dots + F_n$ est le plus petit sous-espace vectoriel contenant chaque F_i .

□

ATTENTION

Ne pas confondre l'union $\bigcup_{i=1}^n F_i$ et la somme $\sum_{i=1}^n F_i$. Dans la majorité des cas, une union de sous-espaces vectoriels n'est pas un sous-espace vectoriel.

Exemple : Dans $E = \mathbb{R}^3$, on note $F_1 = \text{Vect}(u_1)$ et $F_2 = \text{Vect}(u_2)$ où $u_1 = (1, -1, 0)$ et $u_2 = (1, -2, 1)$. Alors $F_1 + F_2 = \{(x, y, z) \in \mathbb{R}^3 \mid x + y + z = 0\}$.

Proposition 2.38

Soit $F_1, F_2, \dots, F_n$ des sous-espaces vectoriels et $\mathcal{F}_1, \dots, \mathcal{F}_n$ des familles qui engendrent ces espaces alors la concaténée des familles engendre $\sum_{i=1}^n F_i$.

Définition 2.39 (Somme directe)

Soit $F_1, F_2, \dots, F_n$ des sous-espaces vectoriels. On dit qu'ils sont en somme directe si

$$\forall (u_1, \dots, u_n) \in F_1 \times \dots \times F_n, \quad u_1 + \dots + u_n = 0 \Rightarrow u_1 = \dots = u_n = 0$$

Dans ce cas on note $F_1 \oplus \dots \oplus F_n$ ou $\bigoplus_{i=1}^n F_i$ la somme.

Remarque : Cela revient à dire que tout vecteur w de la somme se décompose de manière unique comme une somme $w = u_1 + \dots + u_n$ où $u_i \in F_i$.

ATTENTION

1. Dans le cas de deux espaces vectoriels il y a une caractérisation plus simple : F_1 et F_2 sont en somme directe si et seulement si $F_1 \cap F_2 = \{0\}$. Cette caractérisation ne s'entend pas à plus de trois espaces vectoriels
2. Des espaces peuvent être deux à deux en somme directe sans être en somme directe : par exemple trois droites du plan.

Exemple : Soit $n \geq 0$. Soit $Q \in \mathbb{R}_n[X]$ un polynôme non nul. On note $d = \deg(Q)$. Pour tout $k \in \llbracket 0; n-d \rrbracket$ on note $\Delta_k = \text{Vect}(X^k Q)$. Montrons que $\Delta_0, \dots, \Delta_{n-d}$ sont en somme directe.

Soit $0 = P_0 + \dots + P_{n-d} \in \Delta_0 + \dots + \Delta_{n-d}$ où $P_k \in \Delta_k$. Pour tout $k \in \llbracket 0; n-d \rrbracket$, il existe $\alpha_k \in \mathbb{R}$ tel que $P_k = X^k Q$. On en déduit que

$$0 = \left(\sum_{k=0}^{n-d} \alpha_k X^k \right) Q$$

Un produit de deux polynômes n'est nul que si l'un ou l'autre est nul donc, pour tout $k \in \llbracket 0; n-d \rrbracket$, $\alpha_k = 0$.

Cela montre bien que la somme est directe.

Exercice : Montrer que $F_1, \dots, F_n$ sont en somme directe si et seulement si pour tout $i \in \llbracket 1; n \rrbracket$,

F_i est en somme directe avec $\sum_{j \neq i} F_j$.

Proposition 2.40

Soit $F_1, F_2, \dots, F_n$ des sous-espaces vectoriels

1. Soit, pour tout $i \in \llbracket 1; n \rrbracket$ une famille libre $\mathcal{F}_i$ de F_i . Si la somme des $F_1, \dots, F_n$ est directe alors, la famille concaténée des $\mathcal{F}_i$ est libre.
2. Soit, pour tout $i \in \llbracket 1; n \rrbracket$ une base $\mathcal{F}_i$ de F_i . La somme des $F_1, \dots, F_n$ est directe si et seulement si la famille concaténée des $\mathcal{F}_i$ est une base de $\bigoplus_{i=1}^n F_i$.

Démonstration :

1. Pour tout $i \in \llbracket 1; n \rrbracket$, on note $(u_{i,1}, \dots, u_{i,d_i})$ la famille $\mathcal{F}_i$. On note $K = \{(i, j) \in \llbracket 1; n \rrbracket \times \mathbb{N}^*, j \leq d_i\}$. Soit $(\lambda_{i,j})_{(i,j) \in K}$ une famille de scalaire telle que $\sum_{i=1}^n \sum_{j=1}^{d_i} \lambda_{i,j} u_{i,j} = 0$. Pour tout $i \in \llbracket 1; n \rrbracket$ on pose $w_i = \sum_{j=1}^{d_i} \lambda_{i,j} u_{i,j} \in F_i$.

On a donc $w_1 + \dots + w_n = 0$ et donc, pour tout $i \in \llbracket 1; n \rrbracket$, $w_i = 0$ car la famille $F_1, \dots, F_n$ est en somme directe. En utilisant que les familles $\mathcal{F}_i$ sont libres on a bien que $(\lambda_{i,j})_{(i,j) \in K}$ est la famille nulle.

2. On procède par double inclusion

- $\Rightarrow$ Le caractère libre a été démontré ci-dessus. Comme chaque $\mathcal{F}_i$ engendre F_i , on obtient que la famille concaténée engendre bien $\bigoplus_{i=1}^n F_i$.
- $\Leftarrow$ Soit $w_1, \dots, w_n \in F_1 \times \dots \times F_n$ telle que $w_1 + \dots + w_n = 0$. On décompose chaque w_i dans la base $\mathcal{F}_i$. Le fait que la famille concaténée soit libre affirme alors que tous les w_i sont nuls. La somme $F_1, \dots, F_n$ est bien directe.

□

Remarque : En particulier, si on se donne une base $(u_i)_{i \in I}$ de E , que l'on partitionne I en $I_1, \dots, I_p$ et l'on note $F_k = \text{Vect}(u_i)_{i \in I_k}$ pour $k \in \llbracket 1; p \rrbracket$ alors

$$F_1 \oplus \dots \oplus F_p = E$$

Définition 2.41

Soit E un espace vectoriel et $F_1, \dots, F_p$ des sous-espaces vectoriels tels que $E = \bigoplus_{i=1}^p F_i$. Pour tout $i \in \llbracket 1; p \rrbracket$, on notera $\pi_i \in \mathcal{L}(E)$ le projecteur sur F_i parallèlement à $\bigoplus_{j \neq i} F_j$. Pour tout $w \in E$, il existe un unique p -uplet $(u_1, \dots, u_p) \in F_1 \times \dots \times F_p$ tel que $w = u_1 + \dots + u_p$. Pour tout $i \in \llbracket 1; p \rrbracket$, $\pi_i(w) = u_i$.

Exercice : Soit F_1, F_2 et F_3 trois sous-espaces de E . On suppose que F_1 et F_2 sont en somme directe et que $F_1 \oplus F_2$ est en somme directe avec F_3 . Montrer que F_1, F_2 et F_3 sont en somme directe.

Proposition 2.42 (Formule de Grassman)

1. On a $\dim(F + G) = \dim F + \dim G - \dim(F \cap G)$
2. On a $(\dim(F_1 + \cdots + F_n) = \dim F_1 + \cdots + \dim F_n) \iff$ (la somme de $F_1, \dots, F_n$ est directe).

4.2 Parties stables

Proposition 2.43 (Définition d'une application linéaire par son action sur une base)

Soit E et F des espaces vectoriels. Soit $\mathcal{B} = (u_i)_{i \in I}$ une base de E et $(v_i)_{i \in I}$ une famille de vecteurs de F indexées par le même ensemble. Il existe une unique application linéaire $f : E \rightarrow F$ telle que pour tout $i \in I$, $f(u_i) = v_i$. Elle est donnée par

$$f(w) = \sum_{i \in I} \lambda_i v_i \text{ où } w = \sum_{i \in I} \lambda_i u_i.$$

Remarque : Si on suppose que E et F sont de dimension finie. On peut alors définir la matrice de l'application linéaire f . Si on note $\mathcal{C}$ une base de F , on appelle par définition matrice de f dans les bases $\mathcal{B}$ et $\mathcal{C}$ et on note $\text{Mat}_{\mathcal{B}, \mathcal{C}}(f)$ la matrice dont les colonnes sont les coordonnées des $v_i = f(u_i)$ dans la base $\mathcal{C}$.

L'énoncé ci-dessus, revient à exprimer f en connaissant sa valeur sur les droites vectorielle $\Delta_i = \text{Vect}(u_i)$. On peut généraliser cela.

Proposition 2.44

Soit E et F des espaces vectoriels. Soit $E_1, \dots, E_p$ des sous-espaces vectoriels de E tels que

$$E = \bigoplus_{i=1}^p E_i$$

Soit $u_1, \dots, u_p \in \mathcal{L}(E_1, F) \times \cdots \times \mathcal{L}(E_p, F)$.

Il existe une unique application linéaire $u \in \mathcal{L}(E, F)$ telle que pour tout $i \in \llbracket 1 ; p \rrbracket$, $u|_{E_i} = u_i$. Elle est définie par

$$\forall w \in E, u(w) = \sum_{i=1}^p u_i(w_i) \text{ où } w = \underbrace{w_1}_{\in E_1} + \cdots + \underbrace{w_p}_{\in E_p}$$

Remarque : Dans le cas où les espaces vectoriels sont de dimension finie, on peut visualiser ce résultat sur les matrices. En effet si on considère des bases $\mathcal{B}_i$ de E_i et si on note $\mathcal{B}$ la base obtenue en concaténant ces bases.

On obtient alors la matrice $A = \text{Mat}_{\mathcal{B}, \mathcal{C}}(u)$ en « collant » les matrices $A_i = \text{Mat}_{\mathcal{B}_i, \mathcal{C}}(u_i)$:

$$A = \left(A_1 \mid \cdots \mid A_p \right).$$

Définition 2.45

Soit u une application de E dans E . Soit X une partie de E elle est dite stable par u si

$$u(X) \subset X$$

Remarque : On utilisera cela souvent dans le cas où u est linéaire et X un sous-espace vectoriel mais la terminologie existe dans ce cadre plus général (par exemple intervalle stable dans l'étude des suites récurrentes $u_{n+1} = f(u_n)$.)

Définition 2.46

Soit u une application de E dans E . Soit X une partie stable par u de E , on considère souvent l'application induite

$$\begin{aligned} \tilde{u} : X &\rightarrow X \\ x &\mapsto u(x) \end{aligned}$$

ATTENTION

De manière générale, pour $u \in \mathcal{F}(E, E)$ et $X \subset E$, on peut définir $u|_X : X \rightarrow E$. Dans le cas où la partie est stable on peut aussi se limiter à X pour l'espace d'arrivée. De ce fait si u est linéaire et X un sous-espace vectoriel, l'application $\tilde{u}$ est un **endomorphisme** de E .

Exemples :

1. Soit $u : \mathbb{R}^3 \rightarrow \mathbb{R}^3$ dont la matrice dans la base canonique est

$$A = \begin{pmatrix} 2 & 2 & 3 \\ 0 & -1 & -2 \\ -1 & 0 & 0 \end{pmatrix}.$$

Soit $H = \{(x, y, z) \in \mathbb{R}^3 \mid x + y + z = 0\}$. Montrons que H est stable par u et calculons la matrice de l'application induite par u dans H dans une base.

On détermine une base de H . On pose $v_1 = (1, -1, 0)$, $v_2 = (0, 1, -1)$. C'est une base de H (qui est de dimension 2).

On calcule $u(v_1)$ et $u(v_2)$ en faisant le produit :

$$\begin{pmatrix} 2 & 2 & 3 \\ 0 & -1 & -2 \\ -1 & 0 & 0 \end{pmatrix} \begin{pmatrix} 1 & 0 \\ -1 & 1 \\ 0 & -1 \end{pmatrix} = \begin{pmatrix} 0 & -1 \\ 1 & 1 \\ -1 & 0 \end{pmatrix}.$$

On voit donc que $u(v_1) = v_2$ et $u(v_2) = -v_1$. De ce fait H est stable par u et la matrice de $\tilde{u}$ dans la base (v_1, v_2) est

$$\begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}.$$

On reconnaît un quart de tour (rotation d'angle $\pi/2$).

2. Recommencer avec $A = \begin{pmatrix} 1 & 1 & -1 \\ 2 & 1 & 0 \\ 2 & 2 & -2 \end{pmatrix}$ et $H = \{(x, y, z) \in \mathbb{R}^3 \mid 2x - z = 0\}$.

Proposition 2.47

Soit u un endomorphisme de E (de dimension finie). Soit F un sous-espace vectoriel de E . On note $\mathcal{F}$ une base de F et $\mathcal{G}$ une base de G qui est un supplémentaire de F dans E . Soit $\mathcal{B}$ la base obtenue en concaténant $\mathcal{F}$ et $\mathcal{G}$,

$$F \text{ est stable par } u \iff \text{Mat}_{\mathcal{B}}(u) = \left(\begin{array}{ccc|ccc} * & \cdots & * & * & \cdots & * \\ \vdots & & \vdots & \vdots & & \vdots \\ * & \cdots & * & * & \cdots & * \\ \hline & & & 0 & & \\ & & & & * & \cdots & * \\ & & & & \vdots & & \vdots \\ & & & & * & \cdots & * \end{array} \right)$$

Démonstration : On note $\mathcal{F} = (f_1, \dots, f_p)$ et $\mathcal{G} = (g_{p+1}, \dots, g_n)$.

- $\Rightarrow$: On suppose que F est stable par u . De ce fait pour $i \in \llbracket 1; p \rrbracket$, $u(f_i) \in F$ donc il s'exprime comme combinaison linéaire uniquement de f_k .
- $\Leftarrow$: On suppose que la matrice de u est de la forme voulue. On voit donc que pour tout $i \in \llbracket 1; p \rrbracket$, $u(f_i) \in F$. De ce fait, pour tout $w \in F$, il s'écrit $w = \sum_{i=1}^p \lambda_i f_i$ et donc $u(w) = \sum_{i=1}^p \lambda_i u(f_i) \in F$. Donc F est stable par u .

□

4.3 Calcul par blocs

Théorème 2.48

Soit $(A_{ij})_{(i,j) \in \llbracket 1; r \rrbracket \times \llbracket 1; s \rrbracket}$ où $A_{ij} \in \mathcal{M}_{n_i p_j}(\mathbf{K})$ et $(B_{jk})_{(j,k) \in \llbracket 1; s \rrbracket \times \llbracket 1; t \rrbracket}$ où $B_{jk} \in \mathcal{M}_{p_j q_k}(\mathbf{K})$. On note

$$A = \left(\begin{array}{ccc|ccc} A_{11} & \cdots & A_{1s} & & & \\ \vdots & & \vdots & & & \\ A_{r1} & \cdots & A_{rs} & & & \end{array} \right) \text{ et } B = \left(\begin{array}{ccc|ccc} B_{11} & \cdots & B_{1t} & & & \\ \vdots & & \vdots & & & \\ B_{s1} & \cdots & B_{st} & & & \end{array} \right)$$

Le produit AB est donné par

$$AB = \left(\begin{array}{ccc|ccc} \sum_{j=1}^s A_{1j} B_{j1} & \cdots & \sum_{j=1}^s A_{1j} B_{js} & & & \\ \vdots & & \vdots & & & \\ \sum_{j=1}^s A_{rj} B_{j1} & \cdots & \sum_{j=1}^s A_{rj} B_{js} & & & \end{array} \right)$$

Démonstration : Il « suffit » de faire le calcul.

□

Interprétation géométrique

Si on se donne E et F des espaces vectoriels et que l'on a des décomposition en somme directes

$$E = \bigoplus_{j=1}^s E_j \text{ et } F = \bigoplus_{i=1}^r F_i$$

On se donne alors $\mathcal{E}_1, \dots, \mathcal{E}_s$ des bases de $E_1, \dots, E_s$ et $\mathcal{F}_1, \dots, \mathcal{F}_r$ des bases de $F_1, \dots, F_r$. On note $\mathcal{E}$ et $\mathcal{F}$ les bases de E et F obtenues en concaténant les bases ci-dessus. Dans ce cas pour $u \in \mathcal{L}(E, F)$, si on note

$$A = \text{Mat}_{\mathcal{E}, \mathcal{F}}(u) = \left(\begin{array}{c|c|c} A_{11} & \cdots & A_{1s} \\ \vdots & & \vdots \\ A_{r1} & \cdots & A_{rs} \end{array} \right).$$

On peut alors voir A_{ij} comme la matrice (dans les bases $\mathcal{E}_i$ et $\mathcal{F}_j$) de la restriction à E_j de $\pi_i \circ u$ où π_i la projection sur F_i parallèlement à $\bigoplus_{j \neq i} F_j$.

Dans le cas des endomorphismes on prend la même décomposition pour l'espace de départ et d'arrivée

$$E = \bigoplus_{j=1}^s E_j$$

La matrice de u dans la base $\mathcal{E}$ sera dite « diagonale par blocs » si pour tout i , E_i est stable par u . Dans ce cas on obtient

$$\left(\begin{array}{c|ccc|c} A_{11} & 0 & \cdots & 0 & 0 \\ 0 & \ddots & & & 0 \\ \vdots & & \ddots & & \vdots \\ 0 & & & \ddots & 0 \\ \hline 0 & 0 & \cdots & 0 & A_{ss} \end{array} \right).$$

4.4 Calculs par blocs et déterminants

Commençons par un rapide rappel sur le déterminant.

Définition 2.49 (Rappel - Déterminant)

1. Soit $A \in \mathcal{M}_n(\mathbf{K})$. On appelle déterminant de A et on note $\det(A)$ pour

$$\det(A) = \sum_{\sigma \in \mathfrak{S}_n} \varepsilon(\sigma) A[1, \sigma(1)] A[2, \sigma(2)] \cdots A[n, \sigma(n)]$$

2. Soit $u \in \mathcal{L}(E)$, toutes les matrices de u ont le même déterminant. On l'appelle le déterminant de u et on le note $\det(u)$.

Proposition 2.50

1. Soit $A \in \mathcal{M}_n(\mathbf{K})$. La matrice A est inversible si et seulement si $\det(A) \neq 0$.
2. Soit $u \in \mathcal{L}(E)$. L'endomorphisme u est un automorphisme si et seulement si $\det(u) \neq 0$.

Proposition 2.51

Soit M une matrice de $\mathcal{M}_n(\mathbf{K})$ pouvant s'écrire par blocs

$$M = \left(\begin{array}{c|c} A & B \\ \hline 0 & D \end{array} \right)$$

avec $A \in \mathcal{M}_{n'}(\mathbf{K})$, $B \in \mathcal{M}_{n',n-n'}(\mathbf{K})$ et $D \in \mathcal{M}_{n-n'}(\mathbf{K})$. On a

$$\det(M) = \det(A) \det(D)$$

Démonstration : Il y a plusieurs démonstrations de cette proposition.

— La première consiste à reprendre la formule générale du déterminant :

$$\det(M) = \sum_{\sigma \in \mathfrak{S}_n} \varepsilon(\sigma) a_{\sigma(1),1} \cdots a_{\sigma(n),n}$$

Maintenant, pour tout $p > n'$ et $q \leq n'$, $a_{pq} = 0$. De ce fait, dans la somme ci-dessus on peut ne garder que les permutations σ qui laissent $[[n' + 1; n]]$ stable. De ce fait, $[[1; n']]$ est aussi stable. Une telle permutation est donc donnée par une permutation σ_1 de $[[1; n']]$ et une permutation σ_2 de $[[n' + 1; n]]$. On en déduit donc que

$$\det(M) = \left(\sum_{\sigma_1 \in \mathfrak{S}_{n'}} \varepsilon(\sigma_1) a_{\sigma_1(1),1} \cdots a_{\sigma_1(n'),n'} \right) \times \left(\sum_{\sigma_2 \in S} \varepsilon(\sigma_2) a_{\sigma_2(n'+1),n'+1} \cdots a_{\sigma_2(n),n} \right) = \det(A) \times \det(D),$$

où S est l'ensemble des permutations de $[[n' + 1; n]]$.

— Une autre preuve (utile dans certains exercices) consiste à écrire la matrice M comme un produit. On a

$$M = \left(\begin{array}{cc} I & 0 \\ 0 & D \end{array} \right) \times \left(\begin{array}{cc} A & B \\ 0 & I \end{array} \right).$$

Maintenant en développant on trouve que

$$\left| \begin{array}{cc} A & B \\ 0 & I \end{array} \right| = \det(A) \text{ et } \left| \begin{array}{cc} I & 0 \\ 0 & D \end{array} \right| = \det(D).$$

Donc $\det(M) = \det(A) \det(D)$.

□

Corollaire 2.52 (Déterminant d'une matrice triangulaire par blocs)

Soit A une matrice triangulaire par blocs :

$$A = \left(\begin{array}{cccc} A_{11} & * & \cdots & * \\ 0 & \ddots & \ddots & \vdots \\ \vdots & & \ddots & * \\ 0 & \cdots & 0 & A_{nn} \end{array} \right)$$

Alors

$$\det(A) = \prod_{k=1}^n \det(A_{kk})$$

Démonstration : Il suffit de faire une récurrence sur le nombre de blocs. □

Définition 2.53 (Matrice de transvection)

Soit $n \in \mathbb{N}^*$. Soit $(i, j) \in \llbracket 1 ; n \rrbracket^2$ avec $i \neq j$ et $\lambda \in \mathbf{K}$. On appelle matrice de transvection la matrice

$$T_{i,j}(\lambda) = I_n + \lambda E_{i,j}$$

Remarques :

1. Ces matrices ont peut-être été vues en première année. Elles servent à coder matriciellement des opérations élémentaires. En effet si $A \in \mathcal{M}_{n,p}(\mathbf{K})$, la matrice $A' = T_{i,j}(\lambda)A$ est obtenue à partir de A en faisant l'opération $L_i \leftarrow L_i + \lambda L_j$. De même, si $B \in \mathcal{M}_{p,n}(\mathbf{K})$, la matrice $B' = AT_{i,j}(\lambda)$ est obtenue à partir de B en faisant l'opération $C_j \leftarrow C_j + \lambda C_i$.
2. En utilisant les matrices de transvections, on constate que l'on ne modifie pas le déterminant d'une matrice A en réalisant des opérations élémentaires du type $L_i \leftarrow L_i + \lambda L_j$ ou $C_j \leftarrow C_j + \lambda C_i$.

Définition 2.54 (Transvection par blocs)

Soit A une matrice qui s'écrit par blocs

$$A = \left(\begin{array}{c|c|c} A_{11} & \cdots & A_{1s} \\ \vdots & & \vdots \\ \hline A_{r1} & \cdots & A_{rs} \end{array} \right)$$

1. Si on suppose que les blocs ont le même nombre de lignes noté d , on peut considérer la transvection de blocs $L_i \leftarrow L_i + \lambda L_j$ où $(i, j) \in \llbracket 1 ; r \rrbracket^2$ avec $i \neq j$ qui consiste à faire les transvections de lignes $L_{(i-1)d+u} \leftarrow L_{(i-1)d+u} + \lambda L_{(j-1)d+u}$ pour $u \in \llbracket 1 ; d \rrbracket$. On obtient donc la matrice

$$A' = \left(\begin{array}{c|c|c} A_{11} & \cdots & A_{1s} \\ \vdots & & \vdots \\ \hline A_{i1} + \lambda A_{j1} & & A_{is} + \lambda A_{js} \\ \vdots & & \vdots \\ \hline A_{r1} & \cdots & A_{rs} \end{array} \right)$$

2. Si on suppose que les blocs ont le même nombre de colonne noté t , on peut considérer la transvection de blocs $C_i \leftarrow C_i + \lambda C_j$ où $(i, j) \in \llbracket 1 ; s \rrbracket^2$ avec $i \neq j$ qui consiste à faire les transvections de colonnes $C_{(i-1)t+u} \leftarrow C_{(i-1)t+u} + \lambda C_{(j-1)t+u}$ pour $u \in \llbracket 1 ; t \rrbracket$. On obtient donc la matrice

$$A' = \left(\begin{array}{c|c|c|c|c} A_{11} & \cdots & A_{1i} + \lambda A_{1j} & \cdots & A_{1s} \\ \vdots & & & & \vdots \\ \hline A_{r1} & \cdots & A_{ri} + \lambda A_{rj} & \cdots & A_{rs} \end{array} \right)$$

Proposition 2.55

Soit A une matrice carrée. Son déterminant est invariant par transvections par blocs.

Démonstration : Il suffit d'utiliser que les transvections « classiques » ne modifient pas le déterminant.

□

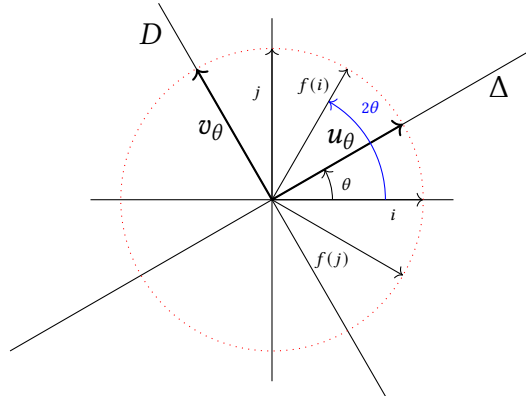
5 Éléments propres d'un endomorphisme

Dans toute cette partie E est un espace vectoriel et $u \in \mathcal{L}(E)$.

Exemple : On considère l'espace vectoriel $E = \mathbb{R}^2$. Pour $\theta \in \mathbb{R}$, on note $u_\theta = (\cos \theta, \sin \theta)$ et $v_\theta = (\cos \varphi, \sin \varphi) = (-\sin(\theta), \cos(\theta))$ le vecteur directement orthogonal à u_θ où $\varphi = \theta + \frac{\pi}{2}$. On regarde la symétrie s par rapport à l'axe $\Delta = \text{Vect}(u_\theta)$ et parallèlement à $D = \text{Vect}(v_\theta)$. Dans la base canonique $\mathcal{C}$ sa matrice est

$$\text{Mat}_{\mathcal{C}}(s) = \begin{pmatrix} \cos(2\theta) & \sin(2\theta) \\ \sin(2\theta) & -\cos(2\theta) \end{pmatrix}$$

Cela peut se voir sur un dessin



En effet $s(u_\theta) = (\cos(2\theta), \sin(2\theta))$ et $s(v_\theta) = (\cos(\psi), \sin(\psi)) = (\sin(2\theta), -\cos(2\theta))$ où $\psi = \theta - (\frac{\pi}{2} - \theta) = 2\theta - \frac{\pi}{2}$.

Pendant si on considère la base $\mathcal{B} = (u_\theta, v_\theta)$ alors

$$\text{Mat}_{\mathcal{B}}(s) = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$$

Cette deuxième matrice est particulièrement simple car les vecteurs u et v qui forment la base ont un image simple par s ; $s(u) = u$ et $s(v) = -v$.

On peut d'ailleurs retrouver plus rigoureusement la matrice dans la base canonique. Si on note P la matrice de passage de la base canonique à $\mathcal{C}$, on a

$$P = \begin{pmatrix} \cos(\theta) & -\sin(\theta) \\ \sin(\theta) & \cos(\theta) \end{pmatrix} \text{ et } P^{-1} = \begin{pmatrix} \cos(\theta) & \sin(\theta) \\ -\sin(\theta) & \cos(\theta) \end{pmatrix}$$

On a alors

$$\text{Mat}_{\mathcal{C}}(s) = P \text{Mat}_{\mathcal{B}}(s) P^{-1} = \begin{pmatrix} \cos^2(\theta) - \sin^2(\theta) & 2 \cos(\theta) \sin(\theta) \\ 2 \cos(\theta) \sin(\theta) & \sin^2(\theta) - \cos^2(\theta) \end{pmatrix} = \begin{pmatrix} \cos(2\theta) & \sin(2\theta) \\ \sin(2\theta) & -\cos(2\theta) \end{pmatrix}$$

On voit sur cet exemple que, pour un même endomorphisme, il peut être intéressant de bien choisir la base pour exprimer une matrice représentant l'endomorphisme.

5.1 Définition

Définition 2.56 (Vecteurs propres)

Soit $x \in E$, on dit que x est un vecteur propre de u si :

- le vecteur x n'est pas nul
- il existe un scalaire $\lambda \in \mathbf{K}$ tel que $u(x) = \lambda x$.

Remarque : Dire que x est un vecteur propre pour u revient à dire que $x \neq 0$ et que la droite $\Delta = \text{Vect}(x)$ est stable par u .

Définition 2.57 (Valeur propre)

Soit $\lambda \in \mathbf{K}$ un scalaire, on dit que c'est une valeur propre de u , s'il existe un vecteur **non nul** x tel que $u(x) = \lambda x$. Dans ce cas, x est un vecteur propre que l'on dit associé à la valeur propre λ .

ATTENTION

Il est important de ne pas oublier que x ne doit pas être nul. En effet pour tout $\lambda \in \mathbf{K}$, $u(0) = 0 = \lambda \cdot 0$.

Exemples :

1. Soit

$$\begin{aligned} \Delta : \mathbf{R}[X] &\rightarrow \mathbf{R}[X] \\ P &\mapsto P' \end{aligned}$$

Soit P un éventuel vecteur propre, $\Delta(P) = P' = \lambda \cdot P$. On voit, pour des raisons de degré que si $\lambda \neq 0$ ce n'est pas possible. Par contre pour $\lambda = 0$, tous les polynômes constants conviennent.

On en déduit que 0 est la seule valeur propre et que l'ensemble des vecteurs propres associés à la valeur propre 0 est $\mathbf{R}_0[X]$ l'ensemble des polynômes constants non nuls.

2. Soit $u = 0_{\mathcal{L}(E)}$. Là encore, l'unique valeur propre est 0. Par contre l'ensemble des vecteurs propres associés à la valeur propre 0 est E en entier (privé de 0).
3. Soit $u = \text{id}_E$. Cette fois, l'unique valeur propre est 1 ; l'ensemble des vecteurs propres associés à la valeur propre 1 est E en entier (privé de 0).
4. Soit p un projecteur (qui vérifie $p^2 = p$) que l'on suppose différent de $0_{\mathcal{L}(E)}$ et de id_E . On note $F = \text{Ker } p$ et $G = \text{Ker}(p - \text{id}_E)$ de telle sorte que ce soit le projecteur sur G parallèlement à F . Comme $p \neq \text{id}_E$ et $p \neq 0_{\mathcal{L}(E)}$ alors ni F ni G ne sont égaux à F donc ils ne sont pas non plus réduits à $\{0\}$.

Pour tout x de F on a donc $p(x) = 0$. De ce fait, 0 est une valeur propre et F est l'ensemble des vecteurs propres associées à 0. De même 1 est une valeur propre et G est l'ensemble des vecteurs propres associées à 1.

Supposons maintenant qu'il existe une autre valeur propre λ . Il existe donc $x \in E$ tel que $x \neq 0$

et $p(x) = \lambda x$. Donc

$$p^2(x) = p(p(x)) = p(\lambda x) = \lambda \cdot p(x) = \lambda^2 \cdot x$$

Mais $p^2(x) = p(x) = \lambda x$ car p est un projecteur. On en déduit que $\lambda^2 \cdot x = \lambda \cdot x$ et donc $\lambda^2 = \lambda$ car $x \neq 0$.

Finalement 0 et 1 sont les seules valeurs propres possibles.

5. Soit u l'endomorphisme de $\mathbb{R}^3$ dont la matrice dans la base canonique est

$$A = \begin{pmatrix} 0 & 2 & -1 \\ 3 & -2 & 0 \\ -2 & 2 & 1 \end{pmatrix}$$

On cherche à quel condition sur λ il existera $x \in E$ tel que $x \neq 0$ et $u(x) = \lambda \cdot x \iff u(x) - \lambda \cdot x = 0 \iff (u - \lambda \cdot \text{id}_E)(x) = 0$. On cherche donc les valeurs de λ telle que $u - \lambda \cdot \text{id}_E$ ne soit pas injective (donc pas bijective).

Il suffit de calculer le déterminant de $A - \lambda \cdot I_3$:

$$\begin{vmatrix} -\lambda & 2 & -1 \\ 3 & -2-\lambda & 0 \\ -2 & 2 & 1-\lambda \end{vmatrix} = -\lambda^3 - \lambda^2 + 10\lambda - 8 = -(\lambda - 1)(\lambda - 2)(\lambda + 4)$$

On en déduit que les valeurs propres de u sont 1, 2 et -4 .

Calculons les vecteurs propres associés à la valeur 1. On résout $u(x) = x \iff (u - \text{id}_E)(x) = 0$.
En matrice

$$\begin{pmatrix} -1 & 2 & -1 \\ 3 & -3 & 0 \\ -2 & 2 & 0 \end{pmatrix}$$

On obtient que les vecteurs propres sont $\text{Vect}(e_1)$ où $e_1 = (1, 1, 1)$.

De même les vecteurs propres associés à 2 sont $\text{Vect}(e_2)$ où $e_2 = (4, 3, -2)$ et ceux associés à -4 sont $\text{Vect}(e_3)$ où $e_3 = (2, -3, 2)$.

Exercice : Déterminer les valeurs propres et vecteurs propres de l'endomorphisme $u \in \mathcal{L}(\mathbb{R}^3)$ dont la matrice dans la base canonique est

$$\begin{pmatrix} 1 & 1 & -1 \\ 0 & 1 & 0 \\ 1 & 0 & 1 \end{pmatrix}.$$

Définition 2.58 (Spectre)

On appelle spectre de u et on note $Sp(u)$ l'ensemble des valeurs propres.

Exercices :

1. Quel est le spectre d'une symétrie ?
2. Déterminer les valeurs propres de l'endomorphisme de $\mathbb{R}^3$ dont la matrice dans la base canonique est :

$$\begin{pmatrix} 1 & 2 & 2 \\ 2 & 1 & 2 \\ 2 & 2 & 1 \end{pmatrix}$$

Proposition 2.59 (Caractérisation des valeurs propres)

Si E est un espace vectoriel de dimension finie et si $u \in \mathcal{L}(E)$. Soit $\lambda \in \mathbf{K}$;

$$\lambda \in \text{Sp}(u) \iff \det(\lambda \text{id}_E - u) = 0$$

Démonstration : Soit $\lambda \in \mathbf{K}$.

$$\begin{aligned} \lambda \in \text{Sp}(u) &\iff \exists x \neq 0_E, u(x) = \lambda x \\ &\iff \exists x \neq 0_E, (u - \lambda \text{id}_E)(x) = 0_E \\ &\iff \text{Ker}(u - \lambda \text{id}_E) \neq \{0_E\} \\ &\iff u - \lambda \text{id}_E \text{ n'est pas injectif} \\ &\iff \det(\lambda \text{id}_E - u) = 0 \end{aligned}$$

□

5.2 Sous-espaces propres

Définition 2.60

Soit $\lambda \in \mathbf{K}$. On appelle sous-espace propre associé à λ et on note $E_\lambda(u)$ le sous-espace vectoriel

$$E_\lambda(u) = \text{Ker}(u - \lambda \text{id}_E) = \{x \in E \mid u(x) = \lambda x\}$$

Remarque : On voit que λ est une valeur propre si et seulement si $E_\lambda(u) \neq \{0\}$. Dans certains cas on trouve la terminologie sous-espace propre réservé aux sous-espaces propres différents de $\{0\}$.

ATTENTION

Le sous-espace propre $E_\lambda(u)$ est constitué des vecteurs propres associés à λ et du vecteur nul. Ne pas dire que le sous-espace propre est l'ensemble des vecteurs propres

Proposition 2.61

Les sous-espaces propres d'un endomorphisme u sont stables par u et l'application induite est une homothétie $x \mapsto \lambda.x$

Démonstration : Soit $E_\lambda(u)$ le sous-espace propre associé à la valeur propre λ . Pour tout $x \in E_\lambda(u)$, on a $u(x) = \lambda.x \in E_\lambda(u)$. L'espace est stable et l'endomorphisme induit est l'homothétie $x \mapsto \lambda.x$

□

Théorème 2.62

Soit $\lambda_1, \dots, \lambda_n$ des valeurs propres deux à deux distinctes. Les sous-espaces propres $E_{\lambda_1}(u), \dots, E_{\lambda_n}(u)$ sont en somme directe.

Démonstration : On procède par récurrence sur n .

- Pour $n = 1$ c'est évident.
- Soit $n \geq 1$. On suppose que si $\lambda_1, \dots, \lambda_n$ sont des valeurs propres deux à deux distinctes alors les sous-espaces propres $E_{\lambda_1}(u), \dots, E_{\lambda_n}(u)$ sont en somme directe. On se donne alors $\lambda_1, \dots, \lambda_{n+1}$ des valeurs propres deux à deux distinctes

On veut montrer que si $x_1 + \dots + x_{n+1} = 0$ avec pour tout $i \in \llbracket 1; n+1 \rrbracket$, $x_i \in E_{\lambda_i}(u)$ alors pour tout i , $x_i = 0$.

On suppose donc

$$x_1 + \dots + x_{n+1} = 0$$

avec pour tout $i \in \llbracket 1; n+1 \rrbracket$, $x_i \in E_{\lambda_i}(u)$. En appliquant u à l'égalité on obtient

$$0 = u(0) = u\left(\sum_{i=1}^{n+1} x_i\right) = \sum_{i=1}^{n+1} \lambda_i x_i.$$

En faisant une combinaison linéaire adaptée on peut alors supprimer le terme en x_{n+1} .

On obtient alors

$$\sum_{i=1}^n (\lambda_i - \lambda_{n+1}) x_i = 0$$

Maintenant comme $E_{\lambda_1}(u), \dots, E_{\lambda_n}(u)$ sont en somme directe, on en déduit que pour tout i , $(\lambda_i - \lambda_{n+1})x_i$ est nul et donc que $x_i = 0$ car $\lambda_i \neq \lambda_{n+1}$. Pour finir, x_{n+1} est aussi nul. On a bien montré que $E_{\lambda_1}(u), \dots, E_{\lambda_{n+1}}(u)$ sont en somme directe.

□

Corollaire 2.63

Soit $(x_i)_{i \in I}$ une famille de vecteurs propres associées à des valeurs propres deux à deux distinctes est libre.

Démonstration : Si I est un ensemble fini, il suffit d'utiliser la propriété précédente.

Dans le cas où I est infini, il suffit d'utiliser qu'une famille (infinie) est libre si toutes ses sous-familles finies sont libres.

□

Exemple : On considère la famille $f_k : t \mapsto \sin(kt)$ pour $k \in \mathbf{N}$. On a alors $f_k'' = -k^2 f_k$. On en déduit que f_k est un vecteur propre pour $\lambda = -k^2$ de l'endomorphisme de $\mathcal{C}^\infty$ défini par $f \mapsto f''$. La famille est donc libre.

Corollaire 2.64

Soit E un espace vectoriel de dimension finie n et $u \in \mathcal{L}(E)$,

$$\sum_{\lambda \in \text{Sp}(u)} \dim(E_\lambda(u)) \leq \dim E$$

En particulier, l'endomorphisme u a au plus n valeurs propres.

$$\#\text{Sp}(u) \leq \dim E$$

ATTENTION

Ce n'est pas une égalité :

- L'identité de $\mathbf{R}^3$ n'a qu'une valeur propre.
- L'application de $\mathbf{R}^2$ dans $\mathbf{R}^2$ définie par $\begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}$ n'a pas de valeurs propres (dans $\mathbf{R}$).
- La matrice $N = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}$, n'a qu'une seule valeur propre qui est 0 et $\dim(E_0(N)) = 1 < 2$.

Proposition 2.65

Soit u et v deux endomorphismes de E qui commutent ($u \circ v = v \circ u$). Les sous-espaces propres de l'un sont stables pour l'autre.

Démonstration : Soit λ une valeur propre de u , on veut montrer que $E_\lambda(u)$ est stable par v .

En effet, soit $x \in E_\lambda(u)$, on a

$$u(v(x)) = v(u(x)) = v(\lambda \cdot x) = \lambda \cdot v(x)$$

donc $v(x) \in E_\lambda(u)$.

□

ATTENTION

Dans le cas général, cela ne veut pas dire que l'image par v d'un vecteur propre pour u est un vecteur propre pour v .

6 Éléments propres d'une matrice

6.1 Définition

Définition 2.66 (Endomorphisme canoniquement associé à une matrice)

Soit $A \in \mathcal{M}_n(\mathbf{K})$. L'endomorphisme canoniquement associé à A est l'unique endomorphisme a de $\mathbf{K}^n$ défini tel que $\text{Mat}_{\text{can}}(a) = A$.

Remarques :

1. Un tel endomorphisme existe et est unique car l'application

$$\begin{aligned} \mathcal{L}(\mathbf{K}^n) &\rightarrow \mathcal{M}_n(\mathbf{K}) \\ u &\mapsto \text{Mat}_{\text{can}}(u) \end{aligned}$$

est un isomorphisme.

2. En identifiant $\mathbf{K}^n$ à $\mathcal{M}_{n,1}(\mathbf{K})$ on a

$$\begin{array}{ccc} a : \mathcal{M}_{n,1}(\mathbf{K}) & \rightarrow & \mathcal{M}_{n,1}(\mathbf{K}) \\ X & \mapsto & AX \end{array}$$

En effet les colonnes de A sont les images par a de la base canoniques. Ce sont bien les AX_i où X_i est le vecteur colonne avec un 1 à la ligne i et des 0 autre part.

3. On peut bien évidemment définir cela plus généralement pour $A \in \mathcal{M}_{n,p}(\mathbf{K})$ et on obtient $a \in \mathcal{L}(\mathbf{K}^n, \mathbf{K}^p)$.

Définition 2.67

Soit $A \in \mathcal{M}_n(\mathbf{K})$ et a l'endomorphisme canoniquement associé à A . Les valeurs propres, espaces propres et spectre de A sont ceux de a . On a donc

$$\forall \lambda \in \mathbf{K}, E_\lambda(A) = \text{Ker}(A - \lambda I_n) = \{X \in \mathcal{M}_{n,1}(\mathbf{K}) \mid AX = \lambda.X\}$$

$$\text{et } \text{Sp}(A) = \{\lambda \in \mathbf{K} \mid \exists X \in \mathcal{M}_{n,1}(\mathbf{K}), X \neq 0 \text{ et } AX = \lambda.X\}.$$

Remarque : les exemples précédents peuvent se réécrire avec cette terminologie, pour

$$A = \begin{pmatrix} 0 & 2 & -1 \\ 3 & -2 & 0 \\ -2 & 2 & 1 \end{pmatrix}$$

on a $\text{Sp}(A) = \{1, 2, -4\}$.

Proposition 2.68

Soit E un espace vectoriel de dimension finie et $u \in \mathcal{L}(E)$. Soit $\mathcal{B} = (e_1, \dots, e_n)$ une base de E . On note $A = \text{Mat}_{\mathcal{B}}(u)$. Les valeurs propres de u sont les valeurs propres de A : $\text{Sp}(u) = \text{Sp}(A)$.

De plus, pour λ une valeur propre de u . Soit $x \in E$, x est un vecteur propre de u associé à λ si et seulement si le vecteur colonne $X = \text{Mat}_{\mathcal{B}}(x)$ des coordonnées de x dans la base $\mathcal{B}$ est un vecteur propre de A

Démonstration : Soit $x \in E$ et $X = \text{Mat}_{\mathcal{B}}(x)$ le vecteur colonne des coordonnées de x dans la base $\mathcal{B}$. On considère $Y = \text{Mat}_{\mathcal{B}}(u(x)) = AX$ le vecteur colonne des coordonnées de $u(x)$ dans la base $\mathcal{B}$.

$$\begin{aligned} x \text{ est un vecteur propre pour } u \text{ associé à } \lambda & \iff u(x) = \lambda x \\ & \iff Y = \lambda X \\ & \iff AX = \lambda X \\ & \iff X \text{ est un vecteur propre pour } A \text{ associé à } \lambda \end{aligned}$$

□

Proposition 2.69

Soit A et B de matrices semblables de $\mathcal{M}_n(\mathbf{K})$. On a $\text{Sp}(A) = \text{Sp}(B)$

Démonstration : Soit $P \in GL_n(\mathbf{K})$ telle que $B = P^{-1}AP$. Soit $\lambda \in \mathbf{K}$,

$$\lambda \in \text{Sp}(B) \iff \text{rg}(B - \lambda I_n) < n \iff \text{rg}(P^{-1}(A - \lambda I_n)P) < n \iff \text{rg}(A - \lambda I_n) < n \iff \lambda \in \text{Sp}(A).$$

□

ATTENTION

Par contre elles n'ont pas les mêmes espaces propres. En effet B peut être obtenue à partir de A par un changement de bases. Pour déterminer les espaces propres de B à partir de ceux de A il faut aussi faire ce changement de base.

Remarque : Pour déterminer les valeurs propres d'un endomorphisme d'un espace vectoriel E de dimension finie, on calcule sa matrice A dans une base de $\mathcal{E}$ puis on détermine les valeurs propres de la matrice obtenue. Le choix de la base n'est pas important. On obtient de même les espaces propres en « traduisant » les espaces propres de A à l'aide de la base $\mathcal{E}$.

Exemple : Soit

$$\begin{aligned} \varphi : \mathbf{K}_2[X] &\rightarrow \mathbf{K}_2[X] \\ P &\mapsto (X^2 - 1)P'' + (2X + 1)P' \end{aligned}$$

On cherche ses valeurs propres / espaces propres.

La matrice de φ dans la base canonique est $A = \begin{pmatrix} 0 & 1 & 0 \\ 0 & 2 & 2 \\ 0 & 0 & 6 \end{pmatrix}$.

On en déduit directement que $\text{Sp}(A) = \text{Sp}(\varphi) = \{0, 2, 6\}$.

Les espaces propres de A sont

$$E_0(A) = \text{Vect} \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} ; \quad E_2(A) = \text{Vect} \begin{pmatrix} 1 \\ 2 \\ 0 \end{pmatrix} \text{ et } E_6(A) = \text{Vect} \begin{pmatrix} 1 \\ 6 \\ 12 \end{pmatrix}$$

Ce qui nous donne pour les espaces propres de φ

$$E_0(\varphi) = \text{Vect}(1) ; \quad E_2(\varphi) = \text{Vect}(2X + 1) \text{ et } E_6(\varphi) = \text{Vect}(12X^2 + 6X + 1)$$

Regardons maintenant une autre base de $\mathbf{K}_2[X]$. On pose $\mathcal{B} = (1, X - 1, (X - 1)^2)$. On a alors

$$B = \text{Mat}_{\mathcal{B}}(\varphi) = \begin{pmatrix} 0 & 3 & 0 \\ 0 & 2 & 10 \\ 0 & 0 & 6 \end{pmatrix}$$

En effet

$$\varphi((X-1)^2) = 2(X^2-1) + 2(2X+1)(X-1) = (X-1)[2(X+1) + 2(2X+1)] = (X-1)(6X+4) = 6(X-1)^2 + 10(X-1)$$

On a encore $\text{Sp}(B) = \{0, 2, 6\}$ et cette fois

$$E_0(B) = \text{Vect} \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} ; \quad E_2(A) = \text{Vect} \begin{pmatrix} 3 \\ 2 \\ 0 \end{pmatrix} \text{ et } E_6(A) = \text{Vect} \begin{pmatrix} 1 \\ 2 \\ \frac{4}{5} \end{pmatrix} = \text{Vect} \begin{pmatrix} 5 \\ 10 \\ 4 \end{pmatrix}$$

On retrouve les mêmes espaces propres pour φ .

Remarque : Le spectre d'une matrice peut dépendre du corps de base. Prenons l'exemple de la matrice $A = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}$ qui est la matrice de la rotation d'angle $\frac{\pi}{2}$. On regarde A comme un élément de $\mathcal{M}_2(\mathbf{R})$ et que l'on cherche son spectre (réel) à savoir les réels λ tels que $\text{rg}(A - \lambda I_2) < 2$. On a alors

$$\text{rg}(A - \lambda I_2) < 2 \iff \lambda^2 + 1 = 0$$

qui n'a pas de solution. La matrice n'a pas de valeurs propres réelles. Maintenant on peut supposer que $A \in \mathcal{M}_2(\mathbf{C})$ ce qui revient à accepter des valeurs complexes de λ . On trouve alors que $\text{Sp}(A) = \{i, -i\}$.

Proposition 2.70

Soit $\mathbf{K}'$ un sous-corps de $\mathbf{K}$ et $A \in \mathcal{M}_n(\mathbf{K}')$. Le spectre de A dans $\mathbf{K}'$ est inclus dans le spectre de A dans $\mathbf{K}$. On note alors

$$\text{Sp}_{\mathbf{K}'}(A) \subset \text{Sp}_{\mathbf{K}}(A).$$

6.2 Endomorphismes et matrices diagonalisables

Il a été vu en première année qu'il était facile de travailler avec des matrices diagonales (pour les puissances par exemples). On va souvent chercher une matrice diagonale d'un endomorphisme (si c'est possible).

Définition 2.71

1. Soit u un endomorphisme de E . Il est dit *diagonalisable* s'il existe une base $\mathcal{E}$ de E tel que la matrice de u dans la base $\mathcal{E}$ soit diagonale.
2. Une matrice A est dite *diagonalisable* si l'endomorphisme canoniquement associé à A est diagonalisable.

Proposition 2.72

L'endomorphisme u est diagonalisable si et seulement si la somme (directe) de ses sous-espaces propres vaut E :

$$\bigoplus_{\lambda \in \text{Sp}(u)} E_{\lambda}(u) = E.$$

Démonstration :

– $\Rightarrow$ On sait que $\bigoplus_{\lambda \in \text{Sp}(u)} E_{\lambda}(u) \subset E$.

On suppose que u est diagonalisable. On choisit une base $\mathcal{E} = (e_1, \dots, e_n)$ telle que la matrice de u de $\mathcal{E}$ soit diagonale. Quitte à réordonner les vecteurs de E on peut supposer que les

coefficients diagonaux égaux de la matrice soient consécutifs.

$$\text{Mat}_{\mathcal{E}}(u) = \begin{pmatrix} \lambda_1 & & & & \\ & \ddots & & & \\ & & \lambda_1 & & \\ & & & \ddots & \\ & & & & \lambda_p & \\ & & & & & \ddots \\ & & & & & & \lambda_p \end{pmatrix}$$

Soit w dans E , si on le décompose dans la base $\mathcal{E}$ puis que l'on regroupe les termes selon la valeur propre, on peut écrire w sous la forme $w = w_1 + \dots + w_p$ où $w_i \in E_{\lambda_i}(u)$. On a donc

$$\bigoplus_{\lambda \in \text{Sp}(u)} E_{\lambda}(u) = E.$$

- $\boxed{\Leftarrow}$ On suppose que la somme des espaces propres vaut E . On choisit alors une base $\mathcal{E}_i$ de chaque $E_{\lambda_i}(u)$. En concaténant ces bases on obtient une base dans laquelle la matrice de u est diagonale.

□

Proposition 2.73

Soit $A \in \mathcal{M}_n(\mathbf{K})$. Elle est diagonalisable si et seulement si elle est semblable à une matrice diagonale.

Démonstration :

- $\boxed{\Rightarrow}$ On suppose que A est diagonalisable, donc son endomorphisme canoniquement associé a de $\mathbf{K}^n$ dans lui même est diagonalisable. Il existe de ce fait une base $\mathcal{B}$ de $\mathbf{K}^n$ tel que $B = \text{Mat}_{\mathcal{B}}(a)$ soit diagonale. Maintenant en notant P la matrice de passage de la base canonique à $\mathcal{B}$ on a

$$B = P^{-1}AP$$

ce qui prouve que A est semblable à B qui est diagonale.

- $\boxed{\Leftarrow}$ On suppose que A est semblable à une matrice diagonale B . Il existe $P \in GL_n(\mathbf{K})$ telle que $B = P^{-1}AP$. Maintenant, si on considère la base $\mathcal{B}$ de $\mathbf{K}^n$ définie par P (dont les vecteurs sont les colonnes de P). Alors $\text{Mat}_{\mathcal{B}}(a) = B$ est diagonale donc a (et de ce fait A) est diagonalisable.

□

Exemples :

1. Soit $A = \begin{pmatrix} 1 & 4 \\ 2 & -1 \end{pmatrix}$. Elle est diagonalisable car elle est semblable à $\begin{pmatrix} 3 & 0 \\ 0 & -3 \end{pmatrix}$.
2. La matrice $A = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}$ est diagonalisable dans $\mathbf{C}$ mais pas dans $\mathbf{R}$.
3. La matrice $A = \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$. N'est pas diagonalisable. En effet son spectre est $\text{Sp}(A) = \{1\}$ et la seule matrice diagonale dont le spectre est juste $\{1\}$ est I_2 mais A n'est pas semblable à I_2 .

4. Soit F et G deux sous-espaces vectoriels supplémentaires de $E : F \oplus G = E$. On note p la projection sur F parallèlement à G et s la symétrie par rapport à F parallèlement à G .

On a déjà vu que pour la projection p , G était le noyau c'est-à-dire l'espace propre associé à la valeur propre 0 et F était l'espace propre associé à la valeur propre 1. On a donc $E = E_0(p) \oplus E_1(p)$. De ce fait p est diagonalisable. De même $E = E_{-1}(s) \oplus E_1(s)$ donc s est aussi diagonalisable.

Proposition 2.74

Soit u un endomorphisme de E . On note $\text{Sp}(u) = \{\lambda_1, \dots, \lambda_p\}$.

1. La somme des dimension des espaces propres est inférieure à la dimension de E :

$$\sum_{i=1}^p \dim E_{\lambda_i} \leq \dim E$$

2. Il y a égalité si et seulement si u est diagonalisable.

Avant de démontrer cette proposition, donnons un corollaire utile

Corollaire 2.75

Si E est de dimension n et que u admet n valeurs propres distinctes alors u est diagonalisable.

ATTENTION

Ce n'est qu'une condition suffisante !

Démonstration de la proposition :

1. On utilise juste que les espaces propres sont en somme directe :

$$\bigoplus_{i=1}^p E_{\lambda_i} \subset E$$

2. Il suffit de voir que

$$\begin{aligned} u \text{ diagonalisable} &\iff \bigoplus_{i=1}^p E_{\lambda_i} = E \\ &\iff \dim \bigoplus_{i=1}^p E_{\lambda_i} = \dim E \\ &\iff \sum_{i=1}^p \dim E_{\lambda_i} = \dim E \end{aligned}$$

□

Démonstration du corollaire : Il suffit de voir que les espaces propres sont de dimension au moins égale à 1. □

Intégration

1	Intégrale généralisée sur un intervalle de la forme $[a, +\infty[$	74
1.1	Définition	74
1.2	Propriétés	78
1.3	Théorème de comparaison pour les fonctions positives sur un intervalle de la forme $[a, +\infty[$	79
1.4	Intégrabilité sur un intervalle de la forme $[a, +\infty[$	80
1.5	Théorèmes de comparaison	83
1.6	Intégrales de référence	84
2	Intégration sur un intervalle quelconque	85
2.1	Intégrales convergentes	85
2.2	Propriétés	87
2.3	Intégrales des fonctions positives et fonctions intégrables	89
2.4	Méthodes de calculs	92
2.5	Intégration des relations de comparaisons	97
3	Les théorèmes de Lebesgue	100
3.1	Limite (simple) d'une suite de fonctions et d'une série de fonctions	100
3.2	Interversion de $\lim$ et $\int$	102
3.3	Théorème de convergence dominée	102
3.4	Théorème de convergence dominée - Cas continu	105
3.5	Intégration terme à terme d'une somme d'une série de fonctions	106

La construction de l'intégrale sur un segment¹ a été vue en première année. Cette année, nous allons généraliser cela à des intervalles non nécessairement fermés et non nécessairement bornés. Nous commencerons par le cas de l'intervalle $[a, +\infty[$ avant de traiter le cas des intervalles $[a, b[$ avec $b \in \mathbf{R}$, $]a, b]$ avec $a \in -\infty \cup \mathbf{R}$ et $]a, b[$ où $a \in -\infty \cup \mathbf{R}$ et $b \in \mathbf{R} \cup \{+\infty\}$.

On s'intéresse aux fonctions à valeurs dans $\mathbf{R}$ ou $\mathbf{C}$. On notera $\mathbf{K}$ pour $\mathbf{R}$ ou $\mathbf{C}$.

Commençons par définir les fonctions continues par morceaux.

1. c'est-à-dire un intervalle qui est de la forme $[a, b]$ avec $(a, b) \in \mathbf{R}^2$. Il est donc fermé et borné

Définition 3.1 (Fonctions continues par morceaux)

1. Soit $J = [\alpha, \beta]$ un segment (intervalle fermé borné) et f une fonction définie sur J à valeurs dans $\mathbf{K}$. Elle est continue par morceaux sur J s'il existe une subdivision ($\alpha = x_0 < x_1 < \dots < x_n = \beta$) telle que pour tout $i \in \llbracket 1; n \rrbracket$
 - La restriction de f à $]x_{i-1}, x_i[$ est continue
 - Les limites $\lim_{x \rightarrow x_{i-1}^+} f$ et $\lim_{x \rightarrow x_i^-} f$ existent et sont finies.
2. Soit I un intervalle quelconque et f une fonction définie sur I à valeurs dans $\mathbf{K}$. Elle est continue par morceaux sur I si pour tout segment J inclus dans I , $f|_J$ est continue par morceaux.

Notation : On note $\mathcal{CM}(I, \mathbf{K})$ l'ensemble des fonctions continues par morceaux sur I à valeurs dans $\mathbf{K}$.

Exemples :

1. La fonction $f : x \mapsto \lfloor x \rfloor$ est continue par morceaux sur $\mathbf{R}$. En effet sur chaque segment $J = [a, b]$, la fonction f n'a qu'un nombre fini de points de discontinuité.
2. La fonction $x \mapsto x^2 \lfloor x \rfloor$ n'est pas continue par morceaux sur $\mathbf{R}$ car sur $[0, 1]$ elle a une infinité de points de discontinuité $\left(x = \frac{1}{2}, \frac{1}{3}, \dots, \frac{1}{n}\right)$.
3. La fonction $x \mapsto \frac{1}{x-1}$ est continue par morceaux sur $[0, 1[$. En effet elle est continue sur tout $[a, b]$ avec $b < 1$.
4. la fonction f définie sur $[0, 1]$ par

$$f : x \mapsto \begin{cases} \frac{1}{x-1} & \text{si } x < 1 \\ 5 & \text{si } x = 1 \end{cases}$$

n'est pas continue par morceaux sur $[0, 1]$.

5. La fonction $x \mapsto \sin\left(\frac{1}{x}\right)$ est continue par morceaux sur $]0, +\infty[$

ATTENTION

Faire attention à la différence entre les fonctions continues par morceaux sur $[a, b]$ et sur $[a, b[$

Remarque culturelle : Le cadre de l'intégration en CPGE est celui de l'intégrale des fonctions continues par morceaux. C'est ce qui s'appelle l'intégrale de Riemann. Ce cadre a l'avantage d'être simple mais peut prêter à certains aspects. On verra par exemple qu'une suite de fonctions continues par morceaux peut converger vers une fonction qui n'est pas continue par morceaux. Il existe une autre théorie de l'intégration due à Henri Lebesgue. Il est fort probable que vous l'étudierez en école.

Matheux (Henri Lebesgue : 1875 - 1941)



Henri-Léon Lebesgue (1875-1941), plus connu sous le nom de Henri Lebesgue, est l'un des grands mathématiciens français de la première moitié du xx^e siècle. Henri Lebesgue a révolutionné et généralisé le calcul intégral. Sa théorie de l'intégration (1902-1904) est extrêmement commode d'emploi, et répond aux besoins des physiciens. En effet, elle permet de rechercher et de prouver l'existence de primitives pour des fonctions « irrégulières » et recouvre différentes théories antérieures qui en sont des cas particuliers.

On lui doit aussi la transformée de Fourier établie dans la fin des années 1930.

1 Intégrale généralisée sur un intervalle de la forme $[a, +\infty[$

Dans ce paragraphe $a \in \mathbb{R}$.

1.1 Définition

Définition 3.2 (Intégrale convergente)

Soit f une fonction continue par morceaux sur $[a, +\infty[$ à valeurs dans $\mathbb{K}$. On dit que l'intégrale $\int_a^{+\infty} f$ converge si $\lim_{x \rightarrow +\infty} \int_a^x f$ existe et est finie. Dans ce cas on note $\int_a^{+\infty} f$ cette limite.

Remarque : Dans la définition ci-dessus, pour tout $x \in [a, +\infty[$, la fonction f est continue par morceaux sur le segment $[a, x]$ et donc $\int_a^x f$ existe.

ATTENTION

Dans le cas où $\lim_{x \rightarrow +\infty} \int_a^x f$ n'existe pas ou n'est pas finie, on dit que $\int_a^{+\infty} f$ diverge mais dans ce cas, le symbole $\int_a^{+\infty} f$ n'est pas un nombre !

Remarque : Le fait d'être convergente ou divergente est la nature de l'intégrale.

Proposition 3.3 (Relation de Chasles)

Soit $f \in \mathcal{CM}([a, +\infty[, \mathbf{K})$.

Pour tout $a' \in [a, +\infty[$, $\int_a^{+\infty} f$ converge si et seulement $\int_{a'}^{+\infty} f$ converge. Dans ce cas,

$$\int_a^{+\infty} f = \int_a^{a'} f + \int_{a'}^{+\infty} f$$

où $\int_a^{a'} f$ existe car f est continue par morceaux sur le segment $[a, a']$.

Remarque : Cela signifie que pour tester la convergence d'une intégrale, on peut se placer « près » de $+\infty$. Ce résultat est analogue à celui qui dit que pour tester la convergence d'une série numériques on peut ne pas considérer les premiers termes.

Démonstration : Pour tout $x \geq a'$ on a $\int_a^x f = \int_a^{a'} f + \int_{a'}^x f$. On en déduit que $x \mapsto \int_a^x f$ a une limite finie en $+\infty$ si et seulement $x \mapsto \int_{a'}^x f$ a une limite finie en $+\infty$. De plus, dans ce cas,

$$\int_a^{+\infty} f = \lim_{x \rightarrow +\infty} \int_a^x f = \lim_{x \rightarrow +\infty} \left(\int_a^{a'} f + \int_{a'}^x f \right) = \int_a^{a'} f + \lim_{x \rightarrow +\infty} \int_{a'}^x f = \int_a^{a'} f + \int_{a'}^{+\infty} f$$

□

Exemples :

1. On regarde la fonction $f : t \mapsto 1$. L'intégrale $\int_0^{+\infty} f$ diverge car $\int_0^x f = x \xrightarrow{x \rightarrow +\infty} +\infty$.

2. Soit $f : t \mapsto e^{-at}$ où $a > 0$ (sinon l'intégrale diverge clairement). On a pour $x \geq 0$,

$$\int_0^x e^{-at} dt = \left[\frac{e^{-at}}{-a} \right]_0^x = \frac{1}{a} (1 - e^{-ax}) \xrightarrow{x \rightarrow +\infty} \frac{1}{a}$$

On en déduit que $\int_0^{+\infty} e^{-at} dt$ converge et $\int_0^{+\infty} e^{-at} dt = \frac{1}{a}$.

3. Soit $f : t \mapsto \frac{1}{t}$ sur $[1, +\infty[$. Pour $x \geq 1$,

$$\int_1^x f = \int_1^x \frac{dt}{t} = [\ln t]_1^x = \ln(x) \xrightarrow{x \rightarrow +\infty} +\infty$$

On en déduit que $\int_1^{+\infty} \frac{dt}{t}$ diverge.

4. Soit $f : t \mapsto t^{-\alpha}$ pour $\alpha \in \mathbf{R} \setminus \{1\}$ sur $[1, +\infty[$. On a

$$\int_1^x \frac{1}{t^\alpha} dt = \left[\frac{t^{1-\alpha}}{1-\alpha} \right]_1^x$$

— Si $\alpha > 1$ on a donc

$$\int_1^x \frac{1}{t^\alpha} dt = \frac{1}{\alpha-1} \left(1 - \frac{1}{x^{\alpha-1}} \right) \xrightarrow{x \rightarrow +\infty} \frac{1}{\alpha-1}$$

cela montre que $\int_1^{+\infty} \frac{dt}{t^\alpha}$ converge et que $\int_1^{+\infty} \frac{dt}{t^\alpha} = \frac{1}{\alpha-1}$.

— Si $\alpha < 1$ on a donc

$$\int_1^x \frac{1}{t^\alpha} dt = \frac{1}{1-\alpha} (x^{1-\alpha} - 1) \xrightarrow{x \rightarrow +\infty} +\infty$$

cela montre que $\int_1^{+\infty} \frac{dt}{t^\alpha}$ diverge.

Remarque : Les trois dernières sont les intégrales de référence les plus souvent utilisées. Elles seront écrites plus loin.

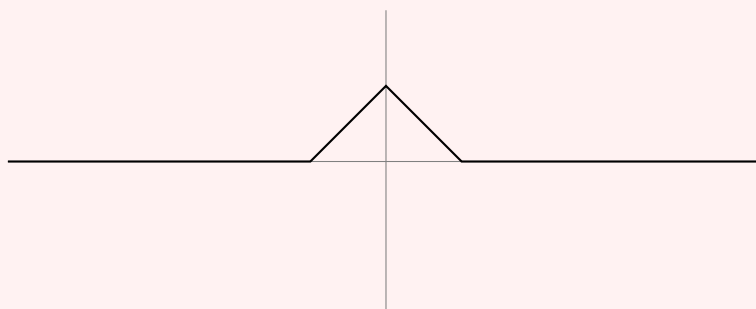
Exercice : Déterminer la nature de $\int_2^{+\infty} \frac{dt}{t \ln t}$

ATTENTION

Contrairement au cas des suites, la convergence de $\int_0^{+\infty} f$ n'implique pas que $\lim_{x \rightarrow +\infty} f(x) = 0$.

On considère la fonction $\delta : \mathbb{R} \rightarrow \mathbb{R}$ définie par

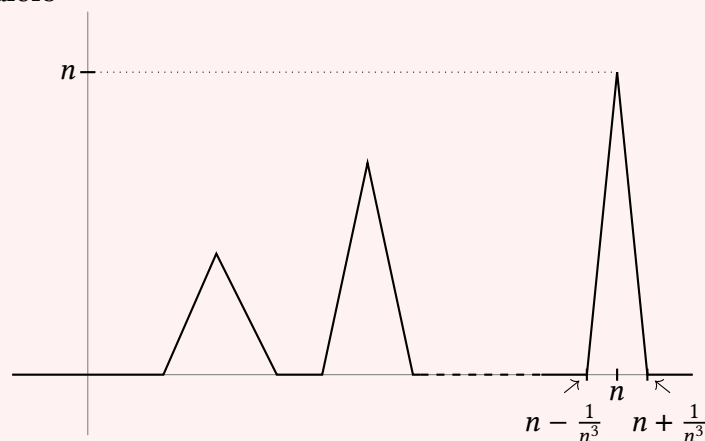
$$\delta : x \mapsto \begin{cases} 0 & \text{si } |x| \geq 1 \\ x + 1 & \text{si } x \in]-1, 0] \\ 1 - x & \text{si } x \in]0, 1[\end{cases}$$



On définit alors la fonction f par

$$f : x \mapsto \sum_{n=2}^{+\infty} n \delta(n^3(x - n))$$

Par définition de la fonction δ , $\delta(n^3(x - n)) \neq 0$ si et seulement si $x \in [n - \frac{1}{n^3}, n + \frac{1}{n^3}]$ ce qui prouve que la série converge car, au plus, un des terme de la série n'est pas nul. La courbe de f est alors



On en déduit que pour tout $x \geq 0$,

$$\int_0^x f \leq \sum_{n=2}^{+\infty} \frac{n \times \frac{2}{n^3}}{2} = \sum_{n=2}^{+\infty} \frac{1}{n^2} < +\infty$$

On a bien que l'intégrale $\int_0^{+\infty} f$ converge alors que f ne tend pas 0 en $+\infty$ (elle n'est même pas bornée au voisinage de $+\infty$).

1.2 Propriétés

Proposition 3.4 (Linéarité)

L'ensemble des fonctions continues par morceaux sur $[a, +\infty[$ à valeurs dans $\mathbf{K}$ dont l'intégrale converge est un sous-espace vectoriel de $\mathcal{CM}([a, +\infty[, \mathbf{K})$ et l'application qui associe à une telle fonction son intégrale est linéaire.

Démonstration : Notons F l'ensemble des fonctions continues par morceaux sur $[a, +\infty[$ dont l'intégrale converge.

- L'ensemble F n'est pas vide car la fonction nulle appartient à F .
- Soit $(f, g) \in F^2$ et $(\lambda, \mu) \in \mathbf{K}^2$. On pose $h = \lambda f + \mu g$. On a alors pour tout x ,

$$\int_a^x h = \lambda \int_a^x f + \mu \int_a^x g$$

Maintenant, par hypothèses, $\lambda \int_a^x f \xrightarrow{x \rightarrow +\infty} \lambda \int_a^{+\infty} f$ et $\mu \int_a^x g \xrightarrow{x \rightarrow +\infty} \mu \int_a^{+\infty} g$.

Donc

$$\int_a^x h \xrightarrow{x \rightarrow +\infty} \lambda \int_a^{+\infty} f + \mu \int_a^{+\infty} g$$

On en déduit que $h \in F$ et que $\int_a^{+\infty} h = \lambda \int_a^{+\infty} f + \mu \int_a^{+\infty} g$.

□

Proposition 3.5 (Positivité - Croissance)

1. Soit f une fonction continue par morceaux sur $[a, +\infty[$ réelle positive dont l'intégrale converge.

On a $\int_a^{+\infty} f \geq 0$.

2. Soit f une fonction **continue** sur $[a, +\infty[$ réelle et positive dont l'intégrale converge.

Si $\int_a^{+\infty} f = 0$ alors f est la fonction nulle.

3. Soit f et g deux fonctions réelles continues par morceaux sur $[a, +\infty[$ dont l'intégrale converge.

Si $f \geq g$ alors $\int_a^{+\infty} f \geq \int_a^{+\infty} g$.

Démonstration :

1. Pour tout $x \in [a, +\infty[$, $\int_a^x f \geq 0$ car f est positive et $a \leq x$. Par passage à la limite on en déduit que $\int_a^{+\infty} f \geq 0$.
2. On sait que comme f est positive, $x \mapsto \int_a^x f$ est croissante. De plus, pour tout x , $\int_a^x f \geq 0$ et $\lim_{x \rightarrow +\infty} \int_a^x f = 0$. On en déduit que pour tout x de $[a, +\infty[$, $\int_a^x f = 0$ et donc, comme f est continue $f|_{[a, x]}$ est la fonction nulle. Cela implique que f est la fonction nulle sur $[a, +\infty[$.

3. Par hypothèse $f - g$ est positive. Cela implique que $\int_a^{+\infty} f - g \geq 0$ d'après le point 1. On peut alors conclure par linéarité.

□

Remarque : L'assertion 2. est fausse si on suppose juste f continue par morceaux. Il suffit par exemple de considérer la fonction f définie sur $\mathbf{R}_+$ par

$$f : x \mapsto \begin{cases} 1 & \text{si } x \in \mathbf{N} \\ 0 & \text{sinon} \end{cases}$$

Corollaire 3.6

Soit f une fonction continue $[a, +\infty[$ telle que $\int_a^{+\infty} f$ converge. On considère $F : x \mapsto \int_x^{+\infty} f$. Elle est dérivable et pour tout x de $[a, +\infty[$, $F(x) = -f(x)$.

Démonstration : On sait que comme f est continue, $x \mapsto \int_a^x f$ est dérivable et de dérivée f . De plus, pour tout x de $[a, +\infty[$, $F(x) + \int_a^x f = \int_a^{+\infty} f$ et donc $F(x) = \int_a^{+\infty} f - \int_a^x f$. □

Remarques :

1. Ce terme est l'analogue du **reste** dans le cas d'une série convergente. En particulier

$$\lim_{x \rightarrow +\infty} \int_x^{+\infty} f = 0$$

2. Cela signifie que $x \mapsto -\int_x^{+\infty} f$ est la primitive de f qui tend vers 0 « en $+\infty$ ». Par exemple pour $f : t \mapsto \frac{1}{t^2}$. Sa primitive qui s'annule en $+\infty$ est $x \mapsto -\frac{1}{x} = -\int_x^{+\infty} \frac{dt}{t^2}$.

1.3 Théorème de comparaison pour les fonctions positives sur un intervalle de la forme $[a, +\infty[$

Comme pour les séries, pour déterminer si une intégrale converge on va souvent se ramener à étudier la convergence absolue. On étudie donc essentiellement des fonctions positives. On peut alors développer des arguments similaires à l'étude de la convergence des séries à termes positifs.

Proposition 3.7 (Convergence de l'intégrale d'une fonction positive)

1. Si f est une fonction **positive** continue par morceaux. La fonction $F : x \mapsto \int_a^x f$ est croissante. De ce fait, $\int_a^{+\infty} f$ converge si et seulement si F est majorée.
2. Si f et g sont deux fonctions **positives** continues par morceaux. On suppose que $0 \leq f \leq g$. Si $\int_a^{+\infty} g$ converge alors $\int_a^{+\infty} f$ converge.

Remarque : C'est une proposition qui est analogue à une proposition sur les séries à termes positifs

Démonstration :

1. Soit $x \geq x'$ on a

$$F(x) = \int_a^x f = \int_a^{x'} f + \underbrace{\int_{x'}^x f}_{\geq 0} \geq F(x')$$

La fonction F est donc croissante. De ce fait, elle est majorée si et seulement si elle admet une limite finie en $+\infty$.

2. Si $0 \leq f \leq g$ alors pour tout x de $[a, +\infty[$,

$$\int_a^x f \leq \int_a^x g \leq \int_a^{+\infty} g$$

On en déduit que $x \mapsto \int_a^x f$ est bornée et donc $\int_a^{+\infty} f$ converge.

□

Exemple : On considère $f : t \mapsto e^{-t^2}$ définie sur $[1, +\infty[$. On voit que pour tout $t \in [1, +\infty[$,

$$0 \leq e^{-t^2} \leq e^{-t}$$

Or on a vu en exemple que l'intégrale $\int_1^{+\infty} e^{-t} dt$ converge donc $\int_1^{+\infty} e^{-t^2} dt$ converge aussi.

ATTENTION

Comme dans le cas des séries, quand on étudie une intégrale $\int_a^b f$ où f est une fonction continue par morceaux réelle **positive**. On peut toujours calculer l'intégrale $\int_a^b f$ dans $[0, +\infty]$ et l'intégrale converge si et seulement si l'intégrale a une valeur finie.

1.4 Intégrabilité sur un intervalle de la forme $[a, +\infty[$

Le fait que l'intégrale $\int_a^{+\infty} f$ converge peut se voir comme le fait qu'une série $\sum u_n$ converge. Dans beaucoup de cas on va vouloir montrer la convergence de l'intégrale via des majorations ce qui implique de travailler avec des fonctions réelles positives. Pour cela nous allons avoir besoin d'une notion qui généralise la convergence absolue des séries.

Définition 3.8 (Fonction intégrable)

Soit f une fonction continue par morceaux sur $[a, +\infty[$. On dit que f est intégrable sur $[a, +\infty[$ ou que $\int_a^{+\infty} f$ converge **absolument** si $\int_a^{+\infty} |f|$ converge.

On note $L^1([a, +\infty[, \mathbb{K})$ l'ensemble des fonctions continues par morceaux sur $[a, +\infty[$ à valeurs dans $\mathbb{K}$ qui sont intégrables.

Remarques :

1. Si f est continue par morceaux sur $[a, +\infty[$ alors $|f|$ aussi.
2. Dans le cas où f est une fonction positive, f intégrable est équivalent au fait que $\int_a^{+\infty} f$ converge.

Exemple : Soit $f : t \mapsto \frac{e^{it}}{t^2}$. La fonction f est intégrable sur $[1, +\infty[$. En effet $|f| : t \mapsto \left| \frac{e^{it}}{t^2} \right| = \frac{1}{t^2}$ et $\int_1^{+\infty} \frac{1}{t^2} dt$ converge car $2 > 1$.

Théorème 3.9

Soit f une fonction continue par morceaux sur $[a, +\infty[$. Si l'intégrale $\int_a^{+\infty} f$ est absolument convergente alors l'intégrale $\int_a^{+\infty} f$ converge.

Remarque : C'est l'analogue du théorème qui dit que si une série est absolument convergente alors elle est convergente.

Démonstration : Il suffit d'adapter la preuve du théorème analogue pour les séries.

- On suppose que f est à valeurs réelles. On pose $f_+ = \frac{f + |f|}{2}$ et $f_- = \frac{|f| - f}{2}$. De sorte que $f = f_+ - f_-$.

Maintenant on a $0 \leq f_+ \leq |f|$ et $0 \leq f_- \leq |f|$ car $-|f| \leq f \leq |f|$. On en déduit que $\int_a^{+\infty} f_+$ et $\int_a^{+\infty} f_-$ convergent et donc par linéarité $\int_a^{+\infty} f$ aussi.

- On suppose maintenant que f est éventuellement à valeurs complexes. On pose $f = f_1 + if_2$ avec f_1 et f_2 à valeurs réelles. On a alors

$$|f_1| \leq |f| \text{ et } |f_2| \leq |f|.$$

En effet soit $z = \alpha + i\beta \in \mathbb{C}$, on a $|\alpha| = \sqrt{\alpha^2} \leq \sqrt{\alpha^2 + \beta^2} = |z|$. De même $|\beta| \leq |z|$. De ce fait, la convergence de $\int_a^{+\infty} |f|$ implique de f_1 et f_2 sont intégrables sur $[a, +\infty[$ et donc, en appliquant le cas réel ci-dessus on obtient que $\int_a^{+\infty} f_1$ et $\int_a^{+\infty} f_2$ convergent. On en déduit que $\int_a^{+\infty} f$ converge par linéarité.

□

ATTENTION

Comme dans le cas des séries, la réciproque du théorème précédent est fausse. Il existe des fonctions f telles que $\int_a^{+\infty} f$ converge mais $\int_a^{+\infty} |f|$ ne converge pas. On dit dans ce cas que l'intégrale est semi-convergente.

Exemple : Étudions, $f : t \mapsto \frac{\sin t}{t}$ sur $[\pi, +\infty[$.

– Montrons que $\int_{\pi}^{+\infty} f$ converge :

On va utiliser le théorème spécial pour les séries alternées. On pose $u_n = \int_{n\pi}^{(n+1)\pi} \frac{\sin t}{t} dt$.

On sait que le signe de $\frac{\sin t}{t}$ sur $[n\pi, (n+1)\pi]$ dépend de la parité de n . On a $u_n \geq 0 \iff n$ pair.

De plus

$$|u_{n+1}| - |u_n| = \int_{(n+1)\pi}^{(n+2)\pi} \frac{|\sin t|}{t} dt - \int_{n\pi}^{(n+1)\pi} \frac{|\sin t|}{t} dt$$

car $t \mapsto \frac{\sin t}{t}$ garde un signe constante sur les intervalles d'intégration. On a donc en posant $v = t - \pi$ dans la première intégrale

$$|u_{n+1}| - |u_n| = \int_{n\pi}^{(n+1)\pi} \frac{|\sin(v + \pi)|}{v + \pi} dv - \int_{n\pi}^{(n+1)\pi} \frac{|\sin t|}{t} dt = \int_{n\pi}^{(n+1)\pi} |\sin t| \left(\frac{1}{t + \pi} - \frac{1}{t} \right) dt \leq 0.$$

Pour finir,

$$|u_n| \leq \int_{n\pi}^{(n+1)\pi} \frac{1}{t} dt = [\ln t]_{n\pi}^{(n+1)\pi} = \ln \left(\frac{n+1}{n} \right) \xrightarrow{n \rightarrow +\infty} 0$$

On peut donc appliquer le théorème spécial pour les séries alternées ce qui nous donne que la série $\sum u_n$ converge.

On remarque que pour tout entier $n \geq 1$: $\sum_{k=1}^n u_k = \int_{\pi}^{(n+1)\pi} f$.

Donc si on pose $F : x \mapsto \int_{\pi}^x f(t) dt$ on a que $(F(n\pi))_{n \in \mathbb{N}^*}$ converge vers une limite que l'on peut noter L .

Maintenant pour tout $x \geq \pi$, on note p l'entier tel que $p\pi \leq x < (p+1)\pi$. C'est-à-dire $p = \lfloor x/\pi \rfloor$.

On peut donc écrire

$$F(x) = \int_{\pi}^{p\pi} f + \int_{p\pi}^x f = F(p\pi) + \int_{p\pi}^x f$$

D'après ce qui précède, quand $x \rightarrow +\infty$, alors $p \rightarrow +\infty$ et donc $F(p\pi) \rightarrow L$.

Quant à l'autre terme, on a

$$\left| \int_{p\pi}^x f \right| \leq \int_{p\pi}^{(p+1)\pi} \frac{1}{t} dt = \ln \left(\frac{p+1}{p} \right)$$

Il tend donc vers 0 quand x (et donc p) tend vers $+\infty$.

En conclusion, $\int_{\pi}^{+\infty} f$ converge (et $\int_{\pi}^{+\infty} f = L$).

Remarque : Nous verrons un peu plus loin une autre méthode un peu plus simple pour retrouver ce résultat.

– Montrons que f n'est pas intégrable sur $[\pi, +\infty[$. En reprenant les notations ci-dessus, on a pour $x \geq \pi$

$$\int_{\pi}^x |f| \geq \int_{\pi}^{p\pi} |f| = \sum_{k=1}^{p-1} v_k \text{ où } v_k = \int_{k\pi}^{(k+1)\pi} \frac{|\sin t|}{t} dt.$$

Or,

$$v_k \geq \frac{1}{(k+1)\pi} \int_{k\pi}^{(k+1)\pi} |\sin t| dt = \frac{1}{(k+1)\pi} \int_0^\pi \sin t dt = \frac{1}{(k+1)\pi}.$$

On en déduit que la série $\sum v_k$ diverge et donc $\int_\pi^{+\infty} |f|$ aussi.

1.5 Théorèmes de comparaison

Proposition 3.10 (Théorème de comparaison)

Soit f et g deux fonctions continues par morceaux sur $[a, +\infty[$.

1. Si $f = O_{+\infty}(g)$ alors l'intégrabilité de g implique l'intégrabilité de f .
2. Si $f \sim_{+\infty} g$ alors l'intégrabilité de g est équivalente à l'intégrabilité de f .

Remarques :

1. Si $f = o_{+\infty}(g)$ alors on a $f = O_{+\infty}(g)$ de ce fait l'intégrabilité de g implique celle de f .
2. Si pour tout $x \in [a, +\infty[$, $|f(x)| \leq |g(x)|$ alors $f = O_{+\infty}(g)$ de ce fait l'intégrabilité de g implique celle de f .

Démonstration :

1. Si $f = O_{+\infty}(g)$, il existe un voisinage $[a', +\infty[$ de $+\infty$ et une constante C tels que sur $[a', +\infty[$ $|f| \leq C \cdot |g|$. Comme g est intégrable, $\int_{a'}^{+\infty} |g|$ converge et donc $\int_{a'}^{+\infty} |f|$ converge par comparaison pour les fonctions positives. On en déduit que $\int_a^{+\infty} |f|$ converge et donc que f est intégrable sur $[a, +\infty[$.
2. Si $f \sim_{+\infty} g$ alors $f = O_{+\infty}(g)$ donc $g \in L^1([a, +\infty[, \mathbb{K}) \Rightarrow f \in L^1([a, +\infty[, \mathbb{K})$. De même $g = O_{+\infty}(f)$ donc $f \in L^1([a, +\infty[, \mathbb{K}) \Rightarrow g \in L^1([a, +\infty[, \mathbb{K})$. On en déduit que l'intégrabilité de g est équivalente à l'intégrabilité de f .

□

ATTENTION

Comme dans le cas des séries, ce résultat porte sur l'absolue convergence et n'est plus vrai dans le cas de la convergence. Posons par exemple $f : t \mapsto \frac{(-1)^{[t]}}{\sqrt{t}}$ et $g : t \mapsto \frac{1}{t}$. On a

$g = o_{+\infty}(f)$ et $\int_1^{+\infty} f$ converge^a mais $\int_1^{+\infty} g$ diverge.

a. Cela se démontre en procédant comme pour $\int_1^{+\infty} \frac{\sin t}{t} dt$.

Exemples :

1. Déterminons la nature de $\int_1^{+\infty} \sin\left(\frac{1}{t^2}\right) dt$.

La fonction $t \mapsto \sin\left(\frac{1}{t^2}\right)$ est continue par morceaux sur $[1, +\infty[$. Comme $\sin x \underset{0}{\sim} x$, $\sin\left(\frac{1}{t^2}\right) \underset{+\infty}{\sim} \frac{1}{t^2}$. Or $t \mapsto \frac{1}{t^2}$ est intégrable sur $[1, +\infty[$ (vu en exemple précédemment et sera repris en fonctions de références) donc l'intégrale $\int_1^{+\infty} \sin\left(\frac{1}{t^2}\right) dt$ est absolument convergente donc convergente.

2. On peut retrouver le fait que $\int_{\pi}^{+\infty} \frac{\sin t}{t} dx$ est convergente.

On sait que l'on peut majorer $|\sin t|$ par 1 mais cela ne permet pas de conclure car $\int_{\pi}^{+\infty} \frac{1}{t} dt$ est une intégrale divergente. On va avoir recours à une intégration par parties.

Pour $x > 0$,

$$\int_{\pi}^x \frac{\sin t}{t} dt = \left[\frac{-\cos t}{t} \right]_{\pi}^x - \int_{\pi}^x \frac{\cos t}{t^2} dt = -\frac{\cos x}{x} - \frac{1}{\pi} - \int_{\pi}^x \frac{\cos t}{t^2} dt$$

Maintenant la fonction $t \mapsto \frac{\cos t}{t^2}$ est intégrable $[\pi, +\infty[$ car $\left| \frac{\cos t}{t^2} \right| \leq \frac{1}{t^2}$ et que $t \mapsto \frac{1}{t^2}$ est intégrable sur $[\pi, +\infty[$. Cela implique que l'intégrale $\int_{\pi}^{+\infty} \frac{\cos t}{t^2} dt$ est absolument convergente donc convergente. Cela signifie que la fonction $x \mapsto \int_{\pi}^x \frac{\cos t}{t^2} dt$ a une limite finie quand x tend vers $+\infty$.

De plus $\frac{\cos x}{x} \xrightarrow{x \rightarrow +\infty} 0$ donc la fonction $x \mapsto \int_{\pi}^x \frac{\sin t}{t} dt$ a une limite finie quand $x \rightarrow +\infty$ c'est-à-dire que $\int_{\pi}^{+\infty} \frac{\sin t}{t} dt$ converge et même

$$\int_{\pi}^{+\infty} \frac{\sin t}{t} dt = -\frac{1}{\pi} - \int_{\pi}^{+\infty} \frac{\cos t}{t^2} dt$$

1.6 Intégrales de référence

Proposition 3.11 (Fonctions de référence)

1. La fonction $t \mapsto e^{-at}$ où $a \in \mathbf{R}$ est intégrable sur $[0, +\infty[$ si et seulement si $a > 0$. Dans ce cas l'intégrale converge absolument et

$$\int_0^{+\infty} e^{-at} dt = \frac{1}{a}$$

2. La fonction $t \mapsto \frac{1}{t^{\alpha}}$ est intégrable sur $[1, +\infty[$ si et seulement si $\alpha > 1$. Dans ce cas

$$\int_1^{+\infty} \frac{1}{t^{\alpha}} dt = \frac{1}{\alpha - 1}$$

Démonstration :

1. Voir ci-dessus.
2. Voir ci-dessus.

□

Remarques :

1. Il faut connaître les conditions de convergence. Les valeurs des intégrales sont moins importantes et se retrouvent facilement.
2. On peut généraliser la première en prenant $\omega \in \mathbb{C}$. En effet pour $t \in \mathbb{R}$, $|e^{\omega t}| = e^{\Re(\omega)t}$. On en déduit que

La fonction $t \mapsto e^{\omega t}$ où $\omega \in \mathbb{C}$ est intégrable sur $[0, +\infty[$ si et seulement $\Re(\omega) < 0$.

De plus, le calcul ci-dessus reste correct et

$$\int_0^{+\infty} e^{\omega t} dt = -\frac{1}{\omega}$$

★ Méthode : Croissance de $x^\alpha f$:

Le but de cette méthode est de montrer la convergence d'une intégrale sur $[a, +\infty[$ en comparant aux intégrales $\int_a^{+\infty} \frac{1}{t^\alpha} dt$.

S'il existe $\alpha > 1$ tel que $t^\alpha f(t)$ est bornée au voisinage de $+\infty$ (ce qui est vrai si $t^\alpha f(t) \xrightarrow[t \rightarrow +\infty]{} \ell \in \mathbb{R}$) alors $f(t) = O\left(\frac{1}{t^\alpha}\right)$. Comme $t \mapsto \frac{1}{t^\alpha}$ est intégrable sur $[1, +\infty[$, la fonction f aussi. Finalement l'intégrale $\int_1^{+\infty} f$ est absolument convergente donc convergente.

Exemple : On cherche la nature de $\int_2^{+\infty} \ln(t)e^{-t} dt$. On voit que $t^2 \ln(t)e^{-t} \xrightarrow[t \rightarrow +\infty]{} 0$ par croissances comparées. De ce fait $\ln(t)e^{-t} = o_{+\infty}\left(\frac{1}{t^2}\right)$ et donc $\ln(t)e^{-t} = O_{+\infty}\left(\frac{1}{t^2}\right)$. Comme de $t \mapsto \frac{1}{t^2}$ est intégrable sur $[2, +\infty[$, $t \mapsto \ln(t)e^{-t}$ aussi et donc $\int_2^{+\infty} \ln(t)e^{-t} dt$ est absolument convergente donc convergente.

Exercice : Déterminer la nature des intégrales suivantes :

$$\int_1^{+\infty} \frac{\cos t}{t^2} dt ; \int_0^{+\infty} x \sin x e^{-x} dx ; \int_0^{+\infty} \frac{t^2 e^{-\sqrt{t}}}{1+t} dt$$

2 Intégration sur un intervalle quelconque

2.1 Intégrales convergentes

On va reprendre les définitions vues ci-dessus dans le cas des intégrales sur des intervalles semi-ouverts de la forme $]a, b]$ où $a \in \mathbb{R} \cup \{-\infty\}$ et $b \in \mathbb{R}$ ou de la forme $[a, b[$ où $a \in \mathbb{R}$ et $b \in \mathbb{R} \cup \{+\infty\}$ ou des intervalles ouverts de la forme $]a, b[$ avec $a \in \mathbb{R} \cup \{-\infty\}$ et $b \in \mathbb{R} \cup \{+\infty\}$.

Définition 3.12 (Intégrale convergente)

1. Soit f une fonction continue par morceaux sur $]a, b]$ où $a \in \mathbf{R} \cup \{-\infty\}$ et $b \in \mathbf{R}$ et à valeurs dans $\mathbf{K}$. On dit que l'intégrale $\int_a^b f$ converge si $\lim_{x \rightarrow a} \int_x^b f$ existe et est finie. Dans ce cas on note $\int_a^b f$ cette limite.
2. Soit f une fonction continue par morceaux sur $[a, b[$ où $a \in \mathbf{R}$ et $b \in \mathbf{R}$ et à valeurs dans $\mathbf{K}$. On dit que l'intégrale $\int_a^b f$ converge si $\lim_{x \rightarrow b} \int_a^x f$ existe et est finie. Dans ce cas on note $\int_a^b f$ cette limite.
3. Soit f une fonction continue par morceaux sur $]a, b[$ avec $a \in \mathbf{R} \cup \{-\infty\}$ et $b \in \mathbf{R} \cup \{+\infty\}$ et à valeurs dans $\mathbf{K}$. On dit que l'intégrale $\int_a^b f$ converge s'il existe $c \in]a, b[$ tel que les intégrales $\int_c^b f$ et $\int_a^c f$ converge.

Remarques :

1. Dans le cas contraire, on dit encore que l'intégrale diverge. On parle encore de la nature de l'intégrale.
2. Dans le cas 2, si $b = +\infty$ on retrouve la définition vue ci-dessus.
3. Dans le cas d'un intervalle ouvert de la forme $]a, b[$, en utilisant la relation de Chasles on voit que la nature et, le cas échéant la valeur, ne dépend pas du réel c que l'on prend. En effet pour $(c, c') \in]a, b[$ on a

$$\left(\int_a^c f \text{ converge ssi } \int_a^{c'} f \text{ converge} \right) \text{ et } \left(\int_c^b f \text{ converge ssi } \int_{c'}^b f \text{ converge} \right)$$

De plus,

$$\int_a^c f + \int_c^b f = \int_a^{c'} f + \int_{c'}^c f + \int_c^{c'} f + \int_{c'}^b f = \int_a^{c'} f + \int_{c'}^b f$$

4. Dans le cas où a, b sont des réels. Si on suppose que f est une fonction continue par morceaux sur le segment $[a, b]$, l'intégrale $\int_a^b f$ converge et

$$\int_a^b f = \int_a^b f$$

La formule ci-dessus affirme que l'intégrale vue en première année coïncide avec l'intégrale généralisée puisque dans ce cas, comme f est continue par morceaux sur le segment $[a, b]$, elle est bornée sur $[a, b]$. Si on note $M = \sup_{x \in [a, b]} |f(x)|$, on a

$$\left| \int_a^b f - \int_a^x f \right| = \left| \int_x^b f \right| \leq M(b-x) \rightarrow 0$$

Ceci est en particulier vrai si la fonction f est continue par morceaux et définie sur $]a, b]$ ou $[a, b[$ et qu'elle se prolonge par continuité en une fonction continue par morceaux sur $[a, b]$.

Exemples :

1. L'intégrale $\int_{-\infty}^{-1} \frac{1}{t^2} dt$ converge puisque pour $x < -1$,

$$\int_x^{-1} \frac{1}{t^2} dt = \left[-\frac{1}{t} \right]_x^{-1} = 1 + \frac{1}{x} \xrightarrow{x \rightarrow -\infty} 1$$

2. L'intégrale $\int_0^1 \frac{1}{t^2} dt$ diverge puisque pour $x \in]0, 1[$,

$$\int_x^1 \frac{1}{t^2} dt = \left[-\frac{1}{t} \right]_x^1 = -1 + \frac{1}{x} \xrightarrow{x \rightarrow 0^+} +\infty$$

Notation : Dans la suite on travaille sur un intervalle I de la forme $]a, b]$, $[a, b[$ ou $]a, b[$.

Pour $f \in \mathcal{CM}(I, \mathbf{K})$ dont l'intégrale converge, on notera $\int_I f$ pour $\int_a^b f$.

2.2 Propriétés

Proposition 3.13 (Linéarité)

L'ensemble des fonctions continues par morceaux sur I à valeurs dans $\mathbf{K}$ dont l'intégrale converge est un sous-espace vectoriel de $\mathcal{CM}(I, \mathbf{K})$ et l'application qui associe à une telle fonction son intégrale est linéaire.

Démonstration : La démonstration est similaire à celle du cas où $I = [a, +\infty[$. □

Proposition 3.14 (Positivité - Croissance - Relation de Chasles)

1. Soit f une fonction continue par morceaux sur I dont l'intégrale converge. Pour tout $c \in]a, b[$,

$$\int_a^b f = \int_a^c f + \int_c^b f$$

2. Soit f une fonction continue par morceaux sur I réelle positive dont l'intégrale converge.

On a $\int_I f \geq 0$.

3. Soit f et g deux fonctions réelles continues par morceaux sur I dont les intégrales convergent.

Si $f \geq g$ alors $\int_I f \geq \int_I g$.

4. Soit f une fonction **continue** sur I réelle et positive dont l'intégrale converge. Si $\int_I f = 0$ et I n'est pas réduit à un point alors f est la fonction nulle.

5. Soit f une fonction continue par morceaux sur I à valeurs dans $\mathbf{K}$ intégrable,

$$\left| \int_I f \right| \leq \int_I |f|$$

Démonstration :

1. La démonstration est similaire à celle du cas où $I = [a, +\infty[$.
2. La démonstration est similaire à celle du cas où $I = [a, +\infty[$.
3. La démonstration est similaire à celle du cas où $I = [a, +\infty[$.
4. On peut adapter la preuve faite sur les intervalles de la forme $[a, +\infty[$, on peut aussi faire une démonstration par l'absurde. Si on suppose que I n'est pas réduit à un point et que $f \neq \tilde{0}$, il existe $x_0 \in]a, b[$ tel que $f(x_0) > 0$. Par continuité de f , il existe $\varepsilon > 0$ tel que pour tout $x \in]x_0 - \varepsilon, x_0 + \varepsilon[$,

$$f(x) \geq \frac{f(x_0)}{2}$$

Par relation de Chasles et positivité on a alors

$$\int_a^b f = \int_a^{x_0-\varepsilon} f + \int_{x_0-\varepsilon}^{x_0+\varepsilon} f + \int_{x_0+\varepsilon}^b f \geq \int_{x_0-\varepsilon}^{x_0+\varepsilon} f \geq (2\varepsilon) \times \frac{f(x_0)}{2} = f(x_0) > 0$$

C'est absurde.

On a donc montré que $f = \tilde{0}$.

5. — Traitons d'abord le cas d'une fonction réelle. On suppose que $f \in \mathcal{CM}(I, \mathbf{R})$.

On sait que $\int_I f$ converge et que $-|f| \leq f \leq |f|$ donc $-\int_I |f| \leq \int_I f \leq \int_I |f|$.

Cela implique que $\left| \int_I f \right| \leq \int_I |f|$.

- Soit f une fonction à valeurs complexes. Notons f_r et f_i les parties réelle et imaginaire de f . On note aussi A et B les parties réelles et imaginaires de $\int_I f$:

$$\int_I f = \int_I f_r + i \int_I f_i = A + iB$$

Si on suppose que $B = 0$ et donc $\int_I f = \int_I f_r$. On en déduit que

$$\left| \int_I f \right| = \left| \int_I f_r \right| \leq \int_I |f_r| \leq \int_I |f|$$

La dernière inégalité venant du fait que pour tout nombre complexe z , $|\Re(z)| \leq |z|$.

Dans le cas général, notons φ un argument de $\int_I f$ (qui n'est pas nul, car sinon, on est dans le cas précédent). Posons $g : t \mapsto f(t)e^{-i\varphi}$. On en déduit que

$$\int_I g = \int_I f(t)e^{-i\varphi} dt = e^{-i\varphi} \int_I f(t) dt$$

D'après la définition de φ , $\int_I g$ est réel. On peut appliquer ce qui précède à g .

$$\left| \int_I f \right| = \left| \int_I g \right| \leq \int_I |g| = \int_I |f|$$

□

2.3 Intégrales des fonctions positives et fonctions intégrables

Proposition 3.15 (Convergence de l'intégrale d'une fonction positive)

Soit f et g deux fonctions réelles **positives** continues par morceaux sur I . On suppose que $0 \leq f \leq g$. Si $\int_a^b g$ converge alors $\int_a^b f$ converge.

Démonstration :

- Le cas où $I = [a, b[$ avec $b \in \mathbf{R}$ est similaire au cas où $I = [a, +\infty[$.
- Si $I =]a, b]$ où $a \in \mathbf{R} \cup \{-\infty\}$ et $b \in \mathbf{R}$. Considère la fonction $F : x \mapsto \int_x^b f$. Comme f est positive elle est **décroissante**. On en déduit que F a une limite en a^+ dans $\mathbf{R} \cup \{+\infty\}$.

Cependant, pour $x \in]a, b]$,

$$F(x) \leq \int_x^b g \leq \int_a^b g$$

La fonction F est donc majorée d'où $\lim_{x \rightarrow a} F(x) \in \mathbf{R}$ et donc $\int_a^b f$ converge.

- Si $I =]a, b[$. Soit $c \in]a, b[$. D'après nos hypothèses, $\int_a^c g$ et $\int_c^b g$ convergent donc $\int_a^c f$ et $\int_c^b f$ convergent d'où $\int_a^b f$ converge.

□

□

ATTENTION

Comme dans le cas des intervalles $[a, +\infty[$, si f est une fonction continue par morceaux et positive sur I , on peut toujours définir $\int_I f$ dans $[0, +\infty]$ avec la convention que $\int_I f = +\infty$ si l'intégrable diverge.

Définition 3.16

Soit f une fonction continue par morceaux sur l'intervalle I et à valeurs dans $\mathbf{K}$. On dit que la fonction est **intégrable** sur I ou que l'intégrale $\int_a^b f$ converge absolument si l'intégrale

$$\int_a^b |f| \text{ converge.}$$

On note $L^1(I, \mathbf{K})$ l'ensemble des fonctions continues par morceaux sur I à valeurs dans $\mathbf{K}$ qui sont intégrables.

Proposition 3.17

Soit $f \in \mathcal{CM}(I, \mathbb{K})$. Si $\int_a^b f$ converge absolument alors $\int_a^b f$ converge.

Démonstration : Il suffit d'adapter la preuve du résultat similaire pour une fonction continue par morceaux sur $[a, +\infty[$.

Proposition 3.18 (Théorème de comparaison)

Soit f et g deux fonctions continues par morceaux sur I .

1. Si $I = [a, b[$ avec $b \in \mathbb{R} \cup \{+\infty\}$.
 - Si $f = O_b(g)$ alors l'intégrabilité de g implique l'intégrabilité de f .
 - Si $f \sim_b g$ alors l'intégrabilité de g est équivalente à l'intégrabilité de f .
2. Si $I =]a, b]$ avec $a \in \mathbb{R} \cup \{-\infty\}$.
 - Si $f = O_a(g)$ alors l'intégrabilité de g implique l'intégrabilité de f .
 - Si $f \sim_a g$ alors l'intégrabilité de g est équivalente à l'intégrabilité de f .

Démonstration : Comme ci-dessus. □

Proposition 3.19 (Fonctions de référence)

1. En $+\infty$ (Rappel) :
 - La fonction $t \mapsto e^{-at}$ où $a \in \mathbb{R}$ est intégrable sur $[0, +\infty[$ si et seulement $a > 0$.
 - La fonction $t \mapsto \frac{1}{t^\alpha}$ est intégrable sur $[1, +\infty[$ si et seulement si $\alpha > 1$.
2. En 0 :
 La fonction $t \mapsto \frac{1}{t^\alpha}$ est intégrable sur $]0, 1]$ si et seulement si $\alpha < 1$. Dans ce cas,

$$\int_0^1 \frac{1}{t^\alpha} dt = \frac{1}{1-\alpha}$$

Démonstration :

1. Déjà vu.
2. – Si $\alpha = 1$ Soit $x \in]0, 1]$,

$$\int_x^1 \frac{dt}{t} = [\ln t]_x^1 = -\ln(x) \xrightarrow{x \rightarrow 0^+} +\infty$$

Donc $t \mapsto \frac{1}{t}$ n'est pas intégrable sur $]0, 1]$.

- Si $\alpha \neq 1$. Soit $x \in]0, 1]$,

$$\int_x^1 \frac{dt}{t^\alpha} = \left[\frac{t^{1-\alpha}}{1-\alpha} \right]_x^1 = \frac{1}{1-\alpha} (1 - x^{1-\alpha})$$

On en déduit que $t \mapsto \frac{1}{t^\alpha}$ est intégrable si et seulement si $\alpha < 1$. Dans ce cas,

$$\int_0^1 \frac{1}{t^\alpha} dt = \frac{1}{1-\alpha}$$

Remarques :

1. On a en particulier que $\int_0^1 \frac{dt}{t^2}$ diverge alors que $\int_0^1 \frac{dt}{\sqrt{t}}$ converge.
2. Nous verrons plus loin les changements de variables qui permettent de ramener l'étude de la nature d'une intégrable à une étude en 0^+ ou en $+\infty$.

ATTENTION

On utilise très souvent les intégrales $\int_0^1 \frac{1}{t^\alpha} dt$ et $\int_1^{+\infty} \frac{1}{t^\alpha} dt$. Il faut bien connaître les conditions de convergence et ne pas les mélanger.

★ Méthode : Croissance de $x^\alpha f$:

Le but de cette méthode est de montrer la convergence d'une intégrale en comparant aux intégrales

$$\int_0^1 \frac{1}{t^\alpha} dt.$$

S'il existe $\alpha < 1$ tel que $t^\alpha f(t)$ est bornée au voisinage de 0 (ce qui est vrai si $t^\alpha f(t) \xrightarrow{t \rightarrow 0} \ell \in \mathbf{R}$)

alors $f(t) = O\left(\frac{1}{t^\alpha}\right)$. Or $t \mapsto \frac{1}{t^\alpha}$ est intégrable car $\alpha < 1$. On en déduit que $\int_0^1 f$ est absolument convergente donc convergente.

La plupart du temps, on applique cela en prenant $\alpha = \frac{1}{2}$ et en étudiant $t \mapsto \sqrt{t}f(t)$.

Exemples :

1. Étudions la nature de $\int_0^1 (\ln t) dt$.

La fonction $t \mapsto \ln t$ est continue par morceaux sur $]0, 1]$. On sait que plus que $\sqrt{t} \ln t \xrightarrow{t \rightarrow 0^+} 0$

donc $\ln t = O\left(\frac{1}{\sqrt{t}}\right)$. On sait que $t \mapsto \frac{1}{\sqrt{t}} = \frac{1}{t^{\frac{1}{2}}}$ est intégrable sur $]0, 1]$ car $\frac{1}{2} < 1$. On en déduit

que $\int_0^1 \ln t dt$ est absolument convergente donc convergente.

Dans ce cas, on pouvait aussi remarquer que pour $x \in]0, 1]$,

$$\int_x^1 \ln t dt = [t \ln t - t]_x^1 = (0 - 1) - (x \ln x - x) \xrightarrow{x \rightarrow 0^+} 1$$

On en déduit que $\int_0^1 (\ln t) dt$ converge et que

$$\int_0^1 \ln t dt = -1$$

2. Soit $\alpha \in \mathbf{R}$, l'intégrale $\int_0^{+\infty} \frac{1}{t^\alpha} dt$ diverge. En effet, si $\alpha \geq 1$ alors $\int_0^1 \frac{1}{t^\alpha} dt$ diverge et si $\alpha \leq 1$ c'est $\int_1^{+\infty} \frac{1}{t^\alpha} dt$ qui diverge.

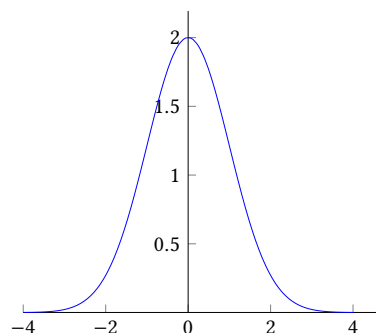
3. Étudions la nature de $\int_{-\infty}^{+\infty} e^{-a|t|} dt$ où $a > 0$. On sait que $\int_0^{+\infty} e^{-at} dt$ converge et $\int_0^{+\infty} e^{-at} dt = \frac{1}{a}$. Par le même calcul, $\int_{-\infty}^0 e^{at} dt$ converge et vaut $\frac{1}{a}$. L'intégrale étudiée converge donc et

$$\int_{-\infty}^{+\infty} e^{-a|t|} dt = \frac{2}{a}.$$

4. La fonction $t \mapsto e^{-t^2}$ est intégrable sur $\mathbb{R}$.

En effet, $t \mapsto e^{-t^2}$ est continue par morceaux sur $\mathbb{R}$. De plus, $t^2 e^{-t^2} \xrightarrow{t \rightarrow \pm\infty} 0$. On a donc $e^{-t^2} = o_{\pm\infty}\left(\frac{1}{t^2}\right)$. Comme $t \mapsto \frac{1}{t^2}$ est intégrable sur $[1, +\infty[$ et que $] -\infty, -1]$ les intégrales $\int_1^{+\infty} e^{-t^2} dt$ et $\int_{-\infty}^{-1} e^{-t^2} dt$ convergent car elles sont absolument convergentes. Finalement $\int_{-\infty}^{+\infty} e^{-t^2} dt$ converge (et donc $t \mapsto e^{-t^2}$ est intégrable car c'est une fonction positive).

On a de plus $\int_{-\infty}^{+\infty} e^{-t^2} dt = \sqrt{\pi}$ (intégrale de Gauss).



(TBC)

2.4 Méthodes de calculs

Nous allons reprendre les méthodes classiques vues en première année - intégration par parties et changement de variables - dans le cadre des intégrales sur un intervalle éventuellement ouvert.

Proposition 3.20 (Intégration par parties)

Soit f et g des fonctions de classe $\mathcal{C}^1$ sur I et $(a, b) \in \bar{I}^2$. On suppose que fg a une limite finie en a et en b et on note

$$[fg]_a^b = \lim_{t \rightarrow b} fg(t) - \lim_{t \rightarrow a} fg(t).$$

Les intégrales $\int_a^b f'g$ et $\int_a^b fg'$ sont de même nature et si elles convergent,

$$\int_a^b f'g = [fg]_a^b - \int_a^b fg'.$$

Remarque : Dans le cas où a et b appartiennent à I , on retrouve l'énoncé classique vue en première année.

Démonstration : On considère $c \in]a, b[$. Soit $x \in]a, c[$, on sait d'après le résultat usuel de l'intégration par parties que

$$\int_x^c f'g = [fg]_x^c - \int_x^c fg'.$$

Comme fg a une limite finie en a par hypothèse, $\int_a^c f'g$ converge si et seulement si $\int_a^c fg'$ converge.

Dans ce cas

$$\int_a^c f'g = [fg]_a^c - \int_a^c fg'.$$

On peut procéder de même pour les intégrales de c à b .

On a bien montré finalement que les intégrales $\int_a^b fg'$ et $\int_a^b f'g$ ont la même nature.

□

Exemple : On pose pour $n \geq 0$, $I_n = \int_0^{+\infty} t^n e^{-t} dt$.

Commençons par voir que pour $n \geq 0$, $t^{n+2}e^{-t} \xrightarrow{t \rightarrow +\infty} 0$ donc $t^n e^{-t} = o\left(\frac{1}{t^2}\right)$ de ce fait, pour tout $n \geq 0$, I_n est convergente.

Maintenant, on sait que pour tout $n \geq 1$, en posant

$$\begin{aligned} u(t) &= t^n & u'(t) &= nt^{n-1} \\ v(t) &= -e^{-t} & v'(t) &= e^{-t} \end{aligned}$$

On a u et v de classe $\mathcal{C}^1$ sur $[0, +\infty[$.

De plus, $\lim_{t \rightarrow +\infty} -t^n e^{-t} = 0$ donc

$$I_n = \int_0^{+\infty} t^n e^{-t} dt = [-t^n e^{-t}]_0^{+\infty} + n \int_0^{+\infty} t^{n-1} e^{-t} dt = nI_{n-1}$$

Or $I_0 = \int_0^{+\infty} e^{-t} dt = [-e^{-t}]_0^{+\infty} = 1$ donc $I_1 = 1$, $I_2 = 2$ et de manière générale, $I_n = n!$.

Exercice : Convergence et calcul de $\int_0^1 \frac{\ln t}{(1+t)^2} dt$. On fera attention au choix de la primitive.

Proposition 3.21 (Changement de variables)

Soit f une fonction continue sur $]a, b[$ et $\varphi :]\alpha, \beta[\rightarrow]a, b[$ qui est :

- bijective
- strictement croissante
- de classe $\mathcal{C}^1$.

Les intégrales

$$\int_{\varphi(\alpha)}^{\varphi(\beta)} f(t) dt \text{ et } \int_{\alpha}^{\beta} f(\varphi(u)) \varphi'(u) du$$

sont de même nature et sont égales en cas de convergence.

Remarque : Ici, $\int_{\varphi(\alpha)}^{\varphi(\beta)} f(t) dt = \int_a^b f(t) dt$

Démonstration : Soit $c \in]\alpha, \beta[$. Comme dans la démonstration du théorème sur l'intégration par partie, on va se contenter de montrer que les intégrales $\int_{\alpha}^c f(\varphi(u)) \varphi'(u) du$ et $\int_{\varphi(\alpha)}^{\varphi(c)} f(t) dt$ ont la

même nature et qu'en cas de convergence elles sont égales. Le cas général s'en déduit en découpant les intégrales en deux.

Soit $x \in]\alpha, c[$, on pose $x' = \varphi(x)$ et $c' = \varphi(c)$. On sait alors que l'on a

$$\int_{x'}^{c'} f(t)dt = \int_x^c f(\varphi(u))\varphi'(u)du$$

Cela se démontre en utilisant que

$$c \mapsto \int_x^c f(\varphi(u))\varphi'(u)du \text{ et } c \mapsto \int_{x'}^{\varphi(c)} f(t)dt$$

ont la même dérivée et d'annulent en x .

Regardons alors la nature de ces intégrales. D'après les hypothèses, $\lim_{x \rightarrow \alpha} \varphi(x) = a$ donc

$$\int_a^{c'} f(t)dt \text{ et } \int_\alpha^c f(\varphi(u))\varphi'(u)du$$

sont de même nature et égale dans le cas de convergence.

On en déduit que

$$\int_a^b f(t)dt \text{ et } \int_\alpha^\beta f(\varphi(u))\varphi'(u)du$$

sont de même nature et sont égales en cas de convergence.

□

Remarques :

1. Il existe une variante avec un changement de variables $\varphi :]\alpha, \beta[\rightarrow]a, b[$ strictement décroissant. Dans ce cas, on a $a = \varphi(\beta)$ et $b = \varphi(\alpha)$ donc on compare

$$\int_\alpha^\beta f(\varphi(u))\varphi'(u)du \text{ et } \int_b^a f(t)dt = \int_{\varphi(\alpha)}^{\varphi(\beta)} f(t)dt.$$

On voit que si f est positive, les deux termes sont bien négatifs car $\varphi'(u) < 0$.

2. Si φ est une fonction de classe $\mathcal{C}^1$ (donc continue) sur $]\alpha, \beta[$ et strictement croissante, elle est automatiquement bijective de $]\alpha, \beta[$ sur son intervalle image. L'énoncé consiste donc à dire que

$$\lim_{x \rightarrow \alpha} \varphi(x) = a \text{ et } \lim_{x \rightarrow \beta} \varphi(x) = b$$

3. Après avoir justifié que le changement de variable φ vérifie les hypothèses (sauf dans les cas usuels : fonctions affines, puissances, exponentielles, logarithme) on pourra utiliser les notations :

$$t = \varphi(u); dt = \varphi'(u)du$$

Remarque : On peut utiliser le changement de variables dans les deux sens. Précisément, si on veut calculer $\int_a^b f(t)dt$. On peut poser $t = \varphi(u)$ avec φ qui vérifie les hypothèses de l'énoncé. Il arrive aussi que l'on pose $u = \psi(t)$, dans quel cas le changement de variables avec les notations de l'énoncé est $t = \psi^{-1}(u)$ ce qui ne pose pas de problèmes car ψ est bijectif. Il faut faire attention au fait que ψ^{-1} soit encore de classe $\mathcal{C}^1$.

Remarque : Ce théorème justifie la méthode vue précédemment

Exemples :

1. Déterminons la nature de $I = \int_0^2 \frac{dx}{\sqrt{6-x-x^2}}$.

On commence par factoriser $x^2 + x - 6 = (x-2)(x+3)$. On en déduit que $x \mapsto \frac{1}{\sqrt{6-x-x^2}}$ est continue sur $[0, 2[$. Comme expliqué, on pose $u = 2 - x$; $du = -dx$ qui est un changement de variable affine. On sait alors que l'intégrale I est de la même nature que

$$J = - \int_2^0 \frac{du}{\sqrt{6-(2-u)-(2-u)^2}} = \int_0^2 \frac{du}{\sqrt{5u-u^2}}$$

Or $\frac{1}{\sqrt{5u-u^2}} \underset{u \rightarrow 0}{\sim} \frac{1}{\sqrt{5u}}$

On sait que $u \mapsto \frac{1}{\sqrt{u}}$ est intégrable sur $]0, 2]$ donc $u \mapsto \frac{1}{\sqrt{5u-u^2}}$. On en déduit que l'intégrale J est absolument convergente donc convergente et donc que I est convergente.

2. Déterminons la nature et calculons $\int_0^{+\infty} \frac{\ln t}{1+t^2} dt$ en posant $t = \frac{1}{u}$.

On voit d'abord que $\frac{\ln t}{1+t^2} \underset{t \rightarrow +\infty}{\sim} \frac{\ln t}{t^2} = o\left(\frac{1}{t^{3/2}}\right)$. Donc l'intégrale $\int_1^{+\infty} \frac{\ln t}{1+t^2} dt$ converge.

De même, $\frac{\ln t}{1+t^2} \underset{t \rightarrow 0}{\sim} \ln t$ or $\int_0^1 \ln t dt$ converge donc $\int_0^1 \frac{\ln t}{1+t^2} dt$ converge.

La fonction $\varphi : u \mapsto \frac{1}{u}$ réalise une bijection strictement décroissante de classe $\mathcal{C}^1$ de $]0, 1[$ dans $]0, +\infty[$. On pose $t = \frac{1}{u}$ et $dt = -\frac{1}{u^2} du$. Donc

$$\int_1^{+\infty} \frac{\ln t}{1+t^2} = \int_1^0 \frac{\ln(1/u)}{1+(1/u)^2} \frac{-1}{u^2} du = - \int_0^1 \frac{\ln u}{1+u^2} du.$$

On en déduit que

$$\int_0^{+\infty} \frac{\ln t}{1+t^2} dt = 0.$$

Remarque : Notons au passage que $\int_1^{+\infty} \frac{\ln t}{1+t^2} \simeq 0,915966$ s'appelle la constante de Catalan (souvent notée K). C'est un nombre dont on ne sait pas s'il est rationnel ou non.

Matheux (Eugène Catalan : 1814 - 1894)



Eugène Charles Catalan naît le 30 mai 1814 à Bruges, en Belgique. Il intègre l'École polytechnique (Promotion X 1833). Il est expulsé de l'école l'année suivante pour son activisme républicain. Il est autorisé en 1835 à reprendre ses études à Polytechnique et en sort diplômé.

La conjecture de Catalan stipule que la seule solution à l'équation $x^p - y^q = 1$ où x, y, p et q sont entiers est $x = 3, y = 2, p = 2$ et $q = 3$. Elle est démontrée en 2002 par Preda Mihăilescu.

3. Calculons $I = \int_0^1 \frac{1+t^2}{1+t^4} dt$ en posant $t = e^{-x}$. C'est un changement de variables dans l'autre sens.

Pour commencer l'intégrale est définie car la fonction est continue sur $[0, 1]$.

On pose $\varphi : t \mapsto e^{-t}$ qui est $\mathcal{C}^1$, bijectif de $]0, 1[$ sur $]0, +\infty[$ et strictement décroissant. On a donc $t = e^{-x}$ et $dt = -e^{-x} dx$. Cela donne :

$$\int_0^1 \frac{1+t^2}{1+t^4} dt = - \int_{+\infty}^0 \frac{1+e^{-2x}}{1+e^{-4x}} e^{-x} dx = \int_0^{+\infty} \frac{e^x + e^{-x}}{e^{2x} + e^{-2x}} dx = \int_0^{+\infty} \frac{\operatorname{ch} x}{\operatorname{ch} 2x} dx$$

On utilisant les formules trigonométriques, on voit que $\frac{\operatorname{ch} x}{\operatorname{ch} 2x} = \frac{\operatorname{ch} x}{1 + 2 \operatorname{sh}^2 x}$

On pose alors $u = \operatorname{sh} x$ qui est un changement de variables valable de $[0, +\infty[$ vers $[0, +\infty[$ avec $du = \operatorname{ch} x dx$.

Cela donne

$$I = \int_0^{+\infty} \frac{du}{1+2u^2} = \frac{1}{2} \int_0^{+\infty} \frac{du}{(1/\sqrt{2})^2 + u^2} = \frac{1}{2} \sqrt{2} \left[\arctan(\sqrt{2}u) \right]_0^{+\infty} = \frac{\pi}{2\sqrt{2}}.$$

Exercice : Calculer $I = \int_0^1 \frac{\ln t}{\sqrt{1-t}} dt$. On pourra faire une intégration par parties, puis un changement de variables.

Solution : Commençons par justifier la convergence de l'intégrale.

- En 0, $\frac{\ln t}{\sqrt{1-t}} \underset{t \rightarrow 0}{\sim} \ln(t)$. Or $\ln t = \underset{t \rightarrow 0}{o} \left(\frac{1}{\sqrt{t}} \right)$ car $\sqrt{t} \ln(t) \xrightarrow{t \rightarrow 0} 0$ donc $\int_0^{1/2} \frac{\ln t}{\sqrt{1-t}} dt$ converge.
- En 1, $\frac{\ln t}{\sqrt{1-t}} \underset{t \rightarrow 1}{\sim} \frac{t-1}{\sqrt{1-t}} = -\sqrt{1-t} = -(1-t)^{1/2}$ donc $\int_{1/2}^1 \frac{\ln t}{\sqrt{1-t}} dt$ converge.

Calculons maintenant la valeur de l'intégrale par intégration par parties.

On pose

$$\begin{aligned} u(t) &= \ln t & u'(t) &= \frac{1}{t} \\ v(t) &= 2 - 2\sqrt{1-t} & v'(t) &= \frac{1}{\sqrt{1-t}} \end{aligned}$$

On a u et v de classe $\mathcal{C}^1$ sur $]0, 1[$. De plus, $\lim_{t \rightarrow 1} u(t)v(t) = 0$ et pour $t \rightarrow 0$, $v(t) \underset{t \rightarrow 0}{\sim} t$ donc $\lim_{t \rightarrow 0} u(t)v(t) = 0$. On en déduit que

$$\int_0^1 \frac{\ln t}{\sqrt{1-t}} dt = \int_0^1 \frac{2(1-\sqrt{1-t})}{t} dt$$

Il reste à faire un changement de variable. On pose $u = \sqrt{1-t} \iff t = 1-u^2$ qui est un changement de variable valide de $[0, 1]$ sur $[0, 1]$ avec $dt = 2udu$. Cela donne

$$I = \int_1^0 \frac{2(1-u)u}{1-u^2} du = - \int_0^1 \frac{2u}{1+u} du = -2 + \int_0^1 \frac{2}{1+u} du = -2 + 2 \ln 2.$$

★ **Méthode :** On a vu des intégrales de références sur $[1, +\infty[$ (et donc sur $[a, +\infty[$ par la relation de Chasles) ainsi que sur $[0, 1]$ (et donc $]0, b]$ par Chasles). En général on se ramène à ces deux cas par un changement de variables affine. Les changements de variables seront vus un peu plus loin mais nous anticipons un peu.

- Pour étudier une intégrale sur $] -\infty, b]$ on peut se ramener à $[-b, +\infty[$ en posant $u = -t$ et on utilise que $\int_{-\infty}^b f(t)dt$ est de la même nature que $\int_{-b}^{+\infty} f(-u)du$. Ici on a échangé les bornes de l'intégrale car $du = -dt$.
- Pour étudier une intégrale sur $[a, b[$ avec $b \in \mathbf{R}$, on peut se ramener à $]0, c]$ en posant $t = b - u$ et on utilise que $\int_a^b f(t)dt$ est de la même nature que $\int_0^{b-a} f(b - u)du$. Là encore on échange les bornes de l'intégrale car $du = -dt$.
- Pour étudier une intégrale sur $]a, b]$ avec $a \in \mathbf{R}$, on peut se ramener à $]0, c]$ en posant $t = a + u$ et on utilise que $\int_a^b f(t)dt$ est de la même nature que $\int_0^{b-a} f(a + u)du$.

Exercice : Déterminer la nature de $\int_1^2 \frac{1}{\ln(t)} dt$.

AJOUTER $\int_a^b \frac{1}{|x - a|^\alpha}$ voir programme.

2.5 Intégration des relations de comparaisons

Comme dans le cas des séries, on va voir que les relations de comparaisons : O , o et $\sim$ sont compatibles avec l'intégration. Il y a (comme pour les séries) deux cas, si les intégrales sont divergentes on va comparer les intégrales, si elles sont convergentes on va comparer les restes.

Théorème 3.22 (intégration des relations de comparaison)

Soit $a \in \mathbf{R}$ et $b \in \mathbf{R} \cup \{+\infty\}$. Soit $g \in \mathcal{CM}([a, b[, \mathbf{R})$ une fonction **positive**.
Soit $f \in \mathcal{CM}([a, b[, \mathbf{C})$

1. Cas des intégrales divergentes.

On suppose que l'intégrable $\int_a^b g$ est divergente.

(a) Si $f(t) = o_{t \rightarrow b}(g(t))$ alors $\int_a^x f = o_{x \rightarrow b}\left(\int_a^x g\right)$

(b) Si $f(t) = O_{t \rightarrow b}(g(t))$ alors $\int_a^x f = O_{x \rightarrow b}\left(\int_a^x g\right)$

(c) Si $f(t) \sim_{t \rightarrow b} g(t)$ alors $\int_a^x f \sim_{x \rightarrow b} \int_a^x g$

2. Cas des intégrales convergentes.

On suppose que l'intégrable $\int_a^b g$ est convergente.

(a) Si $f(t) = o_{t \rightarrow b}(g(t))$ alors $\int_x^b f = o_{x \rightarrow b}\left(\int_x^b g\right)$

(b) Si $f(t) = O_{t \rightarrow b}(g(t))$ alors $\int_x^b f = O_{x \rightarrow b}\left(\int_x^b g\right)$

(c) Si $f(t) \sim_{t \rightarrow b} g(t)$ alors $\int_x^b f \sim_{x \rightarrow b} \int_x^b g$

Démonstration :

1. Cas des intégrales divergentes.

On suppose que $\int_a^b g$ diverge.

- (a) On suppose que $f(t) = o_{t \rightarrow b}(g(t))$. Cela signifie que pour tout $\varepsilon > 0$, il existe $b' \in [a, b[$ tel que pour tout $t \in [b', b[$, $|f(t)| \leq \varepsilon g(t)$.

En particulier, pour $x \in [b', b[$,

$$\begin{aligned} \left| \int_a^x f \right| &\leq \left| \int_a^{b'} f \right| + \left| \int_{b'}^x f \right| \\ &\leq A + \int_{b'}^x |f| \quad \text{où } A = \left| \int_a^{b'} f \right| \text{ ne dépend pas de } x \\ &\leq A + \varepsilon \int_{b'}^x g \\ &\leq A + \varepsilon \int_a^x g \quad \text{car } g \text{ est positive} \end{aligned}$$

Or, par hypothèse, $\lim_{x \rightarrow b} \int_a^x g = +\infty$. Il existe donc b'' (que l'on peut choisir supérieur à b') tel que pour $x \geq b''$, $\int_a^x g \geq \frac{A}{\varepsilon}$.

Pour $x \geq b'' \geq b'$ on a donc : $\left| \int_a^x f \right| \leq 2\varepsilon \int_a^x g$.

Cela montre que $\int_a^x f = o_{x \rightarrow b} \left(\int_a^x g \right)$

- (b) On suppose que $f(t) = O_{t \rightarrow b}(g(t))$. Cela signifie qu'il existe $b' \in [a, b[$ et $M \in \mathbf{R}_+$ tel que pour tout $t \in [b', b[$, $|f(t)| \leq M g(t)$.

En particulier, pour $x \in [b', b[$,

$$\begin{aligned} \left| \int_a^x f \right| &\leq \left| \int_a^{b'} f \right| + \left| \int_{b'}^x f \right| \\ &\leq A + \int_{b'}^x |f| \quad \text{où } A = \left| \int_a^{b'} f \right| \text{ ne dépend pas de } x \\ &\leq A + M \int_{b'}^x g \\ &\leq A + M \int_a^x g \quad \text{car } g \text{ est positive} \end{aligned}$$

Or, par hypothèse, $\lim_{x \rightarrow b} \int_a^x g = +\infty$. Il existe donc b'' (que l'on peut choisir supérieur à b') tel que pour $x \geq b''$, $\int_a^x g \geq A$.

Pour $x \geq b'' \geq b'$ on a donc : $\left| \int_a^x f \right| \leq (M+1) \int_a^x g$.

Cela montre que $\int_a^x f = O_{x \rightarrow b} \left(\int_a^x g \right)$

- (c) On suppose que $f(t) \sim_{t \rightarrow b} g(t)$.

Cela implique que $f(t) = g(t) + o_{t \rightarrow b}(g(t))$ ou que $f(t) - g(t) = o_{t \rightarrow b}(g(t))$. On peut donc appliquer ce qui a été vu en (a).

On obtient que $\int_a^x f - \int_a^x g = o_{x \rightarrow b} \left(\int_a^x f \right)$ et donc $\int_a^x f \underset{x \rightarrow b}{\sim} \int_a^x g$

2. Cas des intégrales convergentes.

On suppose que $\int_a^b g$ converge. On remarque que dans les trois cas, les intégrales $\int_a^b |f|$ et $\int_a^b f$ sont convergentes. On peut donc considérer pour $x \in [a, b[$ les intégrales $\int_x^b |f|$ et $\int_x^b f$

(a) On suppose que $f(t) = o_{t \rightarrow b}(g(t))$. Cela signifie que pour tout $\varepsilon > 0$, il existe $b' \in [a, b[$ tel que pour tout $t \in [b', b[$, $|f(t)| \leq \varepsilon g(t)$.

En particulier, pour $x \in [b', b[$,

$$\left| \int_x^b f \right| \leq \int_x^b |f| \leq \varepsilon \int_x^b g$$

Cela montre que $\int_x^b f = o_{x \rightarrow b} \left(\int_x^b g \right)$

(b) On suppose que $f(t) = O_{t \rightarrow b}(g(t))$. Cela signifie qu'il existe $b' \in [a, b[$ et $M \in \mathbb{R}_+$ tel que pour tout $t \in [b', b[$, $|f(t)| \leq M g(t)$.

En particulier, pour $x \in [b', b[$,

$$\left| \int_x^b f \right| \leq \int_x^b |f| \leq M \int_x^b g$$

Cela montre que $\int_x^b f = O_{x \rightarrow b} \left(\int_x^b g \right)$

(c) On suppose que $f(t) \underset{t \rightarrow b}{\sim} g(t)$.

Cela implique que $f(t) = g(t) + o_{t \rightarrow b}(g(t))$ ou que $f(t) - g(t) = o_{t \rightarrow b}(g(t))$. On peut donc appliquer ce qui a été vu en (a).

On obtient que $\int_x^b f - \int_x^b g = o_{x \rightarrow b} \left(\int_x^b f \right)$ et donc $\int_x^b f \underset{x \rightarrow b}{\sim} \int_x^b g$

□

Remarques :

1. Il y a des énoncés équivalents pour les fonctions définies sur $]a, b]$ avec $a \in \mathbb{R} \cup \{-\infty\}$ et $b \in \mathbb{R}$.
2. Il n'y a pas de conditions sur la fonction f qui peut ne pas être de signe constant ou même complexe, par contre la fonction g doit être réelle positive.

Exemples :

1. On pose pour $x \in]0, 1]$, $f(x) = \int_x^1 \frac{dt}{\sin t}$. Comme on sait que $\sin t \underset{t \rightarrow 0}{\sim} t$, l'intégrale $\int_0^1 \frac{dt}{\sin t}$ diverge car $\int_0^1 \frac{dt}{t}$ diverge. On en déduit que $\lim_{x \rightarrow 0^+} f(x) = +\infty$.

On cherche un équivalent de f en 0^+ . Comme $\frac{1}{\sin t} \underset{t \rightarrow 0}{\sim} \frac{1}{t}$ et que $\int_0^1 \frac{1}{t} dt$ diverge,

$$f(x) \underset{x \rightarrow 0}{\sim} \int_x^1 \frac{dt}{t} = -\ln(x)$$

2. Soit $\alpha > 0$, on cherche un équivalent de $\int_1^x (\ln t)^\alpha dt$.

On intègre par parties,

$$\int_1^x (\ln t)^\alpha dt = [t(\ln t)^\alpha]_1^x - \alpha \int_1^x (\ln t)^{\alpha-1} dt.$$

Maintenant, $\int_1^x (\ln t)^\alpha dt$ diverge de manière évidente et $(\ln t)^{\alpha-1} = o_{+\infty}((\ln t)^\alpha)$ donc

$$\alpha \int_1^x (\ln t)^{\alpha-1} dt = o_{x \rightarrow +\infty} \left(\alpha \int_1^x (\ln t)^\alpha dt \right)$$

De ce fait,

$$\int_1^x (\ln t)^\alpha dt \underset{+\infty}{\sim} [t(\ln t)^\alpha]_1^x \underset{+\infty}{\sim} x(\ln x)^\alpha.$$

3. Cherchons un équivalent de $\int_x^{+\infty} \frac{t + \sin t}{t^3} dt$ pour $t \rightarrow +\infty$.

On a $\frac{t + \sin t}{t^3} \underset{t \rightarrow +\infty}{\sim} \frac{1}{t^2}$ donc l'intégrale converge.

Cela donne

$$\int_x^{+\infty} \frac{t + \sin t}{t^3} dt \underset{+\infty}{\sim} \int_x^{+\infty} \frac{1}{t^2} dt = \frac{1}{x}.$$

3 Les théorèmes de Lebesgue

On veut étudier si on peut intervertir une limite et une intégrale. Précisément si on dispose de fonctions f_n définies et intégrables sur I , a-t-on

$$\lim_{n \rightarrow \infty} \int_I f_n = \int_I \lim_{n \rightarrow \infty} f_n?$$

Par exemple, est-il vrai que

$$\lim_{n \rightarrow \infty} \int_0^1 x^n dx = \int_0^1 \left(\lim_{n \rightarrow \infty} x^n \right) dx$$

3.1 Limite (simple) d'une suite de fonctions et d'une série de fonctions

Dans ce paragraphe X désigne un ensemble

Définition 3.23 (Convergence simple)

Soit $(f_n) \in \mathcal{F}(X, \mathbf{K})^{\mathbf{N}}$ une suite de fonctions et $f \in \mathcal{F}(X, \mathbf{K})$. On dit que la suite (f_n) converge simplement vers f si pour tout x de X la suite $(f_n(x))$ tend vers $f(x)$. On a donc

$$\forall x \in X, \forall \varepsilon > 0, \exists N \in \mathbf{N}, \forall n \in \mathbf{N}, n \geq N \Rightarrow |f_n(x) - f(x)| \leq \varepsilon$$

On dit alors que f est la limite simple de (f_n) et on note

$$(f_n) \xrightarrow{CS} f$$

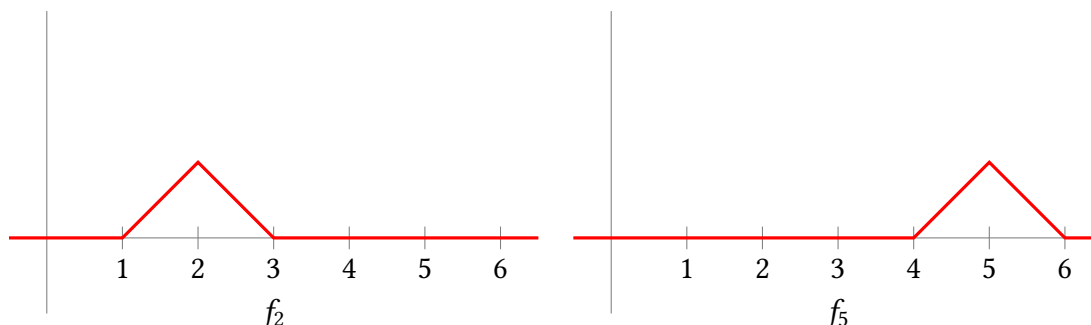
Exemples :

1. Pour $X = \mathbf{R}$, la suite $f_n : x \mapsto x^n$ ne converge pas (simplement) car pour $x = 2$ par exemple $f_n(x) \rightarrow +\infty$.
2. Pour $X = [0, 1]$ la suite $f_n : x \mapsto x^n$ converge vers f définie par

$$f : x \mapsto \begin{cases} 0 & \text{si } x < 1 \\ 1 & \text{si } x = 1 \end{cases}$$

3. Pour $X = [0, 1]$, la suite $f_n : x \mapsto \frac{nx+1}{n}$ converge vers la fonction $f : x \mapsto x$.
4. On considère la fonction f_n définie sur $\mathbf{R}$ par

$$f_n : x \mapsto \begin{cases} 0 & \text{si } x \notin [n-1, n+1] \\ x - n + 1 & \text{si } x \in [n-1, n] \\ -x + n + 1 & \text{si } x \in [n, n+1] \end{cases}$$



La suite de fonctions converge simplement vers la fonction nulle.

Exercice : Pour les suites suivantes définies sur $[0, 1]$, déterminer si elles convergent simplement et si oui déterminer la limite.

$$f_n : x \mapsto \sin(nx) \quad g_n : x \mapsto \frac{n}{1 + n(x+1)}$$

Remarque : Nous verrons plus tard une autre définition plus contraignante de la convergence d'une suite de fonctions.

Proposition 3.24 (Unicité de la limite)

Soit $(f_n) \in \mathcal{F}(X, \mathbf{K})^{\mathbf{N}}$ une suite de fonctions et $(g, h) \in \mathcal{F}(X, \mathbf{K})^2$. Si $(f_n) \xrightarrow{CS} g$ et $(f_n) \xrightarrow{CS} h$ alors $g = h$.

Démonstration : Il suffit d'utiliser l'unicité de la limite des suite d'éléments de $\mathbf{K}$. □

Proposition 3.25 (Linéarité)

Soit λ, μ deux scalaires. Si $(f_n) \xrightarrow{CS} f$ et $(g_n) \xrightarrow{CS} g$ alors $(\lambda f_n + \mu g_n) \xrightarrow{CS} \lambda f + \mu g$.

3.2 Intersion de lim et $\int$

On veut voir si, quand $(f_n) \xrightarrow{CS} f$, la limite des intégrales des fonctions f_n est-elle égale à la l'intégrale de la fonction f .

Exemples :

1. Si $f_n : x \mapsto x^n$ sur $[0, 1]$ qui converge vers $f : x \mapsto \begin{cases} 0 & \text{si } x \neq 1 \\ 1 & \text{si } x = 1 \end{cases}$

On a

$$\int_{[0,1]} f_n = \left[\frac{x^{n+1}}{n+1} \right]_0^1 = \frac{1}{n+1} \longrightarrow 0 = \int_{[0,1]} f.$$

2. Si $f_n : x \mapsto nx^n$ sur $[0, 1[$ qui converge vers la fonction nulle. Cela ne fonctionne plus. En effet, les fonctions f_n sont intégrables sur $[0, 1[$ et $I_n = \int_0^1 f_n = \frac{n}{n+1}$. Par contre $(I_n) \rightarrow 1 \neq 0$.

3. On reprend la suite de fonction (f_n) définie par

$$f_n : x \mapsto \begin{cases} 0 & \text{si } x \notin [n-1, n+1] \\ x - n + 1 & \text{si } x \in [n-1, n] \\ -x + n + 1 & \text{si } x \in [n, n+1] \end{cases}$$

Là encore $(f_n) \xrightarrow{CS} 0$. Par contre, pour tout entier n , $I_n = \int_0^{+\infty} f_n = 1$. On trouve donc que

$$(I_n) \rightarrow 1 \neq 0.$$

3.3 Théorème de convergence dominée

Théorème 3.26 (Théorème de convergence dominée)

Soit (f_n) une suite de fonctions continues par morceaux sur un intervalle I à valeurs dans $\mathbf{K}$. On suppose

- i) La suite (f_n) converge vers une fonction f continue par morceaux
- ii) (Hypothèse de domination) : Il existe une fonction φ **intégrable** sur I telle que pour tout entier n , $|f_n| \leq \varphi$.

Alors f est intégrable sur I et

$$\lim_{n \rightarrow \infty} \int_I f_n = \int_I f$$

ATTENTION

L'hypothèse importante est l'hypothèse de domination. Les hypothèses de continuité par morceaux des fonctions f_n et de la fonction limite f ne sont là que parce que l'on utilise en CPGE l'intégrale de Riemann et non l'intégrale de Lebesgue. Le programme officiel dit : *Pour l'application pratique des énoncés de ce paragraphe, on vérifie les hypothèses de convergence simple et de domination, sans expliciter celles relatives à la continuité par morceaux*

Remarques :

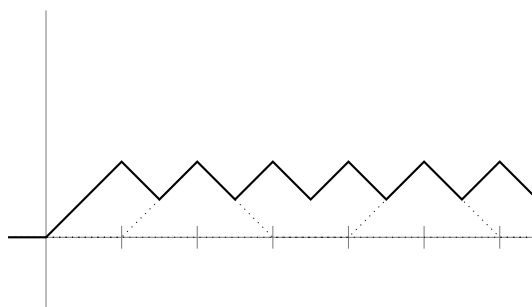
1. Ce théorème est difficile à démontrer et est admis.
2. Le fait que les fonctions f_n soient intégrables découle de l'hypothèse de domination.

Exemples :

1. Reprenons $f_n : x \mapsto x^n$ sur $[0, 1]$. On peut dominer par $\varphi : x \mapsto 1$. Le résultat est donc obtenu par le théorème.
2. Reprenons

$$f_n : x \mapsto \begin{cases} 0 & \text{si } x \notin [n-1, n+1] \\ x - n + 1 & \text{si } x \in [n-1, n] \\ -x + n + 1 & \text{si } x \in [n, n+1] \end{cases}$$

La meilleure majoration possible est par $\varphi(x) = \sup_{n \in \mathbb{N}} |f_n(x)|$ qui n'est pas intégrable.



3. On pose $I_n = \int_0^{+\infty} \frac{dt}{t^n + e^t}$. Pour $n \geq 0$ on considère les fonctions $f_n : t \mapsto \frac{1}{t^n + e^t}$ qui sont continues. On étudie la convergence simple de la suite de fonctions (f_n) :

$$f_n(t) \xrightarrow{n \rightarrow \infty} f(t) = \begin{cases} 1/e^t & \text{si } t < 1 \\ 1/(1 + e^t) & \text{si } t = 1 \\ 0 & \text{si } t > 1 \end{cases}$$

De plus, pour tout t , $|f_n(t)| \leq \frac{1}{e^t} = e^{-t}$. La fonction $\varphi : t \mapsto e^{-t}$ est une fonction intégrable sur $[0, +\infty[$. On en déduit que

$$I_n \rightarrow \int_0^{\infty} f = \int_0^1 e^{-t} dt = 1 - \frac{1}{e}.$$

4. On pose $I_n = \int_0^n \left(1 - \frac{x}{n}\right)^n dx$. On cherche la limite de I_n quand $n \rightarrow \infty$.

Ici, l'intervalle d'intégration dépend de n . Pour cela on pose

$$f_n : x \mapsto \begin{cases} \left(1 - \frac{x}{n}\right)^n & \text{si } x \leq n \\ 0 & \text{sinon} \end{cases}$$

De ce fait,

$$I_n = \int_0^{+\infty} f_n$$

Maintenant, pour x fixé, $f_n(x) \rightarrow e^{-x}$.

On utilise alors que $\ln(1+x)$ est concave et donc $\forall x > -1, \ln(1+x) \leq x$. En obtient alors que pour tout entier n et $x \leq n$,

$$\ln\left(1 - \frac{x}{n}\right) \leq -\frac{x}{n}$$

D'où

$$f_n(x) = \left(1 - \frac{x}{n}\right)^n \leq e^{-x}$$

Cette inégalité est encore vraie si $x > n$. La fonction $\varphi : x \mapsto e^{-x}$ est intégrable sur $[0, +\infty[$ et domine la suite (f_n) .

On en déduit que $(I_n) \rightarrow \int_0^{+\infty} e^{-x} dx = 1$.

5. On veut calculer $\lim_{n \rightarrow +\infty} \int_0^1 nx^n dx$.

On ne peut pas appliquer directement le théorème de convergence dominée. En effet, si on pose $f_n : x \mapsto nx^n$, il est clair que sur $[0, 1[$, $(f_n) \xrightarrow{CS} 0$. Cependant, on ne peut pas obtenir d'hypothèse de domination. En effet on cherche φ une fonction telle que

$$\forall x \in [0, 1[, \forall n \in \mathbb{N}, f_n(x) \leq \varphi(x)$$

Pour cela, à x fixé, on étudie la fonction $t \mapsto tx^t$. Via une dérivation, on vérifie que cette fonction atteint son maximum pour $t = -\frac{1}{\ln x}$. On est donc amené à poser

$$\varphi : x \mapsto -\frac{1}{\ln x} x^{-\frac{1}{\ln x}} = -\frac{1}{e \ln x}$$

Cette fonction, n'est pas intégrable car $\ln x \underset{x \rightarrow 1}{\sim} x - 1$.

Pour cela on réalise le changement de variables $u = x^n$. On en déduit que

$$\int_0^1 nx^n dx = \int_0^1 \sqrt[n]{u} du.$$

Comme pour $u > 0$, $\sqrt[n]{u} \rightarrow 1$ et que les fonctions sont dominées par $\varphi : u \rightarrow 1$, on a

$$\lim_{n \rightarrow +\infty} \int_0^1 nx^n dx = \lim_{n \rightarrow +\infty} \int_0^1 \sqrt[n]{u} du = \int_0^1 1 du = 1.$$

★ **Méthode** : On se donne la suite de fonction (f_n) qui converge simplement vers une fonction f et on cherche une fonction φ qui réalise la domination.

- Il peut y avoir une domination « évidente » : majorer $|\sin(t)|$ ou $|\cos(t)|$ par 1 ; majorer t^n sur $[0, 1]$ par 1... Par exemple si on considère $f_n : t \mapsto \frac{|\sin t|^n}{1+t^2}$ sur $[0, +\infty[$. La suite de fonctions (f_n) converge simplement vers

$$f : t \mapsto \begin{cases} 0 & \text{si } t \neq \frac{\pi}{2}[\pi] \\ \frac{1}{1+t^2} & \text{si } t = \frac{\pi}{2}[\pi] \end{cases}$$

De plus, en posant $\varphi : t \mapsto \frac{1}{1+t^2}$ qui est intégrable on a bien que pour tout entier n , $|f_n| \leq \varphi$. Finalement,

$$\lim_{n \rightarrow \infty} \int_0^\infty \frac{|\sin(t)|^n}{1+t^2} dt = \int_0^\infty f = 0$$

- À t fixé, la suite $(|f_n(t)|)$ peut être monotone. Si elle est décroissante, on peut prendre $\varphi = |f_0|$; si elle est croissante, on peut prendre $\varphi = |f|$.
- Dans le cas général, la « meilleure » fonction à prendre est $\varphi : t \mapsto \sup_{n \in \mathbb{N}} |f_n(t)|$. Elle peut être difficile à calculer. Cela peut nécessiter une étude des variations **par rapport à n** (et à t fixé) du terme $|f_n(t)|$.

3.4 Théorème de convergence dominée - Cas continu

On peut généraliser ce théorème au cas où le paramètre n'est plus entier mais réel.

Théorème 3.27 (Théorème de convergence dominée - cas « continu »)

Soit $(f_y)_{y \in J}$ une famille de fonctions continues par morceaux sur un intervalle I et à valeurs dans $\mathbb{K}$. On suppose

- i) Pour tout x de I , $y \mapsto f_y(x)$ admet une limite finie notée $f(x)$ quand y tend vers $y_0 \in \bar{J}$.
- ii) La fonction f est continue par morceaux.
- iii) (Hypothèse de domination) : Il existe une fonction φ intégrable sur I telle que pour tout $y \in J$, $|f_y| \leq \varphi$.

Alors f est intégrable sur I et

$$\lim_{y \rightarrow y_0} \int_I f_y = \int_I f.$$

Remarques :

1. Les deux premières hypothèses sont les analogues du fait que (f_n) converge simplement vers f continue par morceaux.
2. Là encore, l'hypothèse de domination implique que les f_y sont intégrables sur I .
3. Le programme officiel stipule qu'il n'est pas nécessaire d'évoquer la condition ii).

Démonstration : Considérons une suite (y_n) d'éléments de J tendant vers y_0 (ce qui est possible par définition de $\bar{J}$). On pose $f_n = f_{y_n}$.

Les hypothèses i) et ii) affirment alors que la suite de fonctions (f_n) converge simplement vers f . On peut donc appliquer le théorème de convergence dominée « classique ». On obtient en particulier que f est intégrable et que

$$\lim_{n \rightarrow \infty} \int_I f_{y_n} = \int_I f.$$

Maintenant, on peut faire cela pour toute suite (y_n) tendant vers y_0 . On conclut alors en utilisant la caractérisation séquentielle de la limite. $\square$

Exemple : Pour tout $y \in \mathbb{R}_+$ on pose

$$f_y : \mathbb{R}_+ \rightarrow \mathbb{R} \\ t \mapsto \frac{e^{-yt}}{1+t^2}$$

On veut étudier $\lim_{y \rightarrow +\infty} \int_0^\infty f_y(t) dt$.

- i) Pour tout t de $\mathbf{R}_+$, $y \mapsto f_y(t)$ admet une limite finie notée $f(x)$ quand y tend vers $y_0 \in +\infty$. Il suffit de poser

$$f : t \mapsto \begin{cases} 1 & \text{si } t = 0 \\ 0 & \text{sinon.} \end{cases}$$

- ii) La fonction f est continue par morceaux.

- iii) (Hypothèse de domination) : Il existe une fonction φ intégrable sur I telle que pour tout $y \in J$, $|f_y| \leq \varphi$. On peut prendre $\varphi : t \mapsto \frac{1}{1+t^2}$.

On en déduit que f_y et f sont intégrables et

$$\lim_{y \rightarrow +\infty} \int_0^\infty \frac{e^{-yt}}{1+t^2} dt = \int_0^\infty f = 0.$$

3.5 Intégration terme à terme d'une somme d'une série de fonctions

Définition 3.28 (Convergence simple des séries de fonctions)

Soit (f_n) une suite de fonctions définies sur X et à valeurs dans $\mathbf{K}$.

On note pour tout entier n , $S_n = \sum_{k=0}^n f_k \in \mathcal{F}(X, \mathbf{K})$. On dit que la série de fonctions $\sum f_n$

converge simplement si (S_n) converge simplement. Dans ce cas on note $\sum_{n=0}^{+\infty} f_n$ la limite que l'on appelle la somme de la série de fonctions.

Remarque : Cela signifie que pour tout x de X , la série $\sum f_n(x)$ est convergente. De plus,

$$\left(\sum_{n=0}^{+\infty} f_n \right) (x) = \sum_{n=0}^{+\infty} f_n(x).$$

Exemples :

1. Si on considère $f_n : x \mapsto \frac{x^n}{n!}$. La série $\sum f_n$ converge et la somme est :

$$\sum_{n=0}^{+\infty} f_n : x \mapsto \sum_{n=0}^{+\infty} \frac{x^n}{n!} = \exp(x).$$

2. Si on considère $f_n : x \mapsto x^n$. La série converge si on considère les fonctions définies sur $] -1, 1[$ et la somme est

$$S : x \mapsto \sum_{n=0}^{+\infty} x^n = \frac{1}{1-x}$$

Théorème 3.29 (Intégration terme à terme - cas positif)

Soit (f_n) une suite de fonctions définies sur un intervalle I à valeurs dans $\mathbf{R}_+$. On suppose

- i) La série de fonctions converge (simplement) et on note $S : x \mapsto \sum_{n=0}^{+\infty} f_n(x)$ la somme
- ii) Les fonctions f_n et S sont continues par morceaux.
- iii) Les fonctions f_n sont intégrables.

Alors, dans $[0, +\infty]$

$$\sum_{n=0}^{+\infty} \int_I f_n = \int_I S$$

En particulier, la fonction S est intégrable sur I si et seulement si la série $\sum \int_I f_n$ converge.

Remarques :

1. Ce théorème est admis.
2. Le programme officiel stipule que dans l'application de cet théorème, il n'est pas nécessaire de vérifier que les fonctions f_n et S sont continues par morceaux.

Exemple : Pour $n \geq 1$, on considère $f_n : t \mapsto \frac{1}{(1+t^2)^n}$ définie sur $]0, +\infty[$. On s'intéresse à la série de fonctions $\sum_{n \geq 1} f_n$. On reconnaît une série géométrique de raison $0 < \frac{1}{1+t^2} < 1$ car $t > 0$. On en déduit que la série de fonctions $\sum_{n \geq 1} f_n$ converge vers S où

$$S : t \mapsto \sum_{n=1}^{\infty} \frac{1}{(1+t^2)^n} = \frac{1}{1+t^2} \times \frac{1}{1 - \frac{1}{1+t^2}} = \frac{1}{t^2}$$

Comme on sait que $\int_0^{\infty} \frac{1}{t^2} dt = +\infty$, la série $\sum_{n \geq 1} \int_0^{\infty} \frac{1}{(1+t^2)^n} dt$ diverge.

Théorème 3.30 (Intégration terme à terme)

Soit (f_n) une suite de fonctions définies sur un intervalle I . On suppose

- i) La série de fonctions converge (simplement) et on note $S : x \mapsto \sum_{n=0}^{+\infty} f_n(x)$ la somme
- ii) Les fonctions f_n et S sont continues par morceaux.
- iii) Les fonctions f_n sont intégrables.
- iv) La série $\sum \int_I |f_n|$ converge.

Alors la fonction S est intégrable et

$$\sum_{n=0}^{+\infty} \int_I f_n = \int_I S$$

Remarques :

1. Ce théorème est admis.
2. Le programme officiel stipule que dans l'application de cet théorème, il n'est pas nécessaire de vérifier que les fonctions f_n et S sont continues par morceaux.

Exemple : Pour tout entier naturel n , on pose donc $f_n : x \mapsto x^{2n} \ln x$ définies sur $]0, 1[$. On considère la série de fonctions $\sum f_n$. Elle converge vers S définie par On considère la série de fonctions sur $]0, 1[$:

$$Sx \mapsto \sum_{n=0}^{+\infty} x^{2n} \ln x = \frac{\ln x}{1-x^2}$$

Pour $n > 0$, les fonctions f_n se prolongent par continuité en 0 donc elle sont intégrables sur $[0, 1[$. Pour f_0 , en 0, $f_0(x) = \ln(x) = o\left(\frac{1}{\sqrt{x}}\right)$ car $\sqrt{x} \ln(x) \xrightarrow{x \rightarrow 0} 0$. Donc f_0 est aussi intégrable.

De plus, pour tout entier n ,

$$\int_0^1 x^{2n} \ln x \, dx = \left[\frac{x^{2n+1} \ln x}{2n+1} \right]_0^1 - \int_0^1 \frac{x^{2n} dx}{2n+1} = -\frac{1}{(2n+1)^2}$$

On en déduit que la série $\sum \int_I |f_n|$ converge.

En appliquant le théorème, $S : x \mapsto \frac{\ln x}{1-x^2}$ est intégrable et

$$\int_0^1 \frac{\ln x}{1-x^2} = \sum_{n=0}^{+\infty} -\frac{1}{(2n+1)^2} = -\frac{3\zeta(2)}{4}.$$

Exercice : Calculer $\int_0^1 \frac{\ln(1+x)}{x} dx$. On pourra utiliser que pour $x \in]0, 1[$, $\ln(1+x) = \sum_{k=1}^{\infty} (-1)^{k-1} \frac{x^k}{k}$.

Exemple : Il est possible que le théorème d'intégration terme à terme ne puisse pas s'appliquer mais que l'on peut quand même obtenir une interversion terme à terme en calculant les sommes partielles et en utilisant le théorème de convergence dominée.

Pour $n \geq 1$, on considère $g_n : t \mapsto \frac{(-1)^n}{(1+t^2)^n}$ définie sur $]0, +\infty[$. On s'intéresse à la série de fonctions $\sum_{n \geq 1} g_n$. On reconnaît une série géométrique de raison $-\frac{1}{1+t^2}$ et $\left| -\frac{1}{1+t^2} \right| \leq \frac{1}{2}$ car $t > 0$. On en déduit que la série de fonctions $\sum_{n \geq 1} g_n$ converge vers G où

$$G : t \mapsto \sum_{n=1}^{\infty} \frac{(-1)^n}{(1+t^2)^n} = -\frac{1}{1+t^2} \times \frac{1}{1+\frac{1}{1+t^2}} = -\frac{1}{2+t^2}$$

On a vu précédemment que la série $\sum_{n \geq 1} \int_0^{\infty} \frac{1}{(1+t^2)^n}$ divergeait. On ne peut pas appliquer le théorème d'intégration terme à terme. Par contre, on peut le faire en écrivant les sommes partielles.

Notons $S_N = \sum_{n=1}^N g_n$:

$$S_N(t) = \sum_{n=1}^N \frac{(-1)^n}{(1+t^2)^n} = \frac{1 - \left(\frac{-1}{1+t^2}\right)^{N+1}}{2+t^2}$$

On voit que la suite (S_N) converge simplement vers la fonction $G : t \mapsto \frac{1}{2+t^2}$. De plus, pour tout entier N ,

$$|S_N(t)| \leq \frac{1 + \frac{1}{1+t^2}}{2+t^2} \leq \frac{2}{2+t^2}$$

La fonction $t \mapsto \frac{2}{2+t^2}$ est intégrable. On peut donc utiliser le théorème de convergence dominée :

$$\sum_{n=1}^{\infty} \int_0^{\infty} \frac{(-1)^n}{(1+t^2)^n} dt = \lim_{N \rightarrow +\infty} \sum_{n=1}^N \int_0^{\infty} \frac{(-1)^n}{(1+t^2)^n} dt = \lim_{N \rightarrow +\infty} \int_0^{\infty} S_N(t) dt \stackrel{CD}{=} - \int_0^{\infty} \frac{dt}{2+t^2} = -\frac{\pi}{2\sqrt{2}}$$

Polynôme caractéristique et réduction

1	Polynôme caractéristique	109
1.1	Polynôme caractéristique d'une matrice carrée	109
1.2	Polynôme caractéristique d'un endomorphisme	112
1.3	Multiplicités	113
1.4	Polynôme caractéristique d'une matrice triangulaire	114
1.5	Polynôme caractéristique d'un endomorphisme induit	114
2	Matrices et endomorphismes trigonalisables	116
2.1	Définitions	117
2.2	Critère de trigonalisabilité	117
3	Matrices et endomorphismes nilpotents	122
3.1	Définitions	122
3.2	Critère de nilpotence	123

Dans ce chapitre K désigne R ou C . La plupart des résultats restent vraies pour n'importe quel corps.

1 Polynôme caractéristique

1.1 Polynôme caractéristique d'une matrice carrée

On a vu que si $A \in \mathcal{M}_n(K)$ et $\lambda \in K$,

$$\lambda \in \text{Sp}(A) \iff A - \lambda I_n \notin \text{GL}_n(K) \iff \det(A - \lambda I_n) = 0$$

Définition 4.1 (Polynôme caractéristique)

Soit $A \in \mathcal{M}_n(K)$ une matrice carrée on appelle polynôme caractéristique de A et on note χ_A le polynôme

$$\chi_A = (-1)^n \det(A - XI_n) = \det(XI_n - A).$$

Remarque : Pour définir proprement ce qui est au dessus, on travaille avec $A - XI_n \in \mathcal{M}_n(\mathbb{K}[X])$. On peut construire un déterminant comme dans le cas des matrices à coefficients dans un corps.

Proposition 4.2

Avec les notations précédentes, χ_A est un polynôme unitaire de degré n .
De plus $\chi_A(0) = (-1)^n \det(A)$.

Démonstration : On note $A = (a_{ij})$. On a alors

$$\chi_A = \begin{vmatrix} X - a_{11} & -a_{12} & \cdots & -a_{1n} \\ -a_{21} & X - a_{22} & \cdots & -a_{2n} \\ \vdots & & \ddots & \vdots \\ -a_{n1} & -a_{n2} & \cdots & X - a_{nn} \end{vmatrix}$$

On pose pour tout $(i, j) \in \llbracket 1; n \rrbracket^2$,

$$b_{ij} = \begin{cases} X - a_{ii} & \text{si } i = j \\ -a_{ij} & \text{si } i \neq j \end{cases}$$

On remarque en particulier que si i et j sont différents, b_{ij} est un polynôme de degré 0 alors que pour $i = j$, b_{ii} est un polynôme de degré 1. Par la définition du déterminant d'une matrice,

$$\chi_A = \sum_{\sigma \in \mathfrak{S}_n} \varepsilon(\sigma) b_{\sigma(1)1} \cdots b_{\sigma(n)n}$$

Or, pour tout $\sigma \in \mathfrak{S}_n$, $b_{\sigma(1)1} \cdots b_{\sigma(n)n}$ est un polynôme de degré au plus n (comme produit de n polynômes de degré au plus 1). On en déduit que χ_A est un polynôme de degré au plus n .

De plus, le seul terme qui est de degré n est celui obtenu pour $\sigma = \text{id}$:

$$b_{11} \times \cdots \times b_{nn} = (X - a_{11})(X - a_{22}) \cdots (X - a_{nn}) = X^n + \cdots$$

On en déduit bien que χ_A est unitaire de degré n .

Pour finir, $\chi_A(0) = \det(-A) = (-1)^n \det(A)$. □

Remarque : Il arrive que l'on trouve aussi la définition $\chi_A = \det(A - XI_n)$ (sans le $(-1)^n$). Cela ne change pas beaucoup (à part qu'il n'est plus toujours unitaire) car on va surtout s'intéresser aux racines.

Proposition 4.3

Soit $A \in \mathcal{M}_n(\mathbb{K})$ et $\lambda \in \mathbb{K}$.

$$\lambda \in \text{Sp}(A) \iff \chi_A(\lambda) = 0.$$

Exemples :

1. Soit $A = \begin{pmatrix} 0 & 2 & -1 \\ 3 & -2 & 0 \\ -2 & 2 & 1 \end{pmatrix}$

Son polynôme caractéristique est

$$\chi_A = (-1)^3 \begin{vmatrix} -X & 2 & -1 \\ 3 & -2-X & 0 \\ -2 & 2 & 1-X \end{vmatrix} = X^3 + X^2 - 10X + 8 = (X-1)(X+4)(X-2)$$

On a donc $\text{Sp}(A) = \{-4, 1, 2\}$.

2. Soit $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$. On a

$$\chi_A(X) = \begin{vmatrix} X-a & -b \\ -c & X-d \end{vmatrix} = (X-a)(X-d) - bc = X^2 - (a+d)X + (ad-bc)$$

On a donc

$$\chi_A(x) = X^2 - \text{tr}(A)X + \det(A).$$

Proposition 4.4

Soit $A \in \mathcal{M}_n(\mathbf{K})$ et χ_A son polynôme caractéristique. Le coefficient de degré $n-1$ de χ_A est $-\text{tr}(A)$. On a donc

$$\chi_A = X^n - \text{tr}(A)X^{n-1} + \dots + (-1)^n \det(A)$$

Démonstration : On note $C_1, \dots, C_n$ les colonnes de la matrice A et $E_1, \dots, E_n$ les vecteurs de la bases canonique de $\mathcal{M}_{n,1}(\mathbf{K})$. On a donc

$$E_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, E_2 = \begin{pmatrix} 0 \\ 1 \\ 0 \\ \vdots \\ 0 \end{pmatrix}, \dots, E_n = \begin{pmatrix} 0 \\ \vdots \\ \vdots \\ 0 \\ 1 \end{pmatrix}$$

On a donc

$$\chi_A = (-1)^n \det(C_1 - XE_1 | \dots | C_n - XE_n)$$

Si on développe le déterminant par mutlinéarité on obtient

$$\begin{aligned} \chi_A &= \det(XE_1 - C_1 | \dots | XE_n - C_n) \\ &= \det(XE_1 | \dots | XE_n) + \sum_{i=1}^n \det(XE_1 | \dots | -C_i | \dots | XE_n) + \dots + \det(-C_1 | \dots | -C_n) \\ &= \det(E_1 | \dots | E_n)X^n - \sum_{i=1}^n \det(E_1 | \dots | C_i | \dots | E_n)X^{n-1} + \dots + (-1)^n \det(A) \end{aligned}$$

Maintenant, $\det(E_1 | \dots | E_n) = 1$ car la matrice $(E_1 | \dots | E_n)$ est la matrice identité. On retrouve que χ_A est unitaire.

De même

$$(E_1 | \dots | C_i | \dots | E_n) = \begin{pmatrix} 1 & & a_{1i} & 0 & \dots & 0 \\ 0 & \ddots & \vdots & \vdots & & \vdots \\ \vdots & \ddots & 1 & \vdots & \vdots & \vdots \\ \vdots & & 0 & a_{ii} & 0 & \vdots \\ \vdots & & \vdots & \vdots & 1 & \ddots \\ \vdots & & \vdots & \vdots & & \ddots & 0 \\ 0 & \dots & 0 & a_{ni} & & & 1 \end{pmatrix}$$

Si bien qu'en développant le déterminant par rapport aux colonnes, $\det(E_1 | \cdots | C_i | \cdots | E_n) = a_{ii}$ On en déduit que le terme de degré X^{n-1} vaut $-\sum_{i=1}^n a_{ii} = -\text{tr}(A)$. $\square$

Exercice : En procédant de même exprimer le coefficient de degré 1 à l'aide de la comatrice de A ¹

Exemple : Si on applique cela à la matrice de l'exemple précédent, on a bien $\text{tr}(A) = -1$.

Proposition 4.5

Soit A et B deux matrices semblables,

$$\chi_A = \chi_B.$$

Démonstration : Comme A et B sont semblables, il existe $P \in \text{GL}_n(\mathbf{K})$ telle que $B = P^{-1}AP$.

Maintenant,

$$\chi_B = \det(XI_n - B) = \det(XP^{-1}P - P^{-1}AP) = \det(P^{-1}(XI_n - A)P) = \det(P^{-1})\chi_A \det(P) = \chi_A$$

$\square$

1.2 Polynôme caractéristique d'un endomorphisme

On a vu dans le paragraphe précédent que le polynôme caractéristique de deux matrices semblables était le même. Cela signifie que le polynôme caractéristique dépend de l'endomorphisme sous-jacent et pas de la matrice elle-même.

Dans ce paragraphe on se fixe un espace vectoriel E de dimension finie n .

Définition 4.6 (Polynôme caractéristique d'un endomorphisme)

Soit $u \in \mathcal{L}(E)$ un endomorphisme de E . On appelle polynôme caractéristique de u et on note χ_u le polynôme caractéristique d'une de ses matrices.

$$\chi_u = \chi_{\text{Mat}_{\mathcal{B}}(u)}$$

Remarque : Cette définition est possible du fait que si A et B sont semblables alors $\chi_A = \chi_B$.

Exemple : Soit F un sous-espace vectoriel de dimension p et G un supplémentaire. Soit p la projection sur F par rapport à G . On cherche χ_p . On se place dans une base adaptée $\mathcal{B}$ obtenue en concaténant une base de F avec une base de G . On a alors

$$\text{Mat}_{\mathcal{B}}(p) = \left(\begin{array}{ccc|cc} 1 & & 0 & & \\ & \ddots & & & \\ 0 & & 1 & & 0 \\ \hline & & & 0 & 0 \end{array} \right)$$

1. On trouve $\text{tr}(\text{Com}(A))$

On a donc $\chi_p = (X - 1)^p X^{n-p}$

Exercices :

1. Déterminer le polynôme caractéristique d'une symétrie.
2. Déterminer le polynôme caractéristique de $\Delta : \mathbf{R}_3[X] \rightarrow \mathbf{R}_3[X]$ défini par $\Delta : P \mapsto P'$.

Proposition 4.7

Soit $u \in \mathcal{L}(E)$. Les racines de χ_u sont les valeurs propres de u .

Démonstration : Il suffit de considérer la matrice de u dans une base et d'utiliser le résultat analogue sur les matrices. $\square$

1.3 Multiplicités

Définition 4.8 (Ordre de multiplicité)

Soit $u \in \mathcal{L}(E)$ et $\lambda \in \mathbf{K}$. L'ordre de multiplicité de λ (par rapport à u) est l'ordre de multiplicité de λ dans χ_u .

Remarques :

1. On parle aussi de multiplicité de λ . On dit aussi, « λ est valeur propre d'ordre ... ».
2. Un scalaire de multiplicité 0 n'est pas une valeur propre. Les valeurs propres sont les scalaires de multiplicité strictement positive.
3. De même si $A \in \mathcal{M}_n(\mathbf{K})$ et $\lambda \in \mathbf{K}$. L'ordre de multiplicité de λ (par rapport à A) est l'ordre de multiplicité de λ dans χ_A .

Exemples :

1. Si on reprend notre projecteur de l'exemple ci-dessus. On voit que 1 est valeur propre d'ordre p et 0 est valeur propre d'ordre $n - p$.
2. Pour $A = \begin{pmatrix} 0 & 2 & -1 \\ 3 & -2 & 0 \\ -2 & 2 & 1 \end{pmatrix}$. Les scalaires $-4, 1$ et 2 sont valeurs propres de multiplicité 1. Elles sont dites valeurs propres simples.

Proposition 4.9

Soit $u \in \mathcal{L}(E)$ (resp. $A \in \mathcal{M}_n(\mathbf{K})$). Le nombre de valeur propres comptées avec multiplicités est inférieur ou égal à $n = \dim E$. Si $\mathbf{K} = \mathbf{C}$ le nombre de valeur propres comptées avec multiplicités est égal à $n = \dim E$.

Démonstration : Il suffit de compter les racines de χ_u (resp. χ_A) qui est de degré n . $\square$

1.4 Polynôme caractéristique d'une matrice triangulaire

Corollaire 4.10 (Polynôme caractéristique d'une matrice triangulaire par blocs)

Soit $A = \begin{pmatrix} A_{11} & * & \cdots & * \\ 0 & \ddots & \ddots & \vdots \\ \vdots & & \ddots & * \\ 0 & \cdots & 0 & A_{nn} \end{pmatrix}$ une matrice triangulaire par blocs :

$$\chi_A = \prod_{k=1}^n \chi_{A_{kk}}$$

Démonstration : Il suffit de voir que $XI_n - A$ est encore triangulaire par blocs. $\square$

1.5 Polynôme caractéristique d'un endomorphisme induit

Dans ce paragraphe on se fixe un espace vectoriel E de dimension finie n .

Proposition 4.11

Soit $u \in \mathcal{L}(E)$ et F un sous-espace stable par F , le polynôme caractéristique de l'endomorphisme $u|_F$ induit par u sur F divise le polynôme caractéristique de u :

$$\chi_{u|_F} | \chi_u.$$

Démonstration : Soit G un supplémentaire de F et $\mathcal{B}$ une base adaptée à la décomposition $E = F \oplus G$ obtenue en concaténant une base $\mathcal{F}$ de F et une base $\mathcal{G}$ de G . Dans la base $\mathcal{B}$, la matrice de u est

$$M = \text{Mat}_{\mathcal{B}}(u) = \left(\begin{array}{c|c} A & B \\ \hline 0 & D \end{array} \right)$$

où $A = \text{Mat}_{\mathcal{F}}(u|_F)$. On a donc $\chi_u = \chi_M = \chi_A \times \chi_D$. Donc $\chi_{u|_F} = \chi_A$ divise χ_u . $\square$

Exercice : Soit u l'endomorphisme de $\mathbf{R}^3$ dont la matrice dans la base canonique est

$$A = \begin{pmatrix} 1 & -2 & 1 \\ 4 & 7 & -1 \\ 4 & 1 & -3 \end{pmatrix}$$

1. Montrer que $H = \{(x, y, z) \mid x + y = 0\}$ est stable par u .
2. Calculer χ_A .

on trouve : $X^3 - 5X^2 - 12X + 60$

3. Calculer la matrice de $u|_H$.

Dans la base $((1, -1, 0), (1, -1, 1))$ on trouve $\begin{pmatrix} 0 & 4 \\ 3 & 0 \end{pmatrix}$.

4. Vérifier que $\chi_{u|_H} | \chi_u$.

On a $\chi_{u|_H} = X^2 - 12$.

Théorème 4.12

Soit $u \in \mathcal{L}(E)$ (resp. $A \in \mathcal{M}_n(\mathbf{K})$). Soit $\lambda \in \mathbf{K}$ une valeur propre de u (resp. A) de multiplicité p . La dimension du sous-espace propre $E_\lambda(u)$ (resp. $E_\lambda(A)$) est inférieure ou égale à p .

Démonstration : Traitons le cas de l'endomorphisme u . On sait que $F = E_\lambda(u)$ est stable par u et que $u|_F$ est l'homothétie $x \mapsto \lambda.x$. Sa matrice est donc une matrice scalaire. Notons $q = \dim F$, on a $\chi_{u|_F} = (X - \lambda)^q$.

Comme on sait que $\chi_{u|_F} | \chi_u$ on a bien $q \leq p$. □

Définition 4.13 (Rappels - Polynôme scindé)

Soit P un polynôme sur $\mathbf{K}$. On dit qu'il est scindé (sur $\mathbf{K}$) s'il se décompose en produit de polynômes de degré 1. C'est-à-dire qu'il existe $\lambda \in \mathbf{K}$ et $(\alpha_1, \dots, \alpha_p) \in \mathbf{K}^p$ deux à deux distincts et $(r_1, \dots, r_p) \in (\mathbf{N}^*)^p$ tels que

$$P = \lambda \prod_{i=1}^p (X - \alpha_i)^{r_i}$$

Dans le cas où pour tout $i \in \llbracket 1; p \rrbracket$, $r_i = 1$, on dit que P est scindé à racines simples ou qu'il est simplement scindé.

Exemples :

1. Si $\mathbf{K} = \mathbf{C}$, tout polynôme est scindé.
2. Pour $\mathbf{K} = \mathbf{R}$, $X^2 + 1$ n'est pas scindé.
3. Le polynôme $X^2 - 1 = (X + 1)(X - 1)$ est scindé à racines simples.

Théorème 4.14

Soit $u \in \mathcal{L}(E)$. Il est diagonalisable si et seulement si χ_u est scindé et pour tout $\lambda \in \text{Sp}(u)$ la multiplicité de λ est la dimension de $E_\lambda(u)$.

ATTENTION

Ce théorème exprime qu'il peut exister deux obstructions au fait qu'un endomorphisme soit diagonalisable.

- Il « manque » des valeurs propres. Cela se voit par le fait que le polynôme caractéristique n'est pas scindé. Par exemple, si $\chi_u = X^2 + 1$ et que $\mathbf{K} = \mathbf{R}$.
Cela ne peut pas se produire si $\mathbf{K} = \mathbf{C}$.
- Il y a le bon nombres de valeurs propres mais les espaces propres sont « trop petits ».

Par exemple si u est l'endomorphisme associé à $\begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \end{pmatrix}$.

Démonstration :

- $\Rightarrow$ On suppose que u est diagonalisable. Il existe donc une base $\mathcal{B}$ dans laquelle la matrice est diagonale. On a alors

$$\text{Mat}_{\mathcal{B}}(u) = \begin{pmatrix} \lambda_1 & & & & \\ & \ddots & & & \\ & & \lambda_1 & & \\ & & & \ddots & \\ & & & & \lambda_p & \\ & & & & & \ddots \\ & & & & & & \lambda_p \end{pmatrix}$$

On en déduit que $\chi_u = \prod_{i=1}^p (X - \lambda_i)^{r_i}$ où r_i est la dimension du sous-espace propre associé à λ_i . On a bien χ_u scindé et pour tout i , la multiplicité de λ_i vaut la dimension du sous-espace propre associé.

- $\Leftarrow$ On suppose que χ_u est scindé et pour tout $\lambda \in \text{Sp}(u)$ la multiplicité de λ vaut la dimension du sous-espace propre associé. Si on note $\lambda_1, \dots, \lambda_p$ les valeurs propres et $r_1, \dots, r_p$ les multiplicités / dimensions.

On sait que $\bigoplus_{i=1}^p E_{\lambda_i}(u) \subset E$ mais

$$\dim \left(\bigoplus_{i=1}^p E_{\lambda_i}(u) \right) = r_1 + \dots + r_p = \deg \chi_u = n = \dim E.$$

On a donc $\bigoplus_{i=1}^p E_{\lambda_i}(u) = E$ ce qui signifie que l'endomorphisme est diagonalisable.

□

Corollaire 4.15

Soit $u \in \mathcal{L}(E)$. Si χ_u est simplement scindé alors u est diagonalisable.

Remarque : Nous verrons dans un chapitre ultérieur que cet énoncé est une équivalence.

2 Matrices et endomorphismes trigonalisables

Dans ce paragraphe on se fixe un espace vectoriel E de dimension finie n .

On a vu que certains endomorphismes (certaines matrices) étaient diagonalisables. Cependant il existe des endomorphismes (matrices) qui ne le sont pas. Par exemple

$$A = \begin{pmatrix} 2 & 1 \\ 0 & 2 \end{pmatrix}$$

En effet le spectre de cette matrice est $\text{Sp}(A) = 2$. Donc si elle était diagonalisable elle serait semblable à $\begin{pmatrix} 2 & 0 \\ 0 & 2 \end{pmatrix} = 2I_2$. Cependant ce n'est pas le cas car $2I_2$ est la seule matrice semblable à elle-même (car pour tout $P \in \text{GL}_2(\mathbf{K})$, $P^{-1} \cdot 2I_2 \cdot P = 2P^{-1}P = 2I_2$).

2.1 Définitions

Proposition 4.16 (Matrices et endomorphismes trigonalisables)

1. Soit $A \in \mathcal{M}_n(\mathbf{K})$. Elle est dit trigonalisable si elle est semblable à une matrice triangulaire supérieure. C'est-à-dire s'il existe $P \in \text{GL}_n(\mathbf{K})$ telle que $P^{-1}AP$ soit triangulaire.
2. Soit $u \in \mathcal{L}(E)$. Il est dit trigonalisable s'il existe une base $\mathcal{B}$ tel que $\text{Mat}_{\mathcal{B}}(u)$ soit triangulaire supérieure. C'est-à-dire si ses matrices sont trigonalisables.

Remarque : Une matrice A est trigonalisable si et seulement si son endomorphisme canoniquement associé est trigonalisable.

Interprétation géométrique

Soit $u \in \mathcal{L}(E)$ un endomorphisme trigonalisable et soit $\mathcal{B} = (e_1, \dots, e_n)$ une base telle que $A = \text{Mat}_{\mathcal{B}}(u)$ soit triangulaire. On pose pour tout $k \in \llbracket 1; n \rrbracket$, $F_k = \text{Vect}(e_1, \dots, e_k)$. Les espaces F_k sont stables par u . En effet, du fait que A soit triangulaire supérieure, pour tout $j \in \llbracket 1; n \rrbracket$, il existe $a_{1j}, \dots, a_{jj}$ tels que

$$u(e_j) = \sum_{i=1}^j a_{ij}e_i$$

En particulier, si on fixe k , pour tout $j \leq k$, $u(e_j) \in F_k$ et donc $u(F_k) \subset F_k$. On dit alors que la suite

$$\{0\} \subsetneq F_1 \subsetneq \dots \subsetneq F_{n-1} \subsetneq F_n = E$$

est un drapeau stable.

2.2 Critère de trigonalisabilité

Théorème 4.17

Soit $u \in \mathcal{L}(E)$ (resp. $A \in \mathcal{M}_n(\mathbf{K})$). Il (resp. elle) est trigonalisable si et seulement si son polynôme caractéristique est scindé.

Démonstration :

- $\Leftarrow$ On suppose que u est trigonalisable. Il existe donc une base $\mathcal{B}$ de E telle que $A = \text{Mat}_{\mathcal{B}}(u)$ soit triangulaire supérieure.

$$A = \begin{pmatrix} \lambda_1 & * & \cdots & * \\ & \ddots & \ddots & \vdots \\ & & \ddots & * \\ & & & \lambda_n \end{pmatrix}$$

On en déduit que $\chi_u = \det(XI_n - A) = \prod_{i=1}^n (X - \lambda_i)$. C'est bien un polynôme scindé

- $\Rightarrow$ On procède par récurrence sur n . On pose $\mathcal{P}_n =$ « toute matrice de $\mathcal{M}_n(\mathbf{K})$ ayant un polynôme caractéristique scindé est trigonalisable ».
- Initialisation : Pour $n = 1$, toute matrice de $\mathcal{M}_1(\mathbf{K})$ est triangulaire donc $\mathcal{P}_1$ est vérifié.²

2. On pourrait commencer à 0, la matrice vide de $\mathcal{M}_0(\mathbf{K})$ étant aussi triangulaire.

— Hérédité : Soit $n \geq 1$, on suppose $\mathcal{P}_n$. Soit $A \in \mathcal{M}_{n+1}(\mathbf{K})$ telle que χ_A est scindé. On pose

$$\chi_A = \prod_{i=1}^{n+1} (X - \lambda_i).$$

En particulier λ_{n+1} est une racine de χ_A , c'est donc une valeur propre de A . De ce fait A est semblable à

$$B = \left(\begin{array}{c|c} \lambda_{n+1} & C \\ \hline 0 & B' \\ \vdots & \\ 0 & \end{array} \right).$$

Maintenant $\chi_A = \chi_B = (X - \lambda_{n+1})\chi_{B'}$ donc $\chi_{B'} = \prod_{i=1}^n (X - \lambda_i)$. Il est donc scindé. D'après $\mathcal{P}_n$ la matrice B' est semblable à une matrice triangulaire. Il existe donc $P \in \text{GL}_n(\mathbf{K})$ telle que $P^{-1}B'P = T$ soit triangulaire supérieure. On « prolonge » la matrice P pour obtenir une matrice de changement de bases de taille $(n+1)$ en posant

$$Q = \left(\begin{array}{c|ccc} 1 & 0 & \cdots & 0 \\ \hline 0 & & & \\ \vdots & & P & \\ 0 & & & \end{array} \right).$$

Il est clair que $Q^{-1} = \left(\begin{array}{c|ccc} 1 & 0 & \cdots & 0 \\ \hline 0 & & & \\ \vdots & & P^{-1} & \\ 0 & & & \end{array} \right)$ et donc

$$Q^{-1}BQ = \left(\begin{array}{c|ccc} \lambda_{n+1} & * & \cdots & * \\ \hline 0 & & & \\ \vdots & & P^{-1}B'P & \\ 0 & & & \end{array} \right) = \left(\begin{array}{c|ccc} \lambda_{n+1} & * & \cdots & * \\ \hline 0 & & & \\ \vdots & & T & \\ 0 & & & \end{array} \right)$$

Cette dernière est bien triangulaire donc B et de ce fait A est trigonalisable.

Par récurrence, $\forall n \in \mathbf{N}$, $\mathcal{P}_n$.

□

ATTENTION

Il est important de bien comprendre cette démonstration car c'est un exemple classique de raisonnement par récurrence sur la dimension. De plus cela donne la méthode pour trigonaliser une matrice.

Corollaire 4.18

Pour $\mathbf{K} = \mathbf{C}$, tout endomorphisme (resp. toute matrice) est trigonalisable.

Remarque : Cela concerne aussi les matrices à coefficients réels vue comme élément de $\mathcal{M}_n(\mathbb{C})$ c'est à dire en acceptant des valeurs propres complexes. Par exemple

$$\begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}.$$

Exemples :

1. Pour $n = 2$, si une matrice a une valeur propre elle est trigonalisable. En effet χ_A est de degré 2. S'il y a une valeur propre λ , il est divisible par $X - \lambda$. De ce fait $\chi_A = (X - \lambda)Q$ où Q est unitaire de degré 1 donc de la forme $X - \mu$. Il suffit alors de trouver un vecteur propre u et de se placer dans une base de la forme (u, v) .

Pour $A = \begin{pmatrix} 2 & -1 \\ 7 & -3 \end{pmatrix}$. On trouve $\chi_A = X^2 + X + 1$.

On voit que A n'est pas trigonalisable sur $\mathbb{R}$ car $\text{Sp}_{\mathbb{R}}(A) = \emptyset$. Par contre $\text{Sp}_{\mathbb{C}}(A) = \{j, j^2\}$.

De plus $E_j(A) = \text{Vect} \begin{pmatrix} 1 \\ 2 - j \end{pmatrix}$. On peut faire le changement de base donné par

$$P = \begin{pmatrix} 1 & 0 \\ 2 - j & 1 \end{pmatrix}$$

On a

$$P^{-1} = \begin{pmatrix} 1 & 0 \\ -2 + j & 1 \end{pmatrix}$$

puis

$$P^{-1}AP = \begin{pmatrix} j & -1 \\ 0 & j^2 \end{pmatrix}$$

car

$$A \begin{pmatrix} 0 \\ 1 \end{pmatrix} = \begin{pmatrix} -1 \\ -3 \end{pmatrix} = - \begin{pmatrix} 1 \\ 2 - j \end{pmatrix} + \underbrace{(-1 - j)}_{j^2} \begin{pmatrix} 0 \\ 1 \end{pmatrix}.$$

On aurait aussi pu prendre pour le deuxième vecteur de la base un vecteur propre associé à la valeur propre j^2 et on aurait obtenu une matrice diagonale.

2. Soit $A = \begin{pmatrix} 2 & -1 & -1 \\ 2 & 1 & -2 \\ 3 & -1 & -2 \end{pmatrix}$. On trouve $\chi_A = (X + 1)(X - 1)^2$. La matrice est bien trigonalisable (dans $\mathbb{R}$).

On cherche les sous-espaces propres associés aux valeurs propres 1 et -1 . En résolvant on obtient

$$E_{-1}(A) = \text{Vect} \begin{pmatrix} 1 \\ 1 \\ 2 \end{pmatrix} \text{ et } E_1(A) = \text{Vect} \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}$$

On voit que $E_1(A) \oplus E_{-1}(A) \subsetneq \mathcal{M}_{3,1}(\mathbb{K})$ donc A n'est pas diagonalisable. On considère une base

(u, v, w) où $u = (1, 1, 2)$ et $v = (1, 0, 1)$. En effectuant le changement de base : $P = \begin{pmatrix} 1 & 1 & * \\ 1 & 0 & * \\ 2 & 1 & * \end{pmatrix}$

on obtiendra

$$A' = \begin{pmatrix} -1 & 0 & * \\ 0 & 1 & * \\ 0 & 0 & 1 \end{pmatrix}$$

Le coefficient en bas à droite est obtenu en vérifiant que le déterminant (ou la trace ou le polynôme caractéristique) de A' est celui de A .

Par exemple en prenant $w = (1, 0, 0)$, comme

$$A \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix} = \begin{pmatrix} 2 \\ 2 \\ 3 \end{pmatrix} = 2 \begin{pmatrix} 1 \\ 1 \\ 2 \end{pmatrix} - \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} + \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}$$

Donc

$$A' = \begin{pmatrix} -1 & 0 & 2 \\ 0 & 1 & -1 \\ 0 & 0 & 1 \end{pmatrix}$$

Remarque : On pourrait aussi prendre w tel que $Aw = w + v$ ce qui simplifierait la troisième colonne.

Corollaire 4.19

Soit $u \in \mathcal{L}(E)$ (resp. $A \in \mathcal{M}_n(\mathbf{K})$) un endomorphisme (resp. une matrice) **trigonalisable**. On note $\{\lambda_1, \dots, \lambda_p\}$ les valeurs propres et $r_1, \dots, r_p$ les multiplicités. On a

$$\chi_u = \prod_{i=1}^p (X - \lambda_i)^{r_i}$$

et

$$\text{tr}(u) = \sum_{i=1}^p r_i \lambda_i \text{ et } \det(u) = \prod_{i=1}^p \lambda_i^{r_i}.$$

Remarques :

1. On peut plus rapidement dire que la trace est l'opposée de la somme des valeurs propres et que le déterminant est le produit des valeurs propres **comptées avec leur multiplicité**.
2. On suppose que u est trigonalisable et qu'il a une valeur propre (simple) de module strictement supérieur aux autres. C'est-à-dire que si on note $\lambda_1, \dots, \lambda_n$ les valeurs propres avec multiplicité, on a $\text{tr}(u) = \lambda_1 + \dots + \lambda_n$. Maintenant si on choisit $\mathcal{B}$ tel que

$$\text{Mat}_{\mathcal{B}}(u) = \begin{pmatrix} \lambda_1 & * & \cdots & * \\ & \ddots & \ddots & \vdots \\ & & \ddots & * \\ & & & \lambda_n \end{pmatrix}$$

on a pour les puissance de u ,

$$\text{Mat}_{\mathcal{B}}(u^p) = \begin{pmatrix} \lambda_1^p & * & \cdots & * \\ & \ddots & \ddots & \vdots \\ & & \ddots & * \\ & & & \lambda_n^p \end{pmatrix}.$$

De ce fait $\text{tr}(u^p) = \sum_{i=1}^n \lambda_i^p$. Avec nos hypothèses

$$\frac{\text{tr}(u^{p+1})}{\text{tr}(u^p)} = \frac{\lambda_1^{p+1} + \dots + \lambda_n^{p+1}}{\lambda_1^p + \dots + \lambda_n^p} \underset{p \rightarrow \infty}{\sim} \frac{\lambda_1^{p+1}}{\lambda_1^p} = \lambda_1$$

On voit que l'on peut trouver la valeur propre de plus grand module en étudiant $\lim_{p \rightarrow +\infty} \frac{\text{tr}(u^{p+1})}{\text{tr}(u^p)}$

Exemple : Soit A une matrice de rang 1. On note $t = \text{tr}(A)$.

- Si $t \neq 0$. On sait d'après le théorème du rang que $\dim \text{Ker}(A) = n - 1$ ce qui signifie que 0 est valeur propre de multiplicité (au moins) $n - 1$. Comme la somme des valeurs propres avec multiplicités vaut $t \neq 0$, nécessairement t est valeur propre de multiplicité 1. Cela montre que A est diagonalisable et que $\chi_A = X^{n-1}(X - t)$.

Posons par exemple $H = \begin{pmatrix} 1 & \cdots & 1 \\ \vdots & & \vdots \\ 1 & \cdots & 1 \end{pmatrix}$. Elle est bien de rang 1 et $\text{tr}(H) = n$. On en déduit que

H est diagonalisable et que $\chi_H = X^{n-1}(X - n)$. On peut de plus facilement calculer une base de vecteurs propres de H . Si on note $X_1, \dots, X_n$ la base canonique de $\mathcal{M}_{n,1}(\mathbb{K})$,

$$E_n(A) = \text{Vect}(X_1 + \cdots + X_n) \text{ et } E_0(A) = \text{Vect}(X_2 - X_1, X_3 - X_1, \dots, X_n - X_1)$$

- Si $t = 0$, à l'inverse, 0 est la seule valeur propre qui est donc de multiplicité n . On a donc $\chi_A = X^n$. La matrice A n'est pas diagonalisable.

Exercice : Ecrire un programme python (ou caml) pour l'algorithme qui précède.

Le tester pour

$$A_1 = \begin{pmatrix} 2 & -1 & -1 \\ 2 & 1 & -2 \\ 3 & -1 & -2 \end{pmatrix}; A_2 = \begin{pmatrix} -1 & 2 & -1 \\ 3 & -3 & 0 \\ -2 & 2 & 0 \end{pmatrix}; A_3 = \begin{pmatrix} 11 & 6 & -0 \\ -18 & -10 & 0 \\ 7 & 2 & 2 \end{pmatrix}$$

```
import numpy as np
import numpy.linalg as alg

# Testons les programmes

def trace(a) :
    n = len(a)
    s = 0
    for i in range(n) :
        s += a[i,i]
    return(s)

def quotientTrace(a,p) :
    tr1,tr 2 = 1,1
    mat = np.array(a)
    puisMat = mat
    tr2 = trace(mat)
    for i in range(p) :
        puisMat = np.dot(puisMat,mat)
        tr1 = tr2
        tr2 = trace(puisMat)
    return( float(tr2) / float(tr1))

a1 = np.array([[ -1,2,-1],[3,-3,0],[-2,2,0]])
a2 = np.array([[11,6,0],[-18,-10,0],[7,2,2]])
a3 = np.array([[2,-1,-1],[2,1,-2],[3,-1,-2]])

print(alg.eigvals(a1))
#Le spectre de A1 est -5 , 1 , 0
print(quotientTrace(a1,10))

print(alg.eigvals(a2))
#Le spectre de A2 est -1 , 2 , 2
print(quotientTrace(a2,10))

print(alg.eigvals(a3))
#Le spectre de A2 est -1 , 1 , 1
print(quotientTrace(a3,10))
```

Cela donne

0.333333333333
-4.9999993856
1.99853587116

Ce qui est logique car les spectres sont $\text{Sp}(A_1) = \{-1, 1, 1\}$ (il n'y a pas une valeur propre plus grande que les autres en module); $\text{Sp}(A_2) = \{-5, 1, 0\}$ et $\text{Sp}(A_3) = \{-1, 2, 2\}$ (le fait que 2 soit une valeur propre double ne pose pas de problèmes - adapter la preuve dans ce cas).

3 Matrices et endomorphismes nilpotents

Dans ce paragraphe on se fixe un espace vectoriel E de dimension finie n .

3.1 Définitions

Définition 4.20

Soit $u \in \mathcal{L}(E)$ un endomorphisme. Il est dit nilpotent s'il existe un entier p tel que $u^p = 0_{\mathcal{L}(E)}$.

Remarques :

1. Cette définition s'étend aux endomorphismes d'espaces vectoriel non nécessairement de dimension finie.
2. Si p est tel que $u^p = 0_{\mathcal{L}(E)}$ alors pour tout $q \geq p$ on a encore $u^q = 0_{\mathcal{L}(E)}$

Proposition-Définition 4.21

Soit $u \in \mathcal{L}(E)$ un endomorphisme nilpotent. Il existe un plus petit entier p tel que $u^p = 0_{\mathcal{L}(E)}$ c'est-à-dire, $u^p = 0_{\mathcal{L}(E)}$ et $u^{p-1} \neq 0_{\mathcal{L}(E)}$. On l'appelle l'indice de nilpotence de l'endomorphisme.

Démonstration : Soit $u \neq 0_{\mathcal{L}(E)}$. L'ensemble $C = \{k \in \mathbb{N} \mid u^k = 0_{\mathcal{L}(E)}\}$ est non vide (car u est nilpotent). Il admet donc un plus petit élément noté p . On a alors $p \in C$ donc $u^p = 0_{\mathcal{L}(E)}$ et $(p-1) \notin C$ donc $u^{p-1} \neq 0_{\mathcal{L}(E)}$.

Dans le cas où $u = 0_{\mathcal{L}(E)}$, l'indice de nilpotence vaut 0. □

Exemple : Pour $E = \mathbb{K}_n[X]$. L'endomorphisme de dérivation

$$\begin{aligned} \Delta : E &\rightarrow E \\ P &\mapsto P' \end{aligned}$$

est nilpotent car $\Delta^n = 0$.

Définition 4.22

Soit $A \in \mathcal{M}_n(\mathbb{K})$ une matrice. Elle est dit nilpotente s'il existe un entier p tel que $A^p = 0$. On définit de même l'indice de nilpotence d'une matrice nilpotente.

Exemple : Soit $A = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}$. On a $A^2 = 0$. De manière générale si $A = \begin{pmatrix} 0 & 1 & 0 & \cdots & 0 \\ \vdots & \ddots & \ddots & \ddots & \vdots \\ \vdots & & \ddots & \ddots & 0 \\ \vdots & & & \ddots & 1 \\ 0 & \cdots & \cdots & \cdots & 0 \end{pmatrix}$ est

de taille n , on a $A^n = 0$ et $A^{n-1} \neq 0$.

Exercice : Soit $A = \begin{pmatrix} 3 & 9 & -9 \\ 2 & 0 & 0 \\ 3 & 3 & -3 \end{pmatrix}$. Calculer A^2 et A^3 puis trouver $P \in GL_3(\mathbf{K})$ tel que

$$P^{-1}AP = \begin{pmatrix} 0 & * & * \\ 0 & 0 & * \\ 0 & 0 & 0 \end{pmatrix}$$

Proposition 4.23

Soit $u \in \mathcal{L}(E)$. Les assertions suivantes sont équivalentes

- i) L'endomorphisme u est nilpotent.
- ii) Pour tout base $\mathcal{B}$, $\text{Mat}_{\mathcal{B}}(u)$ est nilpotente.
- iii) Il existe une base $\mathcal{B}$ tel que $\text{Mat}_{\mathcal{B}}(u)$ est nilpotente.

Démonstration :

- $i) \Rightarrow ii)$ On suppose que u est nilpotent. Soit $\mathcal{B}$ une base et $A = \text{Mat}_{\mathcal{B}}(u)$. Pour tout entier p , $A^p = \text{Mat}_{\mathcal{B}}(u^p)$ donc pour p supérieur à l'indice de nilpotence de u on a $A^p = 0$. La matrice est nilpotente.
- $ii) \Rightarrow iii)$ Evident.
- $iii) \Rightarrow i)$ Soit $\mathcal{B}$ une base telle que $A = \text{Mat}_{\mathcal{B}}(u)$ soit nilpotente. Pour tout entier p , $A^p = \text{Mat}_{\mathcal{B}}(u^p)$ donc pour p supérieur à l'indice de nilpotence de A on a $\text{Mat}_{\mathcal{B}}(u^p) = 0$ et donc $u^p = 0$. L'endomorphisme u est nilpotent.

□

Remarque : En particulier si A et B sont deux matrices semblables et que A est nilpotente alors B l'est aussi. Cela se voit directement en utilisant que si $B = P^{-1}AP$ alors $B^p = P^{-1}A^pP$.

Exercices :

1. Soit N une matrice nilpotente, montrer que $A = I + N$ est inversible et calculer A^{-1} .
2. Montrer que si A et B sont nilpotentes et que $AB = BA$ alors $A + B$ est nilpotente. Trouver A et B nilpotente telles que $A + B$ ne soit pas nilpotente.
3. Déterminer les matrices diagonalisables et nilpotentes.

3.2 Critère de nilpotence

Théorème 4.24

Soit $u \in \mathcal{L}(E)$. Les assertions suivantes sont équivalentes.

- i) L'endomorphisme u est nilpotent.
- ii) Il existe une base $\mathcal{B}$ telle que $\text{Mat}_{\mathcal{B}}(u)$ soit triangulaire avec des zéros sur la diagonale (triangulaire stricte)
- iii) L'endomorphisme u est trigonalisable et sa seule valeur propre est 0.
- iv) Le polynôme caractéristique de u est $\chi_u = X^n$.

Démonstration :

- $\boxed{ii) \Leftrightarrow iii) \Leftrightarrow iv)}$ Evident.
- $i) \Rightarrow iv)$ Soit $u \in \mathcal{L}(E)$ un endomorphisme nilpotent. On considère $\mathcal{B}$ une base de E et on note $A = \text{Mat}_{\mathcal{B}}(u)$. On peut considérer A comme une matrice de $\mathcal{M}_n(\mathbb{C})$. Elle est donc trigonalisable. Maintenant soit $\lambda \in \mathbb{C}$ une valeur propre de A et $X \in \mathcal{M}_{n,1}(\mathbb{C})$ un vecteur propre (non nul) associé. On a $AX = \lambda X$ et, pour p l'indice de nilpotence de A , $0 = A^p X = \lambda^p X$. On en déduit que $\lambda = 0$ et donc que $\text{Sp}(A) = \{0\}$. Cela montre que $\chi_u = \chi_A = X^n$.
- $\boxed{ii) \Rightarrow i)}$ Si $\chi_u = X^n$ alors u est trigonalisable et il existe une base $\mathcal{B}$ telle que

$$\text{Mat}_{\mathcal{B}}(u) = \begin{pmatrix} 0 & * & \cdots & * \\ \vdots & \ddots & \ddots & \vdots \\ \vdots & & \ddots & * \\ 0 & \cdots & \cdots & 0 \end{pmatrix}.$$

Soit on note $(e_1, \dots, e_n)$ cette base et que l'on pose $F_i = \text{Vect}(e_1, \dots, e_i)$. On a donc $\{0\} = F_0 \subset F_1 \subset \cdots \subset F_n = E$. La forme de la matrice nous permet d'affirmer que pour tout i , $u(F_i) \subset F_{i-1}$ et donc

$$u^p(F_i) \subset F_{i-p}$$

en posant $F_i = \{0\}$ si $i < 0$. Finalement $u^n(F_n) = \{0\}$ donc $u^n = 0_{\mathcal{L}(E)}$. L'endomorphisme u est bien nilpotent.

□

Corollaire 4.25

Soit u un endomorphisme nilpotent de E , son indice de nilpotent est inférieur ou égal à la dimension de E .

Remarque : Soit $k \in \llbracket 1; n \rrbracket$, il existe $A \in \mathcal{M}_n(\mathbb{K})$ nilpotente donc l'indice de nilpotence vaut k . Pour $k = 1$ on peut prendre $A = 0$ et pour $k \geq 2$

$$A = \sum_{i=1}^{k-1} E_{i, i+n+1-k}$$

ATTENTION

On a montré qu'une matrice nilpotente était **semblable** à une matrice strictement triangulaire supérieure. En pratique, on se ramène souvent à ce dernier cas mais une matrice nilpotente n'est à priori pas **égale** à une matrice strictement triangulaire supérieure

Suites et séries de fonctions

1	Suites de fonctions	125
1.1	Convergence simple et convergence uniforme	126
1.2	Norme infinie	129
1.3	Quelques méthodes	132
2	Continuité et double limite	134
2.1	Continuité	134
2.2	Théorème de la double limite	136
3	Intégration et dérivation	138
3.1	Intégration	138
3.2	Dérivation	141
4	Séries de fonctions	143
4.1	Convergence	143
4.2	Continuité, intégration et dérivation	147
4.3	Exemple d'étude d'une fonction définie par une série	151

Le but est d'étudier les différentes convergences des suites de fonctions puis de considérer le cas des séries.

Commençons par un exemple introductif. On sait que pour tout réel x , $\exp(x) = \sum_{k=0}^{+\infty} \frac{x^k}{k!}$. De ce fait, on peut poser pour tout $n \in \mathbb{N}$,

$$S_n : x \mapsto \sum_{k=0}^n \frac{x^k}{k!}.$$

On veut alors dire que la suite de fonction (S_n) converge simplement vers la fonction $\exp$. On sait que pour tout entier n , les fonctions S_n sont continues, de classe $\mathcal{C}^1$. Peut-on en dire de même de la fonction limite $\exp$?

1 Suites de fonctions

1.1 Convergence simple et convergence uniforme

Dans ce paragraphe A est un ensemble.

Définition 5.1 (Convergence simple)

Soit $(f_n) \in \mathcal{F}(A, \mathbf{K})^{\mathbf{N}}$ une suite de fonctions et $f \in \mathcal{F}(A, \mathbf{K})$. On dit que la suite (f_n) converge simplement vers f si pour tout x de A la suite $(f_n(x))$ tend vers $f(x)$. On a donc

$$\forall x \in A, \forall \varepsilon > 0, \exists N \in \mathbf{N}, \forall n \in \mathbf{N}, n \geq N \Rightarrow |f_n(x) - f(x)| \leq \varepsilon$$

On dit alors que f est la limite simple de (f_n) et on note

$$(f_n) \xrightarrow{CS} f$$

Exemples :

1. Pour $A = \mathbf{R}$, la suite $f_n : x \mapsto x^n$ ne converge pas (simplement) car pour $x = 2$ par exemple $f_n(x) \rightarrow +\infty$.
2. Pour $A = [0, 1]$ la suite $f_n : x \mapsto x^n$ converge vers f définie par

$$f : x \mapsto \begin{cases} 0 & \text{si } x < 1 \\ 1 & \text{si } x = 1 \end{cases}$$

3. Pour $A = [0, 1]$, la suite $f_n : x \mapsto \frac{nx+1}{n}$ converge vers la fonction $x \mapsto x$.

Exercice : Pour les suites suivantes définies sur $[0, 1]$, déterminer si elles convergent simplement et si oui déterminer la limite.

$$f_n : x \mapsto \sin(nx) \quad g_n : x \mapsto \frac{n}{1 + n(x+1)}$$

Remarques :

1. On voit sur l'exemple de la fonction $f_n : x \mapsto x^n$ qu'une limite simple de fonctions continues peut ne pas être continue.
2. Par unicité de la limite des suites $(f_n(x))_{n \in \mathbf{N}}$ si la suite (f_n) converge simplement vers f elle ne peut pas converger simplement vers $g \neq f$.

Définition 5.2 (Convergence uniforme)

Soit $(f_n) \in \mathcal{F}(A, \mathbf{K})^{\mathbf{N}}$ une suite de fonctions et $f \in \mathcal{F}(A, \mathbf{K})$. On dit que la suite (f_n) converge uniformément vers f si

$$\forall \varepsilon > 0, \exists N \in \mathbf{N}, \forall n \in \mathbf{N}, \forall x \in A, n \geq N \Rightarrow |f_n(x) - f(x)| \leq \varepsilon$$

On dit alors que f est la limite uniforme de (f_n) et on note

$$(f_n) \xrightarrow{CU} f$$

ATTENTION

Bien faire la différence entre les deux définitions de convergence. Dans le cas de la limite simple le N est défini après le x , de ce fait il peut en dépendre alors que dans le cas de la limite uniforme le N est le même pour tous les x de A . C'est de là que vient la terminologie uniforme.

Proposition 5.3

Soit (f_n) une suite de fonctions. Si elle converge uniformément vers f alors elle converge simplement vers f . En particulier la limite uniforme d'une suite de fonction est unique (si elle existe).

Démonstration : On suppose que $(f_n) \xrightarrow{CU} f$. Soit $x_0 \in A$ et $\varepsilon > 0$ on sait qu'il existe un entier N tel que $\forall n \in \mathbb{N}, \forall x \in A, n \geq N \Rightarrow |f_n(x) - f(x)| \leq \varepsilon$. C'est en particulier vrai pour x_0 . $\square$

ATTENTION

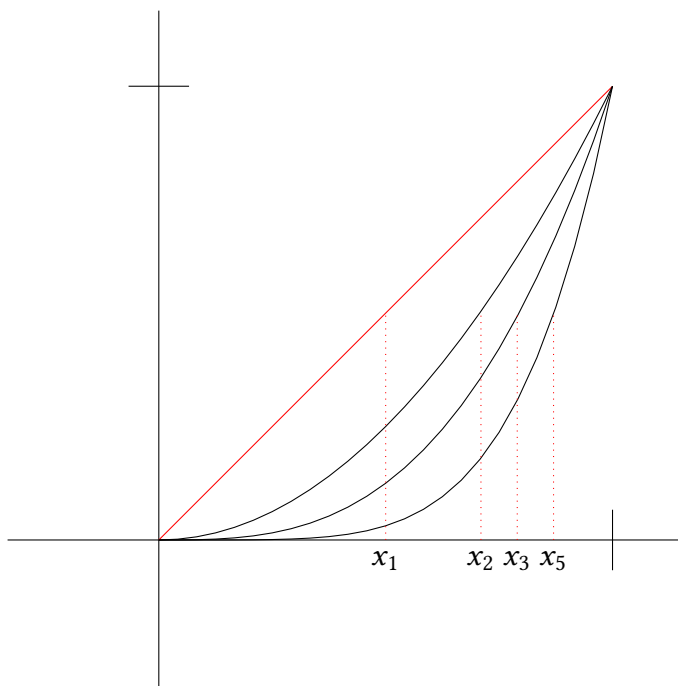
La réciproque est fautive comme on va le voir sur les exemples suivants.

Exemples :

1. Pour $A = \mathbb{R}$, la suite $f_n : x \mapsto x^n$ ne converge pas uniformément car elle ne converge pas simplement.
2. Pour $A = [0, 1]$ la suite $f_n : x \mapsto x^n$ converge simplement vers f définie par

$$f : x \mapsto \begin{cases} 0 & \text{si } x < 1 \\ 1 & \text{si } x = 1 \end{cases}$$

Si elle converge uniformément c'est obligatoirement vers f . Or, on voit bien sur un dessin que pour n aussi grand que l'on veut on va pouvoir trouver x_n (proche de 1 mais strictement inférieur) tel que $|f_n(x_n)| = |f_n(x_n) - f(x_n)| \geq \frac{1}{2}$



On peut prendre $x_n = \frac{1}{2^{1/n}} = \exp\left(-\frac{1}{n} \ln(2)\right)$.

3. Pour $A = [0, 1]$, la suite $f_n : x \mapsto \frac{nx+1}{n}$ converge simplement et uniformément vers la fonction $x \mapsto x$. En effet, on a pour tout $n > 0$ et tout x ,

$$|f_n(x) - f(x)| = \left| \frac{nx+1}{n} - x \right| = \frac{1}{n}.$$

De ce fait on a convergence uniforme.

Exercice : On reprend $f_n : x \mapsto x^n$. Il y a-t-il convergence uniforme vers la fonction nulle sur $[0, \frac{1}{2}]$? Sur $[0, 1[$?

Proposition 5.4

Soit $(f_n) \in \mathcal{F}(A, \mathbb{K})^{\mathbb{N}}$ une suite de fonctions et $f \in \mathcal{F}(A, \mathbb{K})$. Soit $(x_n) \in A^{\mathbb{N}}$ une suite à valeurs dans A . Si $(f_n) \xrightarrow{CU} f$ alors la suite $(f_n(x_n) - f(x_n))_{n \in \mathbb{N}}$ tend vers 0.

Remarques :

1. Cette proposition montre que la convergence uniforme implique une propriété plus forte la seule convergence simple qui impose juste que pour tout x de A , $(f_n(x) - f(x))_{n \in \mathbb{N}}$ tende vers 0.
2. On utilise souvent la contraposée de cette proposition. Si on peut trouver une suite $(x_n) \in A^{\mathbb{N}}$ telle que $(f_n(x_n) - f(x_n))_{n \in \mathbb{N}}$ ne tende pas vers 0 alors (f_n) ne converge pas uniformément vers 0 (voir méthodes plus loin)

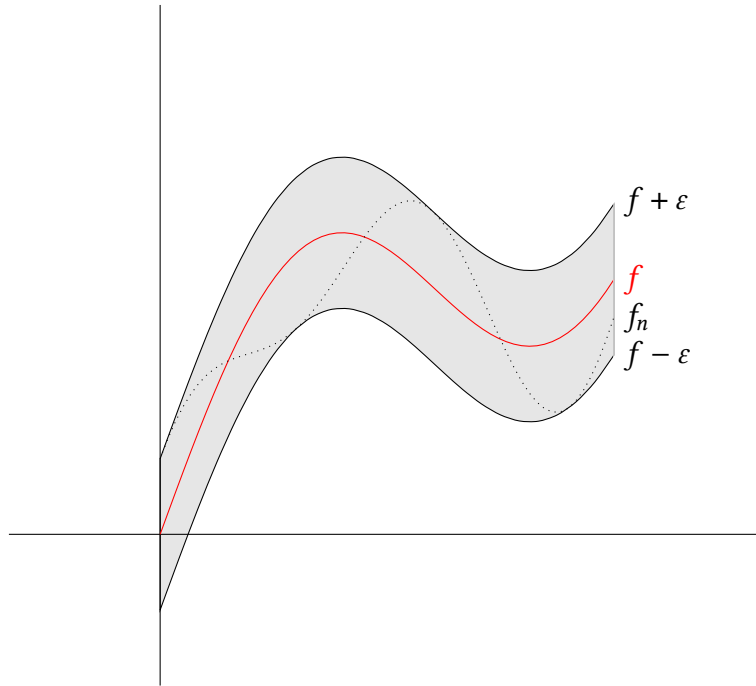
Démonstration : Il suffit de revenir à la définition. Pour tout $\varepsilon > 0$, il existe N tel que pour tout $x \in A$ et tout $n \in \mathbb{N}$, $|f_n(x) - f(x)| \leq \varepsilon$. En particulier, $|f_n(x_n) - f(x_n)| \leq \varepsilon$. $\square$

Exemple : Si on revient à l'exemple 2 ci-dessus. La suite de fonctions (f_n) où $f_n : x \mapsto x^n$ ne converge pas uniformément vers f sur $[0, 1]$ ou $[0, 1[$ car si on pose $x_n = 1 - \frac{1}{n}$, on a

$$f_n(x_n) - f(x_n) = \left(1 - \frac{1}{n}\right)^n \xrightarrow{n \rightarrow +\infty} \frac{1}{e} \neq 0$$

Interpretation graphique de la convergence uniforme pour des fonctions réelles

Le fait qu'une suite de fonction (f_n) converge uniformément vers une fonction f , signifie que pour tout $\varepsilon > 0$ il existe un N tel que pour $n \geq N$, la graphe de la fonction f_n soit inclus dans un « tube » autour du graphe de f de largeur 2ε .



1.2 Norme infinie

Définition 5.5

Soit E un espace vectoriel sur $\mathbf{K} = \mathbf{R}$ ou $\mathbf{C}$. On appelle norme une application $N : E \mapsto \mathbf{R}_+$ telle que

1. $\forall x \in E, N(x) = 0 \iff x = 0$
2. $\forall x \in E, \forall \lambda \in \mathbf{K}, N(\lambda \cdot x) = |\lambda|N(x)$
3. $\forall (x, y) \in E^2, N(x + y) \leq N(x) + N(y)$ (Inégalité triangulaire)

Remarque : Ce sont les propriétés vérifiées par la valeur absolue sur $\mathbf{K}$, le module sur $\mathbf{C}$ et la norme euclidienne sur un espace préhilbertien réel.

Proposition 5.6

L'ensemble $\mathcal{B}(A, \mathbf{K})$ des fonctions bornées sur A est un sous-espace vectoriel de $\mathcal{F}(A, \mathbf{K})$.

Démonstration : On sait que $\mathcal{B}(A, \mathbf{K})$ n'est pas vide car la fonction nulle est bornée. De plus soit $(f, g) \in \mathcal{B}(A, \mathbf{K})^2$ et λ, μ deux scalaires. On pose $h = \lambda f + \mu g$. On suppose que $\forall x \in A, |f(x)| \leq M_f$ et $|g(x)| \leq M_g$. On a donc

$$\forall x \in A, |h(x)| = |\lambda f(x) + \mu g(x)| \leq |\lambda|M_f + |\mu|M_g$$

On a bien que h est bornée donc $\mathcal{B}(A, \mathbf{K})$ est un sous-espace vectoriel de $\mathcal{F}(A, \mathbf{K})$. $\square$

Proposition-Définition 5.7

L'application

$$\begin{aligned} \|\cdot\|_\infty : \mathcal{B}(A, \mathbf{K}) &\rightarrow \mathbf{R}_+ \\ f &\mapsto \sup_{x \in A} |f(x)| \end{aligned}$$

est un norme que l'on appelle norme infinie.

Remarques :

1. Il n'est pas possible de définir une telle norme sur $\mathcal{F}(A, \mathbf{K})$ en entier.
2. Il faut bien utiliser une borne supérieure car elle peut ne pas être atteinte. Par exemple $x \mapsto \arctan(x)$.
3. On va régulièrement s'intéresser à la norme infinie de restrictions d'une fonction. Plus précisément si $f \in \mathcal{F}(A, \mathbf{K})$ et si $B \subset A$ tel que f est bornée sur B on notera

$$\|f\|_{\infty, B} = \|f|_B\|_\infty = \sup_{x \in B} |f(x)|$$

Démonstration :

- Soit $f \in \mathcal{B}(A, \mathbf{K})$. Si $f = 0$ alors $\|f\|_\infty = 0$. Réciproquement si $\|f\|_\infty = 0$ alors pour tout x de A ,

$$0 \leq |f(x)| \leq \|f\|_\infty = 0$$

On a bien $f = 0$.

- Montrons que $\|\lambda f\|_\infty = |\lambda| \|f\|_\infty$.

$$\|\lambda f\|_\infty = \sup_{x \in A} |\lambda f(x)| = \sup_{x \in A} |\lambda| \cdot |f(x)| = |\lambda| \sup_{x \in A} |f(x)| = |\lambda| \|f\|_\infty$$

- Montrons l'inégalité triangulaire. On sait que pour tout x de A , $|f(x)| \leq \|f\|_\infty$ et $|g(x)| \leq \|g\|_\infty$. On a donc

$$|(f+g)(x)| \leq |f(x)| + |g(x)| \leq \|f\|_\infty + \|g\|_\infty$$

De ce fait

$$\|f+g\|_\infty = \sup_{x \in A} |(f+g)(x)| \leq \|f\|_\infty + \|g\|_\infty$$

□

Dans la démonstration ci-dessus on a utilisé le lemme suivant

Lemme 5.8

Soit X une partie non vide de $\mathbf{R}$ et $k \in \mathbf{R}_+$. On note $kX = \{kx, x \in X\}$. On a

$$\sup(kX) = k \sup(X)$$

Dans le cas où X n'est pas majorée et que $k = 0$, on convient que $0 \times +\infty = 0$.

Démonstration :

- Si $k = 0$, on a $\sup(\{0\}) = 0$ et donc l'égalité est vérifiée avec la convention. On suppose ci-dessous que $k \neq 0$.

- Si $\sup(X) = +\infty$ ce qui est équivalent au fait que X n'est pas majoré. Dans ce cas, kX n'est pas majorée puisque pour tout $M \in \mathbf{R}$, il existe $x \in X$ tel que $x \geq \frac{M}{k}$ et donc $kx \geq M$. Cela montre que $\sup(kX) = +\infty$.
- Si $\sup(X) < +\infty$. Pour $x \in X$, $kx \leq k \sup(X)$. Ceci montre que $k \sup(X)$ est un majorant de kX et donc

$$\sup(kX) \leq k \sup(X)$$

On remarque alors que $X = \frac{1}{k}(kX)$ et donc en appliquant ce qui précède,

$$\sup(X) = \sup\left(\frac{1}{k}(kX)\right) \leq \frac{1}{k} \sup(kX)$$

Cela montre l'inégalité dans l'autre sens en multipliant par k .

Finalement $\sup(kX) = k \sup(X)$

□

Proposition 5.9

Soit (f_n) une suite de fonctions et f une fonction .

$$(f_n) \xrightarrow{CU} f \iff \begin{cases} \exists N \in \mathbf{N}, \forall n \geq N, (f_n - f) \text{ est bornée} \\ \|f_n - f\|_\infty \rightarrow 0 \end{cases}$$

Démonstration :

- $\Rightarrow$ On suppose que $(f_n) \xrightarrow{CU} f$ donc pour $\varepsilon = 1$ (par exemple), il existe N tel que pour $n \geq N$, $|f_n - f|$ est majorée (par 1) et donc $(f_n - f)$ est bornée. De plus,

$$\forall \varepsilon > 0, \exists N \in \mathbf{N}, \forall n \in \mathbf{N}, n \geq N \Rightarrow \|f_n - f\|_\infty \leq \varepsilon$$

ce qui est la définition de $\|f_n - f\|_\infty \rightarrow 0$.

- $\Leftarrow$ On suppose l'assertion de droite. On suppose que $(f_n - f)$ est bornée à partir du rang N . Dès lors, comme $\|f_n - f\|_\infty \rightarrow 0$, pour tout $\varepsilon > 0$ il existe un $N' \geq N$ tel que pour $n \geq N'$, $\|f_n - f\|_\infty \leq \varepsilon$. C'est la définition de la convergence uniforme.

□

Corollaire 5.10

Soit λ, μ deux scalaires. Si $(f_n) \xrightarrow{CU} f$ et $(g_n) \xrightarrow{CU} g$ alors $(\lambda f_n + \mu g_n) \xrightarrow{CU} \lambda f + \mu g$.

Démonstration : Comme $f_n - f$ et $g_n - g$ sont bornées à partir d'un certain rang alors $(\lambda f_n + \mu g_n) - (\lambda f + \mu g) = \lambda(f_n - f) + \mu(g_n - g)$ aussi. De plus à partir de ce rang

$$(\lambda f_n + \mu g_n) - (\lambda f + \mu g) = \lambda(f_n - f) + \mu(g_n - g) \Rightarrow \|\lambda(f_n - f) + \mu(g_n - g)\|_\infty \leq |\lambda| \|f_n - f\|_\infty + |\mu| \|g_n - g\|_\infty \rightarrow 0.$$

On a bien $(\lambda f_n + \mu g_n) \xrightarrow{CU} \lambda f + \mu g$.

□

1.3 Quelques méthodes

★ **Méthode** : (Montrer qu'une suite de fonctions converge uniformément vers f) : On commence par déterminer f en étudiant la limite simple de f . Une fois que f est connue, il faut trouver une suite (u_n) tendant vers 0 telle que à partir d'un certain rang, $\|f_n - f\|_\infty \leq u_n$. Pour cela il est souvent commode d'étudier les variations de $f_n - f$ afin de déterminer $\sup_{x \in A} |f_n(x) - f(x)|$.

★ **Méthode** : (Montrer qu'une suite de fonctions ne converge pas uniformément vers f) : Là encore, on commence par déterminer la limite simple. Si la suite de fonction ne converge pas simplement c'est gagné et si elle converge simplement on a la seule valeur possible pour la limite. Il faut ensuite montrer que $\|f_n - f\|_\infty$ ne tend pas vers 0. Il y a deux possibilités.

- On calcule $\sup_{x \in A} |f_n(x) - f(x)|$ à l'aide d'une étude de variations.
- On cherche une suite (x_n) de A telle que $|f_n(x_n) - f(x_n)|$ ne tend pas vers 0. C'est ce que l'on a fait dans l'exemple ci-dessus.

Exemples :

1. On étudie sur $[0, 1]$, $f_n : x \mapsto \frac{ne^{-x} + x^2}{n+x}$.

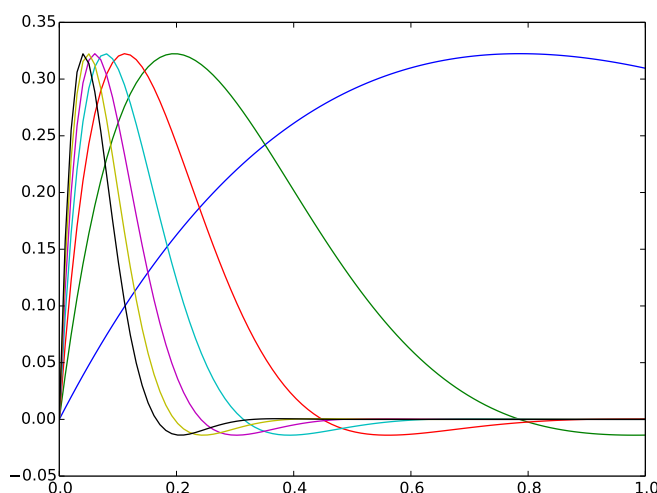
On voit d'abord que pour $x_0 \in [0, 1]$, $f_n(x_0) = \frac{ne^{-x_0} + x_0^2}{n+x_0} \xrightarrow{n \rightarrow +\infty} e^{-x}$. La suite de fonctions converge simplement vers $f : x \mapsto e^{-x}$.

Maintenant pour $x \in [0, 1]$,

$$|f_n(x) - f(x)| = \left| \frac{x^2 - xe^{-x}}{n+x} \right| \leq \frac{1+1}{n} = \frac{2}{n}.$$

Or $\frac{2}{n} \rightarrow 0$. On a donc convergence uniforme.

2. On considère $f_n : x \mapsto \exp(-nx) \sin(nx)$ sur $[0, 1]$. Pour $x = 0$, on a $f_n(0) = 0$ et pour $x > 0$, on a $|f_n(x)| \leq \exp(-nx) \rightarrow 0$. On en déduit que (f_n) converge simplement vers la fonction nulle. Cependant, on voit que $f_n(\frac{1}{n}) = \exp(-1) \sin(1)$ ne tend pas vers 0 donc la convergence n'est pas uniforme.



```
import numpy as np
import pylab as plt
```

```
def f(x,n) :
```

```

return(np.exp(-n*x) * np.sin(n*x))

figure = plt.figure()

for k in range(1,20,3) :
    x = np.linspace(0,1,100)
    y = [f(i,k) for i in x]
    plt.plot(x,y)

plt.show()
figure.savefig("toto.pdf")

```

3. On peut considérer notre exemple introductif. On pose donc $S_n : x \mapsto \sum_{k=0}^n \frac{x^k}{k!}$.

Par définition, $(S_n) \xrightarrow{CS} \exp$. On note de même pour tout entier n ,

$$R_n : x \mapsto \exp(x) - S_n(x) = \sum_{k=n+1}^{\infty} \frac{x^k}{k!}$$

Il est clair que (S_n) ne converge pas uniformément vers $\exp$ sur $\mathbf{R}$ en entier. En effet, pour n fixé, la fonction R_n n'est pas bornée car pour $x \geq 0$, $R_n(x) \geq \frac{x^{n+1}}{(n+1)!}$ et donc

$$\lim_{x \rightarrow +\infty} R_n(x) = +\infty$$

Par contre si on fixe $\alpha \in \mathbf{R}$ et que l'on travaille sur le segment $K = [-\alpha, \alpha]$, on voit que pour $x \in K$,

$$|R_n(x)| = \left| \sum_{k=n+1}^{\infty} \frac{x^k}{k!} \right| \leq \sum_{k=n+1}^{\infty} \frac{|x|^k}{k!} \leq \sum_{k=n+1}^{\infty} \frac{\alpha^k}{k!} = R_n(\alpha)$$

On en déduit que

$$\|S_n - \exp\|_{\infty, K} = \|R_n\|_{\infty, K} \leq R_n(\alpha) \xrightarrow{n \rightarrow +\infty} 0$$

Il y a bien convergence uniforme sur K de la suite (S_n) vers sa limite $\exp$.

ATTENTION

La notion de convergence uniforme dépend de l'ensemble A considéré. On a vu que $f_n : x \mapsto x^n$ ne convergeait pas (même simplement) pour $A = \mathbf{R}_+$. Par contre, cette suite de fonction converge simplement et pas uniformément si on travaille avec les restrictions à $[0, 1]$. Et on peut aussi voir qu'elle convergera uniformément sur tout intervalle de la forme $[0, a]$ pour $a < 1$ (car dans ce cas $|f_n(x)| \leq a^n \rightarrow 0$).

Exercice : Soit (f_n) est une suite de fonctions définie sur $A = X \cup X'$.

1. On suppose que $A = X \cup X'$ et que $((f_n)|_X)$ et $((f_n)|_{X'})$ convergent uniformément. Montrer que (f_n) converge uniformément sur A .
2. Généraliser au cas où $A = X_1 \cup \dots \cup X_p$
3. Montrer que cela n'est plus vrai en général si A est une union infinie.

2 Continuité et double limite

Dans ce paragraphe nous voulons étudier une fonction f qui est limite d'une suite (f_n) . En particulier, on va regarder si la connaissance des limites des f_n quand $x \rightarrow a$ permet de déterminer la limite de f en a .

Dans ce paragraphe les fonctions seront définies sur une partie A de $\mathbf{R}$.

2.1 Continuité

On se demande si quand une suite de fonctions (f_n) tend vers une fonction f et que toutes les fonctions f_n sont continues en $a \in A$ alors f est aussi continue en a . On a déjà vu dans les exemples précédents que si f converge juste simplement cela n'est pas toujours vrai. Il suffit de prendre $f_n : x \mapsto x^n$ et $a = 1$.

Commençons par un rappel

Définition 5.11

1. Soit $a \in \mathbf{R}$ on appelle *voisinage de a* un ensemble de la forme $B(a, \delta) = \{x \in \mathbf{K}, |x-a| \leq \delta\}$ pour $\delta > 0$.
2. Soit $A \subset \mathbf{R}$ et $a \in A$, on appelle *voisinage de a dans A* un ensemble de la forme $A \cap B(a, \delta)$ pour $\delta > 0$.

Remarque : On travaille ici avec des voisinages « fermés » mais on pourrait faire de même avec des voisinages « ouverts » définis de la même manière mais avec une inégalité stricte.

Théorème 5.12

Soit (f_n) une suite de fonctions définies sur $A \subset \mathbf{R}$ vers $\mathbf{K} = \mathbf{R}$ ou $\mathbf{C}$. Soit f une fonction de A dans $\mathbf{K}$ et $a \in A$.

On suppose que

- i) La suite de fonctions (f_n) converge **uniformément** vers f sur A ,
- ii) les fonctions f_n sont continues en a .

Alors la fonction f est continue en a .

Remarques :

1. On peut supposer que les f_n ne sont continues en a qu'à partir d'un certain rang.
2. Il suffit d'avoir la convergence uniforme sur un voisinage V de a dans A .
3. Pour le moment nous ne savons définir la notion de continuité que d'une partie de $\mathbf{R}$ dans une partie de $\mathbf{C}$ mais plus tard dans l'année nous pourrons généraliser ce théorème en prenant pour A une partie d'un espace vectoriel de dimension finie.

Démonstration : On veut montrer :

$$\forall \varepsilon > 0, \exists \eta > 0, \forall x \in A, |x - a| \leq \eta \Rightarrow |f(x) - f(a)| \leq \varepsilon.$$

On se donne donc un $\varepsilon > 0$ et on pose $\varepsilon' = \frac{\varepsilon}{3}$. Maintenant comme $(f_n) \xrightarrow{CU} f$ il existe $n \in \mathbf{N}$ tel que

$\|f_n - f\|_\infty \leq \varepsilon'$. Alors pour tout x de A ,

$$\begin{aligned} |f(x) - f(a)| &= |(f(x) - f_n(x)) + (f_n(x) - f_n(a)) + (f_n(a) - f(a))| \\ &\leq |f(x) - f_n(x)| + |f_n(x) - f_n(a)| + |f_n(a) - f(a)| \\ &\leq 2\varepsilon' + |f_n(x) - f_n(a)| \end{aligned}$$

Maintenant comme f_n est continue, on a $\lim_{x \rightarrow a} f_n(x) = f_n(a)$. De ce fait il existe $\eta > 0$ tel que pour x dans A , $|x - a| \leq \eta \Rightarrow |f_n(x) - f_n(a)| \leq \varepsilon'$. Pour cet η on a donc :

$$|x - a| \leq \eta \Rightarrow |f(x) - f(a)| \leq 3\varepsilon' = \varepsilon.$$

On a bien montré que f était continue en a . □

Corollaire 5.13

Soit (f_n) une suite de fonctions **continues** de A dans $\mathbf{K}$. Si $(f_n) \xrightarrow{CU} f$ alors f est continue.

Remarques :

1. Là encore on peut se contenter du fait que les (f_n) soient continues à partir d'un certain rang.
2. On peut utiliser la contraposée. Si une suite de fonctions continues converge vers une fonction qui n'est pas continue alors la convergence n'est pas uniforme. Comme dans le cas de la suite des $x \mapsto x^n$.

Note

Il arrive qu'une suite de fonctions ne converge pas uniformément sur tout l'ensemble de définition mais on peut quand même justifier la continuité de la limite.

Soit pour $n \geq 1$, $f_n : x \mapsto \sum_{k=0}^n \frac{x^k}{k!}$ définie sur $\mathbf{R}$. La suite (f_n) converge simplement vers la fonction f définie par

$$f : x \mapsto \sum_{k=0}^{+\infty} \frac{x^k}{k!}$$

On a vu que la convergence n'est pas uniforme sur $\mathbf{R}$. Cependant on a aussi vu que la suite (f_n) converge uniformément sur tout intervalle de la forme $[-\alpha, \alpha]$.

De ce fait, soit $x_0 \in \mathbf{R}$ il existe α tel que $x_0 \in]-\alpha, \alpha[$ (par exemple en prenant $\alpha = |x_0| + 1$). En appliquant ce qui précède, la restriction de f à $[-\alpha, \alpha]$ est continue et donc f est continue en x_0 .

On vient de montrer que

$$\exp : x \mapsto \sum_{n=0}^{+\infty} \frac{x^n}{n!}$$

est continue sur $\mathbf{R}$.

Proposition 5.14

Soit (f_n) une suite de fonctions continues. On suppose que pour tout a de A il existe un voisinage V_a de a dans A tel que (f_n) converge uniformément vers f sur V_a alors f est continue.

Démonstration : On procède comme dans l'exemple. Pour tout point $a \in A$, on montre que f est continue en a et donc f est continue. $\square$

2.2 Théorème de la double limite

On peut réécrire l'énoncé du paragraphe précédent sous la forme :

$$\lim_{x \rightarrow a} f(x) = f(a) = \lim_{n \rightarrow +\infty} f_n(a)$$

C'est-à-dire

$$\lim_{x \rightarrow a} \left(\lim_{n \rightarrow +\infty} f_n(x) \right) = \lim_{n \rightarrow +\infty} \left(\lim_{x \rightarrow a} f_n(x) \right).$$

Dans ce paragraphe, nous allons généraliser ce cas où a n'appartient plus nécessairement à A . C'est-à-dire que a peut être une borne d'un intervalle ouvert ou $a = +\infty$ si A n'est pas majoré ($-\infty$ si A n'est pas minoré).

Définition 5.15 (Point adhérent)

Soit $A \subset \mathbf{R}$, un point a de $\mathbf{R}$ est dit adhérent à A si pour tout $\delta > 0$, $B(a, \delta) \cap A \neq \emptyset$.

Remarque : Il est clair que si $a \in A$ alors il est adhérent à A car $a \in B(a, \delta) \cap A$. Cependant il existe des points adhérents à A qui ne sont pas dans A . Ce sont les points qui sont « juste à côté ».

Exemples :

1. Pour $A = [a, b]$, il n'y a pas de points adhérents à A qui ne sont pas dans A .
2. Les points $a = 0$ et $a = 1$ sont adhérents à $A =]0, 1[$.

Théorème 5.16 (Théorème de la double limite)

Soit (f_n) une suite de fonctions de A dans $\mathbf{K}$ et $f \in \mathcal{F}(A, \mathbf{K})$. Soit a un point adhérent à A . On suppose que

- i) La suite (f_n) converge **uniformément** vers f .
- ii) pour tout n , f_n admet en a une limite finie $\ell_n \in \mathbf{K}$.

Alors la suite (ℓ_n) admet une limite finie ℓ et $\lim_{x \rightarrow a} f(x) = \ell$. C'est-à-dire

$$\lim_{n \rightarrow +\infty} \left(\lim_{x \rightarrow a} f_n(x) \right) = \ell = \lim_{x \rightarrow a} \left(\lim_{n \rightarrow +\infty} f_n(x) \right).$$

Remarques :

1. Dans le cas où $a \in A$ on retrouve le théorème du paragraphe précédent.
2. Il suffit d'avoir la convergence uniforme sur un voisinage de a dans A . Par exemple si $A =]0, 1[$ et $a = 0$, il suffit d'avoir la convergence uniforme sur $]0, \varepsilon[$.

Démonstration : (Non exigible) La partie la plus difficile est de montrer que la suite (ℓ_n) admet une limite finie ℓ . Nous le ferons à la fin.

Admettons pour le moment que (ℓ_n) admet une limite finie ℓ . La preuve de la deuxième partie est similaire à celle du théorème du paragraphe précédent. En effet On veut montrer :

$$\forall \varepsilon > 0, \exists \eta > 0, \forall x \in A, |x - a| \leq \eta \Rightarrow |f(x) - \ell| \leq \varepsilon.$$

On se donne donc un $\varepsilon > 0$ et on pose $\varepsilon' = \frac{\varepsilon}{3}$. Maintenant comme $(f_n) \xrightarrow{CU} f$ il existe N tel que implique $n \geq N$, $\|f_n - f\|_\infty \leq \varepsilon'$. De même comme $(\ell_n) \rightarrow \ell$, il existe N' tel que $n \geq N'$ implique $|\ell_n - \ell| \leq \varepsilon'$. On prend alors $n \geq \max(N, N')$. Alors pour tout x de A ,

$$\begin{aligned} |f(x) - \ell| &= |(f(x) - f_n(x)) + (f_n(x) - \ell_n) + (\ell_n - \ell)| \\ &\leq |f(x) - f_n(x)| + |f_n(x) - \ell_n| + |\ell_n - \ell| \\ &\leq 2\varepsilon' + |f_n(x) - \ell_n| \end{aligned}$$

Maintenant on sait que $\lim_{x \rightarrow a} f_n(x) = \ell_n$. De ce fait il existe $\eta > 0$ tel que pour x dans A , $|x - a| \leq \eta \Rightarrow |f_n(x) - \ell_n| \leq \varepsilon'$. Pour cet η on a donc :

$$|x - a| \leq \eta \Rightarrow |f(x) - \ell| \leq 3\varepsilon' = \varepsilon.$$

On a bien montré que $\lim_{x \rightarrow a} f(x) = \ell$.

Il reste donc à montrer que (ℓ_n) tend vers une limite finie ℓ . Commençons par montrer que la suite est bornée. On sait que $\|f_n - f\|_\infty$ tend vers 0. Donc il existe $N \in \mathbb{N}$ tel que pour $n \geq N$, $\|f_n - f\|_\infty \leq 1$ (on aurait pu prendre une autre valeur). Maintenant pour $n \geq N$,

$$\|f_n - f_N\|_\infty \leq \|f_n - f\|_\infty + \|f - f_N\|_\infty \leq 2.$$

On a donc pour tout $x \in A$, $|f_n(x) - f_N(x)| \leq 2$ en passant à la limite quand $x \rightarrow a$ on obtient que $|\ell_n - \ell_N| \leq 2$. Cela signifie que pour $n \geq N$, $\ell_n \in B(\ell_N, 2)$. La suite est bien bornée.

On peut alors appliquer le théorème de Bolzano-Weierstrass. Il existe une suite extraite $(\ell_{\varphi(n)})$ qui converge. On rappelle que φ est une application strictement croissante de $\mathbb{N}$ dans $\mathbb{N}$. Notons ℓ la limite de cette suite extraite. Il reste à montrer que (ℓ_n) converge vers ℓ .

Pour tout $\varepsilon > 0$, on pose encore $\varepsilon' = \frac{\varepsilon}{3}$. En reprenant l'argument ci-dessus on peut montrer qu'il existe N tel que pour n et p supérieurs à N , $\|f_n - f_p\|_\infty \leq 2\varepsilon'$. De ce fait, en passant encore à la limite on a $|\ell_n - \ell_p| \leq 2\varepsilon'$.

Maintenant il existe N' tel que pour $p \geq N'$ et p dans l'image de φ , $|\ell_p - \ell| \leq \varepsilon'$.

En conclusion pour $n \geq \max(N, N')$ on a

$$|\ell_n - \ell| \leq |\ell_n - \ell_p| + |\ell_p - \ell| \leq 3\varepsilon' = \varepsilon$$

où on a choisit $p \geq \max(N, N')$ et dans l'image de φ .

On a donc bien montré que la suite (ℓ_n) convergeait vers un élément ℓ . □

Exemple : On considère pour $n \geq 1$,

$$f_n : x \mapsto \frac{\sum_{k=0}^n x^k/k! - 1}{x} = \sum_{k=1}^n \frac{x^{k-1}}{k!} = 1 + \frac{x}{2} + \frac{x^2}{6} + \cdots + \frac{x^{n-1}}{n!}.$$

Pour $x \in]0, 1]$ on a $f_n(x) \xrightarrow{n \rightarrow \infty} f(x) = \frac{e^x - 1}{x}$.

De plus la convergence est uniforme car pour $x \in]0, 1]$,

$$|f_n(x) - f(x)| = \sum_{k>n} \frac{x^{k-1}}{k!} \leq \sum_{k>n} \frac{1}{k!}$$

et le terme de droite tend vers 0 (reste de la série exponentielle).

De plus, pour tout entier n , $\lim_{x \rightarrow 0} f_n(x) = 1$. On en déduit que $\lim_{x \rightarrow 0} \frac{e^x - 1}{x} = 1$.

Proposition 5.17 (Généralisation)

1. Si $A \subset \mathbb{R}$ et n'est pas majorée, on a le même résultat en remplaçant a par $+\infty$.
2. Si $A \subset \mathbb{R}$ et n'est pas minorée, on a le même résultat en remplaçant a par $-\infty$.

Exemple : On pose $f_n : x \mapsto ne^{-nx}$. Pour tout $x > 0$, $f_n(x) \xrightarrow{n \rightarrow \infty} 0$. Donc (f_n) converge vers la fonction nulle (que l'on note f) sur $]0, +\infty[$. De plus, la convergence est uniforme sur $[1, +\infty[$ car pour tout $x \geq 1$

$$|f_n(x) - 0| = ne^{-nx} \leq ne^{-n} \xrightarrow{n \rightarrow \infty} 0.$$

Maintenant pour tout $n \geq 1$, $\lim_{x \rightarrow +\infty} f_n(x) = 0 = \ell_n$. On retrouve bien que $\lim_{x \rightarrow +\infty} f(x) = 0$.

Par contre, si essaye de faire la même chose en 0^+ , on voit que pour tout entier n , $\lim_{0^+} f_n = n$. On voit que

$$\lim_{x \rightarrow 0^+} f(x) = 0 \neq +\infty = \lim_{n \rightarrow +\infty} n$$

En effet, la suite (f_n) ne converge pas uniformément au voisinage de 0^+ .

3 Intégration et dérivation

3.1 Intégration

On va lier le fait qu'une suite de fonctions (f_n) converge vers f avec des propriétés équivalentes sur les primitives et intégrales.

Théorème 5.18

Soit (f_n) une suite de fonctions définies sur un intervalle I et à valeurs dans $\mathbb{K}$. Soit f une fonction définie sur I et à valeurs dans $\mathbb{K}$. On suppose que

- i) Les fonctions f_n sont continues.
- ii) La suite de fonctions converge uniformément sur tout segment J inclus dans I vers f

Soit $a \in I$, on pose $F_n : x \mapsto \int_a^x f_n(t)dt$ et $F : x \mapsto \int_a^x f(t)dt$. La suite (F_n) converge uniformément vers F sur tout segment J de I .

Remarques :

1. Comme on l'a vu, le fait que la suite (f_n) converge uniformément sur tout segment de I ne signifie pas que (f_n) converge uniformément sur I . De même pour la suite (F_n) .
2. La fonction F est définie car, comme les f_n sont continues et qu'il y a convergence uniforme sur tout segment alors f est continue.
3. On demande que les f_n soient continues et pas juste continues par morceaux. Cette hypothèse est technique (mais elle est dans le programme). Le problème est que si les fonctions (f_n) sont supposées uniquement continues par morceaux, la limite (uniforme) f pourrait ne pas être

continue par morceaux. Par exemple on considère pour tout entier n la fonction f_n définie sur $[0, 1]$ par

$$f_n : x \mapsto \begin{cases} 0 & \text{si } x \notin \mathbf{Q} \\ 1 & \text{si } x = \frac{p}{q} \text{ avec } p \wedge q = 1 \text{ et } q \leq n \\ 0 & \text{si } x = \frac{p}{q} \text{ avec } p \wedge q = 1 \text{ et } q > n \end{cases}$$

On voit aisément que f_n n'a qu'un nombre fini de points de discontinuité mais la limite simple

$$f : x \mapsto \begin{cases} 0 & \text{si } x \notin \mathbf{Q} \\ 1 & \text{sinon} \end{cases}$$

n'est pas continue par morceaux.

4. Les fonctions F_n et F sont les uniques primitives de f_n et f qui s'annulent en a .
5. Même si on suppose que la suite (f_n) converge uniformément sur la totalité de l'intervalle I on n'obtient pas nécessairement que la suite (F_n) converge uniformément sur I vers F . Par exemple si on prend $f_n : x \mapsto \frac{1}{n}$ qui converge uniformément sur $\mathbf{R}$ vers $f = \tilde{0}$. En intégrant (en prenant $a = 0$), on obtient $F_n : x \mapsto \frac{x}{n}$ et $F = \tilde{0}$. On voit que (F_n) ne converge pas uniformément vers F sur $\mathbf{R}$.

Démonstration : Soit J un segment de I . On veut montrer que (F_n) converge uniformément vers F sur J . Or pour $x \in J$,

$$|F_n(x) - F(x)| = \left| \int_a^x f_n(t) dt - \int_a^x f(t) dt \right| = \left| \int_a^x f_n(t) - f(t) dt \right| \leq |x - a| \sup_{t \in [a, x]} |f_n(t) - f(t)|.$$

On considère un segment J' contenant J et a et on note δ son diamètre (c'est-à-dire la distance entre ses deux bornes). On obtient

$$|F_n(x) - F(x)| \leq \delta \|f_n - f\|_{J', \infty}$$

où $\|f_n - f\|_{J', \infty} = \sup_{t \in J'} |f_n(t) - f(t)|$. C'est-à-dire

$$\|F_n - F\|_{J, \infty} \leq \delta \|f_n - f\|_{J', \infty}.$$

Or $(f_n) \xrightarrow{CU} f$ sur J' donc $(F_n) \xrightarrow{CU} f$ sur J . □

Corollaire 5.19

Si (f_n) est une suite de fonctions continues définies sur un segment J et que $(f_n) \xrightarrow{CU} f$ alors

$$\int_J f_n \xrightarrow{n \rightarrow \infty} \int_J f.$$

Démonstration : Si $J = [a, b]$ on applique le théorème précédent. On obtient que $(F_n) \xrightarrow{CU} F$ donc en particulier pour $x = b$ on a

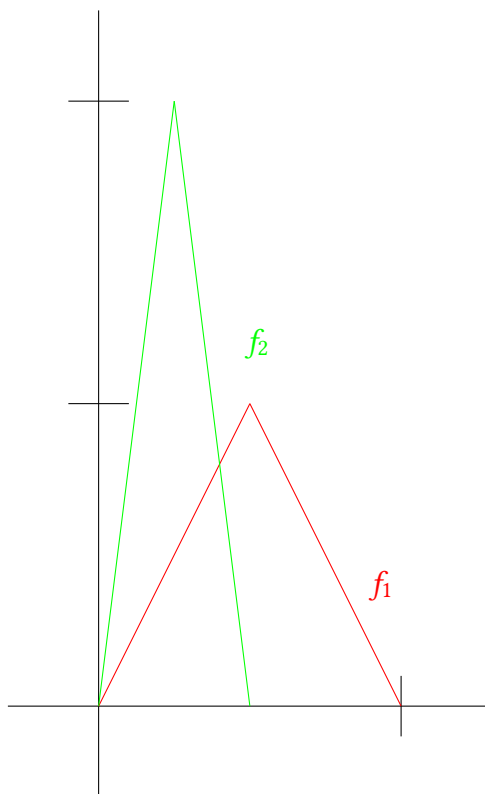
$$F_n(b) = \int_a^b f_n \xrightarrow{n \rightarrow \infty} \int_a^b f = F(b).$$

□

Remarques :

1. Pour ce dernier énoncé, la convergence uniforme est essentielle. Considérons la suite de fonctions (f_n) définie pour $n \geq 1$ sur $[0, 1]$ par :

$$f_n : x \mapsto \begin{cases} 0 & \text{si } x \geq \frac{1}{n} \\ 2n - 2n^2x & \text{si } \frac{1}{2n} \leq x < \frac{1}{n} \\ 2n^2x & \text{si } x < \frac{1}{2n} \end{cases}$$



On peut voir que $(f_n) \xrightarrow{CS} f$ où f est la fonction nulle. En effet pour $x = 0$ on a pour tout $n \in \mathbb{N}$, $f_n(0) = 0$ et pour $x > 0$ on a $f_n(x) = 0$ dès que $x \geq \frac{1}{n}$ c'est-à-dire $n \geq \frac{1}{x}$.

Cependant la convergence n'est pas uniforme. On peut s'en convaincre en voyant que pour tout n , $\int_0^1 f_n(t)dt = \frac{1}{2}$ et $\int_0^1 f(t)dt = 0$.

2. On peut utiliser la contraposée de ce résultat. Si $(f_n) \xrightarrow{CS} f$ sur un segment I et que $\int_I f_n$ ne tend pas vers $\int_I f$ alors la convergence n'est pas uniforme.

ATTENTION

Il faut bien faire la différence entre cet énoncé et le théorème de convergence dominé étudié au chapitre 2. Le corollaire ne peut s'appliquer **que sur un segment** $[a, b]$. Il est plus simple à mettre en place mais **nettement moins puissant**.

Exemple : Posons $I_n = \int_0^1 \frac{1}{1+t^n} dt$. On pose $f_n : t \mapsto \frac{1}{1+t^n}$. Cette suite de fonction converge

simplement vers f définie par

$$f : t \mapsto \begin{cases} 1 & \text{si } t \neq 1 \\ 1/2 & \text{si } t = 1 \end{cases}$$

En constatant que la limite n'est pas continue on sait que la convergence n'est pas uniforme.

Mais on peut appliquer le théorème de convergence dominée (en dominant par la fonction constante égale à 1) pour obtenir que $(I_n) \rightarrow 1$.

3.2 Dérivation

Étudions maintenant si quand (f_n) tend vers f alors la dérivée de la limite est la limite des dérivées. C'est-à-dire, a-t-on

$$\lim_{n \rightarrow \infty} f'_n = f' = \left(\lim_{n \rightarrow \infty} f_n \right)'?$$

Commençons par montrer sur un exemple que la seule convergence (même uniforme) de la suite de fonctions ne suffit pas.

Considérons sur $\mathbf{R}_+$ la suite de fonctions $f_n : x \mapsto \sqrt{x + \frac{1}{n}}$. Ce sont des fonctions dérivables sur $\mathbf{R}_+$.

Maintenant $(f_n) \xrightarrow{CU} f$ où $f : x \mapsto \sqrt{x}$. En effet, pour tout $n \geq 1$ et tout $x \in \mathbf{R}_+$,

$$|f_n(x) - f(x)| = \sqrt{x + \frac{1}{n}} - \sqrt{x} = \frac{\frac{1}{n}}{\sqrt{x + \frac{1}{n}} + \sqrt{x}} \leq \frac{\frac{1}{n}}{\sqrt{\frac{1}{n}}} \leq \frac{1}{\sqrt{n}}$$

Cependant, $f : x \mapsto \sqrt{x}$ n'est pas dérivable en 0.

Théorème 5.20

Soit (f_n) une suite de fonctions définies sur un intervalle I et à valeurs dans $\mathbf{K}$. On suppose que

- i) Pour tout entier n , la fonction f_n est de classe $\mathcal{C}^1$ sur I .
- ii) La suite de fonctions (f_n) converge (simplement) vers une fonction f .
- iii) La suite des fonctions dérivées (f'_n) converge **uniformément** sur tout segment J de I vers une fonction g .

Alors la suite de fonctions (f_n) converge uniformément sur tout segment J de I , f est de classe $\mathcal{C}^1$ sur I et $f' = g$. Dit autrement on a

$$\lim_{n \rightarrow \infty} f'_n = f' = \left(\lim_{n \rightarrow \infty} f_n \right)'.$$

Démonstration : Le principe est d'appliquer le théorème d'intégration de la limite du paragraphe précédent à la suite (f'_n) . On sait que

- i) les fonctions f'_n sont continues sur I ,
- ii) les fonctions f'_n convergent uniformément sur tout segment J de I vers g .

On choisit alors $a \in I$. On note $F_n : x \mapsto \int_a^x f'_n$ et $G : x \mapsto \int_a^x g$. On sait que (F_n) converge uniformément vers G sur tout segment $J \subset I$. Comme f_n est une primitive de f'_n , on sait que pour tout $x \in I$,

$$F_n(x) = \int_a^x f'_n(t) dt = [f_n]_a^x = f_n(x) - f_n(a)$$

On en déduit que

$$f_n(x) = f_n(a) + F_n(x) \xrightarrow{n \rightarrow +\infty} f(a) + G(x)$$

Par unicité de la limite on obtient que

$$\forall x \in I, f(x) = f(a) + G(x)$$

De plus on sait que g est la limite uniforme (sur tout segment) d'une suite de fonctions continues donc elle est continue. D'après le théorème fondamental de l'analyse, f est de classe $\mathcal{C}^1$ et $f' = g$.

Il reste à montrer finalement que (f_n) converge uniformément sur tout segment $J \subset I$ vers f . Pour tout $x \in J$,

$$|f_n(x) - f(x)| = |f_n(a) + F_n(x) - f(a) - G(x)| \leq |f_n(a) - f(a)| + |F_n(x) - G(x)| \leq |f_n(a) - f(a)| + \|F_n - G\|_{\infty, J}$$

On en déduit que

$$\|f_n - f\|_{\infty, J} \leq |f_n(a) - f(a)| + \|F_n - G\|_{\infty, J} \xrightarrow{n \rightarrow +\infty} 0$$

Cela montre la convergence uniforme de la suite (f_n) vers f sur le segment J .

□

Exemple : Reprenons notre exemple de la suite (S_n) .

- i) Pour tout entier n , la fonction S_n est de classe $\mathcal{C}^1$ sur $\mathbf{R}$.
- ii) La suite de fonctions (S_n) converge (simplement) vers la fonction $\exp$.
- iii) On remarque que $S'_0 = 0$ et que pour tout entier $n \geq 1$, $S'_n = S_{n-1}$. On en déduit que la suite des fonctions dérivées (S'_n) converge **uniformément** sur tout segment $[-\alpha, \alpha]$ de $\mathbf{R}$ vers la fonction $\exp$.

On en déduit que la limite $\exp$ est de classe $\mathcal{C}^1$ et que $\exp' = \exp$ (et elle est donc $\mathcal{C}^\infty$ par une récurrence immédiate).

On peut généraliser cela au cas des fonctions de classe $\mathcal{C}^k$.

Théorème 5.21

Soit (f_n) une suite de fonctions définies sur un intervalle I et à valeurs dans $\mathbf{K}$. Soit k un entier. On suppose que

- i) Pour tout entier n , la fonction f_n est de classe $\mathcal{C}^k$ sur I .
- ii) Pour tout $i \in \llbracket 0; k-1 \rrbracket$, la suite de fonctions $(f_n^{(i)})$ converge (simplement) vers une fonction g_i .
- iii) La suite des fonctions dérivées k -ième $(f_n^{(k)})$ converge **uniformément** sur tout segment J de I vers une fonction g_k .

On pose $f = g_0$. Alors f est de classe $\mathcal{C}^k$ et pour tout $i \in \llbracket 0; k \rrbracket$, $f^{(i)} = g_i$ et les suites $(f_n^{(i)})$ convergent uniformément sur tout segment.

Démonstration : On procède par récurrence sur k .

- Pour $k = 0$ c'est juste le théorème sur la continuité de la limite. Pour $k = 1$ c'est le théorème précédent.
- Soit $k \geq 1$, on suppose la propriété vraie au rang k et on veut la montrer au rang $k+1$. On se donne donc une suite de fonctions (f_n) définies sur un intervalle I et à valeurs dans $\mathbf{K}$. On suppose que

- i) Pour tout entier n , la fonction f_n est de classe $\mathcal{C}^{k+1}$ sur I .
- ii) Pour tout $i \in \llbracket 0 ; k \rrbracket$, la suite de fonctions $(f_n^{(i)})$ converge (simplement) vers une fonction g_i
- iii) La suite des fonctions dérivées $k + 1$ -ième $(f_n^{(k+1)})$ converge **uniformément** sur tout segment J de I vers une fonction g_{k+1} .

On pose $f = g_0$. Maintenant on veut appliquer l'hypothèse de récurrence à la suite (h_n) où $h_n = f'_n$. Par définition, cette suite de fonctions vérifie les hypothèses de récurrence au rang k . On en déduit que g_1 est de classe $\mathcal{C}^k$, pour tout $i \in \llbracket 0 ; k \rrbracket$, $g_1^{(i)} = g_{i+1}$ et les suites $(h_n^{(i)}) = (f_n^{(i+1)})$ convergent uniformément sur tout segment de I . Il ne reste plus qu'à appliquer le théorème précédent à la suite (f_n) pour montrer que plus que (f_n) converge uniformément sur tout segment de I et que $f' = g_1$ et donc pour tout $i \in \llbracket 1 ; k + 1 \rrbracket$, $f^{(i)} = g_1^{(i-1)} = g_i$.

Par récurrence la propriété est vraie pour tout entier k . □

Pour finir, on peut aussi obtenir le caractère $\mathcal{C}^\infty$ de la limite.

Théorème 5.22

Soit (f_n) une suite de fonctions définies sur un intervalle I et à valeurs dans $\mathbf{K}$. On suppose que

- i) Pour tout entier n , la fonction f_n est de classe $\mathcal{C}^\infty$ sur I .
- ii) Pour tout $i \in \mathbf{N}$, la suite de fonctions $(f_n^{(i)})$ converge (simplement) vers une fonction g_i
- iii) Pour tout $i \in \mathbf{N}$, la suite des fonctions dérivées i -ième $(f_n^{(i)})$ converge **uniformément** sur tout segment J de I vers la fonction g_i .

On pose $f = g_0$. Alors f est de classe $\mathcal{C}^\infty$ et pour tout $i \in \mathbf{N}$, $f^{(i)} = g_i$.

Remarque : On peut avoir une hypothèse un peu plus faible. Il suffit que les $(f_n^{(i)})$ convergent simplement et qu'il existe un rang N tel que pour $i > N$ les suites de $(f_n^{(i)})$ convergent uniformément sur tout segment alors la limite f est de classe $\mathcal{C}^\infty$ et pour tout i , $f^{(i)} = \lim_{n \rightarrow \infty} f_n^{(i)}$.

4 Séries de fonctions

On va vouloir étudier maintenant des séries de fonctions. Cela revient essentiellement à étudier la suite des sommes partielles.

4.1 Convergence

Dans ce paragraphe A est un ensemble.

Définition 5.23 (Convergence simple et convergence uniforme)

Soit (f_n) une suite de fonctions définies sur A et à valeurs dans $\mathbb{K}$.

On note pour tout entier n , $S_n = \sum_{k=0}^n f_k \in \mathcal{F}(A, \mathbb{K})$.

1. On dit que la série de fonctions $\sum f_n$ converge simplement si (S_n) converge simplement.
2. On dit que la série de fonctions $\sum f_n$ converge uniformément si (S_n) converge uniformément.
3. On dit que la série de fonction $\sum f_n$ converge absolument si la suite des fonctions $\left(\sum_{k=0}^n |f_k|\right)$ converge simplement.

Remarques :

1. Comme la convergence simple / uniforme d'une série de fonctions est définies par la convergence simple / uniforme d'une suite de fonctions, tout ce qui a été vu précédemment s'applique. Par exemple, on essayera souvent de commencer par déterminer la limite simple avant de s'attaquer au caractère uniforme de la convergence.
2. Pour alléger les notations on notera souvent $\sum x^n + 1$ la série de fonctions $\sum (x \mapsto x^n + 1)$.

Proposition 5.24

Avec les notations précédentes,

1. Si $\sum f_n$ converge uniformément, elle converge simplement.
2. Si une série de fonctions converge sa limite est $S : x \mapsto \sum_{n=0}^{+\infty} f_n(x)$.

Proposition 5.25

Soit $\sum f_n$ une série de fonctions.

La série converge uniformément si et seulement si elle converge simplement et la suite des restes converge uniformément vers 0.

Démonstration : Si on suppose que la série converge simplement (afin de pouvoir définir les restes).
On a

$$\begin{aligned} \sum f_n \text{ converge uniformément vers } S &\iff \|S_n - S\|_\infty \rightarrow 0 \\ &\iff \|R_n\|_\infty \rightarrow 0 \\ &\iff (R_n) \xrightarrow{CU} 0 \end{aligned}$$

□

Exemple : Considérons la série de fonctions $\sum (x \mapsto x^{2n})$. Elle converge simplement sur $[0, 1[$ vers $S : x \mapsto \frac{1}{1-x^2}$.

On cherche à étudier l'éventuelle convergence uniforme. On regarde le reste

$$R_p : x \mapsto \sum_{n=p+1}^{+\infty} x^{2n} = \frac{x^{2p+2}}{1-x^2}.$$

- Soit $a < 1$, on a convergence uniforme de (R_p) vers la fonction nulle sur $[0, a]$ car pour $x \leq a$,

$$|R_p(x)| \leq \frac{a^{2p+2}}{1-a^2} \xrightarrow{p \rightarrow \infty} 0.$$

Donc la série de fonctions converge uniformément sur $[0, a]$

- Sur $[0, 1[$ il n'y a pas convergence uniforme. En effet pour tout p , $\lim_{x \rightarrow 1} R_p(x) = +\infty$. De ce fait, il existe $x_p \in [0, 1[$ tel que $R_p(x_p) = 1$ ce qui empêche la convergence uniforme de (R_p) vers la fonction nulle.

Remarque : Dans le cas général, il est difficile d'étudier la convergence uniforme d'une série de fonctions. En effet, il n'est pas aisé de calculer $\|S_n - S\|_\infty$ car $S_n - S$ est une somme d'une infinité de termes et on ne peut pas simplement calculer sa dérivée par exemple.

Définition 5.26 (Convergence normale)

Une série de fonctions $\sum f_n$ définie sur A converge normalement si

- i) toutes les fonctions f_n sont bornées
- ii) la série (numérique) $\sum \|f_n\|_\infty$ est convergente.

Remarques :

1. La notion de convergence normale n'a de sens que pour une série de fonctions (pas pour une suite de fonctions).
2. L'intérêt de la convergence normale est qu'il suffit de montrer la convergence d'une série numérique. Il faut cependant calculer auparavant les normes infinies ou au moins les borner

Théorème 5.27

Soit $\sum f_n$ une série de fonctions définie sur A . On suppose que la série converge normalement

1. La série converge absolument c'est-à-dire que la série $\sum |f_n|$ converge simplement.
2. La série converge uniformément.

Remarque : Dans la plupart des cas, pour montrer la convergence uniforme d'une série de fonctions (pour appliquer les théorèmes relatifs à la continuité, dérivabilité, intégration) on commence par tester la convergence normale.

Démonstration : Soit $\sum f_n$ une série de fonctions définie sur A qui converge normalement.

1. Soit $x \in A$, la série $\sum |f_n(x)|$ est une série à termes positifs. De plus, pour tout entier naturel n , $|f_n(x)| \leq \|f_n\|_\infty$.

Comme la série $\sum \|f_n\|_\infty$ converge par hypothèse donc la série $\sum |f_n(x)|$ converge d'après le théorème de comparaison pour les séries à termes positifs.

On en déduit que la série converge absolument.

2. On sait que la série converge absolument donc elle converge simplement. Pour montrer qu'elle converge uniformément on va s'intéresser aux restes.

Pour $p < q$ et $x \in A$,

$$|f_{p+1}(x) + \cdots + f_q(x)| \leq |f_{p+1}(x)| + \cdots + |f_q(x)| \leq \sum_{k=p+1}^q \|f_k\|_\infty$$

En faisant tendre q vers $+\infty$ (les deux termes convergent) on obtient $|R_p(x)| \leq \sum_{k=p+1}^{+\infty} \|f_k\|_\infty$ et donc $\|R_p\|_\infty \leq \sum_{k=p+1}^{+\infty} \|f_k\|_\infty$. Comme la série $\sum \|f_n\|_\infty$ converge le reste tend vers 0 et donc (R_p) converge uniformément vers 0 d'où la série de fonctions converge uniformément.

□

★ **Méthode** : Pour étudier la nature de la convergence d'une série de fonctions :

- On étudie la convergence simple
- On étudie si possible la convergence normale (afin d'obtenir la convergence uniforme)
- Si la série ne converge pas normalement on essaye de démontrer « à la main » la convergence uniforme.

Exemple : Etudions la série $\sum \frac{e^{-nx^2}}{1+n^2}$. On pose $f_n : x \mapsto \frac{e^{-nx^2}}{1+n^2}$. On a

$$\|f_n\|_\infty = f_n(0) = \frac{1}{1+n^2}$$

On en déduit que la série $\sum \|f_n\|_\infty$ converge car $\frac{1}{1+n^2} \sim \frac{1}{n^2}$. Donc la série de fonctions converge normalement donc uniformément.

★ **Méthode** : (Convergence uniforme d'une série alternée de fonctions)

On considère la série de fonctions $\sum_{n \geq 1} \frac{(-1)^n x^2}{x^4 + n}$.

On pose $f_n : x \mapsto \frac{(-1)^n x^2}{x^4 + n}$ On commence par étudier la convergence simple. Soit $x \in \mathbf{R}$.

- Si $x = 0$ alors la série converge et la somme vaut 0.
- Si $x \neq 0$, on a affaire à une série alternée. En effet

$$|f_n(x)| = \frac{x^2}{x^4 + n} \xrightarrow{n \rightarrow \infty} 0$$

et $\frac{x^2}{x^4 + n}$ décroît quand n augmente (on étudie la variation par rapport à n pas par rapport à x).

On en déduit que la série converge simplement sur tout $\mathbf{R}$.

On peut essayer d'étudier la convergence normale. Pour cela on cherche le maximum de $|f_n|$. En dérivant on a $(|f_n|)'(x) = \frac{2x(n - x^4)}{(x^4 + n)^2}$. Le maximum de la fonction est atteint en $n^{1/4}$. On en déduit que

$$\|f_n\|_\infty = |f_n(n^{1/4})| = \frac{1}{2\sqrt{n}}$$

En particulier, la série $\sum \|f_n\|_\infty$ diverge.

Mais, comme on a affaire à une série alternée, on a une majoration du reste. Pour tout $p \geq 1$

$$|R_p(x)| = \left| \sum_{n \geq p+1} \frac{(-1)^n x^2}{x^4 + n} \right| \leq \frac{x^2}{x^4 + p + 1}.$$

On veut savoir si R_p converge uniformément vers 0. Pour cela on cherche le maximum de la fonction $\theta : x \mapsto \frac{x^2}{x^4 + p + 1}$ sur $\mathbf{R}$. Il suffit de trouver le maximum de $h : x \mapsto \frac{x}{x^2 + p + 1}$ sur $\mathbf{R}_+$. On dérive et on obtient

$$h'(x) = \frac{(p+1) - x^2}{(x^2 + p + 1)^2}$$

On en déduit que le maximum de h est atteint en $x_p = \sqrt{p+1}$ et donc θ atteint son maximum en $\sqrt[4]{p+1}$. De ce fait pour tout x de $\mathbf{R}$,

$$|R_p(x)| \leq \frac{\sqrt{p+1}}{(p+1) + (p+1)} = \frac{1}{2\sqrt{p+1}}$$

On en déduit que R_p converge uniformément vers 0 et donc la série converge uniformément sur $\mathbf{R}$.

ATTENTION

L'exemple ci-dessus, montre qu'il n'y a pas équivalence entre convergence normale et convergence uniforme. Une série peut converger uniformément sans converger normalement.

4.2 Continuité, intégration et dérivation

Nous allons réécrire les résultats démontrés précédemment sur les suites de fonctions dans le cadre des séries de fonctions. Dans ce paragraphe A sera une partie de $\mathbf{R}$.

Théorème 5.28 (Continuité)

Soit $\sum f_n$ une série de fonctions définie sur A . Soit $a \in A$.

On suppose que

- i) pour tout entier n , f_n est continue en a ,
- ii) la série de fonctions $\sum f_n$ converge uniformément

On note $S : x \mapsto \sum_{n=0}^{+\infty} f_n(x)$ la limite de la série. La fonction S est continue en a .

Démonstration : C'est juste le théorème sur les suites de fonctions en remarquant que si toutes les fonctions f_n sont continues en a alors, pour tout entier n , $S_n : x \mapsto \sum_{k=0}^n f_k(x)$ est aussi continue.

□

Corollaire 5.29

Soit $\sum f_n$ une série de fonctions définie sur A .

On suppose que

- i) pour tout entier n , f_n est continue,
- ii) pour tout a de A il existe un voisinage V_a tel que la série de fonctions $\sum f_n$ converge uniformément sur V_a .

On note $S : x \mapsto \sum_{n=0}^{+\infty} f_n(x)$ la limite de la série. La fonction S est continue.

Exemple : On étudie $\sum \frac{x^k}{k!}$ et on pose $f_k : x \mapsto \frac{x^k}{k!}$. On sait que la série converge simplement vers $x \mapsto \exp(x)$ et que toutes les f_k sont continues.

Maintenant soit $a \in \mathbf{R}_+$, on a $\forall x \in [-a, a]$,

$$|f_k(x)| \leq \frac{a^k}{k!}$$

Comme $\sum \frac{a^k}{k!}$ converge, il y a convergence normale (donc uniforme) de la série de fonctions sur $[-a, a]$ (mais pas sur $\mathbf{R}$). On en déduit que la limite $x \mapsto \exp(x)$ est continue.

Adaptons maintenant le théorème de la double limite qui va revenir ici à intervertir $\lim$ avec $\sum$; la somme remplaçant la limite sur n .

Théorème 5.30 (Double limite pour les séries de fonctions)

Soit $\sum f_n$ une série de fonctions définie sur A et a un point adhérent à A . On suppose que

- i) La série de fonctions converge uniformément vers S au voisinage de a
- ii) pour tout n , f_n admet en a une limite finie $\ell_n \in \mathbf{K}$.

Alors la série $\sum \ell_n$ converge (vers ℓ) et $\lim_{x \rightarrow a} S(x) = \ell$. C'est-à-dire :

$$\lim_{x \rightarrow a} \sum_{n=0}^{+\infty} f_n(x) = \sum_{n=0}^{+\infty} \left(\lim_{x \rightarrow a} f_n(x) \right).$$

Démonstration : C'est juste la réécriture du théorème de la double limite dans le cadre des séries.
□

Remarque : Si A n'est pas majoré (resp. minoré) on peut prendre $a = +\infty$ (resp. $a = -\infty$).

Exemple : On étudie $\sum_{n \geq 1} \frac{\arctan(nx)}{n^2}$. On pose $f_n : x \mapsto \frac{\arctan(nx)}{n^2}$. La série est absolument convergente sur $\mathbf{R}$ car pour tout $x \in \mathbf{R}$, $\left| \frac{\arctan(nx)}{n^2} \right| \leq \frac{\pi}{2n^2}$. On note

$$S(x) = \sum_{n=1}^{+\infty} \frac{\arctan(nx)}{n^2}$$

Il est clair que $\|f_n\|_\infty \leq \frac{\pi}{2n^2}$. On en déduit que la série de fonctions converge normalement donc uniformément sur $\mathbf{R}$

On sait que pour tout $n \geq 1$, $\lim_{x \rightarrow +\infty} f_n(x) = \frac{\pi}{2n^2}$. De ce fait,

$$\lim_{x \rightarrow +\infty} S(x) = \sum_{n=1}^{+\infty} \frac{\pi}{2n^2} = \frac{\pi}{2} \frac{\pi^2}{6} = \frac{\pi^3}{12}.$$

Théorème 5.31 (Intégration)

Soit $\sum f_n$ une série de fonctions définies sur un intervalle I et à valeurs dans $\mathbf{K}$. On suppose que

- i) Pour tout entier n , f_n est continue
- ii) La série de fonctions converge uniformément sur tout segment J vers S

Pour tout a de I on pose $F_n : x \mapsto \int_a^x f_n$ la primitive de f_n qui s'annule en a . Alors la série $\sum F_n$ converge uniformément sur tout segment J de I vers $x \mapsto \int_a^x S$. En particulier,

$$\forall x \in I, \sum_{n=0}^{+\infty} \int_a^x f_n = \int_a^x \sum_{n=0}^{+\infty} f_n$$

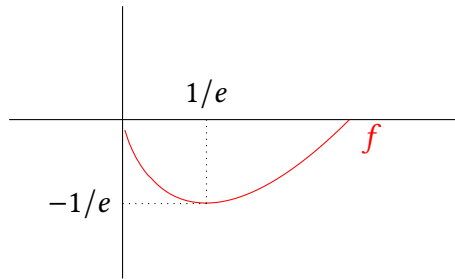
Corollaire 5.32

Si $\sum f_n$ est une série de fonctions continues définies sur un segment J qui converge uniformément vers S ,

$$\sum_{n=0}^{+\infty} \int_J f_n = \int_J \sum_{n=0}^{+\infty} f_n$$

Exemple : On veut trouver une expression de $\int_0^1 x^x dx$.

Pour tout $x \in]0, 1]$, $x^x = \exp(x \ln x) = \sum_{n=0}^{+\infty} \frac{(x \ln x)^n}{n!}$.



On commence en remarquant que la fonction $x \mapsto x \ln x$ se prolonge par continuité en 0 et donc $x \mapsto \frac{(x \ln x)^k}{k!}$ aussi. On peut donc travailler sur le segment $[0, 1]$.

On a donc pour tout n , $\|f_n\|_\infty \leq \frac{1}{e^k k!}$. Or la série $\sum \frac{1}{e^k k!}$ converge donc la série converge normalement (donc uniformément) sur $I = [0, 1]$. On en déduit que

$$\int_0^1 x^x dx = \sum_{n=0}^{+\infty} \int_0^1 \frac{(x \ln x)^n}{n!} dx$$

Il ne reste plus qu'à calculer ces intégrales en intégrant par parties. On obtient que

$$I_n = \int_0^1 \frac{(x \ln x)^n}{n!} dx = \frac{(-1)^n}{(n+1)^{n+1}}$$

En conclusion

$$\int_0^1 x^x dx = \sum_{n=0}^{+\infty} \frac{(-1)^n}{(n+1)^{n+1}}.$$

Théorème 5.33 (Dérivation termes à termes)

Soit $\sum f_n$ une série de fonctions définies sur un intervalle I et à valeurs dans $\mathbf{K}$. On suppose que

- i) Pour tout entier n , f_n est de classe $\mathcal{C}^1$ sur I ,
- ii) la série de fonctions converge (simplement) vers S sur I ,
- iii) la série de fonctions des dérivées $\sum f'_n$ converge uniformément sur tout segment J vers T

Alors la série $\sum f_n$ converge uniformément sur tout segment J , la fonction S est de classe $\mathcal{C}^1$ et $S' = T$. C'est-à-dire

$$\left(\sum_{n=0}^{+\infty} f_n \right)' = \sum_{n=0}^{+\infty} f'_n.$$

Démonstration : On applique encore une fois le théorème vu dans le cas des suites. □

Corollaire 5.34

Soit $\sum f_n$ une série de fonctions définies sur un intervalle I et à valeurs dans $\mathbf{K}$ et $k \in \mathbf{N}$. On suppose que

- i) Pour tout entier n , f_n est de classe $\mathcal{C}^k$ sur I ,
- ii) pour $i < k$ les séries de fonctions $\sum f_n^{(i)}$ convergent (simplement) vers S_i sur I ,
- iii) la série de fonctions des dérivées k -ième $\sum f_n^{(k)}$ converge uniformément sur tout segment J vers S_k

Alors, pour tout $i \in \llbracket 0 ; k \rrbracket$ les séries $\sum f_n^{(i)}$ convergent uniformément sur tout segment J , la fonction $S = S_0$ est de classe $\mathcal{C}^k$ et pour tout $i \leq k$, $S^{(i)} = S_i$. C'est-à-dire

$$\left(\sum_{n=0}^{+\infty} f_n \right)^{(i)} = \sum_{n=0}^{+\infty} f_n^{(i)}.$$

Là encore, on peut généraliser pour les fonctions de classe $\mathcal{C}^\infty$.

Corollaire 5.35

Soit $\sum f_n$ une série de fonctions définies sur un intervalle I et à valeurs dans $\mathbf{K}$. On suppose que

- i) Pour tout entier n , f_n est de classe $\mathcal{C}^\infty$ sur I ,
- ii) pour tout entier i les séries de fonctions $\sum f_n^{(i)}$ convergent (simplement) vers S_i sur I ,
- iii) il existe un entier k tel que pour $i > k$ les séries de fonctions des dérivées i -ième $\sum f_n^{(i)}$ convergent uniformément sur tout segment J

Alors, pour tout $i \in \mathbf{N}$ les séries $\sum f_n^{(i)}$ convergent uniformément sur tout segment J , la fonction $S = S_0$ est de classe $\mathcal{C}^\infty$ et pour tout $i \in \mathbf{N}$, $S^{(i)} = S_i$. C'est-à-dire

$$\left(\sum_{n=0}^{+\infty} f_n \right)^{(i)} = \sum_{n=0}^{+\infty} f_n^{(i)}.$$

4.3 Exemple d'étude d'une fonction définie par une série

Pour finir traitons deux exemples classiques qui permettent de mettre en évidence des méthodes :

★ **Méthode** : Utilisation d'une comparaison série-intégrale (exo [898] ou [895])

On pose

$$f(x) = \sum_{n=1}^{+\infty} \frac{\arctan(nx)}{n^2}$$

1. Déterminer l'ensemble de définition I
2. Montrer que f est continue sur I
3. Montrer que f est de classe $\mathcal{C}^1$ sur deux intervalles à préciser.
4. Déterminer un équivalent de f en 0 et tracer le graphe.

1. La fonction est définie sur $I = \mathbf{R}$.
2. La fonction est continue car elle la série de fonctions normalement, en effet, en posant $u_n : x \mapsto \frac{\arctan(nx)}{n^2}$, on a $\|u_n\|_\infty \leq \frac{\pi}{2n^2}$.
3. Les fonctions u_n sont dérivables et $u'_n : x \mapsto \frac{1}{n(1+n^2x^2)}$. De ce fait $\|u'_n\|_{\infty, [a, +\infty[} \leq \frac{1}{n(1+n^2a^2)}$. On a donc convergence normale (donc uniforme) de la série $\sum_{n \geq 1} u'_n$. On en déduit que f est de classe $\mathcal{C}^1$ sur $[a, +\infty[$ (et de même sur $] -\infty, a]$). Par contre il n'y a pas convergence normale sur $\mathbf{R}$.
4. Pour l'équivalent, on met en place une comparaison série-intégrale. Précisément, pour $x > 0$, on considère

$$g_x : t \mapsto \frac{\arctan(xt)}{t^2}$$

Elle se dérive en

$$g'_x : t \mapsto \frac{g(xt)}{t^3} \text{ où } h : u \mapsto \frac{u}{1+u^2} - 2 \arctan(u)$$

Maintenant, la dérivée de h est négative et $h(0) = 0$ donc on obtient que g_x décroît.

Par comparaison série intégrale on a donc

$$\int_1^{+\infty} g_x(t) dt \leq f(x) = \arctan(x) + \sum_{n=2}^{+\infty} \frac{\arctan(nx)}{n^2} \leq \arctan(x) + \int_1^{+\infty} g_x(t) dt$$

Maintenant, en posant $u = xt$ on obtient que

$$\int_1^{+\infty} g_x(t) dt = \frac{1}{x} \int_x^{+\infty} \frac{\arctan(u)}{\left(\frac{u}{x}\right)^2} du = x \int_x^{+\infty} \frac{\arctan(u)}{u^2} du.$$

Comme $\frac{\arctan(u)}{u^2} \underset{u \rightarrow 0}{\sim} \frac{1}{u}$, par sommation des relations de comparaisons des fonctions positives (divergentes),

$$\int_x^{+\infty} \frac{\arctan(u)}{u^2} du \underset{x \rightarrow 0}{\sim} \int_x^1 \frac{\arctan(u)}{u^2} du \underset{u \rightarrow 0}{\sim} \int_x^0 \frac{1}{u} du = -\ln(x).$$

Finalement, en utilisant que $\arctan(x) \underset{x \rightarrow 0}{\sim} x = o(-x \ln x)$ on obtient que $f(x) \underset{x \rightarrow 0}{\sim} -x \ln(x)$.

★ **Méthode** : Utilisation d'une équation fonctionnelle (voir aussi exo [1830]-CCP)

On veut étudier sur $\mathbf{R}_+^*$ la fonction

$$S : x \mapsto \sum_{n=0}^{+\infty} \frac{(-1)^n}{x+n}$$

1. Définition
2. Dérivabilité
3. Variations
4. Limite en $+\infty$
5. Limite en 0

6. Equivalent en $+\infty$. On pourra commencer par remarquer que $S(x) + S(x+1) = \frac{1}{x}$.

On pose $f_n : x \mapsto \frac{(-1)^n}{x+n}$.

1. Pour tout $x \in \mathbf{R}_+^*$, la série $\sum \frac{(-1)^n}{x+n}$ converge car elle relève du théorème des séries alternées. La fonction est bien définie sur $\mathbf{R}_+^*$.
2. Etudions la dérivabilité. Les fonctions f_n sont de classe $\mathcal{C}^1$ et pour tout entier n et tout $x \in \mathbf{R}_+^*$,

$$f'_n(x) = \frac{(-1)^{n+1}}{(n+x)^2}.$$

On voit que pour $n \geq 1$, $\|f'_n\| = |f'_n(0)| = \frac{1}{n^2}$ et que pour $a > 0$, la fonction f_0 est bornée par $\frac{1}{a}$ sur $[a, +\infty[$. De ce fait la série de fonction $\sum f'_n$ converge normalement (donc uniformément) sur tout segment de $\mathbf{R}_+^*$. On en déduit que la série de fonction converge uniformément sur tout segment de $\mathbf{R}_+^*$ (ce que l'on peut obtenir avec les théorème des séries alternées), que S est dérivable et que pour $x > 0$,

$$S'(x) = \sum_{n=0}^{+\infty} \frac{(-1)^{n+1}}{(n+x)^2}$$

3. D'après le théorème des séries alternées, $S'(x) = \sum_{n=0}^{+\infty} \frac{(-1)^{n+1}}{(n+x)^2}$ est du signe de son premier terme c'est-à-dire négatif donc S décroît.
4. En utilisant encore le théorème des séries alternées, on voit que la série converge uniformément sur $[1, +\infty[$. De plus, $\lim_{x \rightarrow +\infty} f_n(x) = 0$ donc, d'après le théorème de la double limite

$$\lim_{x \rightarrow +\infty} S(x) = \sum_{n=0}^{+\infty} 0 = 0.$$

5. Le premier terme de la somme $f_0 : x \mapsto \frac{1}{x}$ empêche de faire la même chose. On commence par le mettre de côté. On pose $S = f_0 + S_1$ où $S_1 = \sum_{n=1}^{+\infty} f_n$. En procédant de même (double limite), on montre que

$$S_1(x) \xrightarrow{x \rightarrow 0^+} \sum_{n=1}^{+\infty} \frac{(-1)^n}{n} = -\ln(2)$$

On a donc $S(x) \underset{x \rightarrow 0}{\sim} \frac{1}{x}$ donc $\lim_{x \rightarrow 0} S(x) = +\infty$.

6. L'équivalent en $+\infty$ est plus difficile. Il faut d'abord remarquer que $S(x) + S(x+1) = \frac{1}{x}$ (par un décalage) puis que

$$S(x) + S(x+1) \leq 2S(x) \leq S(x) + S(x-1)$$

On en déduit que

$$S(x) \underset{+\infty}{\sim} \frac{1}{2x}.$$

Dénombrabilité et familles sommables

1	Ensembles dénombrables	154
1.1	Définitions	154
1.2	Propriétés et exemples	156
2	Familles sommables	159
2.1	L'ensemble $[0, +\infty]$	159
2.2	Familles sommables d'éléments de $[0, +\infty]$	161
2.3	Familles sommables de nombres complexes	167
2.4	Produit de Cauchy	172

1 Ensembles dénombrables

1.1 Définitions

Commençons par rappeler quelques définitions du cours de première année sur les ensembles finis.

Définition 6.1 (Équipotence)

Soit X et Y deux ensembles. Ils sont dit équipotents s'il existe une bijection $\varphi : X \rightarrow Y$.

Remarque : On vérifie aisément que la relation d'équipotence parfois notée $\equiv$ est une relation d'équivalence¹.

Définition 6.2 (Ensembles finis)

Un ensemble X est dit fini s'il est équipotent à un ensemble de la forme $\llbracket 1 ; p \rrbracket$ où $p \in \mathbb{N}$.

Remarques :

1. Nous passons sous silence qu'il n'y a pas d'ensemble de tous les ensembles et qu'il n'est pas clair sur quel ensemble la relation d'équipotence porte

1. Les ensembles $[[1; p]]$ peuvent être vus comme des « prototypes » d'ensembles finis.
2. Pour $p = 0$, $[[1; 0]] = \emptyset$ qui est un ensemble fini.

Proposition-Définition 6.3 (Cardinal)

Soit X un ensemble fini. Il existe un unique $p \in \mathbb{N}$ tel que $X \equiv [[1; p]]$. On appelle cardinal de X cet unique entier p . On le note

$$|X| \text{ ou } \#X \text{ ou } \text{Card}(X)$$

Définition 6.4 (Ensembles dénombrables)

Un ensemble E est dit dénombrable s'il existe une bijection φ de $\mathbb{N}$ dans E .

Remarques :

1. Cela signifie que l'on peut définir l'ensemble E comme $E = \{x_i \mid i \in \mathbb{N}\}$ en posant $x_i = \varphi(i)$.
2. Soit E un ensemble dénombrable et F un ensemble. S'il existe une bijection de E dans F alors F est aussi dénombrable.
3. Un ensemble fini n'est pas dénombrable.

Exemples :

1. L'ensemble $\mathbb{N}$ est dénombrable.
2. L'ensemble $\mathbb{Z}$ est dénombrable. On peut poser

$$\begin{aligned} \varphi : \mathbb{N} &\rightarrow \mathbb{Z} \\ n &\mapsto \begin{cases} \frac{n}{2} & \text{si } n \text{ est pair} \\ -\frac{n+1}{2} & \text{sinon} \end{cases} \end{aligned}$$

C'est une bijection dont la réciproque est

$$\begin{aligned} \varphi^{-1} : \mathbb{N} &\rightarrow \mathbb{Z} \\ n &\mapsto \begin{cases} 2n & \text{si } n \text{ est positif} \\ -2n - 1 & \text{sinon} \end{cases} \end{aligned}$$

Proposition 6.5

Une partie X de $\mathbb{N}$ ($X \subset \mathbb{N}$) est finie ou dénombrable.

Démonstration : Soit X une partie de $\mathbb{N}$.

- Si elle est majorée et si on note N un majorant de X alors $X \subset [[1; N]]$ donc X est une partie d'un ensemble fini. Elle est finie.
- Si elle n'est pas majorée on peut construire une bijection φ de $\mathbb{N}$ dans X . On pose

$$\begin{aligned} \varphi : \mathbb{N} &\rightarrow X \\ 0 &\mapsto \min(X) \\ 1 &\mapsto \min(X \setminus \{\varphi(0)\}) \\ &\vdots \\ n &\mapsto \min(X \setminus \varphi([0; n-1])) \end{aligned}$$

Le fait que X ne soit pas majorée permet d'être sûr que pour tout entier n , $X \setminus \varphi(\llbracket 0 ; n - 1 \rrbracket)$ n'est pas vide.

De plus, par construction, φ est strictement croissante donc injective.

Il ne reste plus qu'à montrer que φ est surjective. On remarque que pour $x \in X$, $x = \varphi(n)$ où $n = \#\{k \in X \mid k < x\}$.

□

Corollaire 6.6

Soit X un ensemble tel qu'il existe une injection $\iota : X \rightarrow Y$ où Y est fini ou dénombrable. Alors X est fini ou dénombrable. On dit que X est au plus dénombrable.

Démonstration : Il suffit de voir qu'il est en bijection avec $\iota(X)$ qui est une partie de $\mathbb{N}$. □

Corollaire 6.7

Un ensemble X tel qu'il existe une surjection $s : Y \rightarrow X$ où Y est fini ou dénombrable. Alors X est fini ou dénombrable.

Démonstration : On a $f : \mathbb{N} \rightarrow Y \rightarrow X$ surjective et donc on peut construire $\iota : X \hookrightarrow \mathbb{N}$ en associant à tout élément x son plus petit antécédent. □

1.2 Propriétés et exemples

Proposition 6.8

Un produit cartésien d'un nombre fini d'ensembles au plus dénombrables est au plus dénombrable.

Démonstration : (Non exigible) Soit $E_1, \dots, E_r$ des ensembles au plus dénombrables. On note φ_i une injection de E_i dans $\mathbb{N}$ pour $i \in \llbracket 1 ; r \rrbracket$. On considère $p_1, \dots, p_r$ des nombres premiers deux à deux distincts et on construit alors $\iota : E_1 \times \dots \times E_r \rightarrow \mathbb{N}$ par

$$\begin{aligned} \iota : E_1 \times \dots \times E_r &\rightarrow \mathbb{N} \\ (x_1, \dots, x_r) &\mapsto p_1^{\varphi_1(x_1)} \times \dots \times p_r^{\varphi_r(x_r)} \end{aligned}$$

Justifions que ι est injective. Soit $u = (x_1, \dots, x_r)$ et $v = (y_1, \dots, y_r)$ deux éléments de $E_1 \times \dots \times E_r$ tels que $\psi(u) = \psi(v)$. Par unicité de la décomposition en nombres premiers, pour tout $i \in \llbracket 1 ; r \rrbracket$, $\varphi_i(x_i) = \varphi_i(y_i)$. En utilisant l'injectivité des applications $\varphi_1, \dots, \varphi_r$, on obtient que pour tout $i \in \llbracket 1 ; r \rrbracket$, $x_i = y_i$ ce qui signifie que $u = v$.

On en déduit qu'il existe une injection $\iota : E_1 \times \dots \times E_r \hookrightarrow \mathbb{N}$. Cela montre que $E_1 \times \dots \times E_r$ est au plus dénombrable. □

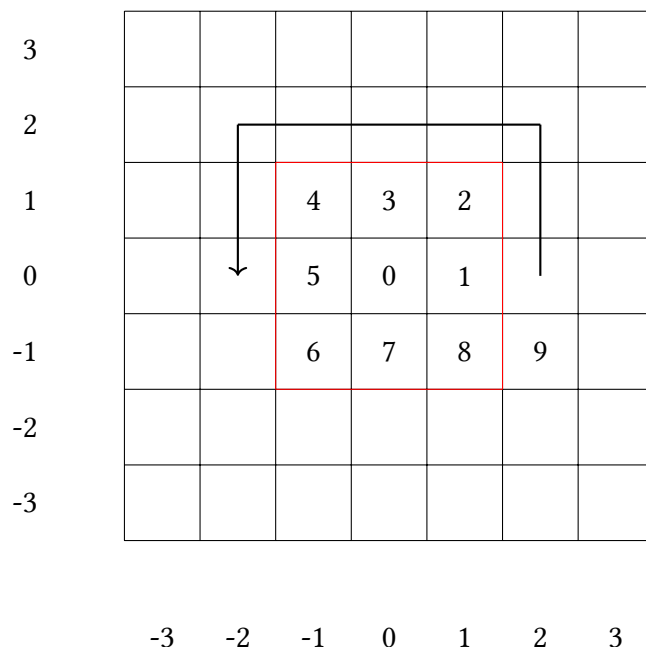
Exemples :

1. On en déduit que $\mathbb{N}^2$, $\mathbb{Z}^2$ et de manière générale $\mathbb{Z}^r$ sont dénombrables.

2. On peut définir une injection de $\mathbb{Q}$ dans $\mathbb{Z}^2$ en utilisant les fractions irréductibles. Pour tout rationnel r , il s'écrit de manière unique sous la forme $r = \frac{a}{b}$ où $a \in \mathbb{Z}$ et $b \in \mathbb{N}^*$ en prenant pour convention que $0 = \frac{0}{1}$. On construit ainsi une injection de $\mathbb{Q}$ dans $\mathbb{Z}^2$ et donc dans $\mathbb{N}$ en composant avec une bijection de $\mathbb{Z}^2$ dans $\mathbb{N}$.

On en déduit que $\mathbb{Q}$ est dénombrable.

3. On peut construire des bijections explicites de $\mathbb{N}$ dans $\mathbb{Z}^2$ par exemple.



Trouver les formules en exercice....

Proposition 6.9

Une réunion finie ou dénombrable d'ensembles finis ou dénombrables est finie ou dénombrable.

Démonstration : (Non exigible) On considère $(E_i)_{i \in I}$ des ensembles finis ou dénombrables indexés par un ensemble I fini ou dénombrable. On se donne donc des injections $\varphi_i : E_i \rightarrow \mathbb{N}$ et une injection $\psi : I \rightarrow \mathbb{N}$. On pose $E = \bigcup_{i \in I} E_i$. Pour tout x de E on note $I_x = \{i \in I \mid x \in E_i\}$ puis i_x l'élément de I_x tel que pour tout $i \in I_x$, $\psi(i_x) \leq \psi(i)$.

On pose alors $\Phi : E \rightarrow \mathbb{N}^2$ définie par $\Phi : x \mapsto (\psi(i_x), \varphi_{i_x}(x))$.

Il est alors élémentaire de vérifier que Φ est injective. En composant avec une injection de $\mathbb{N}^2$ dans $\mathbb{N}$ on obtient que E est fini ou dénombrable.

□

Corollaire 6.10

1. On retrouve que $\mathbb{Z}$ est dénombrable car $\mathbb{Z} = \mathbb{N} \cup -\mathbb{N}$.
2. On retrouve que $\mathbb{Q}$ est dénombrable car $\mathbb{Q} = \bigcup_{p \in \mathbb{N}^*} \frac{1}{p}\mathbb{Z}$.

ATTENTION

Il n'est pas vrai cependant qu'un produit cartésien d'un nombre dénombrable d'ensembles dénombrables (et même fini) est dénombrable.

Exemple : Montrons que $E = \{0, 1\}^{\mathbb{N}}$ n'est pas dénombrable. On peut voir E comme $\{0, 1\} \times \{0, 1\} \times \dots$ ou comme l'ensemble des suites de $\mathbb{N}$ dans $\{0, 1\}$.

Supposons par l'absurde qu'il existe une bijection $\varphi : \mathbb{N} \rightarrow E$. On va aboutir à une absurdité en construisant une suite qui n'appartient pas à $\text{Im}(\varphi)$.

Soit $(u_n) \in E$ définie par

$$\forall n \in \mathbb{N}, u_n = \begin{cases} 1 & \text{si le } n\text{-ième terme de la suite } \varphi(n) \text{ vaut } 0 \\ 0 & \text{si le } n\text{-ième terme de la suite } \varphi(n) \text{ vaut } 1 \end{cases}$$

Supposons par exemple que les premiers termes des suites $\varphi(0), \varphi(1), \dots$ soient donnés par le tableau ci-dessous,

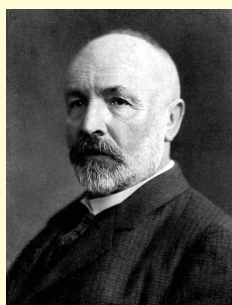
n	$\varphi(n)$						
0	0	1	1	0	1	0	...
1	1	1	1	1	1	1	...
2	1	0	0	1	0	1	...
3	1	1	0	0	0	1	...
4	0	0	0	1	1	1	...
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	

La suite (u_n) est définie par $u_0 = 1, u_1 = 0, u_2 = 1, u_3 = 1, \dots$

La suite (u_n) n'est pas dans $\text{Im}(\varphi)$. En effet pour tout entier n , la suite $(u_n) \neq \varphi(n)$ puisque leur n -ième terme est différent.

Cela montre que φ n'est pas surjective. L'ensemble E n'est pas dénombrable.

Remarque : Ce procédé s'appelle l'extraction diagonale de Cantor.

Matheux (Georg Cantor : 1845 - 1918)

Georg Cantor est un mathématicien allemand, né le 3 mars 1845 à Saint-Petersbourg (Empire russe) et mort le 6 janvier 1918 à Halle (Empire allemand). Il est connu pour être le créateur de la théorie des ensembles.

Il établit l'importance de la bijection entre les ensembles, définit les ensembles infinis et les ensembles bien ordonnés. Il prouva également que les nombres réels sont « plus nombreux » que les entiers naturels. En fait, le théorème de Cantor implique l'existence d'une « infinité d'infinis ». Il définit les nombres cardinaux, les nombres ordinaux et leur arithmétique. Le travail de Cantor est d'un grand intérêt philosophique et a donné lieu à maintes interprétations et à maints débats.

Cantor a été confronté à la résistance de la part des mathématiciens de son époque, en particulier Kronecker. Dans le but de contrer les détracteurs de Cantor, David Hilbert a affirmé : « Nul ne doit nous exclure du Paradis que Cantor a créé. »

Corollaire 6.11

L'ensemble $\mathbf{R}$ n'est pas dénombrable.

Démonstration : (Non exigible) Il suffit de montrer que $[0, 1[$ n'est pas dénombrable. On sait que via l'écriture binaire, $[0, 1[$ est en bijection avec l'ensemble des suites non stationnaires à 1.

En effet, tout nombre réel x dans $[0, 1[$ s'écrit :

$$x = \sum_{n=1}^{+\infty} u_n 2^{-n}$$

Il y a unicité de cette écriture si on supprime les suites (u_n) qui sont stationnaires à 1.

On remarque alors que l'ensemble $\mathcal{S}$ des suites stationnaires à 1 est dénombrable. En effet, notons pour tout entier N , $\mathcal{S}_N$ l'ensemble des suites (u_n) telles que pour $n \geq N$, $u_n = 1$. On voit que $\mathcal{S}_N$ est un ensemble fini. L'ensemble $\mathcal{S}$ est donc une réunion dénombrables d'ensembles finis, c'est un ensemble dénombrable.

Comme on a vu que l'ensemble des suites à valeurs dans $\{0, 1\}$ n'était pas dénombrable et que $\mathcal{S}$ était dénombrable, on obtient que l'ensemble des suites à valeurs dans $\{0, 1\}$ qui ne sont pas stationnaires à 1 n'est pas dénombrable.

Cela permet de montrer que $[0, 1[$ puis $\mathbf{R}$ n'est pas dénombrable.

□

Exercice : Soit $\alpha \in \mathbf{R}$. Il est dit algébrique s'il existe un polynôme $P \in \mathbf{Q}[X]$ tel que $P(\alpha) = 0$. Montrer que l'ensemble des nombres algébriques est dénombrable et en déduire que son complémentaire, l'ensemble des nombres transcendants ne l'est pas.

2 Familles sommables

Nous allons nous intéresser aux familles sommables. L'idée est proche de celle des séries. On va considérer des familles $(u_i)_{i \in I}$ de nombres (réels ou complexes) indexées par un ensemble I (généralisant le cas des suites où $I = \mathbf{N}$). Le but est d'alors de définir la somme de tous ces termes $\sum_{i \in I} u_i$ quand cela a un sens.

2.1 L'ensemble $[0, +\infty]$

Nous allons commencer par rappeler quelques conventions sur l'ensemble $[0, +\infty] = \mathbf{R}_+ \cup \{+\infty\}$.

Proposition-Définition 6.12 (Règles de calculs)

L'addition et le produit de $\mathbf{R}_+$ se prolonge à l'ensemble $[0, +\infty]$. Précisément pour x, y dans $[0, +\infty]$ on pose

$$x + y = \begin{cases} x + y & \text{si } x, y \in \mathbf{R}_+ \\ +\infty & \text{sinon} \end{cases} \quad \text{et } xy = \begin{cases} 0 & \text{si } x = 0 \text{ ou } y = 0 \\ xy & \text{si } x, y \in \mathbf{R}_+ \\ +\infty & \text{sinon} \end{cases}$$

Avec ces conventions, les règles usuelles de calculs (commutativité, associativité, distributivité de la multiplication par rapport à l'addition) restent vraies.

Remarque : Notons en particulier de $0 \times (+\infty) = (+\infty) \times 0 = 0$.

Proposition-Définition 6.13 (Relation d'ordre)

La relation d'ordre usuelle de $\mathbf{R}_+$ se prolonge à l'ensemble $[0, +\infty]$. Pour x, y dans $[0, +\infty]$,

$$x \leq y \iff [(y = +\infty) \text{ ou } (x, y \in \mathbf{R}_+ \text{ et } x \leq y)]$$

Remarque : On vérifie aisément que la relation d'ordre reste compatible aux deux opérations dans le sens où pour x, y, z, t dans $[0, +\infty]$

$$\begin{cases} x \leq y \\ z \leq t \end{cases} \Rightarrow ((x + z) \leq (y + t) \text{ et } xz \leq yt)$$

Proposition 6.14

Soit A une partie non vide d'éléments de $[0, +\infty]$. Elle admet une borne supérieure dans $[0, +\infty]$.

Plus précisément, si $A \subset \mathbf{R}_+$ et A est majorée, la borne supérieure de A vue comme partie de $[0, +\infty]$ est la même que celle de A vue comme partie de $\mathbf{R}_+$. Dans le cas contraire la borne supérieure est égale à $+\infty$.

2.2 Familles sommables d'éléments de $[0, +\infty]$

Définition 6.15

Soit $(u_i)_{i \in I}$ une famille d'éléments de $[0, +\infty]$ indexée par I . On dit qu'elle est sommable si

$$\left\{ \sum_{i \in J} u_i, J \subset I \text{ et } J \text{ fini} \right\}$$

est un ensemble majoré.

Dans ce cas, on appelle somme de la famille et on note $\sum_{i \in I} u_i$ la borne supérieure de l'ensemble ci-dessus :

$$\sum_{i \in I} u_i = \sup_{\substack{J \subset I \\ J \text{ fini}}} \sum_{i \in J} u_i$$

Notation : Dans le cas où la famille n'est pas sommable on note $\sum_{i \in I} u_i = +\infty$.

Remarques :

1. Si I est fini, la famille est sommable et $\sum_{i \in I} u_i$ est la somme des termes.
2. S'il existe $i \in I$ tel que $u_i = +\infty$ alors $\sum_{i \in I} u_i = +\infty$.
3. Comme dans le cas des intégrales (et contrairement au cas des séries) on utilisera la même notation pour la famille et sa somme. Mais il faut cependant faire attention que c'est un abus de notation qui représente deux objets de nature différente.

Proposition 6.16

Soit $(u_i)_{i \in I}$ une famille sommable de réels positifs, le support de la famille c'est-à-dire l'ensemble $T = \{i \in I, u_i \neq 0\}$ est fini ou dénombrable.

Démonstration : Notons $S = \sum_{i \in I} u_i$. Pour tout $p \in \mathbb{N}^*$, on considère $I_p = \left\{ i \in I, u_i > \frac{1}{p} \right\}$.

L'ensemble I_p est fini et on a même $\#I_p \leq pS$. En effet, si on suppose que I_p contient strictement plus de pS éléments. Soit J une partie finie de I_p contenant strictement plus de pS éléments,

$$S \geq \sum_{i \in J} u_i \geq \sum_{i \in J} \frac{1}{p} > pS \times \frac{1}{p} = S$$

C'est donc absurde.

On voit alors que $T = \bigcup_{p \in \mathbb{N}^*} I_p$ est fini ou dénombrable comme union dénombrable d'ensemble finis.

□

Proposition 6.17 (Croissance de la somme)

Soit $(u_i)_{i \in I}$ et $(v_i)_{i \in I}$ deux familles d'éléments de $[0, +\infty]$ indexées par un ensemble I . Si pour tout i de I , $u_i \leq v_i$ alors

$$\sum_{i \in I} u_i \leq \sum_{i \in I} v_i$$

En particulier si $(v_i)_{i \in I}$ est sommable, $(u_i)_{i \in I}$ aussi.

Démonstration :

Soit K une partie finie de I ,

$$\sum_{i \in K} u_i \leq \sum_{i \in K} v_i \leq \sum_{i \in I} v_i$$

On en déduit que $\sum_{i \in I} v_i$ est un majorant de l'ensemble $\left\{ \sum_{i \in J} u_i, J \subset I \text{ et } J \text{ fini} \right\}$ et donc que $\sum_{i \in I} u_i \leq \sum_{i \in I} v_i$

□

Proposition 6.18 (Lien avec les séries)

Soit $(u_i)_{i \in \mathbb{N}}$ une famille de réels positifs indexées par $I = \mathbb{N}$.

$$\sum_{i \in \mathbb{N}} u_i = \sum_{i=0}^{+\infty} u_i$$

En particulier $(u_i)_{i \in I}$ est sommable si et seulement si la série $\sum_{i \geq 0} u_i$ est convergente.

Démonstration : On procède par double inclusion

– $\boxed{\leq}$ Soit I' une partie finie de $I = \mathbb{N}$. On note $N = \max(I')$ son plus grand élément. Alors

$$\sum_{i \in I'} u_i \leq \sum_{i=0}^N u_i \leq \sum_{n=0}^{+\infty} u_i$$

On en déduit que

$$\sum_{i \in \mathbb{N}} u_i \leq \sum_{i=0}^{+\infty} u_i$$

– $\boxed{\geq}$ Soit $n \in \mathbb{N}$. L'ensemble $[[0; n]]$ est une partie finie de $I = \mathbb{N}$ donc

$$\sum_{i=0}^n u_i = \sum_{i \in [[0; n]]} u_i \leq \sum_{i \in \mathbb{N}} u_i$$

En faisant tendre n vers $+\infty$ on obtient que

$$\sum_{i \in \mathbb{N}} u_i \geq \sum_{i=0}^{+\infty} u_i$$

□

Proposition 6.19 (Invariance par permutation des indices)

Soit (u_i) une famille d'éléments de $[0, +\infty]$ indexée par I . Soit J un ensemble en bijection avec I et φ une bijection de J dans I . Pour tout $j \in J$ on pose $v_j = u_{\varphi(j)}$. On a alors

$$\sum_{i \in I} u_i = \sum_{j \in J} v_j$$

En particulier la famille $(v_j)_{j \in J}$ est sommable si et seulement si $(u_i)_{i \in I}$ l'est.

Démonstration : Notons $U = \left\{ \sum_{i \in I'} u_i, I' \subset I \text{ et } I' \text{ fini} \right\}$ et $V = \left\{ \sum_{j \in J'} v_j, J' \subset J \text{ et } J' \text{ fini} \right\}$. Montrons que $U = V$.

Soit $x \in V$. Il existe $J' \subset J$ un ensemble fini tel que $x = \sum_{j \in J'} v_j$. Si on note $I' = \varphi(J')$ alors I' est une partie finie de I et

$$\sum_{j \in J'} v_j = \sum_{j \in J'} u_{\varphi(j)} = \sum_{i \in I'} u_i$$

Cela montre que $x \in U$ et donc $V \subset U$.

On peut considérer φ^{-1} qui est une bijection de I dans J pour montrer que réciproquement $U \subset V$.

Finalement $U = V$ et donc

$$\sum_{i \in I} u_i = \sup(U) = \sup(V) = \sum_{j \in J} v_j$$

En particulier la famille $(v_j)_{j \in J}$ est sommable si et seulement si $(u_i)_{i \in I}$ l'est.

□

Proposition 6.20 (Opérations)

Soit $(u_i)_{i \in I}$ et $(v_i)_{i \in I}$ deux familles d'éléments de $[0, +\infty]$ indexées par I . Soit k un réel positif,

$$\sum_{i \in I} (u_i + v_i) = \sum_{i \in I} u_i + \sum_{i \in I} v_i \text{ et } \sum_{i \in I} (ku_i) = k \sum_{i \in I} u_i$$

Démonstration : En exercice.

□

Théorème 6.21 (Somme par paquets - cas réels positifs)

Soit $(u_i)_{i \in I}$ une famille de réels positifs indexée par I . Soit $(I_j)_{j \in J}$ une partition de I . Dans $[0, +\infty]$,

$$\sum_{j \in J} \left(\sum_{i \in I_j} u_i \right) = \sum_{i \in I} u_i$$

En particulier, $(u_i)_{i \in I}$ est sommable si et seulement si toutes les familles $(u_i)_{i \in I_j}$ sont sommables et que la famille $\left(\sum_{i \in I_j} u_i \right)_{j \in J}$ est sommable.

Démonstration : Ce théorème est admis. On va démontrer cette égalité par double inégalité.

Avant de commencer la preuve, on va avoir besoin d'un lemme.

Lemme 6.22

Soit $(u_i)_{i \in I}$ une famille de réels positifs. Soit $J \subset I$. Dans $[0, +\infty]$

$$\sum_{i \in J} u_i \leq \sum_{i \in I} u_i$$

Démonstration du lemme : Soit J' une partie finie de J . C'est aussi une partie finie de I donc

$$\sum_{i \in J'} u_i \leq \sum_{i \in I} u_i$$

Cela montre que

$$\sum_{i \in J} u_i \leq \sum_{i \in I} u_i$$

□

- $\supseteq$ Soit $K \subset I$ un ensemble fini. Pour tout $j \in J$, on note $K_j = K \cap I_j$ et on considère T l'ensemble des indices j tels que $K_j \neq \emptyset$. L'ensemble T est une partie finie de J car K est un ensemble fini. On a alors

$$\sum_{i \in K} u_i = \sum_{j \in T} \left(\sum_{i \in K_j} u_i \right) \leq \sum_{j \in T} \left(\sum_{i \in I_j} u_i \right) \leq \sum_{j \in J} \left(\sum_{i \in I_j} u_i \right)$$

où la première inégalité est obtenue en montrant que pour tout $j \in J$, $\sum_{i \in K_j} u_i \leq \sum_{i \in I_j} u_i$ d'après le lemme ci-dessus et la croissance de la somme.

Cela montre que

$$\sum_{i \in I} u_i \leq \sum_{j \in J} \left(\sum_{i \in I_j} u_i \right)$$

- $\leq$ Soit J' une partie finie de J , on veut majorer $\sum_{j \in J'} \left(\sum_{i \in I_j} u_i \right)$. On fixe $\varepsilon > 0$ et on note $N = \#J'$. Par définition de la borne supérieure, pour tout $j \in J'$, il existe une partie finie X_j de I_j telle que

$$\sum_{i \in I_j} u_i - \frac{\varepsilon}{N} \leq \sum_{i \in X_j} u_i \leq \sum_{i \in I_j} u_i$$

En faisant la somme sur tous les éléments de J' et en notant $X = \bigcup_{j \in J'} X_j$,

$$\sum_{j \in J'} \left(\sum_{i \in I_j} u_i \right) - \varepsilon \leq \sum_{j \in J'} \left(\sum_{i \in X_j} u_i \right) = \sum_{i \in X} u_i \leq \sum_{i \in I} u_i$$

L'inégalité de droite venant du fait que X est un ensemble fini.

Cette inégalité étant vraie pour tout $\varepsilon > 0$, on a donc

$$\sum_{j \in J'} \left(\sum_{i \in I_j} u_i \right) \leq \sum_{i \in I} u_i$$

Cela implique que

$$\sum_{j \in J} \left(\sum_{i \in I_j} u_i \right) \leq \sum_{i \in I} u_i$$

□

Remarques :

1. Ce théorème est important. C'est celui que l'on utilisera le plus souvent.
2. S'il existe un $j \in J$ tel que $\sum_{i \in I_j} u_i = +\infty$ alors automatiquement $\sum_{j \in J} \left(\sum_{i \in I_j} u_i \right) = +\infty$ et donc $\sum_{i \in I} u_i = +\infty$.

Corollaire 6.23 (Théorème de Fubini)

Soit A, B deux ensembles et $I = A \times B$. Soit $(u_{a,b})_{(a,b) \in I}$ une famille de réels positifs. On a

$$\sum_{a \in A} \left(\sum_{b \in B} u_{a,b} \right) = \sum_{b \in B} \left(\sum_{a \in A} u_{a,b} \right) = \sum_{(a,b) \in A \times B} u_{a,b}$$

Démonstration : Il suffit d'appliquer le théorème de sommation par paquets en voyant que si on pose pour tout $a \in A$, $I_a = \{a\} \times B$, les ensembles I_a forment une partition de I . □

Remarque : On l'utilisera la plupart du temps avec $A = B = \mathbb{N}$ (ou $\mathbb{N}^*$). Une famille de réels positifs $(u_{p,q})_{(p,q) \in \mathbb{N}^2}$ est sommable si et seulement si :

- pour tout $p \in \mathbb{N}$, la famille $(u_{p,q})_{q \in \mathbb{N}}$ est sommable ce qui revient à demander que la série $\sum_{q \geq 0} u_{p,q}$ converge.
- la famille $\left(\sum_{q \in \mathbb{N}} u_{p,q} \right)_{p \in \mathbb{N}}$ est sommable ce qui revient au fait que la série $\sum_{p \geq 0} \left(\sum_{q \in \mathbb{N}} u_{p,q} \right)$ converge.

Matheux (Guido Fubini : 1879 - 1943)



Né à Venise, Guido Fubini (19 janvier 1879 - 6 juin 1943) est poussé vers les mathématiques à un âge précoce, encouragé par ses professeurs et par son père, lui-même professeur de mathématiques. En 1896 il intègre l'École normale supérieure de Pise, où il est l'élève d'Ulisse Dini et de Luigi Bianchi. Il acquiert une certaine notoriété dès 1902, lorsque sa thèse de doctorat (datant de 1900), intitulée *le parallélisme de Clifford dans les espaces elliptiques*, est discutée dans un travail très lu sur la géométrie différentielle publié par Bianchi.

Il est célèbre notamment pour ses travaux sur les intégrales, en particulier le théorème de Fubini.

Exemples :

1. On considère la famille $(u_{p,q})_{(p,q) \in \mathbb{N} \times \mathbb{N}^*}$ définie par

$$\forall (p, q) \in \mathbb{N} \times \mathbb{N}^*, u_{p,q} = \frac{1}{(p+q^2)(p+q^2+1)}$$

On veut calculer sa somme (et donc voir si elle est sommable).

On remarque que pour $p \geq 0$ et $q \geq 1$,

$$u_{p,q} = \frac{1}{(p+q^2)(p+q^2+1)} = \frac{1}{p+q^2} - \frac{1}{p+q^2+1}$$

Donc, pour tout $q \in \mathbb{N}^*$,

$$\sum_{p=0}^{\infty} u_{p,q} = \sum_{p=0}^{\infty} \frac{1}{p+q^2} - \frac{1}{p+q^2+1} = \frac{1}{q^2}$$

Comme $\sum_{q \geq 1} \frac{1}{q^2} = \zeta(2)$ on obtient

$$\sum_{(p,q) \in \mathbb{N} \times \mathbb{N}^*} \frac{1}{(p+q^2)(p+q^2+1)} = \zeta(2)$$

2. On considère la famille $u_{p,q} = \frac{1}{(p+q)^\alpha}$ où $(p, q) \in (\mathbb{N}^*)^2$. On veut savoir pour quels $\alpha > 0$ c'est une famille sommable. On regroupe par paquets selon la valeur de $p+q$. Précisément, on a

$$I = (\mathbb{N}^*)^2 = \bigcup_{k=2}^{+\infty} I_k \text{ où}$$

$$I_k = \{(p, q) \in (\mathbb{N}^*)^2 \mid p+q = k\}.$$

Il est clair que I_k est fini et $\#I_k = (k-1)$ car

$$I_k = \{(1, k-1), (2, k-2), \dots, (k-1, 1)\}.$$

On en déduit que la famille $(u_{p,q})_{(p,q) \in I_k}$ est sommable et

$$\sum_{(p,q) \in I_k} u_{p,q} = \frac{k-1}{k^\alpha}$$

On obtient alors que la famille est sommable si et seulement si la série $\sum_{k \geq 2} \frac{k-1}{k^\alpha}$ est convergente.

C'est une série à termes positifs et $\frac{k-1}{k^\alpha} \underset{k \rightarrow \infty}{\sim} \frac{1}{k^{\alpha-1}}$. La famille est donc sommable si et seulement si $\alpha - 1 > 1 \iff \alpha > 2$.

Remarque : Dans la plupart des cas on travaillera avec $I = \mathbb{N}^2$ ou $\mathbb{Z}^2$. Les partitions utilisées seront souvent

- Les « lignes » : $L_j = \{(i, j) \mid i \in \mathbb{N}\}$ ou $L_j = \{(i, j) \mid i \in \mathbb{Z}\}$.
- Les « colonnes » : $C_i = \{(i, j) \mid j \in \mathbb{N}\}$ ou $C_j = \{(i, j) \mid j \in \mathbb{Z}\}$.
- Les « diagonales » : $\Delta_k = \{(i, j) \in I \mid i+j = k\}$.

FAIRE DES DESSINS

2.3 Familles sommables de nombres complexes

Nous allons nous intéresser dans ce paragraphe aux familles de réels non nécessairement positifs ou de nombres complexes. Les idées développées sont proches de celles des séries absolument convergente et des fonctions absolument intégrables.

Définition 6.24

Soit $(u_i)_{i \in I}$ une famille de nombres complexes indexée par I . Elle est dite sommable si la famille de réels positifs $(|u_i|)_{i \in I}$ est sommable, c'est-à-dire si $\sum_{i \in I} |u_i| < +\infty$.

On note $\ell^1(I)$ l'ensemble des familles sommables indexées par I .

ATTENTION

Quand on travaille avec des réels positifs, on peut toujours donner un sens à $\sum_{i \in I} u_i$ que la famille soit sommable ou non. Dans le cas des nombres réels qui ne sont pas de signe constant ou dans le cas des nombres complexes ce n'est plus le cas. Il **ne faut pas** écrire $\sum_{i \in I} u_i$ avant d'avoir montré la sommabilité de la famille.

Proposition-Définition 6.25

1. Soit $(u_i)_{i \in I}$ une famille sommable de nombres réels indexée par I . Les familles $(u_i^+)_{i \in I}$ et $(u_i^-)_{i \in I}$ sont sommables. On pose alors

$$\sum_{i \in I} u_i = \sum_{i \in I} u_i^+ - \sum_{i \in I} u_i^-.$$

2. Soit $(u_i)_{i \in I}$ une famille sommable de nombres complexes indexée par I . Les familles $(\Re(u_i))_{i \in I}$ et $(\Im(u_i))_{i \in I}$ sont sommables. On pose alors

$$\sum_{i \in I} u_i = \sum_{i \in I} \Re(u_i) + i \sum_{i \in I} \Im(u_i).$$

Rappel : Pour tout réel x on appelle partie positive (resp. négative) et on note x^+ (resp. x^-) les réels positifs :

$$x^+ = \begin{cases} x & \text{si } x \geq 0 \\ 0 & \text{sinon} \end{cases} ; x^- = \begin{cases} 0 & \text{si } x \geq 0 \\ -x & \text{sinon} \end{cases}$$

On a ainsi $x = x^+ - x^-$.

Démonstration : Il suffit de voir que si $x \in \mathbf{R}$, $x^+ \leq |x|$ et $x^- \leq |x|$. De même si $x \in \mathbf{C}$, $|\Re(x)| \leq |x|$ et $|\Im(x)| \leq |x|$. $\square$

Proposition 6.26 (Invariance par permutation des indices)

Soit (u_i) une famille de nombres complexes indexée par I . Soit J un ensemble en bijection avec I et φ une bijection de J dans I . Pour tout $j \in J$ on pose $v_j = u_{\varphi(j)}$.

1. La famille $(v_j)_{j \in J}$ est sommable si et seulement si $(u_i)_{i \in I}$ l'est.
2. Dans ce cas,

$$\sum_{j \in J} v_j = \sum_{i \in I} u_i.$$

Démonstration : On a déjà traité le cas des réels positifs. Montrons que l'on peut s'y ramener dans le cas général

1. Pour la sommabilité

$$(v_j)_{j \in J} \text{ sommable} \iff (|v_j|)_{j \in J} \text{ sommable} \iff (|u_i|)_{i \in I} \text{ sommable} \iff (u_i)_{i \in I} \text{ sommable}$$

2. Dans le cas où les familles sont sommables et que ce sont des familles de réels,

$$\sum_{i \in I} u_i = \sum_{i \in I} u_i^+ - \sum_{i \in I} u_i^- = \sum_{j \in J} v_j^+ - \sum_{j \in J} v_j^- = \sum_{j \in J} v_j$$

Le cas des nombres complexes est similaire.

□

Proposition 6.27 (Sommabilité d'une sous-famille)

Soit (u_i) une famille de nombres complexes indexée par I . Soit $J \subset I$ une partie de I . Si $(u_i)_{i \in I}$ est sommable alors $(u_i)_{i \in J}$ aussi.

Démonstration : On a déjà traité le cas des réels positifs. Le cas général s'y ramène

$$(u_i)_{i \in I} \text{ sommable} \iff (|u_i|)_{i \in I} \text{ sommable} \Rightarrow (|u_i|)_{i \in J} \text{ sommable} \Rightarrow (u_i)_{i \in J} \text{ sommable}$$

□

Proposition 6.28 (Sommabilité par majoration)

Soit (u_i) une famille de nombres complexes indexée par I . Soit $(v_i)_{i \in I}$ une famille sommable de réels positifs. On suppose que pour tout $i \in I$, $|u_i| \leq v_i$ alors $(u_i)_{i \in I}$ est sommable.

Démonstration : Il suffit de voir que

$$(v_i)_{i \in I} \text{ sommable} \Rightarrow (|u_i|)_{i \in I} \text{ sommable} \Rightarrow (u_i)_{i \in I} \text{ sommable}$$

□

Proposition 6.29 (Lien avec les séries)

Dans le cas où $I = \mathbb{N}$, la famille $(u_i)_{i \in \mathbb{N}}$ est sommable si et seulement si la série $\sum_{i \geq 0} u_i$ est absolument convergente. Dans ce cas

$$\sum_{i \in \mathbb{N}} u_i = \sum_{i=0}^{+\infty} u_i.$$

Démonstration :

On a

$$(u_i)_{i \in \mathbb{N}} \text{ sommable} \iff (|u_i|)_{i \in \mathbb{N}} \text{ sommable} \iff \sum_{i \geq 0} |u_i| \text{ converge} \iff \sum_{i \geq 0} u_i \text{ converge absolument}$$

Dans ce cas vérifions que la formule sur les sommes se déduit de celle dans le cas des réels positifs.

— Cas des réels :

$$\sum_{i \in I} u_i = \sum_{i \in I} u_i^+ - \sum_{i \in I} u_i^- = \sum_{i=0}^{+\infty} u_i^+ - \sum_{i=0}^{+\infty} u_i^- = \sum_{i=0}^{+\infty} (u_i^+ - u_i^-) = \sum_{i=0}^{+\infty} u_i.$$

— Cas des complexes :

$$\sum_{i \in I} u_i = \sum_{i \in I} \Re(u_i) + i \sum_{i \in I} \Im(u_i) = \sum_{i=0}^{+\infty} \Re(u_i) + i \sum_{i=0}^{+\infty} \Im(u_i) = \sum_{i=0}^{+\infty} (\Re(u_i) + i \Im(u_i)) = \sum_{i=0}^{+\infty} u_i.$$

□

Corollaire 6.30 (Propriétés)

1. (Linéarité de la somme) Soit $(u_i)_{i \in I}$ et $(v_i)_{i \in I}$ deux familles sommables. Soit (λ, μ) deux scalaires. La famille $(\lambda u_i + \mu v_i)_{i \in I}$ est sommable et

$$\sum_{i \in I} \lambda u_i + \mu v_i = \lambda \sum_{i \in I} u_i + \mu \sum_{i \in I} v_i$$

2. (Inégalité triangulaire) Soit $(u_i)_{i \in I}$ une famille sommable,

$$\left| \sum_{i \in I} u_i \right| \leq \sum_{i \in I} |u_i|$$

Démonstration :

1. On sait que comme $(u_i)_{i \in I}$ et $(v_i)_{i \in I}$ sont sommables les supports Z_u et Z_v de ces familles sont finis ou dénombrables. On peut alors remplacer I par $Z_u \cup Z_v$ qui est encore fini ou dénombrable. Le cas où I est fini étant évident, on peut supposer que I est dénombrable.

Soit φ une bijection entre $\mathbb{N}$ et I . On considère $x_n = u_{\varphi(n)}$ et $y_n = v_{\varphi(n)}$.

On a alors

$$\begin{aligned}
(u_i)_{i \in I} \text{ et } (v_i)_{i \in I} \text{ sommables} &\Rightarrow (x_n)_{n \in \mathbb{N}} \text{ et } (y_n)_{n \in \mathbb{N}} \text{ sommables} \\
&\Rightarrow \sum_{n \geq 0} x_n \text{ et } \sum_{n \geq 0} y_n \text{ absolument convergentes} \\
&\Rightarrow \sum_{n \geq 0} (\lambda x_n + \mu y_n) \text{ absolument convergente} \\
&\Rightarrow (\lambda x_n + \mu y_n)_{n \in \mathbb{N}} \text{ sommable} \\
&\Rightarrow (\lambda u_i + \mu v_j)_{i \in I} \text{ sommable}
\end{aligned}$$

De plus

$$\sum_{i \in I} \lambda u_i + \mu v_i = \sum_{n=0}^{\infty} \lambda x_n + \mu y_n = \lambda \sum_{n=0}^{\infty} x_n + \mu \sum_{n=0}^{\infty} y_n = \lambda \sum_{i \in I} u_i + \mu \sum_{i \in I} v_i$$

2. Il suffit de revenir à la définition de la somme d'une famille sommable.

□

Il y a aussi un théorème de sommation par paquets.

Proposition 6.31 (Théorème d'approximation)

Soit $(u_i)_{i \in I}$ une famille sommable. Soit $\varepsilon > 0$. Il existe une partie $F \subset I$ finie telle que

$$\left| \sum_{i \in F} u_i - \sum_{i \in I} u_i \right| \leq \varepsilon$$

Remarque : On peut même montrer un résultat un peu plus fort qui se rapproche plus de la définition de la limite. Il existe un ensemble F fini tel que pour tout ensemble F' contenant F ,

$$\left| \sum_{i \in F'} u_i - \sum_{i \in I} u_i \right| \leq \varepsilon$$

Démonstration :

□

Théorème 6.32 (Sommation par paquets - cas complexe)

Soit $(u_i)_{i \in I}$ une famille de nombres complexes indexée sur un ensemble I . Soit $(I_j)_{j \in J}$ une partition de I . La famille $(u_i)_{i \in I}$ est sommable si et seulement si

- Pour tout $j \in J$, la famille $(u_i)_{i \in I_j}$ est sommable.
- La famille $\left(\sum_{i \in I_j} |u_i| \right)_{j \in J}$ est sommable.

Dans ce cas,

$$\sum_{i \in I} u_i = \sum_{j \in J} \left(\sum_{i \in I_j} u_i \right)$$

Démonstration : Ce théorème est admis.

- Commençons par la partie sur la sommabilité. D'après le théorème de sommation par paquets pour les réels positifs, dans $[0, +\infty]$,

$$\sum_{j \in J} \left(\sum_{i \in I_j} |u_i| \right) = \sum_{i \in I} |u_i|$$

Comme la famille $(u_i)_{i \in I}$ est sommable si et seulement si la famille $(|u_i|)_{i \in I}$ est sommable, on obtient le résultat voulu.

- On suppose maintenant que la famille $(|u_i|)_{i \in I}$ est sommable. Il suffit de revenir à la définition. Traitons le cas où $K = \mathbb{R}$.

$$\sum_{i \in I} u_i = \sum_{i \in I} u_i^+ - \sum_{i \in I} u_i^- = \sum_{j \in J} \sum_{i \in I_j} u_i^+ - \sum_{j \in J} \sum_{i \in I_j} u_i^- = \sum_{j \in J} \sum_{i \in I_j} (u_i^+ - u_i^-) = \sum_{j \in J} \sum_{i \in I_j} u_i$$

□

ATTENTION

Il est important de noter que c'est la famille $\left(\sum_{i \in I_j} |u_i| \right)_{j \in J}$ qui doit être sommable et pas juste $\left(\sum_{i \in I_j} u_i \right)_{j \in J}$. Par exemple pour $u_{p,q} = \begin{cases} 1 & \text{si } p > q \\ -1 & \text{si } p < q \\ 0 & \text{si } p = q \end{cases}$. Si on considère les paquets $I_n = \{(p, q) \in \mathbb{N}^2, p + q = n\}$, on a que pour tout entier n , $(u_{p,q})_{(p,q) \in I_n}$ est sommable (car finie) et que $\sum_{(p,q) \in I_n} u_{p,q} = 0$. Par contre la famille n'est pas sommable.

Corollaire 6.33 (Théorème de Fubini)

Soit A, B deux ensembles et $I = A \times B$. Soit $(u_{a,b})_{(a,b) \in I}$ une famille indexée par I .

La famille $(u_{a,b})_{(a,b) \in I}$ est sommable si et seulement si

- Pour tout $a \in A$, la famille $(u_{a,b})_{b \in B}$ est sommable.
- La famille $\left(\sum_{b \in B} |u_{a,b}| \right)_{a \in A}$ est sommable

Dans ce cas

$$\sum_{a \in A} \left(\sum_{b \in B} u_{a,b} \right) = \sum_{b \in B} \left(\sum_{a \in A} u_{a,b} \right) = \sum_{(a,b) \in A \times B} u_{a,b}$$

Corollaire 6.34 (Théorème de Fubini pour $I = \mathbb{N}^2$)

Soit $(u_{n,m})_{(n,m) \in \mathbb{N}^2}$ une famille de nombres complexes indexée par $\mathbb{N}^2$. Elle est sommable si et seulement si

- Pour tout entier n , la famille la série $\sum_{m \geq 0} u_{n,m}$ converge absolument.
- La série $\sum_{n \geq 0} \left(\sum_{m=0}^{+\infty} |u_{n,m}| \right)$ converge.

Dans ce cas,

$$\sum_{n=0}^{+\infty} \left(\sum_{m=0}^{+\infty} u_{n,m} \right) = \sum_{m=0}^{+\infty} \left(\sum_{n=0}^{+\infty} u_{n,m} \right).$$

Corollaire 6.35

Soit $(a_i)_{i \in I}$ et $(b_j)_{j \in J}$ deux familles supposées sommables. La famille $(a_i b_j)_{(i,j) \in I \times J}$ est sommable et

$$\sum_{i \in I} a_i \times \sum_{j \in J} b_j = \sum_{(i,j) \in I \times J} a_i b_j$$

Démonstration : On applique le théorème de sommations par paquets à la famille $(a_i b_j)_{(i,j) \in I \times J}$:

- Pour tout $i \in I$, la famille $(a_i b_j)_{j \in J}$ est sommable car c'est juste la multiplication par un scalaire fixé de la famille $(b_j)_{j \in J}$
- La famille $\left(\sum_{j \in J} |a_i b_j| \right)_{i \in I}$ car c'est juste la famille $(|a_i|)_{i \in I}$ multipliée par le scalaire $\sum_{j \in J} |b_j|$.

De plus

$$\sum_{(i,j) \in I \times J} a_i b_j = \sum_{i \in I} \left(\sum_{j \in J} a_i b_j \right) = \sum_{i \in I} \left(a_i \sum_{j \in J} b_j \right) = \sum_{i \in I} a_i \times \sum_{j \in J} b_j$$

2.4 Produit de Cauchy

Définition 6.36

Soit $\sum_{n \geq 0} a_n$ et $\sum_{n \geq 0} b_n$ deux séries numériques (réelles ou complexes). On appelle produit de Cauchy la série $\sum_{p \geq 0} c_p$ où

$$\forall p \in \mathbb{N}, c_p = \sum_{n+m=p} a_n b_m = \sum_{n=0}^p a_n b_{p-n} = \sum_{m=0}^p a_{p-m} b_m.$$

Matheux (Augustin Louis Cauchy : 1789 - 1857)

Augustin Louis, baron Cauchy, né à Paris le 21 août 1789 et mort à Sceaux le 23 mai 1857, est un mathématicien français, membre de l'Académie des sciences et professeur à l'École polytechnique. Il est l'un des mathématiciens les plus prolifiques de l'histoire. On lui doit notamment en analyse l'introduction des fonctions holomorphes et des critères de convergence des suites et des séries entières. Ses travaux sur les permutations sont précurseurs de la théorie des groupes. En optique, on lui doit des travaux sur la propagation des ondes électromagnétiques.

Son œuvre a fortement influencé le développement des mathématiques au XIXe siècle, mais le fait qu'il publie ses résultats dès leur découverte sans y appliquer toute la rigueur souhaitée et la négligence dont il fait preuve concernant les travaux d'Évariste Galois et de Niels Abel entachent son prestige.

Théorème 6.37

Si les séries $\sum_{n \geq 0} a_n$ et $\sum_{n \geq 0} b_n$ convergent absolument alors la série $\sum_{p \geq 0} c_p$ converge absolument et

$$\sum_{p=0}^{+\infty} c_p = \left(\sum_{n=0}^{+\infty} a_n \right) \left(\sum_{m=0}^{+\infty} b_m \right).$$

Démonstration : Les séries $\sum_{n \geq 0} a_n$ et $\sum_{n \geq 0} b_n$ étant absolument convergentes, les familles $(a_n)_{n \in \mathbb{N}}$ et $(b_n)_{n \in \mathbb{N}}$ sont sommables et donc la famille $(a_n b_m)_{(n,m) \in \mathbb{N}^2}$ est sommable.

De plus

$$\sum_{(n,m) \in \mathbb{N}^2} a_n b_m = \left(\sum_{n=0}^{+\infty} a_n \right) \left(\sum_{m=0}^{+\infty} b_m \right).$$

Il suffit maintenant de découper selon les paquets $I_p = \{(n, m) \in \mathbb{N}^2, n + m = p\}$. On a alors

$$\sum_{(n,m) \in \mathbb{N}^2} a_n b_m = \sum_{p \in \mathbb{N}} \left(\sum_{(n,m) \in I_p} a_n b_m \right) = \sum_{p=0}^{\infty} c_p$$

□

Remarque : Ce résultat sera utilisé pour démontrer le produit de Cauchy de deux séries entières.

Exemple : On veut calculer $\sum_{n=0}^{+\infty} (n+1)4^{-n}$. Il est clair que c'est la somme d'une série (absolument) convergente.

— Première méthode : On considère la famille

$$u_{n,m} = \begin{cases} 4^{-n} & \text{si } m \leq n \\ 0 & \text{sinon.} \end{cases}$$

On montre aisément que c'est une famille sommable et

$$\sum_{n=0}^{+\infty} (n+1)4^{-n} = \sum_{n=0}^{+\infty} \sum_{m=0}^{+\infty} u_{n,m} = \sum_{m=0}^{+\infty} \sum_{n=0}^{+\infty} u_{n,m} = \sum_{m=0}^{+\infty} \sum_{n=m}^{+\infty} 4^{-n} \sum_{m=0}^{+\infty} \left(4^{-m} \frac{1}{1 - \frac{1}{4}} \right) = \frac{4}{3} \sum_{m=0}^{+\infty} 4^{-m} = \frac{16}{9}.$$

— Deuxième méthode : On remarque que pour tout entier n ,

$$(n+1) \left(\frac{1}{4} \right)^n = \sum_{k=0}^n \frac{1}{4^k} \frac{1}{4^{n-k}}$$

C'est donc un produit de Cauchy. On a directement,

$$\sum_{n=0}^{+\infty} (n+1)4^{-n} = \left(\sum_{n=0}^{+\infty} 4^{-n} \right) \left(\sum_{k=0}^{+\infty} 4^{-k} \right) = \frac{16}{9}$$

— Troisième méthode : On peut aussi le faire en dérivant $x \mapsto \sum_{n=0}^{+\infty} x^n = \frac{1}{1-x}$. Sur $] -1, 1[$ on a donc

$$\sum_{n=0}^{+\infty} nx^{n-1} = \frac{1}{(1-x)^2}.$$

On en tire pour $x = \frac{1}{4}$

$$A = \sum_{n=0}^{+\infty} n \frac{1}{4}^{n-1} = \frac{1}{(1 - \frac{1}{4})^2} = \frac{16}{9}.$$

Or

$$\sum_{n=0}^{+\infty} (n+1)4^{-n} = \frac{1}{4}A + \sum_{n=0}^{+\infty} 4^{-n} = \frac{4}{9} + \frac{1}{1 - \frac{1}{4}} = \frac{16}{9}.$$

Probabilités

1	Espaces probabilisés	175
1.1	Tribus	175
1.2	Probabilités	177
2	Propriétés élémentaires des probabilités	179
2.1	Continuité	179
2.2	Événements négligeables, presque sûrs	182
3	Indépendance et probabilités conditionnelles	183
3.1	Probabilités conditionnelles	183
3.2	Événements indépendants	186
4	Variables aléatoires discrètes	188
4.1	Généralités	189
4.2	Couples de variables aléatoires et vecteurs aléatoires	192
4.3	Couples et familles de variables aléatoires indépendantes	194
5	Lois usuelles	198
5.1	Loi uniforme	198
5.2	Loi de Bernoulli et loi binomiale	198
5.3	Loi géométrique	199
5.4	La loi de Poisson	201

1 Espaces probabilisés

Nous allons dans ce paragraphe étendre les notions vues en première année au cas où les univers ne sont pas nécessairement finis.

1.1 Tribus

Contrairement au cas des univers Ω finis, il ne sera pas toujours possible de définir les probabilités sur l'intégralité de $\mathcal{P}(\Omega)$ (qui est un ensemble « très grand » qui contient de nombreuses parties qui ne nous intéressent pas). On va donc se restreindre à une à une sous-partie de Ω que l'on appelle tribu (ou σ -algèbre).

Définition 7.1

Soit Ω un ensemble. On appelle tribu sur Ω un sous ensemble $\mathcal{A}$ de $\mathcal{P}(\Omega)$ tel que

1. $\emptyset \in \mathcal{A}$
2. Pour tout A dans $\mathcal{A}$, son complémentaire $\mathbb{C}_{\Omega}A$ appartient à $\mathcal{A}$.
3. Pour toute famille $(A_i)_{i \in I} \in \mathcal{A}^I$ où I est fini ou dénombrable, l'union $\bigcup_{i \in I} A_i$ appartient à $\mathcal{A}$.

Remarques :

1. On dit que $\mathcal{A}$ est stable par complémentaire et par union au plus dénombrable.
2. On obtient que $\Omega = \mathbb{C}_{\Omega}\emptyset \in \mathcal{A}$.
3. On voit aussi que $\mathcal{A}$ est stable par intersection au plus dénombrable car pour $(A_i)_{i \in I} \in \mathcal{A}^I$,

$$\bigcap_{i \in I} A_i = \mathbb{C}_{\Omega} \left(\bigcup_{i \in I} \mathbb{C}_{\Omega} A_i \right)$$

Exemples :

1. L'ensemble $\mathcal{P}(\Omega)$ en entier est une tribu.
2. L'ensemble $\{\emptyset, \Omega\}$ est une tribu.
3. Pour $\Omega = \mathbb{R}$. On peut chercher la plus petite tribu $\mathcal{A}$ qui contient tous les singletons $\{x_0\}$. Par stabilité par union dénombrable, cette tribu, doit contenir tous les ensembles au plus dénombrables car si A est au plus dénombrable,

$$A = \bigcup_{a \in A} \{a\}$$

De plus, par stabilité par complémentaire, $\mathcal{A}$ doit contenir toutes les parties A telle que $\mathbb{C}_{\Omega}A$ est au plus dénombrable.

On peut alors montrer que l'ensemble des parties qui sont au plus dénombrables ou de complémentaire au plus dénombrable est une tribu. C'est la tribu engendrée par les singletons.

Définition 7.2

On appelle espace probabilisable un couple $(\Omega, \mathcal{A})$ où $\mathcal{A}$ est une tribu sur Ω . Les parties A de Ω qui appartiennent à $\mathcal{A}$ s'appellent les événements.

Remarque : Rappelons la terminologie vu en première année :

- Un événement élémentaire est un singleton.
- Un système complet d'événements $(A_i)_{i \in I} \in \mathcal{A}^I$ est une partition de Ω par des parties appartenant à $\mathcal{A}$.
- Soit A, B deux événements, ils sont dit disjoints ou incompatibles si $A \cap B = \emptyset$.

1.2 Probabilités

Définition 7.3

Soit $(\Omega, \mathcal{A})$ un espace probabilisable. Une probabilité sur $(\Omega, \mathcal{A})$ est une application

$$P : \mathcal{A} \rightarrow \mathbb{R}_+$$

telle que

- i) $P(\Omega) = 1$
- ii) la fonction P est σ -additive c'est-à-dire que pour toute suite $(A_n)_{n \in \mathbb{N}}$ d'événements deux à deux disjoints,

$$P\left(\bigcup_{n=0}^{+\infty} A_n\right) = \sum_{n=0}^{+\infty} P(A_n)$$

Remarques :

1. Dans la définition on sous-entend que la série définissant le terme de droite converge.
2. Soit I un ensemble dénombrable et $(A_i)_{i \in I}$ une suite d'événements deux à deux disjoints. On a encore

$$P\left(\bigcup_{i \in I} A_i\right) = \sum_{i \in I} P(A_i)$$

Il suffit de prendre une bijection entre I et $\mathbb{N}$.

Définition 7.4

On appelle espace probabilisé un triplet $(\Omega, \mathcal{A}, P)$ où $\mathcal{A}$ est une tribu sur Ω et P une probabilité sur $(\Omega, \mathcal{A})$.

Proposition 7.5

Soit $(\Omega, \mathcal{A}, P)$ un espace probabilisé

1. On a $P(\emptyset) = 0$.
2. Pour toute famille $(A_1, \dots, A_n)$ d'événements deux à deux disjoints, $P(\bigcup_{i=1}^n A_i) = \sum_{i=1}^n P(A_i)$. En particulier si A et B sont disjoints, $P(A \cup B) = P(A) + P(B)$.
3. Pour tout événement A , $P(\overline{A}) = 1 - P(A)$.
4. Pour tous événements A et B , $P(A \setminus B) = P(A) - P(A \cap B)$.
5. Pour tous événements A et B , $P(A \cup B) = P(A) + P(B) - P(A \cap B)$.
6. Pour tous événements A et B , si $A \subset B$ alors $P(A) \leq P(B)$.
En particulier P est à valeurs dans $[0, 1]$.

Démonstration :

1. Il suffit de prendre $A_n = \emptyset$ pour tout entier n . La série $\sum P(\emptyset)$ converge donc $P(\emptyset) = 0$.
2. Il suffit d'utiliser le deuxième axiome en posant $A_i = \emptyset$ pour $i > n$.

3. On utilise que $A \cap \bar{A} = \emptyset$ et $A \cup \bar{A} = \Omega$ donc $P(A) + P(\bar{A}) = P(A \cup \bar{A}) = P(\Omega) = 1$.
4. On utilise que $A = (A \cap B) \cup (A \setminus B)$.
5. On utilise que $A \cup B = A \cup (B \setminus A)$ et $A \cap (B \setminus A) = \emptyset$.
6. On utilise que si $A \subset B$ alors $B = A \cup (B \setminus A)$

□

Remarque : Dans le cas où Ω est fini et $\mathcal{A} = \mathcal{P}(\Omega)$, on retrouve la définition de probabilité vue en sup où l'axiome ii) est remplacé par

Si A et B sont disjoints alors $P(A \cup B) = P(A) + P(B)$.

Cette condition, implique en effet que si $A = \emptyset$ alors $P(A) = 0$ car $P(A) = 2P(A)$. Ensuite, par récurrence, si $(A_1, \dots, A_n)$ est une famille finie d'ensembles deux à deux disjoints alors $P(\bigcup_{i=1}^n A_i) = \sum_{i=1}^n P(A_i)$.

Finalement, si $(A_n)_{n \in \mathbb{N}}$ est une suite d'événements deux à deux disjoints, comme il n'y en a qu'un nombre fini qui n'est pas vide car $\mathcal{P}(\Omega)$ est fini alors il existe $N \in \mathbb{N}$ tel que pour $n \geq N$, $A_n = \emptyset$. On en déduit que

$$P\left(\bigcup_{n=0}^{+\infty} A_n\right) = P\left(\bigcup_{n=0}^N A_n\right) = \sum_{n=0}^N P(A_n) = \sum_{n=0}^{+\infty} P(A_n)$$

Proposition 7.6

On suppose que Ω est fini ou dénombrable et $\mathcal{A} = \mathcal{P}(\Omega)$.

1. Si P est une probabilité, on pose pour tout $\omega \in \Omega$, $p_\omega = P(\{\omega\})$. La famille $(p_\omega)_{\omega \in \Omega}$ est sommable et de somme 1.
2. Réciproquement si $(p_\omega)_{\omega \in \Omega}$ est une famille sommable et de somme 1, il existe une unique probabilité P telle que pour tout $\omega \in \Omega$, $p_\omega = P(\{\omega\})$. Dans ce cas,

$$\forall A \in \mathcal{P}(\Omega), P(A) = \sum_{\omega \in A} p_\omega.$$

Démonstration : Le cas fini a déjà été vu l'année dernière. On suppose donc que Ω est infini dénombrable. On pose φ une bijection de $\mathbb{N}$ dans Ω .

1. On pose pour tout $n \in \mathbb{N}$, $A_n = \{\varphi(n)\}$. La famille $(A_n)_{n \in \mathbb{N}}$ est bien composée d'événements deux à deux disjoints. On alors

$$1 = P(\Omega) = P\left(\bigcup_{n=0}^{+\infty} A_n\right) = \sum_{n=0}^{+\infty} P(A_n) = \sum_{n=0}^{+\infty} p_{\varphi(n)}.$$

On en déduit que la famille $(p_{\varphi(n)})_{n \in \mathbb{N}}$ est sommable de somme 1 donc $(p_\omega)_{\omega \in \Omega}$ est sommable et de somme 1.

2. On se donne une famille sommable $(p_\omega)_{\omega \in \Omega}$ de somme 1. Soit P la fonction définie par

$$P : A \mapsto \sum_{\omega \in A} p_\omega$$

On a bien que $P(\Omega) = 1$ et que P est à valeurs dans $[0, 1]$. Maintenant, si on se donne une suite $(A_n)_{n \in \mathbb{N}}$ de parties deux à deux disjointes de Ω . On pose $B = \bigcup_{n \in \mathbb{N}} A_n$. La famille $(p_\omega)_{\omega \in B}$ est

sommable et, par sommation par paquets (car les A_n forment une partition de B)

$$\mathbf{P}(B) = \sum_{\omega \in B} p_{\omega} = \sum_{n=0}^{+\infty} \sum_{\omega \in A_n} p_{\omega} = \sum_{n=0}^{+\infty} \mathbf{P}(A_n).$$

□

Exemple : Cherchons une probabilité sur $\mathbb{N}$ telle que pour tout i , $p_{i+1} = \mathbf{P}(\{i+1\}) = \frac{1}{5}p_i$. On en déduit que $p_i = \frac{\lambda}{5^i}$ et comme $\sum_{i=0}^{+\infty} \frac{1}{5^i} = \frac{5}{4}$ on prend $\lambda = \frac{4}{5}$. Finalement, $p_i = \frac{4}{5^{i+1}}$.

ATTENTION

Ce résultat n'est plus vrai dans le cas où Ω n'est pas dénombrable. Or, ce sont des situations que l'on veut étudier. Citons par exemple :

- Une expérience où on fait une suite infinie de lancers à pile ou face.
- On tire un nombre réel au hasard entre 0 et 1 avec une probabilité uniforme.

2 Propriétés élémentaires des probabilités

Dans ce paragraphe on se fixe un espace probabilisé $(\Omega, \mathcal{A}, P)$.

2.1 Continuité

On veut énoncer des théorèmes pour dire que si des événements A_n « tendent » vers un élément A alors $\mathbf{P}(A_n)$ va tendre vers $\mathbf{P}(A)$.

Proposition 7.7 (Continuité croissante)

Soit $(A_n)_{n \in \mathbb{N}}$ une suite croissante (pour l'inclusion) d'événements de $\mathcal{A}$. On a

$$\mathbf{P}(A_n) \xrightarrow{n \rightarrow \infty} \mathbf{P}\left(\bigcup_{n=0}^{+\infty} A_n\right)$$

Remarque : C'est bien un résultat de continuité car, la suite étant croissante, les événements A_n tendent vers $\bigcup_{n=0}^{+\infty} A_n$.

Démonstration : On pose $B_0 = A_0$ et pour $n \geq 1$, $B_n = A_n \setminus A_{n-1}$. De ce fait les B_i sont deux à deux disjoints car pour $i < j$ alors $B_i \subset A_i$ et $A_i \cap (B_j \setminus A_j) = \emptyset$ car $A_i \subset A_j$ (la suite est croissante).

De plus, pour tout entier n ,

$$\bigcup_{i=0}^n B_i = \bigcup_{i=0}^n A_i.$$

En effet si $\omega \in \bigcup_{i=0}^n B_i$ et soit $i_0 \in \llbracket 0; n \rrbracket$ tel que $\omega \in B_{i_0}$ alors $\omega \in A_{i_0} \in \bigcup_{i=0}^n A_i$. Réciproquement, pour $\omega \in \bigcup_{i=0}^n A_i$, l'ensemble des entiers $i \in \llbracket 0; n \rrbracket$ tel que $\omega \in A_i$ n'est pas vide. On pose alors $i_0 \in \llbracket 0; n \rrbracket$ le plus petit entier tel que $\omega \in A_{i_0}$ alors $\omega \in A_{i_0} \setminus A_{i_0-1} = B_{i_0} \in \bigcup_{i=0}^n B_i$.

Maintenant, pour $i \geq 1$, $P(B_i) = P(A_i) - P(A_{i-1})$. De ce fait

$$\sum_{i=0}^n P(B_i) = P(A_0) + \sum_{i=1}^n P(A_i) - P(A_{i-1}) = P(A_n)$$

On en déduit que

$$\lim_{n \rightarrow \infty} P(A_n) = \sum_{i=0}^{+\infty} P(B_i) = P\left(\bigcup_{i=0}^{+\infty} B_i\right) = P\left(\bigcup_{n=0}^{+\infty} A_n\right).$$

□

Exemple : On réalise une suite infinie de lancers à pile ou face. On note pour $n \geq 1$, $A_n = \ll \text{il y a au moins un pile lors des } n \text{ premiers lancers} \gg$. On a de manière évidente que $A_n \subset A_{n+1}$ et donc

$$\lim_{n \rightarrow \infty} P(A_n) = P\left(\bigcup_{n=0}^{+\infty} A_n\right)$$

Il y a un résultat analogue pour les suites décroissantes

Proposition 7.8 (Continuité décroissante)

Soit $(A_n)_{n \in \mathbb{N}}$ une suite décroissante (pour l'inclusion) d'événements de $\mathcal{A}$. On a

$$P(A_n) \xrightarrow{n \rightarrow \infty} P\left(\bigcap_{n=0}^{+\infty} A_n\right)$$

Remarque : Là encore, c'est bien un résultat de continuité car, la suite étant décroissante, les événements A_n tendent vers $\bigcap_{n=0}^{+\infty} A_n$.

Démonstration : Il suffit de passer au complémentaire dans la proposition précédente. Précisément si on pose $B_n = \overline{A_n}$. La suite (B_n) est donc croissante. De ce fait

$$P(B_n) \xrightarrow{n \rightarrow \infty} P\left(\bigcup_{n=0}^{+\infty} B_n\right).$$

Maintenant, on a vu que $P(B_n) = P(\overline{A_n}) = 1 - P(A_n)$. De même

$$P\left(\bigcup_{n=0}^{+\infty} B_n\right) = P\left(\bigcup_{n=0}^{+\infty} \overline{A_n}\right) = P\left(\overline{\bigcap_{n=0}^{+\infty} A_n}\right) = 1 - P\left(\bigcap_{n=0}^{+\infty} A_n\right).$$

□

Exemple : On lance une infinité de fois un dé à 6 faces équilibrés. On note $A_n = \ll \text{on n'a pas obtenu de 6 lors des } n \text{ premiers lancers} \gg$. On peut alors montrer que $P(A_n) = \left(\frac{5}{6}\right)^n$. La suite (A_n) est décroissante et donc

$$P\left(\bigcap_{n=0}^{+\infty} A_n\right) = \lim_{n \rightarrow \infty} P(A_n) = 0$$

Corollaire 7.9

Soit $(A_n)_{n \in \mathbb{N}}$ une suite d'événements (non nécessairement monotone)

$$\lim_{n \rightarrow +\infty} \mathbf{P} \left(\bigcup_{k=0}^n A_k \right) = \mathbf{P} \left(\bigcup_{k=0}^{\infty} A_k \right) \quad \text{et} \quad \lim_{n \rightarrow +\infty} \mathbf{P} \left(\bigcap_{k=0}^n A_k \right) = \mathbf{P} \left(\bigcap_{k=0}^{\infty} A_k \right)$$

Démonstration : On pose pour tout entier n , $B_n = \bigcup_{k=0}^n A_k$. La suite (B_n) est croissante d'où, par continuité croissante,

$$\lim_{n \rightarrow +\infty} \mathbf{P} \left(\bigcup_{k=0}^n A_k \right) = \lim_{n \rightarrow +\infty} \mathbf{P}(B_n) = \mathbf{P} \left(\bigcup_{n=0}^{\infty} B_n \right) = \mathbf{P} \left(\bigcup_{n=0}^{\infty} A_n \right)$$

De même, on pose pour tout entier n , $C_n = \bigcap_{k=0}^n A_k$. La suite (C_n) est décroissante d'où, par continuité décroissante,

$$\lim_{n \rightarrow +\infty} \mathbf{P} \left(\bigcap_{k=0}^n A_k \right) = \lim_{n \rightarrow +\infty} \mathbf{P}(C_n) = \mathbf{P} \left(\bigcap_{n=0}^{\infty} C_n \right) = \mathbf{P} \left(\bigcap_{n=0}^{\infty} A_n \right)$$

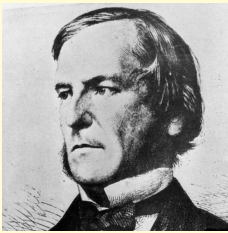
□

Proposition 7.10 (Inégalité de Boole - Sous-additivité de P)

Soit (A_n) une suite d'événements,

$$\mathbf{P} \left(\bigcup_{n=0}^{+\infty} A_n \right) \leq \sum_{n=0}^{+\infty} \mathbf{P}(A_n)$$

Remarque : Cet énoncé reste vrai dans le cas où la série de droite diverge (vers $+\infty$).

Matheux (Georges Boole : 1815-1864)

George Boole, né le 2 novembre 1815 à Lincoln (Royaume-Uni) et mort le 8 décembre 1864 à Ballintemple (Irlande), est un logicien, mathématicien et philosophe britannique. Il est le créateur de la logique moderne, fondée sur une structure algébrique et sémantique, que l'on appelle algèbre de Boole en son honneur. Il a aussi travaillé dans d'autres domaines mathématiques, des équations différentielles aux probabilités en passant par l'analyse.

Démonstration : Soit $(A_i)_{i \in \mathbb{N}}$ des événements, on considère les événements deux à deux disjoints $C_i = A_i \setminus \left(\bigcup_{j < i} A_j \right)$.

On a alors pour tout entier N ,

$$\mathbf{P} \left(\bigcup_{i=0}^N A_i \right) = \mathbf{P} \left(\bigcup_{i=0}^N C_i \right) = \sum_{i=0}^N \mathbf{P}(C_i) \leq \sum_{i=0}^N \mathbf{P}(A_i)$$

car pour tout $i \in \llbracket 0, N \rrbracket$, $C_i \subset A_i$ et donc $P(C_i) \leq P(A_i)$.

Donc

$$P\left(\bigcup_{n=0}^{\infty} A_n\right) = \lim_{N \rightarrow +\infty} P\left(\bigcup_{i=0}^N A_i\right) \leq \lim_{N \rightarrow +\infty} \sum_{i=0}^N P(A_i) = \sum_{i=0}^{\infty} P(A_i)$$

□

2.2 Événements négligeables, presque sûrs

Définition 7.11

Soit A un événement.

1. Il est dit *négligeable* si $P(A) = 0$.
2. Il est dit *presque sûr* si $P(A) = 1$.

ATTENTION

Un événement négligeable n'est pas nécessairement impossible. Si on reprend la suite infinie de lancers de dés à 6 faces et que l'on note toujours A_n « on n'a pas obtenu de 6 lors des n premiers lancers ». L'événement $\bigcup_{n=0}^{+\infty} A_n$ est négligeable mais pas pour autant impossible. Il y a en effet une infinité de possibilités qui réalise cet événement. De même un événement presque sûr n'est pas certain.

Proposition 7.12

Une réunion (resp. une intersection) finie ou dénombrable d'événements négligeables (resp. presque sûrs) est négligeable (resp. presque sûre).

Démonstration : Il suffit d'utiliser l'inégalité de Boole. Soit $(A_n)_{n \in \mathbb{N}}$ une suite d'événements négligeables (on peut supposer que la famille est indexée par $\mathbb{N}$ quitte à utiliser une bijection et à compléter par l'ensemble vide s'il n'y avait qu'un nombre fini d'événements). On a alors

$$P\left(\bigcup_{n=0}^{+\infty} A_n\right) \leq \sum_{n=0}^{+\infty} P(A_n) = 0$$

Donc $\bigcup_{n=0}^{+\infty} A_n$ est bien négligeable.

Le résultat sur l'intersection d'événements presque sûrs s'obtient alors par passage au complémentaire.

□

Remarque : Dans l'énoncé ci-dessus, il est essentiel que la réunion soit au plus dénombrable. Si on considère une suite infinie de lancers à pile ou face. L'univers $\Omega = \{P, F\}^{\mathbb{N}}$. Pour chaque ω dans Ω , $P(\{\omega\}) = 0$ donc $\{\omega\}$ est négligeable mais

$$\Omega = \bigcup_{\omega \in \Omega} \{\omega\}$$

n'est pas négligeable.

Définition 7.13 (Système quasi-complet d'événements)

Soit $(\Omega, \mathcal{A}, P)$ un espace probabilisé. On appelle système quasi-complet d'événements une famille $(A_i)_{i \in I} \in \mathcal{A}^I$ d'éléments de $\mathcal{A}$ telle que :

- Les A_i recouvrent quasiment Ω : l'événement $\bigcup_{i \in I} A_i$ est presque sûr.
- Les A_i sont deux à deux disjoints : $\forall (i, j) \in I^2, i \neq j \Rightarrow A_i \cap A_j = \emptyset$.

Remarque : Soit $(A_i)_{i \in I}$ un système quasi-complet d'événements. Pour tout $p \in \mathbb{N}^*$ l'ensemble $I_p = \{i \in I, P(A_i) > \frac{1}{p}\}$ contient moins de p éléments. On en déduit que $J = \{i \in I, P(A_i) \neq 0\}$ est fini ou dénombrable comme union dénombrable d'ensemble finis. On peut donc considérer $(A_i)_{i \in J}$ et supposer donc que J est fini ou dénombrable.

3 Indépendance et probabilités conditionnelles

3.1 Probabilités conditionnelles

Définition 7.14 (Probabilité conditionnelle)

Soit $(\Omega, \mathcal{A}, P)$ un espace probabilisé. Soit $(A, B) \in \mathcal{A}^2$ avec $P(B) \neq 0$. On appelle probabilité conditionnelle de A sachant B et on note $P(A|B)$ ou $P_B(A)$ la probabilité

$$P(A|B) = P_B(A) = \frac{P(A \cap B)}{P(B)}.$$

Remarque : On a donc $P(A \cap B) = P(B)P(A|B)$. Cette forme à « presque » du sens quand $P(B) = 0$.

Proposition 7.15

Avec les notations précédentes, P_B est une probabilité sur $(\Omega, \mathcal{A})$.

Démonstration : Comme $A \cap B \subset B$, $P(A \cap B) \leq P(B)$ et donc $P_B(A) \in [0, 1]$. De plus, $P_B(\Omega) = \frac{P(A \cap \Omega)}{P(A)} = 1$.

Pour finir, si $(A_n)_{n \in \mathbb{N}}$ est suite d'éléments deux à deux disjoints,

$$P_B\left(\bigcup_{n=0}^{+\infty} A_n\right) = \frac{P\left(B \cap \bigcup_{n=0}^{+\infty} A_n\right)}{P(B)} = \frac{P\left(\bigcup_{n=0}^{+\infty} (B \cap A_n)\right)}{P(B)} = \sum_{n=0}^{+\infty} \frac{P(B \cap A_n)}{P(B)} = \sum_{n=0}^{+\infty} P_B(A_n).$$

On a bien montré que P_B est une probabilité. □

Exercice : Montrer que si A et A' sont deux événements tels que $P(A \cap A') \neq 0$, $(P_A)_{A'} = P_{A \cap A'}$.

Exemple : On lance deux dés à 6 faces équilibrés. On note A l'événement « la somme des deux dés est égale à 8 ». On sait que

$$P(A) = \frac{5}{36}$$

car il y a 5 couples qui font 8 ((2, 6); (3, 5); (4, 4); (5, 3); (6, 2)) parmi les 36 couples possibles (équiprobables). Maintenant, si on lance d'abord le premier dé et que l'on note B l'événement « on a fait 3 ». La probabilité conditionnelle est égale à $\frac{1}{6}$ qui est la probabilité de faire 5 sur le deuxième dé. On a

$$P_B(A) = \frac{1}{6} = \frac{P(A \cap B)}{P(B)} = \frac{\frac{1}{36}}{\frac{1}{6}}.$$

Proposition 7.16 (Formule des probabilités composées)

Soit $(\Omega, \mathcal{A}, P)$ un espace probabilisé et $A_1, \dots, A_n$ des événements tels que $P(A_1 \cap \dots \cap A_{n-1}) \neq 0$. On a

$$P(A_1 \cap \dots \cap A_n) = P(A_1) \times P(A_2|A_1) \times P(A_3|A_1 \cap A_2) \times \dots \times P(A_n|A_1 \cap \dots \cap A_{n-1}).$$

Remarques :

1. Le fait que $P(A_1 \cap \dots \cap A_{n-1})$ soit non nul nous assure que toutes les probabilités conditionnelles sont bien définies.
2. C'est cohérent avec l'intuition : le fait que $A_1, \dots, A_n$ soit réalisé revient à ce que A_1 soit réalisé, A_2 soit réalisé sachant que A_1 l'est, A_3 soit réalisé sachant que A_1 et A_2 le sont, $\dots$, A_n soit réalisé sachant que $A_1, \dots, A_{n-1}$ le sont.

Démonstration :

Le terme de droite vaut

$$P(A_1) \times \frac{P(A_1 \cap A_2)}{P(A_1)} \times \frac{P(A_1 \cap A_2 \cap A_3)}{P(A_1 \cap A_2)} \times \dots \times \frac{P(A_1 \cap \dots \cap A_n)}{P(A_1 \cap \dots \cap A_{n-1})} = P(A_1 \cap \dots \cap A_n).$$

□

Exemple : On considère une urne avec a boules blanches et b boules noires ($a \geq 4$). On suppose que toutes les boules sont semblables. On tire trois boules successivement et sans remise dans l'urne et on note A' l'événement « les trois boules sont blanches ». On note, pour $i \in \{1, 2\}$, B_i l'événement « la i ème boule tirée est blanche ». On a donc $A' = B_1 \cap B_2 \cap B_3$ et

$$P(A') = P(B_1) \cdot P(B_2|B_1) \cdot P(B_3|B_1 \cap B_2).$$

D'où

$$P(A') = \frac{a(a-1)(a-2)}{(a+b)(a+b-1)(a+b-2)}.$$

Proposition 7.17 (Formule des probabilités totales)

Soit $(\Omega, \mathcal{A}, P)$ un espace probabilisé et $(A_i)_{i \in I}$ un système (quasi)-complet d'événements. Pour tout événement B on a

$$\begin{aligned} P(B) &= \sum_{i \in I} P(B \cap A_i) \\ &= \sum_{i \in I} P(A_i) \cdot P(B|A_i) \text{ si aucun des } A_i \text{ n'est négligeable} \end{aligned}$$

Remarques :

1. En théorie, on ne peut appliquer la deuxième partie de la formule que si les événements A_i ont des probabilités non nulles afin de pouvoir déterminer $P(B|A_i)$. Cependant, on voit dans la formule que si pour un (ou plusieurs) i on a $P(A_i) = 0$, certes le terme $P(B|A_i)$ n'est pas défini, mais il est multiplié par 0. On a en effet $P(A_i \cap B) = 0 = 0 \cdot P(B|A_i)$.
2. Cette formule est utile quand on a un système complet d'événements qui permet de "découper" l'univers. Pour connaître la probabilités d'un événement il s'agit de connaître les probabilités de l'événement en supposant être dans l'un cas.

Démonstration : Notons $\Lambda = \bigcup_{i \in I} A_i$. Comme le système est quasi-complet, $P(\bar{\Lambda}) = 0$.

On sait que $B = \left(\bigcup_{i \in I} (A_i \cap B) \right) \cup (\bar{\Lambda} \cap B)$. De plus les événements $(A_i \cap B)$ sont deux à deux disjoints et il sont aussi disjoints de $\bar{\Lambda} \cap B$. On a donc

$$P(B) = \sum_{i \in I} P(B \cap A_i) + P(B \cap \bar{\Lambda}) = \sum_{i \in I} P(B \cap A_i)$$

En effet $P(B \cap \bar{\Lambda}) \leq P(\bar{\Lambda}) = 0$.

□

Exemple : Toujours le même exemple. On considère une urne avec a boules blanches et b boules noires ($a \geq 4$). On suppose que toutes les boules sont semblables. On tire trois boules successivement et sans remise dans l'urne. On note, pour $i \in \{1, 2\}$, B_i l'événement « la i ème boule tirée est blanche ». On cherche à calculer $P(B_2)$. On considère le système complet d'événement $(B_1, \bar{B}_1)$. On a

$$\begin{aligned} P(B_2) &= P(B_1 \cap B_2) + P(\bar{B}_1 \cap B_2) \\ &= P(B_1) \cdot P(B_2|B_1) + P(\bar{B}_1) \cdot P(B_2|\bar{B}_1) \\ &= \frac{a}{a+b} \cdot \frac{a-1}{a+b-1} + \frac{b}{a+b} \cdot \frac{a}{a+b-1} \\ &= \frac{a}{a+b} \end{aligned}$$

On peut présenter cela sous forme d'un arbre.

Proposition 7.18 (Formule de Bayes)

Soit $(\Omega, \mathcal{A}, P)$ un espace probablisé et $(A_i)_{i \in I}$ un système complet d'événements. Soit B un événement, pour tout $j \in I$, on a

$$P_B(A_j) = \frac{P(A_j \cap B)}{P(B)} = \frac{P(A_j) \cdot P(B|A_j)}{\sum_{i \in I} P(A_i) \cdot P(B|A_i)}.$$

Remarques :

1. Cette formule est, dans un sens, la formule des probabilités totales "à l'envers". En effet, si on suppose que l'on sait calculer $P(A_i)$ et $P(B|A_i)$, elle permet de savoir avec quelle probabilité on est dans l'événement A_i en sachant que B est réalisé.
2. On utilise souvent le terme le plus à droite, mais il est plus simple de "retrouver" la formule que d'essayer de s'en souvenir.

Matheux (Thomas Bayes : 1702-1761)

Thomas Bayes (né env. en 1702 à Londres - mort le 7 avril 1761 à Tunbridge Wells, dans le Kent) est un mathématicien britannique et pasteur de l'Église presbytérienne, connu pour avoir formulé le théorème de Bayes.

Exemple : Toujours le même exemple. On considère une urne avec a boules blanches et b boules noires ($a \geq 4$). On suppose que toutes les boules sont semblables. On tire trois boules successivement et sans remise dans l'urne. On note pour $i \in \{1, 2\}$, B_i l'événement « la i ème boule tirée est blanche ». On suppose que la deuxième boule tirée est blanche et on veut savoir quelle est la probabilité que la première le fut. C'est à dire on veut calculer $P(B_1|B_2)$. La formule de Bayes donne alors :

$$P(B_1|B_2) = \frac{P(B_1).P(B_2|B_1)}{P(B_1).P(B_2|B_1) + P(\overline{B_1}).P(B_2|\overline{B_1})} = \frac{a(a-1)}{a(a-1) + a.b} = \frac{a-1}{a+b-1}.$$

Exercice : Faire un exercice sur les chaines de Markov : Mines PC 2017 ; exo 5055

3.2 Événements indépendants

Définition 7.19 (Événements indépendants)

Soit $(\Omega, \mathcal{A}, P)$ un espace probabilisé. Soit A, B deux événements de $\mathcal{A}$. On dit qu'ils sont indépendants si

$$P(A \cap B) = P(A)P(B)$$

On note $A \perp B$.

Remarque : Cela exprime le fait que la réalisation de l'un des événement n'influe sur la probabilité que l'autre se réalise. A savoir, si A n'est pas négligeable

$$P(B|A) = P(B).$$

Exemple : On dispose d'une pièce que l'on sait équilibrée. On la lance deux fois. Si on note P_i l'événement on fait pile au lancer i alors P_1 et P_2 sont indépendants. On a $P(P_1) = P(P_2) = \frac{1}{2}$ et $P(P_1 \cap P_2) = \frac{1}{4}$.

Maintenant, si on ne suppose plus que l'on sait que la pièce est équilibrée, les événements ne sont plus indépendant car le fait que l'on ait fait pile au premier lancer va impliquer que l'on a plus de chance d'avoir pris une pièce qui fait souvent pile. Supposons que l'on dispose de $2n + 1$ pièces et que la pièce $i \in \llbracket -n ; n \rrbracket$ ait une probabilité $p_i = \frac{1}{2} + \frac{i}{4n} = \frac{2n+i}{4n}$ de faire pile.

On suppose que l'on tire une pièce au hasard et on note A_i l'événement « on a pris la pièce i ». Les événements $(A_i)_{i \in \llbracket -n ; n \rrbracket}$ forment un système complet d'événements.

On a $\forall i \in \llbracket -n ; n \rrbracket, P(A_i) = \frac{1}{2n+1}$.

De ce fait, par la formule des probabilités totales

$$P(P_1) = P(P_2) = \sum_{i=-n}^n P(A_i)P(P_1|A_i) = \frac{1}{2n+1} \sum_{i=-n}^n \left(\frac{1}{2} + \frac{i}{4n} \right) = \frac{1}{2} + 0 = \frac{1}{2}.$$

On veut maintenant calculer $P(P_2|P_1) = \frac{P(P_1 \cap P_2)}{P(P_1)}$. Comme ci-dessus,

$$P(P_1 \cap P_2) = \sum_{i=-n}^n P(A_i)P(P_1 \cap P_2|A_i) = \frac{1}{2n+1} \sum_{i=-n}^n \left(\frac{1}{2} + \frac{i}{4n} \right)^2 = \frac{1}{4} + \frac{n+1}{48n}.$$

On en déduit donc

$$P(P_2|P_1) = \frac{P(P_1 \cap P_2)}{P(P_1)} = \frac{13n+1}{24n} > \frac{1}{2}.$$

Exercice : En reprenant la situation de l'exemple précédent, calculer $P_{P_1}(A_i)$ et vérifier que $\sum_{i=-n}^n P_{P_1}(A_i) = 1$.

Maintenant, si on sait que le premier lancer a donné une pile (on suppose P_1). On calcule par la formule de Bayes que

$$\forall i \in \llbracket -n; n \rrbracket, P_{P_1}(A_i) = \frac{P(P_1 \cap A_i)}{P(P_1)} = 2P(A_i)P(P_1|A_i) = 2 \times \frac{1}{2n+1} \times \frac{2n+i}{4n} = \frac{2n+i}{2n(2n+1)}.$$

On vérifie bien évidemment que $\sum_{i=-n}^n P_{P_1}(A_i) = 1$.

ATTENTION

La notion d'indépendance **dépend** de la probabilité.

Définition 7.20 (Famille d'événements indépendants)

Soit $(\Omega, \mathcal{A}, P)$ un espace probabilisé et $(A_i)_{i \in I}$ une famille d'événements. Ces événements sont indépendants si pour toute partie **finie** de I ,

$$P\left(\bigcap_{i \in J} A_i\right) = \prod_{i \in J} P(A_i).$$

ATTENTION

Cette condition est plus forte que juste 2 à 2 indépendants.

Par exemple si on lance deux dés à 6 faces et que l'on note :

- A : « le premier dé fait un résultat pair »
- B : « le deuxième dé fait un résultat impair »
- C : « la somme est paire »

On voit que $P(A) = P(B) = P(C) = \frac{1}{2}$ puis que $P(A \cap B) = P(A \cap C) = P(B \cap C) = \frac{1}{4}$ par contre $P(A \cap B \cap C) = 0$.

Remarque : Dans la majorité des cas, l'indépendance sera une hypothèse de l'énoncé.

Proposition 7.21 (Indépendance et complémentaire)

Soit $(\Omega, \mathcal{A}, P)$ un espace probabilisé

1. Soit A, B deux événements indépendants alors A et $\bar{B}$ sont indépendants.
2. Soit $(A_i)_{i \in I}$ une famille d'événements indépendants. Pour tout J, J' des parties finies et disjointes de I ,

$$P\left(\left(\bigcap_{i \in J} A_i\right) \cap \left(\bigcap_{i \in J'} \bar{A}_i\right)\right) = \prod_{i \in J} P(A_i) \times \prod_{i \in J'} P(\bar{A}_i)$$

Démonstration :

1. On sait que $A = (A \cap B) \cup (A \cap \bar{B})$ d'où $A \cap \bar{B} = A \setminus (A \cap B)$. On en déduit que

$$P(A \cap \bar{B}) = P(A) - P(A \cap B) = P(A) - P(A)P(B) = P(A)(1 - P(B)) = P(A)P(\bar{B})$$

2. On va procéder par récurrence sur le cardinal de J' .

- **I** : si $J' = \emptyset$ la propriété est vraie car les événements sont indépendants.
- **H** : soit $r \in \mathbb{N}$ on suppose la propriété vraie pour tout ensemble fini J et tout ensemble fini J' tel que J' est de cardinal r avec $J \cap J' = \emptyset$.

On considère J, J' deux ensembles finis avec J' de cardinal $r + 1$ avec $J \cap J' = \emptyset$. En particulier, J' n'est pas vide. On note alors i_0 un élément de J' . Posons alors $K = J \cup \{i_0\}$ et $K' = J' \setminus \{i_0\}$. On a encore K, K' fini, $K \cap K' = \emptyset$ et K' est maintenant de cardinal r .

$$\left(\bigcap_{i \in J} A_i\right) \cap \left(\bigcap_{i \in J'} \bar{A}_i\right) = \left(\left(\bigcap_{i \in J} A_i\right) \cap \left(\bigcap_{i \in K'} \bar{A}_i\right)\right) \setminus \left(\left(\bigcap_{i \in K} A_i\right) \cap \left(\bigcap_{i \in K'} \bar{A}_i\right)\right)$$

D'où, par récurrence, en utilisant que

$$\left(\bigcap_{i \in K} A_i\right) \cap \left(\bigcap_{i \in K'} \bar{A}_i\right) \subset \left(\bigcap_{i \in J} A_i\right) \cap \left(\bigcap_{i \in K'} \bar{A}_i\right)$$

$$\begin{aligned} P\left(\left(\bigcap_{i \in J} A_i\right) \cap \left(\bigcap_{i \in J'} \bar{A}_i\right)\right) &= P\left(\left(\bigcap_{i \in J} A_i\right) \cap \left(\bigcap_{i \in K'} \bar{A}_i\right)\right) - P\left(\left(\bigcap_{i \in K} A_i\right) \cap \left(\bigcap_{i \in K'} \bar{A}_i\right)\right) \\ &= \prod_{i \in J} P(A_i) \times \prod_{i \in K'} P(\bar{A}_i) - \prod_{i \in K} P(A_i) \times \prod_{i \in K'} P(\bar{A}_i) \\ &= \prod_{i \in J} P(A_i) \times \prod_{i \in K'} P(\bar{A}_i) \times (1 - P(A_{i_0})) \end{aligned}$$

□

Exercice : Faire un exercice sur Borel-Cantelli

4 Variables aléatoires discrètes

Dans ce paragraphe on se fixe un espace probabilisé $(\Omega, \mathcal{A}, P)$.

4.1 Généralités

Définition 7.22 (Variables aléatoires discrètes)

Soit E un ensemble. Une variable aléatoire discrète à valeurs dans E est une application $X : \Omega \rightarrow E$ telle que

- i) L'image $X(\Omega)$ est finie ou dénombrable.
- ii) Pour tout $x \in E$, l'ensemble $X^{-1}(\{x\})$ des antécédents de x par X appartient à $\mathcal{A}$.

Remarques :

1. La plupart du temps l'ensemble E sera $\mathbf{R}$. On parle alors de variable aléatoire discrète réelle.
2. Si $\mathcal{A} = \mathcal{P}(\Omega)$ la deuxième condition est toujours vérifiée.
3. Nous ne regarderons que des variables discrètes et essentiellement réelles. Le terme « variable aléatoire » sans autre mention signifiera « variable aléatoire discrète réelle ».
4. La condition signifie que pour toute partie B de $X(\Omega)$, $B = \bigcup_{x \in B} \{x\}$ où B est finie et dénombrable.

De ce fait,

$$X^{-1}(B) = \bigcup_{x \in B} X^{-1}(\{x\})$$

est donc un élément de $\mathcal{A}$.

ATTENTION

Il est important de comprendre que même si on se limite au cas où $X(\Omega)$ sera dénombrable, il est nécessaire d'envisager que Ω soit plus « gros » notamment non dénombrable^a. On fera tous les calculs à partir de $X(\Omega)$ et pas à partir de Ω .

^a. Si on veut modéliser une suite infinie de lancers à pile ou face et si on note X le rang du premier pile. On a bien $X(\Omega) = \mathbf{N} \cup \{+\infty\}$ mais par contre $\Omega = \{P, F\}^{\mathbf{N}}$ qui n'est pas dénombrable.

Exemples :

1. Toutes les variables aléatoires définies sur un univers fini, sont des variables aléatoires discrètes.
2. Soit A un événement. La fonction indicatrice

$$\begin{aligned} \mathbb{1}_A : \Omega &\rightarrow \mathbf{R} \\ \omega &\mapsto \begin{cases} 1 & \text{si } \omega \in A \\ 0 & \text{sinon} \end{cases} \end{aligned}$$

est une variable aléatoire discrète car $\mathbb{1}_A(\Omega) = \{0, 1\}$ et $\mathbb{1}_A^{-1}(\{1\}) = A \in \mathcal{A}$ et $\mathbb{1}_A^{-1}(\{0\}) = \bar{A} \in \mathcal{A}$

3. On allume une ampoule et on note X le temps de vie de cette ampoule. L'ensemble image $X(\Omega)$ ne sera pas dénombrable. De ce fait, X n'est pas une variable aléatoire discrète.

Notations :

1. On note $(X = x)$ ou $\{X = x\}$ pour $X^{-1}(\{x\})$ et donc $\mathbf{P}(X = x)$ pour $\mathbf{P}(X^{-1}(\{x\}))$
2. Si $E \subset \mathbf{R}$, on notera $(a \leq X \leq b)$ ou $\{a \leq X \leq b\}$ pour $X^{-1}([a, b])$ et donc $\mathbf{P}(a \leq X \leq b)$ pour $\mathbf{P}(X^{-1}([a, b]))$. De même pour $(X < a)$, $\{X < a\}$...
3. De manière plus générale, si $B \subset X(\Omega)$ on note $(X \in B)$ ou $\{X \in B\}$ pour $X^{-1}(B)$.

Définition 7.23

Soit E un ensemble et X une variable aléatoire discrète à valeurs dans E . Soit $f : E \rightarrow F$. On note $f(X)$ la variable aléatoire

$$f(X) = f \circ X : \Omega \rightarrow F$$

Remarque : Il est clair que c'est bien une variable aléatoire car

- i) L'image de Ω par $f \circ X$ est l'image de $X(\Omega)$ par f et de ce fait est fini ou dénombrable.
- ii) Soit $y \in F$, on cherche à montrer que $(f \circ X)^{-1}(\{y\}) \in \mathcal{A}$. Pour cela voit que

$$(f \circ X)^{-1}(\{y\}) = \bigcup_{x \in B} X^{-1}(\{x\})$$

où $B = \{x \in X(\Omega) \mid f(x) = y\}$. Donc B étant fini ou dénombrable, $(f \circ X)^{-1}(\{y\}) \in \mathcal{A}$.

Proposition 7.24

Soit X une variable aléatoire discrète. La famille $(X = x)_{x \in X(\Omega)}$ est un système complet d'événements.

Démonstration : Cours de première année.

Proposition-Définition 7.25 (Loi d'une variable aléatoire)

Soit X une variable aléatoire à valeurs dans E .

1. La variable aléatoire X induit une probabilité P_X sur $X(\Omega)$ définie par

$$\forall B \subset X(\Omega), P_X(B) = \mathbf{P}(X^{-1}(B)) = \mathbf{P}(X \in B)$$

2. On dit alors que P_X est la loi de X .

Remarque : Comme $X(\Omega)$ est au plus dénombrable, on lui associe la tribu $\mathcal{P}(X(\Omega))$.

Démonstration : On vérifie les axiomes des probabilités.

- Il est clair que pour $B \subset X(\Omega)$, $P_X(B) \in [0, 1]$.
- On a $P_X(\emptyset) = \mathbf{P}(X^{-1}(\emptyset)) = \mathbf{P}(\emptyset) = 0$.
- Soit $(B_i)_{i \in \mathbb{N}}$ une famille de parties de $X(\Omega)$ deux à deux disjointes, On a $X^{-1}(\bigcup_{i \in \mathbb{N}} B_i) = \bigcup_{i \in \mathbb{N}} X^{-1}(B_i)$. De ce fait

$$P_X\left(\bigcup_{i \in \mathbb{N}} B_i\right) = \mathbf{P}\left(\bigcup_{i \in \mathbb{N}} X^{-1}(B_i)\right) = \sum_{i=0}^{+\infty} \mathbf{P}(X^{-1}(B_i)) = \sum_{i=0}^{+\infty} P_X(B_i).$$

En effet les $X^{-1}(B_i)$ sont deux à deux disjoints.

On a bien montré que P_X était une probabilité. □

Remarques :

1. Pour simplifier, on peut se donner la loi d'une variable aléatoire discrète en utilisant un ensemble fini ou dénombrable T contenant $X(\Omega)$. Dans ce cas, $\forall x \in T \setminus X(\Omega), \mathbf{P}(X = x) = 0$.

2. On peut à l'inverse donner la loi de X sur un ensemble U inclus dans $X(\Omega)$ en enlevant les éléments x tels que $P(X = x) = 0$.
3. On a aussi $P_X(B) = \sum_{x \in B} P(\{x\})$ UNIQUEMENT quand l'univers Ω est dénombrable.

ATTENTION

Se donner la loi d'une variable aléatoire discrète revient donc à se donner $X(\Omega)$ qui est fini ou dénombrable et à se donner une probabilité sur $X(\Omega)$, c'est-à-dire une famille sommable $(p_x)_{x \in X(\Omega)}$ de somme 1.

Notations :

1. Soit X une variable aléatoire et $\mathcal{L}$ une probabilité sur $X(\Omega)$. Si $P_X = \mathcal{L}$ on note $X \sim \mathcal{L}$.
2. Si X et Y sont deux variables aléatoires qui ont la même loi on note $X \sim Y$. Notons que les variables X et Y peuvent être définies sur des espaces probabilisés distincts.

Exemple : La loi binomiale de paramètre n, p (notée $\mathcal{B}(n, p)$) a été définie en première année. Une variable aléatoire X suit la loi $\mathcal{B}(n, p)$ ($X \sim \mathcal{B}(n, p)$) si

- $X(\Omega) = \llbracket 0 ; n \rrbracket$
- $\forall k \in \llbracket 0 ; n \rrbracket, P(X = k) = \binom{n}{k} p^k (1-p)^{n-k}$.

ATTENTION

Si X et Y sont deux variables aléatoires sur le même espace probabilisé Ω telles que $X(\Omega) = Y(\Omega)$. Il faut faire la différence entre :

- Le fait que X et Y soient égales ($X = Y$). Ce qui signifie que pour tout ω dans Ω , $X(\omega) = Y(\omega)$.
- Le fait que X et Y soient égales presque sûrement. Ce qui signifie que l'événement $(X \neq Y)$ est négligeable ou que $P(X = Y) = 1$.
- Le fait que X et Y suivent la même loi. Par exemple, si on lance 10 pièces équilibrées, le nombre de pile obtenu et le nombre de face suivent la même loi mais elles ne sont pas égales (même presque sûrement).

Proposition 7.26

Soit X et X' deux variables aléatoires discrètes définies que $(\Omega, \mathcal{A}, P)$ et $(\Omega', \mathcal{A}', P')$. On suppose que $X(\Omega) = X'(\Omega')$ et que $X \sim X'$. Soit $f : X(\Omega) \rightarrow F$. Les variables aléatoires discrètes $f(X)$ et $f(X')$ ont la même loi.

Démonstration : Notons $Y = f(X)$ et $Y' = f(X')$. On voit que $Y(\Omega) = f(X(\Omega)) = f(X'(\Omega')) = Y'(\Omega')$. De plus pour $y \in Y(\Omega)$,

$$P(Y = y) = P\left(\bigcup_{\substack{x \in X(\Omega) \\ f(x)=y}} (X = x)\right) = \sum_{\substack{x \in X(\Omega) \\ f(x)=y}} P(X = x) = \sum_{\substack{x \in X(\Omega) \\ f(x)=y}} P(X' = x) = P(Y' = y)$$

On en déduit que $Y \sim Y'$.

□

Définition 7.27 (Loi conditionnelle)

Soit X une variable aléatoire discrète sur un espace probabilisé $(\Omega, \mathcal{A}, P)$. Soit A un événement de probabilité non nulle. On appelle loi conditionnelle de X sachant A la probabilité sur $(X(\Omega), \mathcal{P}(X(\Omega)))$ définie par

$$\forall B \subset X(\Omega), P_{X|A}(B) = P_A(X^{-1}(B)) = P_A(X \in B)$$

Remarque : Il n'y a pas de notations dans le programme officielle pour cette loi conditionnelle.

4.2 Couples de variables aléatoires et vecteurs aléatoires

Proposition-Définition 7.28 (Loi conjointe)

Soit X et Y deux variables aléatoires discrètes à valeurs dans E et F respectivement. Le couple

$$(X, Y) : \Omega \rightarrow E \times F$$

est une variable aléatoire. Sa loi $P_{X \times Y}$ s'appelle la loi conjointe de du couple (X, Y) .

Démonstration : Il faut montrer que (X, Y) est une variable aléatoire.

- On remarque que $(X, Y)(\Omega) \subset X(\Omega) \times Y(\Omega)$. On en déduit que c'est un ensemble fini ou dénombrable.
- Soit $(x, y) \in X(\Omega) \times Y(\Omega)$, $(X, Y)^{-1}(\{(x, y)\}) = X^{-1}(\{x\}) \cap Y^{-1}(\{y\}) \in \mathcal{A}$.

On a bien montré que (X, Y) était une variable aléatoire. □

Remarques :

1. Quand on voudra décrire un couple (X, Y) , on se contentera souvent de $X(\Omega) \times Y(\Omega)$ à la place de $(X, Y)(\Omega)$ qui peut être plus compliqué à calculer.
2. Se donner la loi conjointe revient donc à se donner, pour tout $x \in X(\Omega)$ et tout $y \in Y(\Omega)$ la probabilité

$$P_{X,Y}(\{(x, y)\}) = P((X = x) \cap (Y = y)) = P((X, Y) = (x, y)).$$

3. On peut utiliser cela pour montrer que si X et Y sont deux variables aléatoires discrètes alors $X + Y$, $\min(X, Y)$ ou $\max(X, Y)$ sont des variables aléatoires discrètes en écrivant par exemple $X + Y = S(X, Y)$ où $S : (x, y) \mapsto x + y$.

ATTENTION

Un couple de variables aléatoires (quand on le considère avec la loi conjointe) est donc une variable aléatoire. On peut de ce fait lui appliquer tout ce qui a été (et qui sera) vu dans le cadre des variables aléatoires.

Exemple : On lance deux fois un dé équilibré. On note X le nombre de chiffres pairs obtenus et Y le nombre de chiffres supérieurs à 4 obtenus. On sait que $X \hookrightarrow \mathcal{B}(2, \frac{1}{2})$ et $Y \hookrightarrow \mathcal{B}(2, \frac{1}{2})$. On en déduit que $X(\Omega) \times Y(\Omega) = \llbracket 0; 2 \rrbracket^2$. Cependant le fait d'obtenir un nombre plus grand que 4 n'est pas indépendant du fait de faire un nombre pair. On va donc calculer "à la main" les nombres

$$p_{i,j} = P((X, Y) = (i, j)) \text{ pour } (i, j) \in \llbracket 0; 2 \rrbracket^2.$$

On a $\Omega = \llbracket 1; 6 \rrbracket^2$. On le munit de la loi uniforme. On a

- $(X = 0) \cap (Y = 0) = \{1, 3\}^2$. De cela on en déduit $p_{0,0} = 1/9$.
- $(X = 0) \cap (Y = 1) = \{(1, 5), (3, 5), (5, 1), (5, 3)\}^2$. De cela on en déduit $p_{0,1} = 1/9$.
- $(X = 0) \cap (Y = 2) = \{(5, 5)\}^2$. De cela on en déduit $p_{0,2} = 1/36$.
- $(X = 1) \cap (Y = 0) = \{(1, 2), (2, 1), (1, 3), (3, 1)\}^2$. De cela on en déduit $p_{1,0} = 1/9$.

En continuant ainsi on trouve finalement :

$X \backslash Y$	0	1	2
0	1/9	1/9	1/36
1	1/9	5/18	1/9
2	1/36	1/9	1/9

Définition 7.29 (Lois marginales)

Soit (X, Y) un couple de variables aléatoires discrètes à valeurs dans $E \times F$. On appelle lois marginales du couples les lois de X et de Y .

Proposition 7.30

Soit (X, Y) un couple de variables aléatoires discrètes à valeurs dans $E \times F$. On peut retrouver les lois marginales à partir de la loi conjointe :

$$\forall x \in X(\Omega), P(X = x) = \sum_{y \in Y(\Omega)} P((X, Y) = (x, y))$$

$$\forall y \in Y(\Omega), P(Y = y) = \sum_{x \in X(\Omega)} P((X, Y) = (x, y))$$

Exemple : Dans notre exemple on obtient

$X \backslash Y$	0	1	2	X
0	1/9	1/9	1/36	1/4
1	1/9	5/18	1/9	1/2
2	1/36	1/9	1/9	1/4
Y	1/4	1/2	1/4	1

Exemple : Toujours sur notre exemple, la loi de X sachant $(Y = 0)$ est donnée par :

x	0	1	2
$P_{(Y=0)}(X = x)$	4/9	4/9	1/9

La loi de X sachant $(Y = 1)$ est donnée par :

x	0	1	2
$P_{(Y=1)}(X = x)$	2/9	5/9	2/9

ATTENTION

- On peut calculer les lois conditionnelles et les lois marginales à partir de la loi conjointe.
- Dans le cas général la connaissance des lois marginales ne suffisent pas à connaître la loi conjointe. Par contre, si on connaît les lois marginales et les lois conditionnelles on retrouve la loi conjointe car

$$\mathbf{P}((X, Y) = (x, y)) = \mathbf{P}(X = x)\mathbf{P}(Y = y|X = x).$$

Définition 7.31 (Vecteur aléatoire discret)

On appelle vecteur aléatoire discret un n -uplet $(X_1, \dots, X_n)$ de variables aléatoires discrètes.

Remarques :

1. Si pour tout $i \in \llbracket 1 ; n \rrbracket$, X_i est à valeurs dans E_i , un vecteur aléatoire est donc la variable aléatoire

$$X : \Omega \rightarrow \prod_{i=1}^n E_i$$

telle que pour tout $\omega \in \Omega$,

$$X(\omega) = (X_1(\omega), \dots, X_n(\omega))$$

2. Les définitions précédentes s'étendent au cas des vecteurs aléatoires discrets :

- La loi conjointe : $\mathbf{P}((X_1, \dots, X_n) = (x_1, \dots, x_n))$
- Les lois marginales qui sont les lois des X_i .
- Les lois conditionnelles qui sont les lois d'une variable X_i en supposant connues les valeurs des autres.
- Là encore, un vecteur aléatoire est juste un cas particulier de variables aléatoires.
- Soit $(X_1, \dots, X_n)$ un vecteur aléatoire discret. Soit $f : \prod_{i=1}^n X_i(\Omega) \rightarrow E$, la composée

$$f(X_1, \dots, X_n) : \Omega \xrightarrow{(X_1, \dots, X_n)} \prod_{i=1}^n X_i(\Omega) \rightarrow E$$

est encore un vecteur aléatoire discret. Il suffit d'utiliser que si X est une variable aléatoire discrète alors $f(X)$ aussi.

Par exemple, si les X_i sont à valeurs réelles, $\max(X_i)$, $\min(X_i)$, $\sum_i X_i$ et $\prod_i X_i$ sont des variables aléatoires discrètes.

4.3 Couples et familles de variables aléatoires indépendantes**Définition 7.32**

Soit (X, Y) un couple de variables aléatoires discrètes à valeurs dans $E \times F$. Elles sont dites indépendantes si

$$\forall A \subset X(\Omega), \forall B \subset Y(\Omega), \mathbf{P}((X \in A) \cap (Y \in B)) = \mathbf{P}(X \in A)\mathbf{P}(Y \in B)$$

On note alors $X \perp\!\!\!\perp Y$.

On peut le caractériser avec les probabilités sur les valeurs de $X(\Omega)$ et $Y(\Omega)$ ou avec les lois conditionnelles.

Proposition 7.33

Soit (X, Y) un couple de variables aléatoires discrètes à valeurs dans $E \times F$. Les assertions suivantes sont équivalentes

- i) X et Y sont indépendantes.
- ii) Pour tout $x \in X(\Omega)$ et tout $y \in Y(\Omega)$, $\mathbf{P}((X, Y) = (x, y)) = \mathbf{P}(X = x)\mathbf{P}(Y = y)$
- iii) Pour tout $A \subset X(\Omega)$ tel que A ne soit pas négligeable la loi conditionnelle Y sachant A est égale à la loi de Y .

Démonstration :

- $i) \Rightarrow ii)$ Evident en prenant $A = \{x\}$ et $B = \{y\}$.
- $ii) \Rightarrow i)$ Soit $A \subset X(\Omega)$ et $B \subset Y(\Omega)$,

$$\begin{aligned}
 \mathbf{P}((X \in A) \cap (Y \in B)) &= \sum_{\substack{x \in A \\ y \in B}} \mathbf{P}((X, Y) = (x, y)) \\
 &= \sum_{\substack{x \in A \\ y \in B}} \mathbf{P}(X = x)\mathbf{P}(Y = y) \\
 &= \left(\sum_{x \in A} \mathbf{P}(X = x) \right) \left(\sum_{y \in B} \mathbf{P}(Y = y) \right) \text{ résultats sur les familles sommables} \\
 &= \mathbf{P}(X \in A)\mathbf{P}(Y \in B).
 \end{aligned}$$

- $i) \Rightarrow iii)$ Soit $A \subset X(\Omega)$ tel que A ne soit pas négligeable, pour tout $B \subset Y(\Omega)$,

$$P_{(X \in A)}(Y \in B) = \frac{\mathbf{P}((X \in A) \cap (Y \in B))}{\mathbf{P}(X \in A)} = \frac{\mathbf{P}(X \in A)\mathbf{P}(Y \in B)}{\mathbf{P}(X \in A)} = \mathbf{P}(Y \in B).$$

Ce résultat est en particulier vrai pour $B = \{y\}$. La loi conditionnelle de Y sachant $(X \in A)$ est égale à la loi de Y .

- $iii) \Rightarrow i)$ Même calcul que ci-dessus mais dans l'autre sens.

□

Définition 7.34

Soit $(X_i)_{i \in I}$ une famille de variables aléatoires. Elles sont dites indépendantes si pour tout $A_i \subset X_i(\Omega)$ les événements $(X_i \in A_i)_{i \in I}$ sont mutuellement indépendants, c'est-à-dire que pour toute partie **finie** J de I ,

$$\mathbf{P}\left(\bigcap_{i \in J} (X_i \in A_i)\right) = \prod_{i \in J} \mathbf{P}(X_i \in A_i)$$

Remarque : Il est souvent très compliqué de montrer par le calcul que des variables sont indépendantes. La plupart du temps, l'indépendance des variables sera supposé ou déduite de la modélisation.

Proposition 7.35

Soit X et Y deux variables aléatoires sur un espace probabilisé $(\Omega, \mathcal{A}, P)$ et $f : X(\Omega) \rightarrow E$, $g : Y(\Omega) \rightarrow F$. Si $X \perp\!\!\!\perp Y$ alors $f(X) \perp\!\!\!\perp g(Y)$.

Démonstration : Soit $x \in E$ et $y \in F$, on sait que

$$(f(X) = x) = \bigcup_{\substack{u \in X(\Omega) \\ f(u)=x}} (X = u) \quad \text{et} \quad (g(Y) = y) = \bigcup_{\substack{v \in Y(\Omega) \\ g(v)=y}} (Y = v)$$

De ce fait, par distributivité de $\cap$ sur $\cup$

$$(f(X) = x) \cap (g(Y) = y) = \bigcup_{\substack{(u,v) \in X(\Omega) \times Y(\Omega) \\ f(u)=x, g(v)=y}} (X = u) \cap (Y = v)$$

Les événements dans cette union étant deux à deux disjoints et les variables X et Y étant supposées indépendantes, on obtient que

$$P\left((f(X) = x) \cap (g(Y) = y)\right) = \sum_{\substack{(u,v) \in X(\Omega) \times Y(\Omega) \\ f(u)=x, g(v)=y}} P(X = u) \times P(Y = v)$$

Or, en procédant de même,

$$P(f(X) = x) = \sum_{\substack{u \in X(\Omega) \\ f(u)=x}} P(X = u) \quad \text{et} \quad P(g(Y) = y) = \sum_{\substack{v \in Y(\Omega) \\ g(v)=y}} P(Y = v)$$

Le corollaire du théorème de sommation par paquets permet alors de conclure :

$$P\left((f(X) = x) \cap (g(Y) = y)\right) = P(f(X) = x) \times P(g(Y) = y)$$

Les variables $f(X)$ et $g(Y)$ sont bien indépendantes. □

Théorème 7.36 (Lemme des coalitions)

Soit $(X_1, \dots, X_n)$ des variables aléatoires discrètes indépendantes. Soit $k \in \llbracket 1 ; n-1 \rrbracket$ et

$$f : \prod_{i=1}^k X_i(\Omega) \rightarrow E ; g : \prod_{i=k+1}^n X_i(\Omega) \rightarrow F$$

Les variables $f(X_1, \dots, X_k)$ et $g(X_{k+1}, \dots, X_n)$ sont des variables aléatoires discrètes indépendantes.

Démonstration :

L'idée est de montrer que $Y = (X_1, \dots, X_k)$ et $Z = (X_{k+1}, \dots, X_n)$ sont des vecteurs aléatoires indépendants, d'utiliser le résultat précédent.

Soit $(x_1, \dots, x_k) \in \prod_{i=1}^k X_i(\Omega)$ et $(x_{k+1}, \dots, x_n) \in \prod_{i=k+1}^n X_i(\Omega)$,

$$(Y = (x_1, \dots, x_k)) = \bigcap_{i=1}^k (X_i = x_i) \text{ et } (Z = (x_{k+1}, \dots, x_n)) = \bigcap_{i=k+1}^n (X_i = x_i).$$

On en déduit, d'après l'indépendance des variables X_i que

$$P\left((Y = (x_1, \dots, x_k)) \cap (Z = (x_{k+1}, \dots, x_n))\right) = P\left(\bigcap_{i=1}^n (X_i = x_i)\right) = \prod_{i=1}^n P(X_i = x_i)$$

De même

$$P(Y = (x_1, \dots, x_k)) \times P(Z = (x_{k+1}, \dots, x_n)) = \prod_{i=1}^k P(X_i = x_i) \times \prod_{i=k+1}^n P(X_i = x_i) = \prod_{i=1}^n P(X_i = x_i).$$

On a bien montré que Y et Z étaient indépendantes. Il suffit alors d'utiliser le résultat précédent.

□

Corollaire 7.37 (Extension a plus de deux coalitions)

Soit $(X_1, \dots, X_n)$ des variables aléatoires discrètes indépendantes. Soit $r \geq 1$ et $1 < a_2 < \dots < a_{r-1} < a_r < n$ et $f_1, \dots, f_r$ sont des fonctions telles que pour tout $k \in \llbracket 1; r \rrbracket$, f_k est définie sur $\prod_{i=a_k}^{a_{k+1}-1} X_k(\Omega)$. Les variables aléatoires $f_1(X_1, \dots, X_{a_2-1}), \dots, f_r(X_{a_r}, \dots, X_n)$ sont indépendantes.

Démonstration : Les démonstrations ci-dessus se généralisent sans problème.

□

Finissons ce paragraphe par un théorème important qui permet de dire que ce que l'on fait a un sens. Dans la plupart des cas on ne va s'intéresser qu'aux variables aléatoires sans regarder précisément les espaces probabilisés sous-jacents. On peut se demander si cela a un sens.

Théorème 7.38 (Théorème d'extension de Kolmogorov)

Soit $(\mathcal{L}_n)_{n \in \mathbb{N}}$ une suite de lois discrètes (c'est-à-dire des probabilités sur des ensembles E_n finis ou dénombrables munis de la tribu $\mathcal{P}(E_n)$). Il existe un espace probabilisé $(\Omega, \mathcal{A}, P)$ et des variables aléatoires discrètes indépendantes telles que

$$\forall n \in \mathbb{N}, X_n \sim \mathcal{L}_n$$

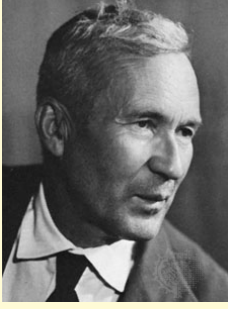
Démonstration : Ce théorème est admis

□

Remarque : Ce théorème est essentiellement là pour justifier les énoncés des exercices / problèmes. Il est à priori utilisé de manière sous-entendue.

Exemple : On veut modéliser un jeu de pile ou face infini. Soit $p \in]0, 1[$. On se donne les lois de probabilités $\mathcal{L}_n$ toutes égales à la loi de Bernoulli $\mathcal{B}(1, p)$. Le théorème d'existence de Kolmogorov assure qu'il existe un espace probabilisé qui modélise cette expérience.

Matheux (Andrei Kolmogorov : 1903-1987)



Andreï Nikolaïevitch Kolmogorov (25 avril 1903 à Tambov – 20 octobre 1987 à Moscou) est un mathématicien russe et soviétique qui a apporté des contributions significatives en mathématiques, notamment en théorie des probabilités, topologie, turbulence, mécanique classique, logique intuitionniste, théorie algorithmique de l'information et en analyse de la complexité des algorithmes

5 Lois usuelles

Nous allons commencer par revoir les lois usuelles vues en première année puis nous étudierons les lois géométriques et lois de Poisson.

5.1 Loi uniforme

Définition 7.39

Soit $N \in \mathbb{N}^*$ et $x_1, \dots, x_N$ des réels distincts. Soit $(\Omega, \mathcal{A}, P)$ un espace probabilisé et X une variable aléatoire réelle. On dit que X suit la loi uniforme sur $\{x_1, \dots, x_N\}$ si :

1. $X(\Omega) = \{x_1, \dots, x_N\}$

2. pour tout i compris entre 1 et N , $P(X = x_i) = \frac{1}{N}$.

On note alors $X \sim \mathcal{U}(\{x_1, \dots, x_N\})$.

Remarques :

1. C'est la loi de l'équiprobabilité. En effet, savoir que toutes les instances sont équiprobables induit que la loi de la variable aléatoire est la loi uniforme.
2. La plupart du temps on regardera - ou en se ramènera à étudier - la loi uniforme sur $[[1; N]]$.

5.2 Loi de Bernoulli et loi binomiale

Ces lois ont été vues en premières années.

Définition 7.40

Soit $(\Omega, \mathcal{A}, P)$ un espace probabilisé et X une variable aléatoire réelle.

1. Soit $p \in [0, 1]$, on dit que X suit la loi de Bernoulli de paramètre p si :

(a) $X(\Omega) = \{0, 1\}$

(b) $P(X = 1) = p$ et donc $P(X = 0) = 1 - p$.

On note alors $X \sim \mathcal{B}(p)$.

2. Soit $p \in [0, 1]$ et $n \in \mathbb{N}^*$. On dit que X suit la loi binomiale de paramètres n, p si :

(a) $X(\Omega) = \{0, 1, \dots, n\}$

(b) Pour tout $k \in \{0, 1, \dots, n\}$, $P(X = k) = \binom{n}{k} p^k q^{n-k}$ où $q = 1 - p$.

On note alors $X \sim \mathcal{B}(n, p)$.

Matheux (Jacques Bernoulli : 1654-1705)



Jacques ou Jakob Bernoulli (1654-1705) est un mathématicien et physicien suisse (né et mort à Bâle), frère de Jean Bernoulli et oncle de Daniel Bernoulli et Nicolas Bernoulli.

Les travaux de Jacques Bernoulli sont de première importance dans de nombreux domaines : calcul différentiel, calcul intégral (le terme est de lui), séries, probabilités et dans l'étude de courbes qui l'amène à résoudre certaines équations différentielles.

La loi binomiale correspond au schéma de Bernoulli : répétition de n expériences **indépendantes** ayant deux issues échec / succès. La probabilité de succès et la même pour toutes les épreuves et elle est notée p .

Plus précisément, on se donne une suite de variables aléatoires $(X_i)_{i \in \llbracket 1; n \rrbracket}$ mutuellement indépendantes suivant toutes la loi de Bernoulli $\mathcal{B}(1, p)$ et on considère $X = \sum_{i=1}^n X_i$. On peut interpréter X comme le nombre de succès lors des n expériences. Un calcul fait en première année montre que $X \sim \mathcal{B}(n, p)$.

On note que c'est bien une loi de probabilité car

$$\sum_{k=0}^n P(X = k) = \sum_{k=0}^n \binom{n}{k} p^k q^{n-k} = (p + 1 - p)^n = 1.$$

5.3 Loi géométrique

Soit $p \in [0, 1]$, on se donne une suite $(X_n)_{n \in \mathbb{N}^*}$ de variables aléatoires indépendantes de même loi de probabilité $\mathcal{B}(p)$. On peut penser à une suite infinie de lancers à pile ou face avec une pièce ayant une probabilité p de faire pile.

On s'intéresse alors au rang du premier succès. Précisément on pose

$$N : \Omega \rightarrow \mathbb{N}^* \cup \{+\infty\}$$

$$\omega \mapsto \begin{cases} +\infty & \text{si } \forall k \in \mathbb{N}, X_k(\omega) = 0 \\ \min\{k \in \mathbb{N}^* \mid X_k(\omega) = 1\} & \text{sinon} \end{cases}$$

C'est une variable aléatoire discrète car $N(\Omega) \subset \mathbf{N}^* \cup \{+\infty\}$ est bien dénombrable. De plus pour tout $k \in \mathbf{N}^*$,

$$N^{-1}(\{k\}) = X_k^{-1}(\{1\}) \cap \bigcap_{i=1}^{k-1} X_i^{-1}(\{0\}) \in \mathcal{A}$$

et

$$N^{-1}(\{+\infty\}) = \overline{\bigcup_{k \in \mathbf{N}^*} N^{-1}(\{k\})} \in \mathcal{A}.$$

Maintenant, pour $k \in \mathbf{N}^*$, on a

$$(N = k) = \bigcap_{i=1}^{k-1} (X_i = 0) \cap (X_k = 1)$$

Par indépendance,

$$\mathbf{P}(N = k) = (1 - p)^{k-1} p.$$

De plus, si on note $A_k = \bigcap_{i=1}^k (X_i = 0)$ l'événement « on n'a pas eu de succès lors des k premiers essais ». On a

$$\mathbf{P}(A_k) = (1 - p)^k$$

La suite (A_k) est décroissante et $\bigcap_{k=0}^{+\infty} A_k = (N = +\infty)$, on en déduit que

$$\mathbf{P}(N = +\infty) = \lim_{k \rightarrow \infty} (1 - p)^k = \begin{cases} 0 & \text{si } p \neq 0 \\ 1 & \text{sinon.} \end{cases}$$

On peut aussi voir que

$$1 = \mathbf{P}(\Omega) = \sum_{k=0}^{+\infty} \mathbf{P}(N = k) + \mathbf{P}(N = +\infty).$$

De ce fait,

$$\mathbf{P}(N = +\infty) = 1 - \sum_{k=1}^{+\infty} \mathbf{P}(N = k) = 1 - \sum_{k=1}^{+\infty} (1 - p)^{k-1} p = 1 - p \sum_{k=1}^{+\infty} (1 - p)^{k-1} = 1 - p \frac{1}{1 - (1 - p)} = 0.$$

Pour éviter ce problème, on supposera que $p \neq 0$ et $p \neq 1$.

Définition 7.41 (Loi géométrique)

Soit $p \in]0, 1[$. On dit qu'une variable aléatoire discrète N à valeurs dans $\mathbf{N}^*$ suit la loi géométrique de paramètre p et on note $N \sim \mathcal{G}(p)$ si

$$\forall k \in \mathbf{N}^*, \mathbf{P}(N = k) = (1 - p)^{k-1} p.$$

Remarques :

1. On peut ajouter $+\infty$ dans l'image de la loi et ajouter alors $\mathbf{P}(N = +\infty) = 0$.
2. Il est souvent plus simple de travailler avec des inégalités. En effet pour $k \in \mathbf{N}$, $(N > k) =$ « les k premiers essais ont été des échecs ». On a donc

$$\mathbf{P}(N > k) = (1 - p)^k.$$

ATTENTION

Ne pas mélanger la loi géométrique avec la loi binomiale. Dans les deux cas on s'intéresse à une suite d'épreuves de Bernoulli indépendantes de même probabilité de succès mais pour la loi binomiale on regarde le **nombre de succès** lors d'un nombre donné d'épreuves alors que pour la loi géométrique on regarde le **rang du premier succès**.

5.4 La loi de Poisson

Commençons par une modélisation.

On considère un péage autoroutier et on veut savoir combien de voiture vont se présenter au péage durant une période T . On considère que si on se donne un instant (très) court δt il ne peut pas en arriver deux en même temps et que la probabilité qu'il en arrive une est proportionnelle à δt . On note α le coefficient de colinéarité.

On découpe donc notre temps T en n instants δt de durée $\frac{T}{n}$. Au total la loi du nombre de voiture est donnée par la loi binomiale $\mathcal{B}(n, \alpha \frac{T}{n})$.

De manière un peu plus générale on considère une suite de variables aléatoires discrètes X_n telles que

$$X_n \sim \mathcal{B}(n, p_n)$$

où (np_n) converge vers une constante notée λ (dans notre exemple $\lambda = \alpha T$).

On a donc $\forall k \in \mathbb{N}$,

$$\mathbf{P}(X_n = k) = \begin{cases} 0 & \text{si } n < k \\ \binom{n}{k} p_n^k (1 - p_n)^{n-k} & \text{sinon.} \end{cases}$$

Si on s'intéresse à $\lim_{n \rightarrow \infty} \mathbf{P}(X_n = k)$ on ne peut s'intéresser qu'à la deuxième formule car n finira par être supérieur à k .

Maintenant

$$\begin{aligned} \binom{n}{k} p_n^k (1 - p_n)^{n-k} &= \frac{n(n-1) \cdots (n-k+1)}{k!} p_n^k (1 - p_n)^{n-k} \\ &\sim \frac{n^k}{k!} p_n^k (1 - p_n)^{n-k} \\ &\sim \frac{n^k}{k!} \left(\frac{\lambda}{n}\right)^k \exp((n-k) \ln(1 - p_n)) \\ &\sim \frac{\lambda^k e^{-\lambda}}{k!} \end{aligned}$$

Définition 7.42 (Loi de Poisson)

Pour tout $\lambda \in \mathbb{R}_+^*$, on dit qu'une variable aléatoire discrète à valeurs dans $\mathbb{N}$ suit la loi de Poisson de paramètre λ ($X \sim \mathcal{P}(\lambda)$) si

$$\forall k \in \mathbb{N}, \mathbf{P}(X = k) = \frac{\lambda^k e^{-\lambda}}{k!}.$$

Remarques :

1. La loi de Poisson s'appelle la loi des événements rares car, comme dans la modélisation ci-dessus, elle permet d'évaluer le nombre d'apparition d'un phénomène rare lors d'un assez grand laps de temps.
2. La loi de Poisson est plus simple à calculer que la loi Binomiale.
3. En étudiant la suite $u_k = \frac{\lambda^k e^{-\lambda}}{k!}$ qui est une suite positive, on voit que $\frac{u_{k+1}}{u_k} = \frac{\lambda}{k+1}$. Cette suite est donc croissante pour $k < \lambda$ et décroissante ensuite. Le maximum de probabilité est donc situé « autour » de λ .
4. On peut vérifier que l'on a bien

$$\sum_{k=0}^{+\infty} \mathbf{P}(X = k) = e^{-\lambda} \sum_{k=0}^{+\infty} \frac{\lambda^k}{k!} = 1$$

Matheux (Siméon Denis Poisson : 1781-1840 (X - 1798))



Siméon Denis Poisson (21 juin 1781 à Pithiviers - 25 avril 1840 à Sceaux) est un mathématicien, géomètre et physicien français. En mathématique, ses travaux les plus importants portent sur la série sur les intégrales définies et sa discussion sur les séries de Fourier, qui préparèrent le terrain des recherches classiques de Dirichlet et Bernhard Riemann sur le même sujet. Il étudia aussi les intégrales de Fourier.

Il faut retenir néanmoins, qu'en tant que membre de l'Académie des sciences, il fut chargé en 1830 avec Lacroix d'examiner le mémoire d'un jeune mathématicien Evariste Galois Conditions pour qu'une équation soit résoluble par radicaux. Poisson rendra un rapport négatif le 4 juillet, jugeant le travail incompréhensible.

Exercice : Soit $X \sim \mathcal{P}(\lambda)$ et $Y \sim \mathcal{P}(\mu)$ des variables indépendantes. Calculer la loi de $X + Y$.

Polynôme minimal et réduction

1	Algèbres	204
1.1	Définition	204
1.2	Sous-algèbre et morphismes d'algèbre	205
2	Polynômes d'endomorphismes et de matrices	206
2.1	Polynômes de matrices	206
2.2	Polynômes d'endomorphismes	208
3	Compléments d'algèbre	210
3.1	Idéaux d'un anneau commutatif	210
3.2	Plus grand commun diviseur	213
3.3	Algorithme d'Euclide	217
3.4	Généralisation à n polynômes	218
4	Polynômes annulateurs	219
4.1	Définitions	219
4.2	Polynôme minimal	220
4.3	Polynômes annulateurs et valeurs propres	224
4.4	Théorème de Cayley-Hamilton	224
4.5	Calcul de puissances	225
5	Critère de diagonalisabilité	226
5.1	Décomposition des noyaux	227
5.2	Polynômes annulateurs et diagonalisabilité	229
5.3	Endomorphismes annulés par un polynôme scindé	231
6	Compléments et applications	234
6.1	Influence du corps de base	234
6.2	Calculs de puissances	234
6.3	Commutant	235
6.4	Décomposition de Dunford	236
6.5	Racines carrées	237
6.6	Sous-espaces stables	237

Dans ce chapitre $\mathbf{K}$ désigne $\mathbf{R}$ ou $\mathbf{C}$. La plupart des résultats restent vraies pour n'importe quel corps. On continue l'étude de la réduction des endomorphismes sur un espace vectoriel E de dimension finie n .

1 Algèbres

1.1 Définition

Définition 8.1

On appelle algèbre sur $\mathbf{K}$ ou $\mathbf{K}$ -algèbre un quadruplet $(E, +, \cdot, \times)$ où E est un ensemble, $+$ et $\times$ des lois internes et $\cdot$ une loi externe :

$$\begin{array}{llll} + : E \times E \rightarrow E & ; & \times : E \times E \rightarrow E & \text{et} & \cdot : \mathbf{K} \times E \rightarrow E \\ (u, v) \mapsto u + v & & (u, v) \mapsto u \times v & & (\lambda, u) \mapsto \lambda \cdot u \end{array}$$

vérifiant

- $(E, +, \cdot)$ est un $\mathbf{K}$ -espace vectoriel.
- $(E, +, \times)$ est un anneau ; son élément neutre pour $\times$ s'appelle l'unité de l'algèbre.
- L'application $(u, v) \mapsto u \times v$ est bilinéaire.

Remarque : On rappelle que dire que $(u, v) \mapsto u \times v$ est bilinéaire signifie que pour tout u de E , l'application $v \mapsto u \times v$ est linéaire et que pour tout v de E l'application $u \mapsto u \times v$ aussi. C'est-à-dire :

- $\forall u \in E, \forall v \in E, \forall w \in E, u \times (v + w) = u \times v + u \times w$
- $\forall u \in E, \forall v \in E, \forall \lambda \in \mathbf{K}, u \times (\lambda \cdot v) = \lambda \cdot (u \times v)$
- $\forall w \in E, \forall u \in E, \forall v \in E, (u + v) \times w = u \times w + v \times w$
- $\forall v \in E, \forall u \in E, \forall \lambda \in \mathbf{K}, (\lambda \cdot u) \times v = \lambda \cdot (u \times v)$

Exemples :

1. Le corps $(\mathbf{K}, +, \cdot, \times)$ est une $\mathbf{K}$ -algèbre. Dans ce cas, $\cdot$ et $\times$ sont les mêmes lois.
2. L'ensemble $(\mathbf{K}[X], +, \cdot, \times)$ est une $\mathbf{K}$ -algèbre. Son unité est le polynôme 1.
3. Soit X un ensemble, $\mathcal{F}(X, \mathbf{K})$ est muni d'une structure de $\mathbf{K}$ -algèbre où si $(f, g) \in \mathcal{F}(X, \mathbf{K})^2$ et $\lambda \in \mathbf{K}$

- L'application $f + g$ est définie par $x \mapsto f(x) + g(x)$
- L'application $\lambda \cdot f$ est définie par $x \mapsto \lambda \cdot f(x)$
- L'application $f \times g$ est définie par $x \mapsto f(x) \times g(x)$

Son unité est la fonction constante égale à 1. En particulier pour $X = \mathbf{N}$ on obtient que l'ensemble des suites à valeurs dans $\mathbf{K}$ admet une structure de $\mathbf{K}$ -algèbre.

4. L'ensemble des matrices carrées $\mathcal{M}_n(\mathbf{K})$ a une structure d'algèbre. Son unité est I_n .
5. Soit E un espace vectoriel, l'ensemble des endomorphismes a une structure d'algèbre. Le produit est la composition. Son unité est id_E .

1.2 Sous-algèbre et morphismes d'algèbre

Définition 8.2 (Sous-algèbre)

Soit $(E, +, \cdot, \times)$ une $\mathbf{K}$ -algèbre. Une sous-algèbre de E est une partie F de E telle que

- L'unité de E appartient à F .
- F est stable par combinaison linéaire (ce qui revient à dire que F est un sous-espace vectoriel)
- F est stable par produit (ce qui revient à dire que F est un sous-anneau car si u, v sont des éléments de F on sait déjà que $u - v \in F$).

Exemples :

1. Pour $E = \mathcal{F}(\mathbf{R}, \mathbf{R})$, on a comme sous-algèbres :
 - l'ensemble des fonctions continues.
 - Pour tout k ($0 \leq k \leq +\infty$), $\mathcal{C}^k(\mathbf{R}, \mathbf{R})$ est une sous-algèbre.
 - L'ensemble des fonctions polynomiales est une sous-algèbre.
 - L'ensemble des fonctions qui s'annulent en 0 n'est pas une sous-algèbre. Il vérifie les conditions de stabilité mais ne contient pas l'unité.
2. Pour $E = \mathbf{R}^{\mathbf{N}}$, l'ensemble des suites convergentes est une sous-algèbre.
3. Pour $E = \mathbf{K}[X]$, l'ensemble $\mathbf{K}[X^2]$ des polynômes pairs est une sous-algèbre.
4. Pour $E = \mathcal{M}_n(\mathbf{K})$, l'ensemble $T_n(\mathbf{K})$ des matrices triangulaires supérieures est une sous-algèbre.

Définition 8.3 (Morphisme d'algèbres)

Soit $(\mathcal{A}, +, \cdot, \times)$ et $(\mathcal{B}, \oplus, \odot, \otimes)$ deux $\mathbf{K}$ -algèbres. Un morphisme d'algèbre de $\mathcal{A}$ vers $\mathcal{B}$ est une application $F : \mathcal{A} \rightarrow \mathcal{B}$ compatible aux structures à savoir :

- L'image de l'unité est l'unité : $F(1_{\mathcal{A}}) = 1_{\mathcal{B}}$.
- C'est une application linéaire de $(\mathcal{A}, +, \cdot)$ dans $(\mathcal{B}, \oplus, \odot)$:

$$\forall (u, v) \in \mathcal{A}^2, \forall (\lambda, \mu) \in \mathbf{K}^2, F(\lambda \cdot u + \mu \cdot v) = \lambda \odot F(u) \oplus \mu \odot F(v)$$

- Elle est compatible à la multiplication :

$$\forall (u, v) \in \mathcal{A}^2, F(u \times v) = F(u) \otimes F(v)$$

Exemples :

1. La fonction qui associe à une suite convergente sa limite est un morphisme d'algèbre.
2. Soit $E = \mathcal{F}(X, \mathbf{K})$ et $x_0 \in X$, la fonction d'évaluation en x_0 définie par $f \mapsto f(x_0)$ est un morphisme d'algèbre.
3. Soit E un espace vectoriel de dimension finie et $\mathcal{B}$ une base. L'application de $\mathcal{L}(E)$ dans $\mathcal{M}_n(\mathbf{K})$ qui associe à un endomorphisme sa matrice dans la base $\mathcal{B}$ est un morphisme d'algèbre.

Exercice : La dérivation de $\mathcal{C}^\infty(\mathbf{R}, \mathbf{R})$ dans lui-même est-il un morphisme d'algèbre ? La transposition de $\mathcal{M}_n(\mathbf{K})$ dans lui-même est-il un morphisme d'algèbre ?

Définition 8.4

Soit $F : \mathcal{A} \rightarrow \mathcal{B}$ un morphisme d'algèbre.

1. On appelle image de F et on note $\text{Im}(F)$ l'ensemble des images des éléments de $\mathcal{A}$ par F :

$$\text{Im}(F) = F(\mathcal{A}) = \{F(u) , u \in \mathcal{A}\} \subset \mathcal{B}.$$

2. On appelle noyau de F et on note $\text{Ker}(F)$ l'image inverse de $\{0\}$.

$$\text{Ker}(F) = F^{-1}(\{0\}) = \{u \in \mathcal{A} , F(u) = 0\}.$$

Proposition 8.5

Avec les notations précédentes, $\text{Im}(F)$ est une sous-algèbre de $\mathcal{B}$.

Remarque : Par contre, dès que $\mathcal{B} \neq \{0\}$, $\text{Ker}(F)$ n'est pas une sous-algèbre car $F(1_{\mathcal{A}}) = 1_{\mathcal{B}} \neq 0_{\mathcal{B}}$. Donc $1_{\mathcal{A}} \notin \text{Ker}(F)$.

Démonstration : On vérifie la caractérisation :

- On a $1_{\mathcal{B}} \in \text{Im}(F)$ car $F(1_{\mathcal{A}}) = 1_{\mathcal{B}}$.
- Soit $(v, v') \in (\text{Im}(F))^2$ et $(\lambda, \mu) \in \mathbf{K}^2$. Par définition, il existe (u, u') dans $\mathcal{A}$ tels que $F(u) = v$ et $F(u') = v'$. Or on a

$$F(\lambda.u + \mu.u') = \lambda.F(u) + \mu.F(u') = \lambda.v + \mu.v'$$

Donc $\lambda.v + \mu.v' \in \text{Im}(F)$.

- On fait pareil pour le produit.

□

2 Polynômes d'endomorphismes et de matrices

2.1 Polynômes de matrices

Définition 8.6

Soit $A \in \mathcal{M}_n(\mathbf{K})$ une matrice et $P \in \mathbf{K}[X]$ un polynôme. On pose

$$P(A) = \sum_{k=0}^d a_k A^k \text{ où } P = \sum_{k=0}^d a_k X^k.$$

Remarque : C'est la structure de $\mathbf{K}$ -algèbre de $\mathcal{M}_n(\mathbf{K})$ qui permet de poser cette définition. On utilise en effet le produit (pour A^k), la multiplication externe (pour la multiplication par a_k) et l'addition (pour la somme).

Exemples :

1. Si $P = X^2 - 1$, $P(A) = A^2 - I_n$ car $A^0 = I_n$.

Dans le cas où $A = \begin{pmatrix} 1 & 1 & 1 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{pmatrix}$, on a $A^2 = \begin{pmatrix} 1 & 2 & 3 \\ 0 & 1 & 2 \\ 0 & 0 & 1 \end{pmatrix}$ et donc

$$P(A) = \begin{pmatrix} 0 & 2 & 3 \\ 0 & 0 & 2 \\ 0 & 0 & 0 \end{pmatrix}.$$

2. Si $P = X$, $P(A) = A$.

Proposition 8.7

Soit $A \in \mathcal{M}_n(\mathbf{K})$, l'application

$$\begin{aligned} \Phi_A : \mathbf{K}[X] &\rightarrow \mathcal{M}_n(\mathbf{K}) \\ P &\mapsto P(A) \end{aligned}$$

est un morphisme d'algèbres ce qui signifie que

- i) $\Phi_A(1) = I_n$
- ii) $\forall (P, Q) \in \mathbf{K}[X]^2, \Phi_A(P + Q) = \Phi_A(P) + \Phi_A(Q)$ c'est-à-dire $(P + Q)(A) = P(A) + Q(A)$
- iii) $\forall (P, Q) \in \mathbf{K}[X]^2, \Phi_A(P \times Q) = \Phi_A(P) \times \Phi_A(Q)$ c'est-à-dire $(P \times Q)(A) = P(A) \times Q(A)$
- iv) $\forall P \in \mathbf{K}[X], \forall \lambda \in \mathbf{K}, \Phi_A(\lambda.P) = \lambda.\Phi_A(P)$ c'est-à-dire $(\lambda.P)(A) = \lambda.P(A)$

Démonstration : Toutes les propositions sont « évidentes ». Il faut juste remarquer pour iii) que si on note $P = \sum_{i=0}^d a_i X^i$ et $Q = \sum_{j=0}^d b_j X^j$ où $d \geq \max(\deg(P), \deg(Q))$,

$$P(A) \times Q(A) = \left(\sum_{i=0}^d a_i A^i \right) \times \left(\sum_{j=0}^d b_j A^j \right) = \sum_{k=0}^{2d} c_k A^k = (P \times Q)(A)$$

où $c_k = \sum_{i+j=k} a_i b_j$ car A^i commute à A^j . □

Remarque : De manière générale si $\mathcal{A}$ est une $\mathbf{K}$ -algèbre on peut poser pour $x \in \mathcal{A}$, $P(x) = \sum_{i=0}^d a_i x^i$

où $P = \sum_{k=0}^d a_k X^k$. On construit ainsi un morphisme d'algèbre

$$\begin{aligned} \Phi_x : \mathbf{K}[X] &\rightarrow \mathcal{A} \\ P &\mapsto P(x) \end{aligned}$$

Définition 8.8

Soit $A \in \mathcal{M}_n(\mathbf{K})$. On note $\mathbf{K}[A] = \{P(A) \mid P \in \mathbf{K}[X]\}$. C'est l'ensemble des polynômes en A .

Proposition 8.9

Avec les notations précédentes,

1. $\mathbf{K}[A]$ est une sous-algèbre de $\mathcal{M}_n(\mathbf{K})$, c'est-à-dire un sous-espace vectoriel, stable par produit contenant I_n .
2. La sous-algèbre $\mathbf{K}[A]$ est commutative.
3. La famille $(I_n, A, A^2, \dots)$ engendre $\mathbf{K}[A]$ en tant qu'espace vectoriel.

Démonstration :

1. On peut le démontrer à la main ou voir que $\mathbf{K}[A] = \text{Im}\Phi_A$ donc c'est une sous-algèbre.
2. pour tout P et Q dans $\mathbf{K}[X]$,

$$P(A) \times Q(A) = (P \times Q)(A) = (Q \times P)(A) = Q(A) \times P(A).$$

3. Par définition, pour tout $B \in \mathbf{K}[A]$, il existe $(a_0, \dots, a_d) \in \mathbf{K}^{d+1}$ tels que $B = \sum_{k=0}^d a_k A^k$ et donc $B \in \text{Vect}(I_n, A, A^2, \dots)$.

Réciproquement, comme, pour tout $n \in \mathbf{N}$, $A^n \in \mathbf{K}[A]$, $\text{Vect}(I_n, A, A^2, \dots) \subset \mathbf{K}[A]$.

□

Remarque : L'algèbre $\mathcal{M}_n(\mathbf{K})$ elle n'est pas commutative sauf si $n = 1$.

2.2 Polynômes d'endomorphismes

On peut de même définir des polynômes d'endomorphismes en utilisant que $\mathcal{L}(E)$ est une $\mathbf{K}$ -algèbre.

ATTENTION

Dans la $\mathbf{K}$ -algèbre $\mathcal{L}(E)$, le « produit » est la composition qui se note $\circ$.

Définition 8.10

Soit $u \in \mathcal{L}(E)$ un endomorphisme et $P \in \mathbf{K}[X]$ un polynôme. On pose

$$P(u) = \sum_{k=0}^d a_k u^k \text{ où } P = \sum_{k=0}^d a_k X^k.$$

Proposition 8.11

Soit $u \in \mathcal{L}(E)$, l'application

$$\begin{aligned} \Phi_u : \mathbf{K}[X] &\rightarrow \mathcal{L}(E) \\ P &\mapsto P(u) \end{aligned}$$

est un morphisme d'algèbres ce qui signifie que

- i) $\Phi_u(1) = \text{id}_E$
- ii) $\forall (P, Q) \in \mathbf{K}[X]^2, \Phi_u(P + Q) = \Phi_u(P) + \Phi_u(Q)$ c'est-à-dire $(P + Q)(u) = P(u) + Q(u)$
- iii) $\forall (P, Q) \in \mathbf{K}[X]^2, \Phi_u(P \times Q) = \Phi_u(P) \circ \Phi_u(Q)$ c'est-à-dire $(P \times Q)(u) = P(u) \circ Q(u)$
- iv) $\forall P \in \mathbf{K}[X], \forall \lambda \in \mathbf{K}, \Phi_u(\lambda.P) = \lambda.\Phi_u(P)$ c'est-à-dire $(\lambda.P)(u) = \lambda.P(u)$

Démonstration : Comme ci-dessus. Là encore l'argument clé est que $u^i \circ u^j = u^j \circ u^i$. $\square$

Définition 8.12

Soit $u \in \mathcal{L}(E)$. On note $\mathbf{K}[u] = \{P(u), P \in \mathbf{K}[X]\}$. C'est l'ensemble des polynômes en u .

Proposition 8.13

Avec les notations précédentes,

1. $\mathbf{K}[u]$ est une sous-algèbre de $\mathcal{L}(E)$, c'est-à-dire un sous-espace vectoriel, stable par produit contenant id_E .
2. La sous-algèbre $\mathbf{K}[u]$ est commutative.
3. La famille $(\text{id}_E, u, u^2, \dots)$ engendre $\mathbf{K}[u]$ en tant qu'espace vectoriel.

ATTENTION

1. L'espace vectoriel E n'est pas à priori une $\mathbf{K}$ -algèbre car il n'y a pas de produit. ON NE PEUT PAS MULTIPLIER DES VECTEURS.

De ce fait, si $P \in \mathbf{K}[X]$ et $x \in E$, $P(x)$ n'a aucun sens.

2. En particulier, si $P \in \mathbf{K}[X]$, $u \in \mathcal{L}(E)$ et $x \in E$, $P(u)(x)$ existe et vaut

$$P(u)(x) = \left(\sum_{k=0}^d a_k u^k \right) (x) = \sum_{k=0}^d a_k u^k(x)$$

alors que $P(u(x))$ ne veut rien dire.

Proposition 8.14

Soit $u \in \mathcal{L}(E)$ et $\mathcal{B}$ une base de E . On note $A = \text{Mat}_{\mathcal{B}}(u)$. Pour tout polynôme P de $\mathbf{K}[X]$,

$$\text{Mat}_{\mathcal{B}}(P(u)) = P(A).$$

Démonstration : Cela découle du fait que

$$\begin{aligned} \text{Mat}_{\mathcal{B}} : \mathcal{L}(E) &\rightarrow \mathcal{M}_n(\mathbf{K}) \\ u &\mapsto \text{Mat}_{\mathcal{B}}(u) \end{aligned}$$

est un morphisme d'algèbre. □

3 Compléments d'algèbre

On va vouloir s'intéresser aux polynômes annulateurs d'un endomorphisme ou d'une matrice. C'est-à-dire les polynôme P tels que $P(u) = 0$ (ou $P(A) = 0$). C'est-à-dire le noyau Φ_u (ou Φ_A). On va donc étudier la structure algébrique de cet ensemble.

3.1 Idéaux d'un anneau commutatif

Définition 8.15

Soit $(A, +, \times)$ un anneau commutatif. Une partie I de A est un idéal si

- i) C 'est un sous-groupe de $(A, +)$
- ii) Pour tout $a \in A$ et $i \in I$, $ai \in I$.

ATTENTION

Ne pas confondre idéal et sous-anneau. La condition de stabilité de l'idéal est plus contraignante, par contre l'élément neutre (multiplicatif) n'appartient pas nécessairement à un idéal.

Exemples :

1. Pour tout anneau commutatif A , les parties $\{0\}$ et A sont des idéaux (dits triviaux)
2. Si $A = \mathbf{Z}$. Pour tout entier $m \in \mathbf{Z}$, $I = m\mathbf{Z}$ est un idéal. De fait l'arithmétique développée en première année avec les sous-groupes de $\mathbf{Z}$ se construit naturellement avec les idéaux (qui dans le cas de $\mathbf{Z}$ sont les mêmes).
3. Si I est un idéal de A qui contient l'élément neutre 1 alors $I = A$ car pour tout $a \in A$, $a = a \times 1 \in I$ car $1 \in I$.
4. Soit $A = \mathcal{F}(X, \mathbf{R})$. Soit $x \in X$. On pose $I_x = \{f \in A, f(x) = 0\}$. Montrons que c'est un idéal.
 - On voit que I_x n'est pas vide car la fonction nulle appartient à I_x . De plus si $(f, g) \in I_x^2$, $f - g \in I_x$ car

$$(f - g)(x) = f(x) - g(x) = 0 - 0 = 0.$$
 Donc $(I_x, +)$ est un sous-groupe de $(A, +)$.
 - Soit $f \in I_x$ et $g \in A$, $fg \in I_x$ car

$$(fg)(x) = f(x)g(x) = 0.$$

On a bien montré que I_x était un idéal de A .

Exercices :

1. Montrer que si A un anneau commutatif, l'ensemble $\text{Nil}(A)$ des éléments nilpotents de A est un idéal

2. Soit A un anneau commutatif et I un idéal de A . On appelle radical de I et on note $\sqrt{I}$ l'ensemble

$$\sqrt{I} = \{x \in A, \exists n \in \mathbb{N}, x^n \in I\}$$

Montrer que $\sqrt{I}$ est un idéal et retrouver le résultat de l'exercice précédent.

Proposition 8.16

Soit A un anneau commutatif et $a \in A$. On note aA défini par

$$aA = \{ab \mid b \in A\}$$

1. L'ensemble aA est un idéal de A
2. L'ensemble aA est le plus petit idéal contenant a .

Définition 8.17 (Idéal engendré)

Avec les notations précédentes, l'idéal aA s'appelle, idéal engendré par a et se note aA ou (a) .

Démonstration :

1. — On voit que $aA \neq \emptyset$ car $0 \in A$. De plus si $(x, x') \in (aA)^2$, il existe $(b, b') \in A^2$ tels que

$$x = ab \text{ et } x' = ab'$$

Dès lors, $x - x' = ab - ab' = a(b - b') \in aA$. On en déduit que aA est bien un sous-groupe de $(A, +)$.

- Soit $x \in aA$ et $y \in A$. Il existe $b \in A$ tel que $x = ab$. De ce fait, $xy = (ab)y = a(by) \in aA$.

Cela montre bien que aA est un idéal de A .

2. Soit I un idéal de A contenant a . Par définition d'un idéal, pour tout $b \in A$, $ab \in I$. Cela montre que $aA \subset I$. Ceci montre que aA est bien le plus petit idéal de A contenant a .

□

Définition 8.18

Soit A un anneau commutatif.

Soit I un idéal de A . Il est dit principal s'il existe $x \in A$ tel que $I = (x)$.

Remarque : On appelle anneau principal un anneau intègre dont tous les idéaux sont principaux.

Exemples :

1. Il a été vu en première année que les idéaux de $\mathbb{Z}$ étaient de la forme $m\mathbb{Z} = (m)$ pour $m \in \mathbb{Z}$. (Donc $\mathbb{Z}$ est un anneau principal). Nous redonnerons la preuve plus loin.
2. Soit $A = \mathcal{F}(\mathbb{R}, \mathbb{R})$ et $x_0 \in \mathbb{R}$. On pose $I_{x_0} = \{f \in A \mid f(x_0) = 0\}$. Montrons que $I_{x_0} = fA$ où $f : x \mapsto x - x_0$.
 - Si $g \in fA$ alors $g(x_0) = 0$ donc $g \in I_{x_0}$.

— Si $g \in I_{x_0}$, on pose

$$u : x \mapsto \begin{cases} \frac{g(x)}{x - x_0} & \text{si } x \neq x_0 \\ 1 & \text{si } x = x_0 \end{cases}$$

On a bien $g = fu$ donc $g \in fA$

On a bien I_{x_0} qui est principal.

3. On se place maintenant dans $\mathcal{C}^0(\mathbf{R}, \mathbf{R})$ et on considère I_0 . Notre argument ne fonctionne plus car rien ne dit que u est encore continue. De fait I_0 n'est pas principal.

Supposons par l'absurde que $I_0 = (f)$ pour $f \in \mathcal{C}^0(\mathbf{R}, \mathbf{R})$. Maintenant $\sqrt{|f|} \in I_0$ donc il existe $g \in \mathcal{C}^0(\mathbf{R}, \mathbf{R})$ tel que $\sqrt{|f|} = fg$. On peut alors montrer qu'il existe un voisinage V de 0 tel que f est nul sur V . Si ce n'était pas vrai on pourrait trouver une suite (x_n) tendant vers 0 telle que $f(x_n) \neq 0$. Dans ce cas, on a

$$\sqrt{|f(x_n)|} = f(x_n)g(x_n) \Rightarrow g(x_n) = \pm \frac{1}{\sqrt{|f(x_n)|}}$$

Or $g(x_n) \rightarrow g(0)$ et l'autre partie de l'égalité n'est pas bornée.

On en déduit que f s'annule sur un voisinage de 0. Mais cela implique que toute fonction qui s'annule en 0, s'annule sur un voisinage ce qui est faux ($x \mapsto x$ par exemple). En conclusion I_0 n'est pas principal.

Exercice : On considère l'anneau $\mathbf{Z}[i\sqrt{5}]$ défini par

$$\mathbf{Z}[i\sqrt{5}] = \{a + bi\sqrt{5} ; (a, b) \in \mathbf{Z}^2\} \subset \mathbf{C}$$

Pour tout $z \in \mathbf{Z}[i\sqrt{5}]$ on pose $N(z) = z\bar{z}$.

1. Montrer que $\mathbf{Z}[i\sqrt{5}]$ est un sous-anneau de $\mathbf{C}$.
2. Montrer que pour tout $z \in \mathbf{Z}[i\sqrt{5}]$, $N(z) \in \mathbf{N}$. Montrer que plus que : $\forall (z, z') \in \mathbf{Z}[i\sqrt{5}]^2, N(zz') = N(z)N(z')$.
3. On considère $I = \{(1 + i\sqrt{5})u + 2v, (u, v) \in \mathbf{Z}[i\sqrt{5}]\}$. Justifier que c'est un idéal et qu'il n'est pas principal. On pourra calculer $N(1 + i\sqrt{5})$ et $N(2)$.

Proposition 8.19

Soit A et B deux anneaux. On suppose que A est commutatif. Soit $\varphi : A \rightarrow B$ un morphisme d'anneaux. Le noyau $\text{Ker}\varphi$ de φ est un idéal de A .

Démonstration :

— Il est clair que $\varphi(0_A) = 0_B$ donc $0_A \in \text{Ker}\varphi$ qui n'est pas vide.

Maintenant si x et y appartiennent à $\text{Ker}\varphi$,

$$\varphi(x - y) = \varphi(x) - \varphi(y) = 0.$$

Donc $x - y \in \text{Ker}\varphi$. De ce fait, $\text{Ker}\varphi$ est un sous-groupe de $(A, +)$.

— Soit $x \in \text{Ker}\varphi$ et $a \in A$, $\varphi(ax) = \varphi(a)\varphi(x) = 0$. Donc $ax \in \text{Ker}\varphi$.

On a bien que $\text{Ker}\varphi$ est un idéal de A . □

Remarques :

1. Ceci est aussi vrai dans le cas de morphismes d'algèbres $\varphi : \mathcal{A} \rightarrow \mathcal{B}$ si $\mathcal{A}$ est commutative car ce sont des morphismes d'anneaux.
2. On retrouve que I_x est un idéal car c'est le noyau de l'application

$$\begin{array}{ccc} \varphi_x : \mathcal{F}(X, \mathbf{R}) & \rightarrow & \mathbf{R} \\ f & \mapsto & f(x) \end{array}$$

3.2 Plus grand commun diviseur

Revoyons rapidement les chapitres d'arithmétique vus en première année. Nous utiliserons la notion d'idéal pour unifier l'arithmétique de $K[X]$ et celle de Z .

Définition 8.20 (Anneau intègre)

Soit A un anneau. Il est dit intègre si

- i) Il est commutatif
- ii) Pour tout $(a, b) \in A^2$, $ab = 0 \Rightarrow a = 0$ ou $b = 0$.

Remarque : On peut aussi considérer la contraposée de la deuxième partie.

$$\forall (a, b) \in A^2, a \neq 0 \text{ et } b \neq 0 \Rightarrow ab \neq 0$$

Définition 8.21 (Divisibilité)

Soit A un anneau intègre. Soit a et b dans A . On dit que a divise b ou que b est un multiple de a et on note $a|b$ s'il existe $c \in A$ tel que $b = ac$

Proposition 8.22

Soit A un anneau intègre. Soit a et b dans A .

$$a|b \iff bA \subset aA$$

Remarques :

1. On voit ici que la relation de divisibilité est juste la relation d'inclusion des idéaux engendrés.
2. De fait l'idéal aA est l'ensemble des multiples de a

Démonstration :

- $\Rightarrow$ On suppose que $a|b$. Il existe donc c tel que $b = ac$. Dès lors soit $x \in bA$, il peut s'écrire by avec $y \in A$, c'est-à-dire

$$x = by = acy$$

On a bien $x \in aA$ et donc $bA \subset aA$.

- $\Leftarrow$ On suppose que $bA \subset aA$. En particulier, $b \in bA$ donc $b \in aA$ ce qui signifie qu'il existe $c \in A$ tel que $b = ac$.

□

Proposition 8.23 (Éléments associés)

1. Soit $(m, m') \in \mathbb{Z}^2$,

$$m\mathbb{Z} = m'\mathbb{Z} \iff (m|m' \text{ et } m'|m) \iff m = \pm m'.$$

2. Soit $(P, Q) \in \mathbb{K}[X]^2$,

$$PK[X] = QK[X] \iff (P|Q \text{ et } Q|P) \iff \exists \alpha \in \mathbb{K}^*, P = \alpha Q.$$

Dans ces cas, les éléments sont dits associés.

Démonstration :

1. La première équivalence est évidente d'après la proposition précédente.

Pour la deuxième, le sens $\Leftarrow$ est évident.

Montrons le sens $\Rightarrow$. On suppose que $m|m'$ et $m'|m$. Il existe alors k, k' dans $\mathbb{Z}$ tels que $m = km'$ et $m' = k'm$. En particulier, $m = kk'm$. De ce fait, si $m = 0$ alors $m' = 0$ et donc $m' = \pm m$ et si $m \neq 0$ alors $kk' = 1$. Cela implique que k est inversible dans $\mathbb{Z}$. On en déduit que $k = \pm 1$.

2. La démonstration est identique. La seule différence est que les éléments inversibles de $\mathbb{K}[X]$ sont les polynômes constants non nuls.

□

Théorème 8.24

1. Tous les idéaux de $\mathbb{Z}$ sont principaux. De plus pour tout idéal I de $\mathbb{Z}$ il existe un unique $m \geq 0$ tel que $I = m\mathbb{Z}$.
2. Tous les idéaux de $\mathbb{K}[X]$ sont principaux. De plus pour tout idéal I de $\mathbb{K}[X]$ il existe un unique P unitaire (ou nul) tel que $I = PK[X]$.

Remarque : Avant de donner la preuve de ce théorème (qui s'appuie sur la division euclidienne dans les deux cas), notons que c'est ce fait qui permet de développer l'arithmétique. En clair, tout ce qui suit reste vrai dans un anneau principal.

Démonstration :

1. Soit I un idéal de $\mathbb{Z}$ que l'on suppose différent de $\{0\} = 0\mathbb{Z}$ et de $\mathbb{Z} = 1\mathbb{Z}$ en entier.

L'ensemble $P = I \cap \mathbb{N}^*$ est une partie non vide de $\mathbb{N}^*$. En effet soit x un élément non nul de I alors x et $-x$ appartient à I (car I est sous-groupe de $\mathbb{Z}$). Dès lors P a un plus petit élément que nous noterons m . Il est clair que $m\mathbb{Z} \subset I$. Montrons par l'absurde que $I \subset m\mathbb{Z}$. Soit x un élément de I qui n'est pas un multiple de m . On fait la division euclidienne de x par m : $x = qm + r \iff x - qm = r$. Or on a vu que $-qm \in I$ donc $r \in I$. Or comme $r \neq 0$ (x ne divise pas m) et $r < m$ c'est absurde.

Pour l'unicité il suffit de voir que si $m\mathbb{Z} = m'\mathbb{Z}$ alors m et m' sont associés et de ce fait $m = \pm m'$ et donc $m = m'$ s'ils sont tous les deux positifs.

2. Soit I un idéal de $\mathbb{K}[X]$. Là encore on suppose $I \neq \{0\}$ (qui est engendré par 0) et $I \neq \mathbb{K}[X]$ (qui est engendré par 1).

On peut donc considérer un polynôme P de degré minimal d parmi les polynômes non nuls

de I . Notons que $d \geq 1$ car si $d = 0$ cela signifie qu'il y a un polynôme constant $P = a$. De ce fait, $\frac{1}{a}P = 1 \in I$ et donc $I = K[X]$.

Montrons alors que $I = PK[X]$. On a clairement que $PK[X] \subset I$. Supposons par l'absurde que $I \subset PK[X]$ soit fausse. Il existe alors un polynôme Q non multiple de P dans I . On fait la division euclidienne :

$$Q = PS + R \text{ et } \deg(R) < d = \deg(P).$$

Là encore on en déduit que $R \in I$ ce qui absurde.

On a bien $I = PK[X]$.

Là encore si on a deux générateurs P et Q comme P et Q sont associés, il existe $\alpha \in K^*$ tel que $P = \alpha Q$. Il n'y a donc qu'un seul générateur unitaire.

□

Proposition 8.25

Soit A un anneau intègre et I, J deux idéaux. La somme $I + J = \{i + j, i \in I, j \in J\}$ est un idéal de A .

Démonstration :

- On sait que $0 \in I$ et $0 \in J$, donc $0 = 0 + 0 \in I + J$. De ce fait, $I + J \neq \emptyset$. Maintenant soit $(x, y) \in (I + J)^2$. Il existe $(x_1, y_1) \in I^2$ et $(x_2, y_2) \in J^2$ tels que

$$x = x_1 + x_2 \text{ et } y = y_1 + y_2$$

De ce fait,

$$x - y = \underbrace{x_1 - y_1}_{\in I} + \underbrace{x_2 - y_2}_{\in J}$$

On a donc $x - y \in I + J$. On en déduit que $I + J$ est un sous-groupe de $(A, +)$.

- Soit $x \in I + J$ et $a \in A$. Il existe donc $(x_1, x_2) \in I \times J$ tels que $x = x_1 + x_2$. On a alors

$$ax = ax_1 + ax_2 \in I + J$$

On a bien montré que $I + J$ était un idéal.

□

Définition 8.26

1. Soit $(a, b) \in \mathbb{Z}^2$, on appelle PGCD de a et de b et on l'on note $a \wedge b$ l'unique générateur positif de l'idéal $a\mathbb{Z} + b\mathbb{Z}$.
2. Soit $(P, Q) \in K[X]^2$, on appelle PGCD de P et de Q et on l'on note $P \wedge Q$ l'unique générateur unitaire (ou nul) de l'idéal $P\mathbb{Z} + Q\mathbb{Z}$.

Remarques :

1. Pour la suite nous nous intéresserons essentiellement à l'arithmétique de $K[X]$ mais tous les résultats restent vrais pour $\mathbb{Z}$.
2. On peut aussi définir la notion de PPCM en vérifiant que si I et J sont des idéaux alors $I \cap J$ est encore un idéal.

3. Si $P = Q = 0$, on a $PK[X] + QK[X] = 0$ donc $P \wedge Q = 0$.

Définition 8.27

Soit $(P, Q) \in K[X]^2$. Ils sont dit premiers entre eux si $P \wedge Q = 1$.

Proposition 8.28

Le PGCD défini ci-dessus coïncide avec la définition « classique » du PGCD. Cela signifie que les diviseurs communs de P et Q sont les diviseurs de $P \wedge Q$.

$$\forall T \in K[X], (T|P \text{ et } T|Q) \iff (T|P \wedge Q).$$

- $\Leftarrow$ On vérifie d'abord que $PK[X] \subset (PK[X] + QK[X]) = (P \wedge Q)K[X]$. On en déduit que $(P \wedge Q)|P$ et de même $(P \wedge Q)|Q$. De ce fait, par transitivité, si T divise $P \wedge Q$ il divise P et il divise Q .
- $\Rightarrow$ Comme $P \wedge Q \in (P \wedge Q)K[X] = PK[X] + QK[X]$, il existe des polynômes U et V tels que

$$P \wedge Q = PU + QV.$$

Maintenant si T divise P et Q , il divise $PU + QV = P \wedge Q$.

□

Corollaire 8.29 (Théorème de Bézout)

Soit $(P, Q) \in K[X]^2$.

1. Il existe $(U, V) \in K[X]^2$ tels que $PU + QV = P \wedge Q$.
2. $(P \wedge Q = 1) \iff \exists (U, V) \in K[X]^2, PU + QV = 1$

ATTENTION

De manière générale on sait que le PGCD de P et de Q peut s'écrire $PU + QV$, cependant ce n'est pas parce qu'il existe U et V tels que $PU + QV = C$ que C est le PGCD de P et Q . Cela n'est vrai que pour $C = 1$.

Démonstration :

1. Cela a été vu dans la démonstration précédente.
2. Le sens $\Rightarrow$ est l'énoncé précédent. Réciproquement, s'il existe $(U, V) \in K[X]^2$ tels que $PU + QV = 1$, on en déduit que $PK[X] + QK[X]$ est un idéal contenant 1, c'est donc $K[X]$ qui est engendré par 1. Cela permet d'obtenir $(P \wedge Q = 1)$.

□

Exercice : Soit P, Q deux polynômes tels que $P \wedge Q = 1$. Soit r, s deux entiers naturels non nuls, montrer que $P^r \wedge Q^s = 1$. En déduire que pour $\alpha \neq \beta$, $(X - \alpha)^r \wedge (X - \beta)^s = 1$.

On peut aussi rappeler le lemme de Gauss

Proposition 8.30 (Lemme de Gauss)

Soit P, Q, R trois polynômes. On suppose que P et Q sont premiers entre eux. On a alors

$$P|QR \Rightarrow P|R$$

Démonstration : On sait qu'il existe U et V tels que $1 = PU + QV$. De ce fait,

$$R = R.1 = RPU + RQV.$$

Maintenant $P|RPU$ et $P|RQV$ (par hypothèse) donc $P|R$. □

3.3 Algorithme d'Euclide

Rappelons que comme il a été vu en premier année, l'algorithme d'Euclide (via la division euclidienne) permet de calculer le PGCD (ne pas oublier de « normaliser » le résultat). De plus, l'algorithme d'Euclide augmenté permet de plus de trouver U et V tels que

$$PU + QV = (P \wedge Q)$$

Exemple : Calculons de PGCD de $P = X^5 + X^4 + 2X^3 + X^2 + X + 2$ et $Q = X^4 + 2X^3 + X^2 - 4$.

On a

$$\begin{aligned} X^5 + X^4 + 2X^3 + X^2 + X + 2 &= (X - 1) \times (X^4 + 2X^3 + X^2 - 4) + 3X^3 + 2X^2 + 5X^2 - 2 \\ X^4 + 2X^3 + X^2 - 4 &= \frac{1}{9}(3X + 4) \times (3X^3 + 2X^2 + 5X^2 - 2) - \frac{14}{9} \boxed{(X^2 + X + 2)} \\ 3X^3 + 2X^2 + 5X^2 - 2 &= (3X - 1) \times (X^2 + X + 2) + 0 \end{aligned}$$

Le PGCD est le dernier reste non nul normalisé. Ici, $P \wedge Q = X^2 + X + 2$.

Pour trouver U et V on « remonte » :

$$X^2 + X + 2 = -\frac{9}{14}(X^4 + 2X^3 + X^2 - 4) + \frac{1}{14}(3X + 4) \times (3X^3 + 2X^2 + 5X^2 - 2)$$

or

$$3X^3 + 2X^2 + 5X^2 - 2 = -(X - 1) \times (X^4 + 2X^3 + X^2 - 4) + (X^5 + X^4 + 2X^3 + X^2 + X + 2)$$

donc

$$X^2 + X + 2 = -\frac{1}{14}(3X^2 + X + 5)Q + \frac{1}{14}(3X + 4)P.$$

Pour calculer les coefficients de Bézout, il est souvent préférable d'utiliser l'algorithme d'Euclide étendu rappelons le principe dans $\mathbb{Z}$ (c'est identique dans $\mathbb{K}[X]$). On se donne deux entiers a et b et on cherche u, v tels que $au + bv = r$. Le principe est de construire des suites récurrentes (r_n) , (u_n) et (v_n) telles que pour tout entier n , $au_n + bv_n = r_n$. On initialise les suites par

$$r_0 = a ; u_0 = 1 ; v_0 = 0 \text{ et } r_1 = b ; u_1 = 0 ; v_1 = 1$$

Ensuite, pour tout entier $n \geq 1$, r_{n+1} est le reste de la division euclidienne de r_{n-1} par r_n : $r_{n+1} = r_{n-1} - q_n r_n$ où q_n est le quotient de r_{n-1} par r_n et on pose donc

$$u_{n+1} = u_{n-1} - q_n u_n \quad ; \quad v_{n+1} = v_{n-1} - q_n v_n$$

ce qui permet d'assurer que, comme $au_{n-1} + bv_{n-1} = r_{n-1}$ et $au_n + bv_n = r_n$ alors $au_{n+1} + bv_{n+1} = r_{n+1}$

On peut présenter les résultats dans un tableau par exemple pour $a = 32$ et $b = 23$

r	u	v	q	
32	1	0		
23	0	1		
9	1	-1	1	$(32 = 1 \times 23 + 9)$
5	-2	3	2	$(23 = 2 \times 9 + 5)$
4	3	-4	1	$(9 = 1 \times 5 + 4)$
1	-5	7	1	$(5 = 1 \times 4 + 1)$

3.4 Généralisation à n polynômes

Définition 8.31

Soit $P_1, \dots, P_n$ des polynômes, on appelle PGCD de $P_1, \dots, P_n$ et on note $P_1 \wedge \dots \wedge P_n$ l'unique polynôme unitaire tel que

$$P_1\mathbf{K}[X] + \dots + P_n\mathbf{K}[X] = (P_1 \wedge \dots \wedge P_n)\mathbf{K}[X]$$

Remarque : Comme précédemment, si $P_1 = P_2 = \dots = P_n = 0$ alors $P_1\mathbf{K}[X] + \dots + P_n\mathbf{K}[X] = 0$ et donc $P_1 \wedge \dots \wedge P_n = 0$ (qui n'est pas unitaire).

Proposition 8.32 (Associativité - Commutativité)

L'opération $\wedge$ est associative et commutative.

Démonstration : Cela découle du fait que la somme d'idéaux est associative et commutative.

□

Proposition 8.33

Soit $P_1, \dots, P_n$ des polynômes et Q un polynôme unitaire.

$$(QP_1) \wedge \dots \wedge (QP_n) = Q(P_1 \wedge \dots \wedge P_n)$$

Démonstration :

On a $QP_1\mathbf{K}[X] + \dots + QP_n\mathbf{K}[X] = Q(P_1\mathbf{K}[X] + \dots + P_n\mathbf{K}[X]) = Q(P_1 \wedge \dots \wedge P_n)\mathbf{K}[X]$. De plus $Q(P_1 \wedge \dots \wedge P_n)$ est unitaire car Q et $P_1 \wedge \dots \wedge P_n$ le sont. □

Proposition 8.34

Soit $P_1, \dots, P_n$ des polynômes. Il existe $(U_1, \dots, U_n)$ tels que

$$P_1U_1 + \dots + P_nU_n = P_1 \wedge \dots \wedge P_n$$

Démonstration : Par définition

□

Remarque : Pour calculer les coefficients de Bézout, pour trois polynômes P_1, P_2, P_3 . On calcule d'abord U_1 et U_2 tels que

$$P_1U_1 + P_2U_2 = P_1 \wedge P_2$$

Ensuite on calcule V et T tels que

$$(P_1 \wedge P_2)V + P_3T = ((P_1 \wedge P_2) \wedge P_3) = P_1 \wedge P_2 \wedge P_3$$

On a donc

$$P_1U_1V + P_2U_2V + P_3T = P_1 \wedge P_2 \wedge P_3$$

Pour plus de 3 polynômes il suffit de recommencer.

Définition 8.35 (Polynômes premiers entre eux)

Soit $P_1, \dots, P_n$ des polynômes. On dit qu'ils sont premiers entre eux (dans leur ensemble) si $P_1 \wedge \dots \wedge P_n = 1$.

Remarques :

1. Par définition, si $P_1, \dots, P_n$ sont premiers entre eux si et seulement s'il existe $U_1, \dots, U_n$ tels que

$$P_1U_1 + \dots + P_nU_n = 1$$

2. Soit $P_1, \dots, P_n$ des polynômes. S'ils sont deux à deux premiers entre eux, ils sont premiers entre eux (dans leur ensemble) car dans ce cas

$$\mathbf{K}[X] = P_1\mathbf{K}[X] + P_2\mathbf{K}[X] \subset P_1\mathbf{K}[X] + \dots + P_n\mathbf{K}[X]$$

La réciproque n'est pas vraie. On peut par exemple poser

$$P_1 = X(X - 1), P_2 = X(X - 2) \text{ et } P_3 = (X - 1)(X - 2)$$

On a alors

$$P_1 \wedge P_2 = X, P_1 \wedge P_3 = X - 1 \text{ et } P_2 \wedge P_3 = X - 2$$

Par contre

$$P_1 \wedge P_2 \wedge P_3 = X \wedge (X - 1)(X - 2) = 1$$

4 Polynômes annulateurs

4.1 Définitions

Définition 8.36 (Polynômes annulateurs)

1. Soit $u \in \mathcal{L}(E)$. Un polynôme P de $\mathbf{K}[X]$ est dit annulateur de u si $P(u) = 0_{\mathcal{L}(E)}$.
2. Soit $A \in \mathcal{M}_n(\mathbf{K})$. Un polynôme P de $\mathbf{K}[X]$ est dit annulateur de A si $P(A) = 0$.

Proposition 8.37

Soit $u \in \mathcal{L}(E)$. L'ensemble des polynômes annulateurs de u est un idéal de $\mathbf{K}[X]$ non réduit à $\{0\}$.

Démonstration : On remarque que pour $P \in \mathbf{K}[X]$,

$$P(u) = 0 \iff \Phi_u(P) = 0$$

où

$$\begin{array}{ccc} \Phi_u : \mathbf{K}[X] & \rightarrow & \mathcal{L}(E) \\ P & \mapsto & P(u) \end{array}$$

On voit donc que l'ensemble des polynômes annulateurs est le noyau d'un morphisme d'anneau (même d'algèbre) dont l'origine est un anneau commutatif. C'est donc un idéal de $\mathbf{K}[X]$.

Il reste à montrer qu'il n'est pas réduit à $\{0\}$. On sait que $\mathcal{L}(E)$ est de dimension n^2 , de ce fait la famille $\text{id}_E = u^0, u^1, u^2, \dots, u^{n^2}$ est liée car elle est de cardinal $n^2 + 1 > n^2$. On en déduit qu'il existe un combinaison linéaire non triviale $\sum_{i=0}^{n^2} \lambda_i u^i = 0_{\mathcal{L}(E)}$. On pose alors $P = \sum_{i=0}^{n^2} \lambda_i X^i$. C'est bien un polynôme annulateur non trivial. $\square$

Remarque : La même propriété est vraie pour les matrices.

ATTENTION

Cela n'est plus vrai si l'espace vectoriel n'est pas de dimension finie. Par exemple pour $E = \mathcal{C}^\infty(\mathbf{R}, \mathbf{R})$ et

$$\begin{array}{ccc} u : E & \rightarrow & E \\ f & \mapsto & f' \end{array}$$

Pour tout polynôme $P \in \mathbf{R}[X]$ non nul, on pose $P = \sum_{k=0}^d a_k X^k$. Alors $P(u) : f \mapsto \sum_{k=0}^d a_k f^{(k)}$.

Pour $f : x \mapsto e^{\alpha x}$ on a donc

$$P(u)(f) : x \mapsto P(\alpha) e^{\alpha x}$$

Il suffit de prendre α qui n'est pas une racine de P pour voir que $P(u)(f) \neq 0$ donc $P(u) \neq 0$.

4.2 Polynôme minimal

Définition 8.38 (Polynôme minimal)

Soit $u \in \mathcal{L}(E)$. On appelle polynôme minimal de u et on note μ_u ou π_u le générateur unitaire de l'idéal des polynômes annulateurs. C'est à dire que μ_u est défini par :

$$\forall P \in \mathbf{K}[X], P(u) = 0_{\mathcal{L}(E)} \iff \mu_u | P.$$

Remarques :

1. On définit de même le polynôme minimal μ_A ou π_A d'une matrice $A \in \mathcal{M}_n(\mathbf{K})$.
2. Le polynôme minimal d'un endomorphisme ou d'une matrice est le polynôme unitaire annulateur de plus bas degré.

Proposition 8.39

Soit $u \in \mathcal{L}(E)$ et $\mathcal{B}$ une base de E . On pose $A = \text{Mat}_{\mathcal{B}}(u)$. Les polynômes minimaux coïncident :

$$\mu_A = \mu_u.$$

Démonstration : Ce sont les mêmes idéaux annulateurs car pour $P \in \mathbf{K}[X]$, $\Phi_u(P) = 0_{\mathcal{L}(E)} \iff \Phi_A(P) = 0_{\mathcal{M}_n(\mathbf{K})}$. $\square$

Exemples :

1. Soit $A = \begin{pmatrix} 2 & 0 \\ 1 & 3 \end{pmatrix}$. On a $A^2 = \begin{pmatrix} 4 & 0 \\ 5 & 9 \end{pmatrix}$. On peut alors remarquer que

$$A^2 - 5A = \begin{pmatrix} -6 & 0 \\ 0 & -6 \end{pmatrix} = -6I_2.$$

On a donc $X^2 - 5X + 6$ qui est un polynôme annulateur. Maintenant, tout polynôme de degré inférieur ou égal à 2 n'annule pas A donc $\mu_A = X^2 - 5X + 6$.

2. Soit A est la matrice nulle de $\mathcal{M}(\mathbf{K})$ alors $\mu_A = X$. De manière générale, si $A = \lambda I_n$ alors $\mu_A = X - \lambda$.

3. Soit $A = \begin{pmatrix} 0 & 1 & 0 & \cdots & 0 \\ \vdots & \ddots & \ddots & \ddots & \vdots \\ \vdots & & \ddots & \ddots & 0 \\ \vdots & & & \ddots & 1 \\ 0 & \cdots & \cdots & \cdots & 0 \end{pmatrix} \in \mathcal{M}_n(\mathbf{K})$

On a vu que $A^n = 0_{\mathcal{M}_n(\mathbf{K})}$, donc $\mu_A | X^n$. On en déduit que $\mu_A = X^k$ pour $k \in \llbracket 1; n \rrbracket$. Or pour $k < n$, $A^k \neq 0$ donc $\mu_A = X^n$.

De manière générale si A est nilpotente et si p est son indice de nilpotence, $\mu_A = X^p$.

4. Soit $A = \begin{pmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 3 \end{pmatrix}$ On a de manière générale $A^p = \begin{pmatrix} (-1)^p & 0 & 0 \\ 0 & (-1)^p & 0 \\ 0 & 0 & 3^p \end{pmatrix}$ Cherchons un polynôme de degré au plus 2 annulant A . On pose $P = aX^2 + bX + c$. On a alors

$$P(A) = \begin{pmatrix} a - b + c & 0 & 0 \\ 0 & a - b + c & 0 \\ 0 & 0 & 9a + 3b + c \end{pmatrix}$$

On résout le système

$$\begin{cases} a - b + c = 0 \\ 9a + 3b + c = 0 \end{cases} \iff \begin{cases} a - b + c = 0 \\ 12b - 8c = 0 \end{cases} \iff \begin{cases} a = -\frac{1}{3}c \\ b = \frac{2}{3}c \end{cases}$$

On en déduit que $\mu_A = X^2 - 2X - 3 = (X + 1)(X - 3)$.

5. On regarde maintenant $A = \begin{pmatrix} -1 & 1 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 3 \end{pmatrix}$ On a de manière générale $A^p = \begin{pmatrix} (-1)^p & (-1)^{p-1}p & 0 \\ 0 & (-1)^p & 0 \\ 0 & 0 & 3^p \end{pmatrix}$

Cherchons un polynôme $P = \sum_{k=0}^d a_k X^k$ annulant A . On a alors

$$P(A) = \begin{pmatrix} P(-1) & \sum_{k=0}^d a_k (-1)^{k-1} k & 0 \\ 0 & P(-1) & 0 \\ 0 & 0 & P(3) \end{pmatrix} = \begin{pmatrix} P(-1) & P'(-1) & 0 \\ 0 & P(-1) & 0 \\ 0 & 0 & P(3) \end{pmatrix}$$

On en déduit que $P(A) = 0 \iff P(-1) = P'(-1) = P(3) = 0 \iff (X-3)(X+1)^2|P$. On en déduit que $\mu_A = (X+1)^2(X-3)$.

6. De manière générale, on suppose que A est diagonalisable. On sait que si B est semblable à A alors $\mu_A = \mu_B$. On se ramène donc à déterminer le polynôme minimal de

$$B = \begin{pmatrix} \lambda_1 & & & & \\ & \ddots & & & \\ & & \lambda_1 & & \\ & & & \ddots & \\ & & & & \lambda_p & \\ & & & & & \ddots \\ & & & & & & \lambda_p \end{pmatrix}$$

Maintenant, on a vu que si $P \in \mathbf{K}[X]$,

$$P(B) = \begin{pmatrix} P(\lambda_1) & & & & \\ & \ddots & & & \\ & & P(\lambda_1) & & \\ & & & \ddots & \\ & & & & P(\lambda_p) & \\ & & & & & \ddots \\ & & & & & & P(\lambda_p) \end{pmatrix}$$

De ce fait

$$P(B) = 0 \iff P(\lambda_1) = \dots = P(\lambda_p) = 0 \iff (X - \lambda_1) \dots (X - \lambda_p) | P.$$

Finalement $\mu_A = \mu_B = \prod_{\lambda \in \text{Sp}(A)} (X - \lambda)$.

Exercice classique à savoir faire : On se donne A une matrice diagonale par blocs :

$$A = \left(\begin{array}{c|c|c} A_1 & 0 & 0 \\ \hline 0 & \ddots & 0 \\ \hline 0 & 0 & A_p \end{array} \right)$$

Pour tout polynôme P on a alors

$$P(A) = \left(\begin{array}{c|c|c} P(A_1) & 0 & 0 \\ \hline 0 & \ddots & 0 \\ \hline 0 & 0 & P(A_p) \end{array} \right)$$

On en déduit que

$$P(A) = 0 \iff P(A_1) = \dots = P(A_p) = 0 \iff \forall i \in \llbracket 1; p \rrbracket, \mu_{A_i} | P.$$

Cela signifie que l'ensemble des polynômes annulateurs de A est

$$\mu_{A_1} \mathbf{K}[X] \cap \dots \cap \mu_{A_p} \mathbf{K}[X] = (\mu_{A_1} \vee \dots \vee \mu_{A_p}) \mathbf{K}[X]$$

Finalement est le PPCM des μ_{A_i} .

$$\mu_A = \mu_{A_1} \vee \dots \vee \mu_{A_p}$$

Proposition 8.40

Soit A et B deux matrices semblables, $\mu_A = \mu_B$.

Démonstration : Il suffit de vérifier que pour $P \in \mathbf{K}[X]$, $P(A) = 0$ si et seulement si $P(B) = 0$.

□

Exercice : Trouver deux matrices (de même taille) ayant le même polynôme annulateur mais qui ne sont pas semblables.

Proposition 8.41

Soit $u \in \mathcal{L}(E)$. On suppose que son polynôme minimal μ_u est de degré d alors la famille $(u^k)_{0 \leq k \leq d-1}$ est une base de $\mathbf{K}[u]$.

Démonstration : On a déjà vu que (par définition) la famille $(u^k)_{k \in \mathbf{N}}$ engendrait $\mathbf{K}[u]$ mais elle n'est pas libre. En effet si $\mu_u = \sum_{k=0}^d a_k X^k$ alors on a une combinaison linéaire non triviale

$$\sum_{k=0}^d a_k u^k = 0.$$

Posons $F = \text{Vect}(\text{id}_E, u, \dots, u^{d-1}) \subset \mathbf{K}[u]$ et montrons que $F = \mathbf{K}[u]$. Soit $k \in \mathbf{N}$, on fait la division euclidienne de X^k par μ_u :

$$X^k = Q \times \mu_u + S \text{ où } \deg(S) < d$$

De ce fait,

$$u^k = Q(u) \circ \mu_u(u) + S(u) = S(u) \in F$$

On a bien $F = \mathbf{K}[u]$ c'est-à-dire que la famille $(u^k)_{0 \leq k \leq d-1}$ engendre $\mathbf{K}[u]$.

Montrons maintenant que la famille $(u^k)_{0 \leq k \leq d-1}$ est libre. Supposons par l'absurde qu'elle est liée. Il existe alors une relation de dépendance linéaire

$$\sum_{k=0}^{d-1} b_k u^k = 0.$$

Si on pose $Q = \sum_{k=0}^{d-1} b_k X^k$, on a donc $Q(u) = 0$ ce qui contredit le caractère minimal de μ_u .

Finalement la famille $(u^k)_{0 \leq k \leq d-1}$ est bien une base de $\mathbf{K}[u]$.

□

ATTENTION

Il faut comprendre et connaître cette démonstration.

4.3 Polynômes annulateurs et valeurs propres

Nous allons maintenant relier les polynômes annulateurs (et donc le polynôme minimal) aux valeurs propres (et donc le polynôme caractéristique).

Proposition 8.42

Soit $u \in \mathcal{L}(E)$, $\lambda \in \mathbf{K}$ une valeur propre et $x \in E_\lambda(u)$ un vecteur propre associé à λ .
Soit $P \in \mathbf{K}[X]$,

$$P(u)(x) = P(\lambda)x$$

ATTENTION

Attention aux parenthèses pour écrire une relation qui a un sens.

Démonstration : Soit $k \in \mathbf{N}$, $u^k(x) = u^{k-1}(u(x)) = u^{k-1}(\lambda.x) = \lambda.u^{k-1}(x)$. A partir de là, par une récurrence immédiate on obtient que $u^k(x) = \lambda^k x$. Maintenant si $P = \sum_{k=0}^d a_k X^k$ on a $P(u) = \sum_{k=0}^d a_k u^k$ et donc

$$P(u)(x) = \sum_{k=0}^d a_k u^k(x) = \sum_{k=0}^d a_k \lambda^k x = P(\lambda).x$$

□

Corollaire 8.43

Soit $u \in \mathcal{L}(E)$, $\lambda \in \mathbf{K}$ une valeur propre. Si P est un polynôme annulateur de u , $P(\lambda) = 0$.

Démonstration : Il suffit de voir que pour un vecteur propre x non nul,

$$0 = P(u)(x) = P(\lambda).x$$

□

Remarque : En particulier, μ_u s'annule en toutes les valeurs propres de u . C'est-à-dire que si on note $Z(P)$ les racines d'un polynôme P ; $\text{Sp}(u) \subset Z(\mu_u)$.

4.4 Théorème de Cayley-Hamilton

Théorème 8.44 (Théorème de Cayley-Hamilton)

Soit $u \in \mathcal{L}(E)$. Le polynôme caractéristique χ_u annule u . C'est-à-dire que $\mu_u | \chi_u$

Démonstration : La démonstration est non exigible. Voir DM.

□

Remarque : On en déduit que $Z(\mu_u) \subset Z(\chi_u)$. Finalement $\text{Sp}(u) \subset Z(\mu_u) \subset Z(\chi_u) = \text{Sp}(u)$ donc

$$\text{Sp}(u) = Z(\mu_u) = Z(\chi_u).$$

Corollaire 8.45

Soit $u \in \mathcal{L}(E)$, si

$$\chi_u = \prod_{i=1}^p (X - \lambda_i)^{r_i}$$

alors

$$\mu_u = \prod_{i=1}^p (X - \lambda_i)^{s_i}$$

où pour tout $i \in \llbracket 1 ; p \rrbracket$, $0 \leq s_i \leq r_i$.

4.5 Calcul de puissances

Il est fréquent de devoir calculer les puissances d'une matrice. Voici deux méthodes

En utilisant une diagonalisation

Soit A une matrice diagonalisable. Il existe alors $P \in \text{GL}_n(\mathbb{K})$ telle que

$$B = P^{-1}AP = \begin{pmatrix} \lambda_1 & 0 & \cdots & 0 \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \cdots & 0 & \lambda_n \end{pmatrix}.$$

Maintenant pour tout entier k ,

$$A^k = (PBP^{-1})^k = PB^kP^{-1} = P \begin{pmatrix} \lambda_1^k & 0 & \cdots & 0 \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \cdots & 0 & \lambda_n^k \end{pmatrix} P^{-1}.$$

Cette méthode peut être assez longue (calcul de P et de P^{-1}).

En utilisant un polynôme annulateur

Soit $A \in \mathcal{M}_n(\mathbb{K})$. On suppose connaître un polynôme annulateur P de A . On peut bien évidemment prendre χ_A (théorème de Cayley-Hamilton). Il est cependant préférable de prendre μ_A qui est de degré moindre (mais souvent plus difficile à calculer).

Pour tout entier k , on peut effectuer la division euclidienne de X^k par P .

$$X^k = P \times Q + R \text{ avec } \deg(R) < d$$

où $d = \deg(P)$.

On en déduit que

$$A^k = P(A) \times Q(A) + R(A) = R(A).$$

Il reste donc à calculer le reste R de la division euclidienne.

- Si P est un polynôme scindé à racines simples : $P = \prod_{i=1}^d (X - \lambda_i)$

On sait que R est de degré strictement inférieur à d , or pour tout λ_i ($i \in \llbracket 1; d \rrbracket$), $\lambda_i^k = P(\lambda_i) \times Q(\lambda_i) + R(\lambda_i) = R(\lambda_i)$. On en déduit que R est le polynôme d'interpolation de Lagrange tel que $\forall i \in \llbracket 1; d \rrbracket, R(\lambda_i) = \lambda_i^k$.

On a donc (voir cours de sup)

$$R = \sum_{i=1}^d \lambda_i^k L_i \text{ où } L_i = \prod_{j \neq i} \frac{X - \lambda_j}{\lambda_i - \lambda_j}$$

- Si P n'a qu'une seule racine $P = (X - \lambda)^d$. Cela signifie que $N = A - \lambda I$ est une matrice nilpotente vérifiant $N^d = 0$.

Il suffit alors d'écrire $A = N + \lambda I$ et d'utiliser le binôme de Newton (car N et I commutent).

On en tire,

$$A^k = (N + \lambda I)^k = \sum_{i=0}^k \binom{k}{i} N^i \lambda^{k-i}$$

Maintenant dans la somme on peut s'arrêter pour $i = d - 1$ car pour $i \geq d$, $N^i = 0$.

- (A NE PAS FAIRE) Si $P = (X - \lambda)^{d-1}(X - \mu)$ avec $\mu \neq \lambda$.

La décomposition euclidienne s'écrit :

$$X^k = Q \times (X - \lambda)^{d-1}(X - \mu) + R$$

On en déduit que

$$\mu^k = R(\mu) \text{ et que } \lambda^k = R(\lambda).$$

On remarque alors que $X^k - R = Q \times (X - \lambda)^{d-1}(X - \mu)$ donc λ est une racine d'ordre $d - 1$ de $X^k - R$. Cela veut dire que pour tout $i \in \llbracket 0; d - 2 \rrbracket$, $(X^k - R)^{(i)}(\lambda) = 0$. C'est-à-dire :

$$R^{(i)}(\mu) = \frac{k!}{(k-i)!} \mu^{k-i}.$$

Pour conclure, on peut alors utiliser la formule de Taylor :

$$R = \sum_{i=0}^{d-1} a_i (X - \mu)^i \text{ où } a_i = \frac{R^{(i)}(\mu)}{i!}.$$

On obtient donc

$$R = \sum_{i=0}^{d-2} \frac{k!}{i!(k-i)!} \mu^{k-i} (X - \mu)^i + a_{d-1} (X - \mu)^{d-1}$$

Il suffit d'utiliser le fait que $\lambda^k = R(\lambda)$ pour déterminer a_{d-1} .

- (A NE PAS FAIRE) Si $P = (X - \lambda)^i (X - \mu)^{d-i}$ avec $\mu \neq \lambda$. En exercice - utiliser Bézout.

5 Critère de diagonalisabilité

On va établir une condition nécessaire et suffisante pour qu'un endomorphisme (resp. une matrice) soit diagonalisable.

5.1 Décomposition des noyaux

Exemple : Soit $u \in \mathcal{L}(E)$. On suppose que $P = X^2 - 3X + 2 = (X - 1)(X - 2)$ annule u . On a donc $u^2 - 3u + 2\text{id}_E = 0_{\mathcal{L}(E)}$.

On sait alors que les seules valeurs propres possibles sont 1 et 2 (racines de P).

On sait que $E_1(u)$ et $E_2(u)$ sont en somme directe, de plus

$$E = E_1(u) \oplus E_2(u).$$

En effet soit $x \in E$, on veut écrire

$$x = x_1 + x_2 \text{ avec } (x_1, x_2) \in E_1(u) \times E_2(u)$$

On procède par analyse synthèse.

- Analyse : On suppose que $x = x_1 + x_2$ avec $x_1 \in E_1(u)$ et $x_2 \in E_2(u)$. En appliquant u , on a $u(x) = u(x_1) + u(x_2) = x_1 + 2x_2$. En résolvant le système cela donne :

$$x_1 = 2x - u(x) \text{ et } x_2 = u(x) - x.$$

- Synthèse : On pose

$$x_1 = 2x - u(x) \text{ et } x_2 = u(x) - x.$$

On a bien $x_1 + x_2 = x$, de plus

$$u(x_1) = 2u(x) - u^2(x) = 2u(x) - (3u(x) - 2x) = 2x - u(x) = x_1$$

et

$$u(x_2) = u^2(x) - u(x) = (3u(x) - 2x) - u(x) = 2u(x) - 2x = 2x_2.$$

En conclusion,

$$E = E_1(u) \oplus E_2(u).$$

De plus, on peut voir que le projecteur sur E_1 parallèlement à E_2 est $\pi_1 = 2\text{id}_E - u$ et que le projecteur sur E_2 parallèlement à E_1 est $\pi_2 = -\text{id}_E + u$. Ce sont des polynômes en u .

Nous allons maintenant essayer de généraliser cet exemple.

Théorème 8.46 (Théorème de décomposition des noyaux)

Soit $u \in \mathcal{L}(E)$, soit $P_1, \dots, P_d$ des polynômes deux à deux premiers entre eux. On pose

$$P = \prod_{i=1}^d P_i.$$

$$\text{Ker}(P(u)) = \bigoplus_{i=1}^d \text{Ker}(P_i(u))$$

En particulier si P est un polynôme annulateur de u ,

$$E = \bigoplus_{i=1}^d \text{Ker}(P_i(u))$$

Bonus : Les projecteurs associés à cette décomposition sont des polynômes en u .

Remarques :

1. C'est bien une généralisation de l'exemple précédent en posant $P_1 = X - 1, P_2 = X - 2$ et $P = X^2 - 3X + 2$.
2. La condition deux à deux premiers entre eux est plus forte que premiers entre eux.

Démonstration : Traitons pour commencer le cas de deux polynômes avant de s'attaquer au cas général.

- Pour $d = 2$. On veut montrer que $\text{Ker}(P(u)) = \text{Ker}(P_1(u)) \oplus \text{Ker}(P_2(u))$. Comme P_1 et P_2 sont premiers entre eux, il existe des polynômes A, B tels que

$$P_1 A + P_2 B = 1$$

c'est-à-dire

$$P_1(u) \circ A(u) + P_2(u) \circ B(u) = \text{id}_E.$$

- Montrons que $\text{Ker}(P_1(u))$ et $\text{Ker}(P_2(u))$ sont en somme directe : Soit $x \in \text{Ker}(P_1(u)) \cap \text{Ker}(P_2(u))$. On sait que

$$x = \text{id}_E(x) = A(u)(P_1(u)(x)) + B(u)(P_2(u)(x)) = 0_E.$$

- Montrons que $\text{Ker}(P_1(u)) \oplus \text{Ker}(P_2(u)) \subset \text{Ker}(P(u))$.

Soit $x \in \text{Ker}(P_1(u)) \oplus \text{Ker}(P_2(u))$, il existe $(\alpha, \beta) \in \text{Ker}(P_1(u)) \times \text{Ker}(P_2(u))$ tels que $x = \alpha + \beta$. On a alors

$$P(u)(x) = P_2(u)(P_1(u)(\alpha)) + P_1(u)(P_2(u)(\beta)) = 0_E$$

On a bien $x \in \text{Ker}(P(u))$.

- Montrons réciproquement que $\text{Ker}(P(u)) \subset \text{Ker}(P_1(u)) \oplus \text{Ker}(P_2(u))$. Soit $x \in \text{Ker}(P(u))$, on a

$$x = \text{id}_E(x) = \underbrace{A(u)(P_1(u)(x))}_{\in \text{Ker}(P_2(u))} + \underbrace{B(u)(P_2(u)(x))}_{\in \text{Ker}(P_1(u))}.$$

Notons que le projecteur sur $\text{Ker}(P_1(u))$ est donc $(B \circ P_2)(u)$. C'est un polynôme en u . De même pour le projecteur sur $\text{Ker}(P_2(u))$.

- Passons au cas général. Pour tout $i \in \llbracket 1 ; d \rrbracket$, on pose $T_i = \prod_{j \neq i} P_j = \frac{P}{P_i}$.

Comme les polynômes $P_1, \dots, P_d$ sont deux à deux premiers entre eux alors $T_1, \dots, T_d$ sont premiers entre eux dans leur ensemble. En effet si D est un polynôme irréductible qui divise tous les T_i . Comme il divise T_1 , il divise le produit $P_2 \times \dots \times P_d$. Comme D est irréductible, il existe $i \in \llbracket 2 ; d \rrbracket$ tel que $D|P_i$. On sait de plus que $D|T_i$ donc il existe $j \in \llbracket 1 ; n \rrbracket \setminus \{i\}$ tel que $D|P_j$. Cela contredit le fait que P_i et P_j soient premiers entre eux.

Maintenant, comme les polynômes $T_1, \dots, T_d$ sont premiers entre eux, il existe des polynômes $A_1, \dots, A_d$ tels que

$$T_1 A_1 + \dots + T_d A_d = 1$$

c'est-à-dire

$$T_1(u) \circ A_1(u) + \dots + T_d(u) \circ A_d(u) = \text{id}_E.$$

- Montrons que $\text{Ker}(P_1(u)), \dots, \text{Ker}(P_d(u))$ sont en somme directe : Soit $(x_1, \dots, x_d) \in \prod_{i=1}^d \text{Ker}(P_i(u))$ tels que $x_1 + \dots + x_d = 0_E$.

Pour tout i, j dans $[[1; n]]$ avec $i \neq j$,

$$T_i(u)(x_j) = \left(\prod_{k \notin \{i, j\}} P_k(u) \right) (P_j(u)(x_j)) = 0_E$$

On en déduit que pour tout $j \in [[1; n]]$,

$$x_j = \text{id}_E(x_j) = \sum_{i=1}^d A_i(u)(T_i(u)(x_j)) = A_j(u)(T_j(u)(x_j)) = -A_j(u) \left(T_j(u) \left(\sum_{k \neq j} x_k \right) \right) = 0_E$$

— Montrons que $\bigoplus_{i=1}^d \text{Ker}(P_i(u)) \subset \text{Ker}(P(u))$.

Soit $x \in \bigoplus_{i=1}^d \text{Ker}(P_i(u))$, il existe $(x_1, \dots, x_d) \in \prod_{i=1}^d \text{Ker}(P_i(u))$ tels que $x = x_1 + \dots + x_d$.

On a alors

$$P(u)(x) = \sum_{i=1}^d P(u)(x_i) = \sum_{i=1}^d (T_i(u) \circ P_i(u))(x) = 0_E$$

On a bien $x \in \text{Ker}(P(u))$.

— Montrons réciproquement que $\text{Ker}(P(u)) \subset \bigoplus_{i=1}^d \text{Ker}(P_i(u))$. Soit $x \in \text{Ker}(P(u))$, on a

$$x = \text{id}_E(x) = \sum_{i=1}^d A_i(u)(T_i(u)(x))$$

Or, pour tout $i \in [[1; d]]$, $P_i(u)(A_i(u)(T_i(u)(x))) = A_i(u)(P(u)(x)) = 0_E$.

Cela signifie que $A_i(u)(T_i(u)(x)) \in \text{Ker}(P_i(u))$.

On a bien montré que $x \in \bigoplus_{i=1}^d \text{Ker}(P_i(u))$.

Notons que le projecteur sur $\text{Ker}(P_i(u))$ parallèlement à $\bigoplus_{j \neq i} \text{Ker}(P_j(u))$ est $(A_i \times T_i)(u)$. C'est un polynôme en u .

□

Exemple : Soit p un projecteur. On a $p^2 = p$. Le polynôme $P = X^2 - X = X(X - 1)$ est un polynôme annulateur. On en déduit que

$$E = \text{Ker}(p) \oplus \text{Ker}(p - \text{Id}_E)$$

Ce que l'on savait déjà.

De même pour une symétrie.

5.2 Polynômes annulateurs et diagonalisabilité

Nous allons établir un nouveau critère de diagonalisabilité d'une matrice / d'un endomorphisme. Rappelons que l'on a déjà vu qu'un endomorphisme qui vérifie que $\sum_{\lambda \in \text{Sp}(u)} \dim E_\lambda(u) = \dim E$ est diagonalisable. Ce critère est difficile à utiliser car il nécessite de calculer la dimension des espaces propres. On s'en sert essentiellement dans le cas où le nombre de valeurs propres distinctes est égal à la dimension de l'espace (dans quel cas tous les espaces propres sont de dimension 1).

Théorème 8.47

Soit $u \in \mathcal{L}(E)$. Les assertions suivantes sont équivalentes.

- i) L'endomorphisme u est diagonalisable.
- ii) Il existe un polynôme scindé à racines simples qui annule u .
- iii) Le polynôme minimal de u , μ_u est scindé à racines simples.

ATTENTION

Il se peut que le polynôme caractéristique ne soit pas à racines simples. Par exemple pour l'identité. On a $\mu_{\text{id}} = X - 1$ et $\chi_{\text{id}} = (X - 1)^n$.

Remarque : Il y a un énoncé analogue pour les matrices.

Démonstration :

- $i) \Rightarrow ii)$ On suppose que u est diagonalisable. Soit $\mathcal{B}$ telle que la matrice de u dans la base $\mathcal{B}$ soit diagonale :

$$A = \text{Mat}_{\mathcal{B}}(u) = \begin{pmatrix} \lambda_1 & & & & \\ & \ddots & & & \\ & & \lambda_1 & & \\ & & & \ddots & \\ & & & & \lambda_p & \\ & & & & & \ddots \\ & & & & & & \lambda_p \end{pmatrix}$$

On pose $P = \prod_{i=1}^p (X - \lambda_i)$ qui est scindé à racines simples et on a $P(A) = 0$ donc $P(u) = 0_{\mathcal{L}(E)}$.

- $ii) \Rightarrow iii)$ On suppose qu'il existe un polynôme scindé à racines simples P qui annule u . Par définition du polynôme minimal, $\mu_u | P$. Donc μ_u est scindé à racines simples.
- $iii) \Rightarrow i)$ On suppose que μ_u est scindé à racines simples. On pose

$$\mu_u = \prod_{i=1}^p (X - \lambda_i).$$

Il suffit d'utiliser le lemme de décomposition des noyaux :

$$E = \text{Ker}(\mu_u(u)) = \bigoplus_{i=1}^p \text{Ker}(u - \lambda_i \text{id}_E)$$

Comme E est somme directe des sous-espaces propres de u , u est diagonalisable. □

Exemples :

1. Soit p un projecteur. Comme $p^2 = p$, le polynôme $X^2 - X = x(X - 1)$ est un polynôme annulateur de p . Cela montre que p est diagonalisable.

De même une symétrie est diagonalisable car elle est annulée par le polynôme $X^2 - 1 = (X - 1)(X + 1)$ qui est scindé à racines simples.

2. Soit $A = \begin{pmatrix} 1 & \cdots & 1 \\ \vdots & & \vdots \\ 1 & \cdots & 1 \end{pmatrix} \in \mathcal{M}_n(\mathbf{K})$. On a vu que $A^2 = nA$. Cela montre que le polynôme $X^2 - 2 - nX = X(X - n)$ annule A . La matrice A est donc diagonalisable.
3. Soit A une matrice nilpotente. On sait que son polynôme minimal est de la forme $\pi_A = X^p$ où p est l'indice de nilpotence. Si $p \neq 1$ (c'est-à-dire si $A \neq 0$) alors A n'est pas diagonalisable.

Proposition 8.48

Soit $u \in \mathcal{L}(E)$ et F un sous-espace vectoriel u -stable. On note $\tilde{u}$ l'endomorphisme de F induit par u . Le polynôme minimal de $\tilde{u}$ divise celui de u :

$$\mu_{\tilde{u}} \mid \mu_u.$$

Démonstration : Il suffit de voir que $\mu_u(\tilde{u}) = 0$. De ce fait, μ_u est un polynôme annulateur de $\tilde{u}$ donc un multiple de $\mu_{\tilde{u}}$. $\square$

Remarque : Cela a déjà été étudié lors du calcul du polynôme minimal d'une matrice diagonale par blocs. Rappelons que si $E = \bigoplus_{i=1}^p F_i$ et que les sous-espaces vectoriels F_i sont stables par u . Alors, en notant μ_i le polynôme minimal de l'endomorphisme induit par u sur F_i , $\mu = \mu_1 \vee \cdots \vee \mu_p$.

Corollaire 8.49 (Diagonalisabilité d'un endomorphisme induit)

Soit $u \in \mathcal{L}(E)$ et F un sous-espace vectoriel u -stable. On note $\tilde{u}$ l'endomorphisme de F induit par u . Si u est diagonalisable alors $\tilde{u}$ aussi.

Démonstration : Comme u est diagonalisable, μ_u est scindé à racines simples. Maintenant, $\mu_{\tilde{u}} \mid \mu_u$ donc $\mu_{\tilde{u}}$ est aussi scindé et à racines simples. De ce fait, $\tilde{u}$ est diagonalisable. $\square$

5.3 Endomorphismes annulés par un polynôme scindé

Définition 8.50 (Sous-espaces caractéristiques)

Soit $u \in \mathcal{L}(E)$. On suppose que son polynôme caractéristique est scindé et on note $\chi_u = \prod_{i=1}^p (X - \lambda_i)^{r_i}$. Pour tout $i \in \llbracket 1; p \rrbracket$ on appelle sous-espace caractéristique et on note $C_{\lambda_i}(u)$ le sous-espace vectoriel $C_{\lambda_i}(u) = \text{Ker}(u - \lambda_i \text{id}_E)^{r_i}$.

Remarques :

1. Les sous-espaces caractéristiques sont des « super » espaces propres. On a $E_{\lambda}(u) \subset C_{\lambda}(u)$.
2. Pour définir les sous-espaces caractéristiques, on peut aussi prendre les multiplicités dans le polynôme minimal (ou dans tout polynôme annulateur scindé). En effet si on note $\mu_u = \prod_{i=1}^p (X - \lambda_i)^{s_i}$ et si on considère $i \in \llbracket 1; p \rrbracket$. On sait d'après le lemme des noyaux que

$$E = \text{Ker}(u - \lambda_i \text{id}_E)^{s_i} \oplus G$$

où $G = \bigoplus_{j \neq i} \text{Ker}(u - \lambda_j \text{id}_E)^{s_j}$. Prenons $t \geq s_i$ et $x \in \text{Ker}(u - \lambda_i \text{id}_E)^t$. On peut décomposer $t = \alpha + \beta$ où $\alpha \in \text{Ker}(u - \lambda_i \text{id}_E)^{s_i}$ et $\beta \in G$. On a alors

$$0 = (u - \lambda_i \text{id}_E)^t(x) = 0 + (u - \lambda_i \text{id}_E)^t(\beta)$$

Or l'endomorphisme induit par $u - \lambda_i \text{id}_E$ sur G est inversible donc cela implique que $\beta = 0$.

On a montré que $\forall t \geq s_i, \text{Ker}(u - \lambda_i \text{id}_E)^{s_i} = \text{Ker}(u - \lambda_i \text{id}_E)^t$.

Théorème 8.51

Soit $u \in \mathcal{L}(E)$. On suppose que son polynôme caractéristique est scindé. On note alors $C_1(u), \dots, C_p(u)$ les sous-espaces caractéristiques de u .

1. On a $E = \bigoplus_{i=1}^p C_i(u)$
2. Les sous-espaces vectoriels $C_i(u)$ sont stables par u .
3. Soit $i \in \llbracket 1; p \rrbracket$, l'endomorphisme $\tilde{u}_i$ induit par u sur $C_i(u)$ peut s'écrire comme somme de l'homothétie de rapport λ_i et d'un endomorphisme nilpotent.

Remarques :

1. Dans le cas où $\mathbf{K} = \mathbf{C}$ on a toujours χ_u scindé.
2. Cela signifie que si on prend une base adaptée à la décomposition $E = \bigoplus_{i=1}^p C_i(u)$, la matrice de u est diagonale par blocs et les différents blocs sont de la forme $\lambda_i I + N_i$ où N_i est une matrice nilpotente.

Démonstration :

1. C'est le lemme de décomposition des noyaux
2. Soit $i \in \llbracket 1; p \rrbracket$, il est clair que u commute avec $(u - \lambda_i \text{id}_E)^{r_i}$, donc $C_i = \text{Ker}(u - \lambda_i \text{id}_E)^{r_i}$ est stable par u .
3. Soit $i \in \llbracket 1; p \rrbracket$, on peut écrire l'endomorphisme $\tilde{u}_i$ induit par u sur C_i ,

$$\tilde{u}_i = \lambda_i \text{id}_{C_i} + v_i \text{ où } v_i = \tilde{u}_i - \lambda_i \text{id}_{C_i}.$$

C'est-à-dire que v_i est l'endomorphisme induit par $u - \lambda_i \text{id}_E$ sur C_i . Par définition de C_i on a donc $v_i^{r_i} = 0$. Cela signifie que v_i est nilpotent.

□

Remarque : Cela s'applique aussi aux matrices. En particulier, si A est une matrice annulée par un polynôme scindé $P = \prod_{i=1}^p (X - \lambda_i)^{r_i}$, il existe une matrice B semblable à A de la forme

$$\left(\begin{array}{c|c|c} \lambda_1 I + N_1 & & \\ \hline & \ddots & \\ \hline & & \lambda_p I + N_p \end{array} \right).$$

De plus on sait que toutes les matrices N_i étant nilpotentes, elles sont semblables à des matrices

strictement triangulaires supérieures. C'est-à-dire que A est semblable à

$$\left(\begin{array}{cccc|c|c} \lambda_1 & * & \cdots & * & & \\ 0 & \ddots & \ddots & \vdots & & \\ \vdots & & \ddots & * & & \\ 0 & \cdots & 0 & \lambda_1 & & \\ \hline & & & & \ddots & \\ \hline & & & & & \lambda_p & * & \cdots & * \\ & & & & & 0 & \ddots & \ddots & \vdots \\ & & & & & \vdots & & \ddots & * \\ & & & & & 0 & \cdots & 0 & \lambda_p \end{array} \right).$$

Exemple : Soit $u \in \mathcal{L}(E)$. Supposons qu'il admette un polynôme annulateur scindé $P = \prod_{i=1}^d (X - \lambda_i)^{r_i}$. On pose alors $C_{\lambda_i}(u) = \text{Ker}(u - \lambda_i \text{Id}_E)^{r_i}$. On a

$$E = \bigoplus_{i=1}^d C_{\lambda_i}(u).$$

Par exemple soit

$$A = \begin{pmatrix} -8 & -3 & -3 & 1 \\ 6 & 3 & 2 & -1 \\ 26 & 7 & 10 & -2 \\ 0 & 0 & 0 & 2 \end{pmatrix}$$

Le polynôme caractéristique est

$$\chi_A = X^4 - 7X^3 + 18X^2 - 20X + 8 = (X - 2)^3(X - 1)$$

On a alors

$$A - I_4 = \begin{pmatrix} -9 & -3 & -3 & 1 \\ 6 & 2 & 2 & -1 \\ 26 & 7 & 9 & -2 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

On en déduit que $\text{Ker}(A - I_4) = \text{Vect} \begin{pmatrix} -2 \\ 1 \\ 5 \\ 0 \end{pmatrix}$. Maintenant,

$$A - 2I_4 = \begin{pmatrix} -10 & -3 & -3 & 1 \\ 6 & 1 & 2 & -1 \\ 26 & 7 & 8 & -2 \\ 0 & 0 & 0 & 0 \end{pmatrix}$$

On a donc que $\text{Ker}(A - 2I_4) = \text{Vect} \begin{pmatrix} -3 \\ 2 \\ 8 \\ 0 \end{pmatrix}$. On en déduit que A n'est pas diagonalisable. Par

contre

$$(A - 2I_4)^3 = \begin{pmatrix} -4 & -6 & 0 & -2 \\ 2 & 3 & 0 & 1 \\ 10 & 15 & 0 & 5 \\ 0 & 0 & 0 & 0 \end{pmatrix}$$

est une matrice de rang 1. On obtient que $\text{Ker}(A - 2I_4)^3 = \text{Vect}\left(\begin{pmatrix} -3 \\ 2 \\ 8 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 0 \\ 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 1 \\ 0 \\ 0 \\ -2 \end{pmatrix}\right)$. On a bien $\mathbf{K}^4 = C_1(A) \oplus C_2(A)$

6 Compléments et applications

Les résultats de ce paragraphe ne sont pas au programme. Ce sont des résultats classiques qu'il est bon de connaître.

6.1 Influence du corps de base

Quand on considère $A \in \mathcal{M}_n(\mathbf{R})$ et que l'on parle de son endomorphisme canoniquement associé, cela peut-être ambiguë car A appartient aussi à $\mathcal{M}_n(\mathbf{C})$. On peut vouloir considérer l'endomorphisme de $\mathbf{R}^n$ ou celui de $\mathbf{C}^n$.

Cela change notamment le spectre. On peut donc être amené à préciser $\text{Sp}_{\mathbf{R}}(A)$ ou $\text{Sp}_{\mathbf{C}}(A)$.

Notons cependant que le polynôme caractéristique ne dépend pas du corps et donc

$$\text{Sp}_{\mathbf{R}}(A) = \text{Sp}_{\mathbf{C}}(A) \cap \mathbf{R}.$$

Proposition 8.52

Soit $(A, B) \in \mathcal{M}_n(\mathbf{R})$ si elles sont semblables dans $\mathcal{M}_n(\mathbf{C})$, elles le sont dans $\mathcal{M}_n(\mathbf{R})$.

Remarque : Ce n'est pas évident que le fait qu'il existe $P \in \text{GL}_n(\mathbf{C})$ telle que $B = P^{-1}AP$ alors il existe $Q \in \text{GL}_n(\mathbf{R})$ telle que $Q^{-1}AQ = B$.

Démonstration : On suppose qu'il existe $P \in \text{GL}_n(\mathbf{C})$ telle que $B = P^{-1}AP$, ce qui implique $PB = AP$.

Maintenant comme A et B sont réelles, On en déduit que

$$\begin{cases} \Re(P)B = A\Re(P) \\ \Im(P)B = A\Im(P) \end{cases}$$

Donc pour tout $t \in \mathbf{R}$, $QB = AQ$ où $Q = \Re(P) + t\Im(P)$.

Il reste à montrer que l'on peut trouver t tel que Q soit inversible. Pour cela on considère $\det(Q)$ qui est un polynôme en t qui est non nul car $\det(P) \neq 0$. Il a un nombre fini de racines donc on peut trouver t tel que $Q \in \text{GL}_n(\mathbf{R})$. $\square$

6.2 Calculs de puissances

Soit $A \in \mathcal{M}_n(\mathbf{K})$. On suppose que A est annulée par un polynôme scindé : $P = \prod_{i=1}^d (X - \lambda_i)^{r_i}$. Il existe alors $P \in \text{GL}_n(\mathbf{K})$ telle que $B = P^{-1}AP$ soit de la forme :

$$B = \left(\begin{array}{c|c|c} \lambda_1 I + N_1 & & \\ \hline & \ddots & \\ \hline & & \lambda_p I + N_p \end{array} \right).$$

Maintenant les puissances de B se calculent par blocs

$$B^k = \left(\begin{array}{c|c|c} (\lambda_1 I + N_1)^k & & \\ \hline & \ddots & \\ \hline & & (\lambda_p I + N_p)^k \end{array} \right).$$

Mais comme $N_1, \dots, N_p$ sont nilpotentes, les puissances se calculent aisément par le binôme de Newton :

$$(\lambda_i I + N_i)^k = \sum_{s=0}^k \binom{k}{s} \lambda_i^{k-s} N_i^s$$

Le choix des matrices nilpotentes strictement triangulaires simplifie (un peu) le calcul.

6.3 Commutant

Soit $u \in \mathcal{L}(E)$ (ou A une matrice). On appelle commutant de u l'ensemble des endomorphismes de E qui commutent avec u .

$$\mathcal{C}(u) = \{v \in \mathcal{L}(E) \mid u \circ v = v \circ u\}.$$

Proposition 8.53 (Structure du commutant)

1. Le commutant $\mathcal{C}(u)$ est une sous-algèbre de $\mathcal{L}(E)$.
2. L'ensemble des polynômes en u appartiennent au commutant :

$$\mathbf{K}[u] \subset \mathcal{C}(u)$$

Démonstration : A écrire

□

Avant de continuer, commençons par rappeler que l'on a vu que si u, v sont deux endomorphismes de E qui commutent alors les espaces propres de l'un sont stable par l'autre.

Proposition 8.54

Soit u un endomorphisme diagonalisable. Soit $v \in \mathcal{L}(E)$. L'endomorphisme v commute avec u si et seulement s'il stabilise tous les espaces propres de u .

Démonstration : On note $\text{Sp}(u) = \{\lambda_1, \dots, \lambda_p\}$ et $\mathcal{B}$ une base telle que

$$A = \text{Mat}_{\mathcal{B}}(u) = \begin{pmatrix} \lambda_1 & & & & \\ & \ddots & & & \\ & & \lambda_1 & & \\ & & & \ddots & \\ & & & & \lambda_p & \\ & & & & & \ddots \\ & & & & & & \lambda_p \end{pmatrix}$$

– $\Rightarrow$ On suppose que $v \in \mathcal{C}(u)$. On sait alors que les espaces propres de u sont stables par v .

- $\boxed{\Leftarrow}$ On suppose réciproquement que les espaces propres de u sont stables par v . Cela signifie que $\text{Mat}_{\mathcal{B}}(v)$ est diagonale par blocs :

$$B = \text{Mat}_{\mathcal{B}}(v) = \begin{pmatrix} A_1 & & \\ & \ddots & \\ & & A_p \end{pmatrix}$$

La formule du produit matriciel par blocs nous donne alors que $AB = BA$ et donc $v \in \mathcal{C}(u)$.

□

Remarques :

1. Cela signifie que la dimension (en tant qu'espace vectoriel) de $\mathcal{C}(u)$ est $\sum_{i=1}^p (\dim E_{\lambda_i}(u))^2$
2. Dans le cas particulier où u a n valeurs propres distinctes, on a alors $\mu_u = \prod_{i=1}^n (X - \lambda_i)$. De ce fait $\dim \mathbf{K}[u] = n$ car $(\text{id}, u, u^2, \dots, u^{n-1})$ est libre. Cependant, la formule ci-dessus nous donne que $\dim \mathcal{C}(u) = n$. On a donc

$$\mathcal{C}(u) = \mathbf{K}[u].$$

Dans ce cas « générique », il n'y a que les polynômes en u qui commutent avec u .

6.4 Décomposition de Dunford

Les résultats précédents permettent de décomposer un endomorphisme en une partie diagonalisable et une partie nilpotente. C'est une décomposition classique et utile mais qui est explicitement hors programme.

Soit $u \in \mathcal{L}(E)$. On suppose que u est annulé par un polynôme scindé (ce qui est toujours vrai si $\mathbf{K} = \mathbf{C}$). On reprend les notations du théorème de décomposition. On a $E = \bigoplus_{i=1}^p C_i$ et $\check{u}_i$ l'endomorphisme induit par u sur C_i (qui est stable par u) peut s'écrire

$$\check{u}_i = \lambda_i \text{id}_{C_i} + n_i$$

où n_i est nilpotent.

On pose alors d l'endomorphisme de E tel que pour tout i , l'endomorphisme induit par d sur C_i soit $\lambda_i \text{id}_{C_i}$. De même on pose n l'endomorphisme de E tel que pour tout i , l'endomorphisme induit par n sur C_i soit n_i .

Il est alors aisé de vérifier que :

- l'endomorphisme d est diagonalisable,
- l'endomorphisme n est nilpotent,
- on a $u = d + n$,
- les endomorphismes d et n commutent.
- l'endomorphisme d est un polynôme en u et donc $n = u - d$ aussi.

On peut montrer (mais c'est un peu plus difficile) qu'il y a unicité d'une telle décomposition.

Notons de plus que l'on sait que pour tout i compris entre 1 et p , le projecteur π_i sur C_i parallèlement à $\bigoplus_{\substack{1 \leq j \leq p \\ j \neq i}} C_j$ est un polynôme en u . Cela implique que

$$d = \sum_{i=1}^p \lambda_i \pi_i \text{ et } n = u - d$$

sont aussi des polynômes en u .

6.5 Racines carrées

On se donne $u \in \mathcal{L}(E)$ et on cherche les endomorphismes v tels que $v^2 = u$.

On remarque d'abord que comme v est un polynôme en u , les endomorphismes u et v commutent. On en déduit que v laisse stable des espaces propres de u .

Supposons que u soit diagonalisable

Soit $\mathcal{B}$ tel que

$$\text{Mat}_{\mathcal{B}}(u) = \begin{pmatrix} \lambda_1 & & & & \\ & \ddots & & & \\ & & \lambda_1 & & \\ & & & \ddots & \\ & & & & \lambda_p & \\ & & & & & \ddots \\ & & & & & & \lambda_p \end{pmatrix}$$

On a

$$B = \text{Mat}_{\mathcal{B}}(v) = \begin{pmatrix} A_1 & & \\ \hline & \ddots & \\ \hline & & A_p \end{pmatrix}$$

avec pour $i \in [1; p]$, $A_i^2 = \lambda_i I$

- Si on suppose que les espaces propres sont tous de dimension 1. Alors A et B sont diagonales. On est ramené à résoudre pour tout i l'équation $x_i^2 = \lambda_i$. On en déduit que
 - Si tous les λ_i sont strictement positifs ($\mathbf{K} = \mathbf{R}$) ou non nuls ($\mathbf{K} = \mathbf{C}$), il y a 2^n solutions.
 - S'il y a un λ_i nul, il y a 2^{n-1} solutions.
 - S'il y a un λ_i strictement négatif ($\mathbf{K} = \mathbf{R}$) il n'y a pas de solutions.
- Dans le cas général on essaye de résoudre $A^2 = \lambda I$.
 - Si $\mathbf{K} = \mathbf{C}$ (ou si $\lambda > 0$) cela se ramène à $\left(\frac{1}{\delta} A\right)^2 = I$ où $\delta^2 = \lambda$. C'est-à-dire les symétries. Il y en a une infinité.
 - Si $\mathbf{K} = \mathbf{R}$ et $\lambda < 0$. C'est plus compliqué mais pas nécessairement impossible. Par exemple la rotation d'angle $\pi/4$ de matrice $B = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}$ vérifie $B^2 = -I$.

6.6 Sous-espaces stables

Exercice [480]

$$\text{Soit } A = \begin{pmatrix} 2 & 1 & 1 \\ 1 & 2 & 1 \\ 0 & 0 & 3 \end{pmatrix}.$$

On cherche les sous-espaces vectoriels stables par A .

On commence en calculant le polynôme caractéristique $\chi_A = X^3 - 7X^2 + 15X - 9 = (X - 1)(X - 3)^2$. On a aussi $\mu_A = \chi_A$ et A n'est pas diagonalisable.

Soit F un sous-espace vectoriel stable par A .

- Si $\dim F = 0$ alors $F = \{0\}$.

- Si $\dim F = 3$ alors $F = \mathbf{R}^3$.
- Si $\dim F = 1$ alors F est une droite vectorielle qui doit être inclus dans un espace propre. On en déduit $F = E_1(A)$ ou $F = E_3(A)$.
- Si $\dim F = 2$. On note a l'endomorphisme associé à A et $u = \check{a}$ l'endomorphisme induit par a sur F . On sait que $\chi_u | \chi_a$ et donc $\chi_u = (X - 3)^2$ ou $\chi_u = (X - 1)(X - 3)$.
 - Si $\chi_u = (X - 3)^2$. D'après Cayley-Hamilton, $(u - 3\text{id})^2 = 0$. On en déduit que $F \subset C_3(a) = \text{Ker}(a - 3\text{id})^2$. Pour des raisons de dimension, on a égalité.
 - Si $\chi_u = (X - 1)(X - 3)$. On doit avoir un vecteur propre associé à chacune des valeurs propres dans F donc $F = E_1(A) \oplus E_3(A)$.

Espaces vectoriels normés I

1	Normes	240
1.1	Définition	240
1.2	Exemples de normes	241
1.3	Distance associée à une norme	248
1.4	Boules	249
1.5	Parties, suites et fonctions bornées	251
1.6	Normes équivalentes	252
2	Exemple de normes équivalentes	253
2.1	Sur $\mathbb{K}^n$	253
2.2	Sur les espaces de fonctions	256
3	Suites à valeurs dans un espace vectoriel normé	256
3.1	Convergence	257
3.2	Propriétés	260
3.3	Suites extraites et valeurs d'adhérence	262
4	Séries à valeurs dans un espace vectoriel de dimension finie	264
4.1	Généralités	264
4.2	Série géométrique de matrices	265
5	Applications linéaires lipschitziennes	266
5.1	Définitions	266

Dans ce chapitre $\mathbb{K}$ désigne $\mathbb{R}$ ou $\mathbb{C}$.

On sait comment caractériser le fait qu'une suite de réels ou de complexes converge. On sait de même comment définir le fait qu'une fonction définie sur $\mathbb{R}$ ou $\mathbb{C}$ soit continue. On veut en faire de même avec les espaces vectoriels. Cela peut permettre par exemple de calculer l'exponentielle d'une matrice par

$$\exp(A) = \lim_{n \rightarrow \infty} \sum_{k=0}^n \frac{A^k}{k!}.$$

Pour cela nous aurons besoin de généraliser la notion de valeur absolue (ou module) qui permet de donner un sens à la distance entre deux vecteurs.

1 Normes

1.1 Définition

Dans ce chapitre, E désigne un espace vectoriel sur $\mathbf{K}$.

Définition 9.1 (Norme)

On appelle norme sur E une application $N : E \rightarrow \mathbf{R}$ telle que

- i) $\forall x \in E, N(x) \geq 0$
- ii) $\forall x \in E, x = 0 \iff N(x) = 0$
- iii) $\forall x \in E, \forall \lambda \in \mathbf{K}, N(\lambda.x) = |\lambda|N(x)$
- iv) $\forall (x, y) \in E^2, N(x + y) \leq N(x) + N(y)$ (inégalité triangulaire)

Terminologie :

1. On appelle espace vectoriel normé, un couple (E, N) où E est un espace vectoriel et N une norme sur E .
2. Dans un espace vectoriel normé, un vecteur est dit unitaire s'il est de norme 1. En particulier si $x \neq 0$, $\frac{1}{N(x)}.x$ est unitaire.

Remarques :

1. Dans l'hypothèse ii), c'est le sens $\Leftarrow$ qu'il faut prouver car le sens $\Rightarrow$ découle de iii).
2. La première hypothèse est parfois donnée en disant que N est à valeurs dans $\mathbf{R}_+$.

Proposition 9.2 (Deuxième inégalité triangulaire)

Soit N une norme sur E .

$$\forall (x, y) \in E^2, |N(x) - N(y)| \leq N(x + y)$$

Remarques :

1. On a aussi

$$\forall (x, y) \in E^2, |N(x) - N(y)| \leq N(x - y)$$

en remplaçant y par $-y$.

2. La deuxième inégalité triangulaire se comprend bien. Si on part de chez soi et que l'on fait 10km. Puis on change de direction et on fait 2km. On se retrouve alors à au moins 8 km de chez soi.

Démonstration : Il suffit de montrer que

$$-N(x + y) \leq N(x) - N(y) \leq N(x + y)$$

c'est-à-dire

$$N(y) \leq N(x) + N(x + y) \text{ et } N(x) \leq N(y) + N(x + y)$$

— Pour la première

$$N(y) = N((x + y) - x) \leq N(x + y) + N(-x) = N(x + y) + N(y)$$

– Pour la deuxième

$$N(x) = N((x+y) - y) \leq N(x+y) + N(-y) = N(x+y) + N(y)$$

□

1.2 Exemples de normes

★ Normes euclidiennes

On suppose que E est un espace préhilbertien réel. C'est-à-dire qu'il y a un produit scalaire $(\cdot|\cdot)$ sur E . Il a été vu en première année que l'on peut lui associer une norme (dite euclidienne).

Définition 9.3

On appelle produit scalaire sur E la donnée d'une forme bilinéaire symétrique définie positive. C'est-à-dire,

1. Pour tout x de E , les applications de E dans $\mathbf{R}$ définies par $(x|\cdot) : y \mapsto (x|y)$ et $(\cdot|x) : y \mapsto (y|x)$ sont linéaires.
2. Pour tout $(x, y) \in E^2$, $(x|y) = (y|x)$. (Symétrique)
3. Pour tout x de E , $(x|x) \geq 0$. (Positive)
4. Pour tout x de E , $(x|x) = 0 \Rightarrow x = 0$. (Définie).

Proposition 9.4 (Inégalité de Cauchy-Schwarz)

Soit $(x, y) \in E^2$,

$$(x|y)^2 \leq (x|x)(y|y)$$

De plus il y a égalité si et seulement si x et y sont colinéaires.

Démonstration : Rappelons la preuve vue en première année. On considère la fonction

$$P : \lambda \mapsto (x + \lambda y|x + \lambda y) = (x|x) + 2\lambda(x|y) + \lambda^2(y|y)$$

- Si $y = 0_E$, l'inégalité est vraie.
- Si $y \neq 0_E$, c'est un trinôme du second degré et par hypothèses, pour tout $\lambda \in \mathbf{R}$, $P(\lambda) \geq 0$ par positivité du produit scalaire. On en déduit que le discriminant est négatif ou nul et donc

$$\Delta = 4(x|y)^2 - 4(x|x)(y|y) \leq 0$$

L'inégalité en découle.

Pour ce qui est des cas d'égalités, là encore si $y = 0$, il y a égalité et x et y sont colinéaires. On suppose donc que $y \neq 0$ par la suite.

- $\Rightarrow$ On suppose qu'il y a égalité dans l'inégalité alors $\Delta = 0$ et donc il existe $\lambda_0 \in \mathbf{R}$ tel que $P(\lambda_0) = 0$. Cela revient à $(x + \lambda_0 y|x + \lambda_0 y) = 0$ et donc (car le produit scalaire est défini) à $x + \lambda_0 y$. On a bien que x et y sont colinéaires.
- $\Leftarrow$ On suppose que x et y sont colinéaires. Comme y n'est pas nul, il existe α tel que $x = \alpha y$. Donc $P(-\alpha) = 0$. On en déduit que le trinôme s'annule et donc que $\Delta = 0$. Il y a bien égalité dans l'inégalité de Cauchy-Schwarz.

□

□

Proposition-Définition 9.5

Soit $(E, (.,.))$ un espace préhilbertien réel. On appelle norme euclidienne et on note $\|.\|$ la norme

$$\begin{aligned}\|.\| : E &\rightarrow \mathbf{R}_+ \\ x &\mapsto \sqrt{(x|x)}\end{aligned}$$

Remarque : On peut définir la racine carrée car, par positivité, $(x|x) \geq 0$.

Démonstration : Vérifions les axiomes

- i) Pour tout x de E , $\|x\| \geq 0$ car c'est une racine carrée.
- ii) Pour x de E , $\|x\| = 0 \Rightarrow (x|x) = 0 \Rightarrow x = 0$ par le caractère défini du produit scalaire.
- iii) Pour x de E et λ de $\mathbf{R}$, $\|\lambda.x\| = \sqrt{(\lambda.x|\lambda.x)} = \sqrt{\lambda^2(x|x)} = |\lambda|\|x\|$.
- iv) Pour $(x, y) \in E^2$ il suffit de montrer que $\|x + y\|^2 \leq (\|x\| + \|y\|)^2 = \|x\|^2 + \|y\|^2 + 2\|x\|\|y\|$.
Or on sait que

$$\begin{aligned}\|x + y\|^2 &= (x + y, x + y) \\ &= \|x\|^2 + \|y\|^2 + 2(x|y) \text{ (Linéarité et symétrie)} \\ &\leq \|x\|^2 + \|y\|^2 + 2\|x\|\|y\| \text{ (Cauchy-Schwarz)}\end{aligned}$$

□

Proposition 9.6 (Autre écriture de l'inégalité de Cauchy-Schwarz)

Soit $(x, y) \in E^2$,

$$|(x|y)| \leq \sqrt{(x|x)}\sqrt{(y|y)} = \|x\| \cdot \|y\|$$

Exemples :

1. Pour $E = \mathbf{R}^n$ avec le produit scalaire usuel défini par

$$(x|y) = \sum_{i=1}^n x_i y_i$$

où $x = (x_1, \dots, x_n)$ et $y = (y_1, \dots, y_n)$. On obtient

$$\|x\| = \sqrt{x_1^2 + \dots + x_n^2}$$

En particulier les vecteurs de la base canonique sont unitaires.

2. De manière générale soit E un $\mathbf{R}$ -espace vectoriel de dimension finie et $\mathcal{B} = (e_1, \dots, e_n)$ une base de E . On peut définir le produit scalaire pour lequel $\mathcal{B}$ est une base orthonormée par

$$(x|y) = \sum_{i=1}^n e_i^*(x) e_i^*(y)$$

où les e_i^* sont les formes coordonnées. C'est-à-dire que si

$$x = \sum_{i=1}^n x_i e_i \quad \text{et} \quad y = \sum_{i=1}^n y_i e_i$$

alors $(x|y) = \sum_{i=1}^n x_i y_i$. Dans ce cas,

$$\|x\| = \sqrt{e_1^*(x)^2 + \cdots + e_n^*(x)^2}.$$

3. Sur $E = \mathcal{M}_n(\mathbf{R})$. On sait que $(A|B) = \text{tr}(A^\top B)$ est un produit scalaire (exercice). On en déduit que

$$\|A\| = \sqrt{\text{tr}(A^\top A)} = \sqrt{\sum_{i=1}^n \sum_{j=1}^n a_{ij}^2}.$$

4. Soit I un intervalle et E l'ensemble des fonctions continues sur I de carré intégrable, c'est-à-dire que f^2 est intégrable ce qui revient à $\int_I f^2$ converge. On pose alors pour tout f et g de E ,

$$(f|g) = \int_I fg$$

On peut vérifier que E est un espace vectoriel et que c'est un produit scalaire. En effet

- Commençons par vérifier que fg est intégrable. Traitons le cas où $I = [a, b[$. On a pour $x < b$,

$$\int_a^x |fg| \leq \int_a^x |f|^2 + \int_a^x |g|^2 \leq \int_a^b |f|^2 + \int_a^b |g|^2$$

car pour tout α, β réels ou complexes,

$$|\alpha\beta| \leq \frac{1}{2}(|\alpha|^2 + |\beta|^2) \leq |\alpha|^2 + |\beta|^2.$$

On pouvait aussi voir que

$$\int_a^x |fg| \leq \sqrt{\int_a^x f^2 \times \int_a^x g^2} \leq \sqrt{\int_a^b f^2 \times \int_a^b g^2}$$

d'après l'inégalité de Cauchy-Schwarz. On en déduit que $x \mapsto \int_a^x |fg|$ est majorée donc

$\int_I |fg|$ converge et fg est intégrable sur I .

- On en déduit que E est un sous-espace vectoriel de $\mathcal{C}^0(I)$. Le point clé est qu'une somme de fonctions de carré intégrable est encore de carré intégrable. En effet, si $(f, g) \in E^2$,

$$(f+g)^2 = f^2 + g^2 + 2fg \leq f^2 + g^2 + 2|fg|.$$

- Il est clair que $(\cdot|\cdot)$ est bilinéaire et symétrique.

- Soit $f \in E$, $(f|f) = \int_I f^2 \geq 0$.

- Soit $f \in E$, comme f est continue, $(f|f) = \int_I f^2 = 0 \Rightarrow f = 0$.

On obtient alors la norme

$$\|f\| = \sqrt{\int_I f^2}$$

Cette norme s'appelle aussi norme L^2 et on appelle $L^2(I)$ l'ensemble des fonctions de carrée intégrable. On appelle aussi cette norme, *norme de la moyenne quadratique*.

Remarque : La plupart du temps, on notera $\|\cdot\|_2$ les normes euclidiennes.

Exercices :

1. Montrer que $(P, Q) \mapsto (P|Q) = \int_{\mathbb{R}} e^{-x^2} \tilde{P}(x) \tilde{Q}(x) dx$ est un produit scalaire sur $\mathbb{R}[X]$. En déduire la norme euclidienne associée.
2. Montrer que $(P, Q) \mapsto (P|Q) = P(-1)Q(-1) + P(0)Q(0) + P(1)Q(1)$ est un produit scalaire sur $\mathbb{R}_2[X]$. En déduire la norme euclidienne associée.
3. Soit E l'ensemble des suites (u_n) telles que la série $\sum u_n^2$ converge. Montrer que E est un espace vectoriel, que $(u, v) \mapsto \sum_{n=0}^{\infty} u_n v_n$ est un produit scalaire et écrire la norme euclidienne associée.

★ Extension à $K = \mathbb{C}$

Tous les exemples ci-dessus s'étendent à des espaces vectoriels définis sur $\mathbb{C}$. Dans ce cas on ne parle plus d'espace préhilbertien réel mais d'espace préhilbertien complexe. De même on parle alors de produit scalaire hermitien.

Exemples :

1. Dans $\mathbb{C}^n$ on pose

$$\|x\|_2 = \sqrt{\sum_{i=1}^n |z_i|^2} \quad \text{ou} \quad x = (z_1, \dots, z_n) \in \mathbb{C}^n.$$

Dans la formule ci-dessus, $|\cdot|$ désigne le module.

Il est clair que la formule ci-dessus désigne une fonction positive, définie et homogène. Il faut vérifier l'inégalité triangulaire.

Pour cela on peut définir la notion de produit scalaire hermitien,

$$((x_1, \dots, x_n), (y_1, \dots, y_n)) \mapsto \sum_{i=1}^n \overline{x_i} y_i$$

et vérifier des propriétés analogues à celles des produits scalaires usuels.

On peut aussi vérifier « à la main » l'inégalité triangulaire :

Pour $(u = (x_1, \dots, x_n), v = (y_1, \dots, y_n)) \in (\mathbb{C}^n)^2$,

$$\|u + v\| = \sqrt{\sum_{i=1}^n |x_i + y_i|^2} \leq \sqrt{\sum_{i=1}^n (|x_i| + |y_i|)^2} \leq \sqrt{\sum_{i=1}^n |x_i|^2} + \sqrt{\sum_{i=1}^n |y_i|^2}$$

La première inégalité vient du fait que $|x_i + y_i| \leq |x_i| + |y_i|$; la deuxième inégalité est l'inégalité triangulaire pour les vecteurs $(|x_1|, \dots, |x_n|)$ et $(|y_1|, \dots, |y_n|)$.

2. Dans un espace vectoriel E sur $\mathbb{C}$ muni d'une base $\mathcal{B}$, on note

$$\|x\|_{2, \mathcal{B}} = \sqrt{\sum_{i=1}^n |e_i^*(x)|^2}$$

où e_i^* désigne la i -ème forme coordonnée de la base $\mathcal{B}$.

3. Dans $\mathcal{C}^0(I, \mathbb{C})$. On considère le sous-espace vectoriel E des fonctions de carré intégrable.

$$E = \{f \in \mathcal{C}^0(I, \mathbb{C}), f^2 \text{ est intégrable}\}.$$

Là encore on peut définir une norme sur E par

$$\|f\|_2 = \int_I |f|^2.$$

★ Sur $\mathbb{K}^n$

Présentons quelques normes classiques sur $\mathbb{K}^n$. Toutes ces normes s'étendent au cas d'un espace vectoriel de dimension finie via le choix d'une base.

Définition 9.7 (Norme 1 et Norme infinie)

1. On appelle norme 1 et on note $\|\cdot\|_1$ la norme

$$\begin{aligned} \|\cdot\|_1 : \quad \mathbb{K}^n &\rightarrow \mathbb{R}_+ \\ (x_1, \dots, x_n) &\mapsto \sum_{i=1}^n |x_i| \end{aligned}$$

2. On appelle norme infinie et on note $\|\cdot\|_\infty$ la norme

$$\begin{aligned} \|\cdot\|_\infty : \quad \mathbb{K}^n &\rightarrow \mathbb{R}_+ \\ (x_1, \dots, x_n) &\mapsto \max_{i \in \llbracket 1; n \rrbracket} |x_i| \end{aligned}$$

Proposition 9.8

La norme 1 et la norme infinie sont des normes.

Avant d'écrire la démonstration, commençons par rappeler un lemme

Lemme 9.9

Soit A une partie non vide de $\mathbb{R}$ et $k \in \mathbb{R}_+$, $\sup(kA) = k \sup(A)$.

Démonstration : A écrire

□

Définition 9.10

Soit E un espace vectoriel de dimension finie et $\mathcal{B} = (e_1, \dots, e_n)$ une base de E . On définit les normes 1 et infinie

1. On appelle norme 1 et on note $\|\cdot\|_1$ la norme

$$\begin{aligned} \|\cdot\|_{1,\mathcal{B}} : E &\rightarrow \mathbf{R}_+ \\ x &\mapsto \sum_{i=1}^n |e_i^*(x)| \end{aligned}$$

2. On appelle norme infinie et on note $\|\cdot\|_\infty$ la norme

$$\begin{aligned} \|\cdot\|_{\infty,\mathcal{B}} : E &\rightarrow \mathbf{R}_+ \\ x &\mapsto \max_{i \in \llbracket 1; n \rrbracket} |e_i^*(x)| \end{aligned}$$

Remarques :

1. Cela signifie que si x se décompose dans la base sous la forme $x = \sum_{i=1}^n x_i e_i$ alors

$$\|x\|_{1,\mathcal{B}} = \sum_{i=1}^n |x_i| \quad \text{et} \quad \|x\|_{\infty,\mathcal{B}} = \max_{i \in \llbracket 1; n \rrbracket} |x_i|$$

2. Il est clair que ces normes dépendent du choix de la base $\mathcal{B}$. Cependant, pour ne pas alourdir les notations nous ne noterons pas systématiquement la base $\mathcal{B}$ en indice.
3. On peut donc définir les normes $\|\cdot\|_1$, $\|\cdot\|_\infty$ et aussi $\|\cdot\|_2$ sur tout espace vectoriel de dimension finie comme $\mathbf{K}_n[X]$ ou $\mathcal{M}_n(\mathbf{K})$.

★ Norme 1 et norme ∞ sur un espace de fonctions

On peut définir les normes 1 et ∞ sur des sous-espaces vectoriels des fonctions.

Définition 9.11

Soit X un ensemble non vide. L'ensemble $\mathcal{B}(X, \mathbf{K})$ des fonctions bornées sur X est un sous-espace vectoriel de $\mathcal{F}(X, \mathbf{K})$. La norme infinie sur $\mathcal{B}(X, \mathbf{K})$ est définie par

$$\|f\|_\infty = \sup_{x \in X} |f(x)|.$$

Remarques :

1. Le fait que la borne supérieure existe vient du fait que f soit bornée.
2. On a déjà vu que c'était une norme. On l'appelle norme de la convergence absolue.

On peut aussi définir la norme 1 sur l'ensemble des fonctions continues intégrables sur un intervalle I ayant une infinité de points.

Proposition-Définition 9.12

Soit E le sous-espace vectoriel de $\mathcal{C}^0(I, \mathbf{K})$ des fonctions intégrables. On pose pour tout f de E .

$$\|f\|_1 = \int_I |f|.$$

On définit ainsi une norme sur E .

Démonstration :

- i) Pour toute f de E , $\|f\|_1 \geq 0$ par définition.
- ii) Soit $f \in E$. On suppose que $\|f\|_1 = \int_I |f| = 0$. Comme $|f|$ est continue et positive, on obtient que $|f|$ est la fonction nulle et donc $f = 0$.
- iii) Soit $f \in E$ et $\lambda \in \mathbf{K}$, $\|\lambda \cdot f\|_1 = \int_I |\lambda \cdot f| = |\lambda| \int_I |f| = |\lambda| \|f\|_1$.
- iv) Soit f et g dans E .

$$\|f + g\|_1 = \int_I |f + g| \leq \int_I |f| + |g| = \int_I |f| + \int_I |g| = \|f\|_1 + \|g\|_1.$$

On a bien montré que $\|\cdot\|_1$ était une norme sur E . □

Remarque : On appelle aussi cette norme, la norme de la convergence en moyenne.

ATTENTION

Il faut bien connaître et ne pas mélanger les définitions des normes $\|\cdot\|_1$, $\|\cdot\|_2$ et $\|\cdot\|_\infty$ dans le cas des espaces vectoriels $\mathbf{K}^n$ (ou tout espace vectoriel de dimension finie avec une base), fonctions à valeurs dans $\mathbf{K}$ et suites numériques.

Espace	$\ \cdot\ _1$	$\ \cdot\ _2$	$\ \cdot\ _\infty$
$\mathbf{K}^n$	$\sum x_i $	$\sqrt{\sum x_i ^2}$	$\max x_i $
E	$\sum e_i^*(x) $	$\sqrt{\sum e_i^*(x) ^2}$	$\max e_i^*(x) $
$\mathcal{M}_n(\mathbf{K})$	$\sum a_{ij} $	$\sqrt{\sum a_{ij} ^2}$	$\max a_{ij} $
fonctions	$\int f $	$\sqrt{\int f ^2}$	$\max f(x) $

★ Produit fini d'espaces vectoriels normés**Proposition-Définition 9.13**

Soit $(E_1, N_1), \dots, (E_n, N_n)$ des espaces vectoriels normés.

L'application

$$N : E_1 \times \dots \times E_n \rightarrow \mathbf{R}_+ \\ (x_1, \dots, x_n) \mapsto \max_{i \in \llbracket 1; n \rrbracket} N_i(x_i)$$

est une norme sur $E_1 \times \dots \times E_n$. On l'appelle la norme produit.

Démonstration :

i) Pour tout $x = (x_1, \dots, x_n) \in E_1 \times \dots \times E_n$ et pour tout $i \in \llbracket 1; n \rrbracket$, $N_i(x_i) \geq 0$ donc

$$N(x) = \max_{i \in \llbracket 1; n \rrbracket} N_i(x_i) \geq 0.$$

ii) Soit $x = (x_1, \dots, x_n) \in E_1 \times \dots \times E_n$. On suppose que $N(x) = \max_{i \in \llbracket 1; n \rrbracket} N_i(x_i) = 0$. Cela

implique que pour tout $i \in \llbracket 1; n \rrbracket$, $N_i(x_i) = 0$ et donc $x_i = 0$. Finalement $x = 0$

iii) Soit $x = (x_1, \dots, x_n) \in E_1 \times \dots \times E_n$ et $\lambda \in \mathbf{K}$. Pour tout $i \in \llbracket 1; n \rrbracket$, $N_i(\lambda x_i) = |\lambda| N_i(x_i)$. On a donc

$$N(\lambda x) = \max_{i \in \llbracket 1; n \rrbracket} |\lambda| N_i(x_i) = |\lambda| N(x).$$

iv) Soit $x = (x_1, \dots, x_n) \in E_1 \times \dots \times E_n$ et $y = (y_1, \dots, y_n) \in E_1 \times \dots \times E_n$. On note $z = x + y = (z_1, \dots, z_n)$. Pour tout $i \in \llbracket 1; n \rrbracket$, $N_i(z_i) \leq N_i(x_i) + N_i(y_i)$. On a donc

$$N(z) = \max_{i \in \llbracket 1; n \rrbracket} N_i(z_i) \leq \max_{i \in \llbracket 1; n \rrbracket} (N_i(x_i) + N_i(y_i)) \leq \max_{i \in \llbracket 1; n \rrbracket} N_i(x_i) + \max_{i \in \llbracket 1; n \rrbracket} N_i(y_i) = N(x) + N(y).$$

On a bien montré que N était une norme sur $E_1 \times \dots \times E_n$. □

Remarque : En prenant sur $\mathbf{R}$ ou $\mathbf{C}$ la norme $x \mapsto |x|$ on obtient ainsi la norme ∞ .

1.3 Distance associée à une norme

Définition 9.14

Soit X un ensemble. On appelle distance sur X une application

$$d : X \times X \rightarrow \mathbf{R}$$

vérifiant les axiomes suivants

- i) Pour tout $(x, y) \in X \times X$, $d(x, y) \geq 0$
- ii) Pour tout $(x, y) \in X \times X$, $d(x, y) = 0 \iff x = y$
- iii) Pour tout $(x, y) \in X \times X$, $d(x, y) = d(y, x)$
- iv) Pour tout $(x, y, z) \in X^3$, $d(x, z) \leq d(x, y) + d(y, z)$. (Inégalité triangulaire)

Proposition-Définition 9.15 (Distance associée à une norme)

Soit $(E, \|\cdot\|)$ un espace vectoriel normé. On définit la fonction

$$\begin{aligned} d_{\|\cdot\|} : E \times E &\rightarrow \mathbf{R} \\ (x, y) &\mapsto \|x - y\| \end{aligned}$$

C'est une distance sur E . Elle s'appelle la distance associée à la norme $\|\cdot\|$.

Démonstration : En exercice □

Remarques :

1. Cela généralise la distance classique dans $\mathbf{R}$ ou $\mathbf{C}$ qui est définie par $d(x, y) = |x - y|$

2. La fonction distance dépend de la norme utilisée.
3. Il existe des distances qui ne proviennent pas de normes. Par exemple sur $\mathbf{R}$ on peut définir

$$\delta(x, y) = \begin{cases} 0 & \text{si } x = y \\ 1 & \text{sinon} \end{cases}$$

Cette distance n'est pas issue d'une norme car $\delta(\lambda.x, \lambda.y) = \delta(x, y) \neq |\lambda|.d(x, y)$ dans le cas où $|\lambda| \neq 1$ et $x \neq y$.

Exemples :

1. On se place dans $\mathbf{R}^2$. On considère $u = (1, 3)$ et $v = (4, -1)$. On a alors

$$d_1(u, v) = \|u - v\|_1 = 7; d_2(u, v) = \|u - v\|_2 = 5; d_\infty(u, v) = \|u - v\|_\infty = 4$$

2. Dans $\mathcal{M}_3(\mathbf{R})$, on considère $A = I_3$ et $B = \begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 1 & 0 & 0 \end{pmatrix}$. On a alors

$$d_1(A, B) = \|A - B\|_1 = 6; d_2(A, B) = \|A - B\|_2 = \sqrt{6}; d_\infty(A, B) = \|A - B\|_\infty = 1$$

3. Pour $E = \mathcal{C}^0([0, 1], \mathbf{R})$, $f : x \mapsto x^2$ et $g : x \mapsto x^3$,

$$d_1(f, g) = \|f - g\|_1 = \frac{1}{12}; d_2(f, g) = \sqrt{\frac{1}{105}}; d_\infty(f, g) = \|f - g\|_\infty = \frac{4}{27}$$

1.4 Boules

Définition 9.16 (Boules)

Soit $(E, \|\cdot\|)$ un espace vectoriel normé. Soit $a \in E$ et $\delta \in \mathbf{R}$.

1. On appelle *boule ouverte* de centre a et de rayon δ et on note $B(a, \delta)$ l'ensemble des éléments de E à distance strictement inférieure à δ de a .

$$B(a, \delta) = \{x \in E \mid d(x, a) < \delta\} = \{x \in E \mid \|x - a\| < \delta\}.$$

2. On appelle *boule fermée* de centre a et de rayon δ et on note $\overline{B}(a, \delta)$ l'ensemble des éléments de E à distance inférieure (ou égale) à δ de a .

$$\overline{B}(a, \delta) = \{x \in E \mid d(x, a) \leq \delta\} = \{x \in E \mid \|x - a\| \leq \delta\}.$$

3. On appelle *sphère* de centre a et de rayon δ et on note $S(a, \delta)$ l'ensemble des éléments de E à distance δ de a .

$$S(a, \delta) = \{x \in E \mid d(x, a) = \delta\} = \{x \in E \mid \|x - a\| = \delta\}.$$

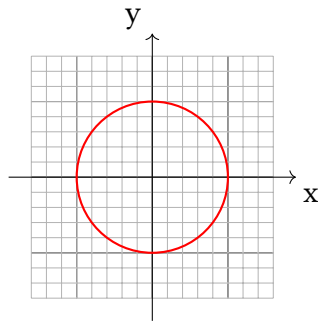
Terminologie : Dans le cas où $a = 0_E$ et $\delta = 1$, on parle de *boule unité* et de *sphère unité*. La sphère unité est l'ensemble des vecteurs unitaires.

Exemple : Dans $E = \mathbf{R}^2$. On cherche la boule unité selon les normes.

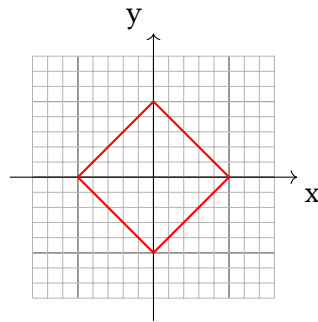
- Pour la norme euclidienne on trouve la boule classique :

$$B_{\|\cdot\|_2}(0, 1) = \{(x, y) \in \mathbf{R}^2 \mid \sqrt{|x|^2 + |y|^2} \leq 1\}$$

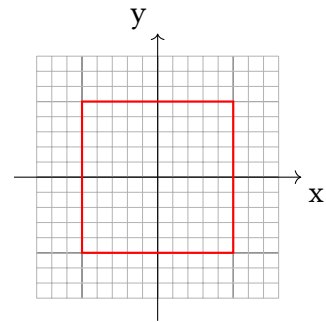
- Pour la norme 1 on a que $(x, y) \in B_{\|\cdot\|_1}(0, 1) \iff \|(x, y)\|_1 \leq 1 \iff |x| + |y| \leq 1$.
- Pour la norme ∞ on a que $(x, y) \in B_{\|\cdot\|_\infty}(0, 1) \iff \|(x, y)\|_\infty \leq 1 \iff |x| \leq 1 \text{ et } |y| \leq 1$.



Norme 2



Norme 1



Norme ∞

Exercice : Déterminer selon $a = (x_0, y_0)$ les boules $B(a, \delta)$ pour ces trois normes.

Définition 9.17

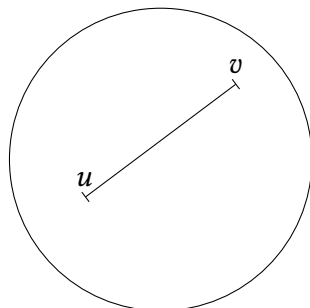
Soit E un espace vectoriel normé.

- Soit u, v deux vecteurs de E . On appelle segment $[u, v]$ l'ensemble

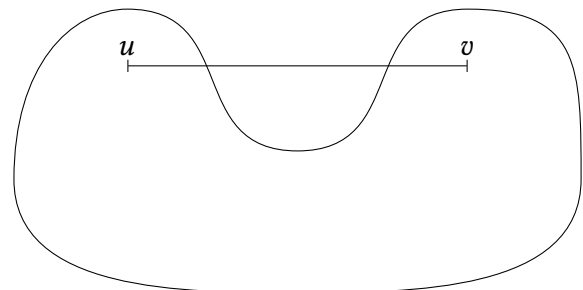
$$[u, v] = \{tu + (1 - t)v, t \in [0, 1]\}$$

- Soit A une partie de E , elle est dite convexe si pour tout couple $(u, v) \in E^2$, $[u, v] \subset A$.

Exemple :



Convexe



Non Convexe

Proposition 9.18

Soit $(E, \|\cdot\|)$ un espace vectoriel normé. Les boules ouvertes (resp. fermées) sont convexes.

Démonstration : Soit $B(a, \delta)$ une boule ouverte. Soit x, y deux éléments de la boule. Ils vérifient donc que $\|x - a\| < \delta$ et $\|y - a\| < \delta$. Soit $z = tx + (1 - t)y$ un élément du segment $[x, y]$ où $t \in [0, 1]$. On a

$$d(z, a) = \|z - a\| = \|t(x - a) + (1 - t)(y - a)\| \leq t\|x - a\| + (1 - t)\|y - a\| < t\delta + (1 - t)\delta = \delta.$$

On en déduit que $z \in B(a, \delta)$ et donc que $[x, y] \subset B(a, \delta)$.

Le raisonnement est similaire pour les boules fermées.

1.5 Parties, suites et fonctions bornées

Définition 9.19

Soit $(E, \|\cdot\|)$ un espace vectoriel normé.

1. Une partie $X \subset E$ est dite bornée s'il existe $M \in \mathbf{R}$ tel que

$$\forall x \in X, \|x\| \leq M.$$

2. Une suite (u_n) à valeurs dans E est dite bornée si son image $\{u_n \mid n \in \mathbf{N}\}$ est une partie bornée de E . C'est-à-dire

$$\exists M \in \mathbf{R}, \forall n \in \mathbf{N}, \|u_n\| \leq M.$$

3. Soit X un ensemble et $f \in \mathcal{F}(X, E)$. Elle est dite bornée si son image $f(X) = \{f(x) \mid x \in X\}$ est bornée. C'est-à-dire

$$\exists M \in \mathbf{R}, \forall x \in X, \|f(x)\| \leq M.$$

Proposition 9.20

Avec les notations précédentes. Soit $X \subset E$. Il y a équivalence entre les assertions suivantes :

- i) La partie X est bornée.
- ii) Il existe $M \in \mathbf{R}$ tel que $X \subset B(0_E, M)$.
- iii) Il existe $a \in E$ et $M \in \mathbf{R}$ tel que $X \subset B(a, M)$.

Démonstration :

- $i) \Rightarrow ii)$ Par définition
- $ii) \Rightarrow iii)$ Evident en prenant $a = 0_E$.
- $iii) \Rightarrow i)$ Soit $a \in E$ et $M \in \mathbf{R}$ tel que $X \subset B(a, M)$. Il suffit de trouver M' tel que $B(a, M) \subset B(0_E, M')$. On pose $M' = M + d(a, 0_E)$ et on a que pour $x \in B(a, M)$,

$$d(0_E, x) \leq d(0_E, a) + d(a, x) \leq d(0_E, a) + M \leq M'.$$

□

ATTENTION

Contrairement à l'intuition, le fait qu'une partie soit bornée dépend de la norme choisie. Dans l'ensemble des fonctions continues bornées (donc intégrables) sur $[0, 1]$, on considère

$$f_n : x \mapsto \begin{cases} n^2 x & \text{si } x < \frac{1}{2n} \\ 2n - n^2 x & \text{si } \frac{1}{2n} \leq x < \frac{1}{n} \\ 0 & \text{sinon.} \end{cases}$$

La suite $(f_n)_{n \geq 1}$ n'est pas bornée pour $\|\cdot\|_\infty$ car $f_n(\frac{1}{n}) = n$ et on a $\|f_n\|_\infty = n$. Par contre elle est bornée pour $\|\cdot\|_1$ car $\|f_n\|_1 = \frac{1}{2}$.

1.6 Normes équivalentes

On a vu que sur un même espace vectoriel on peut définir plusieurs normes. On vient de voir qu'une même partie peut-être bornée pour une norme et pas une autre. Il y a un cas où cela ne peut pas arriver.

Définition 9.21

Soit E un espace vectoriel et N_1 et N_2 deux normes sur E .
Elles sont dites équivalentes s'il existe des constantes C_1 et C_2 strictement positives telles que

$$N_1 \leq C_1 N_2 \text{ et } N_2 \leq C_2 N_1$$

Remarques :

1. Si les normes N_1 et N_2 sont équivalentes on a

$$\frac{1}{C_2} N_2 \leq N_1 \leq C_1 N_2 \text{ et } \frac{1}{C_1} N_1 \leq N_2 \leq C_2 N_1.$$

2. Si on a juste C tel que $N_2 \leq C N_1$ on dit que N_1 est *plus fine* que N_2 .

Proposition 9.22

Le fait d'être équivalentes pour des normes est une relation d'équivalence.

Proposition 9.23 (Invariance du caractère borné)

Soit N_1 et N_2 deux normes équivalentes sur un espace vectoriel E

1. Soit X une partie de E , elle est bornée pour N_1 si et seulement si elle est bornée pour N_2 .
2. Soit $(u_n) \in E^{\mathbb{N}}$, elle est bornée pour N_1 si et seulement si elle est bornée pour N_2 .
3. Soit f une fonction à valeurs dans E , elle est bornée pour N_1 si et seulement si elle est bornée pour N_2 .

Démonstration :

1. On suppose que X est bornée pour N_1 . Il existe alors M_1 telle que $X \subset \overline{B_1}(0_E, M_1)$ (la boule pour la norme N_1) c'est-à-dire

$$\forall x \in X, N_1(x) \leq M_1.$$

Maintenant comme N_1 et N_2 sont équivalentes, il existe C_2 tel que $N_2 \leq C_2 N_1$. On en déduit que

$$\forall x \in X, N_2(x) \leq C_2 N_1(x) \leq C_2 M_1.$$

La partie X est bien bornée pour N_2 car $X \subset \overline{B_2}(0_E, C_2 M_1)$ (la boule pour la norme N_2). Comme la relation « être équivalente » est une relation d'équivalence, on peut échanger les rôles de N_1 et N_2 .

2. Il suffit d'utiliser que (u_n) bornée si et seulement si l'image $\{u_n \mid n \in \mathbb{N}\}$ est bornée.
3. Il suffit d'utiliser que f bornée si et seulement si l'image de f est bornée.

□

★ **Méthode** : Montrer que des normes ne sont pas équivalentes.

Pour montrer que deux normes N_1 ou N_2 ne sont pas équivalentes il faut montrer que l'on ne peut pas trouver un réel K tel que pour tout $x \in E$, $N_1(x) \leq KN_2(x)$ ou dans l'autre sens. Traitons ce cas. Cela revient à montrer que l'ensemble $\left\{ \frac{N_1(x)}{N_2(x)}, x \neq 0 \right\}$ n'est pas majoré.

Exemple : On reprend les fonctions f_n sur $[0, 1]$ telles que pour tout entier $n \geq 1$, $\|f_n\|_\infty = n$ et $\|f_n\|_1 = \frac{1}{2}$. On en déduit qu'il n'existe pas de réel K tel que $\|\cdot\|_\infty \leq K\|\cdot\|_1$. Ces deux normes ne sont pas équivalentes.

Exercice : Dans $\mathbb{R}[X]$ on pose

$$N_1 : P \mapsto \sup_{t \in [-1, 0]} |\tilde{P}(t)| \text{ et } N_2 : P \mapsto \sup_{t \in [0, 1]} |\tilde{P}(t)|.$$

1. Montrer que N_1 et N_2 sont des normes.
2. Justifier qu'elles ne sont pas équivalentes. Considérer $P_n = (X+1)^n$
3. Les normes sont-elles équivalentes sur $\mathbb{R}_n[X]$?

Elles le sont. Considérons par exemple les polynômes $L_0, L_1, \dots, L_n$ d'interpolation aux points $x_0, x_1, \dots, x_n$ où $x_k = \frac{k}{n}$. Pour tout polynôme $P \in \mathbb{R}_n[X]$, $P = \sum P(x_k)L_k$ et donc

$$N_1(P) \leq \sum |P(x_k)| N_1(L_k) \leq K N_2(P) \text{ où } K = \sum N_1(L_k)$$

Théorème 9.24

Soit E un espace vectoriel de dimension finie. Toutes les normes sont équivalentes.

Démonstration : On démontrera ce théorème plus tard. □

Remarque : Ce théorème est très important. Il signifie que sur un espace vectoriel de dimension finie, il n'est souvent pas nécessaire de préciser la norme avec laquelle on travaille ; nous allons y revenir.

2 Exemple de normes équivalentes

2.1 Sur $\mathbb{K}^n$

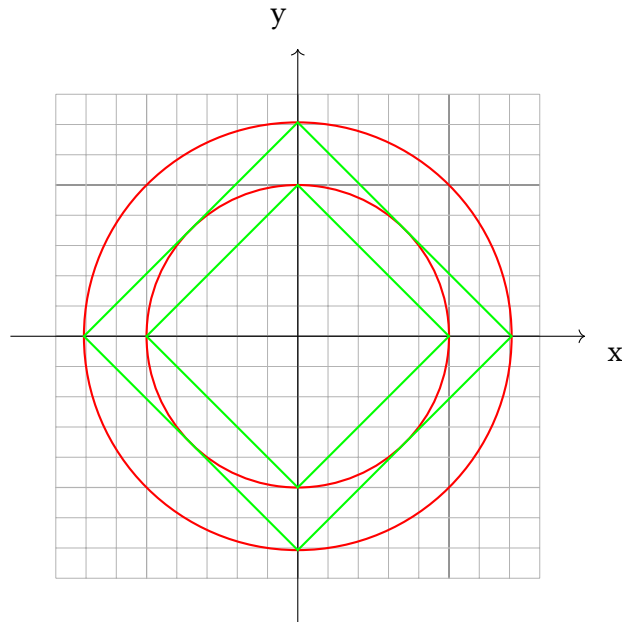
Nous verrons plus loin que sur un espace de dimension finie toutes les normes sont équivalentes. Dans le cas des normes classiques ($\|\cdot\|_1$, $\|\cdot\|_2$ et $\|\cdot\|_\infty$) on peut le démontrer et trouver explicitement les constantes.

Commençons par le cas de $\mathbb{R}^2$. On peut visualiser l'équivalence des normes à l'aide des boules unités

On voit que l'on peut placer une boule pour une norme entre deux boules pour une autre norme.

Précisément :

- Comparaison $\|\cdot\|_1$ et $\|\cdot\|_2$:



Sur le dessin on voit que

$$B_1(0, 1) \subset B_2(0, 1) \subset B_1(0, \sqrt{2})$$

Cela signifie que pour x de E , $\|x\|_1 \leq 1 \Rightarrow \|x\|_2 \leq 1$ et donc $\|x\|_2 \leq \|x\|_1$. De même comme $\|x\|_2 \leq 1 \Rightarrow \|x\|_1 \leq \sqrt{2}$ on obtient que $\|x\|_1 \leq \sqrt{2}\|x\|_2$.

De manière générale sur $\mathbf{K}^n$ on a pour tout x

$$\|x\|_2 \leq \|x\|_1 \leq \sqrt{n}\|x\|_2.$$

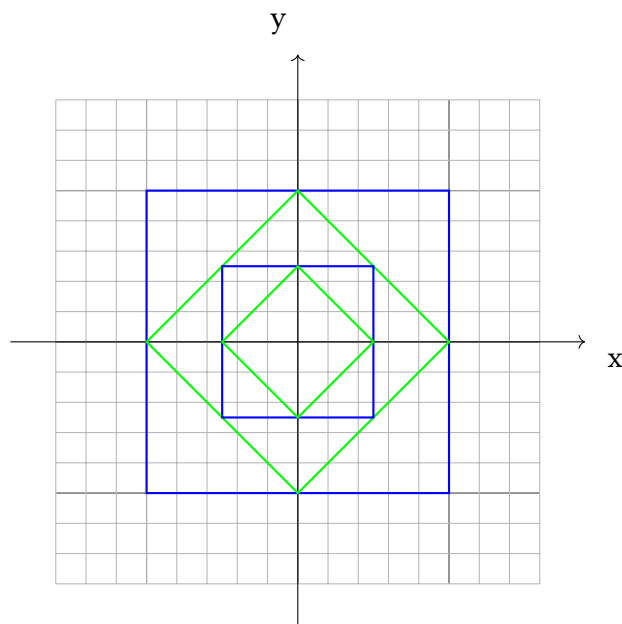
En effet pour $x = (x_1, \dots, x_n)$,

$$\|x\|_2^2 = |x_1|^2 + \dots + |x_n|^2 \leq \sum_{1 \leq i, j \leq n} |x_i| |x_j| = \|x\|_1^2.$$

Dans l'autre sens, on utilise l'inégalité de Cauchy-Schwarz en posant $y_1 = \dots = y_n = 1$, on a

$$\|x\|_1^2 = (x|y)^2 \leq (x|x)(y|y) = n\|x\|_2^2$$

— Comparaison $\|\cdot\|_1$ et $\|\cdot\|_\infty$:



Sur le dessin on voit que

$$B_\infty\left(0, \frac{1}{2}\right) \subset B_1(0, 1) \subset B_\infty(0, 1)$$

Cela signifie que pour x de E , $\|x\|_\infty \leq \frac{1}{2} \Rightarrow \|x\|_1 \leq 1$ et donc $\|x\|_1 \leq 2\|x\|_\infty$. De même comme $\|x\|_1 \leq 1 \Rightarrow \|x\|_\infty \leq 1$ on obtient que $\|x\|_\infty \leq \|x\|_1$.

De manière générale sur $\mathbf{K}^n$ on a pour tout x

$$\|x\|_\infty \leq \|x\|_1 \leq n\|x\|_\infty.$$

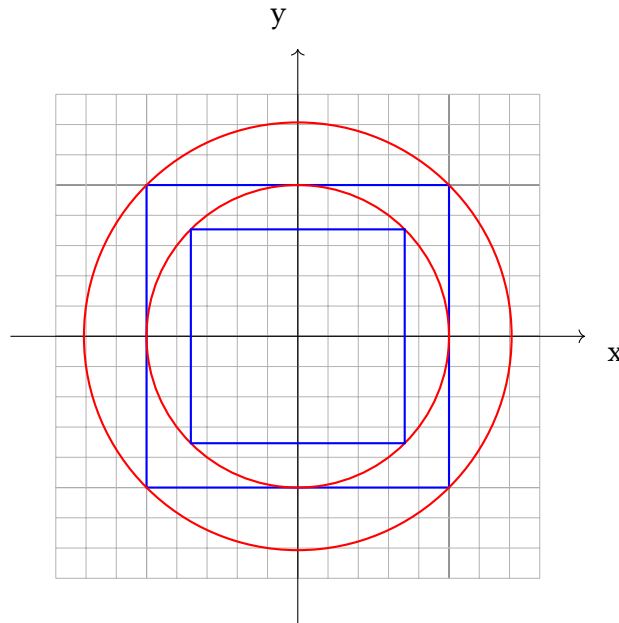
En effet pour $x = (x_1, \dots, x_n)$,

$$\|x\|_1 = |x_1| + \dots + |x_n| \geq \max_{i \in \llbracket 1; n \rrbracket} |x_i| = \|x\|_\infty.$$

Dans l'autre sens,

$$\|x\|_1 = |x_1| + \dots + |x_n| \leq n \cdot \max_{i \in \llbracket 1; n \rrbracket} |x_i| = n\|x\|_\infty.$$

— Comparaison $\|\cdot\|_2$ et $\|\cdot\|_\infty$:



Sur le dessin on voit que

$$B_2(0, 1) \subset B_\infty(0, 1) \subset B_2(0, \sqrt{2})$$

Cela signifie que pour x de E , $\|x\|_2 \leq 1 \Rightarrow \|x\|_\infty \leq 1$ et donc $\|x\|_\infty \leq \|x\|_2$. De même comme $\|x\|_\infty \leq 1 \Rightarrow \|x\|_2 \leq \sqrt{2}$ on obtient que $\|x\|_2 \leq \sqrt{2}\|x\|_\infty$.

De manière générale sur $\mathbf{K}^n$ on a pour tout x

$$\|x\|_\infty \leq \|x\|_2 \leq \sqrt{n}\|x\|_\infty.$$

En effet pour $x = (x_1, \dots, x_n)$,

$$\|x\|_2^2 = |x_1|^2 + \dots + |x_n|^2 \geq \left(\max_{i \in \llbracket 1; n \rrbracket} |x_i| \right)^2 = \|x\|_\infty^2.$$

Dans l'autre sens,

$$\|x\|_2^2 = |x_1|^2 + \dots + |x_n|^2 \leq n \left(\max_{i \in \llbracket 1; n \rrbracket} |x_i| \right)^2 = n\|x\|_\infty^2.$$

On a bien $\|x\|_2 \leq \sqrt{n}\|x\|_\infty$.

2.2 Sur les espaces de fonctions

— Cas où I est borné.

On a déjà vu que pour $[0, 1]$ la norme $\|\cdot\|_1$ et la norme $\|\cdot\|_\infty$ n'étaient pas équivalentes. Plus précisément soit $I = [a, b]$, on a

$$\|\cdot\|_1 \leq (b-a)\|\cdot\|_\infty.$$

Par contre il n'y a pas d'inégalité dans l'autre sens. On peut trouver f telle que $\frac{\|f\|_\infty}{\|f\|_1}$ soit arbitrairement grand.

De même,

$$\|f\|_2 = \sqrt{\int_a^b |f|^2} \leq \sqrt{b-a}\|f\|_\infty.$$

Là encore, en adaptant l'exemple ci-dessus, il n'y a pas d'inégalité dans l'autre sens. On peut trouver f telle que $\frac{\|f\|_\infty}{\|f\|_2}$ soit arbitrairement grand.

Il reste à comparer $\|\cdot\|_1$ et $\|\cdot\|_2$. En utilisant l'inégalité de Cauchy-Schwarz et en posant $g : t \mapsto 1$, on obtient que pour toute fonction f ,

$$\int_I |f| = \int_I |f|g \leq \sqrt{\int_I |f|^2} \sqrt{\int_I g^2} = \sqrt{b-a} \sqrt{\int_I |f|^2}$$

C'est-à-dire que

$$\|f\|_1 \leq \sqrt{b-a}\|f\|_2$$

Inversement, si $I = [0, 1]$, en prenant des fonctions continues qui approchent la fonction caractéristique de $[0, \frac{1}{n}]$ on voit que l'on peut obtenir un rapport aussi grand que l'on veut pour $\frac{\|f\|_2}{\|f\|_1}$.

— Si I n'est pas borné.

Dans ce cas, on peut trouver des fonctions f telles que $\frac{\|f\|_1}{\|f\|_\infty}$ et $\frac{\|f\|_2}{\|f\|_\infty}$ soit aussi grand que l'on veut.

En se plaçant sur $[1, +\infty[$, il suffit de penser aux fonctions $t \mapsto \frac{1}{t^\alpha}$ avec $\alpha \rightarrow 1$ pour le premier rapport et $\alpha \rightarrow \frac{1}{2}$ pour le deuxième.

3 Suites à valeurs dans un espace vectoriel normé

Dans ce paragraphe on se fixe un espace vectoriel normé $(E, \|\cdot\|)$. On va généraliser les notions vues pour les suites à valeurs dans $\mathbf{K}$ au cas des suites à valeurs dans E . Cela permettra de revoir les notions sur les suites de fonctions, étudier les suites de matrices, de polynômes.

3.1 Convergence

Définition 9.25 (Convergence)

Soit (u_n) une suite à valeurs dans E . Soit $\ell \in E$. On dit que la suite (u_n) tend vers ℓ et on note $(u_n) \xrightarrow[n \rightarrow \infty]{} \ell$ si

$$\forall \varepsilon > 0, \exists N \in \mathbb{N}, \forall n \in \mathbb{N}, n \geq N \Rightarrow \|u_n - \ell\| \leq \varepsilon$$

Remarques :

1. On peut remplacer $\|u_n - \ell\| \leq \varepsilon$ par $d(u_n, \ell) \leq \varepsilon$ ou $u_n \in \overline{B}(\ell, \varepsilon)$.
2. On utilise une inégalité large mais on peut aussi utiliser une inégalité stricte. En effet la définition avec l'inégalité stricte implique celle avec l'inégalité large car $B(\ell, \varepsilon) \subset \overline{B}(\ell, \varepsilon)$. Inversement, la définition avec l'inégalité large implique celle avec l'inégalité stricte car $\overline{B}(\ell, \frac{\varepsilon}{2}) \subset B(\ell, \varepsilon)$.

Définition 9.26

1. On dit qu'une suite à valeurs dans E est convergente s'il existe $\ell \in E$, tel que $(u_n) \xrightarrow[n \rightarrow \infty]{} \ell$.
2. On dit qu'une suite à valeurs dans E est divergente si elle n'est pas convergente.

ATTENTION

Cela n'a pas de sens de dire qu'une suite à valeurs dans E tend vers $+\infty$. Dans certains cas, on pourra considérer le fait que $\|u_n\| \xrightarrow[n \rightarrow \infty]{} +\infty$.

Proposition 9.27 (Unicité de la limite)

Soit $(u_n) \in E^{\mathbb{N}}$. Soit ℓ et ℓ' . Si $(u_n) \xrightarrow[n \rightarrow \infty]{} \ell$ et $(u_n) \xrightarrow[n \rightarrow \infty]{} \ell'$ alors $\ell = \ell'$.

Démonstration : Il suffit de réécrire la démonstration usuelle. A faire en exercice. $\square$

Définition 9.28 (Limite)

Soit $(u_n) \in E^{\mathbb{N}}$. Si $(u_n) \xrightarrow[n \rightarrow \infty]{} \ell$ on dit que ℓ est la limite de (u_n) et on note

$$\lim_{n \rightarrow \infty} u_n = \ell \quad \text{ou} \quad \lim(u_n) = \ell.$$

Exemples :

1. Dans le cas où $E = \mathbb{K}$, on retrouve la définition classique de la limite.
2. On considère E l'ensemble des fonctions bornées sur un intervalle I muni de la norme infinie. Dire qu'une suite de fonctions (f_n) converge vers f signifie que (f_n) converge uniformément vers f .

3. On considère la suite de fonctions f_n définies par

$$f_n : x \mapsto \begin{cases} nx & \text{si } x < \frac{1}{2n} \\ 2 - nx & \text{si } \frac{1}{2n} \leq x < \frac{1}{n} \\ 0 & \text{sinon.} \end{cases}$$

La suite (f_n) converge vers la fonction nulle pour la norme 1 car $\|f_n\|_1 \rightarrow 0$. Par contre elle ne converge pas vers la fonction nulle pour la norme ∞ .

4. Soit $A = \begin{pmatrix} \frac{1}{2} & \frac{1}{2} \\ 0 & \frac{1}{2} \end{pmatrix} = \frac{1}{2}B$ où $B = \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$. On considère la suite $U_n = A^n = \frac{1}{2^n}B^n$. Par une récurrence, on montre que pour tout entier n , $B^n = \begin{pmatrix} 1 & n \\ 0 & 1 \end{pmatrix}$ et donc (U_n) tend vers la matrice nulle pour $\|\cdot\|_\infty$.

ATTENTION

Le fait qu'une suite converge dépend de la norme utilisée.

Proposition 9.29

Soit $(u_n) \in E^{\mathbb{N}}$ et $\ell \in E$. Les assertions suivantes sont équivalentes

- i) La suite (u_n) tend vers ℓ .
- ii) La suite $(u_n - \ell)$ tend vers 0_E
- iii) La suite réelle $(\|u_n - \ell\|)$ tend vers 0.

Démonstration :

- $i) \Rightarrow ii)$ Par définition....
- $ii) \Rightarrow iii)$ Par définition....
- $iii) \Rightarrow i)$ Par définition....

□

Proposition 9.30

On considère un produit d'espace vectoriel normé muni de la norme produit $(E = \prod_{i=1}^p E_i)$.

Soit (u_n) une suite de E , on pose les suites $(u_n(i)) \in E_i^{\mathbb{N}}$ telles que pour tout entier n , $u_n = (u_n(1), \dots, u_n(p))$.

La suite (u_n) converge si et seulement si toutes les suites $(u_n(i))$ convergent.

Dans ce cas

$$\lim(u_n) = (\ell_1, \dots, \ell_p)$$

où $\ell_i = \lim((u_n(i)))$.

Démonstration :

- $\Rightarrow$ On suppose que la suite (u_n) converge. Notons $\ell = (\ell_1, \dots, \ell_p)$ la limite. On a donc que $(N(u_n - \ell))$ tend vers 0. On a alors pour tout $k \in \llbracket 1 ; p \rrbracket$ et tout $n \in \mathbb{N}$,

$$N_k(u_n(k) - \ell_k) \leq \max_{i \in \llbracket 1 ; p \rrbracket} N_i(u_n(i) - \ell_i) = N(u_n - \ell)$$

- On en déduit que la suite $(N_k(u_n(k) - \ell_k))$ tend vers 0 et donc la suite $(u_n(k))$ tend vers ℓ_k .
- $\Leftarrow$ Supposons que pour tout $k \in \llbracket 1 ; p \rrbracket$, $(u_n(k))$ tend vers un élément noté ℓ_k de E_k . Posons $\ell = (\ell_1, \dots, \ell_p)$. On a que pour tout entier n ,

$$N(u_n - \ell) = \max_{i \in \llbracket 1 ; p \rrbracket} N_i(u_n(i) - \ell_i) \leq \sum_{i=1}^p N_i(u_n(i) - \ell_i)$$

Maintenant, par hypothèses, toutes les suites $(N_i(u_n(i) - \ell_i))$ tendent vers 0 donc $(N(u_n - \ell))$ tend vers 0 et donc (u_n) tend vers ℓ .

□

Remarque : En particulier, si pour tout i , on prend $E_i = \mathbf{K}$, on obtient $E = \mathbf{K}^p$ et la norme produit est alors la norme infinie. On en déduit que, pour la norme $\|\cdot\|_\infty$, une suite $(u_n) \in (\mathbf{K}^p)^{\mathbf{N}}$ converge si et seulement si toutes ses composantes convergent.

Proposition 9.31 (Invariance de la convergence)

Soit N_1 et N_2 deux normes équivalentes sur un espace vectoriel E . Soit $(u_n) \in E^{\mathbf{N}}$. La suite converge pour la norme N_1 si et seulement si elle converge pour la norme N_2 . Dans ce cas les limites sont les mêmes.

Démonstration : On suppose que (u_n) converge vers ℓ pour N_1 . Dans ce cas $(N_1(u_n - \ell))$ tend vers 0. On

$$\forall n \in \mathbf{N}, N_2(u_n - \ell) \leq C_2 N_1(u_n - \ell)$$

Donc $N_2(u_n - \ell)$ tend aussi vers 0 et de ce fait, (u_n) tend vers ℓ pour N_2 .

□

Exercice : On suppose que l'on a juste N_1 plus fine que N_2 .

- A-t-on X bornée pour N_1 implique X bornée pour N_2 ? Dans l'autre sens?
- A-t-on (u_n) converge pour N_1 implique que (u_n) converge pour N_2 ? Dans l'autre sens?

★ **Méthode :** Montrer que des normes ne sont pas équivalentes à l'aide de suites.

Pour montrer que deux normes ne sont pas équivalentes, il suffit de trouver une suite (u_n) qui soit bornée (resp. convergente) pour une norme et pas pour l'autre.

Exemple : On a construit une suite de fonctions sur $[0, 1]$ qui était bornée pour $\|\cdot\|_1$ mais pas pour $\|\cdot\|_\infty$. Ces normes ne sont pas équivalentes.

Notons que l'on a que

$$\forall f \in \mathcal{B}([0, 1], \mathbf{R}), \|f\|_1 = \int_0^1 |f(t)| dt \leq (1 - 0) \sup_{t \in [0, 1]} |f(t)| = \|f\|_\infty$$

De ce fait, toute partie bornée pour $\|\cdot\|_\infty$ l'est pour $\|\cdot\|_1$.

Corollaire 9.32

Soit E un espace vectoriel de dimension finie. Soit $\mathcal{B} = (e_1, \dots, e_n)$. Soit $(u_n) \in E^{\mathbf{N}}$.

1. Elle est bornée si et seulement si pour tout $i \in \llbracket 1 ; n \rrbracket$, $(e_i^*(u_n))$ est bornée.
2. Elle est convergente si et seulement si pour tout $i \in \llbracket 1 ; n \rrbracket$, $(e_i^*(u_n))$ est convergente. De plus, dans ce cas, si on note ℓ la limite de (u_n) , $e_i^*(\ell) = \lim(e_i^*(u_n))$.

Démonstration : Il suffit de prendre pour norme sur E (car on a le choix) la norme infinie par rapport à la base $\mathcal{B}$. La démonstration est alors similaire à la démonstration de la convergence dans un espace produit. $\square$

Exemples :

1. Soit (A_p) une suite de $\mathcal{M}_n(\mathbf{K})$. On note pour tout p , $A_p = (a_{ij}(p))$.

La suite (A_p) converge vers $M = (m_{ij})$ si et seulement si pour tout $(i, j) \in \llbracket 1; n \rrbracket^2$, $a_{ij}(p) \xrightarrow{p \rightarrow \infty} m_{ij}$.

2. Soit A une matrice de $\mathcal{M}_n(\mathbf{K})$. On la suppose diagonalisable et $\text{Sp}(A) \subset \{z \in \mathbf{C} \mid |z| < 1\}$.

Il existe alors $P \in \text{GL}_n(\mathbf{K})$ telle que $P^{-1}AP = D$ où D est une matrice diagonale dont les coefficients diagonaux sont de module strictement inférieurs à 1. On les note $\lambda_1, \dots, \lambda_n$.

Maintenant, $D^n = (P^{-1}AP)^n = P^{-1}A^nP$ et donc $A^n = PD^nP^{-1}$. On en déduit que chaque coefficient que A^n est une combinaison à coefficients fixés des termes de la forme $\lambda_i^n \xrightarrow{n \rightarrow \infty} 0$. Donc $A^n \xrightarrow{n \rightarrow \infty} 0$. On peut aussi montrer que $\|D^n\|_\infty = (\max |\lambda_i|)^n \rightarrow 0$ et que $\|A^n\|_\infty = \|PD^nP^{-1}\|_\infty \leq n^2 \|D^n\|_\infty$.

3. On reprend la même situation mais on suppose maintenant que 1 est aussi valeur propre. En procédant de même on voit que D^n va tendre vers

$$D = \begin{pmatrix} 1 & 0 & \cdots & \cdots & \cdots & 0 \\ 0 & \ddots & \ddots & & & \vdots \\ \vdots & \ddots & 1 & \ddots & & \vdots \\ \vdots & & \ddots & 0 & \ddots & \vdots \\ \vdots & & & \ddots & \ddots & 0 \\ 0 & \cdots & \cdots & \cdots & 0 & 0 \end{pmatrix}.$$

De ce fait, en procédant, coordonnées par coordonnées, on voit que A^n va tendre vers $M = PDP^{-1}$. C'est un projecteur car $M^2 = M$. C'est le projecteur sur $E_1(A)$ parallèlement à $\bigoplus_{\lambda \in \text{Sp}(A) \setminus \{1\}} E_\lambda(A)$.

4. (CLASSIQUE - A SAVOIR FAIRE) Soit $A \in \mathcal{M}_n(\mathbf{C})$. On veut montrer qu'il existe une suite (A_p) de matrices *inversibles* telle que $(A_p) \rightarrow A$. Il suffit de poser (x_p) une suite de complexes tendant vers 0 dont le support ne rencontre pas $\text{Sp}(A)$. On peut par exemple prendre $x_p = \frac{\delta}{p+1}$ où $\delta \in \mathbf{R}_+$ est strictement inférieur au module des éléments non nuls de $\text{Sp}(A)$. On pose alors $A_p = A - x_p I_n$.

Exercice : Montrer que si $A \in \mathcal{M}_n(\mathbf{C})$ il existe une suite (A_p) de matrices *diagonalisables* telle que $(A_p) \rightarrow A$.

3.2 Propriétés

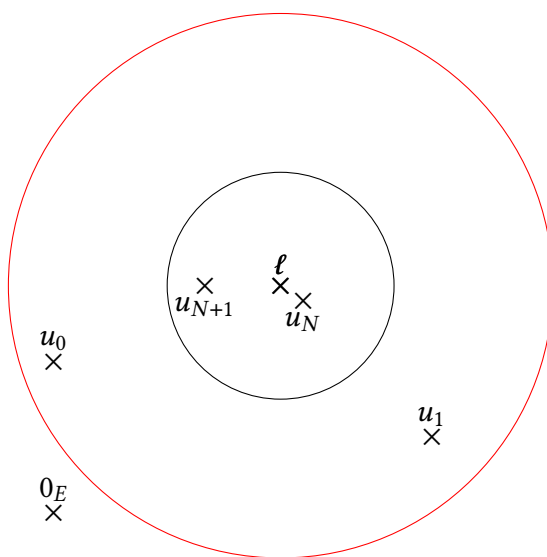
Proposition 9.33

Une suite convergente est bornée.

Démonstration : Soit (u_n) une suite tendant vers ℓ . Pour $\varepsilon = 1$, il existe un rang N tel que pour $n \geq N$, $u_n \in \overline{B}(\ell, \varepsilon)$. Maintenant on considère les éléments $u_0, \dots, u_{N-1}$ qui sont en nombre fini. On note

$$\delta = \max_{i \in \llbracket 1; N-1 \rrbracket} \|u_i - \ell\|.$$

On a alors que tous les termes de la suite appartiennent à $\overline{B}(\ell, \delta')$ où $\delta' = \max(1, \delta)$. La suite (u_n) est donc bornée.



□

Proposition 9.34

L'ensemble des suites convergentes est un sous-espace vectoriel de $E^{\mathbb{N}}$. De plus l'application qui associe à une suite convergente sa limite est linéaire. C'est-à-dire que si λ, μ sont deux scalaires et que (u_n) et (v_n) sont deux suites convergentes alors $(\lambda u_n + \mu v_n)$ converge et

$$\lim(\lambda u_n + \mu v_n) = \lambda \lim(u_n) + \mu \lim(v_n).$$

Démonstration : En reprenant les notations de l'énoncé. On pose $\ell = \lim(u_n)$ et $\ell' = \lim(v_n)$. Maintenant, pour tout $\varepsilon' > 0$, il existe N, N' tels que

$$n \geq N \rightarrow \|u_n - \ell\| \leq \varepsilon' \text{ et } n \geq N' \Rightarrow \|v_n - \ell'\| \leq \varepsilon'.$$

De ce fait, par inégalité triangulaire, pour

$$n \geq \max(N, N') \Rightarrow \|(\lambda u_n + \mu v_n) - (\lambda \ell + \mu \ell')\| \leq |\lambda| \|u_n - \ell\| + |\mu| \|v_n - \ell'\| \leq (|\lambda| + |\mu|) \varepsilon'.$$

De ce fait, pour tout $\varepsilon > 0$, on obtient ce que l'on veut en prenant $\varepsilon' = \frac{\varepsilon}{|\lambda| + |\mu|}$ si $|\lambda| + |\mu| \neq 0$ (ce dernier cas est évident). On a bien $\lim(\lambda u_n + \mu v_n) = \lambda \lim(u_n) + \mu \lim(v_n)$. □

Proposition 9.35

Soit $(u_n) \in E^{\mathbb{N}}$ une suite convergente de limite ℓ .

1. La suite réelle $(\|u_n\|)$ converge vers $\|\ell\|$.
2. Soit (λ_n) une suite de scalaire convergente vers Λ . La suite $(\lambda_n u_n)$ converge vers $\Lambda \ell$.

Démonstration :

1. Il suffit d'utiliser la deuxième inégalité triangulaire :

$$|(\|u_n\| - \|\ell\|)| \leq \|u_n - \ell\|$$

2. Il suffit de montrer que la suite $(\|\lambda_n u_n - \Lambda \ell\|)$ tend vers 0. Maintenant pour tout entier n ,

$$\|\lambda_n u_n - \Lambda \ell\| = \|(\lambda_n u_n - \Lambda u_n) + \Lambda(u_n - \ell)\| \leq |\lambda_n - \Lambda| \|u_n\| + |\Lambda| \|u_n - \ell\|.$$

Or la suite (u_n) étant convergente elle est bornée, de ce fait $(|\lambda_n - \Lambda| \|u_n\|)$ tend vers 0. De même $(|\Lambda| \|u_n - \ell\|)$ tend vers 0. On en déduit que la suite $(\lambda_n u_n)$ converge vers $\Lambda \ell$.

□

Proposition 9.36

Soit $(\mathcal{A}, \|\cdot\|)$ une algèbre muni d'une norme telle qu'il existe $C > 0$ vérifiant que pour tout x, y dans $\mathcal{A}$, $\|x \times y\| \leq C \|x\| \cdot \|y\|$. Soit (u_n) et (v_n) deux suites convergentes vers ℓ et ℓ' . La suite $(u_n v_n)$ converge et

$$\lim(u_n v_n) = \ell \ell'.$$

Démonstration : Il suffit de voir que

$$\|u_n v_n - \ell \ell'\| = \|u_n(v_n - \ell') + (u_n - \ell)\ell'\| \leq \|u_n(v_n - \ell')\| + \|(u_n - \ell)\ell'\|.$$

Maintenant avec l'hypothèse sur la norme, il existe une constante C telle que $\forall (x, y) \in \mathcal{A}^2$, $\|x \cdot y\| \leq C \|x\| \|y\|$. On en déduit que

$$\|u_n v_n - \ell \ell'\| \leq C \cdot M \|v_n - \ell'\| + C \cdot \|\ell'\| \|u_n - \ell\|$$

où M est un majorant de $\|u_n\|$ (la suite (u_n) étant bornée car convergente). On en déduit que $\|u_n v_n - \ell \ell'\| \xrightarrow[n \rightarrow \infty]{} 0$.

□

Exemple : Si on se donne une suite (A_n) de matrice. On suppose que la suite converge vers une matrice A pour une norme qui vérifie l'hypothèse (par exemple la norme infinie avec $C = n$). Soit P une matrice inversible et pour tout entier n on note $B_n = P^{-1} A_n P$. Alors la suite (B_n) converge vers $P^{-1} A P$.

3.3 Suites extraites et valeurs d'adhérence**Définition 9.37**

Soit $(u_n) \in E^{\mathbb{N}}$. On appelle suite extraite de (u_n) ou sous-suite de (u_n) une suite de la forme $(u_{\varphi(n)})$ où φ est une application strictement croissante de $\mathbb{N}$ dans $\mathbb{N}$.

Remarques :

1. La fonction φ s'appelle une extractrice.
2. Cela revient à considérer la suite obtenue en enlevant certains termes (il faut juste en garder une infinité).

Proposition 9.38

Soit $(u_n) \in E^{\mathbb{N}}$. Si (u_n) converge vers ℓ alors toute suite extraite (v_n) converge encore vers ℓ .

Démonstration : Il suffit d'utiliser le résultat analogue sur les suites réelles en voyant que la suite réelle $(\|u_{\varphi(n)} - \ell\|)$ est une suite extraite de la suite $(\|u_n - \ell\|)$ qui tend vers 0. $\square$

Proposition 9.39

Soit $(u_n) \in E^{\mathbb{N}}$ une suite. On considère la suite des termes pairs (u_{2n}) et la suite des termes impairs (u_{2n+1}) . Les suites (u_{2n}) et (u_{2n+1}) convergent vers la même limite ℓ si et seulement si la suite (u_n) converge vers ℓ .

Démonstration :

- $\Leftarrow$ On suppose que (u_n) converge vers ℓ alors (u_{2n}) et (u_{2n+1}) convergent aussi vers ℓ en utilisant la proposition précédente.
- $\Rightarrow$ On suppose que (u_{2n}) et (u_{2n+1}) convergent vers la même limite ℓ . Dès lors si on se donne un réel positif ε . On sait qu'il existe des entiers N_1 et N_2 tels que pour tout entier n , si $n \geq N_1$ on a $\|u_{2n} - \ell\| \leq \varepsilon$ et si $n \geq N_2$ on a $\|u_{2n+1} - \ell\| \leq \varepsilon$. On pose alors $N = 2 \cdot \max(N_1, N_2) + 1$. Soit n un entier, on suppose que $n \geq N$. Si n est pair alors $n = 2p$ et

$$p = \frac{n}{2} \geq \frac{N}{2} \geq N_1.$$

On en déduit que $\|u_n - \ell\| = \|u_{2p} - \ell\| \leq \varepsilon$.

De même si n est impair, on a $n = 2p + 1$ et

$$p = \frac{n-1}{2} \geq \frac{N-1}{2} \geq N_2.$$

On en déduit encore que $\|u_n - \ell\| = \|u_{2p+1} - \ell\| \leq \varepsilon$.

$\square$

Définition 9.40

Soit $(u_n) \in E^{\mathbb{N}}$. Une valeur d'adhérence de la suite (u_n) est un élément x de E tel que :

$$\forall \varepsilon > 0, \forall N \in \mathbb{N}, \exists n \geq N, u_n \in B(x, \varepsilon).$$

Remarque : Cela signifie que dans une boule ouverte $B(x, \varepsilon)$ centrée en x , il existe « aussi loin que l'on veut » un élément de la suite (u_n) ¹. Intuitivement, une valeur d'adhérence d'une suite est un élément x tel que la suite « passera » une infinité de fois près de x .

1. et même une infinité car si on trouve u_n avec $n \geq N$ dans $B(x, \varepsilon)$, on peut recommencer en remplaçant N par $n + 1$ et ainsi de suite

Proposition 9.41 (Caractérisation des valeurs d'adhérence d'une suite)

Soit $(u_n) \in E^{\mathbb{N}}$ et $x \in E$. L'élément x est une valeur d'adhérence si et seulement s'il existe une suite extraite $(u_{\varphi(n)})$ qui converge vers x .

Démonstration :

- $\Leftarrow$ On suppose que x est une valeur d'adhérence et on note (v_n) une suite extraite de (u_n) qui tend vers x . Soit $N \in \mathbb{N}$ et $\varepsilon > 0$, La suite extraite (v'_n) de (v_n) obtenue en ne gardant que les termes de la forme u_k avec $k \geq N$ tend encore x . Elle a donc des termes dans la boule $B(x, \varepsilon)$.
- $\Rightarrow$ On suppose que $\forall \varepsilon > 0, \forall N \in \mathbb{N}, \exists n \geq N, u_n \in B(x, \varepsilon)$. On va construire une suite extraite $(v_n) = (u_{\varphi(n)})$ telle que $(v_n) \rightarrow x$. Pour $\varepsilon = \frac{1}{0+1}$, il existe $n_0 = \varphi(0)$ tel que $u_{n_0} \in B(x, \frac{1}{1})$.

De même pour $\varepsilon = \frac{1}{1+1}$, il existe $n_1 = \varphi(1) > \varphi(0)$ tel que $u_{n_1} \in B(x, \frac{1}{1+1})$. De proche en proche, pour tout entier p , on peut construire $\varphi(p) > \varphi(p-1) > \dots > \varphi(0)$ tel que $u_{\varphi(p)} \in B(x, \frac{1}{p+1})$. La suite construite converge vers x .

□

Exemple : La suite $u_n = (-1)^n + \frac{1}{n+1}$ a deux valeurs d'adhérence $\{\pm 1\}$.

Proposition 9.42

Une suite convergente a une unique valeur d'adhérence : sa limite

Démonstration : C'est juste un corolaire du fait que toute suite extraite d'une suite convergente converge vers cette limite. □

Remarque : On peut utiliser la contraposée pour montrer qu'une suite diverge. En effet toute suite ayant au moins deux valeurs d'adhérence diverge.

Exercice : Trouver une suite divergente ayant une unique valeur d'adhérence. Peut-on en trouver une qui soit bornée ?

4 Séries à valeurs dans un espace vectoriel de dimension finie

4.1 Généralités

Définition 9.43

Soit E un espace vectoriel de dimension finie. Soit $(\sum u_n)$ une série d'éléments de E .

1. On dit que la série converge si la suite des sommes partielles converge.
2. On dit que la série converge absolument si la série numérique $(\sum \|u_n\|)$ converge.

Remarque : Toutes les normes étant équivalentes, l'absolue convergence ne dépend pas de la norme choisie.

Théorème 9.44

Soit $(\sum u_n)$ une série d'éléments de E absolument convergente. Elle converge.

Démonstration : On considère une série $(\sum u_n)$ absolument convergente. On se fixe une base $\mathcal{B} = (e_1, \dots, e_p)$ de E . En particulier, la série $\sum \|u_n\|_\infty$ est convergente. Comme pour tout $i \in \llbracket 1; p \rrbracket$ et tout n , $|e_i^*(u_n)| \leq \|u_n\|_\infty$ les séries $\sum e_i^*(u_n)$ sont absolument convergentes donc convergentes. Si on note (S_n) la suite des sommes partielles, on vient d'établir que pour tout $i \in \llbracket 1; p \rrbracket$, $e_i^*(S_n) = \sum_{k=0}^n e_i^*(u_k)$ converge donc (S_n) converge. $\square$

ATTENTION

Ce théorème n'est pas vrai (en général) dans un espace vectoriel normé de dimension infinie. Par exemple pour $E = \mathbf{K}[X]$ avec la norme infinie $\|\cdot\|_\infty$. On pose pour tout entier n non nul $P_n = \frac{1}{n^2}X^n$. Il est clair que $\|P_n\|_\infty = \frac{1}{n^2}$ et donc la série $\sum_{n \geq 0} \|P_n\|_\infty$ converge. Par contre la série $\sum_{n \geq 0} P_n$ diverge. En effet, supposons par l'absurde qu'elle convergeait et notons Q la somme de la série et $d = \deg(Q)$. Pour tout entier $N > d$,

$$e_{d+1}^* \left(\sum_{n=1}^N P_n - Q \right) = e_{d+1}^*(P_{d+1})$$

On en déduit que

$$\left\| \sum_{n=1}^N P_n - Q \right\|_\infty \geq \frac{1}{(d+1)^2}$$

On a bien une absurdité.

Remarque : On peut généraliser ce théorème en disant que si $(E, \|\cdot\|)$ est *complet* alors l'absolue convergence d'une série implique sa convergence car il est simple de montrer que si $\sum u_n$ est absolument convergente alors la suite des sommes partielles est une suite de Cauchy.

4.2 Série géométrique de matrices

On se place dans $E = \mathcal{M}_n(\mathbf{K})$. On veut étudier les séries géométriques. On veut donc pouvoir « estimer » $\|A^p\|$

Lemme 9.45

Soit $\|\cdot\|$ une norme sur E , il existe C tel que

$$\forall (A, B) \in E^2, \|AB\| \leq C\|A\|\|B\|.$$

Démonstration : On utilise que pour $\|\cdot\|_\infty$, $\|AB\|_\infty \leq n\|A\|_\infty\|B\|_\infty$ et que toutes les normes sont équivalentes. $\square$

Remarque : On verra qu'il existe aussi des normes qui vérifient

$$\forall (A, B) \in \mathcal{M}_n(\mathbf{K})^2, \|AB\| \leq \|A\| \|B\|$$

Dans toute la suite on notera C une constante réelle telle que

$$\forall (A, B) \in E^2, \|AB\| \leq C \|A\| \|B\|.$$

Proposition 9.46

Soit $A \in \mathcal{M}_n(\mathbf{K})$, si $\|A\| < \frac{1}{C}$ alors la série géométrique $(\sum A^p)$ converge absolument. De plus sa somme vaut

$$\sum_{p=0}^{+\infty} A^p = (I_n - A)^{-1}$$

Démonstration : Pour la convergence, on voit que

$$\|A^p\| \leq C \|A^{p-1}\| \|A\| \leq C^2 \|A^{p-2}\| \|A\|^2 \leq \dots \leq C^{p-1} \|A\| \|A\|^{p-1} = C^{p-1} \|A\|^p$$

Or $C^{p-1} \|A\|^p = \frac{1}{C} (C \|A\|)^p$ donc si $\|A\| < \frac{1}{C}$ alors la série $\left(\sum \frac{1}{C} (C \|A\|)^p \right)$ converge et la série $(\sum A^p)$ converge absolument.

Pour la somme, Il suffit de remarquer que

$$(I_n - A) \sum_{p=0}^{+\infty} A^p = \sum_{p=0}^{+\infty} A^p - \sum_{p=0}^{+\infty} A^{p+1} = \sum_{p=0}^{+\infty} A^p - \sum_{p=1}^{+\infty} A^p = I_n.$$

□

Remarques :

1. On montre ainsi que si A a une norme « petite » alors $I_n - A$ est inversible. Ce n'est qu'une condition suffisante, on peut par exemple trouver des matrices nilpotentes de très grande norme et on a encore $I_n - N$ inversible d'inverse $\sum_{p=0}^{+\infty} N^p$ qui converge car la suite stationne à 0.
2. On peut aussi considérer la série $(\sum u^p)$ pour $u \in \mathcal{L}(E)$.

5 Applications linéaires lipschitziennes

5.1 Définitions

Définition 9.47 (Applications lipschitziennes)

Soit (E, N_E) et (F, N_F) deux espaces vectoriels normés. Soit f une application de E dans F .

1. Soit $k \in \mathbf{R}_+$. L'application f est dite k -lipschitzienne si

$$\forall (x, y) \in E^2, N_F(f(x) - f(y)) \leq k N_E(x - y)$$

2. L'application f est dite lipschitzienne s'il existe un réel positif k tel que f soit k -lipschitzienne.

Remarque : Cette définition est compatible avec la définition d'applications lipschitziennes vue en première année sur les fonctions de $\mathbf{R}$ dans $\mathbf{R}$.

ATTENTION

Dire que f est lipschitzienne signifie que

$$\exists k > 0, \forall (x, y) \in E^2, N_F(f(x) - f(y)) \leq k N_E(x - y)$$

ce qu'il ne faut pas confondre avec

$$\forall (x, y) \in E^2, \exists k > 0, N_F(f(x) - f(y)) \leq k N_E(x - y).$$

Proposition 9.48

Soit f une application lipschitzienne de E dans F . Soit (u_n) une suite de E^N . Si (u_n) converge vers ℓ alors $(f(u_n))$ converge vers $f(\ell)$.

Remarque : C'est un premier pas vers la notion de fonctions continues.

Démonstration : On suppose que f est k -lipschitzienne. On a donc pour tout entier n ,

$$0 \leq N_F(f(\ell) - f(u_n)) \leq k N_E(\ell - u_n).$$

Par encadrement, on en déduit que la suite $(N_F(f(\ell) - f(u_n)))$ tend vers 0 et donc $(f(u_n))$ tend vers $f(\ell)$.

□

Nous utiliserons cette notion essentiellement pour les applications linéaires. Nous avons alors une caractérisation plus simple.

Proposition 9.49

Soit (E, N_E) et (F, N_F) deux espaces vectoriels normés. Soit f une application **linéaire** de E dans F .

Elle est lipschitzienne si et seulement s'il existe $k > 0$ tel que pour tout x de E , $N_F(f(x)) \leq k N_E(x)$.

Démonstration :

- $\Rightarrow$ Il suffit d'appliquer la définition de f lipschitzienne en prenant $y = 0_E$ car dans ce cas $f(y) = 0_F$.
- $\Leftarrow$ On suppose qu'il existe $k > 0$ telle que $\forall x \in E, N_F(f(x)) \leq k N_E(x)$. Pour tout couple $(x, y) \in E^2$

$$N_F(f(x) - f(y)) = N_F(f(x - y)) \leq k N_E(x - y).$$

□

Notation : On note $\mathcal{L}_c(E, F)$ l'ensemble des applications linéaires lipschitziennes de E dans F . La notation vient du fait que ces applications seront les applications linéaires continues.

Définition 9.50 (Norme subordonnée)

Soit $f : (E, N_E) \rightarrow (F, N_F)$ une application linéaire lipschitzienne, on appelle norme subordonnée de f et on note $|||f|||$ ou $\|f\|_{op}$ la meilleure constante telle que pour tout x , $N_F(f(x)) \leq k N_E(x)$. On a donc

$$|||f||| = \sup \left\{ \frac{N_F(f(x))}{N_E(x)}, x \in E \setminus \{0\} \right\} = \sup \left\{ N_F(f(x)), x \in \overline{B}(0, 1) \right\} = \sup \{ N_F(f(x)), x \in S(0, 1) \}$$

Remarques :

1. Par définition, pour $x \in E \setminus \{0\}$, $N_F(f(x)) \leq \|f\|_{op} N_E(x)$ et cette inégalité est aussi vérifiée si $x = 0$.

2. Vérifions que les trois bornes supérieures sont les mêmes.

$$- \text{ Comme } S(0, 1) \subset \overline{B}(0, 1), \sup \{ N_F(f(x)), x \in S(0, 1) \} \leq \sup \{ N_F(f(x)), x \in \overline{B}(0, 1) \}.$$

De plus, pour tout $x \in \overline{B}(0, 1) \setminus \{0\}$, on peut poser $x' = \frac{x}{N_E(x)} \in S(0, 1)$. On a alors $x = N_E(x)x'$ et donc par linéarité de f , $f(x) = N_E(x)f(x')$ et donc

$$N_F(f(x)) = N_E(x)N_F(f(x')) \leq N_F(f(x')) \leq \sup \{ N_F(f(x)), x \in S(0, 1) \}$$

L'inégalité est aussi vraie pour $x = 0_E$. On en déduit que

$$\sup \{ N_F(f(x)), x \in \overline{B}(0, 1) \} \leq \sup \{ N_F(f(x)), x \in S(0, 1) \}$$

Par double inégalités,

$$\sup \{ N_F(f(x)), x \in \overline{B}(0, 1) \} = \sup \{ N_F(f(x)), x \in S(0, 1) \}$$

- Soit $x \in S(0, 1)$,

$$N_F(f(x)) = \frac{N_F(f(x))}{N_E(x)} \leq \sup \left\{ \frac{N_F(f(x))}{N_E(x)}, x \in E \setminus \{0\} \right\}$$

$$\text{On en déduit que } \sup \{ N_F(f(x)), x \in S(0, 1) \} \leq \sup \left\{ \frac{N_F(f(x))}{N_E(x)}, x \in E \setminus \{0\} \right\}.$$

Inversement, pour $x \in E \setminus \{0\}$ on peut poser $x' = \frac{x}{N_E(x)} \in S(0, 1)$ et par homogénéité de la norme et linéarité de f ,

$$\frac{N_F(f(x))}{N_E(x)} = N_F \left(f \left(\frac{x}{N_E(x)} \right) \right) = N_F(f(x')) \leq \sup \{ N_F(f(x)), x \in S(0, 1) \}$$

Cela montre l'inégalité inverse et finalement,

$$\sup \left\{ \frac{N_F(f(x))}{N_E(x)}, x \in E \setminus \{0\} \right\} = \sup \{ N_F(f(x)), x \in S(0, 1) \}$$

Exemples :

1. Soit $A \in \mathcal{M}_n(\mathbf{R})$. On considère $a : \mathcal{M}_{n,1}(\mathbf{R}) \rightarrow \mathcal{M}_{n,1}(\mathbf{R})$ l'endomorphisme canoniquement associé à A . Soit $\|\cdot\|$ une norme sur $\mathcal{M}_{n,1}(\mathbf{R})$. La norme subordonnée de a est

$$\|a\|_{op} = \sup \{ \|a(X)\|, X \in \overline{B}(0, 1) \} = \sup \{ \|AX\|, X \in \overline{B}(0, 1) \}$$

Traitons par exemple le cas où $\|\cdot\|$ est la norme ∞ .

Soit $X \in \mathcal{M}_{n,1}(\mathbf{R})$ tel que $\|X\|_\infty = 1$. Pour tout $i \in \llbracket 1 ; n \rrbracket$,

$$|(AX)[i]| = \left| \sum_{j=1}^n A[i, j]X[j] \right| \leq \sum_{j=1}^n |A[i, j]| |X[j]| \leq \left(\sum_{j=1}^n |A[i, j]| \right) \|X\|_\infty \leq \sum_{j=1}^n |A[i, j]|$$

En notant $L_i = \sum_{j=1}^n |A[i, j]|$, on voit que $\|AX\|_\infty \leq \max_i L_i$.

Cela montre que a est lipschitzienne et que $\|a\|_{\text{op}} \leq \max_i L_i$.

Réciproquement, nous allons construire $X \in \mathcal{M}_{n,1}(\mathbf{R})$ tel que $\|X\|_\infty = 1$ et $\|AX\|_\infty = \max_i L_i$.

Notons i_0 un entier compris entre 1 et n tel que $L_{i_0} = \max_i L_i$. On pose

$$X[j] = \begin{cases} 1 & \text{si } A[i_0, j] \geq 0 \\ -1 & \text{sinon} \end{cases}$$

On a alors

$$|(AX)[i_0]| = \left| \sum_{j=1}^n A[i_0, j]X[j] \right| = \sum_{j=1}^n |A[i_0, j]| = L_{i_0} = \max_i L_i$$

On a donc montré que $\|a\|_{\text{op}} = \max_i L_i$.

Exercice : Adapter le calcul au cas où $A \in \mathcal{M}_n(\mathbf{C})$.

2. Soit $A \in \mathcal{M}_n(\mathbf{K})$. On considère l'application linéaire

$$\begin{aligned} \varphi_A : \mathcal{M}_n(\mathbf{K}) &\rightarrow \mathcal{M}_n(\mathbf{K}) \\ M &\mapsto AM - MA \end{aligned}$$

On utilise la norme infinie sur $\mathcal{M}_n(\mathbf{K})$ (aussi bien au départ qu'à l'arrivée). On note $A = (a_{ij})$ et $M = (m_{ij})$. On voit que, en notant $\varphi(M) = AM - MA = (c_{ij})$, pour tout $(i, j) \in \llbracket 1 ; n \rrbracket^2$:

$$|c_{ij}| = \left| \sum_{k=1}^n a_{ik}m_{kj} - \sum_{k=1}^n m_{ik}a_{kj} \right| \leq 2n\|A\|_\infty\|M\|_\infty$$

On a donc $\|\varphi(M)\|_\infty \leq 2n\|A\|_\infty\|M\|_\infty$. L'application est donc lipschitzienne en prenant $k = 2n\|A\|_\infty$ et $\|\varphi\|_{\text{op}} \leq 2n\|A\|_\infty$.

3. Soit Δ défini sur $\mathbf{K}[X]$ par :

$$\begin{aligned} \Delta : \mathbf{K}[X] &\rightarrow \mathbf{K}[X] \\ P &\mapsto P' \end{aligned}$$

On utilise la norme $\|P\| = \sup_{t \in [-1, 1]} |\tilde{P}(t)|$. Il est clair que les vecteurs de la base canonique sont de norme 1. On en déduit alors que

$$\|\Delta(X^k)\| = \|kX^{k-1}\| = k\|X^{k-1}\| = k.$$

De ce fait Δ n'est pas lipschitzienne car $\frac{\|\Delta(P)\|}{\|P\|}$ n'est pas borné.

Proposition 9.51

1. Soit (E, N_E) et (F, N_F) deux espaces vectoriels normés. L'ensemble $\mathcal{L}_c(E, F)$ est un sous-espace vectoriel de $\mathcal{L}(E, F)$.
2. Soit (E, N_E) , (F, N_F) et (G, N_G) trois espaces vectoriels normés. Si $f \in \mathcal{L}_c(E, F)$ et $g \in \mathcal{L}_c(F, G)$ alors $g \circ f \in \mathcal{L}_c(E, G)$.

Démonstration :

En exercice.

□

Proposition 9.52

1. Soit (E, N_E) et (F, N_F) , $\| \cdot \|_{\text{op}}$ est une norme sur $\mathcal{L}_c(E, F)$.
2. Soit (E, N_E) , (F, N_F) et (G, N_G) trois espaces vectoriels normés. Si $f \in \mathcal{L}_c(E, F)$ et $g \in \mathcal{L}_c(F, G)$ alors

$$\|g \circ f\|_{\text{op}} \leq \|g\|_{\text{op}} \times \|f\|_{\text{op}}$$

Démonstration :

1. Vérifions les axiomes

- Soit $f \in \mathcal{L}_c(E, F)$, $\|f\|_{\text{op}} \geq 0$.
- Soit $f \in \mathcal{L}_c(E, F)$, telle que $\|f\|_{\text{op}} = 0$. Pour tout $x \in E \setminus \{0\}$, $\frac{N_F(f(x))}{N_E(x)} = 0$ donc $N_F(f(x)) = 0$ ce qui implique que $f(x) = 0$. Comme on a aussi $f(0) = 0$ alors $f = 0$.
- Soit $f \in \mathcal{L}_c(E, F)$ et $\lambda \in \mathbb{K}$,

$$\begin{aligned} \|\lambda f\|_{\text{op}} &= \sup\{N_F(\lambda f(x)) , x \in S(0, 1)\} \\ &= \sup\{|\lambda| N_F(f(x)) , x \in S(0, 1)\} \\ &= |\lambda| \sup\{N_F(f(x)) , x \in S(0, 1)\} \\ &= |\lambda| \|f\|_{\text{op}} \end{aligned}$$

- Soit f, g dans $\mathcal{L}_c(E, F)$ et $x \in S(0, 1)$,

$$N_F((f+g)(x)) \leq N_F(f(x)) + N_F(g(x)) \leq \|f\|_{\text{op}} + \|g\|_{\text{op}}$$

$$\text{Donc } \|f+g\|_{\text{op}} \leq \|f\|_{\text{op}} + \|g\|_{\text{op}}.$$

On a bien montré que $\| \cdot \|_{\text{op}}$ était une norme sur $\mathcal{L}_c(E, F)$.

2. Soit $x \in S(0, 1)$,

$$N_G(g(f(x))) \leq \|g\|_{\text{op}} N_F(f(x)) \leq \|g\|_{\text{op}} \times \|f\|_{\text{op}}$$

$$\text{Donc } \|g \circ f\|_{\text{op}} \leq \|g\|_{\text{op}} \times \|f\|_{\text{op}}.$$

□

Proposition 9.53

Soit (E, N_1) et (F, N_2) deux espaces vectoriels normés. On suppose que E est de dimension finie.

Toute application linéaire de E dans F est lipschitzienne.

Démonstration : Admis (pour le moment)

□

Remarques :

1. Dans la proposition, on peut choisir comme on veut les normes de E et F . De fait, toutes les normes de E sont équivalentes (car E est de dimension finie). Les normes de F ne sont pas nécessairement équivalentes (car F n'est pas nécessairement de dimension finie) mais elles le sont toutes sur $\text{Im}(f)$.
2. En appliquant cela à id_E on retrouve que toutes les normes sont équivalentes sur un espace vectoriel de dimension finie.

Séries entières

1	Généralités et rayon de convergence	272
1.1	Definition	272
1.2	Rayon de convergence	273
1.3	Utilisation de la règle de d'Alembert	276
1.4	Théorèmes de comparaison	278
1.5	Operations	279
2	Étude de la somme d'une série entière	281
2.1	Convergence normale et continuité	281
2.2	Primitivation et dérivation	282
3	Fonctions développables en séries entières	286
3.1	Définitions	286
3.2	Développements en série entière des fonctions usuelles	288
4	Utilisation des séries entières dans l'étude des équations différentielles	288
4.1	Rappels du cours de première année	289
4.2	Rappels - Équations linéaires du premier ordre	290
4.3	Rappels - Équations linéaires du deuxième ordre	292
4.4	Utilisation des séries entières	293

1 Généralités et rayon de convergence

1.1 Definition

Définition 10.1

On appelle série entière une série de fonctions $\sum_{n \geq 0} a_n z^n$ où les termes a_n sont des nombres complexes que l'on appelle les coefficients de la série entière.

La somme de cette série entière est la fonction $z \mapsto \sum_{n=0}^{+\infty} a_n z^n$.

Remarques :

1. On devrait noter plus rigoureusement $\sum_{n \geq 0} (z \mapsto a_n z^n)$ pour bien montrer qu'une série entière est une série de **fonctions** mais, pour simplifier, on acceptera l'abus de notations
2. Une série entière est donc juste une série de fonctions dont les termes sont d'une forme particulière. Tous les théorèmes vus sur les séries de fonctions vont pouvoir s'appliquer (sous une forme plus simple pour certains).
3. En fonction des cas, on travaillera avec des fonctions de la variable réelle ou des fonctions de la variable complexe.

Exemple : La série entière $\sum_{n \geq 0} \frac{z^n}{n!}$ est une série entière. Sa somme est

$$\exp : z \mapsto \sum_{n=0}^{+\infty} \frac{z^n}{n!}$$

1.2 Rayon de convergence

La première question qui se pose quand on étudie une série entière est celle de l'ensemble de définition de la fonction somme. On veut savoir pour quelles valeurs du paramètre (réel ou complexe) la série converge.

Lemme 10.2 (Lemme d'Abel)

Soit $\sum_{n \geq 0} a_n z^n$ une série entière.

On suppose qu'il existe z_0 tel que $(a_n z_0^n)_{n \geq 0}$ soit une suite bornée. Pour tout $z \in \mathbb{C}$ tel que $|z| < |z_0|$ la série $\sum_{n \geq 0} a_n z^n$ est absolument convergente.

Matheux (Niels Henrik Abel : 1802 - 1829)



Niels Henrik Abel est un mathématicien norvégien. Il est connu pour ses travaux en analyse mathématique sur la semi-convergence des séries numériques, des suites et séries de fonctions, les critères de convergence d'intégrale généralisée, sur la notion d'intégrale elliptique ; et en algèbre, sur la résolution des équations.

Il est en particulier à l'origine de la notion de nombres algébriques. Il a été le premier à donner une preuve de l'impossibilité de résoudre l'équation du cinquième degré par radicaux.

ATTENTION

La condition est que $|z|$ est **strictement** inférieur à $|z_0|$. En effet, si on prend par exemple pour tout entier n , $a_n = 1$ alors pour $z_0 = 1$ la suite $(a_n z_0^n)_{n \geq 0}$ est bornée cependant la série $\sum_{n \geq 0} a_n z^n$ ne converge pas pour $z = 1$.

Démonstration : Soit C tel que $\forall n \in \mathbb{N}, |a_n z_0^n| \leq C$. Pour z tel que $|z| < |z_0|$ on a

$$\forall n \in \mathbb{N}, |a_n z^n| = |a_n z_0^n| \times \left| \frac{z}{z_0} \right|^n \leq C \left| \frac{z}{z_0} \right|^n$$

Maintenant comme $|\frac{z}{z_0}| < 1$ (et ne dépend pas de n), la série $\sum_{n \geq 0} |a_n z^n|$ converge par comparaison avec une série géométrique. On en déduit que la série $\sum_{n \geq 0} a_n z^n$ est absolument convergente.

□

Définition 10.3 (Rayon de convergence)

Soit $\sum_{n \geq 0} a_n z^n$ une série entière.

1. On appelle rayon de convergence de la série entière la borne supérieure de l'ensemble des réels positifs ρ tels que la suite $(a_n \rho^n)_{n \geq 0}$ est bornée :

$$R = \sup\{\rho \in \mathbf{R}_+ \mid (a_n \rho^n)_{n \geq 0} \text{ est bornée}\} \in [0, +\infty]$$

2. On appelle disque de convergence (resp. intervalle de convergence) de la série entière le disque **ouvert** $D(0, R)$ de $\mathbf{C}$ (resp. l'intervalle **ouvert** $] - R, R[$)

Remarques :

1. L'ensemble $\{\rho \in \mathbf{R}_+ \mid (a_n \rho^n)_{n \geq 0} \text{ est bornée}\}$ n'est pas vide car il contient 0. Il a donc bien une borne supérieure dans $[0, +\infty]$.
2. Il se peut que $R = +\infty$ si pour tout $\rho \in \mathbf{R}_+$ la suite $(a_n \rho^n)_{n \geq 0}$ est bornée. Dans ce cas, le disque ouvert de convergence est $\mathbf{C}$ et l'intervalle ouvert de convergence est $\mathbf{R}$ en entier.

ATTENTION

Les coefficients a_n sont souvent complexes par contre ρ est un réel. De plus, dire que la suite $(a_n \rho^n)_{n \geq 0}$ est bornée est équivalent au fait que la suite $(|a_n| \rho^n)_{n \geq 0}$ est bornée.

Proposition 10.4

Soit $\sum_{n \geq 0} a_n z^n$ une série entière et R son rayon de convergence

1. Soit z dans le disque de convergence ($|z| < R$). La série $\sum_{n \geq 0} a_n z^n$ converge absolument.
2. Soit z tel que $|z| > R$, la série $\sum_{n \geq 0} a_n z^n$ diverge grossièrement.

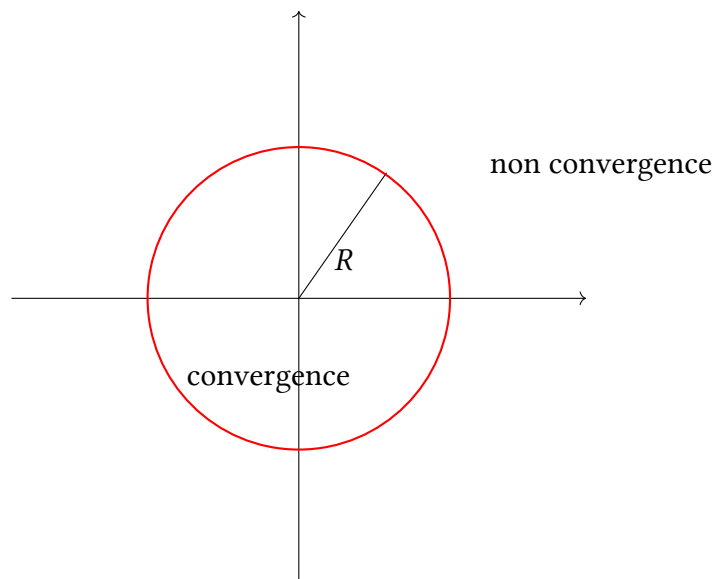
Remarques :

1. La proposition ci-dessus justifie la terminologie « disque / intervalle de convergence ».
2. On ne sait rien dans le cas général sur le comportement d'une série entière « au bord » du disque de convergence c'est-à-dire pour les z tels que $|z| = R$ (dans le cas où $R \neq +\infty$).

Démonstration :

1. On suppose que $|z| < R$. Par définition du rayon de convergence, il existe alors $\rho \in \mathbf{R}^+$ tel que $|z| < \rho$ et tel que la suite $(a_n \rho^n)$ est bornée. Il suffit alors d'utiliser le lemme d'Abel en prenant $z_0 = \rho$.
2. Soit z tel que $|z| > R$. Par définition du rayon de convergence, $(a_n z^n)_{n \geq 0}$ n'est pas bornée. De ce fait la suite $(a_n z^n)$ ne tend pas vers 0 et donc la série $\sum_{n \geq 0} a_n z^n$ diverge.

□



Corollaire 10.5

Soit $\sum_{n \geq 0} a_n z^n$ une série entière. On considère les ensembles :

- $SCA = \{\rho \in \mathbf{R}_+ \mid \text{la série } \sum a_n \rho^n \text{ converge absolument}\}$
- $SC = \{\rho \in \mathbf{R}_+ \mid \text{la série } \sum a_n \rho^n \text{ converge}\}$
- $C_0 = \{\rho \in \mathbf{R}_+ \mid \text{la suite } (a_n \rho^n)_{n \geq 0} \text{ tend vers } 0\}$
- $C = \{\rho \in \mathbf{R}_+ \mid \text{la série } (a_n \rho^n)_{n \geq 0} \text{ converge}\}$
- $B = \{\rho \in \mathbf{R}_+ \mid \text{la série } (a_n \rho^n)_{n \geq 0} \text{ est bornée}\}.$

On a $SCA \subset SC \subset C_0 \subset C \subset B$.

De plus $R = \sup(B) = \sup(C) = \sup(C_0) = \sup(SC) = \sup(SCA)$.

Démonstration :

Les inclusions, $SCA \subset SC \subset C_0 \subset C \subset B$ sont immédiates. On en déduit que

$$\sup(SCA) \leq \sup(SC) \leq \sup(C_0) \leq \sup(C) \leq \sup(B)$$

Supposons par l'absurde que $\sup(SCA) < \sup(B) = R$. Il existe alors $\rho \in \mathbf{R}_+$ tel que $\sup(SCA) < \rho < \sup(B)$. C'est absurde car si $\rho < R$ alors la série $\sum_{n \geq 0} a_n \rho^n$ est absolument convergente d'après la proposition précédente.

□

★ **Méthode :** (Caractérisation du rayon de convergence) Soit $\sum_{n \geq 0} a_n z^n$ une série entière et $\rho \in \mathbf{R}_+$.

- Si la suite $(a_n \rho^n)_{n \geq 0}$ est bornée, $R \geq \rho$. C'est le cas en particulier si $a_n \rho^n \rightarrow 0$ ou si la série $\sum_{n \geq 0} a_n \rho^n$ converge.
- Si la suite $(a_n \rho^n)_{n \geq 0}$ n'est pas bornée, $R \leq \rho$. C'est le cas en particulier si $|a_n| \rho^n \rightarrow +\infty$.
- Si la suite $(a_n \rho^n)_{n \geq 0}$ tend vers une limite finie non nulle, $R = \rho$ car $R \geq \rho$ d'après le premier point et que pour tout $\rho' > \rho$, $|a_n|(\rho')^n = |a_n| \rho^n \left(\frac{\rho'}{\rho}\right)^n \rightarrow +\infty$ et donc, d'après le deuxième point, $R \leq \rho'$.

Exemples :

1. On considère la série entière $\sum_{n \geq 0} \frac{z^n}{n!}$. Il est clair que pour tout réel ρ , la série $\sum_{n \geq 0} \frac{\rho^n}{n!}$ converge donc la suite $\left(\frac{\rho^n}{n!}\right)_{n \geq 0}$ tend vers 0 et donc est bornée. On a donc $R = +\infty$.
2. On considère la série entière $\sum_{n \geq 0} z^n$. Il est clair que pour tout réel $\rho < 1$, la suite $(\rho^n)_{n \geq 0}$ est bornée et pour $\rho > 1$ la suite $(\rho^n)_{n \geq 0}$ ne l'est plus donc $R = 1$.
3. On considère la série entière $\sum_{n \geq 0} \frac{z^n}{n+1}$. Là encore, par croissances comparées, pour tout réel $\rho < 1$, la suite $(\frac{\rho^n}{n+1})_{n \geq 0}$ est bornée car elle tend vers 0 et pour $\rho > 1$ la suite $(\frac{\rho^n}{n+1})_{n \geq 0}$ ne l'est plus donc $R = 1$.

Notons d'ailleurs que pour $z = 1$ qui vérifie $|z| = R = 1$ la série $\sum_{n \geq 0} \frac{1^n}{n+1}$ diverge alors que pour $z = -1$ qui vérifie aussi $|z| = R = 1$, la série $\sum_{n \geq 0} \frac{(-1)^n}{n+1}$. Cela illustre que sur le bord du disque de convergence, il peut y avoir ou ne pas y avoir convergence.

1.3 Utilisation de la règle de d'Alembert

Pour déterminer le rayon de convergence on va souvent utiliser la règle de d'Alembert pour déterminer la convergence ou pas d'une série.

Proposition 10.6 (Règle de d'Alembert - Rappel)

Soit $\sum u_n$ une série à termes strictement positifs. On suppose que $\frac{u_{n+1}}{u_n} \rightarrow \ell \in \overline{\mathbb{R}}_+$.

- Si $\ell > 1$ alors $\sum u_n$ diverge.
- Si $\ell < 1$ alors $\sum u_n$ converge.
- Si $\ell = 1$. On ne peut rien dire....

Commençons par un exemple.

Exemple : Soit $\sum_{n \geq 0} \frac{3^n z^n}{n+2}$ une série entière. On pose $a_n = \frac{3^n}{n+2}$.

Soit $\rho \in \mathbb{R}_+$, on pose $u_n = |a_n| \rho^n > 0$, on a alors

$$\frac{u_{n+1}}{u_n} = \frac{|a_{n+1}|}{|a_n|} \rho = 3 \frac{n+2}{n+3} \rho \xrightarrow{n \rightarrow \infty} 3\rho.$$

On en déduit que si $\rho > \frac{1}{3}$ la série $\sum_{n \geq 0} |a_n| \rho^n$ diverge et donc $R \leq \frac{1}{3}$. De même, si $\rho < \frac{1}{3}$ la série $\sum_{n \geq 0} |a_n| \rho^n$ converge et donc $R \geq \frac{1}{3}$. En conclusion $R = \frac{1}{3}$.

Corollaire 10.7 (Application de la règle de d'Alembert)

Soit $\sum_{n \geq 0} a_n z^n$ une série entière. On suppose que

- i) la suite $(a_n)_{n \geq 0}$ ne s'annule pas au voisinage de $+\infty$
- ii) $\lim_{n \rightarrow +\infty} \frac{|a_{n+1}|}{|a_n|} = \ell \in [0, +\infty]$

alors le rayon de convergence de la série est $R = \frac{1}{\ell}$.

Remarque : Dans l'énoncé ci-dessus, $\frac{1}{\ell}$ s'entend dans $[0, +\infty]$ en posant $\frac{1}{0} = +\infty$ et $\frac{1}{+\infty} = 0$.

Démonstration : On raisonne comme dans l'exemple ci-dessus.

Soit $\rho \in \mathbf{R}_+^*$, on pose $u_n = |a_n|\rho^n > 0$, on a alors. Avec les hypothèses, $\lim_{n \rightarrow +\infty} \frac{u_{n+1}}{u_n} = \ell\rho \in [0, +\infty]$.

- Si $\ell = 0$, pour tout $\rho \in \mathbf{R}_+$, $\lim_{n \rightarrow +\infty} \frac{u_{n+1}}{u_n} = \ell\rho = 0 < 1$. Cela montre que la série $\sum_{n \geq 0} |a_n|\rho^n$ converge toujours donc $R = +\infty$.
- Si $\ell = +\infty$, pour tout $\rho \in \mathbf{R}_+^*$, $\lim_{n \rightarrow +\infty} \frac{u_{n+1}}{u_n} = \ell\rho = +\infty > 1$. Cela montre que la série $\sum_{n \geq 0} |a_n|\rho^n$ ne converge jamais (sauf pour $\rho = 0$) donc $R = 0$.
- Si $\ell \in]0, +\infty[$. Pour $\rho > \frac{1}{\ell}$, $\lim_{n \rightarrow +\infty} \frac{u_{n+1}}{u_n} = \ell\rho > 1$. Donc la série $\sum_{n \geq 0} |a_n|\rho^n$ diverge. Cela implique que $\rho \geq R$. Mais ceci étant vrai pour toute valeur de ρ supérieure strictement à $\frac{1}{\ell}$, on peut « faire tendre ρ vers $\frac{1}{\ell}^+$ » pour obtenir que $\frac{1}{\ell} \geq R$.

De même, pour $\rho < \frac{1}{\ell}$, $\lim_{n \rightarrow +\infty} \frac{u_{n+1}}{u_n} = \ell\rho < 1$. Donc la série $\sum_{n \geq 0} |a_n|\rho^n$ converge. Cela implique que $\rho \leq R$. Mais ceci étant vrai pour toute valeur de ρ supérieure inférieure à $\frac{1}{\ell}$, on peut « faire tendre ρ vers $\frac{1}{\ell}^-$ » pour obtenir que $\frac{1}{\ell} \leq R$.

En conclusion $R = \frac{1}{\ell}$.

□

Exercice : Déterminer les rayons de convergence de

1. $\sum_{n \geq 0} n^2 2^n z^n$

2. $\sum_{n \geq 0} n! z^n$

3. $\sum_{n \geq 0} \frac{n^n}{n!} z^{3n}$.

ATTENTION

Ce résultat ne fonctionne plus sur les séries « lacunaires » comme la troisième de l'exemple ci-dessus

Corollaire 10.8

Soit $\alpha \in \mathbf{R}$. La série entière $\sum_{n \geq 1} n^\alpha z^n$ a un rayon de convergence égal à 1.

Démonstration : On pose $a_n = n^\alpha$. On vérifie bien que n^α ne s'annule pas si $n \geq 1$. De plus

$$\frac{(n+1)^\alpha}{n^\alpha} = \left(1 + \frac{1}{n}\right)^\alpha \xrightarrow{n \rightarrow +\infty} 1$$

On en déduit que le rayon de convergence vaut

$$R = \frac{1}{1} = 1$$

□

1.4 Théorèmes de comparaison

Proposition 10.9

Soit $\sum_{n \geq 0} a_n z^n$ et $\sum_{n \geq 0} b_n z^n$ deux séries entières. On note R_a et R_b leur rayon de convergence respectif.

1. Si $a_n = O(b_n)$ alors $R_a \geq R_b$. Ceci est en particulier vrai si $a_n = o(b_n)$.
2. Si $a_n \sim b_n$ alors $R_a = R_b$.

Démonstration :

1. On suppose que $a_n = O(b_n)$. Donc pour tout $\rho \in \mathbf{R}_+$, $a_n \rho^n = O(b_n \rho^n)$. De ce fait si la suite $(b_n \rho^n)_{n \geq 0}$ est bornée alors la suite $(a_n \rho^n)_{n \geq 0}$ aussi. On en déduit que

$$\{\rho \in \mathbf{R}_+ \mid (b_n \rho^n)_{n \geq 0} \text{ est bornée}\} \subset \{\rho \in \mathbf{R}_+ \mid (a_n \rho^n)_{n \geq 0} \text{ est bornée}\}$$

De ce fait, $R_b \leq R_a$.

2. Si $a_n \sim b_n$ alors $a_n = O(b_n)$ et $b_n = O(a_n)$. Il suffit alors d'utiliser ce qui précède.

□

ATTENTION

Le fait que $\sum_{n \geq 0} a_n z^n$ et $\sum_{n \geq 0} b_n z^n$ ait le même rayon de convergence quand $a_n \sim b_n$ ne signifie par que $\sum_{n \geq 0} a_n z^n$ converge si et seulement si $\sum_{n \geq 0} b_n z^n$ converge. On peut prendre par exemple

$$a_n = \frac{(-1)^n}{\sqrt{n+1}} \text{ et } b_n = \frac{(-1)^n}{\sqrt{n+1}} + \frac{1}{(n+1)}$$

On a bien $a_n \sim b_n$ et donc $R_a = R_b = 1$, mais $\sum_{n \geq 0} a_n = \sum_{n \geq 0} a_n z^n$ converge alors que $\sum_{n \geq 0} b_n = \sum_{n \geq 0} b_n z^n$ ne converge pas.

Proposition 10.10

Soit $\sum_{n \geq 0} a_n z^n$ une série entière. La série $\sum_{n \geq 0} n a_n z^n$ a le même rayon de convergence que $\sum_{n \geq 0} a_n z^n$.

Démonstration : Notons $b_n = n a_n$ et R_a le rayon de convergence de $\sum_{n \geq 0} a_n z^n$ et R_b celui de $\sum_{n \geq 0} n a_n z^n$.

- On remarque d'abord que $a_n = O(n a_n)$ et donc $R_b \leq R_a$.
- C'est un peu plus compliqué pour montrer que $R_a \leq R_b$. Soit $\rho < R_a$. Il existe ρ' tel que $\rho < \rho' < R_a$. Alors pour tout entier n ,

$$|n a_n| \rho^n = n \left(\frac{\rho}{\rho'} \right)^n |a_n| \rho'^n.$$

Or $\frac{\rho}{\rho'} < 1$ donc la suite $\left(n \left(\frac{\rho}{\rho'}\right)^n\right)$ tend vers 0 et donc, à partir d'un certain rang,

$$|na_n|\rho^n \leq |a_n|\rho'^n.$$

Maintenant, on sait que la suite $(a_n\rho'^n)_{n \geq 0}$ est bornée donc la suite $(|na_n|\rho^n)_{n \geq 0}$ aussi ce qui montre que $\rho \leq R_b$. On a donc montré que $\rho < R_a \Rightarrow \rho \leq R_b$ et de ce fait $R_a \leq R_b$.

On a bien $R_a = R_b$. □

Remarque : Dans le cas où on suppose que la suite $\left(\frac{a_{n+1}}{a_n}\right)_{n \geq 0}$ existe et a une limite, on peut faire une preuve beaucoup plus simple en utilisant que

$$\lim_{n \rightarrow +\infty} \frac{(n+1)a_{n+1}}{na_n} = \lim_{n \rightarrow +\infty} \frac{a_{n+1}}{a_n}$$

1.5 Operations

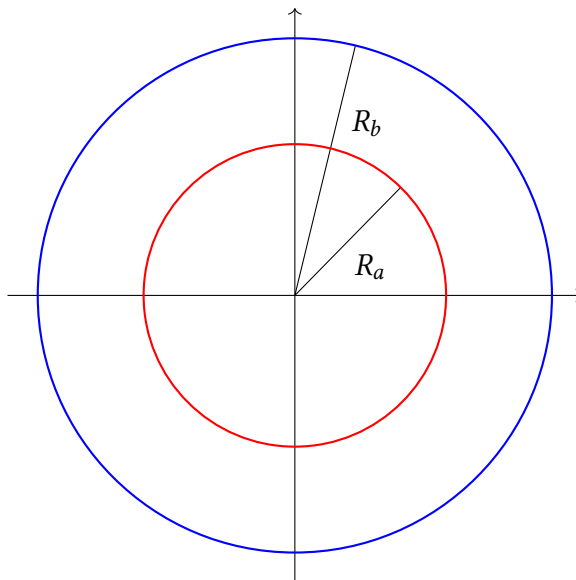
Proposition 10.11 (Somme)

Soit $\sum_{n \geq 0} a_n z^n$ et $\sum_{n \geq 0} b_n z^n$ deux séries entières de rayon de convergence R_a et R_b .

- Si $R_a \neq R_b$ la série entière $\sum_{n \geq 0} (a_n + b_n) z^n$ a pour rayon de convergence $R = \min(R_a, R_b)$.
- Si $R_a = R_b$ la série entière $\sum_{n \geq 0} (a_n + b_n) z^n$ a un rayon de convergence R supérieur ou égal à $R_a = R_b$.

Démonstration :

- On suppose par symétrie que $R_a < R_b$. Soit $\rho \in \mathbf{R}_+$. Si $\rho < R_a$ alors les séries $\sum_{n \geq 0} a_n \rho^n$ et $\sum_{n \geq 0} b_n \rho^n$ convergent donc la série $\sum_{n \geq 0} (a_n + b_n) \rho^n$ converge. On en déduit que $\rho < R$ et donc $R \geq R_a$. Par contre si $\rho \in]R_a, R_b[$, la série $\sum_{n \geq 0} a_n \rho^n$ diverge et $\sum_{n \geq 0} b_n \rho^n$ converge donc la série $\sum_{n \geq 0} (a_n + b_n) \rho^n$ diverge. On en déduit que $\rho \geq R$. De ce fait $R \leq R_a$ et donc $R = R_a = \min(R_a, R_b)$.



- On peut reprendre le début de la preuve précédente pour montrer que $R \geq R_a$. Par contre, la deuxième partie ne s'étend pas à ce cas car la somme de deux séries divergentes peut être une série convergente.

□

Proposition 10.12 (Produit par une constante)

Soit $\sum_{n \geq 0} a_n z^n$ une série entière et λ un scalaire non nul. La série $\sum_{n \geq 0} \lambda a_n z^n$ a le même rayon de convergence.

Voyons maintenant la notion de produit. Le principe est de généraliser les produits de polynômes.

Définition 10.13 (Produit de Cauchy)

Soit $\sum_{n \geq 0} a_n z^n$ et $\sum_{n \geq 0} b_n z^n$ deux séries entières. On pose pour tout n

$$c_n = \sum_{i+j=n} a_i b_j = \sum_{i=0}^n a_i b_{n-i} = \sum_{i=0}^n a_{n-i} b_i.$$

La série entière $\sum_{n \geq 0} c_n z^n$ s'appelle le produit de Cauchy des deux séries entières.

Théorème 10.14

Soit $\sum_{n \geq 0} a_n z^n$ et $\sum_{n \geq 0} b_n z^n$ deux séries entières de rayon de convergence R_a et R_b . Le rayon de convergence du produit de Cauchy $\sum_{n \geq 0} c_n z^n$ est au moins égal au minimum de R_a et R_b . De plus,

$$\forall z \in \mathbb{C}, |z| < \min(R_a, R_b) \Rightarrow \sum_{n=0}^{+\infty} c_n z^n = \left(\sum_{n=0}^{+\infty} a_n z^n \right) \times \left(\sum_{n=0}^{+\infty} b_n z^n \right)$$

Démonstration : Il suffit d'utiliser le résultat sur le produit de Cauchy des séries numériques.

Soit z tel que $|z| < \min(R_a, R_b)$. On pose $u_n = a_n z^n$ et $v_n = b_n z^n$. On sait alors que les séries numériques $\sum_{n \geq 0} u_n$ et $\sum_{n \geq 0} v_n$ sont absolument convergentes. Posons alors pour tout entier naturel n ,

$$w_n = \sum_{p+q=n} u_p v_q = \sum_{p+q=n} a_p z^p b_q z^q = \left(\sum_{p+q=n} a_p b_q \right) z^n = c_n z^n$$

D'après le résultat vu sur le produit de Cauchy des séries numériques, la série $\sum_{n \geq 0} w_n$ est absolument convergente. De plus

$$\sum_{n=0}^{+\infty} w_n = \sum_{n=0}^{+\infty} c_n z^n = \left(\sum_{n=0}^{+\infty} a_n z^n \right) \times \left(\sum_{n=0}^{+\infty} b_n z^n \right)$$

Cela montre bien que la série entière $\sum_{n \geq 0} c_n z^n$ a un rayon de convergence au moins égal à $\min(R_a, R_b)$

et que, pour $z \in \mathbb{C}$ tel que $|z| \leq \min(R_a, R_b)$,

$$\sum_{n=0}^{+\infty} c_n z^n = \left(\sum_{n=0}^{+\infty} a_n z^n \right) \times \left(\sum_{n=0}^{+\infty} b_n z^n \right)$$

□

Application : On sait que la série entière $\sum (n+1)z^n$ a un rayon de convergence égal à 1. On peut le voir comme le produit de Cauchy des séries $\sum a_n z^n$ et $\sum b_n z^n$ où $a_n = b_n = 1$. On en tire que

$$\forall z \in \mathbb{C}, |z| < 1 \Rightarrow \sum_{n=0}^{+\infty} (n+1)z^n = \frac{1}{1-z} \times \frac{1}{1-z} = \frac{1}{(1-z)^2}.$$

2 Étude de la somme d'une série entière

Dans ce paragraphe, on va étudier la fonction $x \mapsto \sum_{n=0}^{+\infty} a_n x^n$ définie sur $] -R_a, R_a[$ ou la fonction $z \mapsto \sum_{n=0}^{+\infty} a_n z^n$ définie sur $D(0, R_a)$ où R_a est le rayon de convergence.

2.1 Convergence normale et continuité

Proposition 10.15

Soit $\sum_{n \geq 0} a_n z^n$ une série entière. Soit R_a son rayon de convergence. La série de fonctions converge normalement (donc uniformément) sur tout disque fermé de rayon **strictement** inférieur à R_a .

Démonstration : Soit $\alpha < R_a$. Sur le disque fermé $\overline{D}(0, \alpha)$ la norme infinie de $z \mapsto a_n z^n$ est $\|z \mapsto a_n z^n\|_{\infty, D} = |a_n| \alpha^n$. Or $\sum_{n \geq 0} a_n \alpha^n$ converge donc il y a bien convergence normale sur $\overline{D}(0, \alpha)$.

De ce fait il y a convergence uniforme sur $\overline{D}(0, \alpha)$

□

ATTENTION

Cela ne signifie pas qu'il y a convergence uniforme sur le disque ouvert de centre 0 et de rayon R_a (et donc encore moins sur le disque fermé).

Exemple : On considère par exemple la série entière $\sum_{n \geq 0} z^n$. On sait que le rayon de convergence vaut 1.

- Pour tout $\alpha < 1$, la série entière $\sum_{n \geq 0} z^n$ converge normalement sur $\overline{D}(0, \alpha)$ ou sur $[-\alpha, \alpha]$ si on travaille avec une variable réelle.
- Travaillons maintenant sur $] -1, 1[$. Pour tout entier $n \geq 0$, $\|z \mapsto z^n\|_{\infty,]-1, 1[} = 1$. Comme la série $\sum_{n \geq 0} 1$ diverge, il n'y a pas convergence normale de la série entière sur $] -1, 1[$.

- On peut même montrer qu'il n'y a pas convergence uniforme sur $] -1, 1[$. En effet, si c'était le cas, on pourrait appliquer le théorème de la double limite ce qui n'est pas possible car la série $\sum_{n \geq 1} (\lim_{x \rightarrow 1^-} x^n)$ diverge.

Corollaire 10.16 (Continuité)

Soit $\sum_{n \geq 0} a_n z^n$ une série entière. La fonction $z \mapsto \sum_{n=0}^{+\infty} a_n z^n$ définie sur le disque ouvert $D(0, R_a)$ est continue.

Démonstration : On remarque d'abord que pour tout entier n , $z \mapsto a_n z^n$ est continue sur $\mathbb{C}$. Il suffit ensuite de voir que pour tout $z \in D(0, R_a)$, on a $|z| < R_a$. Il existe donc α tel que $|z| < \alpha < R_a$. On utilise alors que la série de fonctions converge uniformément car normalement sur $D(0, \alpha)$. $\square$

Remarque : On utilisera la majorité du temps ce résultat avec des fonctions de la variable réelle mais ce résultat est aussi vrai dans le cas des fonctions de la variable complexe. On généralise en effet la définition de la continuité de manière évidente en remplaçant la valeur absolue par le module. Plus précisément, une fonction f définie sur un disque ouvert $D(0, R)$ est continue si

$$\forall a \in D(0, R), \forall \varepsilon > 0, \exists \eta > 0, \forall x \in D(0, R), |x - a| \leq \eta \Rightarrow |f(x) - f(a)| \leq \varepsilon$$

La démonstration des théorèmes de continuité des limites des suites de fonctions et des sommes de séries de fonctions se généralisent sans problème à ce cadre.

Théorème 10.17 (Théorème d'Abel radial)

Soit $\sum_{n \geq 0} a_n x^n$ une série entière de rayon de convergence $R \in \mathbb{R}_+^*$. On suppose de plus que la série numérique $\sum_{n \geq 0} a_n R^n$ converge. On a

$$\lim_{x \rightarrow R^-} \sum_{n=0}^{+\infty} a_n x^n = \sum_{n=0}^{+\infty} a_n R^n$$

Démonstration : Ce théorème est admis; voir DL $\square$

2.2 Primitivation et dérivation

Dans ce paragraphe on travaille avec des séries entières de la variable réelle.

Proposition 10.18 (Primitivation)

Soit $\sum_{n \geq 0} a_n x^n$ une série entière et R_a son rayon de convergence.

Soit $f : x \mapsto \sum_{n=0}^{+\infty} a_n x^n$ la fonction définie sur $] -R_a, R_a[$.

Elle se primitive sur l'intervalle ouvert terme à terme. C'est-à-dire que la fonction $F : x \mapsto \sum_{n=0}^{+\infty} \frac{a_n}{n+1} x^{n+1}$ est la primitive de f sur $] -R_a, R_a[$ qui s'annule en 0.

Démonstration : Les fonctions $x \mapsto a_n x^n$ sont continues sur $] -R_a, R_a[$. De plus, la série converge uniformément sur tout segment J inclus dans $] -R_a, R_a[$ car la série entière converge normalement sur tout $[-\alpha, \alpha]$ pour $\alpha < R_a$. On peut appliquer le théorème d'intégration terme à terme.

La fonction $F : x \mapsto \sum_{n=0}^{+\infty} \frac{a_n}{n+1} x^{n+1}$ est donc la primitive qui s'annule en 0 de f sur $] -R_a, R_a[$. $\square$

Remarque : Le rayon de convergence de la primitive $\sum_{n \geq 0} \frac{a_n}{n+1} x^{n+1}$ est le même que celui de la série

$$\sum_{n \geq 0} (n+1) \frac{a_n}{n+1} x^{n+1} = \sum_{n \geq 0} a_n x^{n+1}. \text{ Ce dernier étant le même que celui de la série entière } \sum_{n \geq 0} a_n x^n.$$

Exemple : On sait que $\sum_{n \geq 0} x^n$ a pour rayon de convergence 1. De plus pour $x \in] -1, 1[$,

$$f(x) = \sum_{n=0}^{+\infty} x^n = \frac{1}{1-x}$$

On peut « primitiver ». On obtient que pour tout $x \in] -1, 1[$,

$$F(x) = -\ln(1-x) = \sum_{n=0}^{+\infty} \frac{1}{n+1} x^{n+1}$$

On en déduit que pour $x \in] -1, 1[$

$$\ln(1+x) = -\sum_{n=0}^{+\infty} \frac{1}{n+1} (-x)^{n+1} = \sum_{n=0}^{+\infty} \frac{(-1)^n}{n+1} x^{n+1}.$$

L'étude du comportement au bord de l'intervalle de convergence peut s'obtenir à l'aide du théorème d'Abel radial.

Comme la série $\sum_{n \geq 0} \frac{(-1)^n}{n+1} x^{n+1}$ converge, on peut appliquer le théorème à la série entière $\sum_{n \geq 1} \frac{(-1)^{n-1}}{n} x^n$ et on obtient

$$\sum_{n=1}^{+\infty} \frac{(-1)^{n-1}}{n} = \lim_{x \rightarrow 1^-} \ln(1+x) = \ln(2)$$

Exercice : Déterminer un développement sur $] -1, 1[$ de $\arctan(x)$. En déduire une formule pour $\frac{\pi}{4}$.

Proposition 10.19 (Dérivation)

Soit $\sum_{n \geq 0} a_n x^n$ une série entière et R_a son rayon de convergence.

Soit $f : x \mapsto \sum_{n=0}^{+\infty} a_n x^n$ la fonction définie sur $] -R_a, R_a[$.

Elle est de classe $\mathcal{C}^\infty$ sur cet intervalle et ses dérivées successives se calculent terme à terme.

$$\forall k \in \mathbb{N}, \forall x \in] -R_a, R_a[, f^{(k)}(x) = \sum_{n=k}^{+\infty} a_n \frac{n!}{(n-k)!} x^{n-k}.$$

Démonstration : On considère la série entière $\sum_{n \geq 1} na_n x^{n-1}$. Cette série a le même rayon de convergence que $f(x) = \sum_{n \geq 0} a_n x^n$. On peut appliquer le théorème précédent. On en déduit que la fonction $x \mapsto f(x) - f(0)$ est la primitive qui s'annule en 0 de $h : x \mapsto \sum_{n \geq 1} na_n x^{n-1}$.

De ce fait, f est dérivable et $\forall x \in]-R_a, R_a[$, $f'(x) = \sum_{n=1}^{+\infty} na_n x^{n-1}$.

Il suffit de reprendre cet argument avec la série entière $\sum_{n \geq 1} na_n x^{n-1}$. □

Corollaire 10.20

Soit $\sum_{n \geq 0} a_n x^n$ une série entière dont le rayon de convergence R_a n'est pas nul. On pose

$$f : x \mapsto \sum_{n=0}^{+\infty} a_n x^n.$$

Pour tout entier n , $a_n = \frac{f^{(n)}(0)}{n!}$

Remarque : La fonction f est égale à sa série de Taylor :

$$f(x) = \sum_{n=0}^{+\infty} \frac{f^{(n)}(0)}{n!} x^n.$$

Corollaire 10.21 (Unicité du développement en série entière)

Soit $\sum_{n \geq 0} a_n x^n$ et $\sum_{n \geq 0} b_n x^n$ deux séries entières. Si elles ont un rayon de convergence strictement supérieur à 0 et qu'elles coïncident sur $[-\delta, \delta]$ où $\delta > 0$ alors pour tout entier n non nul, $a_n = b_n$.

Démonstration : Considérons la fonction f définie sur $[-\delta, \delta]$ par

$$f(x) = \sum_{n=0}^{+\infty} a_n x^n = \sum_{n=0}^{+\infty} b_n x^n$$

D'après le corollaire précédent,

$$\forall n \in \mathbb{N}, a_n = \frac{f^{(n)}(0)}{n!} = b_n.$$

□

Exemple : Étude de la suite de Fibonacci.

On considère la suite (a_n) définie par $a_0 = a_1 = 1$ et pour $n \geq 0$, $a_{n+2} = a_{n+1} + a_n$. C'est la suite de Fibonacci. Il a été vu en première année que

$$\forall n \in \mathbb{N}, a_n = \frac{5 + \sqrt{5}}{10} \varphi^n + \frac{5 - \sqrt{5}}{10} \check{\varphi}^n$$

où $\varphi = \frac{1 + \sqrt{5}}{2}$ et $\check{\varphi} = \frac{1 - \sqrt{5}}{2}$.

Nous allons voir une nouvelle méthode pour démontrer ce résultat. On considère la série entière

$$\sum_{n \geq 0} a_n x^n.$$

On suppose que cette série entière à un rayon de convergence strictement positif R et on pose

$$\forall x \in]-R, R[, f(x) = \sum_{n=0}^{+\infty} a_n x^n.$$

En utilisant la relation de récurrence on a alors pour tout $x \in]-R, R[$,

$$\begin{aligned} f(x) &= 1 + x + \sum_{n=2}^{+\infty} a_n x^n \\ &= 1 + x + \sum_{n=2}^{+\infty} (a_{n-1} + a_{n-2}) x^n \\ &= 1 + x + x \sum_{n=1}^{+\infty} a_n x^n + x^2 \sum_{n=0}^{+\infty} a_n x^n \\ &= 1 + x + x(f(x) - 1) + x^2 f(x) \end{aligned}$$

On en déduit donc que $(1 - x - x^2)f(x) = 1$ et donc $f(x) = \frac{1}{1-x-x^2}$.

Considérons donc la fonction $g : x \mapsto \frac{1}{1-x-x^2}$. C'est une fraction rationnelle, on peut établir sa décomposition en éléments simples. Pour cela on factorise le dénominateur :

$$1 - X - X^2 = -(X - \alpha)(X - \check{\alpha})$$

où $\alpha = \frac{-1 - \sqrt{5}}{2} = -\varphi$ et $\check{\alpha} = \frac{-1 + \sqrt{5}}{2} = -\check{\varphi}$. On obtient alors

$$g : x \mapsto \frac{1}{\sqrt{5}} \frac{1}{x - \alpha} - \frac{1}{\sqrt{5}} \frac{1}{x - \check{\alpha}} = -\frac{1}{\alpha\sqrt{5}} \frac{1}{1 - (x/\alpha)} + \frac{1}{\check{\alpha}\sqrt{5}} \frac{1}{1 - (x/\check{\alpha})}$$

Maintenant on sait que $x \mapsto \frac{1}{1-x}$ est la somme de la série entière $\sum_{n \geq 0} x^n$ de rayon de convergence 1.

On en déduit que $x \mapsto \frac{1}{1-(x/\alpha)}$ est la somme de la série entière $\sum_{n \geq 0} \left(\frac{x}{\alpha}\right)^n$ de rayon de convergence $|\alpha|$

et $x \mapsto \frac{1}{1-(x/\check{\alpha})}$ est la somme de la série entière $\sum_{n \geq 0} \left(\frac{x}{\check{\alpha}}\right)^n$ de rayon de convergence $\check{\alpha}$.

On en déduit que g est la somme d'une série entière $\sum_{n \geq 0} b_n x^n$ de rayon de convergence $\min(|\alpha|, \check{\alpha}) = \check{\alpha}$.

En remontant le calcul ci-dessus, on obtient que pour $x \in]-\check{\alpha}, \check{\alpha}[$,

$$b_0 + (b_1 - b_0) + \sum_{n=2}^{+\infty} (b_n - b_{n-1} - b_{n-2}) x^n = 1$$

Par unicité du la série entière on a bien que

$$b_0 = 1; b_1 = b_0 \text{ et } \forall n \geq 2, b_n = b_{n-1} + b_{n-2}.$$

On a donc $b_n = a_n$ pour tout entier n . La série entière initiale avait bien un rayon de convergence non nul (il vaut $\check{\alpha}$).

On trouve de plus que pour tout entier n ,

$$a_n = b_n = -\frac{1}{\alpha\sqrt{5}} \frac{1}{\alpha^n} + \frac{1}{\check{\alpha}\sqrt{5}} \frac{1}{\check{\alpha}^n}$$

On peut alors vérifier que c'est la formule « usuelle » car $\frac{1}{\alpha} = \check{\varphi}$ et $\frac{1}{\check{\alpha}} = \varphi$.

3 Fonctions développables en séries entières

3.1 Définitions

Définition 10.22

Une fonction f est dite développable en série entière sur l'intervalle ouvert $] -r, r[$ (resp. sur le disque ouvert $D(0, r)$) s'il existe une série entière dont le rayon de convergence est supérieur ou égal à r dont f est la somme. C'est-à-dire, s'il existe $(a_n) \in \mathbb{C}^{\mathbb{N}}$ telle que

$$\forall x \in] -r, r[, f(x) = \sum_{n=0}^{+\infty} a_n x^n \quad \left(\text{resp. } \forall z \in D(0, r), f(z) = \sum_{n=0}^{+\infty} a_n z^n \right)$$

Exemples :

1. On a vu que $z \mapsto \exp(z)$ est développable en série entière sur $\mathbb{C}$ et

$$\forall z \in \mathbb{C}, \exp(z) = \sum_{n=0}^{+\infty} \frac{z^n}{n!}$$

2. On a vu que $z \mapsto \frac{1}{1-z}$ est développable en série entière sur $D(0, 1)$ et

$$\forall z \in D(0, 1), \frac{1}{1-z} = \sum_{n=0}^{+\infty} z^n$$

Proposition 10.23

Soit f une fonction développable en série entière sur $] -r, r[$. Elle est de classe $\mathcal{C}^\infty$ et est égale à son développement de Taylor en 0 :

$$\forall x \in] -r, r[, f(x) = \sum_{n=0}^{+\infty} \frac{f^{(n)}(0)}{n!} x^n$$

Démonstration : Déjà vu. □

ATTENTION

Ce n'est pas une équivalence ! Il ne suffit pas d'être de classe $\mathcal{C}^\infty$ pour être développable en série entière. Il peut se produire deux problèmes :

- La série de Taylor peut avoir un rayon de convergence nulle.
- La série de Taylor peut avoir un rayon de convergence strictement positif mais la somme de la série de Taylor diffère de la fonction de départ - voir ci-dessous

Exemple : On considère la fonction $f : x \mapsto e^{-1/x^2}$ définie sur $\mathbb{R}^*$. Elle se prolonge par continuité en 0 en posant $f(0) = 0$. Notons encore f ce prolongement. Il est clair que la fonction f est de classe $\mathcal{C}^\infty$ sur $\mathbb{R}^*$.

Montrons qu'elle est même de classe $\mathcal{C}^\infty$ sur $\mathbb{R}$ en entier.

Pour $x \neq 0$ on vérifie que :

$$f'(x) = \frac{2}{x^3} e^{-1/x^2} \text{ et } f'' = \frac{-6x^2 + 4}{x^6} e^{-1/x^2}$$

Montrons alors par récurrence que pour tout entier $n \in \mathbb{N}$, il existe un polynôme P_n tel que $f^{(n)}$ est définie sur $\mathbb{R}^*$ par $f^{(n)} : x \mapsto \frac{P_n(x)}{x^{3n}} e^{-1/x^2}$.

- **I** : pour $n = 0$, il suffit de prendre $P_0 : x \mapsto 1$.
- **H** : soit $n \in \mathbb{N}$, on suppose qu'il existe un polynôme P_n tel que pour tout $x \in \mathbb{R}^*$, $f^{(n)}(x) = \frac{P_n(x)}{x^{3n}} e^{-1/x^2}$. En dérivant on obtient que pour tout $x \in \mathbb{R}^*$,

$$f^{(n+1)}(x) = \frac{x^{3n} P_n'(x) - 3n x^{3n-1} P_n(x)}{x^{6n}} e^{-1/x^2} + \frac{P_n(x)}{x^{3n}} \times \frac{2}{x^3} e^{-1/x^2} = \frac{P_{n+1}(x)}{x^{3n+3}} e^{-1/x^2}$$

où

$$P_{n+1} = X^3 P_n' + (2 - 3nX^2) P_n$$

En particulier, pour tout entier naturel n , $f^{(n)}$ a une limite nulle en 0 par croissance comparée. Cela montre que f est de classe $\mathcal{C}^n$ sur $\mathbb{R}$

La fonction f est donc de classe $\mathcal{C}^\infty$ sur $\mathbb{R}$, sa série de Taylor en 0 converge mais la fonction f ne coïncide pas avec la série de Taylor.

Proposition 10.24

Soit f et g deux fonctions développables en série entière sur $] -r, r[$ et $\lambda \in \mathbb{C}$.

- La fonction $f + g$ est développable en série entière sur $] -r, r[$.
- La fonction λf est développable en série entière sur $] -r, r[$.
- La fonction $f \times g$ est développable en série entière sur $] -r, r[$.
- Les dérivées successives de f sont développables en série entière sur $] -r, r[$.
- Les primitives de f sont développables en série entière sur $] -r, r[$.

Exemple : La fonction $x \mapsto \frac{e^x - 1}{x}$ est développable en série entière sur $\mathbb{R}$. Elle est donc $\mathcal{C}^\infty$.

3.2 Développements en série entière des fonctions usuelles

Théorème 10.25

1. (a) Pour tout $z \in \mathbb{C}$, $\exp(z) = \sum_{n=0}^{+\infty} \frac{z^n}{n!}$
(b) Pour tout $x \in \mathbb{R}$, $\sin(x) = \sum_{n=0}^{+\infty} \frac{(-1)^n}{(2n+1)!} x^{2n+1}$.
(c) Pour tout $x \in \mathbb{R}$, $\cos(x) = \sum_{n=0}^{+\infty} \frac{(-1)^n}{2n!} x^{2n}$.
(d) Pour tout $x \in \mathbb{R}$, $\operatorname{sh}(x) = \sum_{n=0}^{+\infty} \frac{1}{(2n+1)!} x^{2n+1}$.
(e) Pour tout $x \in \mathbb{R}$, $\operatorname{ch}(x) = \sum_{n=0}^{+\infty} \frac{1}{2n!} x^{2n}$.
2. (a) Pour tout $z \in \mathbb{C}$ avec $|z| < 1$, $\frac{1}{1-z} = \sum_{n=0}^{+\infty} z^n$.
(b) Pour tout $x \in]-1, 1[$, $\ln(1+x) = \sum_{n=0}^{+\infty} \frac{(-1)^n}{n+1} x^{n+1}$.
(c) Pour tout $x \in]-1, 1[$, $\arctan(x) = \sum_{n=0}^{+\infty} \frac{(-1)^n}{2n+1} x^{2n+1}$.

Théorème 10.26

Soit $\alpha \in \mathbb{R}$, la fonction $x \mapsto (1+x)^\alpha$ est développable en série entière et

$$\forall x \in]-1, 1[, (1+x)^\alpha = \sum_{n=0}^{+\infty} \frac{\alpha(\alpha-1) \cdots (\alpha-n+1)}{n!} x^n.$$

Démonstration : On utilise la méthode de l'équation différentielle – Voir plus loin.

4 Utilisation des séries entières dans l'étude des équations différentielles

4.1 Rappels du cours de première année

Définition 10.27 (Équations différentielles linéaires)

Soit n un entier non nul et I un intervalle.

1. Une équation différentielle linéaire d'ordre n est une équation de la forme :

$$\alpha_n(t)y^{(n)} + \cdots + \alpha_1(t)y' + \alpha_0(t)y = \beta(t), \quad (E)$$

où les $\alpha_0, \alpha_1, \dots, \alpha_n$ et β sont des fonctions **continues** sur I .

- Les fonctions $\alpha_0, \dots, \alpha_n$ s'appellent les coefficients de l'équation différentielle.
- La fonction β s'appelle le second membre.
- L'équation différentielle est dite homogène si β est la fonction nulle.

2. La forme résolue d'une équation différentielle linéaire d'ordre n est

$$y^{(n)} = a_{n-1}(t)y^{(n-1)} + \cdots + a_1(t)y' + a_0(t)y + b(t).$$

Remarques :

1. On associe à une équation

$$\alpha_n(t)y^{(n)} + \cdots + \alpha_1(t)y' + \alpha_0(t)y = \beta(t), \quad (E)$$

son *équation homogène associée*. C'est l'équation différentielle qui a les mêmes coefficients mais dont le second membre est nul

$$\alpha_n(t)y^{(n)} + \cdots + \alpha_1(t)y' + \alpha_0(t)y = 0 \quad (H)$$

2. On considère une équation différentielle linéaire

$$\alpha_n(t)y^{(n)} + \cdots + \alpha_1(t)y' + \alpha_0(t)y = \beta(t) \quad (E)$$

Si la fonction α_n ne s'annule pas sur l'intervalle I on peut lui associer une équation différentielle sous forme résolue

$$y^{(n)} = a_{n-1}(t)y^{(n-1)} + \cdots + a_1(t)y' + a_0(t)y + b(t)$$

en posant pour tout $i \in \llbracket 0 ; n-1 \rrbracket$, $a_i : t \mapsto -\frac{\alpha_i(t)}{\alpha_n(t)}$ et $b : t \mapsto \frac{\beta(t)}{\alpha_n(t)}$.

Définition 10.28 (Solutions d'une équation différentielle linéaire)

Soit n un entier non nul et I un intervalle. On considère l'équation différentielle linéaire

$$\alpha_n(t)y^{(n)} + \cdots + \alpha_1(t)y' + \alpha_0(t)y = \beta(t), \quad (E)$$

où les $\alpha_0, \alpha_1, \dots, \alpha_n$ et β sont des fonctions **continues** sur I .

On appelle *solution* de l'équation (E) toute fonction $y \in \mathcal{C}^n(I, \mathbf{K})$ telle que

$$\forall t \in I, \quad \alpha_n(t)y^{(n)} + \cdots + \alpha_1(t)y' + \alpha_0(t)y = \beta(t)$$

Remarque : On considère une équation différentielle linéaire

$$\alpha_n(t)y^{(n)} + \cdots + \alpha_1(t)y' + \alpha_0(t)y = \beta(t), \quad (E)$$

On peut regarder la fonction

$$\begin{aligned} \Phi : \mathcal{C}^n(I, \mathbf{K}) &\rightarrow \mathcal{C}^0(I, \mathbf{K}) \\ y &\mapsto \sum_{k=0}^n \alpha_k y^{(k)} \end{aligned}$$

Avec ces notations, on voit que $y \in \mathcal{C}^n(I, \mathbf{K})$ vérifie (E) si et seulement si $\Phi(y) = b$.

Théorème 10.29 (Structure de l'ensemble des solutions)

Soit (E) une équation différentielle linéaire d'ordre n définie sur un intervalle I et (H) l'équation homogène associée.

1. L'ensemble $\mathcal{S}_H$ des solutions de (H) est un sous-espace vectoriel de $\mathcal{C}^n(I, \mathbf{K})$.
2. S'il existe une solution f_0 de (E) on a $\mathcal{S}_E = \{f + f_0 \mid f \in \mathcal{S}_H\}$.

On dit alors que $\mathcal{S}_E$ est un espace affine de direction $\mathcal{S}_H$. C'est-à-dire

$$f \in \mathcal{S}_E \iff (f - f_0) \in \mathcal{S}_H.$$

4.2 Rappels - Équations linéaires du premier ordre

Théorème 10.30 (Equation différentielle linéaire d'ordre 1 homogène)

Soit I un intervalle et $a : I \rightarrow \mathbf{C}$ une fonction continue sur I . Les solutions de l'équation différentielle linéaire du premier ordre résolue

$$y'(t) + a(t)y(t) = 0$$

sont les fonctions f définies sur I par :

$$\forall x \in I, \quad f(t) = \lambda \cdot e^{-A(t)},$$

où $\lambda \in \mathbf{C}$ et $x \mapsto A(t)$ est une primitive de a sur I .

Note (Plan de résolution d'une équation différentielle linéaire d'ordre 1)

Voici alors la méthode à appliquer pour résoudre l'équation (E).

$$y'(t) + a(t)y(t) = b(t) \quad (E)$$

1. On exprime l'équation homogène associée (H).
2. On résout (H).
3. On trouve une solution particulière de (E)
4. On exprime la forme générale des solutions.

Proposition 10.31 (Variation de la constante)

Les solutions de l'équation différentielle linéaire résolue du premier ordre sur I

$$y'(t) + a(t)y = b(t)$$

où a et b sont des fonctions continues sur I sont de la forme

$$t \mapsto \left(\int^t b(u)e^{A(u)} du \right) e^{-A(t)} + Ke^{-A(t)}.$$

(On peut supprimer la constante si on fait varier la borne « en bas » de l'intégrale)

Exemple : On veut résoudre

$$y'(t) + y(t) = \cos(t) \quad (E)$$

L'équation homogène associée est

$$y'(t) + y(t) = 0 \quad (H)$$

Les solutions de l'équation homogène sont

$$\mathcal{S}_H = \{t \mapsto \lambda e^{-t} \mid \lambda \in \mathbf{R}\}.$$

On va chercher la solution particulière sous la forme :

$$f : t \mapsto \lambda(t)e^{-t} \text{ où } \lambda : \mathbf{R} \rightarrow \mathbf{R},$$

est une fonction dérivable.

Dés lors on a

$$\begin{aligned} f \text{ vérifie (E)} &\Leftrightarrow \forall t \in \mathbf{R}, f'(t) + f(t) = \cos t \\ &\Leftrightarrow \forall t \in \mathbf{R}, \lambda'(t)e^{-t} - \lambda(t)e^{-t} + \lambda(t)e^{-t} = \cos t \\ &\Leftrightarrow \forall t \in \mathbf{R}, \lambda'(t)e^{-t} = \cos t \\ &\Leftrightarrow \forall t \in \mathbf{R}, \lambda'(t) = e^t \cos t \end{aligned}$$

Il ne reste plus qu'à trouver une primitive de $e^t \cos t$. Pour cela on remarque que

$$\forall x \in \mathbf{R}, e^x \cos x = \Re \left(e^{(1+i)x} \right).$$

On en déduit qu'une primitive de $e^t \cos t$ est donnée par

$$\Re \left(\frac{e^{(1+i)t}}{1+i} \right) = \Re \left(\frac{(1-i)e^t(\cos t + i \sin t)}{2} \right) = \frac{1}{2}e^t(\cos t + \sin t).$$

On en déduit que la fonction $t \mapsto \frac{1}{2}(\cos t + \sin t)$ est une solution particulière et que

$$\mathcal{S}_E = \left\{ t \mapsto \frac{1}{2}(\cos t + \sin t) + \lambda e^{-t} \mid \lambda \in \mathbf{R} \right\}.$$

Exercice :

Résoudre l'équation différentielle : $y'(t) - y(t) = e^t \ln t$

4.3 Rappels - Équations linéaires du deuxième ordre

On va résoudre l'équation différentielle

$$ay''(t) + by'(t) + cy(t) = f(t) \quad (E)$$

où a, b et c sont des **constantes** complexes avec $a \neq 0$ et f une fonction définie sur un intervalle I .

L'équation homogène associée est alors

$$ay''(t) + by'(t) + cy(t) = 0. \quad (H)$$

Définition 10.32

On appelle *équation caractéristique* de (H) l'équation du second degré :

$$aX^2 + bX + c = 0$$

Théorème 10.33 (Cas Complexe)

Avec les notations précédentes, on note $\Delta = b^2 - 4ac$ le discriminant de l'équation caractéristique. On a alors

- Si $\Delta \neq 0$: on note α_1 et α_2 les deux racines de l'équation caractéristique. On a

$$\mathcal{S}_H = \{t \mapsto \lambda e^{\alpha_1 t} + \mu e^{\alpha_2 t} \mid (\lambda, \mu) \in \mathbb{C}^2\}.$$

- Si $\Delta = 0$: on note α la racine double de l'équation caractéristique. On a

$$\mathcal{S}_H = \{t \mapsto (\lambda t + \mu) e^{\alpha t} \mid (\lambda, \mu) \in \mathbb{C}^2\}.$$

En particulier, c'est un $\mathbb{C}$ -espace vectoriel de dimension 2.

Théorème 10.34 (Cas réel)

Avec les notations précédentes, suppose que a, b et c sont des réels. On note $\Delta = b^2 - 4ac$ le discriminant de l'équation caractéristique. On a alors

- Si $\Delta > 0$: on note α_1 et α_2 les deux racines réelles de l'équation caractéristique. On a

$$\mathcal{S}_H = \{t \mapsto \lambda e^{\alpha_1 t} + \mu e^{\alpha_2 t} \mid (\lambda, \mu) \in \mathbb{R}^2\}.$$

- Si $\Delta = 0$: on note α la racine double réelle de l'équation caractéristique. On a

$$\mathcal{S}_H = \{t \mapsto (\lambda t + \mu) e^{\alpha t} \mid (\lambda, \mu) \in \mathbb{R}^2\}.$$

- Si $\Delta < 0$: on note $\alpha + i\beta$ et $\alpha - i\beta$ les deux racines complexes conjuguées de l'équation caractéristique. On a

$$\mathcal{S}_H = \{t \mapsto e^{\alpha t} (\lambda \cos(\beta t) + \mu \sin(\beta t)) \mid (\lambda, \mu) \in \mathbb{R}^2\}.$$

Là encore l'ensemble des solutions est un $\mathbb{R}$ -espace vectoriel de dimension 2.

Théorème 10.35

Soit

$$ay''(t) + by'(t) + cy(t) = Ae^{\lambda t} \quad (\text{E})$$

où A et λ sont des constantes complexes.

Il existe une solution particulière de (E) de la forme $t \mapsto Bt^\alpha e^{\lambda t}$ où

- $\alpha = 0$ si λ n'est pas une racine de l'équation caractéristique,
- $\alpha = 1$ si λ est une racine simple de l'équation caractéristique,
- $\alpha = 2$ si λ est une racine double de l'équation caractéristique.

Exemple : On considère l'équation

$$y''(t) - 4y'(t) + 3y(t) = e^t \cos(2t) \quad (\text{E})$$

On a $e^t \cos(2t) = \Re(e^{(1+2i)t})$. On regarde alors

$$y''(t) - 4y'(t) + 3y(t) = e^{(1+2i)t} \quad (\text{E}')$$

L'équation caractéristique est $X^2 - 4X + 3 = (X - 1)(X - 3)$. On a $(1 + 2i)$ n'est pas une racine de l'équation caractéristique. On cherche une solution sous la forme $x \mapsto Be^{(1+2i)t}$. On trouve que $B \in \mathbb{C}$ vérifie :

$$B((1 + 2i)^2 - 4(1 + 2i) + 3) = 1$$

D'où

$$B = -\frac{1}{4(1 + i)} = -\frac{1 - i}{8}.$$

On en déduit que (E') admet pour solution particulière $t \mapsto -\frac{1 - i}{8}e^t(\cos(2t) + i \sin(2t))$. En prenant la partie réelle, on obtient que :

$$t \mapsto -\frac{1}{8}(\cos(2t) + 2 \sin(2t))e^t$$

est une solution particulière de (E).

$$\mathcal{S} = \{t \mapsto -\frac{1}{8}(\cos(2t) + 2 \sin(2t))e^t + \lambda e^t + \mu e^{3t} \mid (\lambda, \mu) \in \mathbb{R}^2\}.$$

Exercice : Résoudre l'équation différentielle : $y'' - (2 \cos \theta)y' + y = e^t$ où $\theta \in \mathbb{R}$.

4.4 Utilisation des séries entières

Nous allons voir deux liens entre les séries entières et les équations différentielles.

- On peut chercher une solution particulière d'une équation différentielle sous la forme d'une série entière.

Par exemple, résolvons,

$$(\text{E}) \quad xy'' + 2y' + xy = 0$$

On cherche une solution f développable en série entière au voisinage de 0. On suppose donc qu'il existe une série entière $\sum_{n \geq 0} a_n x^n$ de rayon de convergence $R > 0$ telle que

$$\forall x \in]-R, R[, f(x) = \sum_{n=0}^{+\infty} a_n x^n.$$

On a alors

$$\begin{aligned} f \text{ vérifie (E)} &\iff \forall x \in]-R, R[, x f''(x) + 2f'(x) + x f(x) = 0 \\ &\iff \forall x \in]-R, R[, x \sum_{n=0}^{+\infty} n(n-1) a_n x^{n-2} + 2 \sum_{n=0}^{+\infty} n a_n x^{n-1} + x \sum_{n=0}^{+\infty} a_n x^n = 0 \\ &\iff \forall x \in]-R, R[, \sum_{n=1}^{+\infty} n(n-1) a_n x^{n-1} + 2 \sum_{n=0}^{+\infty} n a_n x^{n-1} + \sum_{n=0}^{+\infty} a_n x^{n+1} = 0 \\ &\iff \forall x \in]-R, R[, \sum_{n=0}^{+\infty} (n+2)(n+1) a_{n+1} x^n + \sum_{n=1}^{+\infty} a_{n-1} x^n = 0 \\ &\iff \forall x \in]-R, R[, 2a_1 + \sum_{n=1}^{+\infty} [(n+1)(n+2) a_{n+1} + a_{n-1}] x^n = 0 \end{aligned}$$

On trouve donc, par unicité du développement en série entière que nécessairement, $a_1 = 0$ pour tout entier $n \geq 1$, $(n+2)(n+1) a_{n+1} + a_{n-1} = 0$.

On en déduit que pour n impair, $a_n = 0$ et que pour $n = 2p$ pair,

$$a_{2p} = \frac{(-1)^p}{(2p+1)!}$$

Posons maintenant g la somme de la série entière obtenue

$$g : x \mapsto \sum_{p=0}^{+\infty} \frac{(-1)^p}{(2p+1)!} x^{2p}$$

Son rayon de convergence est infini. On a bien obtenu une solution de (E). Notons pour finir que

$$f : x \mapsto \frac{\sin x}{x}.$$

On peut alors vérifier qu'elle vérifie bien l'équation.

- On peut trouver le développement en série entière d'une fonction en vérifiant qu'elle est solution d'une équation « simple ».

Considérons par exemple

$$f : x \mapsto \frac{\arcsin x}{\sqrt{1-x^2}}.$$

On voit qu'elle est développable en série entière sur $] -1, 1[$ comme produit de deux fonctions développables en séries entières. La fonction f est définie et de classe $\mathcal{C}^\infty$ sur $] -1, 1[$. De plus pour x dans cet intervalle

$$f'(x) = \frac{1 + \frac{x}{\sqrt{1-x^2}} \arcsin x}{1-x^2} = \frac{1}{1-x^2} + \frac{x}{1-x^2} f(x).$$

On en déduit que f vérifie l'équation différentielle :

$$(1 - x^2)y' = xy + 1$$

qui peut aussi s'écrire

$$y' = \frac{x}{1 - x^2}y + \frac{1}{1 - x^2}.$$

On voit de plus que $f(0) = 0$ et donc f est l'unique solution du problème de Cauchy

$$\begin{cases} y' = \frac{x}{1 - x^2}y + \frac{1}{1 - x^2} \\ y(0) = 0 \end{cases}$$

On cherche une solution g développable en série entière au voisinage de 0 de cette équation différentielle. On suppose donc qu'il existe une série entière $\sum_{n \geq 0} a_n x^n$ de rayon de convergence

$R > 0$ telle que

$$\forall x \in]-R, R[, g(x) = \sum_{n=0}^{+\infty} a_n x^n.$$

On a alors

$$\begin{aligned} f \text{ vérifie (E)} &\iff \forall x \in]-R, R[, (1 - x^2)f'(x) - xf(x) = 1 \\ &\iff \forall x \in]-R, R[, (1 - x^2) \sum_{n=0}^{+\infty} n a_n x^{n-1} - x \sum_{n=0}^{+\infty} a_n x^n = 1 \\ &\iff \forall x \in]-R, R[, \sum_{n=1}^{+\infty} n a_n x^{n-1} - \sum_{n=0}^{+\infty} n a_n x^{n+1} - \sum_{n=0}^{+\infty} a_n x^{n+1} = 1 \\ &\iff \forall x \in]-R, R[, \sum_{n=0}^{+\infty} (n+1) a_{n+1} x^n - \sum_{n=1}^{+\infty} n a_{n-1} x^n = 1 \\ &\iff \forall x \in]-R, R[, \sum_{n=0}^{+\infty} [(n+1) a_{n+1} - n a_{n-1}] x^n = 1 \end{aligned}$$

On en déduit, en identifiant les deux séries entières que

- Pour $n = 0$, $a_1 = 1$
- Pour $n \geq 1$, $(n+1)a_{n+1} = n a_{n-1}$

On voit déjà que seuls les coefficients impairs sont non nuls. Ce qui est normal car f est impaire.

Cela donne que

$$\forall p \in \mathbb{N}, a_{2p} = \frac{1 \times 3 \times \cdots \times (2p-1)}{2 \times 4 \times \cdots \times 2p} a_0 \text{ et } a_{2p+1} = \frac{2 \times 4 \times \cdots \times 2p}{1 \times 3 \times \cdots \times (2p+1)} a_1.$$

Cependant, comme f s'annule en 0, on regarde la solution g telle que $g(0) = 0$ c'est-à-dire $a_0 = 0$. On pose donc

$$g : x \mapsto \sum_{p=0}^{+\infty} \frac{(2^p p!)^2}{(2p+1)!} x^{2p+1}.$$

On applique la méthode de D'Alembert. On pose $u_p = \frac{(2^p p!)^2}{(2p+1)!} x^{2p+1}$. On a donc

$$\frac{u_{p+1}}{u_p} = \frac{(2^{p+1}(p+1)!)^2}{(2p+3)!} \frac{(2p+1)!}{(2^p p!)^2} x^2 = \frac{4(p+1)^2}{(2p+2)(2p+3)} x^2 \xrightarrow{p \rightarrow \infty} x^2.$$

On en déduit que la série a un rayon de convergence $R = 1 > 0$. D'où $g = f$.

On peut maintenant revenir sur le théorème vu précédemment :

Théorème 10.36

Soit $\alpha \in \mathbb{R}$, la fonction $x \mapsto (1+x)^\alpha$ est développable en série entière et

$$\forall x \in]-1, 1[, (1+x)^\alpha = \sum_{n=0}^{+\infty} \frac{\alpha(\alpha-1) \cdots (\alpha-n+1)}{n!} x^n.$$

Démonstration : On utilise la méthode de l'équation différentielle :

On pose $f : x \mapsto (1+x)^\alpha$. On vérifie que f est dérivable et $(1+x)f'(x) = \alpha f(x)$ pour tout $x \in]-1, 1[$.

Supposons que f est développable en série entière (avec un rayon de convergence R). On pose $f : x \mapsto \sum_{n=0}^{+\infty} a_n x^n$.

En procédant comme précédemment on trouve que

$$\forall x \in]-R, R[, \sum_{n=0}^{+\infty} [(n+1)a_{n+1} + (n-\alpha)a_n] x^n = 0.$$

Notons que comme on cherche une fonction qui vaut 1 en 0, on va prendre $a_0 = 1$. On alors

$$\forall n \in \mathbb{N}, a_{n+1} = \frac{\alpha - n}{n+1} a_n$$

D'où, par une récurrence immédiate,

$$\forall n \in \mathbb{N}, a_n = \frac{\alpha(\alpha-1) \cdots (\alpha-n+1)}{n!}.$$

Il ne reste plus qu'à prouver que cette série a un rayon de convergence 1 en utilisant la méthode de D'Alembert. On pose $u_n = \frac{\alpha(\alpha-1) \cdots (\alpha-n+1)}{n!} x^n$. On a donc

$$\frac{|u_{n+1}|}{|u_n|} = \frac{|\alpha - n|}{n+1} x \xrightarrow{n \rightarrow \infty} x.$$

Le rayon de convergence vaut bien 1. □

Espaces préhilbertiens réels I

1	Rappels	297
1.1	Définition	298
1.2	Orthogonalité et projection orthogonale	300
2	Adjoint d'un endomorphisme	304
2.1	Définition	304
2.2	Propriétés	305
3	Matrices orthogonales	307
3.1	Définition	307
3.2	Groupe spécial orthogonal et orientation	309
3.3	Produit mixte	311
4	Isométries vectorielles d'un espace euclidien	312
4.1	Définition	312
4.2	Propriétés	313
4.3	Isométries directes et indirectes	315
4.4	Isométries vectorielles du plan euclidien	317
4.5	Réduction des isométries vectorielles	319
4.6	Isométries vectorielles de l'espace	322

Nous allons faire quelques rappels sur les espaces préhilbertiens réels puis étudier quelques endomorphismes particuliers des espaces euclidiens.

1 Rappels

1.1 Définition

Définition 11.1 (Produit scalaire)

Soit E un $\mathbf{R}$ -espace vectoriel. On appelle produit scalaire une forme bilinéaire symétrique définie positive. On a donc

- $\forall u \in E, (\bullet|u) \text{ et } (u|\bullet) \text{ sont linéaires}$
- $\forall (u, v) \in E^2, (u|v) = (v|u)$
- $\forall u \in E, (u|u) \geq 0$
- $\forall u \in E, (u|u) = 0 \iff u = 0.$

Définition 11.2

On appelle espace préhilbertien un espace vectoriel réel muni d'un produit scalaire. Si de plus cet espace est de dimension finie, on parle d'espace euclidien.

Matheux (David Hilbert : 1862 - 1943)



David Hilbert, né en 1862 à Königsberg et mort en 1943 à Göttingen, est un mathématicien allemand. Il est souvent considéré comme un des plus grands mathématiciens du xx^e siècle. Il a créé ou développé un large éventail d'idées fondamentales, que ce soit la théorie des invariants, l'axiomatisation de la géométrie ou les fondements de l'analyse fonctionnelle (avec les espaces de Hilbert).

L'un des exemples les mieux connus de sa position de chef de file est sa présentation, en 1900, de ses fameux problèmes qui ont durablement influencé les recherches mathématiques du xx^e siècle. Hilbert et ses étudiants ont fourni une portion significative de l'infrastructure mathématique nécessaire à l'éclosion de la mécanique quantique et de la relativité générale.

Il a adopté et défendu avec vigueur les idées de Georg Cantor en théorie des ensembles et sur les nombres transfinis. Il est aussi connu comme l'un des fondateurs de la théorie de la démonstration, de la logique mathématique et a clairement distingué les mathématiques des métamathématiques.

Définition 11.3

Soit E un espace préhilbertien. L'application

$$\begin{aligned} \|\cdot\| : E &\rightarrow \mathbf{R} \\ x &\mapsto \|x\| = \sqrt{(x|x)} \end{aligned}$$

est une norme que l'on appelle norme euclidienne.

Exemple : Structure euclidienne canonique de $\mathbf{R}^n$: La structure euclidienne canonique de $\mathbf{R}^n$ (c'est-

à-dire celle pour laquelle la base canonique est une base orthonormée) est définie par

$$\begin{aligned} \langle \bullet | \bullet \rangle : \quad \mathbf{R}^n &\rightarrow \mathbf{R} \\ (x_1, \dots, x_n), (y_1, \dots, y_n) &\mapsto \sum_{i=1}^n x_i y_i \end{aligned}$$

En particulier si on identifie $\mathbf{R}^n$ avec $\mathcal{M}_{n,1}(\mathbf{R})$. Pour X, Y dans $\mathcal{M}_{n,1}(\mathbf{R})$,

$$\langle X | Y \rangle = X^\top Y$$

Ci-dessus, $X^\top Y \in \mathcal{M}_1(\mathbf{R})$ que l'on identifie à $\mathbf{R}$.

Proposition 11.4

Soit E un espace euclidien et $\mathcal{B} = (e_1, \dots, e_n)$ une base orthonormée. Soit u et v deux vecteurs de E . Si on note

$$U = \begin{pmatrix} u_1 \\ \vdots \\ u_n \end{pmatrix} \text{ et } V = \begin{pmatrix} v_1 \\ \vdots \\ v_n \end{pmatrix}$$

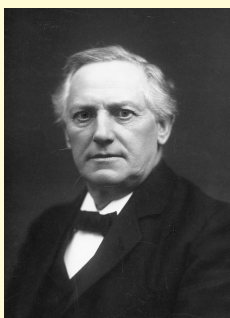
les matrices colonnes des coordonnées de u et v dans la base $\mathcal{B}$. On a

$$(u|v) = \sum_{i=1}^n u_i v_i = U^\top V = \langle U | V \rangle$$

Théorème 11.5 (Procédé d'orthonormalisation de Gram-Schmidt)

Soit $(e_1, \dots, e_p)$ une famille libre de vecteurs de E . Il existe une famille orthonormale $(f_1, \dots, f_p)$ telle que, pour tout $k \in \llbracket 1 ; p \rrbracket$, $\text{Vect}(e_1, \dots, e_k) = \text{Vect}(f_1, \dots, f_k)$.

Matheux (Jørgen Pedersen Gram : 1850 - 1916)



Jørgen Pedersen Gram est le fils de l'agriculteur Peder Jorgensen Gram et de Marie Magdalene Aakjaer. En 1868, il étudie à Copenhague et en 1873 obtient son diplôme. L'année suivante, il écrit un article sur la théorie des invariants.

Gram travailla la majeure partie de sa vie dans le domaine des assurances, c'est ainsi qu'il fonda en 1884 son entreprise, Skjold. Il fut le président du conseil danois de l'assurance de 1910 à 1916.

À partir de 1888, il fait partie de l'Académie des sciences danoise. Le procédé de Gram-Schmidt fut énoncé en 1883 et en 1884, Gram reçut la médaille d'or de l'Académie pour un travail sur les nombres premiers.

Il décéda après avoir été heurté par un vélo, durant un trajet qui le menait vers l'Académie.

Matheux (Erhard Schmidt : 1876 - 1959)



Il a soutenu en 1905 une thèse sur les équations intégrales à l'université de Göttingen, sous la direction de David Hilbert. Avec ce dernier, il est considéré comme l'un des fondateurs de l'analyse fonctionnelle abstraite moderne. En 1948, il fonde la revue Mathematische Nachrichten et en devient son premier rédacteur en chef

Exemple : En Python si on suppose avoir une fonction `scal` qui permet de calculer le produit scalaire de deux vecteurs. On peut écrire une fonction qui, à partir d'une liste de vecteurs renvoie la famille orthonormale obtenue par le procédé d'orthonormalisation de Gram-Schmidt.

```
def GS(l) :  
    res = []  
    for e in l :  
        ftemp = e  
        for f in res :  
            ftemp = ftemp - scal(e,f)*f  
        res.append(ftemp / sqrt(scal(ftemp,ftemp)))  
    return(res)
```

1.2 Orthogonalité et projection orthogonale

Soit E un espace préhilbertien. On le considère comme un espace vectoriel normé par la norme euclidienne notée $\|\cdot\|$

Proposition-Définition 11.6 (Orthogonal)

Soit F un sous-espace vectoriel d'un espace préhilbertien. On appelle orthogonal de F et on note $F^\perp = \{u \in E \mid \forall v \in F, (u|v) = 0\}$ l'ensemble des vecteurs orthogonaux à tous les vecteurs de F . C'est un espace vectoriel.

Théorème 11.7

Si F est un sous-espace vectoriel de dimension finie, $F \oplus F^\perp = E$. On dit que F et $F^\perp$ sont en somme directe orthogonale.

ATTENTION

Ce résultat n'est plus vrai si F n'est pas de dimension finie - voir exemple plus loin

Démonstration :

On a vu que $F \cap F^\perp = \{0\}$. Donnons deux démonstrations du fait que $F + F^\perp = E$.

– Première preuve : On procède par analyse-synthèse. On considère une base $(f_1, \dots, f_p)$ orthonormée de F que l'on peut obtenir par orthonormalisation.

– Analyse : Soit $x \in E$ et $(u, v) \in F \times F^\perp$ tels que $x = u + v$. On a donc $u = \sum_{i=1}^p \lambda_i f_i$ et

$v = x - u = x - \sum_{i=1}^p \lambda_i f_i$. Par hypothèse, $(v|f_1) = 0$ ce qui signifie que

$$\left(x - \sum_{i=1}^p \lambda_i f_i | f_1 \right) = 0 \Rightarrow (x|f_1) = \lambda_1 (f_1|f_1) = \lambda_1.$$

De la même manière, pour tout $i \in \llbracket 1; p \rrbracket$, $(v|f_i) = 0$ et donc $\lambda_i = (x|f_i)$.

– Synthèse : soit $x \in E$, on pose pour tout $i \in \llbracket 1; p \rrbracket$, $\lambda_i = (x|f_i)$. On pose alors $u = \sum_{i=1}^p \lambda_i f_i$ et $v = x - u$. Il est clair que $u \in F$, $u + v = x$ et le calcul précédent montre que pour tout $i \in \llbracket 1; p \rrbracket$, $(v|f_i) = 0$ donc $v \in F^\perp$.

– Deuxième preuve : On considère cette fois une base **non nécessairement orthonormée** $(f_1, \dots, f_p)$ de F . On reprend le même principe on veut écrire $x = u + v$ avec $u \in F$ et $v \in F^\perp$ c'est-à-dire que $u = \sum_{i=1}^p \lambda_i f_i$ et $\forall i \in \llbracket 1; p \rrbracket$, $(v|f_i) = 0$. Cherchons encore les coefficients λ_i .

Le fait que $(x - u|f_1) = 0$ donne :

$$(x|f_1) = \sum_{i=1}^p \lambda_i (f_i|f_1)$$

En prenant toutes les équations on trouve que le vecteur $\Lambda = \begin{pmatrix} \lambda_1 \\ \vdots \\ \lambda_p \end{pmatrix}$ doit vérifier l'équation

$$G\Lambda = X$$

où

$$G = \begin{pmatrix} (f_1|f_1) & \cdots & (f_p|f_1) \\ \vdots & & \vdots \\ (f_1|f_p) & \cdots & (f_p|f_p) \end{pmatrix} \text{ et } X = \begin{pmatrix} (x|f_1) \\ \vdots \\ (x|f_p) \end{pmatrix}$$

La matrice G s'appelle la matrice de Gram de la famille.

Remarquons alors que la matrice G est inversible en effet son noyau est réduit à $\{0\}$. Cela se déduit du fait que si $Y = \begin{pmatrix} y_1 \\ \vdots \\ y_p \end{pmatrix}$ alors pour tout $i \in \llbracket 1; p \rrbracket$ $\sum_{j=1}^p y_j (f_j|f_i) = 0$ de ce fait

$$Y^\top G Y = \sum_{i=1}^p y_i \sum_{j=1}^p y_j (f_j|f_i) = 0$$

or ce terme est juste $(\sum y_i f_i | \sum y_j f_j) = \|\sum y_i f_i\|_2^2$. On en déduit que tous les y_i sont nuls car $(f_1, \dots, f_p)$ est libre.

Finalement ce système a une unique solution Λ qui répond au problème.

□

Exemple : Donnons un exemple d'espace préhilbertien E et de sous espace vectoriel F tel que $F^\perp \oplus F \neq E$. On pose $E = \{(u_n)_{n \geq 0} \mid \sum_{n=0}^{+\infty} u_n^2 < +\infty\}$ l'ensemble des familles de carré sommable.

Commençons par vérifier que c'est un bien un sous-espace vectoriel de $\mathbf{R}^{\mathbf{N}}$. Pour cela, montrons que si $(u_n)_{n \geq 0}$ et $(v_n)_{n \geq 0}$ appartiennent à E alors la famille $(u_n v_n)_{n \geq 0}$ est sommable. En effet d'après l'inégalité arithmético-géométrique, pour tout $n \in \mathbf{N}$,

$$|u_n v_n| \leq \frac{1}{2}(u_n^2 + v_n^2)$$

On en déduit que dans $[0, +\infty]$,

$$\sum_{n=0}^{+\infty} |u_n v_n| \leq \frac{1}{2} \left(\sum_{n=0}^{+\infty} u_n^2 + \sum_{n=0}^{+\infty} v_n^2 \right) < +\infty$$

On en déduit que si $(u_n)_{n \geq 0}$ et $(v_n)_{n \geq 0}$ sont des éléments de E alors, pour tout $(\lambda, \mu) \in \mathbf{R}^2$,

$$(\lambda u_n + \mu v_n)^2 = \lambda^2 u_n^2 + 2\lambda\mu u_n v_n + \mu^2 v_n^2$$

On en déduit que $\sum_{n \geq 0} (\lambda u_n + \mu v_n)^2$ converge. Cela montre que E est bien un sous-espace vectoriel.

On vérifie alors facilement qu'en posant pour $U = (u_n)_{n \geq 0}$ et $V = (v_n)_{n \geq 0}$,

$$(U|V) = \sum_{n=0}^{+\infty} u_n v_n$$

on obtient un produit scalaire.

On considère alors $F \subset E$ l'ensemble des suites stationnaires à 0. Pour tout entier $m \in \mathbf{N}$ on note $\delta(m) \in F$ la suite définie par

$$\forall n \in \mathbf{N}, \delta(m)_n = \delta_{n,m}$$

Maintenant, pour toute suite $U = (u_n)_{n \geq 0}$ appartenant à $F^\perp$, on a, pour tout entier m ,

$$0 = (U|\delta(m)) = u_m$$

Cela montre que $F^\perp = \{0\}$ et donc $F \oplus F^\perp = F \subsetneq E$.

Proposition 11.8

Avec les notations précédentes, si $(f_1, \dots, f_p)$ est une base orthonormée de F alors pour tout vecteur x ,

$$p_F(x) = \sum_{i=1}^p \lambda_i f_i \text{ où } \lambda_i = (x|f_i).$$

Démonstration : Cela a été vu dans la démonstration précédente. □

Définition 11.9

Soit F un sous-espace vectoriel de dimension finie. On appelle projection orthogonale sur F et on note p_F la projection sur F parallèlement à $F^\perp$

Remarque : Les démonstrations ci-dessus donne un moyen de calculer le projeté orthogonal d'un vecteur x (le vecteur u de la démonstration). Il faut utiliser la première si on a déjà une base orthonormée de F et la deuxième dans le cas contraire (à moins de calculer au préalable l'orthonormalisation de la base).

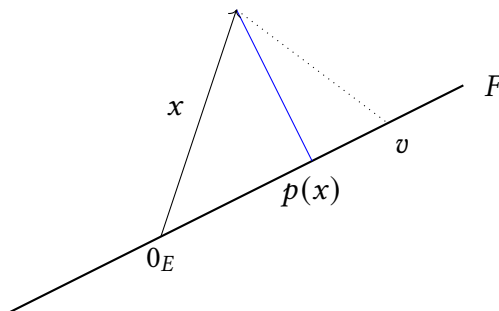
Proposition 11.10

Soit F un sous-espace vectoriel de dimension finie. Pour tout x de E ,

$$d(x, F) = \|x - p(x)\|_2.$$

C'est-à-dire que $p(x)$ est le vecteur de F le plus proche de x .

Démonstration :



Pour tout v de F

$$\begin{aligned} d(v, x)^2 &= (x - v | x - v) \\ &= (x - p_F(x) + p_F(x) - v | x - p_F(x) + p_F(x) - v) \\ &= \|x - p_F(x)\|_2^2 + 2(x - p_F(x) | p_F(x) - v) + \|p_F(x) - v\|_2^2. \end{aligned}$$

Or $p_F(x) - v \in F$ et $x - p_F(x) \in F^\perp$ donc $(x - p_F(x) | p_F(x) - v) = 0$. Finalement,

$$d(v, x)^2 = (x - v | x - v) = \|x - p_F(x)\|_2^2 + \|p_F(x) - v\|_2^2 \geq \|x - p_F(x)\|_2^2 = d(x, p_F(x))^2.$$

□

Exemple : On veut calculer

$$\inf_{(a,b) \in \mathbb{R}^2} \int_0^{+\infty} e^{-t} (t^2 - (at + b))^2 dt.$$

On considère $E = \{f \in \mathcal{C}^0(\mathbb{R}_+, \mathbb{R}) \mid t \mapsto e^{-t}f(t) \text{ est intégrable}\}$. En particulier, les fonctions polynomiales appartiennent à E . C'est un espace vectoriel (sous-espace vectoriel de $\mathcal{C}^0(\mathbb{R}_+, \mathbb{R})$). On peut aussi travailler avec $E = \mathbb{R}[X]$.

On pose alors le produit scalaire

$$\forall (f, g) \in E^2, (f | g) = \int_0^{+\infty} e^{-t} f(t) g(t) dt.$$

Le fait que ce soit un produit scalaire est « classique ». Si on pose $F = \mathbb{R}_1[X]$, la question est de calculer $d(X^2, F)^2$. Pour cela on calcule $P = p_F(X^2)$. Il est de la forme $P = a + bX \in F$ et on sait que $(X^2 - P | 1) = (X^2 - P | X) = 0$. Calculons de manière générale

$$(X^p | X^q) = \int_0^{+\infty} e^{-t} t^{p+q} dt \underset{\text{IPP}}{=} (p+q) \int_0^{+\infty} e^{-t} t^{p+q-1} dt \underset{\text{IPP}}{=} \cdots \underset{\text{IPP}}{=} (p+q)! \int_0^{+\infty} e^{-t} dt = (p+q)!$$

On en déduit que a et b vérifient

$$\begin{cases} a + b = 2 \\ a + 2b = 6 \end{cases}$$

En résolvant, on trouve $a = 4$ et $b = -2$ et donc $P = 4 - 2X$. Pour finir, en utilisant le théorème de Pythagore

$$d(X^2, F)^2 = \|X^2 - p_F(X^2)\|_2^2 = \|X^2\|_2^2 - \|p_F(X^2)\|_2^2 = 24 - 8 = 16.$$

En conclusion,

$$\inf_{(a,b) \in \mathbb{R}^2} \int_0^{+\infty} e^{-t} (t^2 - (at + b))^2 dt = 4.$$

2 Adjoint d'un endomorphisme

Dans ce paragraphe E désigne un espace euclidien

2.1 Définition

Théorème 11.11 (Théorème de représentation)

Soit f une forme linéaire sur E . Il existe un unique vecteur a tel que

$$\forall x \in E, (a|x) = f(x).$$

C'est-à-dire que l'application

$$\begin{aligned} \Phi : E &\mapsto E^* \\ a &\mapsto (x \mapsto (a|x)) \end{aligned}$$

est un isomorphisme.

Démonstration : Soit $a \in E$. Notons $\phi_a = \Phi(a)$ l'application $\phi_a : x \mapsto (a|x)$. Il est clair que c'est une forme linéaire. De plus l'application Φ est bien linéaire car pour $(a_1, a_2) \in E^2$ et $(\lambda_1, \lambda_2) \in \mathbb{R}^2$,

$$\forall x \in E, (\lambda_1 a_1 + \lambda_2 a_2 | x) = \lambda_1 (a_1 | x) + \lambda_2 (a_2 | x)$$

c'est-à-dire que

$$\Phi(\lambda_1 a_1 + \lambda_2 a_2) = \phi_{\lambda_1 a_1 + \lambda_2 a_2} = \lambda_1 \phi_{a_1} + \lambda_2 \phi_{a_2} = \lambda_1 \Phi(a_1) + \lambda_2 \Phi(a_2)$$

De plus Φ est injective car pour $a \in E$, si on suppose que $\Phi(a) = \phi_a$ est l'application nulle alors $\phi_a(a) = (a|a) = 0$ donc $a = 0$.

Comme E est de dimension finie et que $\dim E^* = \dim E$ alors l'application $\Phi : a \mapsto \Phi(a) = \phi_a$ est un isomorphisme. Elle est donc en particulier surjective. $\square$

Proposition-Définition 11.12 (Adjoint d'un endomorphisme d'un espace euclidien)

Soit $u \in \mathcal{L}(E)$. Pour tout $y \in E$, l'application $x \mapsto (u(x)|y)$ est une forme linéaire sur E . Il existe donc un unique vecteur $u_y \in E$ tel que

$$\forall x \in E, (u(x)|y) = (x|u_y)$$

De plus, l'application $y \mapsto u_y$ est un endomorphisme de E . On l'appelle l'adjoint de u et on le note u^* . On a donc

$$\forall (x, y) \in E^2, (u(x)|y) = (x|u^*(y))$$

Démonstration : Soit $y \in E$, $f_y : x \mapsto (u(x)|y)$ est une forme linéaire car c'est la composée de u qui est linéaire avec ϕ_y qui est une forme linéaire. On en déduit en utilisant le théorème de représentation qu'il existe un unique vecteur $u_y \in E$ tel que $\Phi(u_y) = \phi_y \circ u$.

Vérifions maintenant que $y \mapsto u_y$ est linéaire. Soit y_1, y_2 dans E et λ_1, λ_2 dans $\mathbf{R}$, on veut vérifier que

$$u_{\lambda_1 y_1 + \lambda_2 y_2} = \lambda_1 u_{y_1} + \lambda_2 u_{y_2}$$

Or, pour tout $x \in E$,

$$\begin{aligned} (x|u_{\lambda_1 y_1 + \lambda_2 y_2}) &= (u(x)|\lambda_1 y_1 + \lambda_2 y_2) \text{ par définition} \\ &= \lambda_1 (u(x)|y_1) + \lambda_2 (u(x)|y_2) \text{ par linéarité du produit scalaire} \\ &= \lambda_1 (x|u_{y_1}) + \lambda_2 (x|u_{y_2}) \\ &= (x|\lambda_1 u_{y_1} + \lambda_2 u_{y_2}) \end{aligned}$$

On vient de montrer que $\Phi(u_{\lambda_1 y_1 + \lambda_2 y_2}) = \Phi(\lambda_1 u_{y_1} + \lambda_2 u_{y_2})$. On peut conclure en utilisant l'injectivité de Φ . $\square$

Exemple : Dans $E = \mathbf{R}^2$ muni du produit scalaire usuel. On considère l'endomorphisme

$$u : (x_1, x_2) \mapsto (x_1 + x_2, 3x_1 - 2x_2)$$

Pour tout $x = (x_1, x_2)$ et $y = (y_1, y_2)$ dans E on a

$$\begin{aligned} (u(x)|y) &= ((x_1 + x_2, 3x_1 - 2x_2)|(y_1, y_2)) \\ &= (x_1 + x_2)y_1 + (3x_1 - 2x_2)y_2 \\ &= x_1(y_1 + 3y_2) + x_2(y_1 - 2y_2) \\ &= ((x_1, x_2)|(y_1 + 3y_2, y_1 - 2y_2)) \end{aligned}$$

On en déduit que $u^* : (y_1, y_2) \mapsto (y_1 + 3y_2, y_1 - 2y_2)$.

2.2 Propriétés

Proposition 11.13 (Linéarité)

L'application $u \mapsto u^*$ de $\mathcal{L}(E)$ dans $\mathcal{L}(E)$ est linéaire. Pour tout u_1, u_2 dans $\mathcal{L}(E)$ et tout λ_1, λ_2 dans $\mathbf{R}$,

$$(\lambda_1 u_1 + \lambda_2 u_2)^* = \lambda_1 u_1^* + \lambda_2 u_2^*$$

Démonstration : Par unicité de l'adjoint, il suffit de vérifier que pour tout x, y dans E ,

$$\begin{aligned}(x|(\lambda_1 u_1 + \lambda_2 u_2)^*(y)) &= ((\lambda_1 u_1 + \lambda_2 u_2)(x)|y) \\ &= \lambda_1 (u_1(x)|y) + \lambda_2 (u_2(x)|y) \\ &= \lambda_1 (x|u_1^*(y)) + \lambda_2 (x|u_2^*(y)) \\ &= (x|(\lambda_1 u_1^* + \lambda_2 u_2^*)(y))\end{aligned}$$

□

Proposition 11.14 (Propriétés de l'adjoint)

1. Soit u, v dans $\mathcal{L}(E)$, $(u \circ v)^* = v^* \circ u^*$.
2. Soit u dans $\mathcal{L}(E)$, $(u^*)^* = u$

Démonstration :

1. Soit x, y dans E ,

$$(x|(u \circ v)^*(y)) = (u(v(x))|y) = (v(x)|u^*(y)) = (x|v^*(u^*(y)))$$

Cela montre bien que $(u \circ v)^* = v^* \circ u^*$.

2. Soit x, y dans E ,

$$(x|(u^*)^*(y)) = (u^*(x)|y) = (x|u(y))$$

On en déduit que $(u^*)^* = u$.

□

Proposition 11.15 (Matrice de l'adjoint dans une base orthonormée)

Soit $u \in \mathcal{L}(E)$ et $\mathcal{B} = (e_1, \dots, e_n)$ une base orthonormée de E . On a

$$\text{Mat}_{\mathcal{B}}(u^*) = (\text{Mat}_{\mathcal{B}}(u))^{\top}$$

Démonstration : Notons $A = \text{Mat}_{\mathcal{B}}(u)$ et $B = \text{Mat}_{\mathcal{B}}(u^*)$. Comme $\mathcal{B}$ est une base orthonormée, pour tout $(i, j) \in \llbracket 1 ; n \rrbracket^2$,

$$B[i, j] = (u^*(e_j)|e_i) = (e_j|u(e_i)) = A[j, i]$$

On en déduit que $B = A^{\top}$.

□

Remarques :

1. On peut retrouver, via des calculs matriciels que $u \mapsto u^*$ est linéaire, que $(u \circ v)^* = v^* \circ u^*$ et que $(u^*)^* = u$ car ce sont des propriétés de la l'application $A \mapsto A^{\top}$ de $\mathcal{M}_n(\mathbb{R})$ dans lui même.
2. Cette propriété n'est plus vraie si $\mathcal{B}$ n'est pas orthonormée.

Exemple : Si on reprend le calcul fait ci-dessus. La base canonique est orthonormée et on a

$$A = \text{Mat}_{(\text{can})}(u) = \begin{pmatrix} 1 & 1 \\ 3 & -2 \end{pmatrix} \text{ et } B = \text{Mat}_{(\text{can})}(u^*) = \begin{pmatrix} 1 & 3 \\ 1 & -2 \end{pmatrix} = A^{\top}$$

Proposition 11.16 (Sous-espaces stables)

Soit $u \in \mathcal{L}(E)$. Soit F un sous-espace vectoriel de E stable par u . Le sous-espace vectoriel $F^\perp$ est stable par u^* .

Démonstration :

- Preuve matricielle On construit une base orthonormée $\mathcal{B}$ de E en concaténant une base orthonormée $\mathcal{B}_F$ de F avec une base orthonormée $\mathcal{B}_{F^\perp}$ de $F^\perp$.

Comme F est stable par u , la matrice A de u dans la base $\mathcal{B}$ est de la forme

$$A = \text{Mat}_{\mathcal{B}}(u) = \left(\begin{array}{c|c} X & Y \\ \hline 0 & Z \end{array} \right)$$

On en déduit que la matrice de u^* dans cette même base est

$$\text{Mat}_{\mathcal{B}}(u^*) = A^\top = \left(\begin{array}{c|c} X^\top & 0 \\ \hline Y^\top & Z^\top \end{array} \right)$$

Cela montre que $F^\perp$ est stable par u^* .

- Preuve directe : Soit $x \in F^\perp$. On veut montrer que $u^*(x) \in F^\perp$.

Pour tout $y \in F$,

$$(y|u^*(x)) = (u(y)|x) \stackrel{u(y) \in F}{=} 0$$

Le vecteur $u^*(x)$ est orthogonal à tout vecteur de F donc $u^*(x) \in F^\perp$.

On a bien montré que $F^\perp$ est stable par u^* .

□

3 Matrices orthogonales

3.1 Définition

Définition 11.17

Soit $M \in \mathcal{M}_n(\mathbf{R})$. Les conditions suivantes sont équivalentes :

- i) La matrice M est inversible et $M^{-1} = M^\top$.
- ii) Les colonnes de M forment une base orthonormée.
- iii) Les lignes de M forment une base orthonormée.

Remarques :

1. Pour la condition i) il suffit de vérifier que $MM^\top = I_n$ ou que $M^\top M = I_n$.
2. Dans les conditions ii) et iii) on a identifié les colonnes ou les lignes de la matrice à $\mathbf{R}^n$ avec son produit scalaire usuel.

Démonstration : Si on note $M = (a_{ij})_{(i,j) \in \llbracket 1; n \rrbracket^2}$, C_i les colonnes de M et que l'on note $(b_{ij})_{(i,j) \in \llbracket 1; n \rrbracket^2}$ les coefficients de $M^\top M$. On a

$$\forall (i, j) \in \llbracket 1; n \rrbracket^2, b_{ij} = C_i^\top C_j = \sum_{k=1}^n a_{ki} a_{kj} = \langle C_i | C_j \rangle$$

On a donc que les colonnes de M forment une base orthonormée si et seulement si $M^T M = I_n$.

De même les lignes de M forment une base orthonormée si et seulement si $MM^T = I_n$.

On conclut en utilisant que

$$M^T M = I_n \iff MM^T = I_n \iff M^{-1} = M^T.$$

□

Définition 11.18

Une matrice vérifiant les conditions ci-dessus s'appelle une matrice orthogonale. On note $O(n)$ ou $O_n(\mathbf{R})$ l'ensemble des matrices orthogonales de $\mathcal{M}_n(\mathbf{R})$.

Remarque : On voit que $M \in O(n) \iff M^T \in O(n)$.

Exemples :

1. Une matrice diagonale avec des éléments de $\{\pm 1\}$ sur la diagonale est orthogonale.
2. Les matrices $A = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}$ sont orthogonales.

Proposition 11.19

Soit M une matrice orthogonale. Son déterminant appartient à $\{\pm 1\}$.

Démonstration : Il suffit de voir que $M^T M = I_n$ implique $\det(M)^2 = 1$.

□

ATTENTION

Ce n'est pas une équivalence. Par exemple, la matrice $\begin{pmatrix} 2 & 5 \\ 1 & 3 \end{pmatrix}$. Son déterminant vaut 1 mais elle n'est pas orthogonale.

Proposition 11.20

Soit E un espace euclidien. Soit $\mathcal{B}$ une base orthonormée et $\mathcal{B}'$ une autre base de E . La base $\mathcal{B}'$ est orthonormée si et seulement si la matrice de changement de base P est orthogonale.

Démonstration : Notons $\mathcal{B} = (e_1, \dots, e_n)$ et $\mathcal{B}' = (e'_1, \dots, e'_n)$. On pose pour tout $i \in \llbracket 1; n \rrbracket$, $C_i = \text{Mat}_{\mathcal{B}}(e'_i)$ la i -ième colonne de P . Comme $\mathcal{B}$ est une base orthonormée, pour tout $(i, j) \in \llbracket 1; n \rrbracket^2$, $\langle e'_i | e'_j \rangle = \langle C_i | C_j \rangle$. On en déduit que

$$\begin{aligned} \text{la base } \mathcal{B}' \text{ est orthonormée} &\iff \forall (i, j) \in \llbracket 1; n \rrbracket^2, \langle e'_i | e'_j \rangle = \delta_{ij} \\ &\iff \forall (i, j) \in \llbracket 1; n \rrbracket^2, \langle C_i | C_j \rangle = \delta_{ij} \\ &\iff P \in O_n(\mathbf{R}) \end{aligned}$$

□

Définition 11.21 (Matrices orthogonalement semblables)

Soit A, B dans $\mathcal{M}_n(\mathbf{R})$. La matrice B est orthogonalement semblable à A s'il existe $P \in O_n(\mathbf{R})$ telles que $A = P^{-1}BP = P^\top BP$.

Remarque : Quand on travaille sur des endomorphismes d'un espace euclidien, on veut la plupart du temps ne considérer que des bases orthogonales. Contrairement au cas de la réduction des endomorphismes classique où l'on considère des changements de bases en général, on va se restreindre à des changements de bases entre bases orthonormées. On va donc chercher à « remplacer » une matrice A par une matrice orthogonalement semblable le plus simple possible.

Proposition 11.22

La relation binaire $\mathcal{R}$ sur $\mathcal{M}_n(\mathbf{R})$ définie par $A\mathcal{R}B$ si et seulement si B est orthogonalement semblable à A est une relation d'équivalence.

Démonstration : En exercice □

Proposition 11.23

Soit $n \geq 1$. L'ensemble $O_n(\mathbf{R})$ des matrices orthogonales est un sous-groupe (multiplicatif) de $(GL_n(\mathbf{R}), \times)$.
En particulier, le produit de deux matrices orthogonales est orthogonale et l'inverse d'une matrice orthogonale est orthogonale.

Démonstration : On commence par remarquer que, par définition, toute matrice orthogonale est inversible. On a bien $O_n(\mathbf{R}) \subset GL_n(\mathbf{R})$. Montrons que c'est un sous-groupe.

- Comme $I_n \in O_n(\mathbf{R})$, $O_n(\mathbf{R}) \neq \emptyset$.
- Soit A, B dans $O_n(\mathbf{R})$, on pose $C = AB^{-1}$. Montrons que $C \in O_n(\mathbf{R})$. On vérifie que

$$C^\top C = ((AB^{-1})^\top AB^{-1} = (B^{-1})^\top A^\top AB^{-1} = BB^{-1} = I_n$$

On a bien montré que $O_n(\mathbf{R})$ est un sous-groupe de $(GL_n(\mathbf{R}), \times)$. □

3.2 Groupe spécial orthogonal et orientation

Définition 11.24

On appelle *groupe spécial orthogonal* et on note $SO_n(\mathbf{R})$ ou $SO(n)$ le sous-groupe de $O_n(\mathbf{R})$ constitué des matrices orthogonales de déterminant 1.

Les matrices de $SO_n(\mathbf{R})$ s'appellent les *matrices orthogonales positives ou directes*. Les matrices de $O_n(\mathbf{R})$ qui n'appartiennent pas à $SO_n(\mathbf{R})$ s'appellent les *matrices orthogonales négatives ou indirectes*.

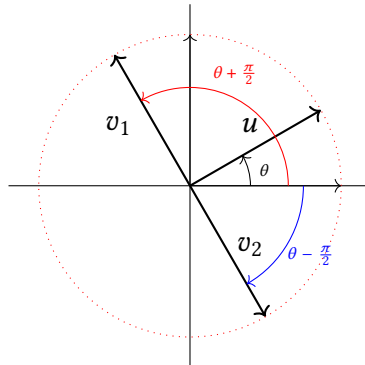
Remarque : On peut justifier que c'est un sous-groupe de $O_n(\mathbf{R})$ en voyant que c'est le noyau du déterminant qui est un morphisme de groupes

$$\det : O_n(\mathbf{R}) \rightarrow \{\pm 1\}$$

Exemple : On considère l'espace euclidien $\mathbf{R}^2$ muni du produit scalaire usuel. On veut construire une base orthonormée (u, v) . Le vecteur $u = (x_u, y_u)$ est nécessairement de norme 1. Cela signifie que $x_u^2 + y_u^2 = 1$. On sait alors qu'il existe $\theta \in \mathbf{R}$ tel que $u = (\cos(\theta), \sin(\theta))$. Si on note alors $v = (x_v, y_v)$. Par hypothèses, v est aussi de norme 1 donc il existe $\varphi \in \mathbf{R}$ tel que $v = (\cos(\varphi), \sin(\varphi))$. De plus

$$0 = (u|v) = \cos(\theta) \cos(\varphi) + \sin(\theta) \sin(\varphi) = \cos(\theta - \varphi)$$

On en déduit que $\theta - \varphi \equiv \frac{\pi}{2} [\pi]$, c'est-à-dire $\varphi \equiv \theta + \frac{\pi}{2} [2\pi]$ ou $\varphi \equiv \theta - \frac{\pi}{2} [2\pi]$. Posons donc $\varphi_1 = \theta + \frac{\pi}{2}$ et $\varphi_2 = \theta - \frac{\pi}{2}$ ainsi que $v_1 = (\cos(\varphi_1), \sin(\varphi_1))$ et $v = (\cos(\varphi_2), \sin(\varphi_2))$



Si on regarde les matrices de changement de bases P_1 et P_2 de la base canonique aux bases $\mathcal{B}_1 = (u, v_1)$ et $\mathcal{B}_2 = (u, v_2)$, on a

$$P_1 = \begin{pmatrix} \cos(\theta) & \cos(\varphi_1) \\ \sin(\theta) & \sin(\varphi_1) \end{pmatrix} = \begin{pmatrix} \cos(\theta) & -\sin(\theta) \\ \sin(\theta) & \cos(\theta) \end{pmatrix} \text{ et } P_2 = \begin{pmatrix} \cos(\theta) & \cos(\varphi_2) \\ \sin(\theta) & \sin(\varphi_2) \end{pmatrix} = \begin{pmatrix} \cos(\theta) & \sin(\theta) \\ \sin(\theta) & -\cos(\theta) \end{pmatrix}$$

Ce sont bien sur des matrices orthogonales car ce sont des matrices de changement de bases entre bases orthonormées. Cependant,

$$\det(P_1) = \cos^2(\theta) + \sin^2(\theta) = 1 \text{ et } \det(P_2) = -\cos^2(\theta) - \sin^2(\theta) = -1$$

On voit que $P_1 \in SO_2(\mathbf{R})$ et $P_2 \notin SO_2(\mathbf{R})$.

Proposition 11.25

On considère sur l'ensemble des bases orthonormées de E la relation binaire suivante :

$$\mathcal{B} R \mathcal{B}' \iff \det(P_{\mathcal{B}'}^{\mathcal{B}}) = 1$$

1. La relation R est une relation d'équivalence.
2. La relation R a deux classes d'équivalence.

Démonstration :

1. Évident
2. Soit $\mathcal{B} = (e_1, \dots, e_n)$ un base orthonormée (il en existe par le procédé de Gram-Schmidt). On considère $\mathcal{B}' = (e_1, \dots, e_{n-1}, -e_n)$ qui est obtenue en prenant l'inverse du dernier vecteur. Il est clair que $\mathcal{B}'$ est encore une base orthonormée. De plus la matrice de passage de $\mathcal{B}$ à $\mathcal{B}'$ est

$$P = \begin{pmatrix} 1 & 0 & \cdots & 0 \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & 1 & 0 \\ 0 & \cdots & 0 & -1 \end{pmatrix}$$

Comme $\det(P) = -1$, les bases $\mathcal{B}$ et $\mathcal{B}'$ ne sont pas en relation. Par contre, pour toute base orthonormée $\mathcal{C}$. Si on considère Q la matrice de changement de base de $\mathcal{C}$ à $\mathcal{B}$ alors PQ est la matrice de changement de base de $\mathcal{C}$ à $\mathcal{B}'$. Comme $\det(PQ) = \det(P) \det(Q) = -\det(Q)$, la base $\mathcal{C}$ est en relation avec $\mathcal{B}$ ou avec $\mathcal{B}'$.

La relation R a bien deux classes d'équivalence.

□

Définition 11.26 (Orientation)

Soit E un espace euclidien. On appelle orientation de l'espace E le choix d'une des deux classes d'équivalence de la relation R .

Les bases qui appartiennent à cette classe s'appellent les bases orthonormées directes.

Remarques :

1. Le choix d'une orientation se fait la plupart du temps en choisissant une base $\mathcal{B}$ et en disant que les bases orthonormées directes sont les bases qui appartiennent à la classe d'équivalence de $\mathcal{B}$. Il est important de comprendre que ce choix est purement arbitraire.
2. Dans le cas où l'espace que l'on considère est $\mathbb{R}^n$ avec le produit scalaire usuel ou $\mathcal{M}_n(\mathbb{R})$ avec le produit scalaire usuel. On oriente la plupart du temps ces espaces en prenant la base canonique comme base directe.

3.3 Produit mixte

Proposition-Définition 11.27 (Rappel)

Soit E un espace vectoriel de dimension finie notée n .

1. L'ensemble des formes n -linéaires alternées est de dimension 1.
2. Pour toute base $\mathcal{B} = (e_1, \dots, e_n)$, il existe une unique forme n -linéaire alternée f telle que $f(e_1, \dots, e_n) = 1$. On l'appelle déterminant en base $\mathcal{B}$ et on la note $\det_{\mathcal{B}}$.

Remarque : Le déterminant en base $\mathcal{B} = (e_1, \dots, e_n)$ est la fonction qui permet de généraliser les notions d'aire et de volume (algébrique). Elle prend comme unité le paralléloétope construit sur les vecteurs de la base $\mathcal{B} : V = \left\{ \sum_{i=1}^n x_i e_i, (x_1, \dots, x_n) \in [0, 1]^n \right\}$.

Proposition 11.28 (Changement de base)

Soit $\mathcal{B}, \mathcal{B}'$ deux bases de E . On note P la matrice de passage de $\mathcal{B}$ à $\mathcal{B}'$.

On a $\det_{\mathcal{B}} = \det_{\mathcal{B}}(\mathcal{B}') \times \det_{\mathcal{B}'} = \det(P) \det_{\mathcal{B}'}$

Démonstration : Il suffit de voir que ce sont deux formes n -linéaires alternées qui valent la même valeur sur $\mathcal{B}'$.

□

Proposition-Définition 11.29 (Produit mixte)

Soit E un espace euclidien orienté. Soit $\mathcal{B}$ et $\mathcal{B}'$ deux bases orthonormées directes. On a $\det_{\mathcal{B}} = \det_{\mathcal{B}'}$.
 On appelle produit mixte le déterminant en base $\mathcal{B}$ où $\mathcal{B}$ est une base orthonormée directe.
 Pour $(x_1, \dots, x_n) \in E^n$ on note alors

$$[x_1, \dots, x_n] = \det_{\mathcal{B}}(x_1, \dots, x_n)$$

4 Isométries vectorielles d'un espace euclidien

Dans cette section E désigne un espace euclidien

4.1 Définition

Définition 11.30

On appelle *isométrie vectorielle* ou *automorphisme orthogonal* un endomorphisme f de E qui conserve la norme. Cela signifie qu'un endomorphisme f de E est une isométrie vectorielle si

$$\forall u \in E, \|f(u)\| = \|u\|$$

On note $O(E)$ l'ensemble des isométries vectorielles de E .

Exemples :

1. Les symétries orthogonales.

Soit s la symétrie par rapport à F et parallèlement à G . Montrons que s est une isométrie vectorielle si et seulement si $G = F^\perp$.

– $\Rightarrow$ On suppose que s est une isométrie. Soit $x \in F$ et $y \in G$, on a

$$(s(x+y)|s(x+y)) = (x-y|x-y) = \|x\|^2 - 2(x|y) + \|y\|^2$$

et

$$(x+y|x+y) = \|x\|^2 + 2(x|y) + \|y\|^2$$

Comme s est une isométrie,

$$\|s(x+y)\|^2 = \|x+y\|^2 \iff -2(x|y) = 2(x|y) \iff x \perp y$$

On obtient que $F \subset G^\perp$ et donc $F = G^\perp$.

– $\Leftarrow$ Soit $u \in E$. Il existe $(x, y) \in F \times G$ tels que $u = x + y$. Comme $G = F^\perp$, $(x|y) = 0$.

Le calcul précédent montre donc que $\|s(u)\|^2 = \|u\|^2$. La symétrie s est bien une isométrie.

En particulier les réflexions qui sont les symétries orthogonales par rapport à un hyperplan sont des isométries vectorielles.

2. Soit p la projection sur F parallèlement à G . Si $G \neq \{0\}$ (c'est-à-dire si $p \neq \text{id}$) alors p n'est pas une isométrie vectorielle. En effet pour x un vecteur non nul dans G ,

$$\|p(x)\| = 0 \neq \|x\|$$

3. Soit E un espace euclidien et $\mathcal{B} = (e_1, \dots, e_n)$ une base orthonormée. Pour toute permutation $\sigma \in \mathfrak{S}_n$, on pose f_σ l'unique endomorphisme de E tel que pour tout $i \in \llbracket 1; n \rrbracket$, $f_\sigma(e_i) = e_{\sigma(i)}$.

On en déduit que pour tout vecteur u , si on note $u = \sum_{i=1}^n x_i e_i$ alors

$$f_\sigma(u) = \sum_{i=1}^n x_i e_{\sigma(i)} = \sum_{i=1}^n x_{\sigma^{-1}(i)} e_i$$

En particulier, comme $\mathcal{B}$ est une base orthonormée,

$$\|f(u)\|^2 = \sum_{i=1}^n x_{\sigma^{-1}(i)}^2 = \sum_{i=1}^n x_i^2 = \|u\|^2$$

On voit que f_σ est une isométrie vectorielle.

ATTENTION

Par définition les isométries vectorielles sont des applications linéaires. On ne travaille pas en géométrie affine. On ne considérera pas par exemple les translations.

4.2 Propriétés

Proposition 11.31

Toute isométrie vectorielle est un automorphisme.

Démonstration : Comme E est de dimension finie, il suffit de montrer qu'il est injectif. Or si $f(u) = 0_E$ alors $\|f(u)\| = 0$ ce qui implique que $\|u\| = 0$ puis $u = 0_E$.

□

Théorème 11.32

Toute isométrie vectorielle préserve le produit scalaire. C'est-à-dire que si f est une isométrie vectorielle,

$$\forall (u, v) \in E^2, (f(u)|f(v)) = (u|v)$$

Démonstration : Il suffit d'utiliser les identités de polarisation.

On suppose que f préserve la norme. Pour tout $(u, v) \in E^2$,

$$\begin{aligned} (f(u)|f(v)) &= \frac{1}{4} (\|f(u) + f(v)\|^2 + \|f(u) - f(v)\|^2) \\ &= \frac{1}{4} (\|f(u+v)\|^2 + \|f(u-v)\|^2) \\ &= \frac{1}{4} (\|u+v\|^2 + \|u-v\|^2) \\ &= (u|v) \end{aligned}$$

□

Exercice : Soit f une application de E dans lui-même qui conserve le produit scalaire :

$$\forall (u, v) \in E^2, (f(u)|f(v)) = (u|v).$$

Montrer que l'application f est linéaire et que c'est une isométrie vectorielle.

Proposition 11.33 (Caractérisation des isométries vectorielles)

Soit $f \in \mathcal{L}(E)$.

Les assertions suivantes sont équivalentes

- i) L'endomorphisme f est une isométrie vectorielle.
- ii) Il existe une base orthonormée telle que son image par f soit encore une base orthonormée.
- iii) Pour toute base orthonormée, son image par f soit encore une base orthonormée.

Démonstration : Soit une et $f \in \mathcal{L}(E)$, On pose

- $\boxed{iii) \Rightarrow ii)}$ Évident.
- $\boxed{ii) \Rightarrow i)}$ Soit $\mathcal{B} = (e_1, \dots, e_n)$ la base orthonormée telle que, si pour tout $i \in \llbracket 1; n \rrbracket$, on pose $e'_i = f(e_i)$, alors $\mathcal{B}' = (e'_1, \dots, e'_n)$ est encore une base orthonormée.

Soit $u \in E$. Il se décompose dans la base $\mathcal{B} : u = \sum_{i=1}^n x_i e_i$. On a alors $f(u) = \sum_{i=1}^n x_i f(e_i) = \sum_{i=1}^n x_i e'_i$.

Dès lors

$$\|f(u)\|^2 = \sum_{i=1}^n x_i^2 = \|u\|^2$$

Cela montre bien que $f \in O(E)$.

- $\boxed{i) \Rightarrow iii)}$ On suppose que $f \in O(E)$. Soit $\mathcal{B} = (e_1, \dots, e_n)$ une base orthonormée. On sait que f préserve le produit scalaire donc, pour tout $(i, j) \in \llbracket 1; n \rrbracket^2$,

$$(f(e_i)|f(e_j)) = (e_i|e_j) = \delta_{ij}$$

Cela montre que la famille $(f(e_1), \dots, f(e_n))$ est une base orthonormée.

□

Corollaire 11.34

Soit $\mathcal{B}$ une base orthonormée. Soit $f \in \mathcal{L}(E)$ on a

$$f \in O(E) \iff \text{Mat}_{\mathcal{B}}(f) \in O_n(\mathbf{R})$$

Démonstration : On note $\mathcal{B} = (e_1, \dots, e_n)$ une base orthonormée. Notons pour tout $i \in \llbracket 1; n \rrbracket$, $e'_i = f(e_i)$ et $C_i = \text{Mat}_{\mathcal{B}}(e'_i)$. Notons pour finir $A = \text{Mat}_{\mathcal{B}}(f)$ de sorte que $C_1, \dots, C_n$ soient les colonnes de A .

- $f \in O(E) \iff$ la famille $(e'_1, \dots, e'_n)$ est une base orthonormée
- $\iff \forall (i, j) \in \llbracket 1; n \rrbracket^2, (e'_i|e'_j) = \delta_{ij}$
- $\iff \forall (i, j) \in \llbracket 1; n \rrbracket^2, \langle C_i|C_j \rangle = \delta_{ij}$ car $(e_i|e_j) = \langle C_i|C_j \rangle$ puisque $\mathcal{B}$ est orthonormée
- $\iff \text{Mat}_{\mathcal{B}}(f) \in O(n)$

□

Proposition 11.35

Soit $f \in \mathcal{L}(E)$. L'endomorphisme f est une isométrie vectorielle si et seulement si $f^* = f^{-1}$.

Démonstration : Soit $\mathcal{B}$ une base orthonormée, notons $A = \text{Mat}_{\mathcal{B}}(f)$ sa matrice dans la base $\mathcal{B}$

$$f \in O(E) \iff A \in O_n(\mathbf{R}) \iff A^T = A^{-1} \iff \text{Mat}_{\mathcal{B}}(f^*) = \text{Mat}_{\mathcal{B}}(f^{-1}) \iff f^* = f^{-1}$$

□

Remarque : On peut aussi faire la démonstration directement sans avoir recours à une base.

Proposition 11.36

L'ensemble $O(E)$ des isométries vectorielles est un sous-groupe (multiplicatif) du groupe linéaire $(\text{GL}(E), \circ)$.

En particulier, la composition de deux isométries vectorielles est une isométrie vectorielle et l'inverse d'une isométrie vectorielle est une isométrie vectorielle.

Démonstration : On utilise la caractérisation usuelle des sous-groupes.

- L'identité est une isométrie vectorielle donc $O(E) \neq \emptyset$.
- Soit f et g deux isométries vectorielles, il faut montrer que $f \circ g^{-1}$ est aussi une isométrie vectorielle.

Soit u dans E , comme g est un automorphisme, on peut poser $v = g^{-1}(u)$ (et donc $u = g(v)$).

$$\|u\| = \|g(v)\| = \|v\| = \|f(v)\| = \|(f \circ g^{-1})(u)\|$$

Cela montre que $f \circ g^{-1}$ appartient à $O(E)$.

On a bien montré que $(O(E), \circ)$ est un sous-groupe de $(\text{GL}(E), \circ)$.

□

4.3 Isométries directes et indirectes

Proposition 11.37

Soit $f \in O(E)$. On a

$$\det(f) \in \{\pm 1\}$$

Démonstration : Soit $\mathcal{B}$ une base orthonormée. On a vu que $A = \text{Mat}_{\mathcal{B}}(f) \in O_n(\mathbf{R})$. On en déduit que

$$\det(f) = \det(A) \in \{\pm 1\}$$

□

Définition 11.38

Soit f une isométrie vectorielle. On dit que c'est une isométrie directe si $\det(f) = 1$ et que c'est une isométrie indirecte si $\det(f) = -1$.

On note $SO(E)$ l'ensemble des isométries vectorielles directes de E .

Exemple : Soit F un sous-espace vectoriel de E et $G = F^\perp$. On considère la symétrie orthogonale par rapport à F (et donc parallèlement à G). Soit $\mathcal{B}$ une base orthonormée de E construite en concaténant une base orthonormée de F avec une base orthonormée de G . On a

$$\text{Mat}_{\mathcal{B}}(f) = \begin{pmatrix} 1 & 0 & \cdots & \cdots & \cdots & 0 \\ 0 & \ddots & \ddots & & & \vdots \\ \vdots & \ddots & 1 & \ddots & & \vdots \\ \vdots & & \ddots & -1 & \ddots & \vdots \\ \vdots & & & \ddots & \ddots & 0 \\ 0 & \cdots & \cdots & \cdots & 0 & -1 \end{pmatrix}$$

Si on note $d = \dim(G)$ la codimension de F , on a donc $\det(f) = (-1)^d$.

En particulier, si f est une réflexion et donc $d = 1$, f est une isométrie indirecte.

Proposition 11.39

L'ensemble $SO(E)$ des isométries vectorielles directes est un sous-groupe (multiplicatif) du groupe linéaire $(GL(E), \circ)$.

On dit que $SO(E)$ est le groupe spécial orthogonal.

Démonstration : Il suffit de voir que c'est le noyau du morphisme de groupe

$$\begin{aligned} \det : O(E) &\rightarrow \{\pm 1\} \\ f &\mapsto \det(f) \end{aligned}$$

□

Proposition 11.40

On suppose que l'espace vectoriel E est orienté. Un endomorphisme f appartient à $SO(E)$ si et seulement s'il transforme une base orthonormée directe en une base orthonormée directe.

Démonstration :

- $\Rightarrow$ On suppose que f est une isométrie vectorielle directe. Soit $\mathcal{B}$ une base orthonormée directe. On note $\mathcal{B}'$ son image par f . On sait que dans ce cas

$$\det(\mathcal{B}') = \det(f) \times \det(\mathcal{B}) = 1$$

Ici, on note $\det$ le produit mixte, c'est-à-dire le déterminant dans une base orthonormée directe quelconque (par exemple $\mathcal{B}$).

On a bien montré que $\mathcal{B}'$ était une base orthonormée directe.

- $\Leftarrow$ On suppose qu'il existe base orthonormée directe $\mathcal{B}$ telle que son image par f est une base orthonormée directe $\mathcal{B}'$. Le même calcul que ci-dessus donne que $\det(f) = 1$ et donc $f \in SO(E)$.

□

4.4 Isométries vectorielles du plan euclidien

Dans tout ce paragraphe, E désigne un plan euclidien orienté c'est-à-dire un espace euclidien de dimension 2 orienté.

Théorème 11.41 (Description de $O_2(\mathbf{R})$ et de $SO_2(\mathbf{R})$)

Soit $\theta \in \mathbf{R}$. On pose

$$R(\theta) = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix} \text{ et } S(\theta) = \begin{pmatrix} \cos \theta & \sin \theta \\ \sin \theta & -\cos \theta \end{pmatrix}$$

On a

$$SO_2(\mathbf{R}) = \{R(\theta), \theta \in \mathbf{R}\} \text{ et } O_2(\mathbf{R}) = \{R(\theta), \theta \in \mathbf{R}\} \cup \{S(\theta), \theta \in \mathbf{R}\}$$

Démonstration : La démonstration a été fait précédemment dans le paragraphe sur l'orientation.

□

Corollaire 11.42

1. L'application de $(\mathbf{R}, +)$ dans $(SO_2(\mathbf{R}), \times)$ qui associe à θ la matrice $R(\theta)$ est un morphisme surjectif de groupe dont le noyau est $2\pi\mathbf{Z}$.
2. Le groupe $(SO_2(\mathbf{R}), \times)$ est isomorphe à $\mathbf{U} = \{z \in \mathbf{C}, |z| = 1\}$. En particulier, c' est un groupe abélien.

Démonstration :

1. On peut vérifier par le calcul que pour θ, θ' dans $\mathbf{R}$, $R(\theta + \theta') = R(\theta) \times R(\theta')$. Comme de plus $R(0) = I_2$, cela montre que $\theta \mapsto R(\theta)$ est bien un morphisme de groupe de $(\mathbf{R}, +)$ dans $(SO_2(\mathbf{R}), \times)$.

Il est surjectif d'après le théorème précédent. Pour finir, pour tout $\theta \in \mathbf{R}$, θ appartient au noyau si et seulement si $\cos(\theta) = 1$ et $\sin(\theta) = 0$ ce qui est équivalent au fait que θ appartienne à $2\pi\mathbf{Z}$.

2. Il suffit de voir qu'un élément z de $\mathbf{U}$ est uniquement déterminé par son argument (à 2π -près) et que de plus $\arg(zz') \equiv \arg(z) + \arg(z') [2\pi]$.

□

Proposition 11.43

Soit f un endomorphisme de $SO(E)$. Il existe un unique réel θ à 2π -près tel que dans toute base orthonormée directe sa matrice soit $R(\theta)$.

Démonstration : Choisissons une base $\mathcal{B}$ orthonormée directe.

On sait que la matrice de f dans la base $\mathcal{B}$ appartient à $SO_2(\mathbb{R})$. Elle est donc de la forme $R(\theta)$ où θ est déterminé de manière unique à 2π -près

Maintenant si $\mathcal{B}'$ est une autre base orthonormée directe alors la matrice de passage P de $\mathcal{B}$ à $\mathcal{B}'$ est aussi un élément de $SO_2(\mathbb{R})$. Comme $SO_2(\mathbb{R})$ est abélien,

$$\text{Mat}_{\mathcal{B}'}(f) = P^{-1}R(\theta)P = R(\theta)P^{-1}P = R(\theta).$$

□

Définition 11.44

Avec les notations précédentes, on dit que f est une rotation d'angle θ ou que θ est une mesure de f .

Remarque : Cela permet de définir la notion de la mesure de l'angle orienté entre deux vecteurs non nuls. Si u et v sont deux vecteurs non nuls, on peut considérer les vecteurs unitaires associés $u' = \frac{u}{\|u\|}$ et $v' = \frac{v}{\|v\|}$. On peut alors voir qu'il existe une unique rotation f qui envoie u' sur v' . Si f est la rotation d'angle θ , on dit que θ est la mesure de l'angle orienté $\widehat{(u, v)}$.

Proposition 11.45

Soit f une isométrie vectorielle indirecte. C'est une réflexion.

Démonstration : On choisit une base orthogonale directe $\mathcal{B}$. On voit que la matrice de f dans la base $\mathcal{B}$ est de la forme $S(\theta)$. On en déduit par le calcul que $f^2 = \text{id}$ car $S(\theta)^2 = I_2$.

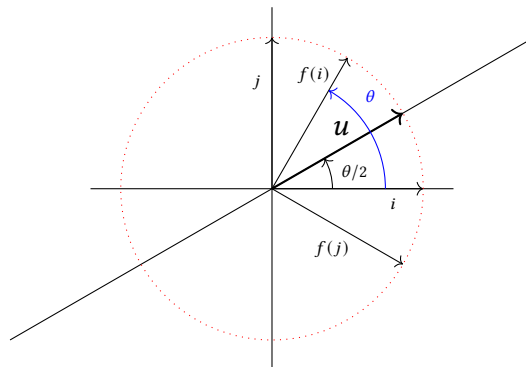
Donc f est une symétrie. C'est de plus une symétrie orthogonale car c'est une isométrie vectorielle. On regarde le noyau de $f - \text{id}$. Comme f n'est pas id ni $-\text{id}$ (qui sont des rotations) on en déduit que le noyau de $f - \text{id}$ est de dimension 1. C'est bien une réflexion.

□

Remarque : Notons $\mathcal{B} = (i, j)$ une base orthonormée directe. On note $S(\theta)$ sa matrice. Si on cherche une équation de $\text{Ker}(f - \text{id})$ il suffit de résoudre :

$$\begin{cases} ((\cos \theta) - 1)x + (\sin \theta)y = 0 \\ (\sin \theta)x + (-(\cos \theta) - 1)y = 0 \end{cases} \iff \begin{cases} \sin(\theta/2) (-\sin(\theta/2)x + \cos(\theta/2)y) = 0 \\ \cos(\theta/2) (\sin(\theta/2)x - \cos(\theta/2)y) = 0 \end{cases}$$

Le vecteur $u = \cos\left(\frac{\theta}{2}\right)i + \sin\left(\frac{\theta}{2}\right)j$ est une solution. On en déduit que f est une symétrie par rapport à $\text{Vect}(u)$.



ATTENTION

Soit f une isométrie vectorielle du plan euclidien.

- Si c'est une isométrie directe, c'est une rotation et il existe θ tel que, pour toute base orthonormée $\mathcal{B}$,

$$\text{Mat}_{\mathcal{B}}(f) = R(\theta)$$

- Si elle n'est pas directe, c'est une réflexion. Pour toute base orthonormée $\mathcal{B}$, il existe θ tel que

$$\text{Mat}_{\mathcal{B}}(f) = S(\theta)$$

Dans ce cas, θ dépend de la base $\mathcal{B}$. En particulier, on peut choisir $\mathcal{B}$ tel que

$$\text{Mat}_{\mathcal{B}}(f) = S(0) = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$$

Remarque : On peut faire ces calculs en identifiant un vecteur (x, y) avec le nombre complexe $z = x + iy$. La matrice $R(\theta)$ correspond alors à la multiplication par $e^{i\theta}$. La matrice $S(\theta)$ à $z \mapsto e^{i\theta}\bar{z}$.

En effet, si on note $X = \begin{pmatrix} x \\ y \end{pmatrix}$.

$$R(\theta)X = \begin{pmatrix} x \cos \theta - y \sin \theta \\ x \sin \theta + y \cos \theta \end{pmatrix} \quad \text{et} \quad S(\theta)X = \begin{pmatrix} x \cos \theta + y \sin \theta \\ x \sin \theta - y \cos \theta \end{pmatrix}$$

On vérifie aisément alors que

$$e^{i\theta}z = (x \cos \theta - y \sin \theta) + i(x \sin \theta + y \cos \theta) \quad \text{et} \quad e^{i\theta}\bar{z} = (x \cos \theta + y \sin \theta) + i(x \sin \theta - y \cos \theta)$$

4.5 Réduction des isométries vectorielles

On considère maintenant un espace euclidien orienté de dimension quelconque. Soit f une isométrie vectorielle, on va chercher une base $\mathcal{B}$ orthogonale, telle que $\text{Mat}_{\mathcal{B}}(f)$ soit la plus simple possible. Remarquons pour commencer que f n'est pas nécessairement diagonalisable (dans $\mathbf{R}$). Par exemple, la rotation d'angle $\frac{\pi}{2}$ dans $\mathbf{R}^2$ a pour matrice dans la base canonique

$$M = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}.$$

Cette matrice a pour polynôme caractéristique $\chi = X^2 + 1$ qui n'a pas de racines réelles.

Afin de construire la base $\mathcal{B}$, la méthode est la suivante :

- Trouver un sous-espace vectoriel F de dimension 1 ou 2 stable par f ; c'est le lemme A. On saura décrire l'endomorphisme induit par f sur F car on connaît les isométries vectorielles sur les espaces vectoriels de dimension 1 ou 2.
- Comme on cherche à construire une base orthonormée, on s'intéresse alors à $F^\perp$. On va montrer que $F^\perp$ est encore stable par f ; c'est le lemme B.
- On peut alors procéder par récurrence sur $\dim E$.

Lemme 11.46 (Lemme A)

Soit f un endomorphisme (quelconque) de E . Si on suppose que $E \neq \{0\}$ alors il existe une droite ou un plan F stable par f .

Remarque : Ce lemme est valable en général pour tout endomorphisme d'un $\mathbf{R}$ -espace vectoriel.

Démonstration :

- Première preuve : Soit A la matrice de f dans une base $\mathcal{B}$ fixée de E . On sait que A admet une valeur propre (éventuellement complexe). C'est-à-dire qu'il existe $\lambda \in \mathbf{C}$ et $X \in \mathcal{M}_{n,1}(\mathbf{C})$ non nul tel que $AX = \lambda X$.
 - Si X est à coefficients réels, alors λ est aussi réel. Soit x tel que $\text{Mat}_{\mathcal{B}}(x) = X$, le vecteur x est un vecteur propre et donc $F = \text{Vect}(x)$ est stable par f .
 - Si X n'est pas à coefficients réels, on pose $X = U + iV$ avec $(U, V) \in \mathcal{M}_{n,1}(\mathbf{R})^2$. Posons aussi $\lambda = \alpha + i\beta$ avec α et β réels. On a alors

$$AX = \lambda X \Rightarrow A(U + iV) = (\alpha + i\beta)(U + iV) \Rightarrow \begin{cases} AU &= \alpha U - \beta V \\ AV &= \beta U + \alpha V \end{cases}$$

Si on note alors u et v les vecteurs de E tels que $\text{Mat}_{\mathcal{B}}(u) = U$ et $\text{Mat}_{\mathcal{B}}(v) = V$. En posant $F = \text{Vect}(u, v)$ on voit que F est stable par f .

De plus, comme v n'est pas nul, F est une droite ou un plan.

Dans tous les cas, il existe une droite ou un plan F stable par f .

- Deuxième preuve : On considère un polynôme P non nul annulateur de f dans $\mathbf{R}[X]$ (que l'on peut choisir unitaire). On peut prendre $P = \mu_f$ ou $P = \chi_f$.

Ce polynôme se factorise dans $\mathbf{R}[X]$ comme un produit de polynômes irréductibles (que l'on peut choisir unitaires) : $P = \prod_{i=1}^r P_i$. Par définition, $P(f) = P_1(f) \circ P_2(f) \circ \cdots \circ P_r(f) = 0_{\mathcal{L}(E)}$.

En particulier, il existe au moins un entier $i \in \llbracket 1 ; r \rrbracket$ tel que $P_i(f)$ ne soit pas injectif.

Par symétrie, supposons que $P_1(f)$ n'est pas injectif et donc il existe un vecteur x non nul dans $\text{Ker}(P_1(f))$. On sait de plus que P_1 est un polynôme irréductible de $\mathbf{R}[X]$

- Si P_1 est de degré 1. On a $P_1 = X - \lambda$. On a alors que x est un vecteur propre associé à λ car $(f - \lambda \text{id})(x) = 0$. On peut poser $F = \text{Vect}(x)$ qui est stable par f .
- Si P_1 est de degré 2, on a $P_2 = X^2 + \alpha X + \beta$. On pose alors $F = \text{Vect}(x, f(x))$ qui est de dimension 1 ou 2. Comme de plus, $f^2(x) = -\alpha f(x) - \beta x$, on en déduit que F est stable par f .

Dans tous les cas, il existe une droite ou un plan F stable par f .

□

Lemme 11.47 (Lemme B)

Soit f une isométrie vectorielle et F un sous-espace stable par f alors $F^\perp$ est aussi stable par f .

Démonstration : Soit $x \in F^\perp$. On veut montrer que $f(x) \in F^\perp$, c'est-à-dire que pour tout $y \in F$, $(f(x)|y) = 0$.

Or f est un isomorphisme car c'est une isométrie vectorielle. Comme $f(F) \subset F$ et que $\dim(f(F)) = \dim(F)$, on a $f(F) = F$. En particulier, pour $y \in F$, il existe $z \in F$ tel que $y = f(z)$. On en déduit que

$$(f(x)|y) = (f(x)|f(z)) = (x|z) = 0$$

On a bien montré que $F^\perp$ était stable par f .

□

Théorème 11.48 (Réduction des isométries - Version endomorphismes)

Soit f une isométrie vectorielle de E . Il existe une base orthonormée $\mathcal{B}$ de E tel que la matrice de f dans la base $\mathcal{B}$ soit diagonale par blocs avec des blocs de la forme

$$R_\theta = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix} \in SO_2(\mathbf{R}); (1) \in O_1(\mathbf{R}); (-1) \in O_1(\mathbf{R}).$$

C'est-à-dire, en regroupant les blocs,

$$\text{Mat}_{\mathcal{B}}(f) = \begin{pmatrix} R_{\theta_1} & & & & \\ & \ddots & & & \\ & & R_{\theta_r} & & \\ & & & I_p & \\ & & & & -I_q \end{pmatrix}$$

Démonstration : On procède par récurrence sur la dimension de E .

- Initialisation : Si $\dim E = 1$, les seules isométries sont id et $-\text{id}$.
- Hérédité : Soit $n \geq 2$. On suppose que la propriété est vraie pour tout les entier $k < n$. Soit E de dimension n . On sait qu'il existe un sous-espace vectoriel F stable par f de dimension 1 ou 2. On sait que $F^\perp$ est aussi stable par f et que $f|_{F^\perp}$ la restriction de f à $F^\perp$ est une isométrie. On peut donc lui appliquer l'hypothèse de récurrence. Maintenant,
 - Si $\dim F = 1$ alors la restriction de f à F qui est une isométrie est id ou $-\text{id}$.
 - Si $\dim F = 2$ alors la restriction de f à F qui est une isométrie vectorielle est soit une rotation de matrice $R_\theta = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}$, soit une symétrie et on peut choisir une base orthonormée telle que sa matrice soit $S_0 = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$.

Il n'y a plus qu'à concaténer les bases.

□

Corollaire 11.49 (Réduction des isométries - Version matrices)

Soit A une matrice orthogonale. Elle est semblable à une matrice D diagonale par blocs avec des blocs de la forme

$$R_\theta = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix} \in SO_2(\mathbf{R}); (1) \in O_1(\mathbf{R}); (-1) \in O_1(\mathbf{R}).$$

C'est-à-dire, en regroupant les blocs,

$$\text{Mat}_{\mathcal{B}}(f) = \begin{pmatrix} R_{\theta_1} & & & & \\ & \ddots & & & \\ & & R_{\theta_r} & & \\ & & & I_p & \\ & & & & -I_q \end{pmatrix}.$$

De plus, cela peut-être obtenu par un changement de bases entre deux bases orthonormales. C'est-à-dire qu'il existe $P \in O_n(\mathbf{R})$ telle que

$$P^{-1}AP = P^T AP = D.$$

4.6 Isométries vectorielles de l'espace

On va appliquer ce qui précède pour $n = 3$.

Par la suite E désigne un espace euclidien orienté de dimension 3.

Proposition 11.50

Soit w un vecteur unitaire de E . Il existe une unique orientation du plan $H = \text{Vect}(w)^\perp$ tel que pour toute base orthonormée directe (u, v) de H , la base (u, v, w) soit directe dans E .

Remarque : On dit alors que H est orienté par le vecteur w . On peut d'ailleurs procéder dans l'autre sens. Si on se donne un plan H , pour l'orienter, il suffit de choisir un des deux vecteur unitaires de $H^\perp$. On va essayer de décrire les éléments de $SO(E)$ et de $SO_3(\mathbf{R})$.

Corollaire 11.51

Soit $A \in SO_3(\mathbf{R})$. Comme $\det A = 1$, elle est semblable a une matrice de la forme

$$\begin{pmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

Démonstration : Il suffit d'appliquer la réduction des isométries.

- S'il n'y a que des blocs de taille 1, on est de la forme voulue en posant $\theta = 0$ ou $\theta = \pi$.
- Sinon, on a un bloc de taille 2 et un bloc de taille 1. Le bloc de taille 1 étant (1) car $\det(u) = 1$.

□

Définition 11.52

Soit w un vecteur non nul de E et θ un réel. On appelle rotation d'axe autour de w et d'angle θ , l'isométrie vectorielle laissant stable $F = \text{Vect}(w)$ et tel que $\tilde{f}$ la restriction de f à $H = F^\perp$ soit la rotation d'angle θ . Si $\mathcal{B} = (u, v, w')$ est une base orthonormée directe avec $w' = w/\|w\|$, la matrice de f dans la base $\mathcal{B}$ est

$$\text{Mat}_{\mathcal{B}}(f) = \begin{pmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

Remarques :

1. La droite vectorielle F s'appelle l'axe de f .
2. Si l'angle vaut π et donc la matrice est $\begin{pmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{pmatrix}$, on dit que c'est un demi-tour.
3. Il faut faire attention que si on change w en $-w$ la rotation change car le plan H est alors orienté dans l'autre sens.
4. La plupart du temps on prendra w unitaire.
5. On voit donc que dans les isométries vectorielles de l'espace il y a les rotations (qui sont les éléments de $SO(E)$). Il y a aussi les réflexions. Mais, contrairement au cas du plan, il y en a d'autre. Par exemple l'application $u = -\text{id}_E$ est un élément de $O(E)$ mais pas de $SO(E)$ car $\det(u) = (-1)^3 = -1$. Mais ce n'est pas une réflexion car aucun vecteur n'est invariant. De fait, c'est la composée du demi-tour de matrice $\begin{pmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{pmatrix}$ et de la réflexion de matrice $\begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -1 \end{pmatrix}$.

★ **Méthode :** Il faut savoir déterminer la matrice d'une rotation dont on connaît w et θ et inversement, connaissant la matrice il faut savoir retrouver w et θ .

— Détermination de la matrice :

Proposition 11.53

Soit w un vecteur unitaire, θ un réel et f la rotation d'angle θ autour de w . Pour tout vecteur u orthogonal à w , on a

$$f(u) = (\cos \theta)u + (\sin \theta)(w \wedge u).$$

Démonstration : Il suffit de voir que, comme $(u, w, u \wedge w)$ est une base orthonormée directe alors $(u, w \wedge u, w)$ aussi. De ce fait, la matrice de f dans cette base est

$$\text{Mat}_{\mathcal{B}}(f) = \begin{pmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

Le résultat en découle en regardant la première colonne. □

Soit w le vecteur $w = (1, 2, 0)$. On cherche la matrice dans la base canonique de $\mathbb{R}^3$ de la rotation autour de w et d'angle θ . Pour tout vecteur $u = (x, y, z)$ on peut le décomposer comme somme d'un vecteur colinéaire à w et d'un vecteur qui lui est orthogonal (car F et $F^\perp$ sont en somme directe si $F = \text{Vect}(w)$). On a

$$u = \frac{x+2y}{5}w + \frac{1}{5}(4x-2y, -2x-3y, 5z).$$

Maintenant $w \wedge \frac{1}{5}(4x-2y, -2x-3y, 5z) = \frac{1}{5}(10z, -5z, -10x+y)$. On a donc

$$f(u) = \frac{x+2y}{5}(1, 2, 0) + \frac{\cos \theta}{5}(4x-2y, -2x-3y, 5z) + \frac{\sin \theta}{5}(10z, -5z, -10x+y).$$

On peut alors en déduire la matrice ...

– Détermination de l'axe et de l'angle : A l'inverse soit $M = \frac{1}{9} \begin{pmatrix} 8 & 1 & -4 \\ -4 & 4 & -7 \\ 1 & 8 & 4 \end{pmatrix}$. On vérifie

que M est bien une matrice orthogonale et que son déterminant vaut 1. L'endomorphisme f canoniquement associé est donc une rotation de $\mathbb{R}^3$.

Pour déterminer l'axe il suffit de chercher les invariants c'est-à-dire le noyau de $f - \text{id}$. On trouve $w = \frac{1}{\sqrt{11}}(3, -1, -1)$ en le prenant normé.

Maintenant on prend un vecteur orthogonal à w et normé. On prend $u = \frac{1}{\sqrt{10}}(1, 3, 0)$. On a

$$\cos \theta = (u, f(u)) = \frac{7}{18} \text{ et } \sin \theta = (f(u), w \wedge u) = \text{Det}(w, u, f(u)) = \frac{5\sqrt{11}}{18}.$$

On peut aussi remarquer que $\text{tr}(f) = 1 + 2 \cos \theta$. On retrouve bien $\cos \theta = \frac{\text{tr}(f) - 1}{2} = \frac{7}{18}$.

On peut alors déterminer le signe de θ en remarquant que $\sin \theta$ est du signe de $[u, f(u), w]$ pour $u \notin \text{Vect}(w)$. En effet, en posant $u = \alpha + kw$ on a

$$[\alpha + kw, f(\alpha) + kw, w] = [\alpha, f(\alpha), w] = \sin \theta [\alpha, w \wedge \alpha, w].$$

Exercice : Soit $A = \frac{1}{4} \begin{pmatrix} 3 & 1 & \sqrt{6} \\ 1 & 3 & -\sqrt{6} \\ -\sqrt{6} & \sqrt{6} & 2 \end{pmatrix}$. Justifier que $A \in SO(3, \mathbb{R})$. Déterminer son axe et son angle.

On trouve $w = (1, 1, 0)$ puis $\text{tr}(A) = 2$ d'où $\cos \theta = \frac{1}{2}$ et donc $\theta = \pm \frac{\pi}{3}$. En prenant $u = (1, 0, 0)$,

$$\begin{vmatrix} 1 & -1 & 1 \\ 0 & 1 & 1 \\ 0 & -\sqrt{6} & 0 \end{vmatrix} = \sqrt{6} > 0$$

donc $\theta = \frac{\pi}{3}$.

Probabilités II

1	Espérance d'une variable aléatoire discrète	325
1.1	Cas des variables positives	326
1.2	Cas des variables réelles ou complexes	330
1.3	Propriétés de l'espérance	331
2	Variance et écart type d'une variable aléatoire réelle	334
2.1	Variables aléatoires réelle de carré sommable	334
2.2	Lois usuelles	337
2.3	Covariance	338
3	Inégalités probabilistes et loi des grands nombres	341
3.1	Inégalités	341
3.2	Loi faible des grands nombres	344
4	Fonctions génératrices	345
4.1	Généralités	345
4.2	Exemples	348

1 Espérance d'une variable aléatoire discrète

Dans ce paragraphe, on se fixe un espace probabilisé $(\Omega, \mathcal{A}, \mathbf{P})$. On ne considère que des variables aléatoires discrètes sur $(\Omega, \mathcal{A}, \mathbf{P})$ à valeurs dans une partie de $\mathbf{R}$.

1.1 Cas des variables positives

Définition 12.1 (Espérance)

Soit X une variable aléatoire discrète à valeurs dans $[0, +\infty]$.

On appelle espérance de X et on note $E(X)$ la somme de la famille $(xP(X = x))_{x \in X(\Omega)}$. On a donc $E(X) \in [0, +\infty]$.

Si cette somme est finie (c'est-à-dire si la famille $(xP(X = x))_{x \in X(\Omega)}$ est sommable), on dit que X est une variable aléatoire d'espérance finie.

On note $X \in L^1$ ou $X \in L^1(\Omega, [0, +\infty])$ pour dire que X est une variable aléatoire d'espérance finie.

Remarques :

1. On pourrait définir $L^1(\Omega, [0, +\infty])$ comme étant l'ensemble des variables aléatoires définies sur $(\Omega, \mathcal{A}, P)$ à valeurs dans $[0, +\infty]$ d'espérance finie. Cependant, on aimerait pouvoir « identifier » deux variables aléatoires étant égales presque sûrement. Pour cela il faudrait une notion d'ensembles quotients qui ne sont pas au programme de CPGE. On se contentera donc de voir $X \in L^1$ comme une notation.
2. Cette année on oubliera la formule

$$E(X) = \sum_{\omega \in \Omega} X(\omega)P(\{\omega\})$$

En effet dans la majorité des cas, Ω ne sera pas dénombrable et $P(\{\omega\})$ sera toujours nul. Pour utiliser une définition de ce genre, il faudrait développer une notion d'intégrale.

Exemple : On considère une variable aléatoire X qui suit la loi $\mathcal{G}(\frac{1}{2})$. On regarde $Y = 2^X$. C'est-à-dire que l'on lance une pièce équilibrée jusqu'à obtenir pile. Si on obtient pile au lancer n (pour la première fois) on gagne 2^n euros. Si on cherche l'espérance de cette variable on trouve

$$E(Y) = \sum_{p=0}^{+\infty} pP(Y = p) = \sum_{n=1}^{+\infty} 2^n P(Y = 2^n) = \sum_{n=1}^{+\infty} 2^n P(X = n) = \sum_{n=1}^{+\infty} 2^n \left(\frac{1}{2}\right)^n = +\infty.$$

Proposition 12.2

Soit X une variable aléatoire discrète à valeurs dans $[[0, +\infty]] = \mathbb{N} \cup \{+\infty\}$. On a dans $[[0, +\infty]]$,

$$E(X) = \sum_{n=1}^{\infty} P(X \geq n)$$

Remarque : Soit X à valeurs entières. Pour tout $n \in \mathbb{N}^*$, $P(X \geq n) = P(X > n - 1)$. Dans le cas où la variable ne prend que des valeurs finies, on a donc

$$E(X) = \sum_{n=1}^{\infty} P(X \geq n) = \sum_{n=1}^{\infty} P(X > n - 1) = \sum_{n=0}^{\infty} P(X > n)$$

Démonstration : Si $P(X = +\infty) = \alpha \neq 0$ alors $E(X) = +\infty$ et, pour tout $n \in \mathbb{N}^*$, $P(X \geq n) \geq \alpha > 0$. Donc $\sum_{n=1}^{\infty} P(X \geq n) = +\infty$.

Supposons maintenant que $P(X = +\infty) = 0$.

On pose pour tout $(n, p) \in (\mathbb{N}^*)^2$,

$$u_{n,p} = \begin{cases} \mathbf{P}(X = n) & \text{si } n \geq p \\ 0 & \text{sinon} \end{cases}$$

n				
$\vdots$				
3	$\mathbf{P}(X = 3)$	$\mathbf{P}(X = 3)$	$\mathbf{P}(X = 3)$	0
2	$\mathbf{P}(X = 2)$	$\mathbf{P}(X = 2)$	0	0
1	$\mathbf{P}(X = 1)$	0	0	0
1	2			p

D'après le théorème de Fubini pour les familles à valeurs positives,

$$\sum_{n \in \mathbb{N}^*} \left(\sum_{p \in \mathbb{N}^*} u_{n,p} \right) = \sum_{p \in \mathbb{N}^*} \left(\sum_{n \in \mathbb{N}^*} u_{n,p} \right)$$

Or pour tout $n \geq 1$, $\sum_{p \in \mathbb{N}^*} u_{n,p} = n\mathbf{P}(X = n)$ et pour tout $p \geq 1$, $\sum_{n \in \mathbb{N}^*} u_{n,p} = \mathbf{P}(X \geq p)$. On obtient donc

$$\mathbf{E}(X) = \sum_{n \in \mathbb{N}^*} n\mathbf{P}(X = n) = \sum_{p \in \mathbb{N}^*} \mathbf{P}(X \geq p)$$

□

Théorème 12.3 (Espérance des lois usuelles)

Soit X une variable aléatoire discrète à valeurs réelles.

1. (Loi de Bernoulli) : Soit $p \in [0, 1]$, si $X \sim \mathcal{B}(p)$ alors $\mathbf{E}(X) = p$
2. (Loi binomiale) : Soit $p \in [0, 1]$ et $n \in \mathbb{N}^*$, si $X \sim \mathcal{B}(n, p)$ alors $\mathbf{E}(X) = np$
3. (Loi géométrique) : Soit $p \in]0, 1]$, si $X \sim \mathcal{G}(p)$ alors $\mathbf{E}(X) = \frac{1}{p}$.
4. (Loi de Poisson) : Soit $\lambda \in \mathbf{R}_+^*$, si $X \sim \mathcal{P}(\lambda)$ alors $\mathbf{E}(X) = \lambda$

Remarques :

1. Dans le cas où $p = 0$, l'espérance de $\mathcal{G}(p)$ est $+\infty$.
2. Il est facile de retenir l'espérance de la loi géométrique. En effet, si on lance un dé équilibré à 10 faces une infinité de fois. Il paraît clair que le temps moyen avant le premier 10 est $10 = \frac{1}{\frac{1}{10}}$.

Démonstration :

1. Voir le cours de première année
2. Voir le cours de première année
3. On utilise le résultat précédent

$$\mathbf{E}(X) = \sum_{n=0}^{+\infty} \mathbf{P}(X > n) = \sum_{n=0}^{+\infty} q^n = \frac{1}{1-q} = \frac{1}{p}$$

4. On considère la série $\sum_{k \in \mathbb{N}} k \frac{\lambda^k e^{-\lambda}}{k!}$.

Pour tout $N \in \mathbb{N}$,

$$\sum_{k=0}^N k \frac{\lambda^k e^{-\lambda}}{k!} = \lambda \sum_{k=1}^N \frac{\lambda^{k-1} e^{-\lambda}}{(k-1)!} = \lambda \sum_{k=0}^N \frac{\lambda^k e^{-\lambda}}{k!} \xrightarrow{N \rightarrow \infty} \lambda.$$

On a donc $E(X) = \lambda$.

□

Théorème 12.4 (Formule de transfert pour les variables positives)

Soit X une variable aléatoire discrète et f une fonction définie sur $X(\Omega)$ et à valeurs dans $\mathbb{R}_+$. On a dans $[0, +\infty]$

$$E(f(X)) = \sum_{x \in X(\Omega)} f(x) P(X = x)$$

En particulier, $f(X)$ est d'espérance finie si et seulement si la famille $(f(x)P(X = x))_{x \in X(\Omega)}$ est sommable.

Démonstration : Notons $Y = f(X)$. Tous les termes étudiés seront positifs. On peut donc faire les calculs dans $[0, +\infty]$

$$\begin{aligned} \sum_{y \in Y(\Omega)} y P(Y = y) &= \sum_{y \in Y(\Omega)} y \sum_{x \in X(\Omega)} P((X, Y) = (x, y)) \\ &= \sum_{(x, y) \in X(\Omega) \times Y(\Omega)} y P((X, Y) = (x, y)) \\ &= \sum_{\substack{(x, y) \in X(\Omega) \times Y(\Omega) \\ y = f(x)}} y P((X, Y) = (x, y)) \\ &= \sum_{\substack{(x, y) \in X(\Omega) \times Y(\Omega) \\ y = f(x)}} f(x) P((X, Y) = (x, y)) \\ &= \sum_{x \in X(\Omega)} f(x) \underbrace{\sum_{\substack{y \in Y(\Omega) \\ y = f(x)}} P((X, Y) = (x, y))}_{P(X=x)} \\ &= \sum_{x \in X(\Omega)} f(x) P(X = x) \end{aligned}$$

□

Exemple : Soit $X \sim \mathcal{G}(p)$ avec $p \in]0, 1[$. On cherche à calculer $E(X^2)$. D'après le théorème

$$E(X^2) = \sum_{k \in \mathbb{N}^*} k^2 (1-p)^{k-1} p$$

On considère la série entière $\sum_{k \geq 0} x^k$. On sait que son rayon de convergence vaut 1 et sa somme est

$$S : x \mapsto \sum_{k=0}^{+\infty} x^k = \frac{1}{1-x}$$

En dérivant deux fois sur $] - 1, 1[$,

$$S''(x) : x \mapsto \sum_{k=0}^{+\infty} k(k-1)x^k = \frac{2}{(1-x)^3}$$

Alors

$$\sum_{k=0}^{+\infty} k^2 x^{k-1} = \sum_{k=0}^{+\infty} k(k-1)x^{k-1} + \sum_{k=0}^{+\infty} kx^{k-1} = x \sum_{k=0}^{+\infty} k(k-1)x^{k-2} + \sum_{k=0}^{+\infty} kx^{k-1} = \frac{2x}{(1-x)^3} + \frac{1}{(1-x)^2}$$

Finalement,

$$E(X^2) = p \sum_{k=0}^{+\infty} k^2 (1-p)^{k-1} = p \left[\frac{2(1-p)}{p^3} + \frac{1}{p^2} \right] = \frac{2}{p^2} - \frac{1}{p}$$

En particulier, $E(X^2) < +\infty$ donc X^2 est d'espérance finie.

Proposition 12.5 (Propriétés de l'espérance)

1. (Linéarité) : Soit X et Y deux variables aléatoires discrètes réelles positives. Dans $[0, +\infty]$:

$$E(X + Y) = E(X) + E(Y).$$

2. Soit X et Y deux variables aléatoires discrètes positives telles que $X \leq Y$. Dans $[0, +\infty]$, $E(X) \leq E(Y)$. C'est-à-dire que si on suppose que $X \leq Y$ et Y d'espérance finie alors X est d'espérance finie et $E(X) \leq E(Y)$.
3. Soit X une variable aléatoire réelle positive. Si $E(X) = 0$ alors X est nulle presque sûrement c'est-à-dire que $P(X > 0) = 0$.

Démonstration :

1. On pose $Z = X + Y = f(X, Y)$ où $f : (x, y) \mapsto x + y$. On a alors par le théorème de transfert dans $[0, +\infty]$.

$$\begin{aligned} E(X + Y) &= \sum_{(x,y) \in X(\Omega) \times Y(\Omega)} (x + y) P((X, Y) = (x, y)) \\ &= \sum_{(x,y) \in X(\Omega) \times Y(\Omega)} x P((X, Y) = (x, y)) + \sum_{(x,y) \in X(\Omega) \times Y(\Omega)} y P((X, Y) = (x, y)) \\ &= \sum_{x \in X(\Omega)} x \sum_{y \in Y(\Omega)} P((X, Y) = (x, y)) + \sum_{y \in Y(\Omega)} y \sum_{x \in X(\Omega)} P((X, Y) = (x, y)) \\ &= E(X) + E(Y) \end{aligned}$$

2. On note $Z = Y - X \geq 0$ de sorte que $Y = X + Z$. Par linéarité,

$$E(Y) = E(X) + E(Z) \geq E(X)$$

3. Supposons par l'absurde qu'il existe $x_0 \in X(\Omega)$ tel que $x_0 > 0$ et $P(X = x_0) > 0$. Alors $x_0 P(X = x_0) > 0$. Cela implique que

$$\sum_{x \in X(\Omega)} x P(X = x) \geq x_0 P(X = x_0) > 0$$

ATTENTION

Dans la démonstration de la linéarité, on décompose $P(X + Y = z)$ en utilisant des termes de la loi conjointe du couple $(X, Y) : P((X, Y) = (x, y))$ mais on ne les décompose pas en $P(X = x)P(Y = y)$. De ce fait, le résultat ne nécessite pas que les variables X et Y soient indépendantes.

1.2 Cas des variables réelles ou complexes

On va généraliser la notion d'espérance au cas des variables aléatoires à valeurs dans $\mathbf{R}$ (mais pas nécessairement positives) et dans $\mathbf{C}$. On note $\mathbf{K} = \mathbf{R}$ ou $\mathbf{C}$.

Définition 12.6 (Espérance)

Soit X une variable aléatoire discrète à valeurs réelles ou complexes.

Elle est dite d'espérance finie si la famille $(xP(X = x))_{x \in X(\Omega)}$ est sommable (c'est-à-dire que $E(|X|) < +\infty$).

Dans ce cas on appelle espérance de X la somme de cette famille :

$$E(X) = \sum_{x \in X(\Omega)} xP(X = x).$$

On note $X \in L^1$ ou $X \in L^1(\Omega, \mathbf{K})$

ATTENTION

Dans le cas où X est une variable aléatoire et que l'on veut étudier son espérance.

- Si X est à valeurs réelles positives, on peut toujours définir son espérance dans $[0, +\infty]$.
- Si X est à valeurs réelles non nécessairement positives ou complexes, il faut toujours justifier que la variable aléatoire est d'espérance finie avant de considérer son espérance. Pour cela on étudie dans $[0, +\infty]$ l'espérance de la variable aléatoire $|X|$.

Remarque : On veut éviter les séries semi-convergentes car il n'y a pas lieu d'imposer un ordre de sommation parmi les éléments de $X(\Omega)$.

Exemple : On considère une variable aléatoire X qui suit la loi $\mathcal{G}(\frac{1}{2})$. On regarde $Y = (-2)^X$. C'est -à-dire que l'on lance une pièce équilibré jusqu'à obtenir pile. Si on obtient pile au lancer n (pour la première fois) on « gagne » 2^n euros si n est pair et on « perd » 2^n si n est impair.

Cette variable n'a pas d'espérance, en effet, pour savoir si la famille $(xP(X = x))_{x \in X(\Omega)}$ on calcule $E(|Y|)$. On peut utiliser le théorème de transfert

$$E(|Y|) = \sum_{n=1}^{+\infty} |(-2)^n|P(X = n) = \sum_{n=1}^{+\infty} 2^n P(X = n) = \sum_{n=1}^{+\infty} 2^n \left(\frac{1}{2}\right)^n = +\infty.$$

Définition 12.7

Soit X une variable aléatoire discrète à valeurs réelles ou complexes. Elle est dite centrée si elle est d'espérance finie et que $E(X) = 0$.

Remarque : Si $X \in L^1$, la variable $X - E(X)$ est centrée.

1.3 Propriétés de l'espérance

La plupart des propriétés vues sur l'espérance des variables aléatoires positives sont encore vraies pour les variables aléatoires réelles ou complexes d'espérance finie.

Théorème 12.8 (Formule de transfert)

Soit X une variable aléatoire discrète et f une fonction définie sur $X(\Omega)$ et à valeurs dans $\mathbf{K}$.

La variable aléatoire discrète $f(X)$ est d'espérance finie si et seulement si la famille $(f(x)P(X = x))_{x \in X(\Omega)}$ est sommable. Dans ce cas,

$$E(f(X)) = \sum_{x \in X(\Omega)} f(x)P(X = x).$$

Démonstration : Par définition,

$$\begin{aligned} (f(x)P(X = x))_{x \in X(\Omega)} \text{ est sommable} &\iff (|f(x)|P(X = x))_{x \in X(\Omega)} \text{ est sommable} \\ &\iff E(|f(X)|) < +\infty \text{ (transfert pour les variables positives)} \\ &\iff f(X) \text{ est d'espérance finie} \end{aligned}$$

Si on est dans ce cas, on peut appliquer les théorème de sommations par paquets. Précisément,

$$\begin{aligned} \sum_{y \in Y(\Omega)} yP(Y = y) &= \sum_{y \in Y(\Omega)} y \sum_{x \in X(\Omega)} P((X, Y) = (x, y)) \\ &= \sum_{(x, y) \in X(\Omega) \times Y(\Omega)} yP((X, Y) = (x, y)) \\ &= \sum_{\substack{(x, y) \in X(\Omega) \times Y(\Omega) \\ y = f(x)}} yP((X, Y) = (x, y)) \\ &= \sum_{\substack{(x, y) \in X(\Omega) \times Y(\Omega) \\ y = f(x)}} f(x)P((X, Y) = (x, y)) \\ &= \sum_{x \in X(\Omega)} f(x) \underbrace{\sum_{\substack{y \in Y(\Omega) \\ y = f(x)}} P((X, Y) = (x, y))}_{P(X=x)} \\ &= \sum_{x \in X(\Omega)} f(x)P(X = x) \end{aligned}$$

□

Proposition 12.9 (Propriétés de l'espérance)

1. (Linéarité) : Soit X et Y deux variables aléatoires discrètes réelles ou complexes d'espérance finie. Soit λ un scalaire, les variables λX et $X + Y$ sont d'espérance finie et

$$E(\lambda X) = \lambda E(X); E(X + Y) = E(X) + E(Y).$$

2. Soit X une variable aléatoire discrète réelle ou complexe et Y une variable aléatoire discrète à valeurs positives. On suppose que $|X| \leq Y$ et Y d'espérance finie alors X est d'espérance finie et $|E(X)| \leq E(|X|) \leq E(Y)$.
3. (Positivité) : Soit X une variable aléatoire positive alors $E(X)$ est positive.
4. (Croissance) : Soit X et Y deux variables aléatoires discrètes réelles d'espérance finie si $X \leq Y$ alors $E(X) \leq E(Y)$.

Démonstration :

1. — D'après l'énoncé du théorème de transfert pour $f : x \mapsto \lambda x$.

$$\begin{aligned} X \text{ est d'espérance finie} &\iff (xP(X=x))_{x \in X(\Omega)} \text{ est sommable} \\ &\implies (\lambda xP(X=x))_{x \in X(\Omega)} \text{ est sommable} \\ &\implies \lambda X \text{ est d'espérance finie} \end{aligned}$$

Dans ce cas,

$$E(\lambda X) = \sum_{x \in X(\Omega)} \lambda x P(X=x) = \lambda E(X)$$

- On pose $Z = X + Y$. On sait que $|Z| = |X + Y| \leq |X| + |Y|$ donc si X et Y sont d'espérance finie alors Z aussi. De plus, dans ce cas, en utilisant le théorème de transfert pour $f : (x, y) \mapsto x + y$,

$$\begin{aligned} E(Z) &= \sum_{(x,y) \in X(\Omega) \times Y(\Omega)} (x+y) P((X,Y) = (x,y)) \\ &= \sum_{x \in X(\Omega)} x \sum_{y \in Y(\Omega)} P((X,Y) = (x,y)) + \sum_{y \in Y(\Omega)} y \sum_{x \in X(\Omega)} P((X,Y) = (x,y)) \\ &= \sum_{x \in X(\Omega)} x P(X=x) + \sum_{y \in Y(\Omega)} y P(Y=y) \\ &= E(X) + E(Y) \end{aligned}$$

2. On sait que $|X| \leq Y$ et Y d'espérance finie donc $E(|X|) < +\infty$ ce qui signifie que X est d'espérance finie. De plus $E(|X|) \leq E(Y)$.

De plus, d'après le théorème de transfert,

$$E(|X|) = \sum_{x \in X(\Omega)} |x| P(X=x) \geq \left| \sum_{x \in X(\Omega)} x P(X=x) \right| = |E(X)|$$

3. Par définition
4. Par linéarité et positivité.

□

Proposition 12.10

Soit X et Y deux variables aléatoires discrètes réelles ou complexes. On suppose que X est presque sûrement égale à Y .

1. La variable $X - Y$ est d'espérance finie et $E(X - Y) = 0$
2. La variable aléatoire X est d'espérance finie si et seulement si Y l'est. Dans ce cas, $E(X) = E(Y)$.

Remarque : Cet énoncé ne figure pas dans le programme ; il faudrait le redémontrer le cas échéant.

Démonstration : On veut calculer $E(X - Y)$. On commence par étudier $E(|X - Y|)$. Dans $[0, +\infty]$, d'après le théorème de transfert,

$$\begin{aligned} E(|X - Y|) &= \sum_{(x,y) \in X(\Omega) \times Y(\Omega)} |x - y| \mathbf{P}(X = x, Y = y) \\ &= \sum_{\substack{(x,y) \in X(\Omega) \times Y(\Omega) \\ x=y}} |x - y| \mathbf{P}(X = x, Y = y) + \sum_{\substack{(x,y) \in X(\Omega) \times Y(\Omega) \\ x \neq y}} |x - y| \mathbf{P}(X = x, Y = y) \\ &= 0 + 0 \end{aligned}$$

On en déduit que $X - Y$ est d'espérance finie. On peut reprendre le calcul sans les valeurs absolues / modules pour obtenir que $E(X - Y) = 0$.

En écrivant $X = (X - Y) + Y$ on obtient que X est d'espérance finie si et seulement si Y est d'espérance finie. Dans ce cas,

$$E(X) = E(X - Y) + E(Y) = E(Y)$$

□

Proposition 12.11

Soit X et Y deux variables aléatoires discrètes **indépendantes** d'espérances finies. La variable XY est d'espérance finie et

$$E(XY) = E(X)E(Y).$$

Démonstration :

— On suppose pour commencer que X et Y sont positives.

On pose $Z = XY$ et $f : (x, y) \mapsto xy$. En appliquant le théorème de transfert on sait que, dans $[0, +\infty]$,

$$\begin{aligned} E(Z) &= \sum_{(x,y) \in X(\Omega) \times Y(\Omega)} xy \mathbf{P}((X, Y) = (x, y)) \\ &= \sum_{(x,y) \in X(\Omega) \times Y(\Omega)} xy \mathbf{P}(X = x) \mathbf{P}(Y = y) \text{ par indépendance} \\ &= \left(\sum_{x \in X(\Omega)} x \mathbf{P}(X = x) \right) \times \left(\sum_{y \in Y(\Omega)} y \mathbf{P}(Y = y) \right) \\ &= E(X) \cdot E(Y) \end{aligned}$$

- En particulier, si X et Y sont d'espérance finie, alors XY aussi et $E(XY) = E(X)E(Y)$.
- Dans le cas général, on applique le résultat ci-dessus aux variables $|X|$ et $|Y|$ ce qui prouve que si X et Y sont d'espérance finie alors XY aussi et que la famille $(xyP((X, Y) = (x, y)))_{x \in X(\Omega), y \in Y(\Omega)}$ est sommable. On peut donc écrire le même calcul que ci-dessus.

□

2 Variance et écart type d'une variable aléatoire réelle

Dans ce paragraphe on se fixe un espace probabilisé $(\Omega, \mathcal{A}, P)$.

2.1 Variables aléatoires réelle de carré sommable

Définition 12.12

Soit X une variable aléatoire discrète réelle. On dit qu'elle est de carré sommable ou qu'elle admet un moment d'ordre 2 si $E(X^2) < +\infty$.
Dans ce cas, on note $X \in L^2$ ou $X \in L^2(\Omega, \mathbb{R})$.

Proposition 12.13

Soit X une variable aléatoire discrète réelle. Si $X \in L^2$ alors $X \in L^1$.

Démonstration :

- Première méthode : Il suffit de voir que

$$|X| = \sqrt{X^2} \leq \frac{X^2 + 1}{2}$$

En effet, de manière générale, pour $(a, b) \in \mathbb{R}_+^2$, $\sqrt{ab} \leq \frac{a+b}{2}$ (comparaison entre la moyenne géométrique et la moyenne arithmétique).

Maintenant, on suppose que X^2 a une espérance finie, or la variable constante égale à 1 a aussi une espérance finie donc $|X|$ aussi.

- Deuxième méthode : on peut procéder par disjonction. On calcule $E(|X|)$ dans $[0, +\infty]$ car $|X|$ est une variable à valeurs positives.

$$\begin{aligned} E(|X|) &= \sum_{x \in X(\Omega)} |x| P(X = x) \text{ (par transfert)} \\ &= \sum_{\substack{x \in X(\Omega) \\ |x| \leq 1}} |x| P(X = x) + \sum_{\substack{x \in X(\Omega) \\ |x| > 1}} |x| P(X = x) \\ &\leq \sum_{\substack{x \in X(\Omega) \\ |x| \leq 1}} P(X = x) + \sum_{\substack{x \in X(\Omega) \\ |x| > 1}} x^2 P(X = x) \\ &\leq 1 + E(X^2) < +\infty \end{aligned}$$

□

ATTENTION

C'est une méthode classique qu'il faut connaître

Exercice : Soit X une variable aléatoire positive. Soit $0 < \alpha \leq \beta$. Montrer que X^β est d'espérance finie alors X^α aussi.

Correction

on peut procéder par disjonction. On calcule $E(|X|^\alpha)$ dans $[0, +\infty]$ car $|X|^\alpha$ est une variable à valeurs positives.

$$\begin{aligned} E(|X|^\alpha) &= \sum_{x \in X(\Omega)} |x|^\alpha \mathbf{P}(X = x) \text{ (par transfert)} \\ &= \sum_{\substack{x \in X(\Omega) \\ |x| \leq 1}} |x|^\alpha \mathbf{P}(X = x) + \sum_{\substack{x \in X(\Omega) \\ |x| > 1}} |x|^\alpha \mathbf{P}(X = x) \\ &\leq \sum_{\substack{x \in X(\Omega) \\ |x| \leq 1}} \mathbf{P}(X = x) + \sum_{\substack{x \in X(\Omega) \\ |x| > 1}} |x|^\beta \mathbf{P}(X = x) \\ &\leq 1 + E(|X|^\beta) < +\infty \end{aligned}$$

Théorème 12.14 (Inégalité de Cauchy-Schwarz)

Soit X et Y deux variables aléatoires discrètes réelles. On suppose $X \in L^2$ et $Y \in L^2$.

1. La variable aléatoire XY est d'espérance finie et

$$E(XY)^2 \leq E(X^2)E(Y^2)$$

2. Il y a égalité si et seulement si X et Y sont proportionnelles presque sûrement c'est-à-dire s'il existe $\lambda \in \mathbf{R}$ tel que $\mathbf{P}(X = \lambda Y) = 1$ ou $\mathbf{P}(Y = \lambda X) = 1$

Démonstration : On utilise la preuve « classique » de l'inégalité de Cauchy-Schwarz.

Remarquons d'abord que $|XY| = \sqrt{X^2 Y^2} \leq \frac{X^2 + Y^2}{2}$ donc XY est d'espérance finie. On en déduit que pour tout λ , $X + \lambda Y \in L^2$ car

$$(X + \lambda Y)^2 = X^2 + 2\lambda XY + \lambda^2 Y^2.$$

Maintenant, on pose

$$P(\lambda) = E((X + \lambda Y)^2) = E(X^2 + 2\lambda XY + \lambda^2 Y^2) = E(X^2) + 2\lambda E(XY) + \lambda^2 E(Y^2)$$

— Si $E(Y^2) = 0$, comme Y est une variable positive, $Y = 0$ presque sûrement. On en déduit que $XY = 0$ presque sûrement et donc $E(XY) = 0$.

Il y a donc égalité : $E(XY) = 0 = E(X^2)E(Y^2)$. On remarque que si on pose $\lambda = 0$, alors $P(Y = 0X) = P(Y = 0) = 1$.

— On suppose que $E(Y^2) \neq 0$. La fonction P est un polynôme du second degré qui ne prend que des valeurs positives donc son discriminant est négatif ou nul. Cela signifie que $\Delta =$

$(2E(XY))^2 - 4E(X^2)E(Y^2) = 4(E(XY)^2 - E(X^2)E(Y^2)) \leq 0$ et donc, en divisant par 4,

$$E(XY)^2 \leq E(X^2)E(Y^2)$$

Regardons les cas d'égalité.

- $\Rightarrow$ On suppose que $E(XY)^2 \leq E(X^2)E(Y^2)$. Cela signifie que $\Delta = 0$ et donc il existe λ_0 tel que $P(\lambda_0) = E((X + \lambda_0 Y)^2) = 0$. On en déduit que $X + \lambda_0 Y = 0$ presque sûrement.
- $\Leftarrow$ On suppose qu'il existe $\lambda \in \mathbb{R}$ tel que $X = \lambda Y$ presque sûrement. Comme $(X = \lambda Y) \subset (XY = \lambda Y^2)$, $P(XY = \lambda Y^2) = 1$ c'est-à-dire que XY est égale à λY^2 presque sûrement. On en déduit que

$$E(XY)^2 = E(\lambda Y^2)^2 = \lambda^2 E(Y^2)E(Y^2) = E((\lambda Y)^2)E(Y^2) = E(X^2)E(Y^2)$$

□

Corollaire 12.15

L'ensemble des variables aléatoires discrètes à valeurs réelles ayant un moment d'ordre 2 est un sous-espace vectoriel de l'ensemble des variables aléatoires discrètes à valeurs réelles

Démonstration : L'ensemble des variables aléatoires discrètes à valeurs réelles ayant un moment d'ordre 2 n'est pas vide car il contient la variable nulle.

Soit X et Y deux variables aléatoires discrètes à valeurs réelles ayant un moment d'ordre 2. Soit λ un réel. On pose $Z = X + \lambda Y$. On a alors $Z^2 = X^2 + 2\lambda XY + \lambda Y^2$. On en déduit que Z^2 est d'espérance finie car c'est la somme de trois variables aléatoires d'espérance finie. On obtient ainsi que Z a un moment d'ordre 2. □

Définition 12.16

Soit X une variable aléatoire discrète réelle telle que $X \in L^2$. On appelle variance de X et on note $V(X)$,

$$V(X) = E((X - E(X))^2).$$

On appelle écart-type et on note $\sigma(X)$,

$$\sigma(X) = \sqrt{V(X)}.$$

Remarque : Comme $X \in L^2$, $X - E(X) \in L^2$.

Proposition 12.17 (Formule de König-Huygens)

Soit X une variable aléatoire discrète réelle telle que $X \in L^2$. On a

$$V(X) = E(X^2) - E(X)^2$$

Remarque : On utilise souvent cette formule pour calculer, avec la formule de transfert la variance :

$$V(X) = \sum_{x \in X(\Omega)} x^2 P(X = x) - \left(\sum_{x \in X(\Omega)} x P(X = x) \right)^2.$$

Démonstration : En utilisant la linéarité de l'espérance :

$$V(X) = E((X - E(X))^2) = E(X^2 - 2E(X)X + E(X)^2) = E(X^2) - 2E(X)^2 + E(X)^2 = E(X^2) - E(X)^2$$

□

Proposition 12.18

Soit X une variable aléatoire discrète réelle telle que $X \in L^2$. Soit a, b deux réels,

$$V(aX + b) = a^2 V(X); \sigma(aX + b) = |a| \sigma(X).$$

Démonstration : Il suffit de faire le calcul.

□

Proposition 12.19

Soit X une variable aléatoire discrète réelle telle que $X \in L^2$. On suppose que $V(X) = 0$, il existe alors $\lambda \in \mathbb{R}$ telle que X est presque sûrement égale à λ .

Démonstration : La variable aléatoire $(X - E(X))^2$ est une variable aléatoire positive d'espérance nulle. Elle est nulle presque sûrement et donc $X = E(X)$ presque sûrement.

□

Proposition-Définition 12.20

Soit X une variable aléatoire discrète réelle admettant un moment d'ordre 2.

1. Elle est dite réduite si $V(X) = 1$.
2. Si $\sigma(X) > 0$, la variable aléatoire $\frac{X - E(X)}{\sigma(X)}$ est centrée est réduite.

2.2 Lois usuelles

Commençons par rappeler les variances des lois usuelles vues en première année

Proposition 12.21

Soit X une variable aléatoire discrète à valeurs réelles.

1. (Loi de Bernoulli) : Soit $p \in [0, 1]$, si $X \sim \mathcal{B}(p)$ alors $V(X) = pq = p(1 - p)$.
2. (Loi binomiale) : Soit $p \in [0, 1]$ et $n \in \mathbb{N}^*$, si $X \sim \mathcal{B}(n, p)$ alors $V(X) = npq$.

Passons aux lois vues cette année,

Théorème 12.22

Soit X une variable aléatoire discrète réelle.

1. (Loi géométrique) : Soit $p \in]0, 1]$, si $X \sim \mathcal{G}(p)$ alors $V(X) = \frac{q}{p^2}$.
2. (Loi de Poisson) : Soit $\lambda \in \mathbb{R}_+^*$, si $X \sim \mathcal{P}(\lambda)$ alors $V(X) = \lambda$

Démonstration :

1. On a déjà vu que si $X \sim \mathcal{G}(p)$ alors $E(X^2) = \frac{2}{p^2} - \frac{1}{p}$. De ce fait,

$$V(X) = E(X^2) - E(X)^2 = \frac{2}{p^2} - \frac{1}{p} - \frac{1}{p^2} = \frac{1}{p^2} - \frac{1}{p} = \frac{q}{p^2}$$

2. Soit $X \sim \mathcal{P}(\lambda)$.

D'après la formule de transfert

$$\begin{aligned} E(X^2) &= \sum_{k=0}^{+\infty} k^2 \frac{\lambda^k e^{-\lambda}}{k!} \\ &= \sum_{k=0}^{+\infty} k(k-1) \frac{\lambda^k e^{-\lambda}}{k!} + \sum_{k=0}^{+\infty} k \frac{\lambda^k e^{-\lambda}}{k!} \\ &= \sum_{k=2}^{+\infty} \frac{\lambda^k e^{-\lambda}}{(k-2)!} + \sum_{k=1}^{+\infty} \frac{\lambda^k e^{-\lambda}}{(k-1)!} \\ &= \lambda^2 \sum_{k=0}^{+\infty} \frac{\lambda^k e^{-\lambda}}{k!} + \lambda \sum_{k=0}^{+\infty} \frac{\lambda^k e^{-\lambda}}{k!} \\ &= \lambda^2 + \lambda \end{aligned}$$

On en déduit que $V(X) = E(X^2) - E(X)^2 = \lambda^2 + \lambda - \lambda^2 = \lambda$.

□

2.3 Covariance

On a vu que $E(XY)$ n'était pas, en général égal à $E(X)E(Y)$.

Proposition-Définition 12.23

Soit X et Y deux variables aléatoires discrètes réelles telles que $X \in L^2$ et $Y \in L^2$.
On appelle covariance de X et Y et on note $\text{Cov}(X, Y)$:

$$\text{Cov}(X, Y) = E((X - E(X))(Y - E(Y))) = E(XY) - E(X)E(Y)$$

Démonstration : Comme $X - E(X) \in L^2$ et $Y - E(Y) \in L^2$ leur produit a une espérance finie et

$$\begin{aligned} E((X - E(X))(Y - E(Y))) &= E(XY - E(X)Y - E(Y)X + E(X)E(Y)) \\ &= E(XY) - 2E(X)E(Y) + E(X)E(Y) \\ &= E(XY) - E(X)E(Y) \end{aligned}$$

□

Proposition 12.24

Soit X et Y deux variables aléatoires discrètes réelles telles que $X \in L^2$ et $Y \in L^2$. Si X et Y sont indépendantes alors $\text{Cov}(X, Y) = 0$.

Démonstration : On a déjà vu que si X et Y étaient indépendantes alors $E(XY) = E(X)E(Y)$ et donc $\text{Cov}(X, Y) = 0$. $\square$

ATTENTION

Ce n'est pas une équivalence, on peut avoir deux variables non indépendantes telles que $\text{Cov}(X, Y) = 0$.

Exemple : On considère une variable aléatoire discrète X dont la loi est définie par $X(\Omega) = \{-1, 0, 1\}$ et

$$P(X = -1) = P(X = 0) = P(X = 1) = \frac{1}{3}$$

On pose alors $Y = X^2$.

On voit que $E(X) = 0$. De plus $X^3 = X$ donc

$$\text{Cov}(X, Y) = E(XY) - E(X)E(Y) = E(X^3) - 0 = E(X) = 0$$

Par contre X et Y ne sont pas indépendantes car $P(X = 1, Y = 0) = 0 \neq P(X = 1)P(Y = 0)$.

Proposition 12.25

Soit E l'espace vectoriel des variables aléatoires discrètes réelles admettant un moment d'ordre 2. La covariance

$$\begin{aligned} \text{Cov} : E \times E &\rightarrow \mathbf{R} \\ (X, Y) &\mapsto \text{Cov}(X, Y) \end{aligned}$$

est une forme bilinéaire positive. De plus pour tout X dans E ,

$$\text{Cov}(X, X) = V(X) = \sigma(X)^2$$

Démonstration : Il suffit d'utiliser la linéarité de l'espérance. $\square$

Remarque : La covariance se comporte presque comme un produit scalaire. Le presque venant du fait qu'elle n'est pas définie, on peut avoir $\text{Cov}(X, X) = V(X) = 0$ sans que X ne soit nul. En effet $V(X) = 0$ implique juste $X = E(X)$ presque sûrement et pas que $X = 0$.

Proposition 12.26

Soit X et Y deux variables aléatoires discrètes réelles telles que $X \in L^2$ et $Y \in L^2$.

$$|\text{Cov}(X, Y)| \leq \sigma(X)\sigma(Y)$$

Démonstration : Il suffit d'utiliser l'inégalité de Cauchy-Schwarz avec la définition $\text{Cov}(X, Y) = E((X - E(X))(Y - E(Y)))$ $\square$

Corollaire 12.27

Soit $X_1, \dots, X_n$ des variables aléatoires discrètes réelles telles que $X_i \in L^2$ pour tout $i \in \llbracket 1; n \rrbracket$.

1. On a

$$V(X_1 + \dots + X_n) = \sum_{i=1}^n V(X_i) + 2 \sum_{1 \leq i < j \leq n} \text{Cov}(X_i, X_j)$$

2. Si on suppose que pour tout $(i, j) \in \llbracket 1; n \rrbracket^2$, $i \neq j \Rightarrow \text{Cov}(X_i, X_j) = 0$ alors

$$V(X_1 + \dots + X_n) = \sum_{i=1}^n V(X_i)$$

Démonstration : Ce sont juste les formules classiques de développement par bilinéarité. $\square$

Remarque : L'hypothèse dans la deuxième partie de l'énoncé est en particulier vérifiée si les variables sont deux à deux indépendantes, ce qui est le cas si la famille $(X_1, \dots, X_n)$ est indépendante.

Exemple : On considère une suite infinie $(X_i)_{i \in \mathbb{N}^*}$ de variables aléatoires indépendantes qui suivent toutes la loi de Bernoulli $\mathcal{B}(p)$ où $p \in]0, 1[$.

Pour tout entier $n \geq 1$, on note Z_n le rang du n -ième succès :

$$Z_n : \omega \mapsto \min\{i \in \mathbb{N}^*, \#\{k \leq i, X_k(\omega) = 1\} = n\}$$

Il est évident que $Z_n(\Omega) \subset \llbracket n, +\infty \rrbracket$. De plus pour $m \geq n$, on veut décrire l'événement $(Z_n = m)$. Pour cela il y a $n - 1$ variables parmi $X_1, \dots, X_{m-1}$ qui prennent la valeur 1 (les autres prennent la valeur 0) et X_m prend la valeur 1. Si on note $\mathcal{K}$ l'ensemble des parties de $\llbracket 1, m-1 \rrbracket$ de cardinal $n-1$ (ensemble qui n'est pas vide car on a supposé $m \geq n$). Pour toute partie $K \in \mathcal{K}$ on note $\bar{K}$ son complémentaire dans $\llbracket 1, m-1 \rrbracket$

On a alors

$$(Z_n = m) = \left(\bigcup_{K \in \mathcal{K}} \left(\bigcap_{i \in K} (X_i = 1) \cap \bigcap_{j \in \bar{K}} (X_j = 0) \right) \right) \cap (X_m = 1)$$

Comme $\mathcal{K}$ a $\binom{m-1}{n-1}$ éléments, on obtient que

$$P(Z_n = m) = \binom{m-1}{n-1} p^n q^{m-n}$$

Considérons maintenant $(Y_i)_{i \geq 1}$ une suite de variables aléatoires indépendantes suivant toutes la loi géométrique de paramètre p . On note pour $n \geq 1$, $T_n = Y_1 + \dots + Y_n$. On cherche à vérifier que T_n suit la même loi que Z_n . On procède par récurrence sur $n \geq 1$.

— Pour $n = 1$ c'est évident

– Soit $n \in \mathbb{N}^*$ on suppose que $T_n \sim Z_n$. On a alors pour tout $m \geq n + 1$,

$$\begin{aligned}
 P(T_{n+1} = m) &= \sum_{i=n}^{m-1} P(T_n = i)P(Y_{n+1} = m - i | T_n = i) \\
 &= \sum_{i=n}^{m-1} P(T_n = i)P(Y_{n+1} = m - i) \text{ car } Y_{n+1} \perp\!\!\!\perp T_n \\
 &= \sum_{i=n}^{m-1} \binom{i-1}{n-1} p^n q^{i-n} q^{m-i-1} p \\
 &= \binom{m-1}{n} p^{n+1} q^{m-(n+1)}
 \end{aligned}$$

où on a utilisé que

$$\sum_{i=n}^{m-1} \binom{i-1}{n-1} = \sum_{i=n}^{m-1} \left(\binom{i}{n} - \binom{i-1}{n} \right) = \binom{m-1}{n}$$

Finalement

On a donc

$$E(Z_n) = E(Y_1 + \dots + Y_n) = np$$

et

$$V(Z_n) = V(Y_1 + \dots + Y_n) = n \frac{1-p}{p^2}$$

3 Inégalités probabilistes et loi des grands nombres

La plupart du temps, il n'est pas possible de calculer les probabilités. On peut cependant obtenir des inégalités.

3.1 Inégalités

Théorème 12.28 (Inégalité de Markov)

Soit X une variable aléatoire discrète réelle **positive** et $a > 0$. On a

$$P(X \geq a) \leq \frac{E(X)}{a}$$

c'est-à-dire, $aP(X \geq a) \leq E(X)$.

Remarque : L'inégalité est évidente si $E(X) = +\infty$.

Démonstration : Donnons deux démonstrations.

– (Première démonstration) : On découpe la somme qui va définir l'espérance selon que x soit

supérieur ou inférieur à a .

$$\begin{aligned}
 E(X) &= \sum_{x \in X(\Omega)} xP(X = x) \\
 &= \sum_{\substack{x \in X(\Omega) \\ x < a}} xP(X = x) + \sum_{\substack{x \in X(\Omega) \\ x \geq a}} xP(X = x) \\
 &\geq \sum_{\substack{x \in X(\Omega) \\ x \geq a}} xP(X = x) \\
 &\geq \sum_{\substack{x \in X(\Omega) \\ x \geq a}} aP(X = x) \\
 &= aP(X \geq a)
 \end{aligned}$$

– (Deuxième démonstration) : On considère la variable aléatoire $a\mathbb{1}_{(X \geq a)}$ de sorte que

$$a\mathbb{1}_{(X \geq a)} \leq X$$

Par croissance,

$$aP(X \geq a) = E(a\mathbb{1}_{(X \geq a)}) \leq E(X).$$

□

Matheux (Andrei Markov : 1856 - 1922)



Sous la tutelle de Tchebychev, il devint membre de l'Académie des sciences de Saint-Petersbourg en 1886.

Ses premiers travaux portent sur la théorie des nombres, les formes quadratiques, les fractions continues, les limites d'intégrales et les convergences de séries.

Ses travaux sur la théorie des probabilités l'ont amené à mettre au point les chaînes de Markov qui peuvent représenter les prémisses de la théorie du calcul stochastique.

Il a officiellement pris sa retraite en 1905, mais a continué à enseigner jusqu'à la fin de sa vie.

Dès 1906, il commence ses recherches sur le calcul de probabilités. Il introduit de façon précise les processus aléatoires et démontre rigoureusement le théorème de la limite centrale.

Exercice : Montrer que si φ est croissante et positive sur I . Soit Y une variable aléatoire discrète réelle telle que $Y(\Omega) \subset I$. Montrer que

$$\forall b \in I \text{ tel que } \varphi(b) > 0, P(Y \geq b) \leq \frac{E(\varphi(Y))}{\varphi(b)}.$$

Théorème 12.29 (Inégalité de Bienaymé-Tchebychev)

Soit X une variable aléatoire discrète réelle telle que $X \in L^2$.

Soit $a > 0$.

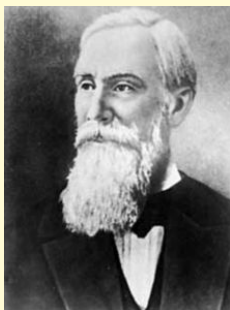
$$P(|X - E(X)| \geq a) \leq \frac{V(X)}{a^2}.$$

Démonstration : On a $P(|X - E(X)| \geq a) = P((X - E(X))^2 \geq a^2)$ or $(X - E(X))^2$ est une variable positive donc on peut appliquer l'inégalité de Markov :

$$P((X - E(X))^2 \geq a^2) \leq \frac{E((X - E(X))^2)}{a^2} = \frac{V(X)}{a^2}.$$

□

Matheux (Pafnouti Tchebychev : 1821-1894)



Il est connu pour ses travaux dans les domaines des probabilités, des statistiques, et de la théorie des nombres.

Tchebychev reprend le vaste programme lancé par Jacques Bernoulli, Abraham de Moivre et Siméon Denis Poisson pour énoncer et démontrer de façon rigoureuse des théorèmes limites, c'est-à-dire pour établir les tendances asymptotiques des phénomènes naturels. Il établit une loi des grands nombres très générale et donne une nouvelle et brillante méthode de démonstration basée sur l'inégalité énoncée par Bienaymé et démontrée par lui-même.

En théorie des nombres, Tchebychev obtint en 1848-1852 des résultats corroborant une conjecture de Gauss et Legendre relative à la raréfaction des nombres premiers. Si ces résultats ne lui permirent pas de démontrer la conjecture (le théorème des nombres premiers), ils lui permirent néanmoins de s'en approcher considérablement, et par ailleurs de démontrer une conjecture énoncée par Bertrand : « Pour tout entier n au moins égal à 2, il existe un nombre premier entre n et $2n$ ».

Matheux (Irénée-Jules Bienaymé : 1796-1878)



Irénée-Jules Bienaymé, né à Paris le 28 août 1796 et mort à Paris le 19 octobre 1878, est un probabiliste et statisticien français. Continuateur de l'œuvre de Laplace dont il généralise la méthode des moindres carrés, il contribue à la théorie des probabilités, au développement de la statistique et à leurs applications aux calculs financiers, à la démographie et aux statistiques sociales.

Il a énoncé en particulier l'inégalité de Bienaymé-Tchebychev concernant la loi des grands nombres (1869). Il est aussi à l'origine de ce que l'on appelle aujourd'hui les arbres de Bienaymé-Galton-Watson

Exemple : On considère une expérience aléatoire décrite par l'espace probabilisé $(\Omega, \mathcal{A}, P)$ et A un événement. On note $p = P(A)$. Si on effectue de nombreuses fois (N) l'expérience de manière indépendante et que l'on note X_N le nombre de fois que l'événement A se réalise et si on note $F_N = \frac{X_N}{N}$ la fréquence à laquelle l'événement A se réalise alors, il semble logique que F_N va être proche de p .

On voit que $X_N \hookrightarrow \mathcal{B}(N, p)$. De ce fait, $E(X_N) = Np$ et $V(X_N) = Npq$. On en déduit que

$$E(F_N) = p \text{ et } V(F_N) = \frac{pq}{N}.$$

Dès lors, l'inégalité de Bienaymé-Tchebychev affirme que, pour $\varepsilon > 0$,

$$P(|F_N - p| \geq \varepsilon) \leq \frac{pq}{N\varepsilon^2}.$$

On voit donc que, quand N tend vers $+\infty$, $P(|F_N - p| \geq \varepsilon)$ tend vers 0 et cela pour toute valeur de ε .

Par exemple si on lance N fois un dé équilibré à 6 faces. Alors le dé fera 6 environ une fois sur 6 et de ce fait, F_N ne devrait pas être loin de $1/6$. D'après l'inégalité de Bienaymé-Tchebychev affirme que, pour $\varepsilon = 5\% = 1/20$

$$P(|F_N - 1/6| \geq 1/20) \leq \frac{5/36}{N \times (1/20)^2} = \frac{1000}{18N}.$$

De ce fait, si on veut être sûr à 90% que F_N ne diffère de pas plus de $1/20$ de $1/6$, on doit prendre N tel que

$$\frac{1000}{18N} \leq 1 - 90/100 = 1/10$$

c'est-à-dire

$$N \geq \frac{10000}{18} = 555,5.$$

Remarque : C'est un cas particulier de la loi faible des grands nombres que nous verrons plus loin.

3.2 Loi faible des grands nombres

On peut généraliser l'exemple vu au moment de l'inégalité de Bienaymé-Tchebychev.

Théorème 12.30 (Loi faible des grands nombres)

Soit $(X_n)_{n \geq 1}$ une suite de variables aléatoires discrètes réelles deux à deux indépendantes, suivant une même loi et telles que pour tout $n \in \mathbb{N}^*$, $X_n \in L^2$.

On pose $S_n = \sum_{i=1}^n X_i$ et $m = E(X_1)$.

Dans ce cas, pour tout $\varepsilon > 0$,

$$P\left(\left|\frac{S_n}{n} - m\right| \geq \varepsilon\right) \xrightarrow{n \rightarrow \infty} 0.$$

Remarques :

1. En statistiques, on s'intéresse souvent à reproduire une même expérience de nombreuses fois. Cela donne naissance pour chaque expérience à une variable aléatoire. De ce fait on récupère une suite de variables aléatoires indépendantes et qui suivent la même loi. On dit que c'est une suite de variables indépendantes et identiquement distribuées (iid).
2. Cela justifie le fait que l'espérance d'une variable aléatoire est bien ce que l'on obtiendra (presque sûrement) en réalisant **un très grand nombre de fois** une expérience.

Démonstration : (A savoir refaire) : Comme ci-dessus, $E(S_n) = nm$ et donc $E\left(\frac{S_n}{n}\right) = m$. D'après l'inégalité de Bienaymé-Tchebychev,

$$P\left(\left|\frac{S_n}{n} - m\right| \geq \varepsilon\right) \leq \frac{V\left(\frac{S_n}{n}\right)}{\varepsilon^2}$$

Maintenant,

$$V\left(\frac{S_n}{n}\right) = \frac{1}{n^2} V(X_1 + \dots + X_n) = \frac{1}{n} V(X_1).$$

On a donc

$$P\left(\left|\frac{S_n}{n} - m\right| \geq \varepsilon\right) \leq \frac{V(X_1)}{n\varepsilon^2} \xrightarrow{n \rightarrow \infty} 0.$$

□

Remarque : On dit que la variable $\frac{S_n}{n}$ converge *en probabilité* vers $E(X)$.

De manière générale quand on se donne une suite de variables aléatoires, (X_k) et une variable aléatoire Y on peut définir plusieurs types de convergence :

– On dit que (X_k) converge *en probabilité* vers Y si, pour tout $\varepsilon > 0$,

$$\lim_{k \rightarrow \infty} P(|X_k - Y| \geq \varepsilon) = 0$$

– On dit que (X_k) converge *presque sûrement* vers Y si,

$$P\left(\lim_{k \rightarrow \infty} X_k = Y\right) = 1$$

Il y a une version « plus forte » qui s'appelle la loi *forte* des grands nombres qui dit que si les X_i ont une espérance finie alors

$$P\left(\lim_{n \rightarrow +\infty} \frac{S_n}{n} = m\right) = 1$$

C'est-à-dire que $\frac{S_n}{n}$ converge *presque sûrement* vers $E(X)$.

4 Fonctions génératrices

4.1 Généralités

Définition 12.31 (Fonctions génératrices)

Soit X une variable aléatoire à valeur dans $\mathbb{N}$, on appelle fonction génératrice de X la fonction

$$G_X : t \mapsto E(t^X) = \sum_{n=0}^{+\infty} P(X = n)t^n$$

Proposition 12.32

La série entière $\sum_{n \geq 0} P(X = n)t^n$ converge normalement sur $[-1, 1]$.

Démonstration : Si on pose $f_n : t \mapsto P(X = n)t^n$ on a $\|f_n\|_\infty = P(X = n)$. Or la série $\sum_n P(X = n)$ converge (et est des somme 1) donc la série entière converge normalement sur $[-1, 1]$. □

Corollaire 12.33

Avec les notations précédentes

1. Le rayon de convergence de la série entière $\sum_{n \geq 0} \mathbf{P}(X = n)t^n$ est au moins 1.
2. La fonction G_X est continue sur $[-1; 1]$
3. La fonction G_X est de classe $\mathcal{C}^\infty$ sur $] - R, R[$ donc sur $] - 1, 1[$

Démonstration :

1. Par définition du rayon de convergence.
2. La fonction G_X est la limite uniforme de fonctions continues qui convergent uniformément car normalement sur $[-1; 1]$. La limite G_X est donc continue.
3. Il suffit d'utiliser que $R \geq 1$.

Remarque : Le rayon de convergence peut être plus grand que 1, par exemple pour une loi de Poisson il est égal à $+\infty$. Il peut de même être juste égal à 1 on verra que c'est en particulier le cas si X n'a pas une espérance finie.

La connaissance de G_X permet d'obtenir de nombreuses informations sur X .

Proposition 12.34

Soit X une variable aléatoire à valeur dans $\mathbf{N}$.

1. $\forall n \in \mathbf{N}, \mathbf{P}(X = n) = \frac{G_X^{(n)}(0)}{n!}$
2. $G_X(1) = 1$.

Exercice : Soit X et Y deux variables aléatoires à valeurs dans $\mathbf{N}$. On suppose que $G_X = G_Y$. A-t-on $X = Y$? A-t-on $X \sim Y$? A-t-on $X = Y$ presque sûrement?

Théorème 12.35

Soit X une variable aléatoire à valeur dans $\mathbf{N}$.

La variable aléatoire admet une espérance finie si et seulement si G_X est dérivable en 1.

Dans ce cas $G'_X(1) = \mathbf{E}(X)$.

Démonstration :

- $\Rightarrow$ On suppose que la variable aléatoire admet une espérance finie. Cela signifie que la série $\sum_{n \geq 0} n\mathbf{P}(X = n)$ est convergente. De ce fait, en posant encore $f_n : t \mapsto \mathbf{P}(X = n)t^n$ la série des dérivées $\sum_{n \geq 0} f'_n$ converge normalement donc uniformément sur $[0, 1]$ car $\|f'_n\|_{\infty, [0, 1]} = n\mathbf{P}(X = n)$. On peut donc appliquer le théorème de dérivation. La fonction G_X est dérivable en 1 et

$$G'_X(1) = \sum_{n=0}^{+\infty} n\mathbf{P}(X = n)1^{n-1} = \mathbf{E}(X).$$

- $\Leftarrow$ (Non exigible)

On suppose G_X dérivable en 1. Pour tout $t \in [0, 1[$, $G'_X(t) = \sum_{n=1}^{\infty} n\mathbf{P}(X = n)t^{n-1}$. Il est clair que la fonction $t \mapsto \sum_{n=0}^{\infty} n\mathbf{P}(X = n)t^{n-1}$ est une fonction croissante sur $[0, 1[$. Elle admet donc une limite (notée ℓ) dans $\mathbf{R}_+ \cup \{+\infty\}$ en 1^- . Comme G_X est dérivable en 1, ℓ est fini.

Maintenant, pour tout entier N et pour tout $t \in [0, 1[$,

$$\sum_{n=0}^N n\mathbf{P}(X = n)t^{n-1} \leq \sum_{n=0}^{+\infty} n\mathbf{P}(X = n)t^{n-1}$$

En faisant tendre t vers 1 on obtient que

$$\sum_{n=0}^N n\mathbf{P}(X = n) \leq \ell$$

Ceci étant vrai pour tout entier N , la série $\sum_{n \geq 0} n\mathbf{P}(X = n)$ converge.

□

Remarques :

1. De même si G_X est de classe $\mathcal{C}^2$ sur $[0, 1]$, par le théorème de transfert on obtient que

$$G''_X(1) = \mathbf{E}(X(X-1)) = \mathbf{E}(X^2) - \mathbf{E}(X)$$

On peut en déduire que

$$\mathbf{V}(X) = G''_X(1) + G'_X(1) - (G'_X(1))^2$$

2. On en déduit que si X n'est pas d'espérance finie alors le rayon de convergence de la série entière qui définit G_X est 1.

Proposition 12.36 (Fonction génératrice d'une somme)

1. Soit X et Y deux variables aléatoires **indépendantes** à valeurs dans $\mathbf{N}$.

$$\forall t \in [-1; 1], G_{X+Y}(t) = G_X(t)G_Y(t)$$

2. Soit $X_1, \dots, X_n$ des variables aléatoires indépendantes à valeurs dans $\mathbf{N}$,

$$\forall t \in [-1, 1], G_{X_1+\dots+X_n}(t) = \prod_{i=1}^n G_{X_i}(t)$$

Démonstration :

1. Pour $t \in [-1, 1]$, les variables e^{tX} et e^{tY} sont indépendantes d'après le lemme des coalitions. On en déduit que

$$G_{X+Y}(t) = \mathbf{E}(e^{t(X+Y)}) = \mathbf{E}(e^{tX}e^{tY}) = \mathbf{E}(e^{tX})\mathbf{E}(e^{tY}) = G_X(t)G_Y(t)$$

2. Le cas général se déduit par récurrence, du résultat précédent, en utilisant que pour tout $i \in \llbracket 2; n \rrbracket$, X_i est indépendante de $X_1 + \dots + X_{i-1}$.

□

4.2 Exemples

Nous allons établir les fonctions génératrices des lois usuelles.

ATTENTION

Il est bon de connaître les formules mais il faut surtout savoir les retrouver.

1. Loi de uniforme : Soit $X \sim \mathcal{U}(\{1, 2, \dots, N\})$. On a pour tout $k \in \llbracket 1; N \rrbracket$, $P(X = k) = \frac{1}{N}$ donc

$$G_X : t \mapsto \frac{1}{N} \sum_{k=1}^N t^k$$

Elle est définie sur $\mathbf{R}$. Si $t \neq 1$, on a $G_X(t) = t \cdot \frac{1 - t^N}{1 - t}$.

2. Loi de Bernoulli : Soit $X \sim \mathcal{B}(p)$ où $p \in [0, 1]$. On a $P(X = 0) = 1 - p$ et $P(X = 1) = p$ donc

$$G_X : t \mapsto 1 - p + pt = q + pt$$

Elle est définie sur $\mathbf{R}$.

3. Loi de binomiale : Soit $X \sim \mathcal{B}(n, p)$ où $n \in \mathbf{N}$ et $p \in [0, 1]$. On sait que l'on peut écrire X comme somme de n variables de Bernoulli indépendantes. On a donc

$$G_X : t \mapsto (1 - p + pt)^n = (q + pt)^n.$$

Elle est définie sur $\mathbf{R}$. Cela se retrouve aussi en utilisant le binôme de Newton.

4. Loi de géométrique : Soit $X \sim \mathcal{G}(p)$ où $p \in [0, 1[$. On calcule :

$$\begin{aligned} G_X(t) &= \sum_{k=1}^{+\infty} p q^{k-1} t^k \\ &= \frac{p}{q} \sum_{k=1}^{+\infty} (qt)^k \\ &= \frac{p}{q} \frac{qt}{1 - qt} = \frac{pt}{1 - qt} \end{aligned}$$

Elle est définie sur $\left] -\frac{1}{q}, \frac{1}{q} \right[$.

5. Loi de Poisson : Soit $X \sim \mathcal{P}(\lambda)$ ou $\lambda \in \mathbf{R}_+^*$.

$$G_X(t) = \sum_{k=0}^{+\infty} \frac{\lambda^k e^{-\lambda}}{k!} t^k = e^{-\lambda} e^{\lambda t} = e^{\lambda(t-1)}.$$

Elle est définie sur $\mathbf{R}$.

Espaces vectoriels normés II

1	Topologie d'un espace vectoriel normé	350
1.1	Ouverts	350
1.2	Voisinages	353
1.3	Fermés	354
1.4	Points intérieurs et intérieur	356
1.5	Points adhérents et adhérence	357
1.6	Caractérisation séquentielle des points adhérents et des fermés	359
1.7	Frontière	360
1.8	Densité	361
1.9	Invariance par changement de normes équivalentes	362
1.10	Topologie induite	363
2	Limite d'une application	364
2.1	Définition	365
2.2	Propriétés	367
2.3	Composition et opérations	369
3	Continuité	369
3.1	Définition	370
3.2	Caractérisation du caractère continu par les images réciproques	373
3.3	Applications uniformément continues	374
3.4	Applications lipschitziennes	375
3.5	Applications linéaires continues	377
4	Parties compactes d'un espace vectoriel normé	379
4.1	Définition	379
4.2	Applications continues sur une partie compacte	381
5	Espaces vectoriels de dimension finie	383
5.1	Equivalence des normes	383
5.2	Topologie des espaces vectoriels de dimension finie	385
5.3	Applications continues	386
6	Séries à valeurs dans un espace vectoriel de dimension finie	387
6.1	Généralités	387

6.2	Série géométrique de matrices	388
6.3	Série exponentielle de matrices	389
7	Parties connexes par arcs	392
7.1	Motivation	392
7.2	Définition	392
7.3	Image d'une partie connexe par arcs par une application continue	394

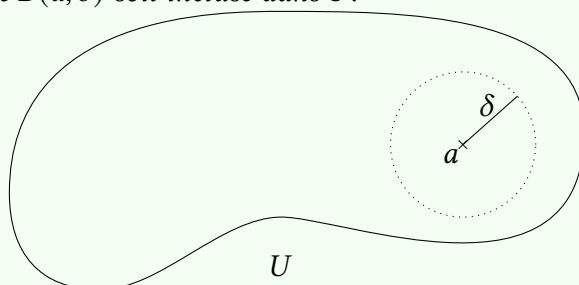
Nous allons continuer à étudier les espaces vectoriels normés. Dans ce chapitre $\mathbf{K}$ désigne $\mathbf{R}$ ou $\mathbf{C}$. Sans précisions autre, E désignera un espace vectoriel normé dont la norme sera notée $\|\cdot\|$.

1 Topologie d'un espace vectoriel normé

1.1 Ouverts

Définition 13.1 (Ouverts)

Soit U une partie de E . On dit que c'est un ouvert de E si pour tout a dans U il existe $\delta > 0$ tel que la boule ouverte $B(a, \delta)$ soit incluse dans U .



On peut dire que X est un ouvert ou que X est ouvert (sous-entendu un ensemble ouvert) voire ouverte (une partie ouverte).

Remarques :

1. La valeur de δ peut dépendre de a .
2. Cela signifie que si U est un ouvert, tout point « à côté » d'un point de U est encore dans U .

Exemples :

1. Dans $\mathbf{R}$, un intervalle ouvert est ouvert. Par exemple $U =]0, 1[$. En effet soit $a \in U$, on pose $\delta = \frac{\min(|a - 1|, |a|)}{2}$ et $B(a, \delta) =]a - \delta, a + \delta[\subset]0, 1[= U$.
2. Dans $\mathbf{R}$, l'ensemble $[0, 1]$ n'est pas un ouvert. En effet, pour $a = 0$ et pour $\delta > 0$, la boule $B(a, \delta)$ n'est pas inclus dans $]0, 1[$.
3. Les parties $\emptyset$ et E sont des ouverts. Ce sont les ouverts *triviaux*.
4. Soit $E = \mathbf{K}_n[X]$. On considère la norme définie par

$$\|\cdot\|_{\infty} : P = \sum_{k=0}^n a_k X^k \mapsto \|P\|_{\infty} = \max_k |a_k|$$

On pose $F = \{P \in E, P(1) > 0\}$. Montrons que F est un ouvert.

Soit $P \in F$. On a donc $P(1) = \sum_{k=0}^n a_k > 0$. On cherche donc $\delta > 0$ tel que pour tout polynôme $Q \in E$, $\|P - Q\|_\infty < \delta$ implique que $Q \in F$. Or $Q(1) = P(1) + (Q - P)(1)$.

On note $Q - P = \sum_{k=0}^n r_k X^k$. On voit $\|P - Q\|_\infty < \delta$ implique que pour tout $k \in \llbracket 0; n \rrbracket$, $|r_k| < \delta$ et donc $|(Q - P)(1)| \leq (n + 1)\delta$.

Finalement, en prenant $\delta = \frac{P(1)}{2(n+1)}$ on obtient que si $\|P - Q\|_\infty < \delta$ alors

$$Q(1) = P(1) + (Q - P)(1) \geq P(1) - (n + 1)\frac{P(1)}{2(n + 1)} = \frac{P(1)}{2} > 0$$

On en déduit donc que $Q \in F$.

5. Soit $E = \mathbf{K}[X]$. On considère encore la norme définie par $\|\cdot\|_\infty : P = \sum_{k=0}^d a_k X^k \mapsto \|P\|_\infty = \max_k |a_k|$.

Montrons cette fois que $F = \{P \in E, P(1) > 0\}$ n'est pas un ouvert.

On doit montrer qu'il existe $P \in F$ tel que pour tout $\delta > 0$, il existe $Q \in E$ tel que $\|Q - P\|_\infty < \delta$ et $Q \notin F$.

Prenons par exemple $P = X$ qui appartient à F car $P(1) = 1 > 0$. Pour tout $\delta > 0$, posons $\delta' \in]0, \delta[$. Il existe $n \in \mathbf{N}$ tel que $n\delta' > 1$. Il suffit alors de poser

$$Q = P - \sum_{k=0}^n \delta' X^k$$

On voit que $\|P - Q\|_\infty = \|\sum_{k=0}^n \delta' X^k\|_\infty = \delta' < \delta$ et que $Q(1) = P(1) - \sum_{k=0}^n \delta' \leq 1 - (n+1)\delta' < 0$. On en déduit que $Q \notin F$.

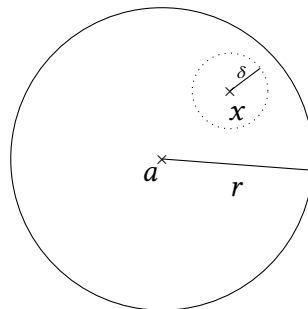
Proposition 13.2

Une boule ouverte $B(a, r)$ de E est un ouvert de E .

Démonstration : Soit $x \in B(a, r)$. On a donc $d(x, a) < r$. Soit ε tel que $d(x, a) + \varepsilon < r$. On a que $B(x, \varepsilon) \subset B(a, r)$.

En effet pour $y \in B(x, \varepsilon)$,

$$\|y - a\| \leq \|(y - x) + (x - a)\| \leq d(x, y) + d(x, a) \leq d(x, a) + \varepsilon \leq r.$$



□

Proposition 13.3

Notons $\mathcal{T}$ l'ensemble des ouverts de E .

1. Les ensembles $\emptyset$ et E appartiennent à $\mathcal{T}$
2. L'ensemble $\mathcal{T}$ est stable par intersection **finie**. C'est-à-dire qu'une intersection finie d'ouverts est encore un ouvert.
3. L'ensemble $\mathcal{T}$ est stable par union quelconque. C'est-à-dire qu'une union d'ouverts est encore un ouvert.

Remarques :

1. Un moyen mnémotechnique pour retenir que c'est une intersection finie est (IFO).
2. L'ensemble des ouverts n'est pas stable par intersection quelconque. Par exemple si on pose $I_n =]0, 1 + \frac{1}{n}[$ qui est un ouvert de $\mathbf{R}$. Alors $I = \bigcap_{n \in \mathbf{N}^*} I_n =]0, 1]$ qui n'est plus ouvert.
3. (Hors programme) L'ensemble $\mathcal{T}$ des ouverts de E est ce que l'on appelle la topologie de E .

Démonstration :

1. Déjà dit.
2. Soit $U_1, \dots, U_n$ des ouverts et $x \in \bigcap_{i=1}^n U_i$. On sait que comme tous les U_i sont des ouverts, pour chaque $i \in \llbracket 1; n \rrbracket$ il existe r_i tel que $B(x, r_i) \subset U_i$. Il ne reste plus qu'à poser $r = \min_{1 \leq i \leq n} r_i$ pour constater que $B(x, r) \subset \bigcap_{i=1}^n U_i$ et donc $\bigcap_{i=1}^n U_i$ est ouvert.
3. Soit $(U_i)_{i \in I}$ une famille d'ouverts. On considère $U = \bigcup_{i=1}^n U_i$. Soit $x \in U$. Il existe $i \in I$ tel que $x \in U_i$. Comme, pour ce i , U_i est ouvert, il existe $r > 0$ tel que $B(x, r) \subset U_i \subset U$. On a bien que U est ouvert.

□

Exemple : On considère $\mathbf{R}^2$ muni de la norme $\|\cdot\|_\infty$ (qui est la norme produit en voyant $\mathbf{R}^2 = \mathbf{R} \times \mathbf{R}$ et en prenant $|\cdot|$ comme norme sur $\mathbf{R}$). Pour $a < b$ et $c < d$, $X =]a, b[\times]c, d[$ est un ouvert.

Proposition 13.4

Soit $(E_1, N_1), \dots, (E_n, N_n)$ des espaces vectoriels normés. On considère l'espace vectoriel normé produit $(E_1 \times \dots \times E_n, N)$ où N est la norme produit. Pour tout $i \in \llbracket 1; n \rrbracket$, on considère un ouvert U_i de E_i . L'ensemble produit $X = U_1 \times \dots \times U_n$ est un ouvert de $(E_1 \times \dots \times E_n, N)$.

Démonstration : Soit $x = (x_1, \dots, x_n) \in X$. Pour tout $i \in \llbracket 1; n \rrbracket$, il existe $\delta_i > 0$ tel que $B_i(x_i, \delta_i) \subset U_i$. Posons $\delta = \min\{\delta_i, 1 \leq i \leq n\}$. Montrons que $B_N(x, \delta) \subset X$.

Soit $y = (y_1, \dots, y_n) \in B_N(x, \delta)$. Par définition, pour tout $i \in \llbracket 1; n \rrbracket$, $N_i(y_i - x_i) \leq N(y - x) < \delta \leq \delta_i$. On en déduit que $y \in X$.

On a bien montré que X était un ouvert de $(E_1 \times \dots \times E_n, N)$.

□

1.2 Voisinages

Définition 13.5

Soit $a \in E$. Un voisinage de a dans E est une partie X de E telle qu'il existe $\delta > 0$ tel que $B(a, \delta) \subset X$.

Remarques :

1. On peut aussi dire que X contient un ouvert contenant a .
2. Dans le cas où $E = \mathbf{R}$ on étend la notion de voisinage pour définir les voisinages de $+\infty$ et de $-\infty$. Un voisinage de $+\infty$ contient un intervalle de la forme $[M, +\infty[$ quand un voisinage de $-\infty$ contient un intervalle de la forme $] -\infty, m]$.

ATTENTION

Ne pas confondre les notions de voisinages et d'ouvert. Quand on parle d'un voisinage on parle toujours d'un voisinage **d'un point**. De fait, un ouvert est une partie de E qui est un voisinage de tous ses points.

Exemple : Dans $\mathbf{R}$, $[0, 1]$ est un voisinage de $\frac{1}{2}$. Par contre, ce n'est pas un voisinage de 1.

Proposition 13.6

Soit $a \in E$.

1. Soit $V_1, \dots, V_n$ une famille **finie** de voisinages de a dans E , l'intersection $\bigcap_{i=1}^n V_i$ est encore un voisinage de a .
2. Soit $(V_i)_{i \in I}$ une famille de voisinages de a dans E . L'union $\bigcup_{i \in I} V_i$ est encore un voisinage de a .

Démonstration : Il suffit de reprendre la démonstration de la proposition analogue dans le cas des ouverts.

Définition 13.7 (Propriété locale)

Soit P une propriété. On dit qu'elle est vraie au voisinage d'un point a de E , s'il existe un voisinage de a où elle est vérifiée. Il est équivalent de dire qu'il existe une boule ouverte contenant a (ou même centrée en a) où la propriété est vraie.

Exemple : La fonction $(x, y) \mapsto x^2 \cos(y)$ est bornée au voisinage de $a = (2, 1)$.

1.3 Fermés

Définition 13.8

Une partie X de E est un fermé si son complémentaire est ouvert.

On peut dire que X est un fermé ou que X est fermé (sous-entendu un ensemble fermé) voire fermée (une partie fermée).

Exemples :

1. Dans $E = \mathbb{R}$, un intervalle fermé $X = [a, b]$ est fermé. En effet son complémentaire $\complement_E X =]-\infty, a[\cup]b, +\infty[$ est ouvert.
2. Les parties $\emptyset$ et E sont fermées.

ATTENTION

La notion d'ouverts et de fermés ne sont pas contraire l'une de l'autre. Une partie peut-être ouverte et fermée (par exemple E en entier) comme elle peut être ni ouverte ni fermée (par exemple $]0, 1[$).

Proposition 13.9 (Propriétés des fermés)

1. Soit X une partie de E , X est un ouvert si et seulement si $\complement_E X$ est un fermé.
2. Soit $(X_1, \dots, X_n)$ une famille finie de fermés. L'union $\bigcup_{i=1}^n X_i$ est un fermé.
3. Soit $(X_i)_{i \in I}$ une famille de fermés. L'intersection $\bigcap_{i \in I} X_i$ est un fermé.

Démonstration :

1. Soit $X \subset E$,

$$\complement_E X \text{ fermé} \iff \complement_E (\complement_E X) \text{ ouvert} \iff X \text{ ouvert}$$

2. Par définition, pour tout $i \in \llbracket 1 ; n \rrbracket$, $\complement_E X_i$ sont des ouverts. On a alors

$$\complement_E \left(\bigcup_{1 \leq i \leq n} X_i \right) = \bigcap_{1 \leq i \leq n} \complement_E X_i$$

Donc $\complement_E \left(\bigcup_{1 \leq i \leq n} X_i \right)$ est un ouvert, d'où $\bigcup_{1 \leq i \leq n} X_i$ est un fermé.

3. Cela se démontre en passant au complémentaire comme ci-dessus.

□

Proposition 13.10

Soit $(E_1, N_1), \dots, (E_n, N_n)$ des espaces vectoriels normés. On considère l'espace vectoriel normé produit $(E_1 \times \dots \times E_n, N)$ où N est la norme produit. Pour tout $i \in \llbracket 1 ; n \rrbracket$, on considère un fermé F_i de E_i . L'ensemble produit $Y = F_1 \times \dots \times F_n$ est un fermé de $(E_1 \times \dots \times E_n, N)$.

On dit qu'un produit fini de fermés est fermé.

Démonstration : Montrons que $U = \complement_E Y$ est un ouvert.

Soit $x = (x_1, \dots, x_n) \in E$,

$$\begin{aligned} x \in \complement_E Y &\iff x \notin Y \\ &\iff \neg(\forall i \in \llbracket 1; n \rrbracket x_i \in F_i) \\ &\iff \exists i \in \llbracket 1; n \rrbracket x_i \notin F_i \end{aligned}$$

Pour tout $i \in \llbracket 1; n \rrbracket$ on note $U_i = \complement_{E_i} F_i$ et on pose $Z_i = E_1 \times \dots \times E_{i-1} \times U_i \times \dots \times E_n$. L'ensemble Z_i est l'ensemble des éléments $(x_1, \dots, x_n) \in E$ tel que $x_i \in U_i$ c'est-à-dire, $x_i \notin F_i$.

On a donc

$$\complement_E Y = \bigcup_{i=1}^n Z_i$$

Comme, pour tout $i \in \llbracket 1; n \rrbracket$, Z_i est un ouvert comme produit d'un nombre fini d'ouverts, $\bigcup_{i=1}^n Z_i$ est ouvert comme union finie d'ouverts et donc Y est un fermé. $\square$

Proposition 13.11

Les boules fermées et les sphères sont fermées.
En particulier les singletons qui sont les sphères de rayon 0 sont des fermés.

Démonstration :

- Commençons par le cas des boules fermées. Soit B la boule (fermée) de centre a et de rayon $r \geq 0$, son complémentaire $\complement_E B = \{x \in E \mid d(x, a) > r\}$ est un ouvert. En effet soit $x \in \complement_E B$, $d(x, a) > r$. Si on pose δ tel que $d(x, a) - \delta > r$ (on peut prendre par exemple $\delta = \frac{d(x, a) - r}{2}$) alors $B(x, \delta) \subset \complement_E B$. En effet, pour tout $y \in B(x, \delta)$,

$$\|y - a\| = \|(y - x) + (x - a)\| \geq \|x - a\| - \|y - x\| \geq d(x, a) - \delta > r.$$

On a bien montré que $\complement_E B$ était un ouvert et donc B est un fermé.

- Traitons le cas de la sphère fermée $S = S(a, r) = \{x \in E \mid d(x, a) = r\}$. Il suffit de remarquer que $S = \overline{B}(a, r) \cap \complement_E B(a, r)$. C'est donc l'intersection de deux fermés.¹

$\square$

Exemples :

1. Toute partie finie de E est un fermé car c'est une réunion finie de singleton qui sont des fermés.
2. Pour tout $n \in \mathbb{N}^*$, $I_n = [\frac{1}{n}, 1]$ est un fermé mais $\bigcup_{n \in \mathbb{N}^*} I_n =]0, 1]$ n'est pas fermé.

1. On peut aussi le faire à la main. On montre que $\complement_E S$ est un ouvert. Soit $x \in \complement_E S$. On a donc $d(a, x) \neq r$.

– Si on est dans le cas où $d(a, x) > r$ l'argument ci-dessus est encore valable.

– Si on est dans le cas où $d(a, x) < r$ on a un argument analogue.

On a donc que $\complement_E S$ est un ouvert et donc S est un fermé.

1.4 Points intérieurs et intérieur

Définition 13.12 (Points intérieurs)

Soit X une partie de E .

1. Un point a de X est dit intérieur à X s'il existe $\delta > 0$ tel que $B(a, \delta) \subset X$.
2. L'ensemble des points intérieurs de X s'appelle l'intérieur de X . Il se note $\overset{\circ}{X}$.

Par définition, $\overset{\circ}{X} \subset X$.

Remarque : Un point a de X est donc un point intérieur à X si X est un voisinage de a .

Exemples :

1. Soit $X =]0, 1]$. Les points a appartenant à $]0, 1[$ sont des points intérieurs à X . Par contre $a = 1$ n'est pas un point intérieur. On a donc $\overset{\circ}{X} =]0, 1[$.
2. Soit B la boule fermée de centre a et de rayon r . Ses points intérieurs sont ceux qui sont dans la boule ouverte de centre a et de rayon r . C'est-à-dire $\overset{\circ}{B}(a, r) = B(a, r)$. En effet,
 - si $x \in B(a, r)$, comme $B(a, r)$ est un ouvert, il existe δ tel que

$$B(x, \delta) \subset B(a, r) \subset B.$$

- si $x \notin B(a, r)$ on a donc $d(a, x) = r$. Maintenant pour tout $\delta > 0$, on pose

$$y = x + \delta \frac{x - a}{2\|x - a\|}$$

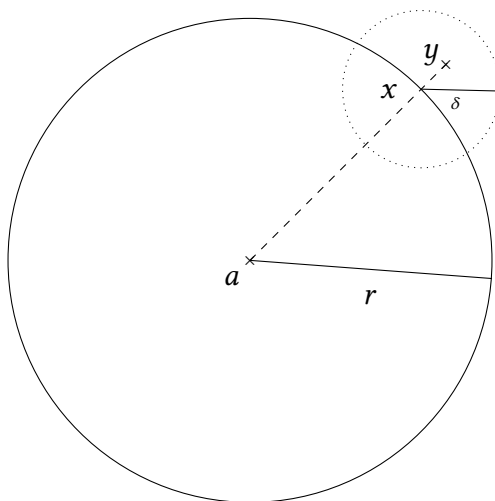
On a donc

$$\|y - x\| = \frac{\delta}{2}$$

donc $y \in B(x, \delta)$ mais

$$y - a = y - x + x - a = \left(1 + \frac{\delta}{2\|x - a\|}\right) \|x - a\| > \|x - a\| = r.$$

Donc $y \notin B$. On en déduit que x n'est pas intérieur.



Proposition 13.13

Soit X une partie de E .

1. Son intérieur $\overset{\circ}{X}$ est ouvert et c'est le plus grand ouvert contenu dans X .
2. On a donc : X ouvert $\Leftrightarrow \overset{\circ}{X} = X$.

Démonstration : Il est clair que $\overset{\circ}{X}$ est inclus dans X . Il reste à montrer que $\overset{\circ}{X}$ est un ouvert et que tout ouvert inclus dans X est inclus dans $\overset{\circ}{X}$.

- Soit $U \subset X$ un ouvert inclus dans X , alors $U \subset \overset{\circ}{X}$. En effet soit $x \in U$, comme U est ouvert il existe une boule $B(x, \delta)$ centrée en x qui est inclus dans U et donc dans X . De ce fait, x est un point intérieur à X . On a bien $U \subset \overset{\circ}{X}$.
- Pour montrer que $\overset{\circ}{X}$ est ouvert on va montrer que

$$\overset{\circ}{X} = \bigcup_{\substack{U \subset X \\ U \text{ ouvert}}} U$$

c'est-à-dire que $\overset{\circ}{X}$ est l'union de tous les ouverts contenu dans X . Posons, $Y = \bigcup_{\substack{U \subset X \\ U \text{ ouvert}}} U$, on

vient de voir que $Y \subset \overset{\circ}{X}$.

Soit $x \in \overset{\circ}{X}$, par définition, il existe $\delta > 0$ tel que $B(x, \delta) \subset X$. En particulier, $B(x, \delta)$ est un ouvert contenu dans X . On en déduit que

$$x \in B(x, \delta) \subset Y$$

Cela montre que $x \in Y$.

On a bien montré que $\overset{\circ}{X}$ est le plus grand ouvert contenant X . □

Exercices :

1. Déterminer $\overset{\circ}{Q}$.
2. Déterminer l'intérieur d'un sous-espace vectoriel strict.

1.5 Points adhérents et adhérence**Définition 13.14**

Soit X une partie de E .

1. Un point a de E est un point adhérent à X si pour tout $\delta > 0$ la boule ouverte centrée en a de rayon δ rencontre X ($X \cap B(a, \delta) \neq \emptyset$).
2. L'ensemble des points adhérents à X s'appelle l'adhérence de X . Il se note $\overline{X}$.

Remarque : Tout point a de X est adhérent à X car $a \in X \cap B(a, \delta)$. On en déduit que $X \subset \overline{X}$.

Exemples :

1. On considère $X =]0, 1]$. Tous les éléments de X sont adhérents à X . De plus 0 est adhérent à X et si $x \notin [0, 1]$ alors x n'est pas adhérent à X . L'ensemble des points adhérents est donc $[0, 1]$.
2. Soit $X = B(a, r)$.
 - Tous les points de $B(a, r)$ sont adhérents à X .

- Soit x tel que $d(a, x) > r$. Le point x n'est pas adhérent à X . En effet, on considère $\delta > 0$ tel que $d(a, x) > \delta + r$. Pour $y \in B(x, \delta)$,

$$d(a, y) \geq d(a, x) - d(x, y) > d(a, x) - r > \delta$$

Cela montre que $B(x, \delta) \cap B(a, r) = \emptyset$ et donc que $x \notin \bar{X}$.

- Les points de la sphère $S(a, r)$ sont adhérents à X . Soit x tel que $\|x - a\| = r$ et $\delta > 0$, il existe un élément y appartenant à $B(a, r) \cap B(x, \delta)$. Quand $\delta < r$, (le cas $\delta \geq r$ est évident) il suffit de prendre

$$y = x + \delta \frac{a - x}{2\|a - x\|}.$$

Il est clair que $y \in B(x, \delta)$ par construction, de plus $\|y - a\| = \left(1 - \frac{\delta}{2\|x - a\|}\right) \|x - a\| < r$.

On a montré que $\overline{B(a, r)} = \bar{B}(a, r)$.

Proposition 13.15

Soit X une partie de E ,

$$\mathbb{C}_E \dot{X} = \overline{\mathbb{C}_E X} \quad \text{et} \quad \mathring{\mathbb{C}_E X} = \mathbb{C}_E \bar{X}.$$

Démonstration :

Soit $a \in E$,

$$\begin{aligned} a \in \mathbb{C}_E \dot{X} &\iff a \notin \dot{X} \\ &\iff \forall \delta > 0, B(a, \delta) \not\subset X \\ &\iff \forall \delta > 0, \exists x \in B(a, \delta) \cap \mathbb{C}_E X \\ &\iff \forall \delta > 0, B(a, \delta) \cap \mathbb{C}_E X \neq \emptyset \\ &\iff a \in \overline{\mathbb{C}_E X} \end{aligned}$$

Pour la deuxième partie, il suffit d'appliquer l'assertion démontrée à $\mathbb{C}_E X$ et de prendre le complémentaire.

□

Proposition 13.16

1. L'adhérence $\bar{X}$ est le plus petit fermé contenant dans X .
2. On a donc : X est fermé $\iff X = \bar{X}$.

Démonstration : D'après ce qui précède, $\mathbb{C}_E \bar{X} = \mathring{\mathbb{C}_E X}$ qui est donc un ouvert d'où $\bar{X}$ est un fermé. De plus $X \subset \bar{X}$; on a donc obtenu que $\bar{X}$ est un fermé contenant X .

Considérons maintenant F un fermé contenant X . En prenant le complémentaire on en déduit que $\mathbb{C}_E F$ est un ouvert contenu dans $\mathbb{C}_E X$ (car le complémentaire est décroissant). On en déduit que $\mathbb{C}_E F \subset \mathring{\mathbb{C}_E X}$.

En reprenant le complémentaire,

$$\overline{X} = \overline{\mathbb{C}_E(\mathbb{C}_E X)} = \mathbb{C}_E \left(\overset{\circ}{\mathbb{C}_E X} \right) \subset \mathbb{C}_E(\mathbb{C}_E F) = F$$

On a bien montré que $\overline{X}$ est le plus petit fermé contenant X . □

Remarque : On peut refaire une preuve « à la main » dans l'idée de celle de la proposition analogue pour l'intérieur.

Exercices :

1. Déterminer $\overline{Q}$.
2. Montrer que l'adhérence et l'intérieur sont croissant

1.6 Caractérisation séquentielle des points adhérents et des fermés

La définition des fermés comme complémentaire des ouverts n'est pas la plus aisée pour montrer qu'un ensemble est fermé. Nous allons mettre en évidence, une méthode pour montrer qu'un point est un point d'adhérent à l'aide de suites. On en déduira une caractérisation des fermés.

Théorème 13.17 (Caractérisation séquentielle des points adhérents et des fermés)

Soit X une partie de E .

1. Soit $a \in E$, il est adhérent à X (c'est-à-dire $a \in \overline{X}$) si et seulement si il existe une suite (x_n) d'éléments de X qui converge vers a .
2. L'ensemble X est fermé si et seulement si toute suite (x_n) d'éléments de X qui converge (dans E) a sa limite dans X .

Remarques :

1. La première assertion illustre bien qu'un point adhérent à X est bien un point qui appartient à X ou qui est juste « à côté » de X .
2. La deuxième assertion illustre la notion de fermé. Cela signifie que l'on ne peut pas « sortir » de X .
3. C'est souvent cette caractérisation que l'on utilise pour montrer qu'un ensemble est fermé.

Démonstration :

1. — $\Rightarrow$ On suppose que a est adhérent à X . De ce fait pour tout entier n , $B(a, \frac{1}{n+1}) \cap X \neq \emptyset$. Il existe donc un élément $x_n \in B(a, \frac{1}{n+1}) \cap X$. La suite (x_n) est bien une suite d'éléments de X . De plus, par construction, $\|x_n - a\| \leq \frac{1}{n+1}$ donc $(x_n) \rightarrow a$.

Notons que l'on peut faire la même preuve avec les boules $B(a, \delta_n)$ où (δ_n) est une suite qui tend vers 0.

- $\Leftarrow$ On suppose qu'il existe une suite $(x_n) \in X^{\mathbb{N}}$ telle que $(x_n) \rightarrow a$. Montrons que a est alors un point adhérent de X . Soit $\delta > 0$, $B(a, \delta)$ rencontre X . En effet il existe un terme de la suite (x_n) tel que $\|x_n - a\| < \delta$.
2. — $\Rightarrow$ On veut montrer que si X est fermé alors toute suite à valeurs dans X convergente (dans E) a sa limite dans X . Soit $(x_n) \in X^{\mathbb{N}}$ et $a = \lim(x_n)$. D'après l'énoncé précédent a est adhérent à X et donc $a \in X$ car X est fermé ($X = \overline{X}$).
- $\Leftarrow$ Soit $a \in \overline{X}$. D'après ce qui précède, il existe une suite (x_n) d'éléments de X tels que $(x_n) \rightarrow a$. Par hypothèse, on obtient donc que $a \in X$. Finalement, $X = \overline{X}$ ce qui montre que X est fermé.

□

Proposition 13.18

Soit E un espace vectoriel et F un sous-espace vectoriel. S'il est de dimension finie alors il est fermé.

Démonstration : Considérons une base $(e_1, \dots, e_p)$ de F . Soit (x_n) une suite d'éléments de F qui converge vers un élément x dans E . Supposons par l'absurde que $x \notin E$ et considérons $F' = F \oplus \text{Vect}(x)$; on note alors $e_{p+1} = x$ car $\mathcal{B} = (e_1, \dots, e_p, x)$ est une base de F' . On note pour tout $i \in \llbracket 1; p+1 \rrbracket$, $e_i^* : F' \rightarrow \mathbf{K}$ la forme linéaire coordonnée. La suite (x_n) est une suite convergente de F' qui est de dimension finie. Cela implique que pour tout $i \in \llbracket 1; p+1 \rrbracket$, $(e_i^*(x_n)) \rightarrow e_i^*(x)$ (ce résultat a été démontré dans le cas où la norme de F' est la norme infinie relativement à la base $(e_1, \dots, e_{p+1})$ et est vrai pour toute norme de E induit sur F' une norme équivalente à $\|\cdot\|_{\infty, \mathcal{B}}$).

C'est absurde car pour tout entier n , $e_{p+1}^*(x_n) = 0$ et $e_{p+1}^*(x) = 1$. □

ATTENTION

Cela n'est plus vrai si F n'est pas de dimension finie. Considérons par exemple $E = \mathcal{C}^0([0, 1], \mathbf{R})$ muni de la norme $\|\cdot\|_1$. Soit $F = \{f \in E, f(0) = 0\}$ qui est un sous-espace vectoriel de E .

Pour tout entier naturel n non nul on pose $f_n : x \mapsto x^{1/n}$ qui appartient à F . La suite (f_n) converge dans E vers la fonction constante égale à 1 notée g car pour tout entier n ,

$$\|f_n - g\|_1 = \int_0^1 |1 - t^{1/n}| dt = \int_0^1 1 - t^{1/n} dt = 1 - \left[\frac{t^{\frac{n+1}{n}}}{\frac{n+1}{n}} \right]_0^1 = 1 - \frac{n}{n+1} = \frac{1}{n+1}$$

Cela montre que $\|f_n - g\|_1 \xrightarrow{n \rightarrow +\infty} 0$. Par contre $g \notin F$.

Exercice : On note $\text{St}_n(\mathbf{R})$ l'ensemble des matrices stochastiques, c'est-à-dire les matrices $M = (m_{ij}) \in \mathcal{M}_n(\mathbf{R})$ telles que

1. $\forall (i, j) \in \llbracket 1; n \rrbracket^2, m_{ij} \geq 0$
2. $\forall j \in \llbracket 1; n \rrbracket, \sum_{i=1}^n m_{ij} = 1$.

Montrer que $\text{St}_n(\mathbf{R})$ est un fermé de $(\mathcal{M}_n(\mathbf{R}), \|\cdot\|_\infty)$

- En montrant que son complémentaire est ouvert.
- En utilisant la caractérisation séquentielle.

1.7 Frontière

Pour toute partie X de E , on a vu que $\overset{\circ}{X} \subset X \subset \overline{X}$. La frontière d'une partie est la différence entre son adhérence et son intérieur.

Définition 13.19

Soit $X \subset E$, on appelle frontière de X et on note $\partial X = \overline{X} \setminus \overset{\circ}{X}$.

Remarque : Comme $\partial X = \overline{X} \cap \mathring{C}_E \dot{X}$ c'est un fermé comme intersection de deux fermés.

Exemples :

1. Si X est une boule (ouverte ou fermée) de centre a et de rayon r alors $\overline{X}$ est la boule fermée et $\mathring{X}$ est la boule ouverte donc ∂X est la sphère $S(a, r)$.
2. Calculons $\partial]0, 1]$.

On a vu que $\overline{]0, 1]} = [0, 1]$ et $\mathring{]0, 1]} =]0, 1[$ donc $\partial]0, 1] = \{0, 1\}$.

3. Calculons $\partial \mathbb{Q}$.

On a vu que $\overline{\mathbb{Q}} = \mathbb{R}$ et $\mathring{\mathbb{Q}} = \emptyset$ donc $\partial \mathbb{Q} = \mathbb{R}$.

1.8 Densité

La notion de parties denses vues en première année dans le cas où $E = \mathbb{R}$ s'étend au cadre général des espaces vectoriels normés.

Définition 13.20 (Densité)

Soit X une partie de E . Elle est dit dense (dans E) si $\overline{X} = E$. C'est-à-dire que tout point de E est adhérent à X .

Proposition 13.21

Soit X une partie de E . Les assertions suivantes sont équivalentes

- i) X est dense
- ii) Toute boule ouverte $B(a, \delta)$ où $\delta > 0$ rencontre X .
- iii) Pour tout point a de E il existe une suite (x_n) d'éléments de X convergeant vers a .

Démonstration : Il suffit d'utiliser les propositions vues sur les points adhérents. $\square$

Exemples :

1. Dans le cours de première année on a vu que $\mathbb{Q}$ l'ensemble des nombres rationnels et $\mathbb{C}_\mathbb{R} \mathbb{Q}$ l'ensemble des nombres irrationnels sont dans $\mathbb{R}$. Il suffit d'approcher tout nombre par un décimal ou par $\pi + r$ où r est un décimal.
2. Montrons que $\text{GL}_n(\mathbb{C})$ est dense dans $\mathcal{M}_n(\mathbb{C})$.

Soit $A \in \mathcal{M}_n(\mathbb{C})$. Son spectre $\text{Sp}(A)$ est un ensemble fini. On considère

$$\delta = \min\{|\lambda|, \lambda \in \text{Sp}(A) \setminus \{0\}\}$$

On convient que $\delta = 1$ si A n'a pas de valeurs propres non nulles. Avec ces notations, pour tout $x \in]0, \delta[$, la matrice $A - xI_n$ est inversible car $x \notin \text{Sp}(A)$. On pose donc pour tout $p \geq 1$, $A_p = A - \frac{\delta}{2^p} I_n$. Par construction, pour tout $p \geq 1$, $A_p \in \text{GL}_n(\mathbb{C})$ et $(A_p) \xrightarrow[n \rightarrow +\infty]{} A$ de manière évidente.

Cela montre bien que $\text{GL}_n(\mathbb{C})$ est dense dans $\mathcal{M}_n(\mathbb{C})$.

Théorème 13.22 (Théorème de Weierstrass)

Soit $[a, b]$ un segment. On considère $E = \mathcal{C}^0([a, b], \mathbf{K})$ l'ensemble des fonctions continues sur $[a, b]$ muni de la norme infinie $\|\cdot\|_\infty$.

On note $\mathcal{P}$ l'ensemble des fonctions polynomiales sur $[a, b]$. Il est dense dans E .

C'est-à-dire que pour toute fonction $f \in \mathcal{C}^0([a, b], \mathbf{K})$, il existe une suite $(P_n)_{n \geq 0} \in \mathcal{P}^{\mathbf{N}}$ telle que $(P_n) \xrightarrow{\|\cdot\|_\infty} f$.

Remarque : On rappelle que le fait que (P_n) converge pour la norme $\|\cdot\|_\infty$ vers f signifie que (P_n) converge uniformément vers f .

Exemple : Avec les notations du théorème. Soit $f \in E$ telle que pour tout $k \in \mathbf{N}$, $\int_a^b f(t)t^k dt = 0$. Montrons que $f = \tilde{0}$.

1.9 Invariance par changement de normes équivalentes

On a déjà vu que si on remplace la norme $\|\cdot\|$ par une norme N qui lui est équivalente, on ne modifie pas le caractère convergent des suites. Nous allons voir que cela ne modifie pas toutes les propriétés topologiques définies précédemment.

Théorème 13.23

Soit $\|\cdot\|$ et N deux normes équivalentes sur E et X une partie de E .

1. La partie X est ouverte pour $\|\cdot\|$ si et seulement si elle est ouverte pour N .
2. La partie X est fermée pour $\|\cdot\|$ si et seulement si elle est fermée pour N .
3. L'adhérence (resp. l'intérieur) de X pour $\|\cdot\|$ est l'adhérence (resp. l'intérieur) pour N .
4. La partie X est dense pour $\|\cdot\|$ si et seulement si elle est dense pour N .

Démonstration :

1. Soit X une partie ouverte pour N . Montrons qu'elle est ouverte pour $\|\cdot\|$. Soit $a \in X$, comme X est ouverte pour N , il existe $\delta > 0$ telle que $B_N(a, \delta) \subset X$. On sait qu'il existe $C > 0$ tel que $N \leq C\|\cdot\|$ donc $B_{\|\cdot\|}(a, \frac{\delta}{C}) \subset B_N(a, \delta)$ donc il existe $\delta' = \frac{\delta}{C}$ tel que $B_{\|\cdot\|}(a, \delta') \subset B_N(a, \delta) \subset X$. On en déduit que X est ouvert pour $\|\cdot\|$.

L'implication inverse est identique par symétrie en utilisant qu'il existe $C' > 0$ tel que $\|\cdot\| \leq C'N$

2. Cela découle de 1. en passant au complémentaire.
3. Pour l'adhérence (resp. l'intérieur) il suffit d'utiliser que c'est le plus petit fermé contenant X (resp. le plus grand ouvert contenu dans X).
4. Cela découle de l'invariance de l'adhérence.

□

ATTENTION

Ce résultat n'est plus vrai pour des normes qui ne sont pas équivalentes. Si on considère $E = \mathcal{C}^0([0, 1], \mathbf{R})$ et $X = \{f \in E \mid f(0) = 0\}$.

- On peut voir que X est un fermé pour $\|\cdot\|_\infty$. En effet soit (f_n) une suite de fonctions qui tend vers f pour $\|\cdot\|_\infty$ alors, en particulier, $(f_n(0))$ tend vers $(f(0))$ et donc $f \in X$.
- Pour $\|\cdot\|_1$, X n'est plus fermé. Il est même dense dans E , on peut « approcher » toute fonction f de E par une fonction g de X telle que $\|f - g\|_1 \leq \varepsilon$.

1.10 Topologie induite

Soit A une partie de E les notions d'ouverts (et de fermés) de E permettent de définir des ouverts (et des fermés) de A .

Exemple : Si on considère $E = \mathbf{R}^2$ et A la boule fermée unité. On veut dire que si $H = \{(x, y) \in \mathbf{R}^2 \mid x > 0\}$ alors $H \cap A$ est un ouvert. En effet, si on oublie ce qui se passe « en dehors » de A , si $\alpha \in H \cap A$ alors tous les points « proches » de α sont dans $H \cap A$.

Définition 13.24

Soit X une partie de A . On dit que X est un ouvert relatif de A si pour tout x de X , il existe $\delta > 0$ tel que $B(x, \delta) \cap A \subset X$.

ATTENTION

Il est important de voir que l'on ne demande pas que l'intégralité de la boule soit dans X mais juste $B(x, \delta) \cap A$ que l'on appelle sa *trace* sur A .

Définition 13.25

1. Soit X une partie de A . On dit que c'est un fermé relatif de A si $\complement_A X$ est un ouvert relatif de A .
2. Soit X une partie de A et $x \in X$. On dit que X est un voisinage relatif de x dans A s'il existe $\delta > 0$ tel que $B(x, \delta) \cap A \subset X$.

Exemples :

1. Pour $A = \mathbf{R}_+^*$, l'intervalle $X =]0, 1]$ est un fermé relatif car son complémentaire (dans A) $] -\infty, 0[\cup]1, +\infty[$ est un ouvert relatif (car c'est un ouvert de $\mathbf{R}$).
2. Dans $A = \mathbf{N}$, tous les singletons sont des fermés relatifs (car ce sont des fermés). Ce sont aussi des ouverts relatifs pour A car pour $n \in \mathbf{N}$, $B(n, \frac{1}{2}) \cap A \subset \{n\}$. De fait, toutes les parties de A sont des ouverts relatifs (et donc des fermés relatifs). On dit alors que A est une partie discrète.

Proposition 13.26 (Caractérisation séquentielle des fermés relatifs)

Soit X une partie de A . C'est un fermé relatif de A si et seulement pour toute suite (x_n) d'éléments de X qui converge dans A a sa limite dans X .

Démonstration : Il suffit de reprendre le principe de la démonstration de la caractérisation séquentielle.

- $\Rightarrow$ On suppose que X est un fermé relatif à A donc $U = \complement_A X$ est un ouvert relatif à A . Maintenant si (x_n) est une suite d'éléments de X qui converge dans A . On note ℓ sa limite. Supposons par l'absurde qu'elle n'est pas dans X , elle est donc dans U qui est ouvert. On peut donc trouver $\delta > 0$ tel que $B(\ell, \delta) \cap A \subset U$ ce qui est absurde car cela implique que $\|x_n - \ell\| \geq \delta$.
- $\Leftarrow$ Par contraposée, montrons que si X n'est pas un fermé relatif (donc si $U = \complement_A X$ n'est pas un ouvert relatif à A) alors on peut trouver une suite (x_n) d'éléments de X convergeant dans A mais pas dans X (c'est-à-dire dont la limite appartient à U). Il suffit d'utiliser que, comme U n'est pas ouvert, il existe $\ell \in U$ tel que pour tout n , $B(\ell, \frac{1}{n+1}) \cap A$ contient un élément qui n'est pas dans U donc qui appartient à X .

□

Exemple : On peut montrer que $X =]0, \frac{1}{2}]$ est un fermé de $]0, 1]$. En effet soit (x_n) une suite de $]0, \frac{1}{2}]$ qui a une limite dans $]0, 1]$ alors sa limite appartient à X .

Proposition 13.27

Soit A une partie de E .

1. La partie X est un ouvert relatif de A si et seulement s'il existe un ouvert U de E tel que $X = U \cap A$. On dit que X est la trace de U .
2. La partie X est un fermé relatif de A si et seulement s'il existe un fermé F de E tel que $X = F \cap A$. On dit que X est la trace de F .

Démonstration :

1. – $\Rightarrow$ On suppose que X est un ouvert relatif. On sait alors que pour tout x de X , il existe δ_x tel que $B(x, \delta_x) \cap A \subset X$. Posons

$$U = \bigcup_{x \in X} B(x, \delta_x)$$

Il est clair que U est un ouvert (de E) comme réunion d'ouverts. De plus, $X \subset U$ par construction, comme $X \subset A$ on a bien $X \subset U \cap A$. Inversement,

$$U \cap A = \bigcup_{x \in X} (B(x, \delta_x) \cap A) \subset X.$$

- $\Leftarrow$ Soit U un ouvert de E tel que $X = U \cap A$. Soit $x \in X$, il existe $\delta > 0$ tel que $B(x, \delta) \subset U$ donc $B(x, \delta) \cap A \subset U \cap A = X$.
2. Découle de 1. en passant au complémentaire.

2 Limite d'une application

2.1 Définition

Nous allons dans ce chapitre généraliser la notion de limite au cas des fonctions f d'un espace vectoriel normé $(E, \|\cdot\|_E)$ dans un autre $(F, \|\cdot\|_F)$. Cela généralise les cas déjà étudié en première année où $E = \mathbf{R}$ et $F = \mathbf{R}$ ou $\mathbf{C}$.

Définition 13.28

Soit f une fonction définie sur une partie A d'un espace vectoriel normé $(E, \|\cdot\|_E)$ à valeurs dans un autre $(F, \|\cdot\|_F)$. Soit a un point adhérent à A et $\ell \in F$. On dit que f tend vers ℓ quand x tend vers a si

$$\forall \varepsilon > 0, \exists \eta > 0, \forall x \in A, \|x - a\|_E \leq \eta \Rightarrow \|f(x) - \ell\|_F \leq \varepsilon$$

On note alors

$$f(x) \xrightarrow{x \rightarrow a} \ell$$

Remarques :

1. C'est la même définition que celle de la limite d'une fonction définie de $\mathbf{R}$ dans $\mathbf{R}$. Il faut juste remplacer la valeur absolue par la norme.
2. Dans le cas des fonctions de $\mathbf{R}$ dans $\mathbf{R}$, la définition de la limite peut aussi s'écrire avec des intervalles. Là on peut l'écrire avec des boules.

$$(f(x) \xrightarrow{x \rightarrow a} \ell) \iff (\forall \varepsilon > 0, \exists \eta > 0, \forall x \in A, x \in \overline{B}(a, \eta) \Rightarrow f(x) \in \overline{B}(\ell, \varepsilon))$$

3. On peut aussi écrire cette définition avec des inégalités strictes (ou des boules ouvertes). Notons (A) et (B) les assertions

- (A) : « $\forall \varepsilon > 0, \exists \eta > 0, \forall x \in A, \|x - a\|_E \leq \eta \Rightarrow \|f(x) - \ell\|_F \leq \varepsilon$ »
- (B) : « $\forall \varepsilon > 0, \exists \eta > 0, \forall x \in A, \|x - a\|_E < \eta \Rightarrow \|f(x) - \ell\|_F < \varepsilon$ »

Vérifions qu'une fonction vérifie (A) si et seulement si elle vérifie (B)

- $\boxed{(A) \Rightarrow (B)}$ Soit $\varepsilon > 0$. On peut utiliser l'assertion (A) pour $\frac{\varepsilon}{2}$. Il existe donc $\eta > 0$ tel que pour tout $x \in A$, $\|x - a\|_E \leq \eta$ implique que $\|f(x) - \ell\|_F \leq \frac{\varepsilon}{2}$. On en déduit que si $\|x - a\|_E < \eta \leq \eta$ alors $\|f(x) - \ell\|_F \leq \frac{\varepsilon}{2} < \varepsilon$.
 - $\boxed{(B) \Rightarrow (A)}$ Soit $\varepsilon > 0$. On peut utiliser l'assertion (B) pour ε . Il existe donc $\eta_1 > 0$ tel que pour tout $x \in A$, $\|x - a\|_E < \eta_1$ implique que $\|f(x) - \ell\|_F < \varepsilon$. Pour $\eta = \frac{\eta_1}{2}$ on a que si $\|x - a\|_E \leq \eta = \frac{\eta_1}{2} < \eta_1$ alors $\|f(x) - \ell\|_F < \varepsilon \leq \varepsilon$.
4. On peut reformuler cette définition avec des voisinages relatifs. En effet cela peut s'écrire que pour tout voisinage $\mathcal{V}$ de ℓ il existe un voisinage relatif $\mathcal{U}$ de a dans A tel que $f(\mathcal{U}) \subset \mathcal{V}$.

Définition 13.29 (Extension aux cas où $a = \infty$ et $\ell = \infty$)

Soit f une fonction définie sur une partie A d'un espace vectoriel normé $(E, \|\cdot\|_E)$ à valeurs dans un autre $(F, \|\cdot\|_F)$.

1. On suppose que A n'est pas borné. On dit que f tend vers $\ell \in F$ quand $\|x\|_E \rightarrow \infty$ si

$$\forall \varepsilon > 0, \exists M \in \mathbf{R}, \forall x \in A, \|x\|_E \geq M \Rightarrow \|f(x) - \ell\|_F \leq \varepsilon.$$

On note alors $f(x) \xrightarrow{\|x\|_E \rightarrow \infty} \ell$.

Noter que c'est $\|x\|_E \rightarrow \infty$ et pas $x \rightarrow \infty$.

2. Si $F = \mathbf{R}$. Soit a un point adhérent à A . On dit que f tend vers $+\infty$ quand $x \rightarrow a$ si

$$\forall M \in \mathbf{R}_+, \exists \eta > 0, \forall x \in A, \|x - a\|_E \leq \eta \Rightarrow f(x) \geq M.$$

On note alors $f(x) \xrightarrow{x \rightarrow a} +\infty$.

On définit de même le fait que f tende vers $-\infty$ en a .

3. On suppose que $E = \mathbf{R}$ et que A n'est pas majoré. On dit que f tend vers $\ell \in F$ quand $x \rightarrow +\infty$ si

$$\forall \varepsilon > 0, \exists M \in \mathbf{R}, \forall x \in A, x \geq M \Rightarrow \|f(x) - \ell\|_F \leq \varepsilon$$

On note alors $f(x) \xrightarrow{x \rightarrow +\infty} \ell$.

Dans le cas où A n'est pas minoré, on définit de même la limite de f en $-\infty$.

Proposition 13.30 (Unicité de la limite)

Soit f une fonction définie sur une partie A d'un espace vectoriel normé $(E, \|\cdot\|_E)$ à valeurs dans un autre $(F, \|\cdot\|_F)$. Soit a un point adhérent à A . On suppose qu'il existe ℓ et ℓ' dans F tels que $f(x) \xrightarrow{x \rightarrow a} \ell$ et $f(x) \xrightarrow{x \rightarrow a} \ell'$. On a alors $\ell = \ell'$.

Démonstration : Il suffit de recopier la démonstration usuelle. □

Définition 13.31 (Limite)

Soit f une fonction définie sur une partie A d'un espace vectoriel normé $(E, \|\cdot\|_E)$ à valeurs dans un autre $(F, \|\cdot\|_F)$. Soit a un point adhérent à A et $\ell \in F$. Si f tend vers ℓ quand x tend vers a on dit que ℓ est la limite de f en a et on note

$$\lim_{x \rightarrow a} f(x) = \ell \text{ ou } \lim_a f = \ell$$

2.2 Propriétés

Proposition 13.32

Soit f une fonction de A dans F . Soit a adhérent à A et $\ell \in F$.

$$\left(\lim_a f = \ell \right) \iff \left(\lim_{x \rightarrow a} \|f(x) - \ell\|_F = 0 \right)$$

Démonstration : Évident □

Exemple : Soit $A \in \mathcal{M}_n(\mathbf{R})$. Soit $f : \mathcal{M}_n(\mathbf{R}) \rightarrow \mathcal{M}_n(\mathbf{R})$ définie par $f : M \mapsto AM$.

Montrons que $\lim_{M \rightarrow I_n} f(M) = A$.

Prenons la norme $\|\cdot\|_\infty$ sur $\mathcal{M}_n(\mathbf{R})$. On sait que si M, N sont deux matrices de $\mathcal{M}_n(\mathbf{R})$, $\|MN\|_\infty \leq n\|M\|_\infty\|N\|_\infty$. On a alors pour $M \in \mathcal{M}_n(\mathbf{R})$,

$$\|f(M) - A\|_\infty = \|A(M - I_n)\|_\infty \leq n\|A\|_\infty\|M - I_n\|_\infty$$

On en déduit que pour tout $\varepsilon > 0$, en prenant $\eta = \frac{\varepsilon}{n\|A\|_\infty}$, $\|M - I_n\|_\infty \leq \eta$, implique $\|f(M) - A\|_\infty \leq \varepsilon$.

Proposition 13.33

Avec les mêmes notations. Si $\lim_a f = \ell$ alors f est bornée au voisinage de a .

Démonstration : Il suffit de prendre $\varepsilon = 1$ par exemple. Il existe alors η tel que si x appartient à $B(a, \eta) \cap A$ alors $f(x) \in B(\ell, 1)$. □

Théorème 13.34 (Caractérisation séquentielle de la limite)

Soit f une fonction de A dans F . Soit a un point adhérent à A et $\ell \in F$.

$$\left(\lim_a f = \ell \right) \iff \left(\forall (u_n) \in A^{\mathbf{N}}, \lim(u_n) = a \Rightarrow \lim(f(u_n)) = \ell \right).$$

Démonstration :

- $\Rightarrow$ On suppose que $\lim_a f = \ell$. On considère une suite $(u_n) \in A^{\mathbf{N}}$ qui converge vers a , on veut montrer que $\lim(f(u_n)) = \ell$, c'est-à-dire

$$\forall \varepsilon > 0, \exists N \in \mathbf{N}, \forall n \in \mathbf{N}, n \geq N \Rightarrow \|f(u_n) - \ell\| \leq \varepsilon$$

Pour $\varepsilon > 0$, comme $\lim_a f = \ell$, il existe $\eta > 0$ tel que pour tout x de A , si $\|x - a\|_E \leq \eta$ alors $\|f(x) - \ell\|_F \leq \varepsilon$. Maintenant pour ce η , comme $\lim(u_n) = a$, il existe $N \in \mathbf{N}$ tel que pour tout entier n , $n \geq N \Rightarrow \|u_n - a\| \leq \eta$. On a donc

$$n \geq N \rightarrow \|f(u_n) - \ell\| \leq \varepsilon.$$

On a bien $\lim(f(u_n)) = \ell$.

- $\boxed{\Leftarrow}$ On procède par contraposée. On montre que si la limite de f en a n'est pas ℓ (ce qui signifie que l'on peut avoir $\lim_a f \neq \ell$ ou que f n'a pas de limite en a) alors il existe une suite $(u_n) \in A^{\mathbb{N}}$ qui tend vers a et telle que $(f(u_n))$ ne tend pas vers ℓ . Le fait que la limite de f en a n'est pas ℓ signifie que

$$\exists \varepsilon > 0, \forall \eta > 0, \exists x \in A, \|x - a\|_E \leq \eta \text{ et } \|f(x) - \ell\|_F > \varepsilon.$$

On prend alors un ε vérifiant l'assertion ci-dessous et on l'applique pour $\eta = \frac{1}{n+1}$. On peut alors construire une suite (x_n) telle que

$$\forall n \in \mathbb{N}, \|x_n - a\| \leq \frac{1}{n+1}$$

donc $(x_n) \rightarrow a$ et $\forall n \in \mathbb{N}, \|f(x_n) - \ell\|_F > \varepsilon$ donc $(f(x_n))$ ne tend pas vers ℓ .

□

Corollaire 13.35

La notion de limite d'une fonction n'est pas modifiée si on remplace la norme de l'espace de départ et / ou de l'espace d'arrivée par une norme équivalente.

Corollaire 13.36

Soit $(F_1, \|\cdot\|_1), \dots, (F_n, \|\cdot\|_n)$ des espaces vectoriels normés. On pose $F = F_1 \times \dots \times F_n$ avec la norme produit notée $\|\cdot\|_F$. Soit f une application d'une partie A d'un espace vectoriel normé $(E, \|\cdot\|_E)$ dans F . On note pour tout $i \in \llbracket 1; n \rrbracket$ la i -ème composante telle que

$$\forall x \in A, f(x) = (f_1(x), \dots, f_n(x)).$$

Soit a un point adhérent à A et $\ell = (\ell_1, \dots, \ell_n) \in F$.

$$\left(\lim_a f = \ell \right) \iff \left(\forall i \in \llbracket 1; n \rrbracket, \lim_a f_i = \ell_i \right).$$

Démonstration : Il suffit d'utiliser la caractérisation séquentielle de la limite et le résultat analogue sur les limites de suites. □

Remarque : En particulier, si F est un espace de dimension finie et si $\mathcal{B} = (e_1, \dots, e_n)$ est une base de F . Pour $f : A \rightarrow F$, on peut noter $f_i = e_i^* \circ f$ la i -ème composante de f dans la base $\mathcal{B}$. On a alors pour a adhérent à A et $\ell = \ell_1 e_1 + \dots + \ell_n e_n$:

$$\left(\lim_a f = \ell \right) \iff \left(\forall i \in \llbracket 1; n \rrbracket, \lim_a f_i = \ell_i \right).$$

2.3 Composition et opérations

Proposition 13.37 (Limite d'une composée)

Soit $(E, \|\cdot\|_E)$, $(F, \|\cdot\|_F)$ et $(G, \|\cdot\|_G)$ trois espaces vectoriels normés. Soit A une partie de E et B une partie de F . On considère f une application de A dans F telle que $f(A) \subset B$ et g une application de B dans G . Soit a un point adhérent de A et $b \in F$. On suppose que $\lim_a f = b$. Alors

1. Le point b est adhérent à B .
2. S'il existe $\ell \in G$ tel que $\lim_b g = \ell$ alors $\lim_a g \circ f = \ell$.

Démonstration :

1. Comme a est adhérent à A , il existe une suite $(x_n) \in A^{\mathbb{N}}$ telle que $(x_n) \rightarrow a$. Comme, de plus $\lim_a f = b$ alors la suite $(f(x_n))$ tend vers b . Mais comme $(f(x_n))$ est une suite d'éléments de B on obtient que $b \in \overline{B}$.
2. Utilisons (pour changer) la définition de la limite par des voisinages. On sait que pour tout voisinage $\mathcal{V}_\ell$ de ℓ dans G , il existe un voisinage relatif $\mathcal{V}_b$ de b dans B tel que $g(\mathcal{V}_b) \subset \mathcal{V}_\ell$. Maintenant pour le voisinage $\mathcal{V}_b$ est la trace dans B d'un voisinage $\mathcal{V}'_b$ de b dans F . Pour ce voisinage, il existe un voisinage relatif $\mathcal{V}_a$ de a dans A tel que $f(\mathcal{V}_a) \subset \mathcal{V}'_b$. On a même $f(\mathcal{V}_a) \subset \mathcal{V}_b$ car $f(A) \subset B$. Finalement $(g \circ f)(\mathcal{V}_a) \subset \mathcal{V}_\ell$.

□

Proposition 13.38 (Combinaisons linéaires)

Soit f_1, f_2 deux fonctions définies sur une partie A d'un espace vectoriel normé $(E, \|\cdot\|_E)$ à valeurs dans un autre $(F, \|\cdot\|_F)$. Soit a un point adhérent à A . Soit $\lambda \in \mathbf{K}$ (où $\mathbf{K}$ est le corps de base de F). Si $\lim_a f_1 = \ell_1$ et $\lim_a f_2 = \ell_2$ alors $\lim_a (f_1 + f_2) = \ell_1 + \ell_2$ et $\lim_a \lambda f_1 = \lambda \ell_1$.

Démonstration : Il suffit d'utiliser la caractérisation séquentielle et la linéarité des limites des suites.

□

Proposition 13.39 (Produit)

Soit f_1, f_2 deux fonctions définies sur une partie A d'un espace vectoriel normé $(E, \|\cdot\|_E)$ à valeurs dans un autre $(F, \|\cdot\|_F)$. Soit a un point adhérent à A . On suppose que F est une algèbre et que $\|\cdot\|_F$ vérifie qu'il existe $C > 0$, tel que pour $(y, y') \in F^2$, $\|y \times y'\|_F \leq C \|y\|_F \|y'\|_F$.

Si $\lim_a f_1 = \ell_1$ et $\lim_a f_2 = \ell_2$ alors $\lim_a (f_1 \times f_2) = \ell_1 \times \ell_2$.

Démonstration : Il suffit d'utiliser la caractérisation séquentielle et la propriété analogue sur les suites.

□

3 Continuité

3.1 Définition

Là encore, on généralise la notion de fonction continue en un point (où f est définie) et de fonctions continues.

Définition 13.40

Soit f une fonction définie sur une partie A d'un espace vectoriel normé E à valeurs dans F . Soit $a \in A$. On dit que f est continue en a si f a une limite en a .

Remarque : La limite est alors nécessairement $f(a)$. On peut donc dire

$$(f \text{ est continue en } a) \iff \left(\lim_{x \rightarrow a} f(x) = f(a) \right)$$

Proposition 13.41 (Caractérisation séquentielle de la continuité en a)

Soit f une fonction définie sur une partie A d'un espace vectoriel normé E à valeurs dans F . Soit $a \in A$.

$$(f \text{ est continue en } a) \iff \left(\forall (u_n) \in A^{\mathbb{N}}, \lim(u_n) = a \Rightarrow \lim(f(u_n)) = f(a) \right).$$

On peut de même réécrire les opérations algébriques ainsi que les composées de limites

Proposition 13.42 (Combinaisons linéaires)

Soit f_1, f_2 deux fonctions définies sur une partie A d'un espace vectoriel normé $(E, \|\cdot\|_E)$ à valeurs dans un autre $(F, \|\cdot\|_F)$. Soit a un point de A . Soit $\lambda \in \mathbf{K}$ (où $\mathbf{K}$ est le corps de base de F). Si f_1 et f_2 sont continues en a alors $f_1 + f_2$ et λf_1 aussi.

Proposition 13.43 (Produit)

Soit f_1, f_2 deux fonctions définies sur une partie A d'un espace vectoriel normé $(E, \|\cdot\|_E)$ à valeurs dans un autre $(F, \|\cdot\|_F)$. Soit a un point de A . On suppose que F est une algèbre et que $\|\cdot\|_F$ vérifie qu'il existe $C > 0$, tel que pour $(y, y') \in F^2$, $\|y \times y'\|_F \leq C \|y\|_F \|y'\|_F$. Si f_1 et f_2 sont continues en a alors $f_1 \times f_2$

Proposition 13.44 (Composition)

Soit $(E, \|\cdot\|_E)$, $(F, \|\cdot\|_F)$ et $(G, \|\cdot\|_G)$ trois espaces vectoriels normés. Soit A une partie de E et B une partie de F . On considère f une application de A dans F telle que $f(A) \subset B$ et g une application de B dans G . Soit a un de A et $b = f(a)$. On suppose que f est continue en a et g est continue en b alors $g \circ f$ est continue en a .

Définition 13.45

Soit f une fonction définie sur une partie A d'un espace vectoriel normé E à valeurs dans F . Elle est dite continue si elle est continue en tout point de A .
On note $\mathcal{C}^0(A, F)$ l'ensemble des fonctions continues de A dans F .

Proposition 13.46

1. L'ensemble $\mathcal{C}^0(A, F)$ des fonctions continues de A dans F est un sous-espace vectoriel de F^A .
2. Soit E, F et G trois espaces vectoriels normés, soit $f : E \rightarrow F$ et $g : F \rightarrow G$ des applications continues. L'application, $g \circ f$ est continue.
3. Soit E un espace vectoriel normé et F une algèbre normée par une norme vérifiant qu'il existe $C > 0$ tel que pour y, y' dans F , $\|y \times y'\| \leq C\|y\| \|y'\|$.
Soit f_1 et f_2 des applications continues de A dans F . Le produit $f_1 \times f_2$ est encore continue.
4. Soit f une application définie sur une partie A d'un espace vectoriel normé E et à valeurs dans F . On ne modifie pas le caractère continu en remplaçant la norme de E et/ou de F par une norme équivalente.

Démonstration : Il suffit d'utiliser les propriétés vues précédemment en tout point a de A . $\square$

Exemples :

1. Soit $E = \mathbf{K}_n[X]$. On note $\|\cdot\|_\infty$ la norme définie par

$$\|\cdot\|_\infty : \sum_{k=0}^n a_k X^k \mapsto \max_{k \in \llbracket 0; n \rrbracket} |a_k|$$

On pose Δ l'endomorphisme de E défini par $\Delta : P \mapsto P'$. Montrons que Δ est continue.

Il est clair que $\|\Delta(P)\| \leq n\|P\|$. De ce fait pour $P \in E$,

$$\lim_{Q \rightarrow P} \Delta(Q) = \lim_{H \rightarrow 0} \Delta(P + H) = \Delta(P) + \lim_{H \rightarrow 0} \Delta(H) = \Delta(P)$$

2. Si on travaille pour $E = \mathbf{K}[X]$, l'application de dérivation Δ n'est plus continue (toujours pour la norme infinie). En effet, si on considère la suite $P_n = \frac{1}{n}X^n$, elle tend vers 0 mais $\Delta(P_n) = X^{n-1}$ ne tend pas vers 0 = $\Delta(0)$ car $\|X^{n-1}\| = 1$.

ATTENTION

Soit f une application entre deux espaces vectoriels normés E et F . En remplaçant la norme de E ou celle de F par une norme équivalente, on ne change pas les « topologies » de ce fait on ne modifie pas le caractère continu de f . Cependant, si les normes ne sont plus équivalentes, ce n'est plus vrai.

Par exemple, soit $E = \mathcal{C}^0([0, 1], \mathbb{R})$ et $F = \mathbb{R}$. On considère $\delta_0 : f \mapsto f(0)$.

- Si on utilise la norme infinie sur E alors f est continue. En effet, $|\delta_0(f)| = |f(0)| \leq \|f\|_\infty$. On peut alors démontrer que δ_0 est continue comme ci-dessus (voir aussi plus loin).
- Si on utilise la norme 1 sur E alors δ_0 n'est pas continue. On peut le montrer en utilisant la suite de fonctions

$$f_n : x \mapsto \begin{cases} 0 & \text{si } x > \frac{1}{n} \\ 1 - nx & \text{sinon} \end{cases}$$

La suite de fonctions (f_n) tend vers la fonction nulle mais pour autant $(\delta_0(f_n)) \rightarrow 1 \neq 0$.

Définition 13.47 (Fonction polynomiale)

Soit E un espace vectoriel de dimension n .

On appelle fonction polynomiale une application $f : E \rightarrow \mathbb{K}$ de la forme

$$f : x \mapsto \sum_{(i_1, \dots, i_n) \in \llbracket 0; d \rrbracket^n} a_{i_1, \dots, i_n} e_1^*(x)^{i_1} \times \dots \times e_n^*(x)^{i_n}$$

où $\mathcal{B} = (e_1, \dots, e_n)$ est une base de E et d est un entier naturel.

Exemples :

1. Dans le cas où $E = \mathbb{K}^n$, une fonction polynomiale est juste une fonction de la forme

$$(x_1, \dots, x_n) \mapsto \sum_{(i_1, \dots, i_n) \in \llbracket 0; d \rrbracket^n} a_{i_1, \dots, i_n} x_1^{i_1} \times \dots \times x_n^{i_n}$$

Par exemple

$$f : (x, y, z) \mapsto xy + 3x^2z - xy^3z^2$$

2. Soit $n \in \mathbb{N}^*$. La fonction déterminant $\det : \mathcal{M}_n(\mathbb{K}) \rightarrow \mathbb{K}$ est une fonction polynomiale car

$$\det : M \mapsto \sum_{\sigma \in \mathfrak{S}_n} \varepsilon(\sigma) M[1, \sigma[1]] \times \dots \times M[n, \sigma[n]]$$

Remarque : La notion de fonction polynomiale ne dépend pas de la base choisie. Soit $\mathcal{B} = (e_1, \dots, e_n)$ et $\mathcal{B}' = (\varepsilon_1, \dots, \varepsilon_n)$ deux bases de E et f une fonction polynomiale relativement à la base $\mathcal{B}$. Elle est donc de la forme

$$f : x \mapsto \sum_{(i_1, \dots, i_n) \in \llbracket 0; d \rrbracket^n} a_{i_1, \dots, i_n} e_1^*(x)^{i_1} \times \dots \times e_n^*(x)^{i_n}$$

Si on note P la matrice de passage de la base $\mathcal{B}$ à la base $\mathcal{B}'$. Soit $x \in E$. On note $X = \text{Mat}_{\mathcal{B}}(x)$ et $X' = \text{Mat}_{\mathcal{B}'}(x)$. On sait que $X = PX'$ c'est-à-dire que pour tout $i \in \llbracket 1; n \rrbracket$,

$$e_i^*(x) = \sum_{j=1}^n P[i, j] \varepsilon_j(x)$$

Il suffit alors de « faire le calcul » pour voir que f est aussi polynomiale par rapport à la base $\mathcal{B}'$.

Proposition 13.48

Les applications polynomiales sont continues.

Démonstration : Ce sont des sommes et des produits des e_i^* qui sont continues. $\square$

Proposition 13.49

Soit f et g deux applications continues définies sur A et à valeurs dans F . On suppose qu'il existe une partie H dense dans A telle que f et g coïncident sur H alors $f = g$.

Démonstration : Soit $a \in A$, comme H est dense dans A , il existe une suite (x_n) d'éléments de H qui tend vers a . Maintenant, pour tout entier n , $f(x_n) = g(x_n)$. Or, par continuité,

$$(f(x_n)) \rightarrow f(a) \quad \text{et} \quad (g(x_n)) \rightarrow g(a)$$

On en déduit que $f = g$. $\square$

Exemple : Soit f une application de $\mathbf{R}$ dans $\mathbf{R}$. On la suppose continue et qu'elle vérifie que

$$\forall (x, y) \in \mathbf{R}^2, f(x + y) = f(x) + f(y).$$

En posant $\alpha = f(1)$. On montre aisément par récurrence que

- $\forall n \in \mathbf{N}, f(n) = n\alpha$
- $\forall n \in \mathbf{Z}, f(n) = n\alpha$
- $\forall r \in \mathbf{Q}, f(r) = r\alpha$

On en déduit que f coïncide avec $g : x \mapsto x\alpha$ sur $\mathbf{Q}$ qui est dense dans $\mathbf{R}$, comme g est aussi continue, $f = g$.

3.2 Caractérisation du caractère continu par les images réciproques

Théorème 13.50

Soit f une application d'une partie A de E dans F . Les assertions suivantes sont équivalentes :

- i) L'application f est continue.
- ii) L'image réciproque de tout ouvert U de F est un ouvert relatif de A .
- iii) L'image réciproque de tout fermé X de F est un fermé relatif de A .

Démonstration :

- En remarquant que si X est une partie de F , $f^{-1}(\bigcup_F X) = \bigcup_A f^{-1}(X)$ il est clair que ii et iii sont équivalents.
- $i) \Rightarrow ii)$ On suppose que f est continue, on se donne un ouvert U de F et on note $V = f^{-1}(U)$. Montrons que V est un ouvert. Pour cela on se donne $a \in V$ et on pose $y = f(a) \in U$. Comme U est un ouvert, il existe ε tel que $B(y, \varepsilon) \subset U$.

Maintenant, comme f est continue, pour $\varepsilon > 0$, il existe η tel que pour tout x de A , $x \in B(a, \eta) \Rightarrow f(x) \in B(f(a), \varepsilon)$. Dit autrement, $B(a, \eta) \cap A \subset f^{-1}(B(f(a), \varepsilon)) \subset f^{-1}(U) = V$. On a bien montré que V était un ouvert relatif.

- ii) $\Rightarrow$ i) On suppose que l'image réciproque de tout ouvert U de F est un ouvert relatif de A . Soit $a \in A$, on pose $y = f(a)$. Pour tout $\varepsilon > 0$, la boule ouvert de centre y et de rayon ε est un ouvert. Son image réciproque V est donc un ouvert qui contient a . On en déduit qu'il existe $\eta > 0$ tel que $B(a, \eta) \cap A \subset V$. En appliquant f on obtient que

$$\forall x \in A, \|x - a\| < \eta \Rightarrow f(x) \in B(y, \varepsilon).$$

L'application f est donc continue.

□

Exemples :

1. Soit $f : \mathbb{R}^2 \rightarrow \mathbb{R}$ telle que $f(x, y) = x^2 + y^2$. Comme c'est une fonction polynômiale elle est continue. On en déduit que l'image réciproque du fermé $\{1\}$ c'est-à-dire la sphère $S = \{(x, y) \in \mathbb{R}^2 \mid x^2 + y^2 = 1\}$ est un fermé.
2. Soit E un espace vectoriel de dimension finie. Soit f une application polynomiale (par exemple une application linéaire). On considère les ensembles

$$A = \{x \in E, f(x) = 0\}; A^+ = \{x \in E, f(x) \geq 0\} \text{ et } A^{++} = \{x \in E, f(x) > 0\}$$

L'ensemble $A = f^{-1}(\{0\})$ est un fermé. En effet, comme $\{0\}$ est un fermé, A est l'image réciproque d'un fermé par une application continue.

L'ensemble $A^+ = f^{-1}([0, +\infty[)$ est un fermé. En effet, comme $[0, +\infty[$ est un fermé, A est l'image réciproque d'un fermé par une application continue.

L'ensemble $A^{++} = f^{-1}(]0, +\infty[)$ est un ouvert. En effet, comme $]0, +\infty[$ est un ouvert, A est l'image réciproque d'un ouvert par une application continue.

3. On peut aussi utiliser la méthode ci-dessus pour des ensembles définis par plusieurs équations. Par exemple si on considère $E = \mathbb{R}_n[X]$ et A l'ensemble des polynômes dont tous les coefficients sont strictement inférieurs à 1 en valeur absolue.

On note $(e_0^*, \dots, e_n^*)$ les formes coordonnées de sorte que

$$A = \{P \in E, \forall i \in \llbracket 0; n \rrbracket, |e_i^*(P)| < 1\}$$

Pour tout $i \in \llbracket 0; n \rrbracket$, l'ensemble $A_i = \{P \in E, e_i^*(P) < 1\}$ est un ouvert d'après ce qui précède. De même l'ensemble $B_i = \{P \in E, e_i^*(P) > -1\}$ est un ouvert.

On en déduit que A est un ouvert comme intersection d'un nombre fini d'ouverts.

Exercices :

1. Retrouver que l'ensemble des matrices stochastiques est fermé.
2. On considère $A = \text{GL}_n(\mathbb{K})$. Montrer que A est ouvert.

3.3 Applications uniformément continues

Définition 13.51

Soit f une application définie sur une partie A d'un espace vectoriel normé E et à valeurs dans F . On dit que f est uniformément continue si

$$\forall \varepsilon > 0, \exists \eta > 0, \forall (x, y) \in A^2, d(x, y) \leq \eta \Rightarrow d(f(x), f(y)) \leq \varepsilon.$$

Remarques :

1. On peut remplacer les inégalités larges par des inégalités strictes.
2. En changeant la norme de E et/ou de F par une norme équivalente, on ne modifie pas le caractère uniformément continue d'une application.

Proposition 13.52

Soit f une application de $A \subset E$ dans F .
Si elle est uniformément continue alors elle est continue.

Démonstration : Il suffit de recopier la démonstration classique.

□

ATTENTION

Ce n'est pas une équivalence. L'application $x \mapsto x^2$ est continue sur $\mathbf{R}$ mais pas uniformément continue. En effet, si on suppose par l'absurde qu'elle est uniformément continue. Pour $\varepsilon = 1$, il existe $\eta > 0$ tel que $|x - y| \leq \eta \Rightarrow |x^2 - y^2| \leq \varepsilon = 1$. Maintenant en posant $y = x + \eta$, $|y^2 - x^2| = 2x\eta + \eta^2$ donc, quand $x \rightarrow \infty$, $|x^2 - y^2| \rightarrow \infty$ ce qui contredit l'hypothèse.

3.4 Applications lipschitziennes

Rappelons, ce qui a déjà été vu sur les applications lipschitziennes.

Définition 13.53 (Applications lipschitziennes)

Soit $(E, \|\cdot\|_E)$ et $(F, \|\cdot\|_F)$ deux espaces vectoriels normés. Soit f une application de $A \subset E$ dans F .

1. Soit $k \in \mathbf{R}_+$. L'application f est dite k -lipschitzienne si

$$\forall (x, y) \in A^2, d(f(x), f(y)) \leq kd(x, y)$$

2. L'application f est dite lipschitzienne s'il existe un réel positif k tel que f soit k -lipschitzienne.

Exemple : L'application norme est 1-lipschitzienne. En effet soit E un espace vectoriel normé,

$$\forall (x, y) \in E^2, \left| \|x\|_E - \|y\|_E \right| \leq \|x - y\|_E$$

d'après la deuxième inégalité triangulaire.

Proposition 13.54

Soit f une application lipschitzienne. Elle est uniformément continue (et donc continue).

Démonstration : Si f est k -lipschitzienne. Pour tout $\varepsilon > 0$, on peut prendre dans la définition de l'uniforme continuité $\eta = \frac{\varepsilon}{k}$.

□

Exemple : L'application norme est donc continue.

ATTENTION

Là encore, ce n'est pas une équivalence. La fonction $x \mapsto \sqrt{x}$ sur $[0, 1]$ n'est pas lipschitzienne car

$$\forall x > 0, \frac{f(x) - f(0)}{x - 0} = \frac{\sqrt{x}}{x} = \frac{1}{\sqrt{x}}$$

et que $\frac{1}{\sqrt{x}} \xrightarrow{x \rightarrow 0^+} +\infty$

Cependant, elle est uniformément continue, en effet pour tout $\varepsilon > 0$, on peut prendre $\eta = \varepsilon^2$ pour obtenir que

$$|x - y| \leq \eta \Rightarrow |\sqrt{x} - \sqrt{y}| \leq \varepsilon$$

En effet :

- si $\sqrt{x}$ et $\sqrt{y}$ sont inférieurs à ε c'est vrai
- si $\sqrt{x}$ ou $\sqrt{y}$ sont supérieurs à ε alors $\sqrt{x} + \sqrt{y} \geq \varepsilon$ et donc

$$|\sqrt{x} - \sqrt{y}| = \frac{|x - y|}{\sqrt{x} + \sqrt{y}} \leq \frac{\varepsilon^2}{\varepsilon} = \varepsilon.$$

Proposition 13.55

L'ensemble des applications lipschitziennes de $A \subset E$ dans F est un espace vectoriel.

Voici un exemple classique (à connaître).

Définition 13.56 (Distance à une partie)

Soit $A \subset E$ un ensemble non vide. On appelle distance à A et on note $d(., A)$ l'application

$$\begin{aligned} d(., A) : E &\mapsto \mathbf{R} \\ x &\mapsto d(x, A) = \inf_{a \in A} d(x, a) \end{aligned}$$

Remarque : La partie $\{d(x, a) \mid a \in A\}$ est une partie non vide de $\mathbf{R}_+$. Elle est donc minorée par 0. De ce fait elle admet une borne inférieure qui peut (ou pas) être atteinte.

Exemple : Soit $A =]0, 1[$, $d(2,]0, 1[) = 1$ mais il n'existe pas d'élément a de A tel que $d(2, a) = 1$.

Proposition 13.57

Soit $A \subset E$ une partie non vide. L'application $d(., A)$ est lipschitzienne car

$$\forall (x, y) \in E^2, |d(x, A) - d(y, A)| \leq |x - y|$$

Démonstration : Il suffit de voir que pour tout $a \in A$,

$$d(x, a) - d(x, y) \leq d(y, a)$$

En utilisant que $d(x, A) \leq d(x, a)$ on obtient que

$$d(x, A) - d(x, y) \leq d(y, a)$$

Ceci étant vrai pour tout a de A ,

$$d(x, A) - d(x, y) \leq d(y, A).$$

On a donc $d(x, A) - d(y, A) \leq d(x, y)$. Maintenant, par symétrie,

$$|d(x, A) - d(y, A)| \leq d(x, y)$$

□

Exercice : Montrer que $\{x \in E \mid d(x, A) = 0\} = \overline{A}$.

3.5 Applications linéaires continues

Nous nous intéresserons la majorité du temps aux applications linéaires. Il y a alors un critère simple pour montrer qu'une application est continue.

Théorème 13.58

Soit f une application **linéaire** d'un espace vectoriel normé E dans un espace vectoriel normé F . Les assertions suivantes sont équivalentes :

- i) L'application f est continue.
- ii) L'application f est continue en 0.
- iii) L'application f est lipschitzienne
- iv) Il existe $C \in \mathbb{R}$ telle que

$$\forall x \in E, \|f(x)\|_F \leq C\|x\|_E$$

Remarque : En pratique on utilise quasiment toujours le critère iv) pour prouver qu'une application linéaire est (ou n'est pas continue).

Démonstration :

- $i) \Rightarrow ii)$ est évidente.
- $iii) \Rightarrow i)$ est déjà vu.
- $iv) \Rightarrow iii)$ a déjà été vu. Rappelons l'argument. On suppose $iv)$. Pour tout $(x, y) \in E^2$,

$$\|f(x) - f(y)\|_F = \|f(x - y)\|_F \leq C\|x - y\|_E.$$

- Montrons $ii) \Rightarrow iv)$. On suppose donc que f est continue en 0. De ce fait pour $\varepsilon = 1$, il existe η tel que

$$\forall x \in E, \|x\|_E \leq \eta \Rightarrow \|f(x)\|_F \leq 1$$

Maintenant pour tout vecteur x non nul,

$$x = \lambda y$$

où $\|y\|_E = \eta$. Il suffit de prendre $\lambda = \frac{\|x\|_E}{\eta}$ et $y = \eta \frac{x}{\|x\|_E}$. Par linéarité, on en déduit que

$$\|f(x)\|_F = |\lambda| \|f(y)\|_F \leq |\lambda| = C\|x\|_E$$

en posant $C = \frac{1}{\eta}$.

□

Notation : On note $\mathcal{L}_c(E, F)$ l'ensemble des applications linéaires continues de E dans F .

On rappelle que pour $f \in \mathcal{L}_c(E, F)$ on note $\|f\|_{\text{op}}$ ou $\|f\|$ la norme subordonnée (ou norme d'opérateurs) :

$$\|f\|_{\text{op}} = \|f\| = \sup \left\{ \frac{\|f(x)\|_F}{\|x\|_E}, x \neq 0_E \right\} = \sup \{ \|f(x)\|_F, x \in S(0, 1) \}$$

On obtient ainsi une norme sur $\mathcal{L}_c(E, F)$ qui est sous-multiplicative :

$$\forall (f, g) \in \mathcal{L}_c(E, F), \|f \circ g\|_{\text{op}} \leq \|f\|_{\text{op}} \times \|g\|_{\text{op}}$$

Proposition 13.59

Soit E, F et G trois espaces vectoriels normés. On note $\|\cdot\|_E, \|\cdot\|_F$ et $\|\cdot\|_G$ leur norme et on considère la norme produit pour $E \times F$. Soit $B : E \times F \rightarrow G$ une application bilinéaire. Elle est continue si et seulement si

$$\exists C \in \mathbf{R}, \forall (x, y) \in E \times F, \|B(x, y)\|_G \leq C\|x\|_E\|y\|_F$$

Démonstration :

- On suppose la condition. Montrons que B est continue pour la norme produit $\|(x, y)\| = \max(\|x\|, \|y\|)$. En effet soit (α_n) une suite tendant vers (x, y) . On pose $\alpha_n = (x_n, y_n)$. On a alors $(x_n) \rightarrow x$ car $\|x_n - x\|_E \leq \|\alpha_n - (x, y)\|$. De même, $(y_n) \rightarrow y$. On a alors

$$\begin{aligned} \|B(x_n, y_n) - B(x, y)\| &= \|B(x_n, y_n) - B(x, y_n) + B(x, y_n) - B(x, y)\| \\ &\leq \|B(x_n - x, y_n)\| + \|B(x, y_n - y)\| \\ &\leq C\|x_n - x\|_E K' + C\|y_n - y\|_F \|x\|_E \end{aligned}$$

où K' est un majorant de $\|y_n\|_F$ (la suite convergeant elle est bornée). On a bien $B(x_n, y_n) \rightarrow B(x, y)$ donc, par caractérisation séquentielle, B est continue.

- On suppose maintenant que B est continue. En particulier, elle est continue en 0. Donc il existe η tel que $\|(x, y)\| = \max(\|x\|, \|y\|) \leq \eta$ implique $\|B(x, y)\| \leq 1$. Dès lors pour tout $(x, y) \in E \times F$ avec $x \neq 0$ et $y \neq 0$,

$$B(x, y) = \frac{\|x\| \cdot \|y\|}{\eta^2} B\left(\eta \frac{x}{\|x\|}, \eta \frac{y}{\|y\|}\right) \leq \frac{\|x\| \cdot \|y\|}{\eta^2}.$$

□

Proposition 13.60 (Généralisation aux applications multilinéaires)

Soit $E_1, \dots, E_p$ des espaces vectoriels normés et G un vectoriel normé. Pour $i \in \llbracket 1; p \rrbracket$ on note $\|\cdot\|_i$ la norme sur E_i . On note $\|\cdot\|_G$ la norme de G . Soit $\Phi : \prod_{i=1}^p E_i \rightarrow G$ une application multilinéaire. Elle est continue (en normant $\prod_{i=1}^p E_i$ avec la norme produit) si et seulement si

$$\exists C \in \mathbf{R}, \forall (x_1, \dots, x_p) \in E_1 \times \dots \times E_p, \|\Phi(x_1, \dots, x_p)\|_G \leq C \prod_{i=1}^p \|x_i\|_i$$

4 Parties compactes d'un espace vectoriel normé

4.1 Définition

Certains théorèmes d'analyse ne sont vrais que sur des segments : théorème de Bolzano-Weierstrass, toute fonction continue sur un segment est bornée et atteint ses bornes. Le but de ce chapitre est de généraliser cela à des parties d'un espace vectoriel normé.

Définition 13.61

Soit A une partie d'un espace vectoriel normé E . On dit que c'est un compact (ou que c'est une partie compacte) si on peut extraire de toute suite de A une sous-suite **convergente dans A** .

Remarques :

1. Il est important que la limite de la suite soit dans A .
2. Cette définition s'appelle *propriété de Bolzano-Weierstrass*. Il existe une autre définition *propriété de Borel-Lebesgue* qui n'est pas au programme.
3. On peut aussi dire : « toute suite a une valeur d'adhérence ».
4. La définition de compact est *absolue*. Il n'y a pas de notion de compact relatif à X .
5. On peut dire que A est un compact ou que A est compact.

Exemples :

1. Dans $\mathbf{R}$, les segments $[a, b]$ sont des compacts. En effet si (u_n) est une suite à valeurs dans $[a, b]$ elle est bornée donc, d'après le théorème de Bolzano-Weierstrass, on peut extraire une sous-suite $(u_{\varphi(n)})$ convergente de (u_n) . De plus, comme $\forall n \in \mathbf{N}, a \leq u_n \leq b$ alors c'est encore vrai pour $(u_{\varphi(n)})$ et donc la limite de la suite extraite est bien dans $[a, b]$.

Cela montre bien que de toute suite $(u_n)_{n \geq 0}$ à valeurs dans $[a, b]$ on peut extraire une sous-suite $(u_{\varphi(n)})_{n \geq 0}$ qui converge vers une limite $\ell \in [a, b]$.

2. L'ensemble vide est un compact de $\mathbf{R}$.

3. Soit $A = \bigcup_{i=1}^p K_i$ une réunion finie de segment de $\mathbf{R}$. Montrons que K est compact.

Soit $(u_n)_{n \geq 0}$ une suite à valeurs dans A . Pour tout $i \in \llbracket 1 ; p \rrbracket$, notons $N_i = \{n \in \mathbf{N}, u_n \in K_i\}$ l'ensemble des indices n tels que le terme u_n de la suite appartienne à K_i . Il existe alors $i_0 \in \llbracket 1 ; p \rrbracket$ tel que N_{i_0} soit infini. On peut donc construire une extractrice φ telle que la sous-suite $(v_n)_{n \geq 0} = (u_{\varphi(n)})_{n \geq 0}$ soit à valeurs dans K_{i_0} . En utilisant maintenant que K_{i_0} est un compact, on peut extraire de $(v_n)_{n \geq 0} = (u_{\varphi(n)})_{n \geq 0}$ une sous-suite $(v_{\psi(n)})_{n \geq 0} = (u_{\varphi \circ \psi(n)})_{n \geq 0}$ convergeant dans $K_{i_0} \subset A$.

Cela montre que A est un compact.

4. L'ensemble $\mathbf{R}$ n'est pas compact.

On considère la suite $(u_n)_{n \geq 0}$ définie par $u_n = n$. Soit $(u_{\varphi(n)})_{n \geq 0}$ une suite extraite. Pour tout entier $n \geq 0$, $u_{\varphi(n)} = \varphi(n) \geq n$. On en déduit que $(u_{\varphi(n)})_{n \geq 0}$ n'est pas bornée donc pas convergente.

L'ensemble $\mathbf{R}$ n'est pas compact.

Exercice : Montrer qu'une intersection de parties compactes est compacte

Proposition 13.62

La caract re compact ne change pas si on remplace la norme de E par une norme  quivalente.

D monstration : Il suffit de voir que l'on ne modifie pas la convergence des suites. $\square$

Proposition 13.63

Soit A une partie compacte de E . Elle est ferm e et born e.

D monstration :

- On commence par montrer que A est ferm e. On utilise la caract risation s quentielle. Il suffit donc de montrer que toute suite de A convergente (dans E) converge dans A . En effet soit (u_n) une suite convergente vers $\ell \in E$. D'apr s la d finition, il existe une suite extraite qui converge dans A . Maintenant, on sait qu'une suite extraite d'une suite convergente converge vers la limite de la suite initiale donc $\ell \in A$.
- Montrons maintenant que A est born e. On proc de par l'absurde. Si A n'est pas born e, alors elle n'est incluse dans aucune boule, en particulier,

$$\forall n \in \mathbb{N}, \exists x_n \in A, \|x_n - 0\| \geq n.$$

On en d duit que toute suite extraite de la suite (x_n) n'est pas born e donc pas convergente.

$\square$

ATTENTION

Dans le cas g n ral ce n'est pas une  quivalence. Pour $E = \mathbb{K}[X]$ avec la norme infinie. La boule unit  ferm e $\overline{B}(0, 1)$ est born e et ferm e. Mais elle n'est pas compacte. En effet, si on pose $u_n = X^n$, on ne peut pas extraire une sous-suite convergente car $\forall (n, m) \in \mathbb{N}^2, n \neq m \Rightarrow d(X^n, X^m) = 1$. Cette propri t  reste vraie pour toute suite extraite.

Proposition 13.64

Soit A une partie compacte de E . Une partie B de A est compacte si et seulement si elle est ferm e (dans A).

D monstration :

- $\Rightarrow$ On suppose que B est compacte. Elle est donc ferm e (dans E) et donc ferm e dans A .
- $\Leftarrow$ On suppose que B est un ferm  relatif de A . Toute suite (u_n) d' l ments de B est une suite d' l ments de A . En particulier, on peut en extraire une sous-suite $(u_{\varphi(n)})$ qui converge (dans A). Maintenant, comme B est un ferm  relatif, la suite $(u_{\varphi(n)})$ qui est une suite d' l ments de B qui converge dans A a sa limite dans B . On a bien montr  que B  tait compacte.

$\square$

Proposition 13.65

Soit A une partie compacte de E et (u_n) une suite de A . Elle est convergente si et seulement si elle a une unique valeur d'adhérence.

Démonstration :

- $\Rightarrow$ Déjà vu. Une suite convergente a une unique valeur d'adhérence : sa limite.
- $\Leftarrow$ Procédons par contraposée. Montrons que si (u_n) ne converge pas alors elle a plusieurs valeurs d'adhérence (car elle ne peut pas en avoir aucune ; par définition toute suite dans un compact a au moins une valeur d'adhérence). Notons ℓ_1 une valeur d'adhérence de (u_n) (qui existe car A est compacte). Maintenant comme (u_n) ne converge pas, il existe ε tel que pour tout $N \in \mathbb{N}$, il existe $n \in \mathbb{N}$ tel que $n \geq N$ et $u_n \notin B(\ell_1, \varepsilon)$. On peut donc ainsi construire une suite extraite $(u_{\varphi(n)})$ telle que

$$\forall n \in \mathbb{N}, u_{\varphi(n)} \notin B(\ell_1, \varepsilon).$$

Maintenant, cette suite extraite a nécessairement une valeur d'adhérence ℓ_2 qui ne peut pas être ℓ_1 . Au final, (u_n) a bien au moins deux valeurs d'adhérence.

□

Proposition 13.66

Soit $(E_1, \|\cdot\|_1), \dots, (E_p, \|\cdot\|_p)$ des espaces vectoriels normés. On se donne pour tout $i \in \llbracket 1 ; p \rrbracket$ une partie A_i compacte de E_i . Alors $A_1 \times \dots \times A_p$ est une partie compacte de $E_1 \times \dots \times E_p$ muni de la norme produit.

Démonstration : Soit (u_n) une suite à valeur dans $\prod_{i=1}^p A_i$. Pour tout $n \in \mathbb{N}$, notons $u_n = (u_n(1), \dots, u_n(p))$.

On peut considérer une extractrice φ_1 telle que $(u_{\varphi_1(n)}(1))$ converge vers ℓ_1 car $(u_n(1))$ est une suite à valeurs dans A_1 qui est compact. Maintenant, on peut extraire de $(u_{\varphi_1(n)}(2))$ une suite $(u_{\varphi_1 \circ \varphi_2(n)}(2))$ qui converge vers ℓ_2 . Notons que $(u_{\varphi_1 \circ \varphi_2(n)}(1))$ converge encore vers ℓ_1 comme suite extraite d'une suite convergente.

En continuant ainsi on peut construire des extractrices $\varphi_1, \dots, \varphi_p$ telles que, en posant $\varphi = \varphi_p \circ \dots \circ \varphi_1$ on ait que pour tout $i \in \llbracket 1 ; p \rrbracket$, $(u_{\varphi(n)}(i))$ converge vers ℓ_i .

De par la définition de la norme produit, la suite $(u_{\varphi(n)})$ converge vers $\ell = (\ell_1, \dots, \ell_p)$.

□

4.2 Applications continues sur une partie compacte

On a défini les parties compactes comme celles qui vérifient la propriété de Bolzano-Weierstrass. On voulait aussi généraliser la fait qu'une fonction **continue** sur un segment est bornée et atteint ses bornes.

Théorème 13.67

Soit E et F des espaces vectoriels normés. Soit A une partie de E et f une application **continue** de A dans F . Soit B une partie compacte de A alors $f(B)$ est compacte.

Remarque : On dit « l'image d'un compact par une application continue est compact »

Démonstration : Soit (y_n) une suite de $f(B)$. Par définition, il existe une suite (x_n) de B telle que $\forall n \in \mathbb{N}, y_n = f(x_n)$. Il suffit en effet de « choisir » un antécédent par chaque y_n . Maintenant, comme (x_n) est une suite de B qui est compact, il existe une extractrice φ telle que $(x_{\varphi(n)})$ converge. Notons ℓ sa limite. Comme f est continue $(y_{\varphi(n)})$ converge vers $f(\ell)$. Cela prouve bien que $f(B)$ est compact. $\square$

ATTENTION

Ne pas confondre le fait que l'image **directe** d'un compact par une application continue est compact avec le fait que l'image **réciroque** d'un fermé / ouvert par une application continue est fermé / ouvert.

Nous pouvons alors retrouver le théorème vu en première année.

Corollaire 13.68 (Théorème des bornes atteintes)

Soit f une application continue de $A \subset E$ dans $\mathbb{R}$. On suppose que A est compact et non vide alors $f(A)$ est bornée et les bornes sont atteintes.

Démonstration : On sait que $f(A)$ est un compact de $\mathbb{R}$. Il est donc borné et fermé. On peut donc poser $M = \sup f(A)$ sa borne supérieure. Par définition, on peut trouver une suite (α_n) d'éléments de $f(A)$ qui tendent vers M car pour tout $\varepsilon = \frac{1}{n+1}$ il existe $\alpha_n \in f(A)$ tel que $M - \varepsilon \leq \alpha_n \leq M$. Mais comme $f(A)$ est fermé, $M = \lim(\alpha_n)$ appartient à $f(A)$.

On peut faire de même pour la borne inférieure. $\square$

Remarque : Quand on considère une application d'une partie compacte A et une application continue f de A dans un espace vectoriel normé $(F, \|\cdot\|_F)$ on peut appliquer ce théorème à

$$\begin{aligned} \|f\|_F : A &\rightarrow \mathbb{R} \\ x &\mapsto \|f(x)\|_F \end{aligned}$$

En effet $\|f\|_F$ est continue en tant que composée de deux applications continues.

Exercice : Soit K une partie compacte non vide. Montrer que pour tout élément x de E , il existe un élément $\alpha \in K$ tel que $d(x, K) = \|x - \alpha\|$.

Exemple : C'est une des étapes qui permet de démontrer le théorème de D'Alembert-Gauss : Supposons par l'absurde qu'il existe un polynôme $P \in \mathbb{C}[X]$ non constant qui ne s'annule pas. On considère alors

$$\begin{aligned} f : \mathbb{C} &\rightarrow \mathbb{R} \\ z &\mapsto |P(z)| \end{aligned}$$

Il est clair que si $\lim_{|z| \rightarrow \infty} f(z) = +\infty$. En effet $|P(z)| \sim |a_d||z|^d$ où d est le degré de P .

Cela signifie que si on considère $z_0 \in \mathbb{C}$, il existe $R \in \mathbb{R}_+$ tel que

$$\forall z \in \mathbb{C}, |z| > R \Rightarrow f(z) \geq f(z_0)$$

On en déduit que pour chercher le maximum de la fonction on peut se restreindre au compact $\overline{B}(0, R)$.

Finalement, la fonction f est minorée et atteint son maximum. Il existe donc un nombre complexe α tel que

$$\forall z \in \mathbb{C}, f(z) \geq f(\alpha)$$

C'est-à-dire

$$\forall z \in \mathbb{C}, |P(z)| \geq |P(\alpha)|.$$

Il ne reste plus qu'à trouver la contradiction....

On peut aussi reformuler le théorème de Heine

Théorème 13.69 (Théorème de Heine)

Soit A une partie compacte de E . Toute application continue f de A dans F est uniformément continue.

Démonstration : Il suffit de recopier la démonstration du cours de première année. On suppose par l'absurde qu'il existe une fonction f de A dans F qui soit continue mais pas uniformément continue. De ce fait, il existe $\varepsilon > 0$ tel que pour tout $\eta > 0$, il existe $(x, y) \in A^2$ tels que $d(x, y) \leq \eta$ et $d(f(x), f(y)) > \varepsilon$. En particulier en prenant $\eta = \frac{1}{n+1}$ on peut construire deux suites (x_n) et (y_n) telles que

$$\forall n \in \mathbb{N}, d(x_n, y_n) \leq \frac{1}{n+1} \text{ et } d(f(x_n), f(y_n)) > \varepsilon.$$

On peut alors réaliser une double extraction avec de construire φ telle que $(x_{\varphi(n)})$ et $(y_{\varphi(n)})$ convergent. En notant ℓ_x et ℓ_y les limites et passant à la limite dans

$$d(x_{\varphi(n)}, y_{\varphi(n)}) \leq \frac{1}{\varphi(n)+1}$$

on obtient que $\ell_x = \ell_y$. En utilisant alors la continuité de f on obtient que $(f(x_{\varphi(n)}))$ et $(f(y_{\varphi(n)}))$ convergent toutes les deux vers $f(\ell_x) = f(\ell_y)$. Ce qui est absurde car $\forall n \in \mathbb{N}, d(f(x_{\varphi(n)}), f(y_{\varphi(n)})) > \varepsilon$.
 $\square$

Exemple : On a montré précédemment que $x \mapsto \sqrt{x}$ était uniformément continue sur $[0, 1]$. Cela découle directement de ce théorème.

Exercice : Montrer que $x \mapsto \sqrt{x}$ est uniformément continue sur $\mathbb{R}_+$.

5 Espaces vectoriels de dimension finie

Nous allons étudier plus particulièrement les notions topologiques étudiées précédemment dans le cas des espaces vectoriels de dimension finie.

5.1 Equivalence des normes

Commençons par un lemme

Lemme 13.70

Soit E un espace vectoriel de dimension finie. Soit $\mathcal{B} = (e_1, \dots, e_n)$ une base de E , on note $\|\cdot\|_\infty$ la norme infinie relative à cette base.
La boule unité fermée $\overline{B}_\infty(0, 1)$ est compacte pour la norme $\|\cdot\|_\infty$.

Démonstration : Notons D la boule unité fermée de $\mathbb{K}$, c'est-à-dire que $D = [-1, 1]$ si $\mathbb{K} = \mathbb{R}$ et $D = \{z \in \mathbb{C}, |z| \leq 1\}$ si $\mathbb{K} = \mathbb{C}$.

Regardons la boule unité fermée

$$\overline{B}_\infty(0, 1) = \{x \in E \mid \|x\|_\infty \leq 1\} = \{\lambda_1 e_1 + \cdots + \lambda_n e_n \mid (\lambda_1, \dots, \lambda_n) \in D^n\}.$$

On considère l'application

$$\begin{aligned} \Phi : (\mathbf{K}^n, \|\cdot\|_\infty) &\rightarrow (E, \|\cdot\|_\infty) \\ (x_1, \dots, x_n) &\mapsto x_1 e_1 + \cdots + x_n e_n \end{aligned}$$

C'est une application linéaire continue (car elle est 1-lipschizienne) et $\overline{B}_\infty(0, 1) = \Phi(D^n)$.

On sait que D est un compact de $\mathbf{K}$ et donc D^n est un compact comme produit fini de compact.

On en déduit que $\overline{B}_\infty(0, 1) = \Phi(D^n)$ est compact comme image d'un compact par une application continue. $\square$

Théorème 13.71

Soit E un espace vectoriel de dimension finie. Toutes les normes sur E sont équivalentes.

Démonstration : (Non exigible) On se fixe une base $\mathcal{B}$ de E et on note encore $\|\cdot\|_\infty$ la norme infinie associée à cette base.

On considère une norme $\|\cdot\|$ sur E . On va montrer quelle est équivalente à $\|\cdot\|_\infty$. Par transitivité on aura alors obtenu que toutes les normes sont équivalentes.

Commençons par montrer que $\|\cdot\| : (E, \|\cdot\|_\infty) \rightarrow \mathbf{R}$ est continue².

Pour tout $x \in E$, on le décompose en $x = \sum_{i=1}^n x_i e_i$ et on a

$$\|x\| \leq \sum_{i=1}^n |x_i| \|e_i\| \leq nC \|x\|_\infty$$

où $C = \max(\|e_1\|, \dots, \|e_n\|)$.

Pour (x, y) dans E on a donc

$$\left| \|x\| - \|y\| \right| \leq \|x - y\| \leq nC \|x - y\|_\infty$$

ce qui prouve bien que $\|\cdot\| : (E, \|\cdot\|_\infty) \rightarrow \mathbf{R}$ est 1-lipschitzienne donc continue.

De plus, pour $\|\cdot\|_\infty$, la sphère unité $S_\infty(0, 1) \subset \overline{B}_\infty(0, 1)$ est un fermé d'une partie compacte (d'après le lemme précédent) c'est donc aussi un compact.

On en déduit que $\|\cdot\|$ est bornée et atteint ses bornes sur $S_\infty(0, 1)$. Il existe donc α et β tels que

$$\forall x \in S_\infty(0, 1), \beta \leq \|x\| \leq \alpha.$$

Notons que comme la valeur minimale est atteinte et que $0 \notin S_\infty(0, 1)$, $\beta > 0$. Maintenant pour tout $x \in E$ avec $x \neq 0$, $\frac{x}{\|x\|_\infty} \in S_\infty(0, 1)$ et donc

$$\beta \leq \left\| \frac{x}{\|x\|_\infty} \right\| \leq \alpha \Rightarrow \beta \|x\|_\infty \leq \|x\| \leq \alpha \|x\|_\infty.$$

2. Faire attention que ce n'est pas évident; il est évident que $\|\cdot\| : (E, \|\cdot\|) \rightarrow \mathbf{R}$ mais ce n'est plus le cas quand la norme de l'espace de départ n'est plus $\|\cdot\|$.

Les normes sont bien équivalentes car

$$\forall x \in E, \|x\| \leq \alpha \|x\|_\infty \text{ et } \|x\|_\infty \leq \frac{1}{\beta} \|x\|.$$

□

Remarque : Dans le cas de l'étude topologique d'un espace vectoriel de dimension finie, il n'est pas nécessaire de préciser la norme utilisée. En effet en modifiant une norme pas une autre qui lui sera équivalente on ne modifie pas :

- La convergence des suites
- Le caractère ouvert et donc l'intérieur
- Le caractère fermé et donc l'adhérence
- Le caractère compact
- La continuité des fonctions.

5.2 Topologie des espaces vectoriels de dimension finie

Nous allons voir certaines propriétés topologiques qui ne sont vraies que pour les espaces vectoriels de dimension finie. Dans tout ce paragraphe, E est un espace vectoriel normé de dimension finie. On note $\|\cdot\|$ sa norme.

Proposition 13.72

Une partie A de E est compacte si et seulement si elle est fermée et bornée.

Démonstration : On a déjà vu que les parties compactes étaient fermées et bornées. Montrons l'implication inverse.

Soit A une partie bornée. Il existe donc $R > 0$ tel que $A \subset \overline{B}(0, R)$. On a vu que la boule unité fermée était compacte pour une norme infinie. La même preuve montre que $\overline{B}_\infty(0, R)$ est aussi compacte. C'est donc aussi vrai pour la norme $\|\cdot\|$.

On en déduit que A est un fermé d'un compact. C'est donc bien un compact. □

Exemples :

1. Les boules fermées et les sphères sont compactes. Rappelons que ce résultat n'est plus vrai en dimension infinie.
2. L'ensemble $O_n(\mathbf{R})$ des matrices orthogonales est une partie compacte de $\mathcal{M}_n(\mathbf{R})$.
 - Pour tout $A \in O_n(\mathbf{R})$, $\|A\|_\infty \leq 1$ car pour tout $i \in \llbracket 1 ; n \rrbracket$, $\sum_{j=1}^n A[i, j]^2 = 1$. On en déduit que $O_n(\mathbf{R})$ est borné.
 - Soit $\Phi : \mathcal{M}_n(\mathbf{R}) \rightarrow \mathcal{M}_n(\mathbf{R})$ définie par $\Phi : A \mapsto A^\top A$. Chaque coordonnée de $\Phi(A)$ est polynomiale en les coordonnées de A . On en déduit que Φ est continue et donc $\Phi^{-1}(\{I_n\})$ est un fermé comme image réciproque d'un fermé par une application continue.

Corollaire 13.73

Soit (u_n) une suite de E bornée.

1. Elle admet une sous-suite convergente.
2. Elle converge si et seulement si elle a une unique valeur d'adhérence.

Démonstration : Il suffit d'utiliser que (u_n) est une suite à valeurs dans le compact $\overline{B}(0, R)$ où $\forall n \in \mathbb{N}, \|u_n\| \leq R$. $\square$

Proposition 13.74 (Rappel)

Les sous-espaces vectoriels sont fermés (mais pas compacts s'ils ne sont pas réduits à $\{0\}$).

Exemple : L'ensemble des matrices de trace nulle est fermé.

5.3 Applications continues

Théorème 13.75

Soit f une application linéaire de E dans F . Si E est de dimension finie alors f est continue.

Démonstration : On considère encore la norme infinie associée à une base $\mathcal{B} = (e_1, \dots, e_n)$. On note $C = \max_{1 \leq i \leq n} \|f(e_i)\|_F$. Pour tout $x = \sum \lambda_i e_i$ de E

$$\|f(x)\|_F \leq \sum_{i=1}^n \|\lambda_i f(e_i)\|_F \leq nC \|x\|_\infty$$

L'application f est donc continue. $\square$

Exemples :

1. Dans $E = \mathcal{M}_n(\mathbb{K})$. Pour toute matrice $P \in \text{GL}_n(\mathbb{K})$, l'application $\Phi : M \mapsto P^{-1}MP$ est linéaire donc continue.

On en déduit que si $(A_n) \rightarrow A$ alors si on pose $B_n = P^{-1}A_nP$, la suite (B_n) converge vers $P^{-1}AP$.

2. Les projections e_i^* sont continues car linéaires.

Proposition 13.76

Soit Φ est un application multilinéaire, il existe une constante C telle que

$$\forall (x_1, \dots, x_n), \|\Phi(x_1, \dots, x_n)\| \leq C \prod_{i=1}^n \|x_i\|.$$

En particulier, Φ est continue.

Démonstration : La démonstration est similaire à la précédente. Traitons le cas des formes bilinéaires. On se donne $(e_1, \dots, e_n)$ et $(f_1, \dots, f_p)$ des bases de E_1 et de E_2 . On sait alors que pour tout $(x, y) \in E_1 \times E_2$ on a

$$\Phi(x, y) = \sum_{i=1}^n \sum_{j=1}^p \lambda_i \mu_j \Phi(e_i, f_j)$$

où $\lambda_i = e_i^*(x)$ et $\mu_j = f_j^*(y)$. En particulier, si on note $C = \max_{\substack{1 \leq i \leq n \\ 1 \leq j \leq p}} \|\Phi(e_i, f_j)\|_F$ alors

$$|\Phi(x, y)| \leq npC \|x\| \|y\| = K \|x\| \|y\|.$$

Le résultat découle alors de la caractérisation de la continuité des applications multilinéaires. $\square$

Exemples :

1. L'application $\Phi : \mathcal{M}_n(\mathbf{K}) \times \mathcal{M}_n(\mathbf{K}) \rightarrow \mathcal{M}_n(\mathbf{K})$ définie par $\Phi : (A, B) \mapsto AB$ est bilinéaire donc continues
2. Soit E un espace euclidien, le produit scalaire $(\bullet, \bullet) : E \times E \rightarrow \mathbf{R}$ est bilinéaire donc continue.
3. Soit E de dimension n , le déterminant est n -linéaire. donc continue
4. Soit $\varphi : \mathcal{M}_n(\mathbf{K}) \rightarrow \mathcal{M}_{n,1}(\mathbf{K}) \rightarrow \mathcal{M}_{n,1}(\mathbf{K})$ l'application définie par $\varphi : (A, X) \mapsto AX$. Elle est bilinéaire donc continue.

6 Séries à valeurs dans un espace vectoriel de dimension finie

6.1 Généralités

Définition 13.77

Soit E un espace vectoriel de dimension finie. Soit $(\sum u_n)$ une série d'éléments de E .

1. On dit que la série converge si la suite des sommes partielles converge.
2. On dit que la série converge absolument si la série numérique $(\sum \|u_n\|)$ converge.

Remarque : Toutes les normes étant équivalente, l'absolue convergence ne dépend pas de la norme choisie.

Théorème 13.78

Soit $(\sum u_n)$ une série d'éléments de E absolument convergente. Elle converge.

Démonstration : On considère une série $(\sum u_n)$ absolument convergente. On se fixe une base $\mathcal{B} = (e_1, \dots, e_p)$ de E . En particulier, la série $\sum \|u_n\|_\infty$ est convergente. Comme pour tout $i \in \llbracket 1 ; p \rrbracket$ et tout n , $|e_i^*(u_n)| \leq \|u_n\|_\infty$ les séries $\sum e_i^*(u_n)$ sont absolument convergentes donc convergentes. Si on note (S_n) la suite des sommes partielles, on vient d'établir que pour tout $i \in \llbracket 1 ; p \rrbracket$, $e_i^*(S_n) = \sum_{k=0}^n e_i^*(u_k)$ converge donc (S_n) converge. $\square$

ATTENTION

Ce théorème n'est pas vrai (en général) dans un espace vectoriel normé de dimension infinie. Par exemple pour $E = \mathbf{K}[X]$ avec la norme infinie $\|\cdot\|_\infty$. On pose pour tout entier n non nul $P_n = \frac{1}{n^2}X^n$. Il est clair que $\|P_n\|_\infty = \frac{1}{n^2}$ et donc la série $\sum_{n \geq 0} \|P_n\|_\infty$ converge.

Par contre la série $\sum_{n \geq 0} P_n$ diverge. En effet, supposons par l'absurde qu'elle convergeait et notons Q la somme de la série et $d = \deg(Q)$. Pour tout entier $N > d$,

$$e_{d+1}^* \left(\sum_{n=1}^N P_n - Q \right) = e_{d+1}^*(P_{d+1})$$

On en déduit que

$$\left\| \sum_{n=1}^N P_n - Q \right\|_\infty \geq \frac{1}{(d+1)^2}$$

On a bien une absurdité.

Remarque : On peut généraliser ce théorème en disant que si $(E, \|\cdot\|)$ est *complet* alors l'absolue convergence d'une série implique sa convergence car il est simple de montrer que si $\sum u_n$ est absolument convergente alors la suite des sommes partielles est une suite de Cauchy.

6.2 Série géométrique de matrices

On se place dans $E = \mathcal{M}_n(\mathbf{K})$. On veut étudier les séries géométriques. On veut donc pouvoir « estimer » $\|A^p\|$

Lemme 13.79

Soit $\|\cdot\|$ une norme sur E , il existe C tel que

$$\forall (A, B) \in E^2, \|AB\| \leq C\|A\|\|B\|.$$

Démonstration : On utilise que pour $\|\cdot\|_\infty$, $\|AB\|_\infty \leq n\|A\|_\infty\|B\|_\infty$ et que toutes les normes sont équivalentes.

□

Remarque : On verra qu'il existe aussi des normes qui vérifient

$$\forall (A, B) \in \mathcal{M}_n(\mathbf{K})^2, \|AB\| \leq \|A\|\|B\|$$

Dans toute la suite on notera C une constante réelle telle que

$$\forall (A, B) \in E^2, \|AB\| \leq C\|A\|\|B\|.$$

Proposition 13.80

Soit $A \in \mathcal{M}_n(\mathbf{K})$, si $\|A\| < \frac{1}{C}$ alors la série géométrique $(\sum A^p)$ converge absolument. De plus sa somme vaut

$$\sum_{p=0}^{+\infty} A^p = (I_n - A)^{-1}$$

Démonstration : Pour la convergence, on voit que

$$\|A^p\| \leq C\|A^{p-1}\| \|A\| \leq C^2\|A^{p-2}\| \|A\|^2 \leq \dots \leq C^{p-1}\|A\| \|A\|^{p-1} = C^{p-1}\|A\|^p$$

Or $C^{p-1}\|A\|^p = \frac{1}{C}(C\|A\|)^p$ donc si $\|A\| < \frac{1}{C}$ alors la série $\left(\sum \frac{1}{C}(C\|A\|)^p\right)$ converge et la série $(\sum A^p)$ converge absolument.

Pour la somme, Il suffit de remarquer que

$$(I_n - A) \sum_{p=0}^{+\infty} A^p = \sum_{p=0}^{+\infty} A^p - \sum_{p=0}^{+\infty} A^{p+1} = \sum_{p=0}^{+\infty} A^p - \sum_{p=1}^{+\infty} A^p = I_n.$$

□

Remarques :

1. On montre ainsi que si A a une norme « petite » alors $I_n - A$ est inversible. Ce n'est qu'une condition suffisante, on peut par exemple trouver des matrices nilpotentes de très grande norme et on a encore $I_n - N$ inversible d'inverse $\sum_{p=0}^{+\infty} N^p$ qui converge car la suite stationne à 0.
2. On peut aussi considérer la série $(\sum u^p)$ pour $u \in \mathcal{L}(E)$.

6.3 Série exponentielle de matrices

Regardons maintenant la série exponentielle $\left(\sum \frac{A^p}{p!}\right)$.

Proposition 13.81

Pour toute matrice $A \in \mathcal{M}_n(\mathbf{K})$ la série exponentielle est absolument convergente. On note $\exp(A)$ sa somme.

Démonstration : Soit $A \in \mathcal{M}_n(\mathbf{K})$, on a vu que pour tout entier p , $\|A\|^p \leq C^{p-1}\|A\|^p$, comme la série $\left(\sum \frac{1}{C} \frac{(C\|A\|)^p}{p!}\right)$ converge, la série exponentielle est absolument convergente donc convergente. □

Remarque : De même pour $u \in \mathcal{L}(E)$ on peut définir $\exp(u)$.

Proposition 13.82

Soit A et B deux matrices telles que $AB = BA$ alors

$$\exp(A + B) = \exp(A) \exp(B).$$

Démonstration :

— Première preuve :

Le résultat découle des résultats sur les produits de Cauchy en les étendant au cas des espaces vectoriels normés de dimension finie . En effet si on pose $u_p = \frac{A^p}{p!}$ et $v_q = \frac{B^q}{q!}$ alors

$$\forall n \in \mathbb{N}, n! w_n = n! \sum_{p+q=n} u_p v_q = \sum_{p=0}^n \binom{n}{p} A^p B^{n-p} = (A+B)^n$$

La dernière égalité découle du fait que A et B commutent.

Comme on sait que $(\sum u_p)$ et $(\sum v_q)$ sont absolument convergentes alors $(\sum w_n)$ aussi et

$$\exp(A+B) = \sum_{n=0}^{+\infty} \frac{(A+B)^n}{n!} = \sum_{n=0}^{+\infty} w_n = \left(\sum_{p=0}^{+\infty} u_p \right) \left(\sum_{q=0}^{+\infty} v_q \right) = \exp(A) \exp(B)$$

— Deuxième preuve : On pose $u_p = \frac{A^p}{p!}$, $v_q = \frac{B^q}{q!}$ et $w_n = n! \sum_{p+q=n} u_p v_q = (A+B)^n$.

Pour tout $N \in \mathbb{N}$,

$$\sum_{n=0}^N \frac{\|w_n\|}{n!} = \sum_{n=0}^N \left\| \sum_{p+q=n} u_p v_q \right\| \leq \sum_{n=0}^N \sum_{p+q=n} C \|u_p\| \|v_q\|$$

Comme tous les termes sont positifs, la somme sur le « triangle $p+q \leq N$ » est inférieure à la somme sur « le carré $(p, q) \in [[0; N]]^2$ » :

$$\sum_{n=0}^N \|w_n\| \leq C \sum_{p=0}^N \sum_{q=0}^N \|u_p\| \|v_q\| = C \left(\sum_{p=0}^N \|u_p\| \right) \left(\sum_{q=0}^N \|v_q\| \right) \leq C \left(\sum_{p=0}^{+\infty} \|u_p\| \right) \left(\sum_{q=0}^{+\infty} \|v_q\| \right).$$

On en déduit que la série $(\sum w_n)$ est absolument convergente.

Maintenant, si on calcule

$$\sum_{n=0}^{2N} w_n = \sum_{n=0}^{2N} \sum_{p+q=n} u_p v_q.$$

On fait la somme sur le « triangle $p+q \leq 2N$ qui contient le « le carré $(p, q) \in [[0; N]]^2$ » et deux autres morceaux. Précisément,

$$\sum_{n=0}^{2N} w_n = \left(\sum_{p=0}^N u_p \right) \left(\sum_{q=0}^N u_q \right) + \sum_{(p,q) \in T_N} u_p v_q$$

faire un dessin

Maintenant,

$$\left\| \sum_{(p,q) \in T_N} u_p v_q \right\| \leq C \sum_{(p,q) \in T_N} \|u_p\| \|v_q\| \leq \left(\sum_{p=N+1}^{+\infty} \|u_p\| \right) \left(\sum_{q=0}^{+\infty} \|v_q\| \right) + \left(\sum_{p=0}^{+\infty} \|u_p\| \right) \left(\sum_{q=N+1}^{+\infty} \|v_q\| \right).$$

On en déduit que $\lim_{N \rightarrow \infty} \sum_{(p,q) \in T_N} u_p v_q = 0$ et donc

$$\exp(A+B) = \sum_{n=0}^{+\infty} \frac{(A+B)^n}{n!} = \sum_{n=0}^{+\infty} w_n = \left(\sum_{p=0}^{+\infty} u_p \right) \left(\sum_{q=0}^{+\infty} v_q \right) = \exp(A) \exp(B)$$

ATTENTION

En particulier, on voit que $I_n = \exp(0) = \exp(A - A) = \exp(A) \exp(-A)$ donc $\exp(A)$ est inversible et $\exp(-A) = \exp(A)^{-1}$.

Nous verrons lors du chapitre sur les équations différentielles l'intérêt de l'exponentielle de matrice. Voyons déjà quelques méthodes de calculs.

1. Si A est une matrice diagonale $A = \begin{pmatrix} \lambda_1 & 0 & \cdots & 0 \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \cdots & 0 & \lambda_n \end{pmatrix}$.

Pour tout $p \in \mathbb{N}$, $A^p = \begin{pmatrix} \lambda_1^p & 0 & \cdots & 0 \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \cdots & 0 & \lambda_n^p \end{pmatrix}$.

On en déduit que

$$\exp(A) = \begin{pmatrix} \exp(\lambda_1) & 0 & \cdots & 0 \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \cdots & 0 & \exp(\lambda_n) \end{pmatrix}.$$

2. Soit $A \in \mathcal{M}_n(\mathbb{K})$ et $P \in \text{GL}_n(\mathbb{K})$. On pose $B = P^{-1}AP$. On a donc que pour tout $p \in \mathbb{N}$, $B^p = P^{-1}A^pP$ et donc

$$\forall n \in \mathbb{N}, \sum_{p=0}^n \frac{B^p}{p!} = P^{-1} \left(\sum_{p=0}^n \frac{A^p}{p!} \right) P.$$

On sait que le produit matriciel est continu donc

$$\exp(B) = P^{-1} \exp(A) P.$$

Dit autrement, pour calculer l'exponentielle d'une matrice on peut la « remplacer » par une matrice semblable puis faire le changement de bases.

3. Dans le cas où A est annulé par un polynôme scindé (toujours vrai sur $\mathbb{C}$) on sait qu'elle est semblable à une matrice diagonale par blocs dont les blocs sont de la forme $\lambda I_p + N$ où N est nilpotente. D'après ce qui précède, on peut donc se ramener à calculer $\exp(\lambda I_p + N)$. De plus λI_p et N commutent on a donc

$$\exp(\lambda I_p + N) = \exp(\lambda) \exp(N).$$

Or le calcul de $\exp(N)$ est aisé car N^k est stationnaire à 0.

4. Dans le cas où A est diagonalisable. Plutôt que de faire le changement de bases qui nécessite de calculer P^{-1} , on peut utiliser les polynômes interpolateurs de Lagrange. Précisément, on

note $\pi_A = \prod_{i=1}^d (X - a_i)$ le polynôme minimal et, pour tout $j \in \llbracket 1; d \rrbracket$, $R_j = \frac{\prod_{i \neq j} (X - a_i)}{\prod_{i \neq j} (a_j - a_i)}$. On

sait alors que pour tout entier n ,

$$X^n = Q\pi_A + S$$

où

$$S = \sum_{j=1}^d a_j^d R_j$$

Dès lors,

$$\exp(A) = \sum_{n=0}^{\infty} \frac{A^n}{n!} = \sum_{n=0}^{\infty} \frac{1}{n!} \sum_{j=1}^d a_j^n R_j(A) = \sum_{j=1}^d \left(\sum_{n=0}^{\infty} \frac{1}{n!} a_j^n \right) R_j(A) = \sum_{j=1}^d \exp(a_j) R_j(A)$$

Exercices :

1. Calculer $\exp(A)$ pour $A = \begin{pmatrix} -1 & 1 & 1 \\ 1 & -1 & 1 \\ 1 & 1 & -1 \end{pmatrix}$. On peut exprimer A en fonction de I et de $M = \begin{pmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \\ 1 & 1 & 1 \end{pmatrix}$.
2. Montrer que pour toute matrice $A \in \mathcal{M}_n(\mathbb{C})$ il existe $P \in \mathbb{C}[X]$ tel que $\exp(A) = P(A)$. On pourra utiliser que $\mathbb{C}[A]$ est un sous-espace vectoriel fermé.
3. Montrer que pour $M \in \mathcal{M}_n(\mathbb{C})$, $\det(\exp(M)) = \exp(\text{tr}(M))$. On pourra trigonaliser M .

7 Parties connexes par arcs

On veut maintenant généraliser le théorème des valeurs intermédiaires.

7.1 Motivation

Le théorème des valeurs intermédiaires n'est plus vérifié dans $\mathbb{C}$ par exemple. Pour $f : t \mapsto e^{it}$ de $\mathbb{R}$ dans $\mathbb{C}$. On a $1 \in f(\mathbb{R})$, $-1 \in f(\mathbb{R})$ mais le segment $[-1, 1]$ n'est pas inclus dans $f(\mathbb{R})$.

7.2 Définition

Définition 13.83

Soit E un espace vectoriel normé, A une partie de E .

1. On appelle chemin dans A (ou arc dans A) une application **continue** $\gamma : [0, 1] \rightarrow A$.
2. Soit x et y dans A . On appelle chemin dans A de x vers y un chemin γ tel que $\gamma(0) = x$ et $\gamma(1) = y$.

Remarque : Il faut imaginer l'arc (ou le chemin) comme l'image $\gamma([0, 1])$ de l'application.

Proposition 13.84

Soit A une partie de E . On construit une relation binaire sur A en posant

$$\forall x \in A^2, x \mathcal{R} y \iff (\text{il existe un chemin dans } A \text{ de } x \text{ vers } y)$$

C'est une relation d'équivalence.

Démonstration :

- La relation est reflexive. En effet, pour tout $x \in A$, il existe un chemin reliant x à lui même. Il suffit de prendre la fonction constante égale à x .
- La relation est symétrique. En effet si pour x, y dans A on suppose que $x \mathcal{R} y$. Il existe donc un chemin γ allant de x à y . Il suffit de considérer $\tilde{\gamma}$ défini par

$$\begin{aligned} \tilde{\gamma} : [0, 1] &\rightarrow A \\ t &\mapsto \gamma(1 - t) \end{aligned}$$

C'est bien une application continue et $\tilde{\gamma}(0) = \gamma(1) = y$, $\tilde{\gamma}(1) = \gamma(0) = x$.

- La relation est transitive. Soit x, y et z dans A . On suppose que $x \mathcal{R} y$ et que $y \mathcal{R} z$. Il existe donc γ_1 et γ_2 des chemins de x vers y et de y vers z . On va construire un chemin de x en z en mettant les deux chemins bout à bout. Précisément, on pose

$$\begin{aligned} \gamma : [0, 1] &\rightarrow A \\ t &\mapsto \begin{cases} \gamma_1(2t) & \text{si } t < \frac{1}{2} \\ \gamma_2(2t - 1) & \text{sinon.} \end{cases} \end{aligned}$$

On vérifie bien que γ est continu car $\gamma_1(1) = \gamma_2(0) = y$ puis que $\gamma(0) = \gamma_1(0) = x$, $\gamma(1) = \gamma_2(1) = z$.

Faire un dessin

La relation $\mathcal{R}$ est bien une relation d'équivalence.

□

Remarque : On va alors regarder les classes d'équivalences de pour cette relation.

Définition 13.85

Une partie A est dite connexe par arcs si pour tout x et y dans A il existe un chemin de x vers y .

Remarque : Cela revient à dire qu'il n'y a qu'une classe d'équivalence pour la relation précédente.

Proposition 13.86

Les parties connexes par arcs de $\mathbf{R}$ sont les intervalles.

Démonstration :

- Il est clair que les intervalles sont connexes par arcs. En effet, pour tout $(x, y) \in A^2$ on peut poser

$$\gamma : t \mapsto (1 - t)x + ty$$

- Réciproquement, montrons que les parties connexes par arcs sont les intervalles. Soit A une partie connexe par arcs. On pourrait faire une démonstration directe avec une disjonction selon que A soit majorée ou non, minorée ou non et selon que l'éventuelle borne supérieure / inférieure appartienne ou non à A . Il est plus simple de montrer que si A est connexe par arcs alors elle est convexe et d'utiliser le résultat prouvé précédemment.

Soit x, y dans A , comme A est connexe par arcs, il existe $\gamma : [0, 1] \rightarrow \mathbf{R}$ une application continue telle que $\gamma(0) = x$, $\gamma(1) = y$ et γ prend ses valeurs dans A . On peut utiliser le théorème des valeurs intermédiaires pour montrer que toutes les valeurs z comprise entre x et y sont atteintes par γ et sont donc dans A . On en déduit que A est convexe. C'est donc un intervalle.

□

Définition 13.87 (Parties étoilées)

Soit $A \subset E$. Elle est dite étoilée s'il existe $x_0 \in A$ tel que pour tout y de A , le segment $[x_0, y]$ est inclus dans A .

Faire un dessin

Proposition 13.88

Soit A une partie de E non vide. On a
 A convexe $\Rightarrow A$ étoilée $\Rightarrow A$ connexe par arcs

Démonstration : Evident

□

ATTENTION

Quand on n'est plus dans $\mathbf{R}$, les implications ci-dessus ne sont pas des équivalences.

Faire un dessin

Définition 13.89 (Composantes connexes par arcs)

Soit $A \subset E$, on appelle composantes connexes par arcs les classes d'équivalences de la relation $\mathcal{R}$. En particulier, elles sont connexes par arcs et forment une partition de A .

Exemples :

1. Dans $A = \mathbf{R}^* \subset \mathbf{R}$ la composante connexe de 1 est $]0, +\infty[$
2. Dans $A = \mathbf{C}^* \subset \mathbf{C}$ la composante connexe de 1 est $\mathbf{C}^*$ en entier.
3. L'ensemble $A = \mathbf{R} \setminus \mathbf{Z} \subset \mathbf{R}$ a une infinité de composantes connexes.

Exercice : Soit A une partie connexe par arcs. Soit X une partie de A qui est ouverte et fermée. Montrer que $X = \emptyset$ ou $X = A$.

7.3 Image d'une partie connexe par arcs par une application continue

Nous allons pouvoir généraliser le théorème des valeurs intermédiaires.

Théorème 13.90

Soit A une partie connexe par arcs et f une application continue de A dans F . Son image $f(A)$ est connexe par arcs.

Démonstration : Soit y_1, y_2 deux éléments de $f(A)$. Il existe x_1 et x_2 tels que $f(x_i) = y_i$. Comme A est connexe par arcs, il existe un chemin γ qui va de x_1 vers x_2 . La composée $f \circ \gamma$ est alors un chemin de y_1 vers y_2 . $\square$

Remarque : C'est bien la généralisation du théorème des valeurs intermédiaires car dans le cas de $\mathbf{R}$, connexe par arcs est équivalent à intervalle.

Exercice : Montrer que $GL_n(\mathbf{R})$ n'est pas connexe par arcs et que $GL_n(\mathbf{C})$ est connexe par arcs.

Espaces préhilbertiens réels 2

1	Endomorphismes autoadjoints d'un espace euclidien	396
1.1	Définition	396
1.2	Matrice d'un endomorphisme autoadjoint	398
1.3	Théorème spectral	398
2	Endomorphismes autoadjoints positifs, définis positifs	401
2.1	Définitions	401
2.2	Caractérisation spectrale	403

Nous allons continuer l'étude des espaces euclidiens. Dans tout ce chapitre E désigne un espace euclidien.

1 Endomorphismes autoadjoints d'un espace euclidien

1.1 Définition

Définition 14.1

Soit $u \in \mathcal{L}(E)$. Il est dit autoadjoint si $u = u^*$.

C'est-à-dire si

$$\forall (x, y) \in E^2, (u(x)|y) = (x|u(y))$$

On dit aussi que u est un endomorphisme symétrique. On note $\mathcal{S}(E)$ l'ensemble des endomorphismes autoadjoints (symétriques) de E .

Exemple : Dans $E = \mathcal{M}_n(\mathbf{R})$ avec le produit scalaire usuelle $(A|B) = \text{tr}(A^\top B)$. Le morphisme $f : M \mapsto M^\top$ est autoadjoint. En effet

$$\forall (A, B) \in E^2, (f(A)|B) = (A^\top|B) = \text{tr}(AB) \text{ et } (A|f(B)) = (A|B^\top) = \text{tr}(A^\top B^\top) = \text{tr}((BA)^\top) = \text{tr}(BA).$$

On a bien $(f(A)|B) = (A|f(B))$. L'endomorphisme f est autoadjoint.

Proposition 14.2

L'ensemble $\mathcal{S}(E)$ est un sous-espace vectoriel de $\mathcal{L}(E)$.

Démonstration : On a vu que $\Phi : u \mapsto u^*$ est un endomorphisme de $\mathcal{L}(E)$. Il suffit alors d'utiliser que $\mathcal{S}(E) = \text{Ker}(\Phi - \text{id}) = \{u \in \mathcal{L}(E), \Phi(u) = u\}$. $\square$

Proposition 14.3

Soit p un projecteur. Il est autoadjoint si et seulement s'il est orthogonal.

Remarque : Rappelons qu'un projecteur est orthogonal si c'est la projection sur F parallèlement à $F^\perp$.

Démonstration :

- $\boxed{\Leftarrow}$ On suppose que p projette sur un sous-espace vectoriel F parallèlement à $F^\perp$. Pour tout couple (x, y) de vecteurs de E . Si on note

$$x = x_1 + x_2 \text{ et } y = y_1 + y_2$$

avec x_1 et x_2 dans F et y_1 et y_2 dans $F^\perp$. On a

$$(p(x)|y) = (x_1|y_1 + y_2) = (x_1|y_1) + (x_1|y_2) = (x_1|y_1).$$

De même,

$$(x|p(y)) = (x_1|y_1).$$

On a bien $p \in \mathcal{S}(E)$.

- $\boxed{\Rightarrow}$ On suppose que $p \in \mathcal{S}(E)$. Soit $x \in \text{Ker}(p)$ et $y \in \text{Im}(p)$ alors $(x|y) = (x|p(y)) = (p(x)|y) = (0|y) = 0$. On en déduit que $\text{Im}(p) \subset (\text{Ker}(p))^\perp$ mais comme $\dim \text{Im}(p) = \dim E - \dim \text{Ker}(p) = \dim(\text{Ker}(p))^\perp$ on a bien $\text{Im}(p) = (\text{Ker}(p))^\perp$.

$\square$

Remarque : On a le même résultat pour les symétries (mais il n'est pas au programme donc à savoir refaire).

- $\boxed{\Leftarrow}$ On considère la symétrie orthogonale s par rapport à F (et parallèlement à $F^\perp$). Pour tout couple (x, y) de vecteurs de E . Si on note

$$x = x_1 + x_2 \text{ et } y = y_1 + y_2$$

avec x_1 et x_2 dans F et y_1 et y_2 dans $F^\perp$. On a

$$(s(x)|y) = (x_1 - x_2|y_1 + y_2) = (x_1|y_1) + (x_1|y_2) - (x_2|y_1) - (x_2|y_2) = (x_1|y_1) - (x_2|y_2).$$

De même,

$$(x|s(y)) = (x_1|y_1) - (x_2|y_2).$$

On a bien $s \in \mathcal{S}(E)$.

- $\boxed{\Rightarrow}$ Soit s une symétrie par rapport à $F = \text{Ker}(s - \text{id})$ et parallèlement à $G = \text{Ker}(s + \text{id})$. On suppose que $s \in \mathcal{S}(E)$. Soit $x \in F$ et $y \in G$ alors $(x|y) = (s(x)|y) = (x|s(y)) = (x|-y) = -(x|y)$. On en déduit que $(x|y) = 0$. Finalement $G \subset F^\perp$ mais comme $\dim G = \dim E - \dim F = \dim F^\perp$ on a bien $G = F^\perp$.

1.2 Matrice d'un endomorphisme autoadjoint

Proposition 14.4

Soit $\mathcal{B}$ une base orthonormale de E et soit $u \in \mathcal{L}(E)$.
L'endomorphisme u est autoadjoint si et seulement si $\text{Mat}_{\mathcal{B}}(u)$ est symétrique.

Démonstration : Notons $A = \text{Mat}_{\mathcal{B}}(u)$. Comme $\mathcal{B}$ est une base orthonormée, on sait que $\text{Mat}_{\mathcal{B}}(u^*) = A^\top$. On a alors

$$u \in \mathcal{S}(E) \iff u = u^* \iff A = A^\top \iff A \text{ est symétrique}$$

□

Notation : On rappelle que l'on note $\mathcal{S}_n(\mathbf{R})$ l'ensemble des matrices symétriques de $\mathcal{M}_n(\mathbf{R})$.

ATTENTION

L'ensemble $\mathcal{S}_n(\mathbf{R})$ n'est pas stable par produit ; l'ensemble $\mathcal{S}(E)$ n'est pas stable par composition.

Exemple : Si on considère une symétrie orthogonale s . On considère une base de E obtenue en concaténant une base orthonormale de F avec une base orthonormale de $F^\perp$. Dans cette base

$$\text{Mat}(s) = \begin{pmatrix} 1 & 0 & \cdots & \cdots & \cdots & 0 \\ 0 & \ddots & \ddots & & & \vdots \\ \vdots & \ddots & 1 & \ddots & & \vdots \\ \vdots & & \ddots & -1 & \ddots & \vdots \\ \vdots & & & \ddots & \ddots & 0 \\ 0 & \cdots & \cdots & \cdots & 0 & -1 \end{pmatrix} \in \mathcal{S}_n(\mathbf{R}).$$

ATTENTION

Le résultat est faux si la base n'est pas orthonormée. Par exemple dans la base (u, v) où $u = (1, 1)$ et $v = (0, 1)$ la projection orthogonale sur (Ox) a pour matrice

$$\begin{pmatrix} 1 & 0 \\ -1 & 0 \end{pmatrix}.$$

Corollaire 14.5

On a

$$\dim \mathcal{S}(E) = \dim \mathcal{S}_n(\mathbf{R}) = \frac{n(n+1)}{2}$$

1.3 Théorème spectral

Nous allons dans ce paragraphe étudier la diagonalisabilité d'un endomorphisme autoadjoint (ou d'une matrice symétrique réelle).

Commençons par deux lemmes qui sont analogues à ceux de la réduction des isométries vectorielles.

Lemme 14.6

Soit $u \in \mathcal{S}(E)$ et F un sous-espace vectoriel stable par u . L'orthogonal $F^\perp$ est encore stable par u .

Démonstration : Soit $u \in \mathcal{S}(E)$ et F un sous-espace vectoriel stable par u . On sait que $F^\top$ est alors stable par $u^* = u$.

□

Lemme 14.7

Soit $u \in \mathcal{S}(E)$. Son polynôme caractéristique est scindé sur $\mathbf{R}$. En particulier si $\dim E > 0$ l'endomorphisme u a au moins une valeur propre réelle.

Démonstration : Notons A la matrice de u dans une base orthonormée. On sait donc que $A \in \mathcal{S}_n(\mathbf{R})$. On veut montrer que χ_A est scindé sur $\mathbf{R}$. On sait qu'il est scindé sur $\mathbf{C}$ car $\mathbf{C}$ est algébriquement clos.

Maintenant si λ est une valeur propre de A (ou de u) a priori complexe. Il existe alors $X = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} \in \mathcal{M}_{n,1}(\mathbf{C})$

non nul tel que $AX = \lambda X$. En conjuguant, on obtient $A\bar{X} = \bar{\lambda}\bar{X}$. De ce fait,

$$\bar{X}^T AX = \bar{X}^T (\lambda X) = \lambda \|X\|^2$$

où

$$\|X\|^2 = \sum_{i=1}^n \bar{x}_i x_i = \sum_{i=1}^n |x_i|^2 > 0.$$

De même,

$$\bar{X}^T AX = (\bar{X}^T A)X = (\bar{X}^T A^T)X = (A\bar{X})^T X = \bar{\lambda}\bar{X}^T X = \bar{\lambda}\|X\|^2$$

On en déduit que $\lambda = \bar{\lambda}$ et donc $\lambda \in \mathbf{R}$.

Ceci étant vrai pour toutes les valeurs propres, χ_A est bien scindé sur $\mathbf{R}$

□

Remarque : Il existe une autre preuve. On considère

$$f : x \mapsto (u(x)|x).$$

C'est une application continue. En effet c'est la composée de

$$\begin{aligned} \varphi : E &\rightarrow E \times E \\ x &\mapsto (u(x), x) \end{aligned}$$

et de

$$\begin{aligned} (|. |) : E \times E &\rightarrow \mathbf{R} \\ (x, y) &\mapsto (u(x)|y) \end{aligned}$$

Or l'application φ est continue car linéaire (et E de dimension finie) et $(|. |)$ est continue car bilinéaire.

Maintenant, comme E est de dimension finie la sphère unité $S = \{x \in E \mid \|x\| = 1\}$ est un fermé et bornée. Elle est donc compacte.

On en déduit que f est bornée et atteint ses bornes sur S .

Soit $x_0 \in S$ tel que $f(x_0)$ soit maximale. On va montrer que x_0 est un vecteur propre de u . Ce qui revient à dire que $\text{Vect}(u(x_0)) \subset \text{Vect}(x_0)$ ¹. Pour montrer que ces deux espaces vectoriels coïncident on va montrer qu'ils ont le même orthogonal.

Soit h tel que $(x_0|h) = 0$. On sait que pour tout $t \in \mathbb{R}$, $\|x_0 + th\|^2 = \|x_0\|^2 + t^2\|h\|^2 \neq 0$ d'après le théorème de Pythagore.

On peut alors regarder la fonction

$$g: \mathbb{R} \rightarrow \mathbb{R} \\ t \mapsto f\left(\frac{x_0 + th}{\|x_0 + th\|}\right)$$

D'après la définition de x_0 , cette fonction atteint son maximum en $t = 0$. Maintenant,

$$\begin{aligned} g(t) &= \left(u\left(\frac{x_0 + th}{\|x_0 + th\|}\right) \middle| \frac{x_0 + th}{\|x_0 + th\|} \right) \\ &= \frac{1}{\|x_0 + th\|^2} (u(x_0) + tu(h)|x_0 + th) \\ &= \frac{1}{1 + t^2\|h\|^2} [(u(x_0)|x_0) + t(u(x_0)|h) + t(u(h)|x_0) + t^2(u(h)|h)] \\ &= ((u(x_0)|x_0) + 2t(u(x_0)|h) + o(t))(1 + 0t + o(t)) \\ &= (u(x_0)|x_0) + 2t(u(x_0)|h) + o(t) \end{aligned}$$

On en déduit que g est dérivable en 0 et que $g'(0) = 2(u(x_0)|h) = 0$.

On a montré que $\text{Vect}(x_0)^\perp \subset \text{Vect}(u(x_0))^\perp$ ce qui implique en prenant l'orthogonal que

$$\text{Vect}(u(x_0)) = (\text{Vect}(u(x_0))^\perp)^\perp \subset (\text{Vect}(x_0)^\perp)^\perp = \text{Vect}(x_0).$$

Ce qui implique bien que x_0 est un vecteur propre.

C'est ce que l'on voulait.

Théorème 14.8 (Théorème spectral - Version endomorphisme)

Si u est un endomorphisme de E . Il est autoadjoint si et seulement si l'espace E est somme directe orthogonale des sous-espaces propres de u . C'est-à-dire si et seulement s'il existe une base orthonormale diagonalisant u .

Démonstration : On procède par double inclusion

- $\Leftarrow$ On suppose qu'il existe une base orthonormée $\mathcal{B}$ telle que $\Delta = \text{Mat}_{\mathcal{B}}(u)$ soit diagonale. En particulier Δ est symétrique donc u est autoadjoint.
- $\Rightarrow$

On procède par récurrence sur la dimension de E . On veut montrer que si u est un endomorphisme autoadjoint de E où $\dim E = n$ alors il existe une base orthonormale diagonalisant u .

1. l'inclusion est là pour traiter le cas où la valeur propre serait nulle

- Initialisation : Pour $n = 0$, il suffit de considérer la famille vide. Pour $n = 1$, il suffit de prendre un vecteur normé engendrant l'espace.
 - Hérédité : Soit $n \in \mathbb{N}^*$. On suppose que la propriété est vraie si $\dim E < n$ et on veut le montrer si $\dim E = n$. On sait maintenant que u admet une valeur propre réelle λ . On note alors e_1 un vecteur propre associé à λ que l'on peut supposer normé. On pose $F = \text{Vect}(e_1)$ qui est stable par u . En utilisant le lemme on obtient que $F^\perp$ est stable par u . Maintenant, la restriction $\tilde{u}$ de u à $F^\perp$ est un endomorphisme autoadjoint de $F^\perp$ (avec le produit scalaire induit par celui de E). Donc il existe une base $(e_2, \dots, e_n)$ orthonormée de $F^\perp$ qui diagonalise $\tilde{u}$.
- Il ne reste qu'à prendre $(e_1, \dots, e_n)$ qui est bien orthonormée car pour tout $i \geq 2$, $(e_1 | e_i) = 0$.

□

On peut aussi en donner une version matricielle :

Théorème 14.9 (Théorème spectral - Version matrice)

Soit A une matrice réelle. Elle est symétrique si et seulement s'il existe une matrice orthogonale $P \in O(n)$ telle que

$$P^{-1}AP = P^T AP = D$$

soit diagonale.

Remarque : On dit que A est orthogonalement semblable à une matrice diagonale.

ATTENTION

On ne peut pas généraliser aux matrices symétriques complexes. Par exemple $A = \begin{pmatrix} 1 & i \\ i & -1 \end{pmatrix}$ n'est pas diagonalisable car son polynôme minimal est $\mu_A = X^2$.

Il existe une extension de cette théorie à $\mathbb{C}$ pour les matrices hermitiennes qui vérifient $A^T = \bar{A}$.

2 Endomorphismes autoadjoints positifs, définis positifs

2.1 Définitions

Définition 14.10

Soit $u \in \mathcal{S}(E)$.

1. On dit que u est positif si pour tout $x \in E$, $(u(x)|x) \geq 0$.
2. On dit que u est défini positif si pour tout $x \in E \setminus \{0\}$, $(u(x)|x) > 0$.

On note $\mathcal{S}^+(E)$ et $\mathcal{S}^{++}(E)$ l'ensemble des endomorphismes autoadjoints positifs et définis positifs respectivement.

On a aussi une définition analogue pour les matrices.

Définition 14.11

Soit $A \in \mathcal{S}_n(\mathbf{R})$.

1. On dit que A est positive si pour tout $X \in \mathcal{M}_{n,1}(\mathbf{R})$, $X^\top A X \geq 0$.
2. On dit que u est définie positive si pour tout $X \in \mathcal{M}_{n,1}(\mathbf{R}) \setminus \{0\}$, $X^\top A X > 0$.

On note $\mathcal{S}_n^+(\mathbf{R})$ et $\mathcal{S}_n^{++}(\mathbf{R})$ l'ensemble des matrices symétriques positives et définies positives respectivement.

Remarque : Soit $A \in \mathcal{S}_n(\mathbf{R})$, on note a l'endomorphisme canoniquement associé à A qui est autoadjoint (ou $E = \mathcal{M}_{n,1}(\mathbf{R})$ est muni de la structure euclidienne canonique). On a alors

$$A \in \mathcal{S}_n^+(\mathbf{R}) \iff a \in \mathcal{S}^+(E) \text{ et } A \in \mathcal{S}_n^{++}(\mathbf{R}) \iff a \in \mathcal{S}^{++}(E)$$

Il suffit de recopier la définition en remarquant que pour $X \in E$,

$$(a(X)|X) = (AX|X) = (AX)^\top X = X^\top A^\top X = A^\top A X$$

Exemple : Soit $A = \begin{pmatrix} a & b \\ b & d \end{pmatrix} \in \mathcal{S}_2(\mathbf{R})$. On cherche des conditions sur les coefficients pour que $A \in \mathcal{S}_2^+(\mathbf{R})$. Soit $X = \begin{pmatrix} x \\ y \end{pmatrix}$, $X^\top A X = ax^2 + 2bxy + dy^2$.

- Si $a < 0$, la matrice A n'est pas symétrique positive car pour $X = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$, $X^\top A X = a < 0$.
- Si $a = 0$. Si $b = 0$ et $d \geq 0$ la matrice est symétrique positive car $X^\top A X = dy^2$. Par contre si $b \neq 0$ la matrice n'est pas symétrique positive car pour $X = \begin{pmatrix} x \\ 1 \end{pmatrix}$, $X^\top A X = 2bx + d$ qui n'est pas de signe constant quand x varie.
- Si $a > 0$, on peut factoriser

$$X^\top A X = a \left(x^2 + \frac{2b}{a}xy + \frac{d}{a}y^2 \right) = a \left(\left(x + \frac{b}{a}y \right)^2 + \frac{ad - b^2}{a^2}y^2 \right)$$

On voit que ce terme est toujours positif si et seulement si $ad - b^2 = \det(A)$ est positif.

Finalement $A \in \mathcal{S}_2^+(\mathbf{R})$ si et seulement si $(a = b = 0 \text{ et } d \geq 0)$ ou $(a > 0 \text{ et } \det(A) \geq 0)$.

De même ce calcul montre que $A \in \mathcal{S}_2^{++}(\mathbf{R})$ si et seulement si $(a > 0 \text{ et } \det(A) > 0)$.

ATTENTION

Comme on le voit ci-dessus, dire qu'une matrice est symétrique positive ne revient pas à dire que ses coefficients sont positifs. Par exemple

$$\begin{pmatrix} 2 & -1 \\ -1 & 2 \end{pmatrix} \in \mathcal{S}_2^+(\mathbf{R}) \text{ et } \begin{pmatrix} 1 & 2 \\ 2 & 1 \end{pmatrix} \notin \mathcal{S}_2^+(\mathbf{R})$$

2.2 Caractérisation spectrale

Proposition 14.12

Soit $u \in \mathcal{S}(E)$.

1. L'endomorphisme u est autoadjoint positif si et seulement si $\text{Sp}(u) \subset \mathbf{R}_+$.
2. L'endomorphisme u est autoadjoint défini positif si et seulement si $\text{Sp}(u) \subset \mathbf{R}_+^*$.

Démonstration :

1. – $\Rightarrow$ On suppose que $u \in \mathcal{S}^+(E)$. Soit $\lambda \in \text{Sp}(u)$. On considère x un vecteur propre (non nul) associé à la valeur propre λ .

$$(u(x)|x) = (\lambda x|x) = \lambda(x|x) = \lambda \|x\|^2$$

On en déduit (car $\|x\|^2 > 0$) que $\lambda = \frac{(u(x)|x)}{\|x\|^2} \geq 0$.

Cela montre que $\text{Sp}(u) \subset \mathbf{R}_+$.

- $\Leftarrow$ On suppose que $\text{Sp}(u) \subset \mathbf{R}_+$. D'après le théorème spectral, on sait que u est diagonalisable. Soit $\mathcal{B} = (e_1, \dots, e_n)$ une base orthonormée de vecteurs propres de u . Pour tout $i \in \llbracket 1 ; n \rrbracket$, on note $\lambda_i \geq 0$ la valeur propre associée à e_i .

Pour $x \in E$, on peut le décomposer dans la base $\mathcal{B} : x = \sum_{i=1}^n x_i e_i$ où pour tout $i \in \llbracket 1 ; n \rrbracket$, $x_i = e_i^*(x)$. On a alors

$$(u(x)|x) = \left(\sum_{i=1}^n \lambda_i x_i e_i \middle| \sum_{j=1}^n x_j e_j \right) = \sum_{i=1}^n \sum_{j=1}^n \lambda_i x_i x_j (e_i | e_j) = \sum_{i=1}^n \lambda_i x_i^2 \geq 0$$

La dernière égalité vient du fait que $\mathcal{B}$ est une base orthonormée.

Cela montre bien que $u \in \mathcal{S}^+(E)$.

2. – $\Rightarrow$ On suppose que $u \in \mathcal{S}^{++}(E)$. Soit $\lambda \in \text{Sp}(u)$. On considère x un vecteur propre (non nul) associé à la valeur propre λ . Le même calcul que ci-dessus donne $\lambda = \frac{(u(x)|x)}{\|x\|^2} > 0$.

Cela montre que $\text{Sp}(u) \subset \mathbf{R}_+^*$.

- $\Leftarrow$ On suppose que $\text{Sp}(u) \subset \mathbf{R}_+^*$. Avec les mêmes notations que ci-dessus, pour $x \in E \setminus \{0\}$,

$$(u(x)|x) = \sum_{i=1}^n \lambda_i x_i^2 \geq 0$$

Comme le vecteur x n'est pas nul, il existe i_0 tel que $x_{i_0} \neq 0$. Comme $\lambda_{i_0} > 0$ on a alors

$$(u(x)|x) \geq \lambda_{i_0} x_{i_0}^2 > 0$$

Cela montre bien que $u \in \mathcal{S}^{++}(E)$.

□

Corollaire 14.13

Soit $A \in \mathcal{S}_n(\mathbf{R})$.

1. La matrice A est symétrique positive si et seulement si $\text{Sp}(A) \subset \mathbf{R}_+$.
2. La matrice A est symétrique définie positive si et seulement si $\text{Sp}(A) \subset \mathbf{R}_+^*$.

Corollaire 14.14

Soit $A \in \mathcal{S}_2(\mathbf{R})$.

1. La matrice A est positive si et seulement si $\text{tr}(A) \geq 0$ et $\det(A) \geq 0$.
2. La matrice A est définie positive si et seulement si $\text{tr}(A) > 0$ et $\det(A) > 0$.

Fonctions à valeurs vectorielles

1	Dérivabilité	405
1.1	Dérivabilité en un point	406
1.2	Opérations	408
1.3	Dérivées successives	411
2	Intégration sur un segment	412
2.1	Définitions	412
2.2	Propriétés de l'intégrale	414
3	Intégrale fonction de sa borne supérieure	416
3.1	Théorème fondamental de l'analyse	416
3.2	Inégalités des accroissements finis	417
4	Formules de Taylor	417
4.1	Formule de Taylor avec reste intégral	417
4.2	Formule de Taylor-Young	418
5	Suites et séries de fonctions à valeurs vectorielles	419
5.1	Généralités	419
5.2	Continuité et théorème de la double limite	421
5.3	Intégration et dérivation	423

Dans tout ce chapitre I désignera un intervalle (non trivial) de $\mathbf{R}$ et F un $\mathbf{K}$ -espace vectoriel normé de dimension finie (notée p). On note $||.||$ la norme.

Nous allons définir la dérivabilité et l'intégration pour les fonctions de I dans F .

1 Dérivabilité

1.1 Dérivabilité en un point

Définition 15.1

Soit f une fonction de I dans F et $t_0 \in I$. On dit que f est dérivable en t_0 si la fonction

$$\begin{aligned} \varphi_{t_0} : I \setminus \{t_0\} &\rightarrow F \\ t &\mapsto \frac{f(t) - f(t_0)}{t - t_0} \end{aligned}$$

a une limite en t_0 .

Dans ce cas, la fonction f est dite dérivable en t_0 , la limite s'appelle dérivée de f en t_0 et on note

$$f'(t_0) = \lim_{t \rightarrow t_0} \frac{f(t) - f(t_0)}{t - t_0} = \lim_{h \rightarrow 0} \frac{f(t_0 + h) - f(t_0)}{h}.$$

Remarques :

1. L'élément $f'(t_0)$ est un élément de F .
2. Dans le cas où $F = \mathbf{R}^2$ ou $F = \mathbf{R}^3$, une application $f : t \rightarrow F$ peut être imaginée comme le déplacement d'un point dans F en fonction du paramètre t (en voyant F comme un espace affine et plus un espace vectoriel). Dans ce cas, $f'(t_0)$ s'appelle le vecteur vitesse.

Proposition 15.2 (Expression en coordonnées)

Avec les mêmes notations. Si on note $\mathcal{B} = (e_1, \dots, e_p)$ une base de F , on pose pour tout $i \in \llbracket 1 ; p \rrbracket$,

$$f_i = e_i^* \circ f : I \rightarrow \mathbf{K}$$

(c'est-à-dire que pour tout t dans I , $f(t) = \sum_{i=1}^p f_i(t) e_i$)

1. La fonction f est dérivable en t_0 si et seulement si toutes les f_i sont dérivables en t_0 .

2. Dans ce cas, $f'(t_0) = \sum_{i=1}^p f'_i(t_0) e_i$

Démonstration : Ce n'est qu'un cas particulier du théorème analogue sur les limites. $\square$

Exemple : On considère l'espace vectoriel $F = \mathcal{M}_2(\mathbf{R})$ et $\mathcal{B}$ la base canonique. Pour savoir si une application est dérivable, on peut regarder coordonnée par coordonnée. Par exemple si on étudie

$$f : t \mapsto f(t) = \begin{pmatrix} \cos t & -\sin t \\ \sin t & \cos t \end{pmatrix}.$$

Chaque application coordonnée est dérivable en $t_0 \in \mathbf{R}$ donc f est dérivable en t_0 et

$$f'(t_0) = \begin{pmatrix} -\sin t_0 & -\cos t_0 \\ \cos t_0 & -\sin t_0 \end{pmatrix} = f\left(t_0 + \frac{\pi}{2}\right)$$

Nous allons donner une définition « à la Taylor-Young » de la dérivabilité, il faut au préalable donner un sens à $o(h)$.

Définition 15.3

On note, pour tout entier k , $o_{h \rightarrow 0}(h^k)$ une fonction définie d'un voisinage de 0 dans F de la forme $h \mapsto h^k \varepsilon(h)$ où ε tend vers 0 quand h tend vers 0.

Remarque : On peut de même définir $o_{t \rightarrow t_0}((t - t_0)^k)$. La plupart du temps, on omettra de préciser au voisinage de quel point on travaille en indice.

Proposition 15.4 (Interprétation à la Taylor-Young)

Avec les notations précédentes, f est dérivable en t_0 si et seulement s'il existe $\alpha \in F$ tel que

$$f(t_0 + h) = f(t_0) + h\alpha + o_0(h).$$

Dans ce cas, $f'(t_0) = \alpha$.

Démonstration : La démonstration est analogue au cas réel.

$$\begin{aligned} \exists \alpha \in F, f(t_0 + h) = f(t_0) + h\alpha + o_0(h) &\iff \exists \alpha \in F, \frac{f(t_0 + h) - f(t_0)}{h} = \alpha + o_0(1) \\ &\iff \lim_{h \rightarrow 0} \frac{f(t_0 + h) - f(t_0)}{h} = \alpha \end{aligned}$$

□

Exemple : En reprenant notre exemple ci-dessus,

$$\begin{aligned} f(t_0 + h) &= \begin{pmatrix} \cos(t_0 + h) & -\sin(t_0 + h) \\ \sin(t_0 + h) & \cos(t_0 + h) \end{pmatrix} \\ &= \begin{pmatrix} \cos t_0 - h \sin t_0 + o(h) & -\sin t_0 - h \cos t_0 + o(h) \\ \sin t_0 + h \cos t_0 + o(h) & \cos t_0 - h \sin t_0 + o(h) \end{pmatrix} \\ &= \begin{pmatrix} \cos t_0 & -\sin t_0 \\ \sin t_0 & \cos t_0 \end{pmatrix} + h \begin{pmatrix} -\sin t_0 & -\cos t_0 \\ \cos t_0 & -\sin t_0 \end{pmatrix} + o(h) \end{aligned}$$

Définition 15.5 (Dérivabilité à droite et à gauche)

Avec les mêmes notations

1. Si t_0 n'est pas la borne inférieure de I , f est dite dérivable à gauche en t_0 si φ_{t_0} a une limite quand $t \rightarrow t_0^-$. La dérivée à gauche est notée $f'_g(t_0)$.
2. Si t_0 n'est pas la borne supérieure de I , f est dite dérivable à droite en t_0 si φ_{t_0} a une limite quand $t \rightarrow t_0^+$. La dérivée à droite est notée $f'_d(t_0)$.

Définition 15.6

Soit f une fonction de I dans F .

1. Elle est dite dérivable si elle est dérivable en tout point t_0 de I . On note alors f' sa fonction dérivée définie sur I par :

$$f' : t \mapsto f'(t).$$

2. Elle est dite de classe $\mathcal{C}^1$ si elle est dérivable et que sa fonction dérivée est continue.

1.2 Opérations

Proposition 15.7

Une combinaison linéaire de fonction dérivables est dérivable. L'ensemble des fonctions dérivables est donc un sous-espace vectoriel de l'ensemble des fonctions.

Démonstration : Il suffit de recopier la preuve usuelle.

Proposition 15.8

Soit G un espace vectoriel normé de dimension finie. Soit f une fonction de I dans F et $L : F \rightarrow G$ une application linéaire.

1. Si f est dérivable en t_0 alors $L \circ f$ est dérivable en t_0 et

$$(L \circ f)'(t_0) = L(f'(t_0))$$

2. Si f est dérivable alors $L \circ f$ est dérivable et

$$(L \circ f)' = L \circ f'$$

Démonstration :

1. On étudie

$$\frac{L(f(t_0 + h)) - L(f(t_0))}{h} = L\left(\frac{f(t_0 + h) - f(t_0)}{h}\right) \xrightarrow{h \rightarrow 0} L(f'(t_0))$$

car L est continue (application linéaire entre deux espaces vectoriels normés de dimension finie). On en déduit bien que $L \circ f$ est dérivable en t_0 et que $(L \circ f)'(t_0) = L(f'(t_0))$.

2. Il suffit d'appliquer 1. en tout t_0 de I .

□

Remarque : Ce résultat sera généralisée dans le chapitre sur le calcul différentiel (chain rule)

Exemple : On reprend notre application $f : I \rightarrow F = \mathcal{M}_2(\mathbb{R})$ définie par :

$$f : t \mapsto f(t) = \begin{pmatrix} \cos t & -\sin t \\ \sin t & \cos t \end{pmatrix}.$$

Pour toute matrice $A \in \mathcal{M}_2(\mathbb{R})$, on peut regarder

$$\begin{aligned} L : F &\rightarrow F \\ M &\mapsto AM \end{aligned}$$

Elle est linéaire ; on en déduit que

$$L \circ f : t \mapsto Af(t)$$

est dérivable et de dérivée : $t \mapsto Af'(t)$.

Proposition 15.9

Soit G, H des espaces vectoriels normés de dimension finie. Soit f et g des fonctions de I dans F et de I dans G . Soit $B : F \times G \rightarrow H$ une application bilinéaire.

1. Si f et g sont dérivables en t_0 alors $t \mapsto B(f(t), g(t))$ est dérivable en t_0 et

$$(B(f, g))'(t_0) = B(f'(t_0), g(t_0)) + B(f(t_0), g'(t_0)).$$

2. Si f et g sont dérivables alors $t \mapsto B(f(t), g(t))$ est dérivable et

$$(B(f, g))' = B(f', g) + B(f, g').$$

Démonstration : La preuve est semblable à la preuve de la dérivabilité d'un produit. Précisément, pour tout h proche de 0

$$B(f(t_0 + h), g(t_0 + h)) = B(f(t_0) + hf'(t_0) + h\varepsilon_1(h), g(t_0) + hg'(t_0) + h\varepsilon_2(h))$$

où ε_1 et ε_2 tendent vers 0 quand $h \rightarrow 0$. On a donc, par bilinéarité :

$$B(f(t_0 + h), g(t_0 + h)) = B(f(t_0), g(t_0)) + h[B(f(t_0), g'(t_0)) + B(f'(t_0), g(t_0))] + h\alpha(h)$$

où

$$\alpha(h) = B(f(t_0), \varepsilon_2(h)) + hB(f'(t_0), g'(t_0)) + hB(f'(t_0), \varepsilon_2(h)) + hB(\varepsilon_1(h), g'(t_0)) + hB(\varepsilon_1(h), \varepsilon_2(h)) + B(\varepsilon_1(h), g(t_0))$$

Comme B est bilinéaire et que F, G sont de dimension finie, il existe $C > 0$ tel que

$$\forall (x, y) \in F \times G, \|B(x, y)\|_H \leq C\|x\|_F \cdot \|y\|_G.$$

Montrons que les 6 termes de α tendent vers 0 quand h tend vers 0.

- $\|B(f(t_0), \varepsilon_2(h))\|_H \leq C\|f(t_0)\|_F \cdot \|\varepsilon_2(h)\|_G \xrightarrow{h \rightarrow 0} 0$ car $\varepsilon_2(h) \xrightarrow{h \rightarrow 0} 0$.
- $\|hB(f'(t_0), g'(t_0))\|_H \leq |h| \cdot \|B(f'(t_0), g'(t_0))\|_H \xrightarrow{h \rightarrow 0} 0$
- les autres termes se traitent de manière similaire.

On en déduit que $\lim_{h \rightarrow 0} \alpha(h) = 0$, c'est-à-dire que

$$B(f(t_0 + h), g(t_0 + h)) = B(f(t_0), g(t_0)) + h[B(f(t_0), g'(t_0)) + B(f'(t_0), g(t_0))] + o(h)$$

ce qui signifie que $B(f, g)$ est dérivable en t_0 et que

$$(B(f, g))'(t_0) = B(f'(t_0), g(t_0)) + B(f(t_0), g'(t_0)).$$

□

On peut étendre cela aux applications multilinéaires.

Proposition 15.10

Soit $p \in \mathbb{N}^*$. Soit $\varphi : \prod_{i=1}^n F_i \rightarrow F$ une application n -linéaire. On considère $f_1, \dots, f_n$ des applications dérivables de I dans F_i et on pose

$$\begin{aligned} \Phi : I &\rightarrow F \\ t &\mapsto \varphi(f_1(t), \dots, f_n(t)) \end{aligned}$$

L'application Φ est dérivable et pour tout $t \in I$,

$$\Phi'(t) = \sum_{k=1}^n \varphi(f_1(t), \dots, f_{k-1}(t), f'_k(t), f_{k+1}(t), \dots, f_n(t))$$

Démonstration : La démonstration est identique à celle du cas bilinéaire en plus technique. $\square$

Exemples :

1. Soit $t \mapsto A(t)$ et $t \mapsto B(t)$ deux applications dérivables de I dans $\mathcal{M}_n(\mathbf{K})$, l'application $\varphi : t \mapsto A(t)B(t)$ est dérivable et

$$\forall t \in I, \varphi'(t) = A'(t)B(t) + A(t)B'(t).$$

2. Dans le cas où F est un espace euclidien, on peut regarder le produit scalaire $(u, v) \mapsto (u|v)$ qui est bilinéaire. On en déduit que si f et g sont deux fonctions dérivables,

$$(f|g)' = (f'|g) + (f|g').$$

Dans le cas où $f = g$ on obtient donc

$$(\|f\|^2)' = 2(f'|f).$$

Dans le cas où $\|f\|$ ne s'annule pas, on peut alors composer par $\sqrt{\cdot}$. On obtient que

$$\|f\|' = \frac{(\|f\|^2)'}{2\sqrt{\|f\|^2}} = \frac{(f|f')}{\|f\|}.$$

En particulier si $\|f\|$ est constante, on obtient que f et f' sont orthogonaux. *résultat connu en physique.*

3. Si on considère une fonction $M : I \rightarrow \mathcal{M}_n(\mathbf{C})$ dérivable. Pour tout $i \in \llbracket 1; n \rrbracket$, on note $C_i(t)$ la i -ème colonne de la matrice $M(t)$. La fonction $\Phi : t \mapsto \det(M) = \det(C_1(t), \dots, C_n(t))$ est dérivable et

$$\forall t \in I, \Phi'(t) = \sum_{k=1}^n \det(C_1(t), \dots, C_{k-1}(t), C'_k(t), C_{k+1}(t), \dots, C_n(t))$$

Exercice : Pour tout n on pose

$$D_n(x) = \begin{vmatrix} x & 1 & 0 & \cdots & 0 \\ \frac{x^2}{2!} & \ddots & \ddots & & \vdots \\ \frac{x^3}{3!} & \frac{x^2}{2!} & \ddots & \ddots & 0 \\ \vdots & & \ddots & \ddots & 1 \\ \frac{x^n}{n!} & \frac{x^{n-1}}{(n-1)!} & \cdots & \frac{x^2}{2!} & x \end{vmatrix}$$

Calculer $D_n(x)$. On pourra commencer par calculer sa dérivée.

Proposition 15.11

Soit f une fonction de I dans F . Soit φ une fonction de J dans I où J est un intervalle de $\mathbf{R}$.

1. Soit $x_0 \in J$, on pose $t_0 = \varphi(x_0)$. Si φ est dérivable en x_0 et f est dérivable en t_0 alors $f \circ \varphi$ est dérivable en x_0 et

$$(f \circ \varphi)'(x_0) = \varphi'(x_0)f'(\varphi(x_0))$$

2. Si f et φ sont dérivables alors $f \circ \varphi$ aussi et

$$(f \circ \varphi)' = \varphi' \times (f' \circ \varphi).$$

Démonstration : La encore, la preuve usuelle. Pour h proche de 0

$$\begin{aligned} f(\varphi(x_0 + h)) &= f(\varphi(x_0) + h\varphi'(x_0) + h\varepsilon_1(h)) \\ &= f(\varphi(x_0)) + (h\varphi'(x_0) + h\varepsilon_1(h))f'(\varphi(x_0)) + (h\varphi'(x_0) + h\varepsilon_1(h))\varepsilon_2(h\varphi'(x_0) + h\varepsilon_1(h)) \\ &= f(\varphi(x_0)) + h\varphi'(x_0)f'(\varphi(x_0)) + h\alpha(h) \end{aligned}$$

où $\alpha(h) \xrightarrow{h \rightarrow 0} 0$.

□

1.3 Dérivées successives

Définition 15.12

Soit f une fonction de I dans F . On définit le fait que f soit de classe $\mathcal{C}^n$ par récurrence :

- On dit que f est de classe $\mathcal{C}^0$ si f est continue.
- Pour tout $n \geq 1$, on dit que f est de classe $\mathcal{C}^n$ si f est dérivable et que f' est de classe $\mathcal{C}^{n-1}$.

On dit alors que f est de classe $\mathcal{C}^\infty$, si pour tout entier n elle est de classe $\mathcal{C}^n$.

On notera $\mathcal{C}^n(I, F)$ et $\mathcal{C}^\infty(I, F)$ l'ensemble des fonctions de classes $\mathcal{C}^n$ (resp. $\mathcal{C}^\infty$) de I dans F .

Notation : Soit f de classe $\mathcal{C}^n$, on pose

$$f^{(0)} = f; f^{(1)} = f' \text{ et } \forall k \in \llbracket 1; n \rrbracket f^{(k)} = (f^{(k-1)})'$$

Remarque : Si $\mathcal{B} = (e_1, \dots, e_p)$ est une base de F . On pose pour tout $i \in \llbracket 1; p \rrbracket$, $f_i = e_i^* \circ f$. On a alors que f est de classe $\mathcal{C}^n$ si et seulement si tous les f_i sont de classe $\mathcal{C}^n$.

Proposition 15.13

On reprend les notations précédentes.

1. Soit $(f, g) \in \mathcal{C}^n(I, F)^2$ et $(\lambda, \mu) \in \mathbf{K}^2$, $\lambda f + \mu g \in \mathcal{C}^n(I, F)$ et $(\lambda f + \mu g)^{(n)} = \lambda f^{(n)} + \mu g^{(n)}$.
En particulier, $\mathcal{C}^n(I, F)$ et $\mathcal{C}^\infty(I, F)$ sont des espaces vectoriels.
2. Soit $f \in \mathcal{C}^n(I, F)$ et $L \in \mathcal{L}(F, G)$, $L \circ f \in \mathcal{C}^n(I, G)$ et $(L \circ f)^{(n)} = L \circ f^{(n)}$
3. Soit B une application bilinéaire de F dans G , $(f, g) \in \mathcal{C}^n(I, F)$ alors $B(f, g) \in \mathcal{C}^n(I, G)$ et

$$(B(f, g))^{(n)} = \sum_{p=0}^n \binom{n}{p} B(f^{(p)}, g^{(n-p)}) \quad (\text{Formule de Leibniz})$$

4. Si $f \in \mathcal{C}^n(I, F)$ et $\varphi \in \mathcal{C}^n(J, I)$ alors $f \circ \varphi \in \mathcal{C}^n(J, F)$ et

$$(f \circ \varphi)^{(n)} = (\text{formule de Faà di Bruno})$$

2 Intégration sur un segment

On s'intéresse maintenant à l'intégration. On se contente de l'intégration des fonctions continues par morceaux sur **un segment**. On notera donc $I = [a, b]$

2.1 Définitions

Définition 15.14 (Fonctions continues par morceaux)

Une fonction f de I dans F est dite *continue par morceaux* s'il existe une subdivision $\sigma = (t_0 = a < t_1 < \dots < t_n = b)$ telle que

- i) Pour tout $i \in \llbracket 1; n \rrbracket$, $f|_{[t_{i-1}, t_i[}$ est continue
- ii) La fonction f à des limites à gauches en tout les t_i (pour $i \in \llbracket 0; n-1 \rrbracket$) et à droite en tout les t_i (pour $i \in \llbracket 1; n \rrbracket$).

Remarques :

1. Soit f une fonction continue par morceaux de I dans F . Soit $\mathcal{B} = (e_1, \dots, e_p)$ une base de F . Les fonctions $e_i^* \circ f$ sont continues par morceaux sur I .
2. On notera $\mathcal{CM}(I, F)$ l'espace vectoriel des fonctions continues par morceaux.

Proposition-Définition 15.15

Soit f une fonction continue par morceaux de I dans F . Soit $\mathcal{B} = (e_1, \dots, e_p)$ une base de F ,

$$\sum_{i=1}^p \left(\int_a^b e_i^* \circ f \right) e_i$$

ne dépend pas du choix de la base. On appelle intégrale de f sur $[a, b]$ et on note $\int_a^b f$ ce terme.

Remarque : Cela correspond à ce qui a été fait pour définir l'intégrale d'une fonction complexe en posant

$$\int_a^b f = \left(\int_a^b \Re(f) \right) + i \left(\int_a^b \Im(f) \right).$$

Démonstration : Fixons les notations. On se donne $\mathcal{B} = (e_1, \dots, e_p)$ et $\mathcal{B}' = (e'_1, \dots, e'_p)$ deux bases de F . Soit P la matrice de changement de base de $\mathcal{B}$ à $\mathcal{B}'$ et $Q = P^{-1}$. On note (α_{ij}) les coefficients de Q .

Maintenant, on note $f_i = e_i^* \circ f$ et $g_i = (e'_i)^* \circ f$.

On considère pour tout $t \in I$, $X(t) = \text{Mat}_{\mathcal{B}}(f(t)) = \begin{pmatrix} f_1(t) \\ \vdots \\ f_n(t) \end{pmatrix}$ et $Y(t) = \text{Mat}_{\mathcal{B}'}(f(t)) = \begin{pmatrix} g_1(t) \\ \vdots \\ g_n(t) \end{pmatrix}$.

Les formules de changement de base donnent que pour tout $t \in I$, $Y(t) = P^{-1}X(t) = QX(t)$. C'est-à-dire que pour tout t dans I et tout i dans $\llbracket 1; p \rrbracket$

$$g_i(t) = \sum_{j=1}^p \alpha_{ij} f_j(t)$$

En intégrant on obtient

$$\int_a^b g_i = \int_a^b \sum_{j=1}^p \alpha_{ij} f_j = \sum_{j=1}^p \int_a^b \alpha_{ij} f_j.$$

En faisant la somme on obtient

$$\begin{aligned} \sum_{i=1}^p \left(\int_a^b g_i \right) e'_i &= \sum_{i=1}^p \left(\sum_{j=1}^p \int_a^b \alpha_{ij} f_j \right) e'_i \\ &= \sum_{j=1}^p \left[\left(\int_a^b f_j \right) \sum_{i=1}^p \alpha_{ij} e'_i \right] \\ &= \sum_{j=1}^p \left(\int_a^b f_j \right) e_j \end{aligned}$$

□

Notation : On notera

$$\int_{[a,b]} f \text{ ou } \int_a^b f \text{ ou } \int_a^b f(t) dt.$$

2.2 Propriétés de l'intégrale

Théorème 15.16 (Propriétés)

1. L'intégrale $\int_a^b$ est une application linéaire de l'espace vectoriel $\mathcal{CM}(I, F)$ dans F .
2. Soit $c \in [a, b]$ et $f \in \mathcal{CM}(I, F)$,

$$\int_a^b f = \int_a^c f + \int_c^b f \quad (\text{Relation de Chasles})$$

3. Soit $f \in \mathcal{CM}(I, F)$, $\|f\| \in \mathcal{CM}(I, \mathbf{R})$ et

$$\left\| \int_a^b f \right\| \leq \int_a^b \|f\| \quad (\text{Inégalité triangulaire})$$

4. Soit $f \in \mathcal{CM}(I, F)$, $\|f\| \in \mathcal{CM}(I, \mathbf{R})$ et

$$\left\| \int_a^b f \right\| \leq (b-a) \sup_{t \in [a, b]} \|f(t)\| \quad (\text{Formule de la moyenne})$$

Démonstration :

1. On considère une base $\mathcal{B} = (e_1, \dots, e_p)$. Soit $(f, g) \in \mathcal{CM}(I, F)^2$ et $\lambda \in \mathbf{K}$, il suffit de voir que coordonnées par coordonnées, on a

$$\int_a^b e_i^*(\lambda f + g) = \int_a^b \lambda e_i^*(f) + e_i^*(g) = \lambda \int_a^b e_i^*(f) + \int_a^b e_i^*(g).$$

2. Là encore, il suffit de décomposer dans une base et de revenir la propriété analogue pour l'intégrale des fonctions scalaires.
3. On ne peut faire une preuve en « décomposant ». En effet si on pose encore $\mathcal{B} = (e_1, \dots, e_p)$ et $f_i = e_i^* \circ f$. On a

$$\begin{aligned} \left\| \int_a^b f \right\| &= \left\| \int_a^b \sum_i f_i e_i \right\| \\ &= \left\| \sum_i \left(\int_a^b f_i \right) e_i \right\| \\ &\leq \sum_i \left| \int_a^b f_i \right| \|e_i\| \\ &\leq \sum_i \int_a^b |f_i| \|e_i\| = \int_a^b \sum_i |f_i| \|e_i\| \end{aligned}$$

Mais $\sum_i |f_i| \|e_i\| \neq \|f\|$ et l'inégalité n'est pas dans le bon sens.

La preuve sera établie après les sommes de Riemann.

□

Définition 15.17 (Sommes de Riemann)

Soit $f : I \rightarrow F$ et $\sigma = (t_0 = a < t_1 < \dots < t_n = b)$ une subdivision de I . On se donne un élément $\alpha = (\alpha_1, \dots, \alpha_n) \in \prod_{i=1}^n [t_{i-1}, t_i]$. On appelle somme de Riemann associée à f , σ et α et on note $R(f, \sigma, \alpha)$,

$$R(f, \sigma, \alpha) = \sum_{i=1}^n (t_i - t_{i-1}) f(\alpha_i) \in F.$$

Remarque : Dans le cadre du programme, on ne considère que des subdivisions à pas réguliers avec α_i pris « à gauche ». C'est-à-dire, σ_n définie par $t_i = a + i \frac{b-a}{n}$, et pour tout $i \in \llbracket 1; n \rrbracket$, $\alpha_i = t_{i-1}$.

C'est-à-dire

$$R_n(f) = \frac{b-a}{n} \sum_{i=0}^{n-1} f\left(a + i \frac{b-a}{n}\right)$$

Théorème 15.18

Soit $f \in \mathcal{C}^0(I, F)$,

$$\lim_{n \rightarrow \infty} R_n(f) = \int_a^b f$$

Démonstration : Il suffit de décomposer dans une base. Précisément, soit $\mathcal{B} = (e_1, \dots, e_p)$ une base. On voit que pour tout entier n ,

$$R_n(f) = \frac{b-a}{n} \sum_{i=0}^{n-1} f\left(a + i \frac{b-a}{n}\right) = \frac{b-a}{n} \sum_{i=0}^{n-1} \sum_{j=1}^p f_j\left(a + i \frac{b-a}{n}\right) e_j = \sum_{j=1}^p R_n(f_j) e_j$$

Dans le cours de première année, on a vu le résultat analogue pour les fonctions à valeurs réelles, c'est-à-dire que, pour tout $i \in \llbracket 1; p \rrbracket$,

$$\lim_{n \rightarrow \infty} R_n(f_i) = \int_a^b f_i$$

Cela implique que

$$\lim_{n \rightarrow \infty} R_n(f) = \sum_{i=1}^p \left(\int_a^b f_i \right) e_i = \int_a^b f$$

□

Démonstration de l'inégalité triangulaire : Avec les notations précédentes, on sait que

$$R_n(f) \xrightarrow{n \rightarrow \infty} \int_a^b f$$

et que

$$R_n(\|f\|) \xrightarrow{n \rightarrow \infty} \int_a^b \|f\|$$

Il suffit alors de remarquer que

$$\|R_n(f)\| = \left\| \frac{b-a}{n} \sum_{i=1}^n f\left(a + i \frac{b-a}{n}\right) \right\| \leq \frac{b-a}{n} \sum_{i=1}^n \|f\| \left(a + i \frac{b-a}{n}\right) = R_n(\|f\|).$$

En passant à la limite on a bien,

$$\left\| \int_a^b f \right\| \leq \int_a^b \|f\|.$$

□

3 Intégrale fonction de sa borne supérieure

3.1 Théorème fondamental de l'analyse

Théorème 15.19 (Théorème fondamental de l'analyse)

Soit I un intervalle (non nécessairement un segment) et f une fonction de I dans F et $a \in I$. Si f est continue, la fonction $H : x \mapsto \int_a^x f(t)dt$ définie sur I est de classe $\mathcal{C}^1$. De plus, $H' = f$, de ce fait, H est l'unique primitive de f qui s'annule en a .

Démonstration : Il suffit, via une base de F , de revenir à la démonstration classique dans le cas des fonctions scalaires. □

Remarque : Avec les mêmes hypothèses, $x \mapsto \int_x^a f(t)dt$ est une primitive de $-f$.

Corollaire 15.20

Soit f une fonction continue sur I et H une primitive de f sur I ,

$$\int_a^b f = [H]_a^b$$

Démonstration : Il suffit d'utiliser que H étant une primitive de f , il existe une constante C telle que

$$\forall x \in I, H(x) = \int_a^x f(t)dt + C.$$

On en déduit que

$$H(b) - H(a) = \int_a^b f(t)dt + C - \int_a^a f(t)dt - C = \int_a^b f(t)dt.$$

□

3.2 Inégalités des accroissements finis

ATTENTION

Le lemme de Rolle ou le théorème des accroissements finis ne sont plus vraie quand $F \neq \mathbf{R}$. En effet on peut « tourner » autour d'un éventuel zéro. Par exemple

$$f : t \mapsto e^{it} - 1$$

s'annule en 0 et en 2π , pourtant, $\forall t \in [0, 2\pi], f'(t) = ie^{it} \neq 0$.

Proposition 15.21 (Inégalités des accroissements finis)

Soit f une fonction de classe $\mathcal{C}^1$ de I dans F . Soit $a, b \in I$ avec $a \leq b$.

$$\forall (a, b) \in I^2, \|f(b) - f(a)\| \leq |b - a| \sup_{t \in [a, b]} \|f'(t)\|$$

Remarque : La borne supérieure existe (c'est même un maximum) car la fonction f' est continue sur le segment $[a, b]$

Démonstration : Avec les notations de l'énoncé,

$$f(b) - f(a) = \int_a^b f'(t) dt.$$

Par inégalité triangulaire,

$$\|f(b) - f(a)\| \leq \int_a^b \|f'(t)\| dt$$

On peut alors utiliser l'inégalité de la moyenne

$$\|f(b) - f(a)\| \leq \int_a^b \|f'(t)\| dt \leq |b - a| \sup_{t \in [a, b]} \|f'(t)\|.$$

□

4 Formules de Taylor

4.1 Formule de Taylor avec reste intégral

Proposition 15.22

Soit $n \in \mathbf{N}$. Soit f une fonction de classe $\mathcal{C}^{n+1}$ de I vers F . Pour tout $(a, b) \in I^2$,

$$f(b) = \sum_{k=0}^n \frac{f^{(k)}(a)(b-a)^k}{k!} + R_n \quad \text{où} \quad R_n = \int_a^b \frac{(b-t)^n}{n!} f^{(n+1)}(t) dt.$$

Démonstration : Il suffit, via une base de F , de revenir à la démonstration classique dans le cas des fonctions scalaires.

Remarques :

1. On peut aussi énoncer le théorème avec $f \in \mathcal{C}^{n+1}([a, b], F)$.
2. Dans le cas $n = 0$, on retrouve le théorème fondamental de l'analyse : $f(b) = f(a) + \int_a^b f'(t)dt$.
3. En posant $b = a + h \iff h = b - a$ on obtient

$$f(a+h) = \sum_{k=0}^n \frac{f^{(k)}(a)h^k}{k!} + R_n \quad \text{où} \quad R_n = \int_a^b \frac{(a+h-t)^n}{n!} f^{(n+1)}(t)dt = \int_a^b \frac{(h-u)^n}{n!} f^{(n+1)}(a+u)du$$

en posant $t = a + u$ dans l'intégrale.

Proposition 15.23 (Inégalité de Taylor-Lagrange)

Soit f une fonction de classe $\mathcal{C}^{n+1}$ de I vers F . Pour tout $(a, b) \in I^2$,

$$\left\| f(b) - \sum_{k=0}^n \frac{f^{(k)}(a)(b-a)^k}{k!} \right\| = \|R_n\| \leq \frac{|b-a|^{n+1}}{(n+1)!} \sup_{t \in [a,b]} \|f^{(n+1)}\|.$$

Remarque : La fonction $f^{(n+1)}$ est continue. Elle est donc bien bornée sur le segment $[a, b]$.

Démonstration : On utilise l'inégalité triangulaire puis l'inégalité de la moyenne

$$\left\| \int_a^b \frac{(b-t)^n}{n!} f^{(n+1)}(t)dt \right\| \leq \int_a^b \left\| \frac{(b-t)^n}{n!} f^{(n+1)}(t) \right\| dt \leq \frac{|b-a|^{n+1}}{(n+1)!} \sup_{t \in [a,b]} \|f^{(n+1)}\|.$$

□

4.2 Formule de Taylor-Young

Théorème 15.24 (Formule de Taylor-Young)

Soit $n \in \mathbb{N}$. Soit f une fonction de classe $\mathcal{C}^n$ de I dans F . Pour tout $a \in I$,

$$f(x) = f(a) + (x-a)f'(a) + \cdots + (x-a)^n \frac{f^{(n)}(a)}{n!} + o_{x \rightarrow a}((x-a)^n)$$

en posant $h = x - a$,

$$f(a+h) = \sum_{k=0}^n \frac{f^{(k)}(a)}{k!} h^k + o_{h \rightarrow 0}(h^n)$$

Démonstration : On revient encore au cas scalaire. En effet soit $\mathcal{B} = (e_1, \dots, e_p)$ une base de F . Pour tout $i \in \llbracket 1; p \rrbracket$, on pose $f_i = e_i^* \circ f$ de fait que

$$f : x \mapsto \sum_{i=1}^p f_i(x) e_i$$

On peut appliquer la formule de Taylor-Young à chaque f_i qui est de classe $\mathcal{C}^n$.

$$\forall x \in I, f_i(x) = f_i(a) + (x-a)f'_i(a) + \cdots + (x-a)^n \frac{f_i^{(n)}(a)}{n!} + (x-a)^n \varepsilon_i(x)$$

où $\varepsilon_i(x) \xrightarrow{x \rightarrow a} 0$.

On en déduit que

$$\forall x \in I, f(x) = f(a) + (x-a)f'(a) + \cdots + (x-a)^n \frac{f^{(n)}(a)}{n!} + (x-a)^n \sum_{i=1}^p \varepsilon_i(x) e_i.$$

Maintenant, en posant $\varepsilon : x \mapsto \sum_{i=1}^p \varepsilon_i(x) e_i$ qui tend vers 0 quand x tend vers a , on a bien,

$$\forall x \in I, f(x) = f(a) + (x-a)f'(a) + \cdots + (x-a)^n \frac{f^{(n)}(a)}{n!} + (x-a)^n \varepsilon(x).$$

□

ATTENTION

La formule de Taylor avec reste intégral et l'inégalité de Taylor-Lagrange sont des formules globales. Elles donnent des informations sur les valeurs de f dans tout l'intervalle. La formule de Taylor-Young en a est par contre locale. On n'en tire que des informations au voisinage du point a .

5 Suites et séries de fonctions à valeurs vectorielles

Précédemment, on a étudié les suites et séries de fonctions définies sur une partie X d'un espace vectoriel de dimension finie (souvent $\mathbf{R}$ ou $\mathbf{C}$) et à valeurs dans un corps $\mathbf{K}$. Nous allons étendre les résultats de ce chapitre au cas des fonctions à valeurs dans F .

5.1 Généralités

Dans ce paragraphe, les fonctions sont définies sur un ensemble X .

Définition 15.25 (Convergence des suites de fonctions)

Soit (f_n) une suite de fonction définie sur un ensemble X et à valeurs dans F . Soit $f : X \rightarrow F$.

1. On dit que la suite (f_n) converge simplement vers f et on note $(f_n) \xrightarrow{CS} f$ si et seulement si, pour tout x de X la suite $(f_n(x))$ tend vers $f(x)$.
2. On dit que la suite (f_n) converge uniformément vers f et on note $(f_n) \xrightarrow{CU} f$ si et seulement si,

$$\forall \varepsilon > 0, \exists N \in \mathbf{N}, \forall n \geq N \forall x \in X, \|f_n(x) - f(x)\| \leq \varepsilon$$

Cela revient à dire qu'à partir d'un certain rang, les fonctions $f_n - f$ sont bornées et $\|f_n - f\|_\infty \rightarrow 0$.

Remarque : Ces deux notions ne dépendent pas de la norme sur F .

Définition 15.26 (Convergence des séries de fonctions)

Soit (f_k) une suite de fonction définie sur un ensemble X et à valeurs dans F . On considère la série de fonctions $\left(\sum_{k \geq 0} f_k\right)$

1. On dit que la série $\left(\sum_{k \geq 0} f_k\right)$ converge simplement vers F si la suite des sommes partielles $S_n : x \mapsto \sum_{k=0}^n f_k(x)$ converge simplement vers F .
2. On dit que la série $\left(\sum_{k \geq 0} f_k\right)$ converge uniformément vers F si la suite des sommes partielles $S_n : x \mapsto \sum_{k=0}^n f_k(x)$ converge uniformément vers F .
3. On dit que la série $\left(\sum_{k \geq 0} f_k\right)$ converge normalement si les f_k sont bornées et que la série numérique à termes positifs $\left(\sum_{k \geq 0} \|f_k\|_\infty\right)$ converge.

Remarque : On peut souvent travailler avec la convergence normale en enlevant un nombre fini de termes qui ne seraient pas bornées « au début ». Par exemple si on pose $f_k : x \mapsto x^{k-1}$ sur $]0, \frac{1}{2}]$. On a alors pour tout n ,

$$\forall x \in]0, \frac{1}{2}], S_n(x) = \frac{1}{x} + \sum_{k=1}^n x^{k-1}$$

et la série $\left(\sum_{k \geq 1} x^{k-1}\right)$ converge normalement car

$$\|x \mapsto x^{k-1}\|_\infty \leq \frac{1}{2^{k-1}}$$

Proposition 15.27

Une série de fonctions converge uniformément si et seulement si elle converge simplement et que la suite des restes converge uniformément vers 0.

Exemple : Soit $n \in \mathbb{N}$, on pose

$$\begin{aligned} f_n : \mathcal{M}_n(\mathbb{K}) &\rightarrow \mathcal{M}_n(\mathbb{K}) \\ A &\mapsto \frac{A^n}{n!} \end{aligned}$$

On considère alors la série de fonctions $\left(\sum_{n \geq 0} f_n\right)$. On considère une norme $\|\cdot\|$ sur $\mathcal{M}_n(\mathbb{K})$ et on pose C tel que

$$\forall (A, B) \in \mathcal{M}_n(\mathbb{K})^2, \|AB\| \leq C\|A\|\|B\|.$$

On en déduit que pour tout $A \in \mathcal{M}_n(\mathbb{K})$,

$$\|f_n(A)\| = \frac{1}{n!} \|A^n\| \leq \frac{C^{n-1} \|A\|^n}{n!}.$$

De ce fait, si X est une partie bornée de $\mathcal{M}_n(\mathbf{K})$ et si $K \in \mathbf{R}$ vérifie que $X \subset B(0_E, K)$, on a que

$$\|A \mapsto f_n(A)\|_{\infty, X} \leq \frac{C^{n-1}K^n}{n!} = \alpha_n.$$

Maintenant la série $(\sum_{n \geq 0} \alpha_n)$ converge donc la série de fonctions converge normalement sur X .

Notons cependant qu'elle ne converge pas normalement sur $\mathcal{M}_n(\mathbf{K})$ en entier.

Théorème 15.28

Soit $(\sum_{n \geq 0} f_n)$ une série de fonctions. Si elle converge normalement alors elle converge uniformément.

Démonstration : Pour montrer que la série converge uniformément, on va montrer qu'elle converge simplement et que la série des restes converge uniformément vers 0.

- Soit $x \in X$, on a que pour tout entier n , $\|f_n(x)\| \leq \|f_n\|_{\infty}$. Par comparaison de série à termes positifs, on en déduit que la série $(\sum_{n \geq 0} \|f_n(x)\|)$ converge car la série $(\sum_{n \geq 0} \|f_n\|_{\infty})$ converge. De ce fait, la série $(\sum_{n \geq 0} f_n(x))$ est absolument convergente donc convergente. La série de fonctions converge simplement.
- On considère alors le reste de la série. Pour tout $x \in X$:

$$\left\| \sum_{k=n+1}^{+\infty} f_k(x) \right\| \leq \sum_{k=n+1}^{+\infty} \|f_k\|_{\infty}.$$

On en déduit que en notant $x \mapsto R_n(x)$ le reste de la série de fonctions

$$\|x \mapsto R_n(x)\|_{\infty} \leq \sum_{k=n+1}^{+\infty} \|f_k\|_{\infty}.$$

De ce fait, la série des fonctions des restes tend uniformément vers 0.

En conclusion la série de fonctions $(\sum_{n \geq 0} f_n)$ converge uniformément. □

Remarque : C'est la même preuve que celle du chapitre 5

Exemple : La série exponentielle converge uniformément sur toute partie bornée de $\mathcal{M}_n(\mathbf{K})$.

5.2 Continuité et théorème de la double limite

On suppose maintenant que X est une partie d'un espace vectoriel normé E . On note $\|\cdot\|_E$ la norme de E .

Les théorèmes sur la continuité des suites et des séries de fonctions se généralisent directement.

Théorème 15.29 (Continuité de la limite d'une suite de fonctions)

Soit (f_n) une suite de fonctions définies sur X et à valeurs dans F qui converge (simplement) vers une fonction f .

1. Soit $x_0 \in X$. On suppose que

- i) Les fonctions f_n sont continues en x_0 .
- ii) Il existe un voisinage de x_0 dans X tel que la suite de fonctions converge uniformément.

Alors la fonction f est continue en x_0 .

2. On suppose que

- i) Les fonctions f_n sont continues sur X
- ii) Pour tout x_0 de X , il existe un voisinage de x_0 dans X tel que la suite de fonctions converge uniformément.

Alors f est continue sur X .

Théorème 15.30 (Continuité de la somme d'une série de fonctions)

Soit $\left(\sum_{n \geq 0} f_n\right)$ une suite de fonctions définies sur X et à valeurs dans F qui converge (simplement) vers une fonction S .

1. Soit $x_0 \in X$. On suppose que

- i) Les fonctions f_n sont continues en x_0 .
- ii) Il existe un voisinage de x_0 dans X tel que la série de fonctions converge uniformément.

Alors la fonction S est continue en x_0 .

2. On suppose que

- i) Les fonctions f_n sont continues sur X
- ii) Pour tout x_0 de X , il existe un voisinage de x_0 dans X tel que la série de fonctions converge uniformément.

Alors S est continue sur X .

Exemple : On reprend la série exponentielle de matrice. Pour tout k , $f_k : A \mapsto \frac{A^k}{k!}$ est continue sur $\mathcal{M}_n(\mathbf{K})$. De plus, pour tout $A_0 \in \mathcal{M}_n(\mathbf{K})$. La boule ouverte $B(0_E, \|A_0\| + 1)$ est un voisinage de A_0 qui est borné. De ce fait, la série de fonctions converge normalement donc uniformément sur cette boule. On en déduit que $A \mapsto \exp(A)$ est une fonction continue sur $\mathcal{M}_n(\mathbf{K})$.

Exercices :

1. Montrer que pour toute matrice M ,

$$\det(\exp(M)) = \exp(\operatorname{tr}(M)).$$

2. Montrer $A \mapsto (I - A)^{-1}$ est continue sur un voisinage ouvert de 0.

On peut aussi généraliser le théorème de la double limite.

Théorème 15.31 (Théorème de la double limite)

Soit (f_n) une suite de fonctions de X dans F et $f \in \mathcal{F}(X, F)$. Soit a un point adhérent à X . On suppose que

- i) La suite (f_n) converge **uniformément** vers f sur un voisinage de a dans X .
- ii) pour tout n , f_n admet en a une limite $\ell_n \in F$.

Alors la suite (ℓ_n) admet une limite ℓ et $\lim_{x \rightarrow a} f(x) = \ell$. C'est-à-dire

$$\lim_{n \rightarrow +\infty} \left(\lim_{x \rightarrow a} f_n(x) \right) = \ell = \lim_{x \rightarrow a} \left(\lim_{n \rightarrow +\infty} f_n(x) \right).$$

Remarque : La démonstration est la même que dans le cas étudié au chapitre 5. Il est important que l'espace d'arrivée soit de dimension finie. En effet, pour montrer que la suite (ℓ_n) a une limite finie, on utilise que c'est une suite bornée et on en extrait alors une sous-suite convergente d'après le théorème de Bolzano-Weierstrass. On peut généraliser cela en utilisant qu'une boule fermée bornée en dimension finie est compacte.

Théorème 15.32 (Double limite pour les séries de fonctions)

Soit $\sum f_n$ une série de fonctions définie sur X et a un point adhérent à A . On suppose que

- i) La série de fonctions converge uniformément vers S au voisinage de a
- ii) pour tout n , f_n admet en a une limite $\ell_n \in F$.

Alors la série $\sum \ell_n$ converge (vers ℓ) et $\lim_{x \rightarrow a} S(x) = \ell$. C'est-à-dire :

$$\lim_{x \rightarrow a} \sum_{n=0}^{+\infty} f_n(x) = \sum_{n=0}^{+\infty} \left(\lim_{x \rightarrow a} f_n(x) \right).$$

5.3 Intégration et dérivation

Nous allons maintenant généraliser l'intégration et la dérivation des suites et séries de fonctions dans le cas des fonctions à valeurs vectorielles.

Théorème 15.33

Soit (f_n) une suite de fonctions définies sur un intervalle I et à valeurs dans F . Soit f une fonction définie sur I et à valeurs dans F . On suppose que

- i) Les fonctions f_n sont continues.
- ii) La suite de fonction converge uniformément sur tout segment J inclus dans I .

Soit $a \in I$, on pose $H_n : x \mapsto \int_a^x f_n(t) dt$ et $H : x \mapsto \int_a^x f(t) dt$. La suite (H_n) converge uniformément vers H sur tout segment J de I .

Démonstration : Il suffit de le démontrer coordonnées par coordonnées sur une base de F . $\square$

Il existe une « version série ».

Théorème 15.34 (Intégration des séries de fonctions)

Soit $\sum f_n$ une série de fonctions définies sur un intervalle I et à valeurs dans F . On suppose que

- i) Pour tout entier n , f_n est continue
- ii) La série de fonction converge uniformément sur tout segment J vers S

Pour tout a de I on pose $H_n : x \mapsto \int_a^x f_n$ la primitive de f_n qui s'annule en a . Alors la série $\sum H_n$ converge uniformément sur tout segment J de I vers $x \mapsto \int_a^x S$. C'est-à-dire,

$$\forall x \in I, \sum_{n=0}^{+\infty} \int_a^x f_n = \int_a^x \sum_{n=0}^{+\infty} f_n$$

Remarque : En particulier si $I = [a, b]$ on obtient $\lim_{n \rightarrow \infty} \left(\int_a^b f_n \right) = \int_a^b f$ dans le premier cas et

$$\sum_{n=0}^{\infty} \int_a^b f_n = \int_a^b f \text{ dans le deuxième cas.}$$

Les théorèmes sur la dérivation se généralisent aussi au cas des fonctions vectorielles.

Théorème 15.35

Soit (f_n) une suite de fonctions définies sur un intervalle I et à valeurs dans F . On suppose que

- i) Pour tout entier n , la fonction f_n est de classe $\mathcal{C}^1$ sur I .
- ii) La suite de fonctions (f_n) converge (simplement) vers une fonction f .
- iii) La suite des fonctions dérivées (f'_n) converge **uniformément** sur tout segment J de I vers une fonction g .

Alors la suite de fonction (f_n) converge uniformément sur tout segment J de I , f est de classe $\mathcal{C}^1$ sur I et $f' = g$. Dit autrement on a

$$\lim_{n \rightarrow \infty} f'_n = f' = \left(\lim_{n \rightarrow \infty} f_n \right)'.$$

Là encore, il existe un « version série ».

Théorème 15.36 (Dérivation termes à termes)

Soit $\sum f_n$ une série de fonctions définies sur un intervalle I et à valeurs dans F . On suppose que

- i) Pour tout entier n , f_n est de classe $\mathcal{C}^1$ sur I ,
- ii) la série de fonctions converge (simplement) vers S sur I ,
- iii) la série de fonctions des dérivées $\sum f'_n$ converge uniformément sur tout segment J vers T

Alors la série $\sum f_n$ converge uniformément sur tout segment J , la fonction S est de classe $\mathcal{C}^1$ et $S' = T$. C'est-à-dire

$$\left(\sum_{n=0}^{+\infty} f_n \right)' = \sum_{n=0}^{+\infty} f'_n.$$

Exemple : Soit A une matrice de $\mathcal{M}_n(\mathbf{K})$, on pose pour tout entier k ,

$$f_k : t \mapsto \frac{(tA)^k}{k!}$$

A t fixé, on sait que la série converge. La série de fonctions $t \mapsto \sum_{k \geq 0} \frac{(tA)^k}{k!}$ converge donc vers

$$S : t \mapsto \exp(tA).$$

Maintenant, à k fixé, f_k est de classe $\mathcal{C}^1$ et $f'_k : t \mapsto \frac{t^{k-1}A^k}{(k-1)!}$ si $k \geq 1$ et $f'_0 : t \mapsto 0$.

Maintenant, pour tout segment J de $\mathbf{R}$, il existe $M \in \mathbf{R}$ tel que $J \subset [-M, M]$. On en déduit que

$$\|t \mapsto f'_k(t)\|_{\infty, J} \leq \frac{M^{k-1}}{(k-1)!} \|A^k\| \leq \frac{(CM)^{k-1} \|A\|^k}{(k-1)!}$$

en notant C tel que

$$\forall (A, B) \in \mathcal{M}_n(\mathbf{K})^2, \|AB\| \leq C \|A\| \|B\|.$$

On en déduit donc que la série de fonctions $\left(\sum_{k \geq 0} f'_k \right)$ converge normalement sur J donc uniformément.

De ce fait, la fonction S est de classe $\mathcal{C}^1$ et

$$\forall t \in \mathbf{R}, S'(t) = \sum_{k=1}^{+\infty} \frac{t^{k-1}A^k}{(k-1)!} = A \exp(tA).$$

En utilisant la continuité du produit matriciel pour dire que

$$\sum_{k=0}^{+\infty} A \frac{t^k A^k}{k!} = A \sum_{k=0}^{+\infty} \frac{t^k A^k}{k!}.$$

On a donc montré que $S'(t) = AS(t)$.

Remarque : On peut encore étendre ces théorèmes au cas des fonctions de classe $\mathcal{C}^p$ et de classe $\mathcal{C}^\infty$.

Équations différentielles linéaires

1	Généralités	426
1.1	Définitions	427
1.2	Structure des solutions	428
1.3	Problème de Cauchy	429
1.4	Équation différentielle scalaire d'ordre n	430
2	Théorème de Cauchy linéaire	432
2.1	Théorème de Cauchy linéaire	432
2.2	Applications du théorème de Cauchy linéaire	434
2.3	Exemples d'équations différentielles non normalisées	435
3	Équations différentielles à coefficients constants	437
3.1	Exponentielle des matrices et endomorphismes	437
3.2	Généralités	439
3.3	Exemples de calculs	441
4	Equations différentielles scalaires du second ordre	443
4.1	Wronskien d'un couple de solution	444
4.2	Variation de la constante	446
4.3	Résolutions d'une équation scalaire du second ordre en connaissant une solution	448

Dans ce chapitre I désigne un intervalle de $\mathbf{R}$ et E un espace vectoriel de dimension finie.

1 Généralités

1.1 Définitions

Définition 16.1

On appelle *équation différentielle linéaire d'ordre 1 (sous forme normale)* une équation

$$x' = a(t)(x) + b(t)$$

où a est une application continue sur I dans $\mathcal{L}(E)$ et b est une application continue de I dans E .

Remarques :

1. Dans le cas où $E = \mathbf{K}$, on retrouve la notion d'équation différentielle scalaire vue en première année. En effet, les applications linéaires de $\mathbf{K}$ dans $\mathbf{K}$ sont les multiplication par un scalaire. Il existe donc des applications continues $\alpha : I \rightarrow \mathbf{K}^*$ et $\beta : I \rightarrow \mathbf{K}$ telles que l'équation différentielle s'écrive

$$x'(t) = \alpha(t)x(t) + \beta(t)$$

2. On appellera équation différentielle linéaire *vectorielle* une équation donc les solutions sont des fonctions à valeurs dans l'espace vectoriel E pour les différentier des équations différentielles linéaires *scalaires* vues en première année.

Définition 16.2 (Solution d'une équation différentielle)

Soit

$$x'(t) = a(t)(x(t)) + b(t) \quad (E)$$

une équation différentielle linéaire d'ordre 1. Une solution de (E) est une fonction $x : t \mapsto x(t)$ définie de classe $\mathcal{C}^1$ sur I et à valeurs dans E telle que

$$\forall t \in I, x'(t) = a(t)(x(t)) + b(t).$$

On notera $\mathcal{S}_E$ l'ensemble des solutions.

ATTENTION

Pour tout $t \in I$, $a(t)$ est une application linéaire de E dans E . Le terme $a(t)(x(t))$ est bien l'évaluation de l'application linéaire $a(t)$ en le vecteur $x(t)$.

Interprétation matricielle : Soit $\mathcal{B} = (e_1, \dots, e_n)$ une base de E . On note

$$X(t) = \text{Mat}_{\mathcal{B}}(x(t)) \in \mathcal{M}_{n,1}(\mathbf{K}), A(t) = \text{Mat}_{\mathcal{B}}(a(t)) \in \mathcal{M}_n(\mathbf{K}) \text{ et } B(t) = \text{Mat}_{\mathcal{B}}(b(t)) \in \mathcal{M}_{n,1}(\mathbf{K}).$$

L'équation différentielle s'écrit alors

$$X'(t) = A(t)X(t) + B(t).$$

Si on note

$$X(t) = \begin{pmatrix} x_1(t) \\ \vdots \\ x_n(t) \end{pmatrix}, A(t) = (a_{ij}(t)) \text{ et } B(t) = \begin{pmatrix} b_1(t) \\ \vdots \\ b_n(t) \end{pmatrix}.$$

En projetant coordonnées par coordonnées on obtient

$$X'(t) = A(t)X(t) + B(t) \iff \forall i \in \llbracket 1; n \rrbracket, x'_i(t) = \sum_{j=1}^n a_{ij}(t)x_j(t) + b_i(t).$$

En particulier, si A est une matrice diagonale, on est juste ramené à n équations différentielles scalaires :

$$\forall i \in \llbracket 1; n \rrbracket, x'_i(t) = a_{ii}(t)x_i(t) + b_i(t).$$

Terminologie :

1. On appelle système différentiel d'ordre 1 une équation

$$X'(t) = A(t)X(t) + B(t)$$

où A est une application continue sur I dans $\mathcal{M}_n(\mathbf{K})$ et B est une application continue de I dans $\mathcal{M}_{n,1}(\mathbf{K})$.

2. L'équation différentielle (E) est dite à coefficients constants si a est une fonction constante. Ce qui revient à ce que la matrice A (et donc ses coefficients) ne dépende pas de t .
3. Le terme b dans l'équation différentielle (E) s'appelle le second membre. Si b est la fonction constante égale à 0, l'équation est dite homogène (ou sans second membre).
4. On associe à l'équation (E) son équation homogène associée

$$x'(t) = a(t)(x(t)) \tag{H}$$

1.2 Structure des solutions

Théorème 16.3

1. Si (H) est une équation linéaire homogène, l'ensemble de ses solutions est un sous-espace vectoriel de l'ensemble des fonctions de classe $\mathcal{C}^1$ de I dans E .
2. Si (E) est une équation différentielle et (H) l'équation homogène associée, soit $\mathcal{S}_E$ est vide, soit $\mathcal{S}_E = \{f_0 + h \mid h \in \mathcal{S}_H\}$ où f_0 est une solution de (E). C'est donc un sous-espace affine de direction $\mathcal{S}_H$.

Démonstration :

1. Vérifions les axiomes

- La fonction nulle $\tilde{0} : I \rightarrow E$ vérifie l'équation (H) car c'est une équation homogène. On en déduit que $\mathcal{S}_E \neq \emptyset$.
- Soit x_1, x_2 deux solutions de (E) et soit λ_1, λ_2 deux scalaires. Pour tout $t \in I$,

$$\begin{aligned} (\lambda_1 x_1 + \lambda_2 x_2)'(t) &= \lambda_1 x'_1(t) + \lambda_2 x'_2(t) \\ &= \lambda_1 a(t)(x_1(t)) + \lambda_2 a(t)(x_2(t)) \\ &= a(t)(\lambda_1 x_1(t) + \lambda_2 x_2(t)) \text{ car } a(t) \text{ est linéaire} \end{aligned}$$

On en déduit que $\lambda_1 x_1 + \lambda_2 x_2 \in \mathcal{S}_H$

On a bien montré que $\mathcal{S}_H$ était un sous-espace vectoriel de $\mathcal{C}^1(I, E)$.

2. On considère $\Phi : \mathcal{C}^1(I, E) \rightarrow \mathcal{C}^0(I, E)$ l'application définie par

$$\Phi : x \rightarrow (t \mapsto x'(t) - a(t)(x(t)))$$

On montre aisément comme ci-dessus que la fonction Φ est linéaire. On a alors pour $x \in \mathcal{C}^1(I, E)$, x vérifie l'équation (E) si et seulement si $\Phi(x) = b$. On en déduit que

- Si b n'est pas dans l'image de Φ , $\mathcal{S}_E = \emptyset$.
- Si b est dans l'image de Φ et si f_0 vérifie $\Phi(f_0) = b$ alors

$$\mathcal{S}_E = \Phi^{-1}(b) = f_0 + \text{Ker}(\Phi) = \{f_0 + h, h \in \mathcal{S}_H\}$$

□

Proposition 16.4 (Principe de superposition)

Soit a une fonction continue de I dans $\mathcal{L}(E)$ et b_1, b_2 deux fonctions continues de I dans E . Si f_1 est solution de l'équation $x' = ax + b_1$ et f_2 est une solution de $x' = ax + b_2$ alors $f_1 + f_2$ est une solution de $x' = ax + (b_1 + b_2)$.

Démonstration : En exercice.

□

Exemple : Résoudre

$$x'(t) + x(t) = 1 + e^{-t}.$$

1.3 Problème de Cauchy

Définition 16.5 (Problème de Cauchy linéaire d'ordre 1)

1. Un problème de Cauchy linéaire d'ordre 1 est la donnée de :

- Une équation différentielle linéaire (vectorielle) normalisée d'ordre 1 : $x'(t) = a(t)(x(t)) + b(t)$ où $a : I \rightarrow \mathcal{L}(E)$ et $b : I \rightarrow E$ sont **continues**.
- Une condition initiale : un couple $(t_0, x_0) \in I \times E$.

On le note

$$(PC) \begin{cases} x' = ax + b \\ x(t_0) = x_0 \end{cases}$$

2. Une solution d'un problème de Cauchy est une solution de l'équation différentielle qui vérifie la condition initiale.

Proposition 16.6 (Forme intégrale d'un problème de Cauchy)

Soit

$$(PC) \begin{cases} x' = ax + b \\ x(t_0) = x_0 \end{cases}$$

un problème de Cauchy et $x : I \rightarrow E$ une fonction de classe $\mathcal{C}^1$. La fonction x est une solution du produit de Cauchy si et seulement si

$$\forall t \in I, x(t) = x_0 + \int_{t_0}^t (a(u)x(u) + b(u)) du.$$

Démonstration :

- $\Rightarrow$ On suppose que x est une solution de classe $\mathcal{C}^1$ du problème de Cauchy, on a alors pour tout $t \in I$.

$$\int_{t_0}^t x'(u) du = x(t) - x(t_0) = x(t) - x_0.$$

On a bien

$$\forall t \in I, x(t) = x_0 + \int_{t_0}^t a(u)x(u) + b(u) du.$$

- $\Leftarrow$ Soit x une fonction qui vérifie,

$$\forall t \in I, x(t) = x_0 + \int_{t_0}^t a(u)x(u) + b(u) du.$$

La fonction x est de classe $\mathcal{C}^1$ car $u \mapsto a(u)x(u) + b(u)$ est continue. De plus on vérifie que $x(t_0) = x_0$ puis, en dérivant, que

$$\forall t \in I, x'(t) = a(t)x(t) + b(t).$$

□

1.4 Équation différentielle scalaire d'ordre n

Définition 16.7 (Équation différentielle scalaire d'ordre n)

Soit $n \geq 1$. On appelle *équation différentielle linéaire scalaire d'ordre n normalisée sur I* une équation différentielle de la forme

$$x^{(n)}(t) = a_{n-1}(t)x^{(n-1)}(t) + \cdots + a_1(t)x'(t) + a_0(t)x(t) + b(t)$$

où les fonctions a_i ($i \in \llbracket 0 ; n \rrbracket$) et b sont des fonctions continues de I dans $\mathbb{K}$. Ses solutions, sont les fonctions x de classe $\mathcal{C}^n(I, \mathbb{K})$ telles que

$$\forall t \in I, x^{(n)}(t) = a_{n-1}(t)x^{(n-1)}(t) + \cdots + a_1(t)x'(t) + a_0(t)x(t) + b(t)$$

Théorème 16.8

Soit

$$x^{(n)}(t) = a_{n-1}(t)x^{(n-1)}(t) + \cdots + a_1(t)x'(t) + a_0(t)x(t) + b(t) \quad (E)$$

une équation différentielle linéaire d'ordre n scalaire normalisée et $x \in \mathcal{C}^n(I, \mathbf{K})$. On pose

$$X : t \mapsto \begin{pmatrix} x(t) \\ x'(t) \\ \vdots \\ x^{(n-1)}(t) \end{pmatrix} \in \mathcal{M}_{n,1}(\mathbf{K}),$$

$$A : t \mapsto \begin{pmatrix} 0 & 1 & 0 & \cdots & 0 \\ \vdots & \ddots & \ddots & & \vdots \\ \vdots & & \ddots & \ddots & \vdots \\ 0 & \cdots & 0 & 0 & 1 \\ a_0(t) & \cdots & \cdots & a_{n-2}(t) & a_{n-1}(t) \end{pmatrix} \in \mathcal{M}_n(\mathbf{K})$$

et

$$B : t \mapsto \begin{pmatrix} 0 \\ 0 \\ \vdots \\ b(t) \end{pmatrix} \in \mathcal{M}_{n,1}(\mathbf{K}).$$

La fonction x vérifie l'équation différentielle (E) si et seulement si X est solution du système différentiel

$$X' = AX + B \quad (S)$$

Démonstration : Soit $x \in \mathcal{C}^n(I, \mathbf{K})$ et $X : I \rightarrow \mathcal{M}_{n,1}(\mathbf{K})$ associée.

$$\begin{aligned} X \text{ vérifie (X)} &\iff \forall t \in I, X'(t) = A(t)X(t) + B \\ &\iff \begin{pmatrix} x'(t) \\ x''(t) \\ \vdots \\ x^{(n)}(t) \end{pmatrix} = \begin{pmatrix} x'(t) \\ x''(t) \\ \vdots \\ a_0(t)x(t) + \cdots + a_{n-1}(t)x^{(n-1)}(t) \end{pmatrix} + \begin{pmatrix} 0 \\ 0 \\ \vdots \\ b(t) \end{pmatrix} \\ &\iff x^{(n)}(t) = a_0(t)x(t) + \cdots + a_{n-1}(t)x^{(n-1)}(t) + b(t) \end{aligned}$$

□

Exemple : Pour résoudre l'équation

$$x''(t) = (\cos t)x'(t) - (\sin t)x(t) + e^t$$

on se ramène à

$$X' = AX + B$$

où

$$X(t) = \begin{pmatrix} x(t) \\ x'(t) \end{pmatrix}, B(t) = \begin{pmatrix} 0 \\ e^t \end{pmatrix} \text{ et } A(t) = \begin{pmatrix} 0 & 1 \\ \cos t & -\sin t \end{pmatrix}.$$

Remarque : On voit donc que résoudre une équation linéaire scalaire d'ordre n revient à résoudre une équation linéaire d'ordre 1 à valeurs dans $\mathbf{K}^n$.

Définition 16.9 (Problème de Cauchy linéaire scalaire d'ordre n)

1. Un problème de Cauchy linéaire scalaire d'ordre n est la donnée de :

— Une équation différentielle linéaire normalisée d'ordre n :

$$x^{(n)}(t) = a_{n-1}(t)x^{(n-1)}(t) + \cdots + a_1(t)x'(t) + a_0(t)x(t) + b(t)$$

— Une condition initiale : un $(n+1)$ -uplet $(t_0, x_0, x'_0, \dots, x_0^{(n-1)}) \in I \times \mathbb{K}^n$.

On le note

$$(PC) \begin{cases} x^{(n)}(t) = a_{n-1}(t)x^{(n-1)}(t) + \cdots + a_1(t)x'(t) + a_0(t)x(t) + b(t) \\ x(t_0) = x_0, x'(t_0) = x'_0, \dots, x^{(n-1)}(t_0) = x_0^{(n-1)} \end{cases}$$

2. Une solution d'un problème de Cauchy est une solution de l'équation différentielle qui vérifie la condition initiale.

2 Théorème de Cauchy linéaire

2.1 Théorème de Cauchy linéaire

Théorème 16.10

Soit

$$(PC) \begin{cases} x'(t) = a(t)x(t) + b(t) \\ x(t_0) = x_0 \end{cases}$$

un problème de Cauchy linéaire d'ordre 1 vectoriel en particulier les fonctions a et b sont **continues**. Il admet une unique solution.

Démonstration : (Non exigible) : Nous allons faire la démonstration dans le cas où I est un segment (donc compact). Le cas général s'en déduit (exercice).

Rappelons pour commencer la forme intégrale du problème de Cauchy :

$$\forall t \in I, x(t) = x_0 + \int_{t_0}^t a(u)x(u) + b(u)du.$$

— Existence : On procède par approximations. Précisément on pose $x_0 : t \mapsto x_0$ et pour tout entier n ,

$$x_{n+1} : t \mapsto x_0 + \int_{t_0}^t a(u)x_n(u) + b(u)du.$$

Le but est alors de démontrer que la suite de fonctions (x_n) converge uniformément vers une fonction x qui sera la solution cherchée. Pour cela, évaluons $x_{n+1}(t) - x_n(t)$.

$$\begin{aligned} \forall n \in \mathbb{N}^*, \forall t \in I, x_{n+1}(t) - x_n(t) &= x_0 + \int_{t_0}^t a(u)x_n(u) + b(u)du - x_0 - \int_{t_0}^t a(u)x_{n-1}(u) + b(u)du \\ &= \int_{t_0}^t a(u)(x_n(u) - x_{n-1}(u))du. \end{aligned}$$

La fonction $a : I \rightarrow \mathcal{L}(E)$ étant continue, elle est bornée sur I qui est compact. On se donne $\|\cdot\|$ une norme de E et on note $\|\cdot\|_{\text{op}}$ la norme subordonnée associée

$$\|f\|_{\text{op}} = \sup_{x \in E, \|x\|=1} \|f(x)\|.$$

C'est une norme sur $\mathcal{L}(E)$. Soit $M \in \mathbf{R}$, tel que

$$\forall t \in I, \|a(t)\|_{\text{op}} \leq M.$$

On a donc

$$\forall t \in I, \|a(t)(x_n - x_{n-1}(t))\| \leq M \|x_n(t) - x_{n-1}(t)\|$$

Par intégration on obtient alors

$$\forall t \geq t_0, \|x_2(t) - x_1(t)\| \leq (t - t_0)M \|x_1 - x_0\|_{\infty}$$

puis

$$\forall t \geq t_0, \|x_3(t) - x_2(t)\| \leq \int_{t_0}^t (u - t_0)M^2 \|x_1 - x_0\|_{\infty} \leq \frac{M^2(t - t_0)^2}{2} \|x_1 - x_0\|_{\infty}$$

Par récurrence, on peut alors montrer que

$$\forall t \geq t_0, \|x_{n+1} - x_n(t)\| \leq \frac{M^n |t - t_0|^n}{n!} \|x_1 - x_0\|_{\infty}.$$

On notera que $\|x_1 - x_0\|_{\infty}$ existe bien car ce sont des fonctions continues sur un compact. On procède de même pour $t \leq t_0$, si bien qu'en notant ℓ la longueur de l'intervalle I on obtient que

$$\|x_{n+1} - x_n\|_{\infty} \leq \frac{M^n \ell^n}{n!} \|x_1 - x_0\|_{\infty}.$$

Cela permet de montrer que la série numérique $\sum_{n \geq 0} \|x_{n+1} - x_n\|_{\infty}$ converge, ce qui signifie que la série de fonctions $\sum_{n \geq 0} (x_{n+1} - x_n)$ converge normalement donc uniformément sur I . Maintenant, pour tout N , par télescopage

$$\sum_{n=0}^N x_{n+1} - x_n = x_{N+1} - x_0$$

ce qui montre que la suite de fonctions (x_n) converge uniformément. Notons x sa limite, il reste à montrer que x est bien la solution cherchée. Par construction, pour tout entier n et tout $t \in I$,

$$x_{n+1}(t) = x_0 + \int_{t_0}^t a(u)x_n(u) + b(u)du$$

on peut alors passer à la limite quand $n \rightarrow \infty$ (en remarquant que $u \mapsto a(u)x_n(u)$ converge uniformément vers $u \mapsto a(u)x(u)$). On a donc bien,

$$\forall t \in I, x(t) = x_0 + \int_{t_0}^t a(u)x(u) + b(u)du.$$

L'existence est démontrée.

- Unicité : On suppose avoir deux fonctions x_1 et x_2 vérifiant notre problème de Cauchy (sous sa forme intégrale). On a alors

$$\forall t \in I, x_1(t) - x_2(t) = \int_{t_0}^t a(u)(x_1(u) - x_2(u))du.$$

On peut répéter l'argument précédent pour montrer que

$$\forall n \in \mathbf{N}, \|x_1 - x_2\|_{\infty} \leq \frac{M^n \ell^n}{n!} \|x_1 - x_2\|_{\infty}$$

En faisant tendre n vers ∞ , on obtient bien que $\|x_1 - x_2\|_{\infty} = 0$ et donc $x_1 = x_2$.

□

2.2 Applications du théorème de Cauchy linéaire

Corollaire 16.11

Soit (H) une équation différentielle vectorielle linéaire d'ordre 1 **homogène**

$$x' = a(t)x(t) \quad (\text{H})$$

où $a : I \rightarrow \mathcal{L}(E)$ est continue. L'ensemble $\mathcal{S}_H$ est un espace vectoriel de dimension égale à la dimension de E . De plus, pour tout $t_0 \in I$,

$$\begin{aligned} \varphi_{t_0} : \mathcal{S}_H &\rightarrow E \\ x &\mapsto x(t_0) \end{aligned}$$

est un isomorphisme.

Remarques :

1. La « nouvelle » information de ce corollaire est la dimension de l'espace vectoriel $\mathcal{S}_H$.
2. Ce résultat avait été établi en première année pour les équations différentielles scalaires linéaires d'ordre 1 et les équations différentielles scalaires linéaires d'ordre 2 à coefficients constants en calculant explicitement $\mathcal{S}_H$.

ATTENTION

Notons que si x est une solution de (H) qui s'annule, alors x est la fonction constante égale à 0.

Démonstration : Il est clair que pour $t_0 \in I$, φ_{t_0} est une application linéaire (évident). Le fait qu'elle soit bijective vient du fait que, d'après le théorème de Cauchy linéaire, pour tout $x_0 \in E$ il existe une unique solution x de (H) telle que $x(t_0) = x_0$. $\square$

Corollaire 16.12 (Solutions des équations linéaires scalaires d'ordre n)

Soit $n \geq 1$, a_i ($i \in \llbracket 0; n \rrbracket$) et b des fonctions **continues** de I dans $\mathbf{K}$. Soit $(t_0, x_0, \dots, x_0^{(n-1)}) \in I \times E^n$. L'équation différentielle linéaire scalaire d'ordre n

$$x^{(n)}(t) = a_{n-1}(t)x^{(n-1)}(t) + \dots + a_1(t)x'(t) + a_0(t)x(t) + b(t)$$

a une unique solution $x : I \rightarrow \mathbf{K}$ telle que

$$\forall k \in \llbracket 0; n-1 \rrbracket, x^{(k)}(t_0) = x_0^{(k-1)}$$

Démonstration : Il suffit d'utiliser l'interprétation d'une équation linéaire scalaire d'ordre n à l'aide d'une équation linéaire vectorielle d'ordre 1 et le théorème de Cauchy linéaire. $\square$

Corollaire 16.13 (Solutions des équations linéaires scalaires d'ordre n)

Soit $n \geq 1$, a_i ($i \in \llbracket 0 ; n \rrbracket$) et b des fonctions **continues** de I dans $\mathbf{K}$. On considère l'équation différentielle linéaire scalaire homogène d'ordre n

$$x^{(n)}(t) = a_{n-1}(t)x^{(n-1)}(t) + \cdots + a_1(t)x'(t) + a_0(t)x(t) \quad (\text{H})$$

Pour tout t_0 ,

$$\begin{aligned} \varphi_{t_0} : \mathcal{S}_H &\rightarrow \mathbf{K}^n \\ x &\mapsto (x(t_0), x'(t_0), \dots, x^{(n-1)}(t_0)) \end{aligned}$$

est un isomorphisme. En particulier, $\mathcal{S}_H$ est un espace vectoriel de dimension n .

ATTENTION

Quand on considère un problème de Cauchy, il est **essentiel** de prendre toutes les dérivées au même point. Par exemple, on sait que les solutions de l'équation

$$x''(t) = -x(t)$$

sont les fonctions de la forme :

$$x : t \mapsto \alpha \cos t + \beta \sin t$$

En particulier il y en a plusieurs qui vérifient $x(0) = x(\pi) = 0$ ou $x(0) = x'(\frac{\pi}{2}) = 0$.

Exemple : Pour l'équation $x'' = x$. L'ensemble des solutions est de dimension 2. Les fonctions $t \mapsto e^t$ et $t \mapsto e^{-t}$ forment une base de l'ensemble des solutions. Les fonctions $t \mapsto \operatorname{ch} t$ et $t \mapsto \operatorname{sh} t$ en forment une autre.

Remarque : Si on a une solution $x : I \rightarrow \mathbf{K}$ d'une équation scalaire d'ordre n telle qu'il existe $t_0 \in I$ vérifiant

$$x(t_0) = x'(t_0) = \cdots = x^{(n-1)}(t_0) = 0$$

alors x est la fonction nulle.

2.3 Exemples d'équations différentielles non normalisées

On va considérer des équations différentielles linéaires scalaires d'ordre 1 ou 2 **non normalisées**. Ce sont des équations différentielles de la forme

$$\alpha(t)x'(t) = \beta(t)x(t) + \gamma(t) \text{ ou } \alpha(t)x''(t) = \beta(t)x'(t) + \gamma(t)x(t) + \delta(t)$$

où les fonctions α, β, γ et éventuellement δ sont continues de I dans $\mathbf{K}$.

Dans le cas où α ne s'annule pas, on se ramène à des équations normalisées en divisant tous les membres par $\alpha(t)$. Il reste à étudier ce qui se passe quand la fonction α s'annule.

Cas des équations d'ordre 1

On considère l'équation

$$\alpha(t)x'(t) = \beta(t)x(t) + \gamma(t) \quad (\text{E})$$

Si on suppose que α ne s'annule qu'en un point t_0 de I . On peut résoudre l'équation différentielle sur $I_- = I \cap]-\infty, t_0[$ et sur $I_+ = I \cap]t_0, +\infty[$. Sur ces deux intervalles on peut résoudre l'équation en se ramenant au cas d'une équation normalisée. On peut donc déterminer

$$\mathcal{S}_+ = \{x : I_+ \rightarrow \mathbf{K} \mid x \text{ vérifie } E\} = \{x_0^+ + \lambda x_h^+ \mid \lambda \in \mathbf{K}\}$$

et

$$\mathcal{S}_- = \{x : I_- \rightarrow \mathbf{K} \mid x \text{ vérifie } E\} = \{x_0^- + \mu x_h^- \mid \mu \in \mathbf{K}\}.$$

Il ne reste plus qu'à raisonner par analyse synthèse.

- Analyse : Soit x une solution sur I de l'équation, les restrictions des x à I_- et à I_+ appartiennent respectivement à $\mathcal{S}_-$ et à $\mathcal{S}_+$. On détermine ensuite (souvent en regardant x et x' en t_0 s'il faut ajouter des conditions entre λ et μ .)
- Synthèse : on vérifie que les fonctions trouvées vérifient (E) sur I en testant sur I_+ , sur I_- et en t_0 .

On peut avoir de nombreux cas : λ et μ quelconques, pas de solutions, λ déterminé par μ , une unique solution...

Exemples :

1. Résolvons l'équation (E)

$$tx' - 2x = 0$$

sur $\mathbf{R}$. On pose $I_+ =]0, +\infty[$ et $I_- =]-\infty, 0[$. On va résoudre l'équation sur chacun de ces intervalles. On pose (E') l'équation normalisée

$$x' - \frac{2}{t}x = 0$$

Une primitive de $t \mapsto -\frac{2}{t}$ sur I_+ ou I_- est $A : t \mapsto -2 \ln |t|$. On en déduit que les solutions sont

$$\mathcal{S}_\pm = \{t \mapsto \lambda t^2\}$$

Soit x une éventuelle solution de (E) sur $\mathbf{R}$, sa restriction à $\mathbf{R}_+$ et à $\mathbf{R}_-$ vérifie l'équation. On en déduit qu'il existe λ_+ et λ_- tels que

$$\forall t \in \mathbf{R}^*, x(t) = \begin{cases} \lambda_+ t^2 & \text{si } t > 0 \\ \lambda_- t^2 & \text{si } t < 0 \end{cases}$$

Réciproquement toutes ces fonctions sont bien des fonctions de classe $\mathcal{C}^1$ sur $\mathbf{R}$ et elles sont solutions de l'équation sur $\mathbf{R}$.

2. Résolvons maintenant l'équation (E)

$$tx' - x = t^2$$

sur $\mathbf{R}$. Par des calculs similaires aux précédents on en déduit que les solutions sont

$$\mathcal{S}_\pm = \{t \mapsto \lambda t + t^2\}$$

Soit x une éventuelle solution de (E) sur $\mathbf{R}$, sa restriction à $\mathbf{R}_+$ et à $\mathbf{R}_-$ vérifie l'équation. On en déduit qu'il existe λ_+ et λ_- tels que

$$\forall t \in \mathbf{R}^*, x(t) = \begin{cases} \lambda_+ t + t^2 & \text{si } t > 0 \\ \lambda_- t + t^2 & \text{si } t < 0 \end{cases}$$

Cette fois, la fonction n'est dérivable que si $\lambda_- = \lambda_+$. Dans ce cas la fonction $x : t \mapsto \lambda t + t^2$ est bien de classe $\mathcal{C}^1$ et vérifie (E).

Exercice : Résoudre sur $]0, +\infty[$ l'équation, $t \ln(t)x' + x = 0$.

Remarque : Si α s'annule en plusieurs points (isolés), on peut « recoller » en tous les points.

Cas des équations d'ordre 2

On considère l'équation

$$\alpha(t)x''(t) = \beta(t)x'(t) + \gamma(t)x(t) + \delta \quad (E)$$

Si on suppose que α ne s'annule qu'un point t_0 de I . On peut résoudre l'équation différentielle sur $I_- = I \cap]-\infty, t_0[$ et sur $I_+ = I \cap]t_0, +\infty[$. Sur ces deux intervalles l'équation peut se ramener à une équation normalisée. On peut donc déterminer

$$\mathcal{S}_+ = \{x : I_+ \rightarrow \mathbf{K} \mid x \text{ vérifie } E\} = \{x_0^+ + \lambda u^+ + \mu^+ v^+ \mid \lambda^+, \mu^+ \in \mathbf{K}^2\}$$

et

$$\mathcal{S}_- = \{x : I_- \rightarrow \mathbf{K} \mid x \text{ vérifie } E\} = \{x_0^- + \lambda^- u^- + \mu^- v^- \mid \lambda^-, \mu^- \in \mathbf{K}^2\}.$$

Il ne reste plus qu'à raisonner de la même manière par analyse synthèse.

3 Équations différentielles à coefficients constants

On va généraliser l'étude des équations différentielles linéaire à coefficients constants d'ordre inférieur ou égal à 2 fait en première année.

On veut donc résoudre :

- Les systèmes différentiels à coefficients constants d'ordre 1 :

$$X'(t) = A.X(t) + B(t)$$

où $A \in \mathcal{M}_n(\mathbf{K})$, $X : I \rightarrow \mathcal{M}_{n,1}(\mathbf{K})$ et $B : I \rightarrow \mathcal{M}_{n,1}(\mathbf{K})$. Cette dernière étant continue.

- Les équations différentielles vectorielles à coefficients constants d'ordre 1 :

$$x'(t) = a(x(t)) + b(t)$$

où $A \in \mathcal{L}(E)$, $x : I \rightarrow E$ et $b : I \rightarrow E$. Cette dernière étant continue.

- Les équations différentielles scalaires d'ordre n :

$$x^{(n)}(t) = \sum_{i=0}^{n-1} a_i x^{(i)}(t) + b(t)$$

où $(a_0, \dots, a_{n-1}) \in \mathbf{K}^n$, $x : I \rightarrow \mathbf{K}$ et $b : I \rightarrow \mathbf{K}$. Cette dernière étant continue.

3.1 Exponentielle des matrices et endomorphismes

L'exponentielle des matrices et des endomorphismes a été vue précédemment. Rappelons les résultats déjà obtenus.

Définition 16.14

1. Soit a un endomorphisme de E . La série $\sum_{k \geq 0} \frac{a^k}{k!}$ converge. On appelle exponentielle de a sa somme et on note

$$\exp(a) = \sum_{k=0}^{+\infty} \frac{a^k}{k!}$$

2. Soit $A \in \mathcal{M}_n(\mathbf{K})$ une matrice carrée de taille n . La série $\sum_{k \geq 0} \frac{A^k}{k!}$ converge. On appelle exponentielle de A sa somme et on note

$$\exp(A) = \sum_{k=0}^{+\infty} \frac{A^k}{k!}$$

Proposition 16.15 (Propriétés de l'exponentielle de matrices)

1. La fonction $\exp : \mathcal{M}_n(\mathbf{K}) \rightarrow \text{GL}_n(\mathbf{K})$ est continue.
 2. Soit $A \in \mathcal{M}_n(\mathbf{K})$, la fonction $S : t \mapsto \exp(tA)$ est dérivable et

$$\forall t \in \mathbf{R}, S'(t) = A \exp(tA) = \exp(tA)A$$

3. Si A et B commutent alors $\exp(A+B) = \exp(A) \exp(B)$
 4. Soit $A \in \mathcal{M}_n(\mathbf{K})$ et $P \in \text{GL}_n(\mathbf{K})$. On note $B = P^{-1}AP$. On a $\exp(B) = P^{-1} \exp(A)P$.
 5. Soit $A \in \mathcal{M}_n(\mathbf{C})$, $\text{Sp}(\exp(A)) = \{\exp(\lambda), \lambda \in \text{Sp}(A)\}$.

Démonstration :

1. Déjà vu.
 2. Déjà vu.
 3. Déjà vu.
 4. Pour tout $N \in \mathbf{N}$,

$$\sum_{k=0}^N \frac{B^k}{k!} = P^{-1} \sum_{k=0}^N \frac{A^k}{k!} P$$

On sait que $M \mapsto P^{-1}MP$ est continue car linéaire sur un espace vectoriel de dimension finie. En faisant tendre $N \rightarrow +\infty$ on obtient,

$$\exp(B) = P^{-1} \exp(A)P$$

5. On peut trigonaliser la matrice A . Il existe T triangulaire supérieure et $P \in \text{GL}_n(\mathbf{C})$ telles que $P^{-1}AP = T$. Si on note $\lambda_1, \dots, \lambda_n$ les coefficients diagonaux de T (qui sont les valeurs propres de A). On en déduit que

$$P^{-1} \exp(A)P = \exp(P^{-1}AP) = \begin{pmatrix} e^{\lambda_1} & * & \dots & * \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & * \\ 0 & \dots & 0 & e^{\lambda_n} \end{pmatrix}$$

On en déduit que le spectre de $\exp(A)$ qui est le même que celui de $P^{-1} \exp(A)P$ est égal à $\{e^\lambda, \lambda \in \text{Sp}(A)\}$.

□

Remarque : On peut établir des résultats analogue pour l'exponentielle des endomorphismes.

3.2 Généralités

Théorème 16.16 (Éq. diff. vectorielle d'ordre 1 homogène à coefficients constants)

On considère une équation différentielle vectorielle d'ordre 1 homogène à coefficients constants

$$X'(t) = AX(t)$$

où $A \in \mathcal{M}_n(\mathbf{K})$.

L'ensemble des solutions est l'espace vectoriel de dimension n :

$$\mathcal{S}_H = \{t \mapsto \exp(tA)\Lambda \mid \Lambda \in \mathcal{M}_{n,1}(\mathbf{K})\}$$

ATTENTION

- Dans le cas où $n = 1$ on retrouve la formule usuelle. En effet, pour $t = 0$, on résout $x' = ax$ alors l'ensemble des solutions est

$$\mathcal{S}_H = \{t \mapsto \lambda e^{at}, \lambda \in \mathbf{K}\}$$

- Dans le cas scalaire, on écrit les solutions sous la forme $t \mapsto \lambda e^{at}$ avec le scalaire devant, dans le cas matriciel, le paramètre Λ doit être à droite.

Démonstration : Soit $\Lambda \in \mathcal{M}_{n,1}(\mathbf{K})$. On pose $X : t \mapsto \exp(tA)\Lambda$. Notons $S : I \rightarrow \mathcal{M}_n(\mathbf{K})$ définie par $S : t \mapsto \exp(tA)$ et $L : \mathcal{M}_n(\mathbf{K}) \rightarrow \mathcal{M}_{n,1}(\mathbf{K})$ définie par $L : M \mapsto M\Lambda$. Avec ces notations on a donc $X = L \circ S$ puisque pour tout $t \in I$, $X(t) = L(S(t))$.

On sait que S est dérivable et pour tout $t \in I$, $S'(t) = A \exp(tA) = AS(t)$. On en déduit que X est aussi dérivable et pour $t \in I$,

$$X'(t) = L(S'(t)) = S'(t)\Lambda = AS(t)\Lambda = AX(t)$$

Comme on sait que $\mathcal{S}_H$ est de dimension n , il suffit alors pour conclure de montrer que l'ensemble $T = \{t \mapsto \exp(tA)\Lambda \mid \Lambda \in \mathcal{M}_{n,1}(\mathbf{K})\}$ est un sous-espace vectoriel de dimension n . Pour cela on remarque que T est l'image de l'application linéaire $\Phi : \mathcal{M}_{n,1}(\mathbf{K}) \rightarrow \mathcal{C}^1(I, \mathcal{M}_{n,1}(\mathbf{K}))$ définie par $\Phi : \Lambda \mapsto (t \mapsto S(t)\Lambda)$. On peut voir que Φ est injective en effet, considérons $\Lambda \in \text{Ker}(\Phi)$, pour tout $t \in I$, $S(t)\Lambda = 0$ ce qui implique $\Lambda = 0$ car $S(t)$ est inversible.

L'application Φ étant injective, son image est de dimension n et donc

$$\mathcal{S}_H = \{t \mapsto \exp(tA)\Lambda \mid \Lambda \in \mathcal{M}_{n,1}(\mathbf{K})\}$$

□

Remarque : Pour tout t , on note $C_1(t), \dots, C_n(t)$ les colonnes de la matrice $\exp(tA)$ de sorte que pour tout $i \in \llbracket 1; n \rrbracket$, $C_i(t) = \exp(tA)E_i$ où E_i est le i -ème vecteur de la base canonique de $\mathcal{M}_{n,1}(\mathbf{K})$. Le théorème ci-dessus stipule que $(C_1, \dots, C_n)$ est une base de $\mathcal{S}_H$.

Corollaire 16.17 (Problème de Cauchy vectoriel d'ordre 1 à coefficients constants)

Soit $X_0 \in \mathcal{M}_{n,1}(\mathbf{K})$ et $t_0 \in \mathbf{R}$. L'unique solution du problème de Cauchy

$$(PC) \begin{cases} X'(t) = AX(t) \\ x(t_0) = X_0 \end{cases}$$

est $t \mapsto \exp((t - t_0)A)X_0 = \exp(tA) \exp(-t_0A)X_0$.

En particulier, si $t_0 = 0$, la solution est $t \mapsto \exp(tA)X_0$.

Démonstration : D'après le théorème de Cauchy linéaire on sait qu'il existe une unique solution du problème de Cauchy. D'après le théorème précédent, la fonction $X : t \mapsto \exp(tA) \exp(-t_0A)X_0$ vérifie l'équation différentielle. Il suffit alors de vérifier que $X(t_0) = x_0$.

□

Remarques :

1. On peut l'écrire sous la forme vectorielle.

La solution du problème de Cauchy

$$(PC) \begin{cases} x'(t) = a(x(t)) \\ x(t_0) = x_0 \end{cases}$$

où $t_0 \in \mathbf{R}$ est donc $x : t \mapsto \exp((t - t_0)a)x_0$.

2. Ce résultat ne se généralise pas au cas d'un équation $X'(t) = A(t)X(t)$ où $t \mapsto A(t)$ est une fonction continue à valeurs dans $\mathcal{M}_n(\mathbf{K})$ comme dans le cas des équations scalaires. En effet, si $t \mapsto Z(t)$ est une fonction telle que $Z'(t) = A(t)Z(t)$: une primitive de A . Si on pose $X : t \mapsto \exp(Z(t))$. Il n'est pas vrai que $X'(t) = Z'(t) \exp(Z(t))$. En effet, la dérivée de $t \mapsto Z^k(t)$ est

$$t \mapsto Z'(t)Z^{k-1}(t) + Z(t)Z'(t)Z^{k-2}(t) + \dots + Z^{k-1}(t)Z'(t)$$

et pas $kZ'(t)Z^{k-1}(t)$ car $Z(t)$ et $Z'(t)$ ne commutent pas toujours

3. On peut redémontrer que si $AB = BA$ alors $\exp(A + B) = \exp(A) \exp(B)$. On pose

$$X : t \mapsto \exp(A + tB) \in \mathcal{M}_n(\mathbf{K})$$

On écrit que X est la somme de la série de fonctions $\sum_{k \geq 0} \left(t \mapsto \frac{(A+tB)^k}{k!} \right)$. En utilisant que $AB = BA$, on peut appliquer le théorème de dérivabilité des séries de fonctions à valeurs vectorielles. On obtient que X est de classe $\mathcal{C}^1$ et que

$$X' : t \mapsto BX(t)$$

La fonction X est donc une solution du problème de Cauchy

$$(PC) \begin{cases} X'(t) = BX(t) \\ X(0) = \exp(A) \end{cases}$$

Ici, $X(t)$ est une matrice carrée et pas une matrice colonne. On travaille avec l'endomorphisme de $\mathcal{M}_n(\mathbf{K})$ défini par $X \mapsto BX$. Maintenant, $t \mapsto \exp(A) \exp(tB)$ est aussi solution du même problème de Cauchy. Elles sont égales.

3.3 Exemples de calculs

ATTENTION

Dans la majorité des cas, on ne calculera pas les exponentielles de matrices. On se ramènera via diagonalisation et trigonalisation aux résultats vus en première année.

Exemples :

1. Soit $A = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$.

On s'intéresse à l'équation différentielle matricielle homogène d'ordre 1

$$X'(t) = AX(t) \quad (\text{H})$$

D'après le résultat précédent, les solutions sont les fonctions de la forme $X : t \mapsto \exp(tA)\Lambda$ où $\Lambda \in \mathcal{M}_{2,1}(\mathbf{K})$.

On peut calculer directement $\exp(tA)$ en remarquant que $(tA)^2 = t^2 I_2$. De ce fait, pour tout entier $p \geq 0$,

$$(tA)^{2p} = t^{2p} I_2 \text{ et } (tA)^{2p+1} = t^{2p+1} A$$

On en déduit que

$$\exp(tA) = \left(\sum_{p=0}^{\infty} \frac{t^{2p}}{(2p)!} \right) I_2 + \left(\sum_{p=0}^{\infty} \frac{t^{2p+1}}{(2p+1)!} \right) A = \text{ch}(t) I_2 + \text{sh}(t) A = \begin{pmatrix} \text{ch}(t) & \text{sh}(t) \\ \text{sh}(t) & \text{ch}(t) \end{pmatrix}$$

En notant $\Lambda = \begin{pmatrix} \alpha \\ \beta \end{pmatrix}$, on obtient que les solutions sont les fonctions de la forme

$$X : t \mapsto \begin{pmatrix} \alpha \text{ch}(t) + \beta \text{sh}(t) \\ \alpha \text{sh}(t) + \beta \text{ch}(t) \end{pmatrix}, \quad \alpha, \beta \in \mathbf{K}$$

On remarque que l'équation (H) est l'écriture matricielle de l'équation scalaire d'ordre 2 : $x''(t) = x(t)$. On retrouve que les solutions de cette équation sont les fonctions de la forme

$$x : t \mapsto \alpha \text{ch}(t) + \beta \text{sh}(t)$$

La deuxième ligne de la matrice colonne X étant alors la dérivée de la première.

2. Si on considère $\mathbf{K} = \mathbf{C}$ et l'équation différentielle

$$ax''(t) + bx'(t) + cx(t) = 0 \quad (\text{H})$$

où $a \in \mathbf{C} \setminus \{0\}$ et $(b, c) \in \mathbf{C}^2$. On pose

$$X(t) = \begin{pmatrix} x(t) \\ x'(t) \end{pmatrix}$$

et

$$A = \begin{pmatrix} 0 & 1 \\ -\frac{c}{a} & -\frac{b}{a} \end{pmatrix}$$

L'équation peut alors s'écrire sous la forme du système

$$X'(t) = \begin{pmatrix} x'(t) \\ x''(t) \end{pmatrix} = AX(t) \quad (\text{S})$$

Le polynôme caractéristique de A est

$$\chi_A = X \left(X + \frac{b}{a} \right) + \frac{c}{a} = \frac{1}{a} (aX^2 + bX + c)$$

En particulier si on note Δ le discriminant du trinôme $aX^2 + bX + c$.

- si $\Delta \neq 0$: il y a deux valeurs propres distinctes α et β . De ce fait, A est semblable à la matrice diagonale

$$D = \begin{pmatrix} \alpha & 0 \\ 0 & \beta \end{pmatrix}.$$

Si on note P la matrice de changement de base, telle que $D = P^{-1}AP$, on a donc

$$t \mapsto \exp(tA) = P \exp(tD) P^{-1} = P \begin{pmatrix} e^{\alpha t} & 0 \\ 0 & e^{\beta t} \end{pmatrix} P^{-1}$$

On en déduit que

$$\mathcal{S}_S = \{t \mapsto P \exp(tD) P^{-1} \Lambda \mid \Lambda \in \mathcal{M}_{2,1}(\mathbb{C})\} = \{t \mapsto P \exp(tD) U \mid U \in \mathcal{M}_{2,1}(\mathbb{C})\}$$

La deuxième égalité vient du fait que $\Lambda \mapsto U = P^{-1}\Lambda$ est une bijection de $\mathcal{M}_{2,1}(\mathbb{C})$ dans lui même.

On écrit alors

$$P = \begin{pmatrix} p_{11} & p_{12} \\ p_{21} & p_{22} \end{pmatrix}$$

En notant $U = \begin{pmatrix} u \\ v \end{pmatrix}$ et en prenant la première ligne de l'ensemble des solutions ci-dessus on obtient

$$S_H = \{t \mapsto p_{11}ue^{\alpha t} + p_{12}ve^{t\beta}, (u, v) \in \mathbb{C}^2\}$$

On remarque alors que $p_{11} \neq 0$ et $p_{12} \neq 0$ car $\begin{pmatrix} 0 \\ 1 \end{pmatrix}$ n'est pas un vecteur propre de A et donc, en utilisant que $u \mapsto u' = p_{11}u$ et $v \mapsto v' = p_{12}v$ sont des bijections de $\mathbb{C}$ dans $\mathbb{C}$, on obtient finalement,

$$S_H = \{t \mapsto u'e^{\alpha t} + v'e^{t\beta}, (u', v') \in \mathbb{C}^2\} = \text{Vect}\left(t \mapsto e^{\alpha t}, t \mapsto e^{t\beta}\right)$$

C'est le résultat vu en première année.

- Si $\Delta = 0$. Il y a alors une racine double α . Comme $A \neq \alpha I_2$ alors elle n'est pas diagonalisable. Elle est donc semblable à

$$T = \begin{pmatrix} \alpha & 1 \\ 0 & \alpha \end{pmatrix}$$

Si on note P la matrice de changement de base, telle que $T = P^{-1}AP$, on a donc

$$t \mapsto \exp(tA) = P \exp(tT) P^{-1}$$

Maintenant, on voit que $tT = t\alpha I_2 + N$ où $N = \begin{pmatrix} 0 & t \\ 0 & 0 \end{pmatrix}$ vérifie $N^2 = 0$. Or $t\alpha I_2$ et N commutent et donc

$$\exp(tT) = \exp(t\alpha I_2) \exp(N) = \exp(t\alpha) I_2 \times (I_2 + N) = \begin{pmatrix} e^{t\alpha} & te^{\alpha t} \\ 0 & e^{t\alpha} \end{pmatrix}$$

En procédant comme dans le cas précédent, on obtient finalement que les solutions de (H) sont des combinaisons linéaires des fonctions $t \mapsto e^{\alpha t}$ et $t \mapsto te^{\alpha t}$ comme il a été vu en première année.

3. Cas d'un système diagonalisable. On considère le système

$$\begin{cases} x' &= -3x + 3y + 2z \\ y' &= -4x + 4y + 2z \\ z' &= -4x + 3y + 3z \end{cases}$$

c'est-à-dire

$$X' = AX$$

en posant $X = \begin{pmatrix} x \\ y \\ z \end{pmatrix}$ et $A = \begin{pmatrix} -3 & 3 & 2 \\ -4 & 4 & 2 \\ -4 & 3 & 3 \end{pmatrix}$.

On commence par diagonaliser (ou trigonaliser A). Ici $\chi_A = X^3 - 4X^2 + 5X - 2 = (X-1)^2(X-2)$.
On regarde donc les espaces propres

$$E_1(A) = \text{Vect}\left(\begin{pmatrix} 1 \\ 0 \\ 2 \end{pmatrix}, \begin{pmatrix} 0 \\ 2 \\ -3 \end{pmatrix}\right) \text{ et } E_2(A) = \text{Vect}\left(\begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}\right).$$

On pose alors

$$P = \begin{pmatrix} 1 & 0 & 1 \\ 0 & 2 & 1 \\ 2 & -3 & 1 \end{pmatrix} \text{ et } D = P^{-1}AP = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 2 \end{pmatrix}$$

En posant alors $Y = P^{-1}X$, on obtient

$$Y'(t) = P^{-1}X'(t) = P^{-1}AX(t) = P^{-1}APP^{-1}X(t) = DY(t).$$

On en déduit que

$$Y(t) = \begin{pmatrix} \alpha e^t \\ \beta e^t \\ \gamma e^{2t} \end{pmatrix}$$

De ce fait

$$X = PY = \begin{pmatrix} \alpha e^t + \gamma e^{2t} \\ 2\beta e^t + \gamma e^{2t} \\ 2\alpha e^t - 3\beta e^t + \gamma e^{2t} \end{pmatrix} = \alpha \begin{pmatrix} e^t \\ 0 \\ 2e^t \end{pmatrix} + \beta \begin{pmatrix} 0 \\ 2e^t \\ -3e^t \end{pmatrix} + \gamma \begin{pmatrix} e^{2t} \\ e^{2t} \\ e^{2t} \end{pmatrix}.$$

Remarque : Avec cette méthode 2, on n'a pas eu besoin de calculer P^{-1} ni $\exp(A)$.

Exercice : Résoudre le système différentiel,

$$\begin{cases} x' &= 2y + 2z \\ y' &= -x + 2y + 2z \\ z' &= -x + y + 3z \end{cases}$$

On fera attention au fait que cette fois la matrice du système est juste trigonalisable.

4 Equations différentielles scalaires du second ordre

Nous allons voir comment généraliser aux équations différentielles scalaires d'ordre 2 les méthodes vues en première année pour les équations différentielles scalaires d'ordre 1.

Dans tout ce paragraphe, on considère une équation différentielle scalaire d'ordre 2 normalisée

$$x''(t) = a(t)x'(t) + b(t)x(t) + c(t) \quad (E)$$

où a, b et c sont continues de I dans $\mathbf{K}$. L'équation homogène associée est

$$x''(t) = a(t)x'(t) + b(t)x(t) \quad (H)$$

On peut les représenter de manière vectorielle en posant

$$X(t) = \begin{pmatrix} x(t) \\ x'(t) \end{pmatrix}; A(t) = \begin{pmatrix} 0 & 1 \\ b(t) & a(t) \end{pmatrix} \text{ et } B(t) = \begin{pmatrix} 0 \\ c(t) \end{pmatrix}$$

On considère alors les systèmes différentiels

$$X'(t) = A(t)X(t) + B(t) \quad (S_E)$$

et

$$X'(t) = A(t)X(t) \quad (S)$$

4.1 Wronskien d'un couple de solution

Définition 16.18 (Wronskien)

Soit (x_1, x_2) des solutions de (H) , on appelle wronskien du couple (x_1, x_2) et on note $w_{(x_1, x_2)} : I \rightarrow \mathbf{K}$ la fonction :

$$w_{(x_1, x_2)} : t \mapsto \begin{vmatrix} x_1(t) & x_2(t) \\ x_1'(t) & x_2'(t) \end{vmatrix}$$

Remarque : On notera le plus souvent juste w le wronskien d'un couple de solutions.

Exemple : Pour l'équation $x'' = -x$ on peut prendre $x_1 : t \mapsto \cos t$ et $x_2 : t \mapsto \sin t$. On obtient alors

$$w : t \mapsto \cos^2(t) + \sin^2(t) = 1.$$

Exercice : Calculer le wronskien pour un couple de solutions de $x'' = x$.

Proposition 16.19

Le wronskien d'un couple de solutions l'équation différentielle

$$x''(t) = a(t)x'(t) + b(t)x(t)$$

vérifie l'équation du premier ordre

$$w'(t) = a(t)w(t)$$

Démonstration : Il suffit de calculer $w'(t)$. Pour tout $t \in I$,

$$\begin{aligned} w'(t) &= x_1'(t)x_2'(t) + x_1(t)x_2''(t) - x_2'(t)x_1'(t) - x_2(t)x_1''(t) \\ &= x_1(t)(a(t)x_2'(t) + b(t)x_2(t)) - x_2(t)(a(t)x_1'(t) + b(t)x_1(t)) \\ &= a(t)w(t) \end{aligned}$$

Remarque : Dans le cas particulier où $a = \tilde{0}$ on voit que le wronskien est constant comme sur l'exemple traité ci-dessus.

Définition 16.20 (Système fondamental de solutions)

On appelle système fondamental de solutions d'une équation différentielle linéaire scalaire d'ordre 2 homogène, un couple (x_1, x_2) de solutions de (H) qui forme une base de l'espace des solutions $\mathcal{S}_H$.

Proposition 16.21

Avec les notations précédentes, pour (x_1, x_2) un couple de solutions de (H) et w le wronskien associé, les assertions suivantes sont équivalentes

- i) Le couple (x_1, x_2) est un système fondamental de solutions de (H) .
- ii) Il existe $t_0 \in I$ tel que $w(t_0) \neq 0$.
- iii) Pour tout $t \in I$, $w(t) \neq 0$.

Démonstration :

- $\boxed{iii) \Rightarrow ii)}$ Evident
- $\boxed{ii) \Rightarrow i)}$ Supposons par contraposée que le couple (x_1, x_2) n'est pas un système fondamental. Cela signifie qu'il existe λ_1, λ_2 des scalaires non tous nuls tels que

$$\lambda_1 x_1 + \lambda_2 x_2 = \tilde{0}$$

On peut alors dériver cette relation

$$\lambda_1 x_1' + \lambda_2 x_2' = \tilde{0}$$

En particulier pour tout $t \in I$,

$$\lambda_1 x_1(t) + \lambda_2 x_2(t) = \lambda_1 x_1'(t) + \lambda_2 x_2'(t) = 0$$

On en déduit que

$$\lambda_1 \begin{pmatrix} x_1(t) \\ x_1'(t) \end{pmatrix} + \lambda_2 \begin{pmatrix} x_2(t) \\ x_2'(t) \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \end{pmatrix}$$

On en déduit que $w(t) = 0$.

- $\boxed{i) \Rightarrow iii)}$ On note $X_1 : t \mapsto \begin{pmatrix} x_1(t) \\ x_1'(t) \end{pmatrix}$ et $X_2 : t \mapsto \begin{pmatrix} x_2(t) \\ x_2'(t) \end{pmatrix}$

Supposons par l'absurde qu'il existe t_0 tel que $w(t_0) = 0$ ce qui signifie que $(X_1(t_0), X_2(t_0))$ est liée. Il existe alors des scalaires λ_1 et λ_2 tels que

$$\lambda_1 X_1(t_0) + \lambda_2 X_2(t_0) = 0$$

On considère $Y : t \mapsto \lambda_1 X_1(t) + \lambda_2 X_2(t)$. C'est une fonction vérifiant (S_H) et telle que $Y(t_0) = 0$ donc d'après le théorème de Cauchy linéaire c'est la fonction nulle. Cela signifie que

$$\forall t \in I, \lambda_1 X_1(t) + \lambda_2 X_2(t) = 0$$

En particulier, $\lambda_1 x_1 + \lambda_2 x_2 = \tilde{0}$ ce qui est absurde car (x_1, x_2) est libre.

□

Remarque : On pouvait aussi directement montrer que $\boxed{ii) \Rightarrow iii)}$. On sait que w vérifie l'équation différentielle linéaire d'ordre 1, $w'(t) = a(t)w(t)$. Si on note α une primitive de a , il existe $\lambda \in \mathbf{K}$ telle que w soit de la forme

$$w : t \mapsto \lambda \exp(\alpha(t))$$

Comme pour tout $t \in I$, $\exp(\alpha(t)) \neq 0$,

$$(\exists t_0 \in I, w(t_0) = 0) \iff \lambda = 0 \iff (\forall t \in I, w(t) = 0)$$

4.2 Variation de la constante

On se donne un système fondamental de solutions de (H) : (x_1, x_2) . On note encore

$$X_1 : t \mapsto \begin{pmatrix} x_1(t) \\ x_1'(t) \end{pmatrix} ; X_2 : t \mapsto \begin{pmatrix} x_2(t) \\ x_2'(t) \end{pmatrix} \text{ et } Q : t \mapsto \begin{pmatrix} x_1 & x_2(t) \\ x_1'(t) & x_2'(t) \end{pmatrix} = (X_1(t) | X_2(t))$$

On veut résoudre le système différentiel (S). Soit $X \in \mathcal{C}^1(I, \mathcal{M}_{2,1}(\mathbf{K}))$, comme pour tout $t \in I$, $w(t) \neq 0$, la matrice $Q(t)$ est inversible et on peut donc poser $X(t) = Q(t)\Lambda(t)$ ce qui est équivalent à $\Lambda(t) = Q^{-1}(t)X(t)$.

De plus, on voit que $Q^{-1}(t) = \frac{1}{w(t)} \begin{pmatrix} x_2'(t) & -x_2(t) \\ x_1'(t) & x_1(t) \end{pmatrix}$ ce qui montre que la fonction $t \mapsto Q^{-1}(t)$ est de classe $\mathcal{C}^1(I, \mathcal{M}_2(\mathbf{K}))$. Finalement l'application $X \mapsto Q^{-1}X$ est une bijection de $\mathcal{C}^1(I, \mathcal{M}_{2,1}(\mathbf{K}))$ dans lui-même.

Si on pose $\Lambda(t) = \begin{pmatrix} \lambda_1(t) \\ \lambda_2(t) \end{pmatrix}$, on est ramené à chercher les solutions de (S) sous la forme

$$t \mapsto X(t) = Q(t)\Lambda(t) = \begin{pmatrix} x_1(t) & x_2(t) \\ x_1'(t) & x_2'(t) \end{pmatrix} \begin{pmatrix} \lambda_1(t) \\ \lambda_2(t) \end{pmatrix} = \begin{pmatrix} \lambda_1(t)x_1(t) + \lambda_2(t)x_2(t) \\ \lambda_1(t)x_1'(t) + \lambda_2(t)x_2'(t) \end{pmatrix}$$

Le produit matriciel étant bilinéaire on a donc

$$X'(t) = Q'(t)\Lambda(t) + Q(t)\Lambda'(t)$$

Or, en travaillant colonne par colonne, on obtient que

$$Q'(t) = (X_1'(t) | X_2'(t)) = A(t) (X_1(t) | X_2(t)) = A(t)Q(t)$$

On en déduit que

$$Q'(t) = A(t)Q(t)\Lambda(t) + Q(t)\Lambda'(t) = AX(t) + Q(t)\Lambda'(t)$$

On obtient alors

$$\begin{aligned} X \text{ vérifie (S)} &\iff \forall t \in I, X'(t) = A(t)X(t) + B(t) \\ &\iff \forall t \in I, AX(t) + Q(t)\Lambda'(t) = A(t)X(t) + B(t) \\ &\iff \forall t \in I, Q(t)\Lambda'(t) = B(t) \quad (\dagger) \\ &\iff \forall t \in I, \Lambda'(t) = Q^{-1}(t)B(t) \end{aligned}$$

On voit que, comme dans la méthode classique pour les équations différentielles scalaires d'ordre 1, on est ramené à calculer des primitives pour déterminer une solution de (S) et donc de (E).

Exemples :

1. Considérons l'équation différentielle

$$x''(t) = -x(t) + \cos(t) \quad (E)$$

Avec les notations ci-dessus, on pose donc $A = \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}$.

On sait que $x_1 : t \mapsto \cos(t)$ et $x_2 : t \mapsto \sin(t)$ est une base de $\mathcal{S}_H$. On a donc $Q(t) = \begin{pmatrix} \cos(t) & \sin(t) \\ -\sin(t) & \cos(t) \end{pmatrix}$. La matrice $Q(t)$ est orthogonale, on en déduit que $Q^{-1}(t) = Q^T(t) = \begin{pmatrix} \cos(t) & -\sin(t) \\ \sin(t) & \cos(t) \end{pmatrix} = Q(-t)$.

On cherche une solution du système différentiel associé sous la forme $X(t) = Q(t)\Lambda(t)$. On posant $\Lambda(t) = \begin{pmatrix} \lambda_1(t) \\ \lambda_2(t) \end{pmatrix}$ et on est ramené à résoudre

$$\begin{cases} \lambda_1'(t) &= -\cos(t) \sin(t) \\ \lambda_2'(t) &= \cos^2(t) = \frac{1+\cos(2t)}{2} \end{cases}$$

On peut donc prendre $\lambda_1 : t \mapsto \frac{1}{2} \cos^2(t)$ et $\lambda_2 : t \mapsto \frac{1}{2}t + \frac{1}{4} \sin(2t)$. On en déduit une solution particulière de (E)

$$x : t \mapsto \frac{1}{2} \cos^2(t) \cos(t) + \frac{1}{2}t \sin(t) + \frac{1}{4} \sin(2t) \sin(t)$$

2. On peut présenter les calculs différemment sans faire intervenir les systèmes différentiels mais en utilisant une « ruse ».

On cherche les solutions de (E) sous la forme

$$x : t \mapsto \lambda_1(t)x_1(t) + \lambda_2(t)x_2(t)$$

On peut calculer les dérivées de x et réinjecter dans l'équation. On a

$$x' : t \mapsto \lambda_1'(t)x_1(t) + \lambda_1(t)x_1'(t) + \lambda_2'(t)x_2(t) + \lambda_2(t)x_2'(t)$$

On ajoute la condition $\lambda_1'x_1 + \lambda_2'x_2 = 0$. Cette condition semble ici artificielle mais elle découle du calcul ci-dessus avec les système différentiels. C'est la première ligne de l'équation (†).

On peut alors redériver. On obtient donc

$$x'' = \lambda_1'x_1' + \lambda_1x_1'' + \lambda_2'x_2' + \lambda_2x_2''$$

En injectant dans l'équation cela donne

$$\begin{aligned} x \text{ vérifie (E)} &\iff \lambda_1'x_1' + \lambda_1x_1'' + \lambda_2'x_2' + \lambda_2x_2'' = a(\lambda_1x_1' + \lambda_2x_2') + b(\lambda_1x_1 + \lambda_2x_2) + c \\ &\iff \lambda_1'x_1' + \lambda_2'x_2' = c \end{aligned}$$

On trouve donc le système (†)

$$\forall t \in I, \begin{cases} \lambda_1'(t)x_1(t) + \lambda_2'(t)x_2(t) &= 0 \\ \lambda_1'(t)x_1'(t) + \lambda_2'(t)x_2'(t) &= c(t) \end{cases}$$

C'est-à-dire $Q(t)\Lambda'(t) = B(t)$.

4.3 Résolutions d'une équation scalaire du second ordre en connaissant une solution

On considère une équation linéaire scalaire d'ordre 2 (pas nécessairement normalisée)

$$\alpha(t)x''(t) = \beta(t)x'(t) + \gamma(t)x(t) + \delta(t) \quad (E)$$

On considère alors son équation homogène associée

$$\alpha(t)x''(t) = \beta(t)x'(t) + \gamma(t)x(t) \quad (H)$$

et on suppose que l'on connaît une solution x_0 de (H) qui ne s'annule pas sur I .

Remarque : La solution x_0 pourra être déterminée par exemple à l'aide d'un développement en série entière.

On peut alors chercher à résoudre (E) à l'aide d'un **changement de fonction inconnue** en posant $x : t \mapsto \lambda(t)x_0(t)$. On a en effet

$$\begin{aligned} x \text{ vérifie (E)} &\iff \alpha x'' = \beta x' + \gamma x + \delta \\ &\iff \alpha(\lambda'' x_0 + 2\lambda' x_0' + \lambda x_0'') = \beta(\lambda' x_0 + \lambda x_0') + \gamma \lambda x_0 + \delta \\ &\iff \alpha \lambda'' x_0 + 2\alpha \lambda' x_0' = \beta \lambda' x_0 + \delta \\ &\iff \lambda' \text{ vérifie } \alpha z' = (\beta x_0 - 2\alpha x_0')z + \delta \end{aligned}$$

Cette dernière équation étant d'ordre 1.

Exemple : On veut résoudre sur $\mathbf{R}$

$$tx''(t) + 2x'(t) - tx(t) = 0$$

— On commence par chercher une solution développable en série entière au voisinage de 0. Soit

$$x : t \mapsto \sum_{n=0}^{+\infty} a_n t^n.$$

$$\begin{aligned} x \text{ vérifie (E)} &\iff \forall t \in \mathbf{R}, \sum_{n=2}^{+\infty} a_n n(n-1)t^{n-1} + \sum_{n=1}^{+\infty} 2a_n n t^{n-1} - \sum_{n=0}^{+\infty} a_n t^{n+1} = 0 \\ &\iff \forall t \in \mathbf{R}, \sum_{n=1}^{+\infty} a_{n+1} n(n+1)t^n + \sum_{n=0}^{+\infty} 2a_{n+1}(n+1)t^n - \sum_{n=1}^{+\infty} a_{n-1} t^n = 0 \\ &\iff \forall t \in \mathbf{R}, 2a_1 + \sum_{n=1}^{+\infty} [a_{n+1}(n+1)(n+2) - a_{n-1}] t^n = 0 \\ &\iff \begin{cases} a_1 = 0 \\ \forall n \geq 1, a_{n+1} = \frac{a_{n-1}}{(n+1)(n+2)} \end{cases} \end{aligned}$$

On en déduit par une récurrence immédiate que pour tout entier p , $a_{2p+1} = 0$ et $a_{2p} = \frac{1}{(2p+1)!} a_0$. En prenant $a_0 = 1$ on en déduit que

$$x : t \mapsto \sum_{p=0}^{+\infty} \frac{1}{(2p+1)!} t^{2p} = \frac{\text{sh } t}{t}.$$

vérifie l'équation.

- La fonction $x_0 : t \mapsto \frac{\text{sh } t}{t}$ ne s'annule pas sur $\mathbf{R}$. De ce fait, pour toute fonction x définie sur $\mathbf{R}$ on peut poser

$$x = \lambda \cdot x_0.$$

Maintenant

$$\begin{aligned} x \text{ vérifie (E)} &\iff \forall t \in \mathbf{R}, tx''(t) + 2x'(t) - tx(t) = 0 \\ &\iff (...) \\ &\iff \forall t \in \mathbf{R}, \text{sh } t \lambda''(t) = -2 \text{ch } t \lambda'(t) \end{aligned}$$

On est donc ramené à résoudre l'équation

$$\text{sh } tz' = -2 \text{ch } tz.$$

Travaillons d'abord sur $\mathbf{R}_+^*$ et $\mathbf{R}_-^*$. On a donc l'équation résolue :

$$z' = -2 \frac{\text{ch } t}{\text{sh } t} z.$$

Comme $t \mapsto -2 \ln |\text{sh } t|$ est une primitive de $t \mapsto -2 \frac{\text{ch } t}{\text{sh } t}$ on en déduit que les solutions de cette équation sont (sur $\mathbf{R}_+^*$ et $\mathbf{R}_-^*$) :

$$t \mapsto C e^{-2 \ln |\text{sh } t|} = \frac{C}{\text{sh }^2 t}.$$

Note

On vient de trouver λ' et pas λ , il faut donc intégrer.

Quitte à modifier la constante, on trouve donc que sur $\mathbf{R}_+^*$ et $\mathbf{R}_-^*$, $\lambda : t \mapsto C \frac{\text{ch } t}{\text{sh } t} + C'$. Finalement, les solutions sur $\mathbf{R}_+^*$ et $\mathbf{R}_-^*$ de (E) sont

$$t \mapsto \alpha \frac{\text{sh } t}{t} + \beta \frac{\text{ch } t}{t}.$$

Il est alors clair que les seules fonctions qui « s'étendent » à $\mathbf{R}$ sont celles où $\beta = 0$.

Intégrales à paramètres

1	Théorèmes généraux	450
1.1	Limite et continuité	451
1.2	Dérivabilité	454
2	Exemples	459

Nous allons étudier dans ce chapitre les intégrales à paramètres. Précisément, dans tout ce chapitre, on supposera avoir une fonction f définie sur $A \times I$ où A est une partie d'un espace vectoriel de dimension finie et I un intervalle de $\mathbf{R}$. On suppose que f est à valeurs dans $\mathbf{K} = \mathbf{R}$ ou $\mathbf{C}$.

On pose alors

$$\begin{aligned} F : A &\rightarrow \mathbf{K} \\ x &\mapsto \int_I f(x, t) dt \end{aligned}$$

Notre but est d'étudier la fonction F : définition, continuité, dérivabilité,...

Commençons par donner quelques exemples qui seront développés à la fin du chapitre :

1. La fonction Γ : on pose

$$\Gamma(x) = \int_0^{+\infty} t^{x-1} e^{-t} dt$$

2. Transformé de Laplace : soit f une fonction continue par morceaux sur $]0, +\infty[$ on pose

$$\mathcal{L}(f) : p \mapsto \int_0^{+\infty} f(t) e^{-pt} dt.$$

3. Transformée de Fourier : soit f une fonction continue sur $\mathbf{R}$, on pose

$$\mathcal{F}(f) = \hat{f} : \omega \mapsto \int_{-\infty}^{+\infty} f(t) e^{-i\omega t} dt$$

1 Théorèmes généraux

1.1 Limite et continuité

Commençons par généraliser le théorème de convergence dominée au cas où le paramètre n n'est plus entier.

Théorème 17.1 (Rappel : Théorème de convergence dominée)

Soit (f_n) une suite de fonctions continues par morceaux sur un intervalle I à valeurs dans $\mathbf{K}$. On suppose

- i) La suite (f_n) converge vers une fonction f continue par morceaux
- ii) (Hypothèse de domination) : Il existe une fonction φ **intégrable** sur I telle que pour tout entier n , $|f_n| \leq \varphi$.

Alors f est intégrable sur I et

$$\lim_{n \rightarrow +\infty} \int_I f_n = \int_I f$$

En notant $f(n, t)$ pour $f_n(t)$, la conclusion du théorème ci-dessus est que :

$$\lim_{n \rightarrow +\infty} \int_I f(n, t) dt = \int_I g(t) dt$$

où, pour tout $t \in I$, $\lim_{n \rightarrow +\infty} f(n, t) = g(t)$.

Théorème 17.2 (Théorème de convergence dominée - version continue)

Soit f une fonction à valeurs dans $\mathbf{K} = \mathbf{R}$ ou $\mathbf{C}$ définie sur $A \times I$ où A est une partie d'un espace vectoriel de dimension finie et I un intervalle de $\mathbf{R}$.

Soit a un point adhérent à A .

On suppose qu'il existe un voisinage V de a dans A tel que

- i) Pour tout x de V , la fonction $t \mapsto f(x, t)$ est continue par morceaux
- ii) Il existe une fonction $g : I \rightarrow \mathbf{K}$ continue par morceaux telle que :

$$\forall t \in I, \lim_{x \rightarrow a} f(x, t) = g(t)$$

- iii) Hypothèse de domination : Il existe une fonction $\varphi : V \rightarrow \mathbf{R}_+$ continue par morceaux et **intégrable** sur I telle que

$$\forall x \in V, \forall t \in I, |f(x, t)| \leq \varphi(t)$$

Alors, $F : x \mapsto \int_I f(x, t) dt$ est définie sur V et

$$\lim_{x \rightarrow a} F(x) = \int_I g(t) dt.$$

Remarques :

1. C'est une version « continue » du théorème de convergence dominée. En effet, plutôt que d'écrire $f(x, t)$, on pourrait voir la fonction f comme une famille $(f_x)_{x \in A}$ de fonctions définies sur I .

2. Les deux premières hypothèses sont les analogues du fait que (f_n) converge simplement vers g continue par morceaux.
3. Là encore, l'hypothèse de domination implique que pour tout x de V , $t \mapsto f(x, t)$ est intégrable.
4. C'est un résultat d'interversion $\lim$ et $\int_I$. En effet la conclusion est que

$$\lim_{x \rightarrow a} \int_I f(x, t) dt = \int_I \lim_{x \rightarrow a} f(x, t) dt.$$

5. Dans le cas où A est une partie non bornée de $\mathbf{R}$ on peut prendre $a \in \{\pm\infty\}$.

Démonstration : Le principe est juste de se ramener au cas « classique » du théorème de convergence dominée en utilisant la caractérisation séquentielle.

- Comme dit dans la remarque, le fait que F soit définie sur V découle directement de l'hypothèse de domination et des théorèmes de comparaisons pour les intégrales.
- Par la caractérisation séquentielle de la limite, pour montrer que $\lim_{x \rightarrow a} F(x) = \int_I g(t) dt$ il suffit de montrer que pour toute suite (x_n) d'éléments de V tendant vers a on a $\lim_{n \rightarrow \infty} F(x_n) = \int_I g(t) dt$.

Soit (x_n) une telle suite, on pose $f_n : t \mapsto f(x_n, t)$.

Les hypothèses *i)* et *ii)* affirment alors que la suite de fonctions (f_n) converge simplement vers g . On peut donc appliquer le théorème de convergence dominée « classique ». On obtient que

$$\lim_{n \rightarrow \infty} F(x_n) = \lim_{n \rightarrow \infty} \int_I f_n = \int_I g.$$

Maintenant, on peut faire cela pour toute suite (x_n) tendant vers a . On conclut alors en utilisant la caractérisation séquentielle de la limite.

□

Exemple : On pose

$$\begin{aligned} f : \mathbf{R}_+ \times \mathbf{R}_+ &\rightarrow \mathbf{R} \\ (x, t) &\mapsto \frac{e^{-xt}}{1+t^2} \end{aligned}$$

On veut étudier $\lim_{x \rightarrow +\infty} \int_0^{+\infty} f(x, t) dt$.

- i)* Pour tout t de $\mathbf{R}_+$, $f(., t) : x \mapsto f(x, t)$ admet une limite finie notée $g(t)$ quand x tend vers $a = +\infty$. Il suffit de poser

$$g : t \mapsto \begin{cases} 1 & \text{si } t = 0 \\ 0 & \text{sinon.} \end{cases}$$

- ii)* La fonction g est continue par morceaux.

- iii)* (Hypothèse de domination) : Il existe une fonction φ intégrable sur $I = \mathbf{R}_+$ telle que pour tout $x \in \mathbf{R}_+$ et tout $t \in \mathbf{R}_+$, $|f(x, t)| \leq \varphi(t)$. On peut prendre $\varphi : t \mapsto \frac{1}{1+t^2}$.

On en déduit que pour tout x de $\mathbf{R}_+$, $t \mapsto f(x, t)$ et g sont intégrables et

$$\lim_{x \rightarrow +\infty} \int_0^{+\infty} \frac{e^{-xt}}{1+t^2} dt = \int_0^{+\infty} g = 0.$$

Corollaire 17.3 (Théorème de continuité (en un point) des intégrales à paramètres)

Soit f une fonction à valeurs dans $\mathbf{K} = \mathbf{R}$ ou $\mathbf{C}$ définie sur $A \times I$ où A est une partie d'un espace vectoriel de dimension finie et I un intervalle de $\mathbf{R}$.

Soit $a \in A$. On suppose qu'il existe un voisinage V de a (indépendant de t) tel que

- i) Pour tout x de V , la fonction $t \mapsto f(x, t)$ est continue par morceaux
- ii) Pour tout $t \in I$, $f(., t) : x \mapsto f(x, t)$ est continue en a
- iii) Hypothèse de domination : Il existe une fonction $\varphi : I \rightarrow \mathbf{R}_+$ continue par morceaux et **intégrable** sur I telle que

$$\forall x \in V, \forall t \in I, |f(x, t)| \leq \varphi(t)$$

Alors, $F : x \mapsto \int_I f(x, t) dt$ est définie sur V et

$$\lim_{x \rightarrow a} F(x) = F(a)$$

c'est-à-dire que F est continue en a .

Corollaire 17.4 (Théorème de continuité des intégrales à paramètres)

Soit f une fonction à valeurs dans $\mathbf{K} = \mathbf{R}$ ou $\mathbf{C}$ définie sur $A \times I$ où A est une partie d'un espace vectoriel de dimension finie et I un intervalle de $\mathbf{R}$.

- i) Pour tout x de A , la fonction $f(x, .) : t \mapsto f(x, t)$ est continue par morceaux
- ii) Pour tout $t \in I$, $f(., t) : x \mapsto f(x, t)$ est continue sur A
- iii) Hypothèse de domination : pour tout a de A , il existe un voisinage K de a dans A et une fonction $\varphi_K : I \rightarrow \mathbf{R}_+$ continue par morceaux et **intégrable** sur I telle que

$$\forall x \in K, \forall t \in I, |f(x, t)| \leq \varphi_K(t)$$

Alors, $F : x \mapsto \int_I f(x, t) dt$ est définie sur A et continue

Remarques :

1. Dans le cas où A est un intervalle de $\mathbf{R}$. L'hypothèse de domination peut s'écrire : « pour tout segment K de A ... ».
2. Le programme officiel dit : « Pour l'application pratique des énoncés, on vérifie les hypothèses de régularité par rapport à x et de domination, sans expliciter celles relatives à la continuité par morceaux par rapport à t »

Exemple : La fonction Γ .

On pose pour $x \in \mathbf{R}$,

$$\Gamma(x, t) = \int_0^{+\infty} t^{x-1} e^{-t} dt.$$

On pose $f(x, t) = t^{x-1} e^{-t}$.

— Définition : Pour tout $x \in \mathbf{R}$, la fonction $t \mapsto f(x, t)$ est continue et positive sur $]0, +\infty[$.

- Pour tout $x \in \mathbb{R}$, $t^{x-1+2}e^{-t} \xrightarrow[t \rightarrow +\infty]{} 0$ donc $t^{x-1}e^{-t} = o(t^{-2})$ comme $t \mapsto t^{-2}$ est intégrable sur $[1, +\infty[$ la fonction $t \mapsto f(x, t)$ est aussi intégrable sur $[1, +\infty[$.
- Pour tout $x \in \mathbb{R}$, $t^{x-1}e^{-t} \underset{t \rightarrow 0}{\sim} t^{x-1}$. Or on sait que t^α est intégrable sur $]0, 1]$ si et seulement si $\alpha > -1$. On en déduit que $t \mapsto f(x, t)$ est intégrable sur $]0, 1]$ si et seulement si $x - 1 > -1 \iff x > 0$.

La fonction Γ est définie sur $\mathbb{R}_+^*$.

- Continuité : Montrons que Γ est continue sur $\mathbb{R}_+^*$.

- Pour $x \in \mathbb{R}_+^*$, $t \mapsto t^{x-1}e^{-t}$ est continue par morceaux sur $\mathbb{R}_+^*$.
- Pour $t \in \mathbb{R}_+^*$, $x \mapsto t^{x-1}e^{-t} = \exp((x-1)\ln t)e^{-t}$ est continue sur $\mathbb{R}_+^*$.
- Hypothèse locale de domination : Pour tout segment $[a, b]$ de $]0, +\infty[$, on considère

$$\varphi : t \mapsto \begin{cases} t^{a-1}e^{-t} & \text{si } t < 1 \\ t^{b-1}e^{-t} & \text{si } t \geq 1 \end{cases}$$

On en déduit que pour $x \in [a, b]$ et $t \in]0, +\infty[$,

$$|f(x, t)| \leq \varphi(t)$$

car $x \mapsto t^{x-1}$ est décroissante quand $t < 1$ et croissante quand $t \geq 1$. De plus, φ est continue par morceaux intégrable sur $]0, +\infty[$ (raisonnement similaire à celui fait pour la définition de Γ).

La fonction Γ est continue

- Equation fonctionnelle : Il est clair que $\Gamma(1) = \int_0^{+\infty} e^{-t} dt = 1$. Maintenant pour $x > 0$, via une intégration par parties :

$$\Gamma(x) = \int_0^{+\infty} t^{x-1}e^{-t} dt = \left[\frac{t^x}{x} e^{-t} \right]_0^{+\infty} + \frac{1}{x} \int_0^{+\infty} t^x e^{-t} dt = \frac{1}{x} \Gamma(x+1)$$

car $\lim_{t \rightarrow 0} \frac{t^x}{x} e^{-t} = 0$ et $\lim_{t \rightarrow +\infty} \frac{t^x}{x} e^{-t} = 0$. On a donc

$$\Gamma(x+1) = x\Gamma(x).$$

En particulier,

$$\Gamma(2) = 1; \Gamma(3) = 2; \Gamma(4) = 3 \times 2 = 6.$$

et

$$\forall n \in \mathbb{N}^*, \Gamma(n) = (n-1)!$$

On voit donc que la fonction Γ généralise les factorielles.

Exercice : Exprimer les coefficients du développement en série entière de $(1+x)^\alpha$ en fonction de la fonction Γ .

Exercice : [1280] Soit $F : x \mapsto \int_0^{+\infty} \frac{e^{-xt}}{\sqrt{t+1}} dt$. Étudier l'ensemble de définition et la continuité de F .

1.2 Dérivabilité

On suppose maintenant que f est définie sur un produit de deux intervalles $J \times I$. On se pose alors la question de la dérivabilité de F

Théorème 17.5 (Théorème de dérivation des intégrales à paramètres)

Soit f une fonction à valeurs dans $\mathbf{K} = \mathbf{R}$ ou $\mathbf{C}$ définie sur $J \times I$ où I et J sont des intervalles de $\mathbf{R}$.

On suppose que

- i) Pour tout x de J , la fonction $f(x, \cdot) : t \mapsto f(x, t)$ est continue par morceaux
- ii) Pour tout x de J , la fonction $f(x, \cdot) : t \mapsto f(x, t)$ est intégrable sur I
- iii) Pour tout $t \in I$, $f(\cdot, t) : x \mapsto f(x, t)$ est dérivable, c'est-à-dire que f admet une dérivée partielle par rapport à sa première variable.
- iv) La fonction $(x, t) \mapsto \frac{\partial f}{\partial x}$ vérifie les hypothèses du théorème de continuité à savoir :
 - Pour tout x de J , la fonction $\frac{\partial f}{\partial x}(x, \cdot) : t \mapsto \frac{\partial f}{\partial x}(x, t)$ est continue par morceaux
 - Pour tout $t \in I$, $\frac{\partial f}{\partial x}(\cdot, t) : x \mapsto \frac{\partial f}{\partial x}(x, t)$ est continue sur J (c'est-à-dire que $x \mapsto f(x, t)$ est de classe $\mathcal{C}^1$).
 - Hypothèse de domination : pour tout segment K de J il existe une fonction $\varphi_K : I \rightarrow \mathbf{R}_+$ continue par morceaux et **intégrable** sur I telle que

$$\forall x \in K, \forall t \in I, \left| \frac{\partial f}{\partial x}(x, t) \right| \leq \varphi_K(t)$$

Alors, pour tout $x \in J$, $t \mapsto \frac{\partial f}{\partial x}(x, t)$ est intégrable sur I , la fonction $F : x \mapsto \int_I f(x, t) dt$ est dérivable et

$$\forall x \in J, F'(x) = \int_I \frac{\partial f}{\partial x}(x, t) dt \quad (\text{Formule de Leibniz})$$

De plus F est de classe $\mathcal{C}^1$.

Remarques :

1. L'hypothèse de domination portant cette fois sur la dérivée, il faut vérifier que la fonction $f(x, \cdot)$ est bien intégrable. De fait, les deux premières conditions sont juste les conditions nécessaires pour que F soit définie.
2. Il est important de faire **clairement** apparaître les hypothèses quand on utilise ce théorème.
3. L'hypothèse principale est l'hypothèse de domination. On voit en particulier que l'on peut se restreindre à montrer que la fonction $\frac{\partial f}{\partial x}$ est dominée sur tout segment par une fonction intégrable.

Exemple : Appliquons ce théorème à la fonction Gamma.

- On a déjà vu que pour $x \in]0, +\infty[$, $t \mapsto t^{x-1}e^{-t}$ est continue et intégrable sur $I =]0, +\infty[$.
- Soit $t \in I$, la fonction $x \mapsto t^{x-1}e^{-t}$ est dérivable et

$$\forall x \in]0, +\infty[, \frac{\partial f}{\partial x}(x, t) = (\ln t)t^{x-1}e^{-t}$$

- Soit $x \in]0, +\infty[$, la fonction $t \mapsto (\ln t)t^{x-1}$ est continue par morceaux.
- Soit $t \in]0, +\infty[$ la fonction $x \mapsto (\ln t)t^{x-1}$ est continue.

— Hypothèse locale de domination : Pour tout segment $[a, b]$ de $]0, +\infty[$, on considère

$$\varphi : t \mapsto \begin{cases} (|\ln t|)t^{a-1}e^{-t} & \text{si } t < 1 \\ (\ln t)t^{b-1}e^{-t} & \text{si } t \geq 1 \end{cases}$$

On en déduit que pour $x \in [a, b]$ et $t \in]0, +\infty[$,

$$\left| \frac{\partial f}{\partial x}(x, t) \right| \leq \varphi(t)$$

car $x \mapsto t^{x-1}$ est décroissante quand $t < 1$ et croissante quand $t \geq 1$. De plus, φ est continue par morceaux intégrable sur $]0, +\infty[$ (raisonnement similaire à celui fait pour la définition de Γ mais à détailler à cause du $\ln t$ - voir plus bas).

On en déduit que pour tout $x \in]0, +\infty[$, $t \mapsto (\ln t)t^{x-1}e^{-t}$ est intégrable et

$$F'(x) = \int_0^{+\infty} (\ln t)t^{x-1}e^{-t}dt.$$

Démonstration : Soit $x_0 \in J$ fixé.

On remarque d'abord, que l'hypothèse de domination sur $\frac{\partial f}{\partial x}$ permet d'assurer que $t \mapsto \frac{\partial f}{\partial x}(x_0, t)$ est intégrable sur I .

Maintenant pour tout x dans J différent de x_0

$$\frac{F(x) - F(x_0)}{x - x_0} = \int_I \frac{f(x, t) - f(x_0, t)}{x - x_0} dt$$

On pose $\theta(x, t) = \frac{f(x, t) - f(x_0, t)}{x - x_0}$. Cette fonction est définie sur $(J \setminus \{x_0\}) \times I$. On veut alors appliquer la version continue du théorème de convergence dominée en x_0 qui est un point adhérent à $V_0 = J \setminus \{x_0\}$.

i) Pour tout $x \in V_0$, la fonction $\theta(x, \cdot) : t \mapsto \theta(x, t)$ est continue par morceaux.

ii) Il existe une fonction $t \mapsto \frac{\partial f}{\partial x}(x_0, t)$ tel que

$$\forall t \in I, \lim_{x \rightarrow x_0} \theta(x, t) = \lim_{x \rightarrow x_0} \frac{f(x, t) - f(x_0, t)}{x - x_0} = \frac{\partial f}{\partial x}(x_0, t)$$

iii) Hypothèse de domination : on a supposé que pour x_0 de J , il existe un voisinage V de x_0 dans J tel qu'il existe une fonction $\varphi : I \rightarrow \mathbf{R}_+$ continue par morceaux et **intégrable** sur I telle que

$$\forall x \in V, \forall t \in I, \left| \frac{\partial f}{\partial x}(x, t) \right| \leq \varphi(t)$$

En particulier, en utilisant l'inégalité des accroissements finis (à t fixé),

$$\forall x \in V \setminus \{x_0\}, \forall t \in I, |\theta(x, t)| = \left| \frac{f(x, t) - f(x_0, t)}{x - x_0} \right| \leq \varphi(t).$$

On en déduit que

$$F'(x_0) = \lim_{x \rightarrow x_0} \frac{F(x) - F(x_0)}{x - x_0} = \lim_{x \rightarrow x_0} \int_I \theta(x, t) dt = \int_I \frac{\partial f}{\partial x}(x_0, t) dt.$$

Il ne reste plus qu'à appliquer le théorème de continuité des intégrales à paramètres pour obtenir que F est de classe $\mathcal{C}^1$.

□

Exercice : [1280] On reprend $F : x \mapsto \int_0^{+\infty} \frac{e^{-xt}}{\sqrt{t+1}} dt$. Est-elle de classe $\mathcal{C}^1$ sur $\mathcal{D}_F$?

On peut en déduire des théorèmes qui permettent de montrer qu'une intégrale à paramètre est de classe $\mathcal{C}^k$ ou $\mathcal{C}^\infty$.

Théorème 17.6 (Caractère $\mathcal{C}^k$ d'une intégrale à paramètre)

Soit f une fonction à valeurs dans $\mathbf{K} = \mathbf{R}$ ou $\mathbf{C}$ définie sur $J \times I$ où I et J sont des intervalles de $\mathbf{R}$.

Soit $k \in \mathbf{N}$ On suppose que

- i) Pour tout $t \in I$, $f(., t) : x \mapsto f(x, t)$ est de classe $\mathcal{C}^k$ sur J
- ii) Pour tout x de J et tout $i \in \llbracket 0 ; k \rrbracket$ la fonction $t \mapsto \frac{\partial^i f}{\partial x^i}(x, t)$ est continue par morceaux
- iii) Pour tout x de J et tout $i \in \llbracket 0 ; k - 1 \rrbracket$ la fonction $t \mapsto \frac{\partial^i f}{\partial x^i}(x, t)$ est intégrable sur I
- iv) Hypothèse de domination : pour tout segment K de J il existe une fonction $\varphi_K : I \rightarrow \mathbf{R}_+$ continue par morceaux et **intégrable** sur I telle que

$$\forall x \in K, \forall t \in I, \left| \frac{\partial^k f}{\partial x^k}(x, t) \right| \leq \varphi_K(t)$$

Alors, pour tout $x \in J$, $t \mapsto \frac{\partial^k f}{\partial x^k}(x, t)$ est intégrable sur I , la fonction $F : x \mapsto \int_I f(x, t) dt$ est de classe $\mathcal{C}^k$ et

$$\forall x \in J, \forall i \in \llbracket 0 ; k \rrbracket F^{(i)}(x) = \int_I \frac{\partial^i f}{\partial x^i}(x, t) dt$$

Théorème 17.7 (Caractère $\mathcal{C}^\infty$ d'une intégrale à paramètre)

Soit f une fonction à valeurs dans $\mathbf{K} = \mathbf{R}$ ou $\mathbf{C}$ définie sur $J \times I$ où I et J sont des intervalles de $\mathbf{R}$.

Soit $k \in \mathbf{N}$ On suppose que

- i) Pour tout $t \in I$, $f(., t) : x \mapsto f(x, t)$ est de classe $\mathcal{C}^\infty$ sur J
- ii) Pour tout x de J et tout $i \in \mathbf{N}$ la fonction $t \mapsto \frac{\partial^i f}{\partial x^i}(x, t)$ est continue par morceaux
- iii) Hypothèse de domination : Pour tout $i \in \mathbf{N}$, pour tout segment K de J il existe une fonction $\varphi_{i,K} : I \rightarrow \mathbf{R}_+$ continue par morceaux et **intégrable** sur I telle que

$$\forall x \in K, \forall t \in I, \left| \frac{\partial^i f}{\partial x^i}(x, t) \right| \leq \varphi_{i,K}(t)$$

Alors, pour tout $x \in J$, pour tout $i \in \mathbf{N}$, $t \mapsto \frac{\partial^i f}{\partial x^i}(x, t)$ est intégrable sur I , la fonction $F : x \mapsto \int_I f(x, t) dt$ est de classe $\mathcal{C}^\infty$ et

$$\forall i \in \mathbf{N}, \forall x \in J, F^{(i)}(x) = \int_I \frac{\partial^i f}{\partial x^i}(x, t) dt$$

Exemple : Quand on reprend encore la fonction Gamma.

On a bien que pour tout $k \in \mathbf{N}$, la fonction $x \mapsto t^{x-1}e^{-t}$ est de classe $\mathcal{C}^k$ et

$$\forall x \in]0, +\infty[, \forall t \in]0, +\infty[, \frac{\partial^k f}{\partial x^k}(x, t) = (\ln t)^k t^{x-1} e^{-t}.$$

De ce fait, sur tout segment $[a, b] \subset]0, +\infty[$ et tout k on pose

$$\varphi_k : t \mapsto \begin{cases} (|\ln t|)^k t^{a-1} e^{-t} & \text{si } t < 1 \\ (\ln t)^k t^{b-1} e^{-t} & \text{si } t \geq 1 \end{cases}$$

Vérifions que φ_k est bien intégrable sur $]0, +\infty[$

- Sur $]0, 1]$, on a $(|\ln t|)^k t^{a-1} e^{-t} \underset{t \rightarrow 0}{\sim} |\ln t|^k t^{a-1}$. Comme $a > 0$, il existe α tel que $-1 < \alpha < a - 1$.

On en déduit que

$$|\ln t|^k t^{a-1} = o(t^\alpha)$$

Comme $t \mapsto t^\alpha$ est intégrable sur $]0, 1]$, la fonction φ_k aussi.

- Sur $[1, +\infty[$. On a $(\ln t)^k t^{b-1} e^{-t} = o(t^{-2})$ donc φ_k est intégrable sur $[1, +\infty[$.

On en déduit que la fonction Γ est de classe $\mathcal{C}^\infty$ et que pour tout $k \in \mathbf{N}$,

$$\forall x \in]0, +\infty[, \Gamma^{(k)}(x) = \int_0^{+\infty} (\ln t)^k t^{x-1} e^{-t} dt.$$

On au passage on peut en déduire que Γ est convexe. En voyant que $\Gamma(1) = \Gamma(2) = 1$, le théorème de Rolle permet de justifier qu'il existe $c \in]1, 2[$ tel que $\Gamma'(c) = 0$. On a donc le tableau de variations :

faire le tableau

Remarquons aussi qu'en utilisant que $\Gamma(x+1) = x\Gamma(x)$ et en faisant tendre x vers 0 on obtient que

$$\Gamma(x) \underset{x \rightarrow 0}{\sim} \frac{1}{x}$$

Comme on a aussi que $\forall x \geq n \geq 2, \Gamma(x) \geq \Gamma(n) = (n-1)!$ on obtient $\lim_{x \rightarrow \infty} \Gamma(x) = +\infty$.

2 Exemples

Transformée de Laplace (CCP MP 2011) Transformée de Fourier (Centrale MP 2012)

Calcul différentiel

1	Différentielle et dérivées partielles	461
1.1	Dérivées selon un vecteur	461
1.2	Différentielle	463
1.3	Expression dans une base	467
1.4	Matrice jacobienne	468
1.5	Gradient	470
1.6	Applications de classe $\mathcal{C}^1$	472
2	Opérations sur les applications différentiables et les applications de classe $\mathcal{C}^1$	474
2.1	Combinaison linéaire et fonctions multilinéaires	474
2.2	Composition et règle de la chaîne	477
3	Dérivée le long d'un arc	482
3.1	Définition	482
3.2	Théorème fondamental de l'analyse	483
3.3	Vecteurs tangents à une partie	484
4	Fonctions de classe $\mathcal{C}^k$	488
4.1	Définitions	488
4.2	Opérations algébriques sur les fonctions de classe $\mathcal{C}^k$	491
5	Exemple d'équations aux dérivées partielles	493
5.1	Premier ordre	493
5.2	Second ordre	494
6	Optimisation	495
6.1	Point critique et extremum	495
6.2	Matrice hessienne	497
6.3	Optimisation sous contrainte	501

- Dans ce chapitre nous allons étudier le calcul différentiel. Toutes les applications considérées seront définies sur un ouvert U d'un $\mathbf{R}$ -espace vectoriel E de dimension finie (le plus souvent $\mathbf{R}^n$) et à valeurs dans un $\mathbf{R}$ -espace vectoriel de dimension finie F .
- Sauf mention explicite du contraire, tous les espaces vectoriels considérés seront de dimension finie.

- On a déjà étudié la notion de continuité de telles fonctions nous allons nous atteler à définir une notion de « dérivée ».

1 Différentielle et dérivées partielles

1.1 Dérivées selon un vecteur

Définition 18.1

Soit $f : U \rightarrow F$, soit $a \in U$ et $v \in E$ non nul.

1. On dit que f est dérivable en a selon le vecteur v si $\varphi : t \mapsto f(a + tv)$ est dérivable en $t = 0$.
2. Dans ce cas, on appelle dérivée de f en a **selon le vecteur** v et on note $D_v f(a)$ la valeur de la dérivée :

$$D_v f(a) = \lim_{t \rightarrow 0} \frac{f(a + tv) - f(a)}{t}$$

Définition 18.2

Si f est dérivable en tout point de U selon le vecteur v on dit que f est dérivable selon le vecteur v et on pose

$$\begin{aligned} D_v f : U &\rightarrow F \\ a &\mapsto D_v f(a) \end{aligned}$$

Exemples :

1. Soit $f : (x, y) \rightarrow x^2 + y^2$. Elle est dérivable selon $v = (\alpha, \beta)$ en tout point $a = (x, y)$. En effet,

$$\varphi : t \mapsto f(a + tv) = (x + t\alpha)^2 + (y + t\beta)^2 = x^2 + y^2 + 2t(x\alpha + y\beta) + t^2(\alpha + \beta)$$

est dérivable en 0. On a alors

$$D_v f(a) = \varphi'(0) = 2(x\alpha + y\beta).$$

2. Soit f de $\mathcal{M}_n(\mathbb{R})$ dans $\mathcal{M}_n(\mathbb{R})$ définie par $f(M) = \exp(M)$. Pour $a = A \in \mathcal{M}_n(\mathbb{R})$ et $v = I_n$, la fonction est dérivable en A selon le vecteur I_n . En effet

$$\varphi : t \mapsto f(A + tI_n) = \exp(A + tI_n)$$

Or comme A et tI_n commutent on a $\varphi(t) = \exp(A) \exp(tI_n)$.

Cette fonction est dérivable en 0. On a alors

$$D_{I_n}(A) = \varphi'(0) = \exp(A) \cdot I_n \cdot \exp(0I_n) = \exp(A)$$

Définition 18.3 (Dérivées partielles)

Soit $f : U \rightarrow F$. Soit $\mathcal{B} = (e_1, \dots, e_n)$ une base de E . Lorsqu'elles existent, les dérivées selon les vecteurs e_i s'appellent les dérivées partielles de f . On note alors

$$\frac{\partial f}{\partial x_i}(a) = D_{e_i} f(a) = \partial_i f(a).$$

Note

Dans la majorité des cas on se ramènera au cas où $E = \mathbf{R}^n$. En effet, en fixant une base $\mathcal{B} = (e_1, \dots, e_n)$ de E , on peut considérer l'application $\varphi : \mathbf{R}^n \rightarrow E$ définie par

$$\varphi : (x_1, \dots, x_n) \mapsto \sum_{i=1}^n x_i e_i$$

Cette application est linéaire donc continue. On pose alors

$$V = \varphi^{-1}(U) = \{(x_1, \dots, x_n) \in \mathbf{R}^n \mid x_1 e_1 + \dots + x_n e_n \in U\}$$

On voit que V est un ouvert comme image réciproque d'un ouvert par une application continue.

On considère alors $g = f \circ \varphi : V \rightarrow F$ définie par

$$g : (x_1, \dots, x_n) \mapsto f(x_1 e_1 + \dots + x_n e_n)$$

On voit alors (par exemple) que pour tout $i \in \llbracket 1 ; n \rrbracket$ et tout $a \in V$,

$$\frac{\partial g}{\partial x_i}(a) = \frac{\partial f}{\partial e_i}(\varphi(a))$$

Exemple : On considère l'application $f : (x, y) \mapsto x^2 + y^2$. On a alors

$$\frac{\partial f}{\partial x} : (x, y) \mapsto 2x \text{ et } \frac{\partial f}{\partial y} : (x, y) \mapsto 2y.$$

ATTENTION

Contrairement au cas des fonctions réelles, l'existence des dérivées partielles (et même selon tous vecteurs) n'implique pas la continuité. Par exemple la fonction

$$f : (x, y) \mapsto \begin{cases} \frac{y^2}{x} & \text{si } x \neq 0 \\ y & \text{sinon.} \end{cases}$$

Elle possède en $(0, 0)$ des dérivées partielles selon tout vecteur v mais n'est pas continue en $(0, 0)$.

— Si v est « vertical » c'est-à-dire de la forme $v = (0, \beta)$.

$$\varphi(t) = f(0 + tv) = f(0, t\beta) = t\beta \text{ qui est dérivable ; } D_v f(0) = \beta.$$

— Si v n'est pas « vertical », de la forme $v = (\alpha, \beta)$.

$$\varphi(t) = f(0 + tv) = f(t\alpha, t\beta) = \frac{t^2 \beta^2}{t\alpha} = \frac{t\beta^2}{\alpha} \text{ qui est dérivable ; } D_v f(0) = \frac{\beta^2}{\alpha}.$$

— La fonction f n'est pas continue en 0 car $f(0) = 0$ et $f(x^2, x) = \frac{x^2}{x^2} = 1$.

1.2 Différentielle

On veut généraliser la formule de Taylor-Young qui permet de définir la dérivée « classique » d'une fonction :

$$f(x_0 + h) = f(x_0) + hf'(x_0) + o(h)$$

Définition 18.4

On note $o(h)$ (ou $o(\|h\|)$) une fonction d'un voisinage de 0 dans E à valeurs dans F qui est « négligeable devant h ». C'est-à-dire de la forme

$$h \mapsto \|h\|\varepsilon(h)$$

où ε est une fonction tendant vers 0 quand h tend vers 0.

On peut de même définir $O(h)$ en imposant que la fonction ε soit bornée.

Remarque : Dans le cas classique d'une application vectorielle de $\mathbf{R}$ dans F , tous les termes de la formule de Taylor sont dans F . Le terme $hf'(x_0)$ appartient à F car $f(x_0) \in F$ et $h \in \mathbf{R}$ est un scalaire.

Maintenant on suppose que f est définie de $U \subset E$ dans F , on ne peut plus définir le terme $hf'(x_0)$ car $h \in E$ et $f'(x_0) \in F$.

Prenons un exemple d'une fonction de $\mathbf{R}^2$ dans $\mathbf{R}$. Soit $f : (x, y) \rightarrow x^2 + y^2$. On pose $a = (x_0, y_0)$. Pour $h = (h_1, h_2)$ on a

$$f(a + h) = (x_0 + h_1)^2 + (y_0 + h_2)^2 = \underbrace{x_0^2 + y_0^2}_{f(a)} + 2x_0h_1 + 2y_0h_2 + \underbrace{h_1^2 + h_2^2}_{o(h)}.$$

On pose alors $\varphi_a : \mathbf{R}^2 \rightarrow \mathbf{R}$ l'application **linéaire** définie par

$$\varphi_a : (h_1, h_2) \mapsto 2x_0h_1 + 2y_0h_2$$

On a finalement,

$$f(a + h) = f(a) + \varphi_a(h) + o(h)$$

Définition 18.5

Soit $f : U \rightarrow F$ et $a \in U$. L'application f est dite **différentiable** en a , s'il existe une application **linéaire** $\varphi_a : E \rightarrow F$ telle que

$$\forall h \in V, f(a + h) = f(a) + \varphi_a(h) + o(h)$$

où V est le voisinage de 0 défini par $V = \{h \in E \mid a + h \in U\}$.

Remarque : Si f est une fonction à valeurs réelles, la différentielle est une forme linéaire.

Exemples :

1. Soit f de $\mathcal{M}_n(\mathbf{R})$ dans $\mathcal{M}_n(\mathbf{R})$ définie par $f(M) = M^2$. Pour $a = A \in \mathcal{M}_n$ on a

$$f(A + H) = (A + H)^2 = A^2 + AH + HA + H^2$$

On pose $\varphi_A : H \mapsto AH + HA$ qui est linéaire. On a bien f qui est différentiable.

2. Soit E un espace euclidien et $f : u \mapsto \|u\|^2 = (u|u)$. Elle est différentiable en tout $a \in E$. En effet

$$f(a+h) = (a+h|a+h) = (a|a) + 2(a|h) + (h|h).$$

On pose

$$\varphi_a : h \mapsto 2(a|h).$$

Proposition 18.6 (Unicité de la différentielle)

Si f est différentiable en a . Il existe une unique application linéaire φ_a telle que

$$\forall h \in V, f(a+h) = f(a) + \varphi_a(h) + o(h)$$

Démonstration : Supposons qu'il existe deux applications linéaires φ_a et ψ_a . On écrit alors que pour tout $h \in V$,

$$f(a+h) = f(a) + \varphi_a(h) + \|h\|\varepsilon_1(h) = f(a) + \psi_a(h) + \|h\|\varepsilon_2(h)$$

où ε_1 et ε_2 tendent vers 0 quand $h \rightarrow 0$. On a alors

$$(\psi_a - \varphi_a)(h) = \|h\|(\varepsilon_1(h) - \varepsilon_2(h))$$

On fixe le vecteur h et on applique cela à th (en faisant tendre t vers 0). On a

$$(\psi_a - \varphi_a)(th) = \|th\|(\varepsilon_1(th) - \varepsilon_2(th))$$

Comme $(\psi_a - \varphi_a)$ est linéaire cela donne

$$(\psi_a - \varphi_a)(h) = \|h\|(\varepsilon_1(th) - \varepsilon_2(th)) \xrightarrow{t \rightarrow 0} 0.$$

On a bien $\psi_a = \varphi_a$. □

Définition 18.7

Avec les notations précédentes, l'application φ_a (qui est unique) s'appelle la différentielle de f en a et se note $df(a)$. Pour tout vecteur h on écrit alors $df(a).h$ pour $\varphi_a(h)$.

Note

Pour essayer de montrer qu'une fonction f est différentiable en a à la main, la méthode consiste à expliciter $f(a+h)$ et à regrouper les termes en trois catégories :

- Les termes qui ne dépendent pas de h ; ce sont ceux qui vont faire $f(a)$
- Les termes qui dépendent linéairement de h ; ce sont ceux qui vont faire la différentielle
- Les autres dont il faut montrer qu'ils sont $o(h)$

Exercice : Vérifier que $M \mapsto M^3$ est différentiable. On vérifiera soigneusement que « le $o(h)$ est bien négligeable devant h ».

Généraliser à $M \mapsto M^k$.

Proposition 18.8

Soit $f : U \rightarrow F$. Soit $\mathcal{F} = (\varepsilon_1, \dots, \varepsilon_p)$ une base de F . Pour tout $i \in \llbracket 1 ; p \rrbracket$ on note $f_i = \varepsilon_i^* \circ f$ de sorte que pour tout $x \in U$,

$$f(x) = f_1(x)\varepsilon_1 + \dots + f_p(x)\varepsilon_p$$

Soit $a \in U$, l'application f est différentiable en a si et seulement si pour tout $i \in \llbracket 1 ; p \rrbracket$, l'application f_i est différentiable en a .

Démonstration : Comme ci-dessus on note V le voisinage de 0 dans E défini par $h \in V \iff a + h \in U$.

- $\Rightarrow$ Comme f est différentiable en a , il existe $\varphi \in \mathcal{L}(E, F)$ et $\varepsilon : V \rightarrow F$ vérifiant que $\lim_{h \rightarrow 0_E} \varepsilon(h) = 0_F$ tels que

$$\forall h \in V, f(a + h) = f(a) + \varphi(h) + ||h||\varepsilon(h)$$

On peut projeter sur chaque coordonnée. Pour tout $i \in \llbracket 1 ; p \rrbracket$,

$$\forall h \in V, f_i(a + h) = f_i(a) + \varepsilon_i^* \circ \varphi(h) + ||h||\varepsilon_i^* \circ \varepsilon(h)$$

Cela montre que f_i est bien différentiable en a car $\varepsilon_i^* \circ \varphi$ est linéaire et $\lim_{h \rightarrow 0_E} \varepsilon_i^* \circ \varepsilon(h) = 0_{\mathbb{R}}$.

- $\Leftarrow$ On suppose que pour tout $i \in \llbracket 1 ; p \rrbracket$, f_i est différentiable en a , il existe $\varphi_i \in \mathcal{L}(E, \mathbb{R})$ et $\alpha_i : V \rightarrow \mathbb{K}$ vérifiant que $\lim_{h \rightarrow 0_E} \alpha_i(h) = 0_{\mathbb{R}}$ tels que

$$\forall h \in V, f_i(a + h) = f_i(a) + \varphi_i(h) + ||h||\alpha_i(h)$$

On en déduit que pour tout $h \in V$,

$$f(a + h) = \sum_{i=1}^p f_i(a + h)\varepsilon_i = f(a) + \Phi(h) + ||h||\alpha(h)$$

où $\Phi(h) = \sum_{i=1}^p \varphi_i(h)\varepsilon_i$ et $\alpha(h) = \sum_{i=1}^p \alpha_i(h)\varepsilon_i$. Comme $\Phi \in \mathcal{L}(E, F)$ et que α tend vers 0_F quand h tend vers 0_E , on obtient bien que f est différentiable en a . □

Proposition 18.9

Si f est différentiable en a alors f est continue en a .

Démonstration : Par définition, $f(a + h) = f(a) + \underbrace{\varphi_a(h) + o(h)}_{\xrightarrow{h \rightarrow 0} 0}$. □

Note

On va voir qu'il y a un lien entre l'existence des dérivées partielles et le fait d'être différentiable. On voit déjà que le fait d'être différentiable est « plus fort » car cela implique la continuité contrairement à l'existence des dérivées partielles.

Théorème 18.10

Si f est différentiable en a , alors pour tout vecteur v non nul, elle admet une dérivée selon v en a et

$$D_v f(a) = df(a).v$$

Démonstration : On suppose que f est différentiable en a . On en déduit que

$$f(a + tv) = f(a) + df(a).(tv) + \|tv\|\varepsilon(tv)$$

De ce fait

$$\frac{f(a + tv) - f(a)}{t} = \frac{1}{t} (df(a).(tv) + \|tv\|\varepsilon(tv)) = df(a).v + \|v\|\varepsilon(tv) \xrightarrow{t \rightarrow 0} df(a).v.$$

On en déduit que f admet une dérivée selon v en a et que

$$D_v f(a) = df(a).v$$

□

Définition 18.11

Si $f : U \rightarrow F$ est différentiable en tout point a de U elle est dite différentiable. On appelle alors différentielle de f et on note df la fonction

$$\begin{aligned} df : U &\rightarrow \mathcal{L}(E, F) \\ a &\mapsto df(a) \end{aligned}$$

Exemples :

1. Soit f une fonction constante. Elle est différentiable et df est la fonction nulle.
2. Soit f une fonction linéaire de U sur F . Elle est différentiable et df est la fonction constante égale à f . En effet, pour tout $a \in U$ et tout $h \in V$

$$f(a + h) = f(a) + f(h)$$

On peut donc poser $df(a) = f$.

3. Dans le cas où $E = \mathbf{R}$, f est dérivable si et seulement si elle est différentiable et

$$\begin{aligned} df : U &\rightarrow \mathcal{L}(\mathbf{R}, F) \\ a &\mapsto (x \mapsto xf'(a)) \end{aligned}$$

car

$$\forall h \in V, f(a + h) = f(a) + hf'(a) + o(h).$$

En particulier, $f'(a) = df_a.1$.

4. Si on dispose d'une base $(\varepsilon_1, \dots, \varepsilon_p)$ de F et que l'on note $f_i = \varepsilon_i^* \circ f$. Dans le cas où f est différentiable,

$$df = \sum_{i=1}^p df_i \varepsilon_i$$

5. Soit $B : E \times F \rightarrow G$ une application bilinéaire. Elle est différentiable. En effet pour tout $a = (x, y) \in F \times G$ et tout $(h_1, h_2) \in F \times G$,

$$B(x + h_1, y + h_2) = B(x, y) + B(h_1, y) + B(x, h_2) + B(h_1, h_2)$$

On sait qu'il existe une constante $C > 0$ telle que

$$\|B(h_1, h_2)\|_G \leq C\|h_1\|_E\|h_2\|_F \leq C\|(h_1, h_2)\|_{E \times F}^2$$

Cela montre que

$$B(x + h_1, y + h_2) = B(x, y) + B(h_1, y) + B(x, h_2) + o(\|(h_1, h_2)\|_{E \times F})$$

Finalement B est bien différentiable en $a = (x, y)$ et

$$dB(x, y) : (h_1, h_2) \mapsto B(h_1, y) + B(x, h_2)$$

1.3 Expression dans une base

Proposition 18.12

On suppose que $\mathcal{B} = (e_1, \dots, e_n)$ est une base de E et que $f : U \rightarrow F$ est différentiable. Soit $a \in U$.

Pour tout vecteur $h = \sum_{i=1}^n h_i e_i$,

$$df(a).h = \sum_{i=1}^n \frac{\partial f}{\partial x_i}(a) h_i$$

C'est-à-dire que

$$df(a) = \sum_{i=1}^n \frac{\partial f}{\partial x_i}(a) e_i^*$$

Exemple : Si on reprend $f : (x, y) \mapsto x^2 + y^2$. On a vu que $\frac{\partial f}{\partial x}(x, y) = 2x$ et $\frac{\partial f}{\partial y}(x, y) = 2y$. On en déduit que pour $a = (x, y)$,

$$df(a) : (h_1, h_2) \mapsto 2xh_1 + 2yh_2$$

En particulier, si $v = (\alpha, \beta)$ la dérivée selon v en a est

$$D_v f(a) = df(a).v = 2x\alpha + 2y\beta$$

Démonstration : On se fixe $a \in U$. On sait que $df(a)$ est une application linéaire. Elle est donc uniquement déterminée par son action sur une base. On regarde sur la base $\mathcal{B} = (e_1, \dots, e_n)$. On a

$$df(a).e_i = D_{e_i} f(a) = \frac{\partial f}{\partial x_i}(a).$$

Par linéarité on obtient bien que

$$df(a).h = \sum_{i=1}^n \frac{\partial f}{\partial x_i}(a) h_i.$$

□

Remarque : Il arrive que l'on note dx_i l'application constante égale à e_i^* définie sur U . En effet, c'est la différentielle de la fonction $e_i^* : (x_1, \dots, x_n) \mapsto x_i$. Avec cette notation on obtient :

$$df = \frac{\partial f}{\partial x_1} dx_1 + \dots + \frac{\partial f}{\partial x_n} dx_n = \sum_{i=1}^n \frac{\partial f}{\partial x_i} dx_i.$$

ATTENTION

Il se peut que les dérivées partielles existent mais que la fonction ne soit pas différentiable !
C'est le cas de la fonction

$$f : (x, y) \mapsto \begin{cases} \frac{y^2}{x} & \text{si } x \neq 0 \\ y & \text{sinon.} \end{cases}$$

Elle n'est pas différentiable car elle n'est pas continue.

Exercice : Calculer les dérivées partielles de $f : M \mapsto M^2$ en $A \in \mathcal{M}_n(\mathbb{R})$. Retrouver la formule de df_A .

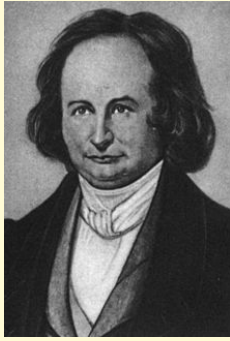
1.4 Matrice jacobienne

Définition 18.13 (Matrice Jacobienne)

Soit f une application différentiable de U dans F . Soit $\mathcal{B} = (e_1, \dots, e_n)$ une base de E et $\mathcal{F} = (\varepsilon_1, \dots, \varepsilon_p)$ une base de F . Pour tout $a \in U$, on appelle matrice jacobienne de f en a relative aux bases $\mathcal{B}$ et $\mathcal{F}$ la matrice de $df(a)$. On note

$$J_f(a) = \text{Mat}_{\mathcal{B}, \mathcal{F}}(df(a)) = \begin{pmatrix} \frac{\partial f_1}{\partial x_1}(a) & \dots & \frac{\partial f_1}{\partial x_n}(a) \\ \vdots & & \vdots \\ \frac{\partial f_p}{\partial x_1}(a) & \dots & \frac{\partial f_p}{\partial x_n}(a) \end{pmatrix}.$$

où pour tout $j \in \llbracket 1; p \rrbracket$, $f_j = \varepsilon_j^* \circ f$



Charles Gustave Jacob Jacobi né le 10 décembre 1804 et mort le 18 février 1851, est un mathématicien allemand surtout connu pour ses travaux sur les intégrales elliptiques, les équations aux dérivées partielles et leur application à la mécanique analytique. Jacobi est le premier mathématicien à appliquer les fonctions elliptiques à la théorie des nombres, prouvant par exemple le théorème des nombres polygonaux annoncé sans preuve par Fermat. Il donne de nouvelles preuves de la loi de réciprocité quadratique, et y apporte des généralisations; pour ce faire, il introduit ce qui aujourd'hui est connu sous le nom de sommes de Jacobi. La fonction thêta de Jacobi, si fréquemment appliquée dans l'étude des séries hypergéométriques, porte son nom. Il en a donné l'équation fonctionnelle.

Il est l'un des fondateurs de la théorie des déterminants. En particulier, il invente le déterminant de la matrice (dite jacobienne) formée par les n^2 dérivées partielles de n fonctions données de n variables indépendantes. Son déterminant, le déterminant jacobien est crucial dans le calcul infinitésimal.

Remarques :

1. La matrice jacobienne dépend des bases mais on ne les précise pas pour alléger.
2. On a vu que les dérivées partielles peuvent exister sans que la différentielle existe. Dans ce cas, la matrice jacobienne peut être définie mais ce n'est pas la matrice d'une application linéaire qui apparaît naturellement.
3. L'application $df(a)$ est une application de E dans F . Le nombre de lignes de la matrice jacobienne est la dimension de F , le nombre de colonnes de la matrice jacobienne est la dimension de E .
4. Si f est une fonction à valeurs réelles, la matrice jacobienne est la matrice d'une forme linéaire, c'est donc une matrice ligne.

$$J_f = \left(\frac{\partial f}{\partial x_1} \cdots \frac{\partial f}{\partial x_n} \right).$$

Exemples :

1. On considère l'application $f : (x, y) \mapsto (xe^y, \sin(x + y))$. On a alors

$$\frac{\partial f_1}{\partial x} = e^y; \frac{\partial f_1}{\partial y} = xe^y$$

et

$$\frac{\partial f_2}{\partial x} = \cos(x + y); \frac{\partial f_2}{\partial y} = \cos(x + y)$$

On a donc

$$J_f(x, y) = \begin{pmatrix} e^y & xe^y \\ \cos(x + y) & \cos(x + y) \end{pmatrix}.$$

On peut donc en déduire que si $v = (2, 1)$,

$$D_v f(x, y) = J_f(x, y) \begin{pmatrix} 2 \\ 1 \end{pmatrix} = \begin{pmatrix} (2 + x)e^y \\ 3 \cos(x + y) \end{pmatrix}$$

2. Dans le cas du changement de variable en polaire, $f : (\rho, \theta) \mapsto (\rho \cos \theta, \rho \sin \theta)$. On obtient

$$J_f(\rho, \theta) = \begin{pmatrix} \cos \theta & -\rho \sin \theta \\ \sin \theta & \rho \cos \theta \end{pmatrix}.$$

En particulier le déterminant de la matrice jacobienne (que l'on appelle le jacobien) vaut

$$|J_f(x, y)| = \rho.$$

Notons que cela justifie le calcul en physique lors d'un changement de variables en polaires :

$$dx dy = \rho d\rho d\theta$$

Exercice : On considère l'application du changement de variables sphérique

$$f : (\rho, \theta, \phi) \mapsto (\rho \sin \theta \cos \phi, \rho \sin \theta \sin \phi, \rho \cos \theta)$$

Calculer la jacobienne de f .

Correction

On trouve

$$J_f(\rho, \theta, \phi) = \begin{pmatrix} \sin \theta \cos \phi & \rho \cos \theta \cos \phi & -\rho \sin \theta \sin \phi \\ \sin \theta \sin \phi & \rho \cos \theta \sin \phi & \rho \sin \theta \cos \phi \\ \cos \theta & -\rho \sin \theta & 0 \end{pmatrix}$$

1.5 Gradient

Dans cette section on considère toujours une fonction définie sur un ouvert U de E et à valeurs dans F . On suppose de plus que E est un espace euclidien et que $F = \mathbf{R}$. Dans ce cas, on peut « représenter » la différentielle de f en a par un vecteur de E .

Commençons par rappeler le théorème de représentation vu dans le chapitre sur les espaces préhilbertiens.

Théorème 18.14 (Théorème de représentation)

Soit E un espace euclidien. Soit f une forme linéaire sur E . Il existe un unique vecteur a tel que

$$\forall x \in E, (a|x) = f(x)$$

C'est-à-dire que l'application

$$\begin{aligned} E &\mapsto E^* \\ a &\mapsto (x \mapsto (a|x)) \end{aligned}$$

est un isomorphisme.

Définition 18.15

Soit f une fonction définie sur un ouvert U d'un espace euclidien E à valeur dans $\mathbf{R}$. Soit $a \in U$. On suppose que f est différentiable en a , on appelle alors gradient de f en a et on note $\nabla f(a)$ le vecteur de E tel que

$$\forall x \in E, (\nabla f(a)|x) = df(a).x$$

Remarques :

1. Le gradient se note aussi $\text{Grad}f(a)$. C'est le vecteur qui « représente » la forme linéaire $df(a)$.
2. Cela signifie géométriquement qu'au voisinage de a , la fonction « évolue » le plus quand on déplace le long du gradient. Elle augmente quand on va dans le sens du gradient et elle diminue quand on va dans l'autre sens. Précisément, si $\nabla f(a) \neq 0$, il est colinéaire au vecteur unitaire v tel que la dérivée $D_v f(a)$ soit maximale.

Proposition 18.16 (Expression du gradient dans une base orthonormée)

Soit $\mathcal{B} = (e_1, \dots, e_n)$ une base orthonormée de E . Avec les notations précédentes,

$$\text{Mat}_{\mathcal{B}}(\nabla f(a)) = \begin{pmatrix} \frac{\partial f}{\partial x_1}(a) \\ \vdots \\ \frac{\partial f}{\partial x_n}(a) \end{pmatrix}$$

Démonstration : On a vu que $df(a)$ était donnée par le fait que pour tout $i \in \llbracket 1 ; p \rrbracket$,

$$df(a).e_i = \frac{\partial f}{\partial x_i}(a)$$

Cela signifie que

$$df(a).(x_1 e_1 + \dots + x_n e_n) = \sum_{i=1}^n \frac{\partial f}{\partial x_i}(a) x_i$$

On en déduit que

$$\text{Mat}_{\mathcal{B}}(\nabla f(a)) = \begin{pmatrix} \frac{\partial f}{\partial x_1}(a) \\ \vdots \\ \frac{\partial f}{\partial x_n}(a) \end{pmatrix}$$

□

Exemple : On considère la fonction $f : (x, y) \mapsto x^2 + y^2$. On a alors

$$\nabla f(x, y) = \begin{pmatrix} 2x \\ 2y \end{pmatrix}$$

La fait que $\nabla f(a)$ soit colinéaire à a est conforme à l'intuition. Quand on avance dans le sens de a , l'augmentation de f est maximale.

Exercice : On considère sur $\mathbf{R}_+^* \times \mathbf{R}$ la fonction $(x, y) \mapsto \arctan(\frac{y}{x})$. Déterminer $\nabla f(a)$ et interpréter.

Correction

On calcule les dérivées partielles

$$\frac{\partial f}{\partial x}(x, y) = -\frac{y}{x} \frac{1}{1 + (\frac{y}{x})^2} = -\frac{y}{x^2 + y^2}$$

et

$$\frac{\partial f}{\partial y}(x, y) = \frac{1}{x} \frac{1}{1 + (\frac{y}{x})^2} = \frac{x}{x^2 + y^2}$$

On en déduit que $\nabla f(x, y)$ est colinéaire à $(-y, x)$ qui est le vecteur directement orthogonal à (x, y) .

1.6 Applications de classe $\mathcal{C}^1$

Définition 18.17

Soit f une fonction définie sur un ouvert U à valeurs dans un espace vectoriel F . Elle est dite de classe $\mathcal{C}^1$ si elle est différentiable et que la différentielle $df : U \rightarrow \mathcal{L}(E, F)$ est continue. On note $\mathcal{C}^1(U, F)$ l'ensemble des fonctions de classe $\mathcal{C}^1$ définies sur U et à valeurs dans F .

Proposition 18.18

Soit $\mathcal{B} = (e_1, \dots, e_n)$ une base de E et $f : U \rightarrow F$.

La fonction f est de classe $\mathcal{C}^1$ si et seulement si elle admet des dérivées partielles par rapport à la base $\mathcal{B}$ et que pour tout $i \in \llbracket 1; n \rrbracket$ la fonction $\frac{\partial f}{\partial e_i} : U \rightarrow F$ est continue.

Démonstration :

- $\Rightarrow$ On suppose qu'elle est de classe $\mathcal{C}^1$. Elle est donc différentiable et de ce fait les dérivées partielles existent. De plus on suppose que $df : a \mapsto df(a)$ est continue ce qui implique que pour tout $i \in \llbracket 1; n \rrbracket$, $\frac{\partial f}{\partial e_i} : a \mapsto \frac{\partial f}{\partial e_i}(a) = df(a).e_i$ est aussi continue. En effet c'est la composée de df qui est continue par hypothèses et de l'application d'évaluation

$$\begin{array}{ccc} \delta_i : \mathcal{L}(E, F) & \rightarrow & F \\ \varphi & \mapsto & \varphi(e_i) \end{array}$$

qui est continue car linéaire entre espaces vectoriels de dimension finie.

- $\Leftarrow$ (Non exigible) Afin de simplifier les notations, on suppose que $E = \mathbf{R}^n$ et que $\mathcal{B}$ est la base canonique ce qui permet d'écrire $f(x_1, \dots, x_n)$ plutôt que $f(x_1 e_1 + \dots + x_n e_n)$. De plus, en utilisant une base de F , on se ramène au cas des fonctions à valeurs dans $\mathbf{R}$.

Pour finir, on traite le cas où $n = 2$, le cas général est similaire mais avec des notations plus lourdes.

On considère donc $f : U \rightarrow \mathbf{R}$ où U est un ouvert de $\mathbf{R}^2$. On suppose que pour tout $i \in \llbracket 1; 2 \rrbracket$ la fonction $\partial_i f : U \rightarrow \mathbf{R}$ est continue. Soit $a = (a_1, a_2) \in U$, on veut montrer que f est différentiable en a . On pose V le voisinage de 0 défini par $h \in V$ si et seulement si $a + h \in U$. On sait que si f est différentiable alors pour $h = (h_1, h_2) \in V$,

$$df(a).h = h_1 \partial_1 f(a) + h_2 \partial_2 f(a)$$

On pose donc $\Phi : V \rightarrow \mathbf{R}$ définie par

$$\Phi : h \mapsto f(a+h) - f(a) - h_1 \partial_1 f(a) - h_2 \partial_2 f(a)$$

On doit montrer que $\Phi(h) = o(h)$. On remarque pour commencer que

$$f(a+h) - f(a) = [f(a_1+h_1, a_2+h_2) - f(a_1+h_1, a_2)] + [f(a_1+h_1, a_2) - f(a_1, a_2)]$$

En utilisant le théorème des accroissements finis pour la fonction $t \mapsto f(t, a_2)$. Il existe c_1 (qui dépend de h_1) compris entre a_1 et a_1+h_1 tel que $f(a_1+h_1, a_2) - f(a_1, a_2) = \partial_1 f(c_1, a_2)h_1$.

De même en utilisant le théorème des accroissements finis pour la fonction $t \mapsto f(a_1+h_1, t)$. Il existe c_2 (qui dépend de h_1 et de h_2) compris entre a_2 et a_2+h_2 tel que $f(a_1+h_1, a_2+h_2) - f(a_1+h_1, a_2) = \partial_2 f(a_2+h_2, c_2)h_2$.

On en déduit que

$$\Phi(h) = h_1(\partial_1 f(c_1, a_2) - \partial_1 f(a)) + h_2(\partial_2 f(a_2+h_2, c_2) - \partial_2 f(a))$$

Quand h tend vers 0, c_1 tend vers a_1 donc (c_1, a_2) tend vers a . De même c_2 tend vers a_2 donc (a_2+h_2, c_2) tend vers a . Comme $\partial_1 f$ et $\partial_2 f$ sont continues,

$$\lim_{h \rightarrow (0,0)} \partial_1 f(c_1, a_2) - \partial_1 f(a) = \lim_{h \rightarrow (0,0)} \partial_2 f(a_2+h_2, c_2) - \partial_2 f(a) = 0$$

Cela implique bien que $\Phi(h) = o(h)$ et donc que f est différentiable.

□

Remarque : C'est surtout la partie $\Leftarrow$ qui est intéressante. Cela permet de montrer qu'une application est différentiable uniquement avec les dérivées partielles alors que le résultat analogue sans la continuité est faux.

Exemples :

1. Les fonctions polynomiales de E à valeurs dans $\mathbf{R}$ sont de classe $\mathcal{C}^1$.

2. Les applications linéaires sont de classe $\mathcal{C}^1$.

3. Soit $f : (x, y) \mapsto \begin{cases} \frac{xy}{x^2+y^2} & \text{si } (x, y) \neq (0, 0) \\ 0 & \text{sinon.} \end{cases}$

— En dehors de $(0, 0)$ elle a des dérivées partielles continues :

$$\frac{\partial f}{\partial x}(x, y) = \frac{y(x^2 - y^2)}{(x^2 + y^2)^2} \text{ et } \frac{\partial f}{\partial y}(x, y) = \frac{x(y^2 - x^2)}{(x^2 + y^2)^2}.$$

— En $(0, 0)$, pour étudier $\frac{\partial f}{\partial x}$ on regarde $x \mapsto f(x, 0) = 0$. On en déduit que la dérivée partielle existe et que

$$\frac{\partial f}{\partial x}(0, 0) = 0$$

De même,

$$\frac{\partial f}{\partial y}(0, 0) = 0$$

- Cependant f (et donc ses dérivées partielles) ne sont pas continues en $(0, 0)$. On peut par exemple regarder $\lim_{x \rightarrow 0} f(x, x)$. La fonction n'est pas différentiable.

4. Soit $f : (x, y) \mapsto \begin{cases} \frac{x^2 y^2}{x^2 + y^2} & \text{si } (x, y) \neq (0, 0) \\ 0 & \text{sinon.} \end{cases}$

- En dehors de $(0, 0)$ elle a des dérivées partielles continues :

$$\frac{\partial f}{\partial x}(x, y) = \frac{2xy^4}{(x^2 + y^2)^2} \text{ et } \frac{\partial f}{\partial y}(x, y) = \frac{2yx^4}{(x^2 + y^2)^2}.$$

- En $(0, 0)$, pour étudier $\frac{\partial f}{\partial x}$ on regarde $x \mapsto f(x, 0) = 0$. On en déduit que la dérivée partielle existe et que

$$\frac{\partial f}{\partial x}(0, 0) = 0$$

De même,

$$\frac{\partial f}{\partial y}(0, 0) = 0$$

- Les dérivées partielles sont continues aussi en 0. En effet, comme $|x| \leq \sqrt{x^2 + y^2}$ et que $|y| \leq \sqrt{x^2 + y^2}$, on a

$$\left| \frac{2xy^4}{(x^2 + y^2)^2} \right| \leq 2 \frac{(x^2 + y^2)^{5/2}}{(x^2 + y^2)^2} \leq 2\sqrt{x^2 + y^2}.$$

On en déduit que f est de classe $\mathcal{C}^1$. En particulier, elle est différentiable.

2 Opérations sur les applications différentiables et les applications de classe $\mathcal{C}^1$

2.1 Combinaison linéaire et fonctions multilinéaires

Proposition 18.19 (Combinaisons linéaires)

1. Soit f_1 et f_2 deux fonctions différentiables définies sur U et à valeurs dans F en $a \in U$. Soit $(\lambda_1, \lambda_2) \in \mathbb{R}^2$ la fonction $g = \lambda_1 f_1 + \lambda_2 f_2$ est différentiable en a et

$$dg(a) = \lambda_1 df_1(a) + \lambda_2 df_2(a).$$

2. Une combinaison linéaire de deux fonctions différentiables sur U est différentiable sur U .
3. Une combinaison linéaire de deux fonctions de classe $\mathcal{C}^1$ sur U est de classe $\mathcal{C}^1$ sur U .

Démonstration :

1. Il suffit d'utiliser la définition. On sait que pour tout $h \in V$,

$$f_1(a + h) = f_1(a) + df_1(a).h + \|h\|_{\mathcal{E}_1}(h) \text{ et } f_2(a + h) = f_2(a) + df_2(a).h + \|h\|_{\mathcal{E}_2}(h)$$

où ε_1 et ε_2 tendent vers 0 quand $h \rightarrow 0$. On a donc

$$g(a+h) = g(a) + (\lambda_1 df_1(a) + \lambda_2 df_2(a)).h + \|h\| \underbrace{(\lambda_1 \varepsilon_1(h) + \lambda_2 \varepsilon_2(h))}_{\xrightarrow{h \rightarrow 0} 0}$$

On en déduit bien que g est différentiable en a et

$$dg(a) = \lambda_1 df_1(a) + \lambda_2 df_2(a)$$

2. Si on suppose que f_1 et f_2 sont différentiables sur U , on peut appliquer ce qui précède en tout point $a \in U$. On en déduit que $g = \lambda_1 f_1 + \lambda_2 f_2$ est différentiable sur U et que

$$dg = g = \lambda_1 df_1 + \lambda_2 df_2$$

3. D'après ce qui précède, si f_1 et f_2 sont de classe $\mathcal{C}^1$ alors df_1 et df_2 sont continues et donc dg aussi. Cela montre que $g \in \mathcal{C}^1(U, F)$. □

Proposition 18.20

Soit $f_1 : U \rightarrow F_1$ et $f_2 : U \rightarrow F_2$ et $B : F_1 \times F_2 \rightarrow G$ une application bilinéaire. On considère $g = B(f_1, f_2) : U \rightarrow G$ définie par $g : a \mapsto B(f_1(a), f_2(a))$.

1. Soit $a \in U$. On suppose que f_1 et f_2 sont différentiables en a . L'application g alors est différentiable en a et

$$dg(a) : h \mapsto B(df_1(a).h, f_2(a)) + B(f_1(a), df_2(a).h)$$

2. Si on suppose que f_1 et f_2 sont de classe $\mathcal{C}^1$ alors g est de classe $\mathcal{C}^1$.

Démonstration :

1. On sait que pour tout $h \in V$,

$$f_1(a+h) = f_1(a) + df_1(a).h + \|h\|\varepsilon_1(h) \text{ et } f_2(a+h) = f_2(a) + df_2(a).h + \|h\|\varepsilon_2(h)$$

où ε_1 et ε_2 tendent vers 0 quand $h \rightarrow 0$. On a alors

$$\begin{aligned} g(a+h) &= B(f_1(a+h), f_2(a+h)) \\ &= B(f_1(a) + df_1(a).h + \|h\|\varepsilon_1(h), f_2(a) + df_2(a).h + \|h\|\varepsilon_2(h)) \\ &= g(a) + B(df_1(a).h, f_2(a)) + B(f_1(a), df_2(a).h) + \alpha(h) \end{aligned}$$

où

$$\alpha(h) = B(f_1(a) + df_1(a).h, \|h\|\varepsilon_2(h)) + B(\|h\|\varepsilon_1(h), f_2(a) + df_2(a).h) + B(df_1(a).h, df_2(a).h) + B(\|h\|\varepsilon_1(h), \|h\|\varepsilon_2(h))$$

Comme $h \mapsto B(df_1(a).h, f_2(a)) + B(f_1(a), df_2(a).h)$ est linéaire, il reste à montrer que $\alpha(h) = o(h)$.

Commençons par remarquer que, comme B est une application bilinéaire sur des espaces vectoriels de dimension finie, il existe C tel que

$$\forall (x, y) \in F_1 \times F_2, \|B(x, y)\| \leq C \|x\| \|y\|$$

où on considère des normes sur F_1 , F_2 et G .

On en déduit que $B(f_1(a) + df_1(a).h, \|h\|_{\varepsilon_2}(h)) = \|h\|B(f_1(a) + df_1(a).h, \varepsilon_2(h))$ et $B(f_1(a) + df_1(a).h, \varepsilon_2(h)) \xrightarrow{h \rightarrow 0} 0$ car

$$\|B(f_1(a) + df_1(a).h, \varepsilon_2(h))\| \leq C\|f_1(a) + df_1(a).h\|\|\varepsilon_2(h)\|$$

Or $\|\varepsilon_2(h)\| \xrightarrow{h \rightarrow 0} 0$ et $\|f_1(a) + df_1(a).h\| \xrightarrow{h \rightarrow 0} \|f_1(a)\|$ donc est bornée car $df_1(a)$ est continue.

De même $B(\|h\|_{\varepsilon_1}(h), f_2(a) + df_2(a).h)$ est négligeable devant h ainsi que $B(\|h\|_{\varepsilon_1}(h), \|h\|_{\varepsilon_2}(h))$.

Il reste à considérer $B(df_1(a).h, df_2(a).h)$. Les applications $df_1(a)$ et $df_2(a)$ sont continues donc elles sont bornées sur la boule unité (qui est compacte car E est de dimension finie). On en déduit que si on note M_1 et M_2 des majorants,

$$\|B(df_1(a).h, df_2(a).h)\| \leq C\|h\|^2 M_1 M_2 = o(h).$$

Au final on a bien,

$$g(a + h) = g(a) + B(df_1(a).h, f_2(a)) + B(f_1(a), df_2(a).h) + o(h)$$

donc g est différentiable en a et

$$dg(a) : h \mapsto B(df_1(a).h, f_2(a)) + B(f_1(a), df_2(a).h)$$

2. On suppose que f_1 et f_2 sont de classe $\mathcal{C}^1$. En appliquant ce qui précède en tout point $a \in U$ on voit que g est différentiable sur U . Il reste donc à montrer que dg est continue. La fonction dg est définie de U dans $\mathcal{L}(E, G)$. On sait que $\mathcal{L}(E, G)$ est de dimension finie. Soit $(h_1, \dots, h_n)$ une base de E . On pose pour $i \in \llbracket 1; n \rrbracket$, $\theta_i : U \mapsto G$ la fonction définie par $\theta_i : a \mapsto dg(a).h_i$. On sait que pour montrer que dg est continue, il suffit de montrer que pour tout $i \in \llbracket 1; n \rrbracket$, θ_i est continue (cela se prouve par exemple en écrivant les matrices de θ_i et de dg). Or, pour tout $i \in \llbracket 1; n \rrbracket$, $a \mapsto df_1(a).h_i$, $a \mapsto f_2(a)$, $a \mapsto f_1(a)$ et $a \mapsto df_2(a).h_i$ sont continues. Cela montre bien que g est de classe $\mathcal{C}^1$.

□

Exemples :

1. Si B est le produit scalaire (pour $F_1 = F_2$ euclidien), on a par exemple dans le cas où $f_1 = f_2$, que si f est différentiable alors $g = \|f\|^2$ aussi et

$$dg(a) : h \mapsto 2(df(a).h|f(a)).$$

2. On peut prendre pour B le produit matriciel. On retrouve que la différentielle de $A \mapsto A^2$ en A est $H \mapsto AH + HA$.

Proposition 18.21 (Généralisation aux fonctions multilinéaires)

Si $f_1, \dots, f_p$ sont des fonctions de U dans F_i et si Φ est une fonction p -linéaire définie de $\prod_{i=1}^p F_i$ dans G . On pose $g : U \mapsto G$ définie par

$$g : a \mapsto \Phi(f_1(a), \dots, f_p(a))$$

Si pour tout $i \in \llbracket 1; p \rrbracket$, les fonctions f_i sont différentiables (resp. de classe $\mathcal{C}^1$) alors g aussi. De plus

$$\forall a \in U, \forall h \in E, dg(a).h = \sum_{i=1}^p \Phi(f_1(a), \dots, f_{i-1}(a), df_i(a).h, f_{i+1}(a), \dots, f_p(a))$$

2.2 Composition et règle de la chaîne

Proposition 18.22 (Composition)

Soit $f : U_1 \rightarrow F$ et $g : U_2 \rightarrow G$ où U_2 est un ouvert de F tel que $f(U_1) \subset U_2$. Soit $a \in U_1$ et $b = f(a)$. On suppose que f est différentiable en a et g différentiable en b alors $g \circ f$ est différentiable en a et

$$d(g \circ f)(a) = dg(f(a)) \circ (df)(a).$$

Remarque :

Cette formule généralise la formule usuelle $(g \circ f)'(a) = f'(a) \times g'(f(a))$ dans les cas des fonctions dérivables de $\mathbf{R}$ dans $\mathbf{R}$. Dans ce cadre, on a vu que pour tout a

$$df(a) : h \mapsto f'(a)h$$

De même pour $b = f(a)$,

$$dg(b) : h \mapsto g'(b)h = g'(f(a))h$$

Le terme de droite de l'égalité est donc

$$dg(f(a)) \circ df(a) : h \mapsto dg(f(a)).(df(a).(h)) = dg(f(a)).(f'(a)h) = (g'(f(a)) \times f'(a))h$$

Démonstration : On écrit les développements limités à l'ordre 1 de f et g . On sait que pour tout $h \in V_1$,

$$f(a+h) = f(a) + df(a).h + \|h\|_{\varepsilon_1}(h)$$

et pour tout $k \in V_2$,

$$g(b+h) = g(b) + dg(b).k + \|k\|_{\varepsilon_2}(k).$$

On en déduit en posant $k = df(a).h + \|h\|_{\varepsilon_1}(h)$:

$$\begin{aligned} g(f(a+h)) &= g(f(a) + df(a).h + \|h\|_{\varepsilon_1}(h)) \\ &= g(f(a)) + dg(b).(k) + \|k\|_{\varepsilon_2}(k) \\ &= g(f(a)) + dg(b).(df(a).h) + dg(b).(\|h\|_{\varepsilon_1}(h)) + \|k\|_{\varepsilon_2}(k) \end{aligned}$$

On a donc

$$g(f(a+h)) = g(f(a)) + (dg(b) \circ df(a)).h + \beta(h)$$

où

$$\beta(h) = dg(b).(\|h\|_{\varepsilon_1}(h)) + \|k\|_{\varepsilon_2}(k).$$

Maintenant, comme dg_b est continue et que $\varepsilon_1(h)$ tend vers 0 quand h tend vers 0 alors en utilisant la linéarité de dg_b , $dg_b(\|h\|_{\varepsilon_1}(h)) = \|h\|_{\varepsilon_1}(h) dg(b).$ est négligeable devant h .

De même, $\|df(a).h + \|h\|_{\varepsilon_1}(h)\| \leq (K + K')\|h\|$ où K est un majorant de $df(a)$ sur la boule unité et K' un majorant de $\varepsilon_1(h)$ au voisinage de 0 (qui existe car ε_1 tend vers 0 pour $h \rightarrow 0$). Comme de plus $h \mapsto df(a).h + \|h\|_{\varepsilon_1}(h)$ tend vers 0 pour $h \rightarrow 0$, on a bien que $h \rightarrow \|df(a).h + \|h\|_{\varepsilon_1}(h)\|_{\varepsilon_2}(df(a).h + \|h\|_{\varepsilon_1}(h))$ est négligeable devant h en 0.

Finalement,

$$g(f(a+h)) = g(f(a)) + (dg(b) \circ df(a)).h + o(h)$$

ce qui signifie que $g \circ f$ est différentiable en a et

$$d(g \circ f)(a) = dg(f(a)) \circ (df)(a).$$

□

ATTENTION

On peut retrouver grâce à ce résultat, les énoncés sur la linéarité des combinaisons linéaires et des « produits ». En effet si f, g sont deux fonctions de classe $\mathcal{C}^1$ de U dans F , la fonction $\Phi : U \rightarrow F \times F$ définie par

$$\Phi : x \mapsto (f(x), g(x))$$

est encore de classe $\mathcal{C}^1$ et sa différentielle est donnée par

$$\forall x \in U, \forall h \in E, d\Phi_x.h = (df_x.h, dg_x.h)$$

Comme de plus $(u, v) \mapsto \lambda u + \mu v$ est de classe $\mathcal{C}^1$ car linéaire, on obtient bien que $x \mapsto \lambda f(x) + \mu g(x)$ est de classe CCC^1 .

De même, pour le « produit ».

Corollaire 18.23 (Règle de la chaîne)

Soit (n, m, p) trois entiers non nuls, U un ouvert de $\mathbf{R}^n$ et V un ouvert de $\mathbf{R}^m$. Soit $f : U \rightarrow \mathbf{R}^m$ et $g : V \rightarrow \mathbf{R}^p$ avec $f(U) \subset V$. On note $(x_1, \dots, x_n)$ les coordonnées dans $\mathbf{R}^n$ et $(u_1, \dots, u_m)$ celles dans $\mathbf{R}^m$. Soit $a \in U$ et $b = f(a)$. On suppose que f est différentiable en a et g en b . Alors

$$\forall j \in \llbracket 1; n \rrbracket \quad \frac{\partial(g \circ f)}{\partial x_j}(a) = \sum_{k=1}^m \frac{\partial g}{\partial u_k}(b) \times \frac{\partial f_k}{\partial x_j}(a)$$

où pour tout $k \in \llbracket 1; m \rrbracket$, $f_k = u_k^* \circ f$.

De même, en projetant sur la base de $\mathbf{R}^p$,

$$\forall j \in \llbracket 1; n \rrbracket, \forall i \in \llbracket 1; p \rrbracket \quad \frac{\partial(g \circ f)_i}{\partial x_j}(a) = \sum_{k=1}^m \frac{\partial g_i}{\partial u_k}(b) \times \frac{\partial f_k}{\partial x_j}(a)$$

Démonstration : Il suffit d'écrire le résultat précédent avec mes matrices jacobiniennes. En effet on a en travaillant dans les bases canoniques,

$$J_{g \circ f}(a) = J_g(b) \times J_f(a).$$

C'est-à-dire

$$\begin{pmatrix} \frac{\partial(g \circ f)_1}{\partial x_1}(a) & \dots & \frac{\partial(g \circ f)_1}{\partial x_n}(a) \\ \vdots & & \vdots \\ \frac{\partial(g \circ f)_p}{\partial x_1}(a) & \dots & \frac{\partial(g \circ f)_p}{\partial x_n}(a) \end{pmatrix} = \begin{pmatrix} \frac{\partial g_1}{\partial u_1}(b) & \dots & \frac{\partial g_1}{\partial u_m}(b) \\ \vdots & & \vdots \\ \frac{\partial g_p}{\partial u_1}(b) & \dots & \frac{\partial g_p}{\partial u_m}(b) \end{pmatrix} \times \begin{pmatrix} \frac{\partial f_1}{\partial x_1}(a) & \dots & \frac{\partial f_1}{\partial x_n}(a) \\ \vdots & & \vdots \\ \frac{\partial f_m}{\partial x_1}(a) & \dots & \frac{\partial f_m}{\partial x_n}(a) \end{pmatrix}.$$

La formule voulue est celle du produit matriciel. □

Remarque : Cette propriété s'appelle aussi « chain rule » d'après sa dénomination anglo-saxonne.

Exemple : On utilise cette formule quand on veut faire des changements de variable. Par exemple on veut déterminer les fonctions $f : \mathbf{R}^2 \rightarrow \mathbf{R}$ telles que

$$3 \frac{\partial f}{\partial x} - 2 \frac{\partial f}{\partial y} = 0. \quad (\text{E})$$

Pour cela on pose le changement de variables

$$\begin{cases} u &= x + y \\ v &= 2x + 3y \end{cases}$$

On peut l'inverser et obtenir

$$\begin{cases} x &= 3u - v \\ y &= v - 2u \end{cases}$$

C'est-à-dire que l'on considère $\varphi : (u, v) \mapsto (3u - v, v - 2u)$ et on note alors $x = \varphi_1(u, v) = 3u - v$ et $y = \varphi_2(u, v) = v - 2u$. On considère alors

$$F = f \circ \varphi : (u, v) \mapsto f(\varphi(u, v)) = f(3u - v, v - 2u).$$

Dès lors d'après le corollaire ci-dessus,

$$\frac{\partial F}{\partial u}(u, v) = \frac{\partial f}{\partial x}(x, y) \times \frac{\partial x}{\partial u}(u, v) + \frac{\partial f}{\partial y}(x, y) \times \frac{\partial y}{\partial u}(u, v)$$

Ici on a remplacé φ_1 et φ_2 par x et y , ce qui fait apparaître « la chaîne » :

$$\frac{\partial F}{\partial u} = \frac{\partial f}{\partial x} \times \frac{\partial x}{\partial u} + \frac{\partial f}{\partial y} \times \frac{\partial y}{\partial u}.$$

Or

$$\frac{\partial x}{\partial u} = 3 \text{ et } \frac{\partial y}{\partial u} = -2.$$

On en déduit que si f vérifie (E) alors

$$\frac{\partial F}{\partial u} = 3 \frac{\partial f}{\partial x} - 2 \frac{\partial f}{\partial y} = 0.$$

Cela implique que F ne dépend pas de u et donc qu'il existe une fonction $\theta : \mathbf{R} \rightarrow \mathbf{R}$ telle que

$$\forall (u, v) \in \mathbf{R}^2, F(u, v) = \theta(v).$$

En utilisant alors que $f = F \circ \varphi^{-1}$ on obtient que $f : (x, y) \rightarrow \theta(2x + 3y)$.

Réciproquement, on montre que si θ est dérivable, alors toute fonction de la forme

$$f : (x, y) \rightarrow \theta(2x + 3y)$$

vérifie (E).

Exercice : Déterminer les fonctions définies sur $\mathbf{R}^2 \setminus \mathbf{R}_-$ vérifiant

$$y \frac{\partial f}{\partial x} - x \frac{\partial f}{\partial y} = 0. \quad (\text{E})$$

On pourra utiliser un changement de variable polaire en posant

$$\begin{cases} x &= \rho \cos \theta \\ y &= \rho \sin \theta \end{cases}$$

Correction

On pose $F(\rho, \theta) = f(x, y)$; c'est-à-dire que $F = f \circ \varphi$ où $\varphi : (\rho, \theta) \mapsto (\rho \cos \theta, \rho \sin \theta)$.
 Grace à la règle de la chaîne, on peut calculer les dérivées partielles de F en fonction de celles de f .

$$\frac{\partial F}{\partial \rho}(\rho, \theta) = \frac{\partial f}{\partial x}(x, y) \times \frac{\partial x}{\partial \rho}(\rho, \theta) + \frac{\partial f}{\partial y}(x, y) \times \frac{\partial y}{\partial \rho}(\rho, \theta) = \frac{\partial f}{\partial x}(x, y) \times \cos \theta + \frac{\partial f}{\partial y}(x, y) \times \sin \theta$$

et

$$\frac{\partial F}{\partial \theta}(\rho, \theta) = \frac{\partial f}{\partial x}(x, y) \times \frac{\partial x}{\partial \theta}(\rho, \theta) + \frac{\partial f}{\partial y}(x, y) \times \frac{\partial y}{\partial \theta}(\rho, \theta) = -\frac{\partial f}{\partial x}(x, y) \times \rho \sin \theta + \frac{\partial f}{\partial y}(x, y) \times \rho \cos \theta$$

On remarque alors que si f vérifie (E),

$$\frac{\partial F}{\partial \theta}(\rho, \theta) = -y \frac{\partial f}{\partial x}(x, y) + x \frac{\partial f}{\partial y}(x, y) = 0$$

Cela signifie que F ne dépend que de ρ .

On en déduit qu'il existe une fonction g telle que $F : (\rho, \theta) \mapsto g(\rho)$.

Finalement $f : (x, y) \mapsto g(\sqrt{x^2 + y^2})$.

On pouvait aussi voir que la règle de la chaîne nous donne :

$$J_F = J_f \times J_\varphi$$

La matrice $J_\varphi = \begin{pmatrix} \cos \theta & -\rho \sin \theta \\ \sin \theta & \rho \cos \theta \end{pmatrix}$ est inversible et

$$(J_\varphi)^{-1} = \frac{1}{\rho} \begin{pmatrix} \rho \cos \theta & \rho \sin \theta \\ -\sin \theta & \cos \theta \end{pmatrix}$$

En utilisant que $J_f = J_F \times (J_\varphi)^{-1}$ on obtient alors que

$$\frac{\partial f}{\partial x}(x, y) = \cos \theta \frac{\partial F}{\partial \rho}(\rho, \theta) - \frac{1}{\rho} \sin \theta \frac{\partial F}{\partial \theta}(\rho, \theta)$$

et

$$\frac{\partial f}{\partial y}(x, y) = \sin \theta \frac{\partial F}{\partial \rho}(\rho, \theta) + \frac{1}{\rho} \cos \theta \frac{\partial F}{\partial \theta}(\rho, \theta)$$

On en déduit que

$$y \frac{\partial f}{\partial x}(x, y) - x \frac{\partial f}{\partial y}(x, y) = -(\sin^2 \theta + \cos^2 \theta) \frac{\partial F}{\partial \theta}(\rho, \theta)$$

On retrouve que si f vérifie (E) alors $\frac{\partial F}{\partial \theta}(\rho, \theta) = 0$.

Exemple : On peut aussi calculer le gradient dans d'autres systèmes de paramètres. En particulier en polaire.

Précisément soit f une fonction définie sur un ouvert U de $\mathbf{R}^2 \setminus \{(0, 0)\}$ à valeurs dans $\mathbf{R}$. On lui

associe la fonction

$$F : (\rho, \theta) \mapsto f(\rho \cos \theta, \rho \sin \theta)$$

C'est-à-dire que $F = f \circ \phi$ où

$$\begin{aligned} \phi :]0, +\infty[\times [0, 2\pi[&\rightarrow \mathbf{R}^2 \\ (\rho, \theta) &\mapsto (\rho \cos \theta, \rho \sin \theta) \end{aligned}$$

On peut alors expliciter les relations entre les dérivées partielles de f (par rapport à x et y) et les dérivées partielles de F (par rapport à ρ et θ).

En utilisant la règle de la chaîne on a $J_F = J_f \times J_\phi$ c'est à-dire :

$$\begin{pmatrix} \frac{\partial F}{\partial \rho} & \frac{\partial F}{\partial \theta} \end{pmatrix} = \begin{pmatrix} \frac{\partial f}{\partial x} & \frac{\partial f}{\partial y} \end{pmatrix} \times \begin{pmatrix} \frac{\partial x}{\partial \rho} & \frac{\partial x}{\partial \theta} \\ \frac{\partial y}{\partial \rho} & \frac{\partial y}{\partial \theta} \end{pmatrix} = \begin{pmatrix} \frac{\partial f}{\partial x} & \frac{\partial f}{\partial y} \end{pmatrix} \times \begin{pmatrix} \cos \theta & -\rho \sin \theta \\ \sin \theta & \rho \cos \theta \end{pmatrix}$$

Ecrit autrement on a au point b et en notant $a = \phi(b)$,

$$\frac{\partial F}{\partial \rho}(b) = \cos \theta \frac{\partial f}{\partial x}(a) + \sin \theta \frac{\partial f}{\partial y}(a)$$

et

$$\frac{\partial F}{\partial \theta}(b) = -\rho \sin \theta \frac{\partial f}{\partial x}(a) + \rho \cos \theta \frac{\partial f}{\partial y}(a)$$

On peut aussi écrire ces formules dans « l'autre sens » en inversant J_ϕ . En effet $\det(J_\phi) = \rho(\cos^2 \theta + \sin^2 \theta) = \rho > 0$. Elle est donc inversible et on a alors $J_f = J_F \times (J_\phi)^{-1}$. C'est-à-dire

$$\begin{pmatrix} \frac{\partial f}{\partial x} & \frac{\partial f}{\partial y} \end{pmatrix} = \begin{pmatrix} \frac{\partial F}{\partial \rho} & \frac{\partial F}{\partial \theta} \end{pmatrix} \times \frac{1}{\rho} \begin{pmatrix} \rho \cos \theta & \rho \sin \theta \\ -\sin \theta & \cos \theta \end{pmatrix}$$

On en déduit les formules

$$\frac{\partial f}{\partial x}(a) = \cos \theta \frac{\partial F}{\partial \rho}(b) - \frac{1}{\rho} \sin \theta \frac{\partial F}{\partial \theta}(b)$$

et

$$\frac{\partial f}{\partial y}(a) = \sin \theta \frac{\partial F}{\partial \rho}(b) + \frac{1}{\rho} \cos \theta \frac{\partial F}{\partial \theta}(b)$$

On peut donc exprimer le gradient de f en un point en fonction des coordonnées polaires. Si on note (e_1, e_2) la base canonique de $\mathbf{R}^2$,

$$\begin{aligned} \nabla f(a) &= \frac{\partial f}{\partial x}(a)e_1 + \frac{\partial f}{\partial y}(a)e_2 \\ &= \left(\cos \theta \frac{\partial F}{\partial \rho}(b) - \frac{1}{\rho} \sin \theta \frac{\partial F}{\partial \theta}(b) \right) e_1 + \left(\sin \theta \frac{\partial F}{\partial \rho}(b) + \frac{1}{\rho} \cos \theta \frac{\partial F}{\partial \theta}(b) \right) e_2 \\ &= \frac{\partial F}{\partial \rho}(b)(\cos \theta e_1 + \sin \theta e_2) + \frac{1}{\rho} \frac{\partial F}{\partial \theta}(b)(-\sin \theta e_1 + \cos \theta e_2) \\ &= \frac{\partial F}{\partial \rho}(b)u(\theta) + \frac{1}{\rho} \frac{\partial F}{\partial \theta}(b)v(\theta) \end{aligned}$$

où $u(\theta)$ est le vecteur $(\cos \theta e_1 + \sin \theta e_2)$ et $v(\theta)$ le vecteur directement orthogonal. C'est la « base tournante » du repère polaire.

3 Dérivée le long d'un arc

3.1 Définition

Proposition 18.24

Soit $\gamma : I \rightarrow E$ une fonction dérivable. Soit $f : U \rightarrow F$ une fonction. Soit $t_0 \in I$, on suppose que

- i) La fonction γ est dérivable en t_0 .
- ii) $\gamma(t_0) \in U$ et f différentiable en $\gamma(t_0)$

Alors $t \mapsto f(\gamma(t))$ est dérivable en t_0 et

$$(f \circ \gamma)'(t_0) = df(\gamma(t_0)).\gamma'(t_0)$$

Démonstration : Il suffit d'appliquer le théorème de composition. En effet comme γ est dérivable en t_0 alors elle est différentiable et

$$\begin{aligned} d\gamma(t_0) : \mathbf{R} &\rightarrow E \\ x &\mapsto x\gamma'(t_0) \end{aligned}$$

De ce fait, $f \circ \gamma$ est différentiable en t_0 et

$$d(f \circ \gamma)(t_0) = df(\gamma(t_0)) \circ d\gamma(t_0)$$

On a donc

$$\begin{aligned} d(f \circ \gamma)(t_0) : \mathbf{R} &\rightarrow E \\ x &\mapsto df(\gamma(t_0)).(x\gamma'(t_0)) = xdf(\gamma(t_0)).\gamma'(t_0) \end{aligned}$$

Cela signifie que $f \circ \gamma$ est dérivable en t_0 et que

$$(f \circ \gamma)'(t_0) = df(\gamma(t_0)).\gamma'(t_0).$$

□

Remarque : On peut aussi faire le calcul « à la main ». En effet on a

$$\gamma(t_0 + h) = \gamma(t_0) + h\gamma'(t_0) + o(h)$$

donc

$$\begin{aligned} (f \circ \gamma)(t_0 + h) &= f(\gamma(t_0)) + df(\gamma(t_0)).(h\gamma'(t_0) + o(h)) + o(h) \\ &= f(\gamma(t_0)) + h.df(\gamma(t_0)).\gamma'(t_0) + o(h) \end{aligned}$$

Note

Si on considère $a \in U$ et $v \in E$ et que l'on pose $\gamma : t \mapsto a + tv$. On retrouve la dérivée selon un vecteur en effet

$$D_v f(a) = (f \circ \gamma)'(0) = df(a).v$$

car $\gamma'(0) = v$.

Proposition 18.25 (Expression en coordonnées)

Soit $x_1, \dots, x_n$ des fonctions définies au voisinage de t_0 et à valeurs dans $\mathbf{R}$.
 On pose $a = (x_1(t_0), \dots, x_n(t_0))$. Soit f une fonction définie dans un voisinage de a dans $\mathbf{R}^n$. On suppose que les fonctions x_i sont dérivables en t_0 et que f est différentiable en a .
 Alors $F : t \mapsto f(x_1(t), \dots, x_n(t))$ est dérivable en t_0 et

$$F'(t_0) = df(a) \cdot (x'_1(t_0), \dots, x'_n(t_0)) = \sum_{i=1}^n \frac{\partial f}{\partial x_i}(a) x'_i(t_0)$$

3.2 Théorème fondamental de l'analyse

Théorème 18.26 (Théorème fondamental de l'analyse)

Soit $f : U \rightarrow F$ une fonction de classe $\mathcal{C}^1$. Soit $\gamma : [0, 1] \rightarrow U$ une fonction de classe $\mathcal{C}^1$.
 On pose $a = \gamma(0)$ et $b = \gamma(1)$. Alors

$$f(b) - f(a) = \int_0^1 df(\gamma(t)) \cdot \gamma'(t) dt$$

Démonstration : Il suffit de poser $g = f \circ \gamma : [0, 1] \rightarrow F$. On sait que g est de classe $\mathcal{C}^1$ et que

$$\forall t \in [0, 1], g'(t) = df(\gamma(t)) \cdot \gamma'(t).$$

On a donc

$$f(b) - f(a) = g(1) - g(0) = \int_0^1 g'(t) dt = \int_0^1 df(\gamma(t)) \cdot \gamma'(t) dt.$$

□

Note

De manière plus générale, on appelle 1-forme différentielle une fonction $\omega : U \rightarrow \mathcal{L}(E, F)$.
 De ce fait, on peut définir la circulation d'une forme différentielle le long d'un arc $\gamma : [0, 1] \rightarrow U$ par

$$\int_\gamma \omega := \int_0^1 \omega(\gamma(t)) \cdot \gamma'(t) dt.$$

Ce que dit le théorème ci-dessus, c'est que quand la forme différentielle est la différentielle d'une fonction (on dit que ω est une forme exacte) alors la circulation ne dépend pas de l'arc en particulier mais juste de ses extrémités.

Proposition 18.27

Soit U un ouvert connexe par arcs. Une fonction $f : U \rightarrow F$ de classe $\mathcal{C}^1$. Elle est constante si et seulement si sa différentielle est nulle.

Démonstration :

- $\Rightarrow$ déjà vu (et évident)
- $\Leftarrow$ On ne le démontre que si U est convexe (mais le résultat énoncé est juste) afin de pouvoir simplement relier deux points de U par un arc non seulement continu mais de classe $\mathcal{C}^1$. On fixe un point $x_0 \in U$. Maintenant pour tout $a \in U$ il existe un arc γ de classe $\mathcal{C}^1$ tel que $\gamma(0) = x_0$ et $\gamma(1) = a$. Dès lors, d'après le théorème fondamental de l'analyse,

$$f(a) - f(x_0) = \int_0^1 df(\gamma(t)) \cdot \gamma'(t) dt = 0$$

On a bien que f est constante.

□

3.3 Vecteurs tangents à une partie

Définition 18.28

Soit X une partie de E et $x \in X$. Un vecteur v de E est tangent à X s'il existe un voisinage I de 0 et un arc γ défini sur I dérivable en 0 tel que $\gamma(0) = x$ et $\gamma'(0) = v$.
On note $T_x X$ l'ensemble des vecteurs tangents à X en x .

Exemples :

1. Dans $\mathbb{R}^3$ on considère la boule unité fermée $X = \bar{B} = \{x \in \mathbb{R}^3 \mid \|x\| \leq 1\}$.

- Soit a un point intérieur à $\bar{B}$ alors tout vecteur est un vecteur tangent. En effet on sait qu'il existe $\delta > 0$ tel que $B(a, \delta) \subset \bar{B}$. Dès lors pour tout vecteur v , on peut poser

$$\gamma : t \mapsto a + tv$$

défini sur $] -\varepsilon, \varepsilon[$ où $\varepsilon\|v\| \leq \delta$. Il est clair que $\gamma(0) = a$ et que $\gamma'(0) = v$.

- Soit a un point de la sphère $S = \{x \in \mathbb{R}^3 \mid \|x\| = 1\}$. On veut montrer que les vecteurs tangents à a à $\bar{B}$ (ou à S) sont les vecteurs orthogonaux à a .

Soit $\gamma : t \mapsto \gamma(t)$ un arc à valeurs dans $\bar{B}$ avec $\gamma(0) = a$. On regarde $\phi : t \mapsto \|\gamma(t)\|^2 = (\gamma(t)|\gamma(t))$. Elle est dérivable et

$$\phi'(t) = 2(\gamma(t)|\gamma'(t)).$$

En particulier, comme 0 est un maximum de la fonction (car $\phi(0) = 1$) on a $\phi'(0) = 0$ donc $\gamma'(0) \perp \gamma(0)$.

Réciproquement, montrons que tout vecteur v orthogonal à a est un vecteur tangent. On pose l'arc

$$\gamma : t \mapsto \frac{a + tv}{\|a + tv\|} = \frac{a + tv}{\sqrt{\|a\|^2 + t^2\|v\|^2}}$$

Il est bien à valeur dans $\bar{B}$ (ou dans S) car pour tout $t \in I$, $\|\gamma(t)\| = 1$ par construction. En dérivant et en utilisant que $\|a\| = 1$, on obtient

$$\gamma'(t) = \frac{(1 + t^2\|v\|^2)v - t\|v\|^2(a + tv)}{(1 + t^2\|v\|^2)^{3/2}}$$

En particulier $\gamma'(0) = v$. On en déduit que v est un vecteur tangent.

2. Soit X un sous-espace affine de E . Il existe $a \in E$ et G un sous-espace vectoriel de E tel que $X = a + G = \{a + g, g \in G\}$. Soit $x \in X$, montrons que $T_x X = G$.
- Soit $\gamma : I \rightarrow X$ un arc tracé sur X . Par définition,

$$\gamma'(0) = \lim_{h \rightarrow 0} \frac{\gamma(h) - \gamma(0)}{h}$$

Or, pour tout $h \in I \setminus \{0\}$, $\frac{1}{h}(\gamma(h) - \gamma(0)) \in G$. Comme G est fermé en tant qu'espace vectoriel de dimension finie, $\gamma'(0) \in G$. On a montré que pour tout $x \in X$, $T_x X \subset G$.

- Réciproquement, soit $x \in X$ et $g \in G$. On pose $\gamma : t \mapsto x + tg$. On voit que γ est à valeurs dans G , que $\gamma(0) = x$ et que $\gamma'(0) = g$. Cela montre que $g \in T_x X$.

Exercice : On cherche les vecteurs tangents à $\mathcal{O}_n(\mathbf{R})$. On note $\mathcal{A}_n(\mathbf{R})$ l'ensemble des matrices antisymétriques.

1. Soit $s \mapsto \gamma(s)$ un arc de classe $\mathcal{C}^1$ à valeurs dans $\mathcal{O}_n(\mathbf{R})$. Montrer que $\gamma^{-1}(0)\gamma'(0)$ est une matrice antisymétrique. En déduire que les vecteurs tangents à $\mathcal{O}_n(\mathbf{R})$ en M s'écrivent MA où A est antisymétrique.
2. Réciproquement, soit $M \in \mathcal{O}_n(\mathbf{R})$ et $A \in \mathcal{A}_n(\mathbf{R})$. Montrer que MA est un vecteur tangent à $\mathcal{O}_n(\mathbf{R})$ en M . On pourra considérer $s \mapsto M \exp(sA)$.

Correction

1. Soit $s \mapsto \gamma(s)$ un arc de classe $\mathcal{C}^1$ à valeurs dans $\mathcal{O}_n(\mathbf{R})$. Par hypothèse, $\gamma(s)^T \gamma(s) = I_n$. On veut dériver cette relation.
On utilise que la dérivée de $s \mapsto \gamma(s)^T$ est $s \mapsto (\gamma'(s))^T$ car $M \mapsto M^T$ est linéaire.
On obtient donc

$$\gamma'(s)^T \gamma(s) + \gamma(s)^T \gamma'(s) = 0$$

Cela peut aussi s'écrire

$$\gamma(s)^T \gamma'(s) = -\gamma'(s)^T \gamma(s) = -\left(\gamma(s)^T \gamma'(s)\right)^T$$

En particulier pour $s = 0$, on obtient que $\gamma(0)^T \gamma'(0) = \gamma(0)^{-1} \gamma'(0)$ est antisymétrique.

Soit v est un vecteur tangent à $\mathcal{O}_n(\mathbf{R})$ en M alors, si γ est un arc tracé dans $\mathcal{O}_n(\mathbf{R})$ tel que $\gamma(0) = M$ et $\gamma'(0) = v$, la relation ci-dessus nous dit qu'il existe une matrice A antisymétrique telle que $M^{-1}v = A$ et donc $v = MA$.

2. On remarque que si $A \in \mathcal{A}_n(\mathbf{R})$ alors pour $s \in \mathbf{R}$,

$$(\exp(sA))^T = \exp(sA^T) = \exp(-sA) = \exp(sA)^{-1}$$

Donc $\gamma : s \mapsto M \exp(sA)$ est un arc tracé dans $\mathcal{O}_n(\mathbf{R})$ (car un produit de deux matrices orthogonales est orthogonal). Maintenant, $\gamma(0) = M$ et $\gamma'(s) = MA \exp(sA)$ et donc $\gamma'(0) = MA$. Cela montre que MA est un vecteur tangent à $\mathcal{O}_n(\mathbf{R})$ en M .

Définition 18.29

Soit f une fonction définie d'un ouvert U de $\mathbf{R}^2$ et à valeurs dans $\mathbf{R}$. On appelle surface représentative de f (ou graphe de f) la partie

$$\Gamma_f = \{(x, y, z) \mid (x, y) \in U \text{ et } z = f(x, y)\}$$

Proposition 18.30

Soit f une fonction définie d'un ouvert U de $\mathbf{R}^2$ et à valeurs dans $\mathbf{R}$. Soit $a = (x_0, y_0) \in U$. On suppose que f est différentiable en a . L'ensemble des vecteurs tangents à Γ_f en $(x_0, y_0, f(a))$ sont les vecteurs (x, y, z) vérifiant $z = df(a).(x, y)$

Démonstration :

- Montrons que si (x, y, z) est un vecteur tangent à Γ_f alors $z = df(a).(x, y)$. Soit $v = (x, y, z)$ un vecteur tangent à Γ_f en $(x_0, y_0, f(a))$

On considère un arc paramétré γ de Γ_f tel que $\gamma(0) = (x_0, y_0, f(a))$ et $\gamma'(0) = (x, y, z)$. Il est donc de la forme $t \mapsto \gamma(t) = (x(t), y(t), z(t))$ où $z(t) = f(x(t), y(t))$.

En dérivant cette relation on obtient que

$$z'(t) = df(x(t), y(t)).(x'(t), y'(t)).$$

En particulier, en 0, on obtient $z = z'(0) = df(x_0, y_0).(x, y)$.

- Réciproquement, soit $v = (x, y, z)$ vérifiant que $z = df(a).(x, y)$. On pose

$$\gamma : t \mapsto (x_0 + tx, y_0 + ty, f(x_0 + tx, y_0 + ty))$$

Par construction, γ est à valeurs dans Γ_f et $\gamma(0) = (x_0, y_0, f(a))$. En dérivant γ on obtient

$$\gamma'(t) = (x, y, df(x_0 + tx, y_0 + ty).(x, y))$$

En particulier,

$$\gamma'(0) = (x, y, df(a).(x, y)).$$

On en déduit que $v = (x, y, z)$ est bien un vecteur tangent à Γ_f .

Remarque : Le plan affine passant par $(x_0, y_0, f(a))$ et de direction $H = \{(x, y, z) \mid z = df(a).(x, y)\}$ s'appelle le plan tangent à Γ_f .

Exemple : Si on considère l'application $f : (x, y) \mapsto \sqrt{1 - (x^2 + y^2)}$ définie sur $\{(x, y) \in \mathbf{R}^2 \mid x^2 + y^2 < 1\}$. On sait que la surface représentative est l'ensemble des $(x, y, z) \in \mathbf{R}^3$ tels que $z = \sqrt{1 - x^2 - y^2}$ c'est donc la demi-sphère

$$\Gamma_f = \{(x, y, z) \in \mathbf{R}^3, x^2 + y^2 + z^2 = 1 \text{ et } z > 0\}$$

Soit $a = (x_0, y_0)$ avec $x_0^2 + y_0^2 < 1$ et $m = (x_0, y_0, z_0)$ le point de Γ_f correspondant. On calcule df . On a

$$\frac{\partial f}{\partial x}(x_0, y_0) = \frac{-x_0}{\sqrt{1 - x_0^2 - y_0^2}} \text{ et } \frac{\partial f}{\partial y}(x_0, y_0) = \frac{-y_0}{\sqrt{1 - x_0^2 - y_0^2}}.$$

On en déduit que

$$df(a) : (x, y) \mapsto -\frac{x_0 x}{z_0} - \frac{y_0 y}{z_0}$$

De ce fait, le plan tangent à la demi-sphère au point m a pour équation

$$z = -\frac{x_0 x}{z_0} - \frac{y_0 y}{z_0} \iff x_0 x + y_0 y + z_0 z = 0$$

C'est-à-dire que le vecteur (x_0, y_0, z_0) est un vecteur normal au plan.

Théorème 18.31

Soit g une fonction de classe $\mathcal{C}^1$ définie d'un ouvert U de E dans $\mathbf{R}$.

On pose $X = g^{-1}(\{0\}) = \{x \in U, g(x) = 0\}$. Soit $x \in X$, si $dg(x) \neq 0$, $T_x X = \text{Ker}(dg(x))$.

Démonstration : Ce théorème est admis.

Il est simple de montrer que $T_x X \subset \text{Ker}(dg(x))$. Il suffit pour cela de considérer un vecteur v tangent à X en x . Il existe un arc γ de classe $\mathcal{C}^1$ à valeur dans X tel que $\gamma(0) = x$ et $\gamma'(0) = v$. Par construction $g \circ \gamma$ est constante (égale à 0) et, par dérivation

$$0 = (g \circ \gamma)'(0) = dg(x) \cdot \gamma'(0) = dg(x) \cdot v$$

Cela montre que $v \in \text{Ker}(dg(x))$.

La réciproque est plus compliquée. Le principe est d'utiliser le théorème des fonctions implicites qui permet, en utilisant que $dg(x) \neq 0$, de paramétrer l'espace X en déterminant une fonction φ telle que, au voisinage de x ,

$$(x_1, \dots, x_n) \in X \iff x_1 = \varphi(x_2, \dots, x_n)$$

On peut alors conclure comme dans l'énoncé précédent. □

Remarques :

1. C'est une généralisation de l'énoncé précédent. En effet dans le cas de l'énoncé précédent, on peut considérer $\Omega = U \times \mathbf{R}$ et poser $g : \Omega \rightarrow \mathbf{R}$ définie par $g(x, y, z) = z - f(x, y)$. La surface Γ_f est bien égale à $X = g^{-1}(\{0\})$. De plus pour tout $a = (x_0, y_0) \in U$ et $z = f(a)$, le point $\omega = (x_0, y_0, z) \in X$ et

$$dg(\omega) = -\frac{\partial f}{\partial x}(a)dx - \frac{\partial f}{\partial y}(a)dy + dz \neq 0$$

On en déduit que l'espace tangent à X en ω , $T_\omega X$ est égal à $\text{Ker}(dg(\omega))$. Ce sont les vecteurs (x, y, z) de $\mathbf{R}^3$ tels que

$$-x \frac{\partial f}{\partial x}(a) - y \frac{\partial f}{\partial y}(a) + z = 0 \iff z = df(a) \cdot (x, y)$$

2. Cela signifie que lorsque l'on se déplace dans l'hyperplan défini par le noyau de la différentielle, la valeur de g ne change pas trop; c'est donc l'hyperplan tangent à X .

Exemples :

1. Dans le cas où $E = \mathbf{R}^n$ (ou plus généralement quand E est un espace euclidien). On obtient avec les notations du théorème que pour tout $\alpha = (\alpha_1, \dots, \alpha_n) \in \mathbf{R}^n$,

$$\begin{aligned} \alpha \in T_x X &\iff \alpha \in \text{Ker}(dg(x)) \\ &\iff dg(x) \cdot \alpha = 0 \\ &\iff (\text{grad } g(x) | \alpha) = 0 \\ &\iff \sum_{i=1}^n \frac{\partial g}{\partial x_i}(x) \alpha_i = 0 \end{aligned}$$

L'hyperplan tangent est défini comme l'ensemble des vecteurs orthogonaux au gradient de g en x .

2. On considère une surface $X \subset \mathbb{R}^3$ définie par une équation g , c'est-à-dire que l'on a une fonction $g : \mathbb{R}^3 \rightarrow \mathbb{R}$ telle que pour tout $a \in \mathbb{R}^3$, $a \in X \iff g(a) = 0$. Le plan tangent à X en a (dans le cas où g est de classe $\mathcal{C}^1$ au voisinage de a et que $dg(a) \neq 0$) est donné par l'équation

$$T_a X = \{(x, y, z) \in \mathbb{R}^3, \frac{\partial g}{\partial x}(a)x + \frac{\partial g}{\partial y}(a)y + \frac{\partial g}{\partial z}(a)z = 0\}$$

Si on considère la sphère unité X qui est définie par l'équation $g : (x, y, z) \mapsto x^2 + y^2 + z^2 - 1$. Soit $a = (x_0, y_0, z_0) \in X$. Le plan tangent à X en a est le plan d'équation $2x_0x + 2y_0y + 2z_0z = 0$. C'est l'ensemble des vecteurs orthogonaux à $\text{grad } g(a) = (2x_0, 2y_0, 2z_0)$.

4 Fonctions de classe $\mathcal{C}^k$

4.1 Définitions

Définition 18.32

Soit $f : U \rightarrow F$ une fonction et $\mathcal{B} = (e_1, \dots, e_n)$ une base de E . On dit que f admet des dérivées partielles secondes lorsque les dérivées partielles premières admettent des dérivées partielles. On note alors

$$\forall i \in \llbracket 1; n \rrbracket^2, \frac{\partial^2 f}{\partial x_i \partial x_j} = \frac{\partial \left(\frac{\partial f}{\partial x_j} \right)}{\partial x_i}.$$

Dans le cas où $i = j$ on notera juste

$$\frac{\partial^2 f}{\partial x_i^2}.$$

Remarque : Pour alléger les notations on notera aussi $\partial_i \partial_j f$ ou $\partial_{i,j} f$ pour $\frac{\partial^2 f}{\partial x_i \partial x_j}$.

Exemple : Soit $f : (x, y) \mapsto (x + y^2) \sin(xy)$. On a

$$\frac{\partial f}{\partial x}(x, y) = \sin(xy) + y(x + y^2) \cos(xy) \text{ et } \frac{\partial f}{\partial y}(x, y) = 2y \sin(xy) + x(x + y^2) \cos(xy).$$

On en déduit que f admet des dérivées partielles secondes :

$$\frac{\partial^2 f}{\partial x^2}(x, y) = y \cos(xy) + y \cos(xy) - y(x + y^2) \sin(xy); \quad \frac{\partial^2 f}{\partial y^2}(x, y) = 2 \sin(xy) + 4xy \cos(xy) - x^2(x + y^2) \sin(xy)$$

et

$$\frac{\partial^2 f}{\partial x \partial y}(x, y) = \frac{\partial^2 f}{\partial y \partial x}(x, y) = x \cos(xy) + (x + 3y^2) \cos(xy) - xy(x + y^2) \sin(xy).$$

Définition 18.33

Soit $f : U \rightarrow F$ une fonction et $\mathcal{B} = (e_1, \dots, e_n)$ une base de E . On peut comme précédemment définir les dérivées successives partielles d'ordre r

$$\forall r \in \mathbb{N}^*, \forall (i_1, \dots, i_r) \in \llbracket 1; n \rrbracket^r, \frac{\partial^r f}{\partial x_{i_1} \partial x_{i_2} \cdots \partial x_{i_r}} = \frac{\partial \left(\frac{\partial^{r-1} f}{\partial x_{i_2} \cdots \partial x_{i_r}} \right)}{\partial x_{i_1}}$$

Remarque : De même on notera $\partial_{i_1} \cdots \partial_{i_r} f$ ou $\partial_{i_1, \dots, i_r} f$ pour $\frac{\partial^r f}{\partial x_{i_1} \partial x_{i_2} \cdots \partial x_{i_r}}$.

Définition 18.34

Soit $k \in \mathbb{N}^*$. Une fonction $f : U \rightarrow F$ est dite de classe $\mathcal{C}^k$, si toutes les dérivées partielles d'ordre k existent (il y en a n^k) et sont continues.

Proposition 18.35

Soit $k \geq 2$. Soit f une fonction de classe $\mathcal{C}^k$. Par définition les dérivées partielles d'ordre $k-1$ existent, de plus elles sont de classe $\mathcal{C}^1$ donc continues. De ce fait f est de classe $\mathcal{C}^{k-1}$.

Théorème 18.36 (Théorème de Schwarz)

Soit $f : U \rightarrow F$ et $\mathcal{B} = (e_1, \dots, e_n)$ une base de E . On suppose que f est classe $\mathcal{C}^2$ alors pour tout $(i, j) \in \llbracket 1; n \rrbracket^2$,

$$\partial_{i,j} f = \partial_{j,i} f$$

Remarque : Il existe une version pour les dérivées d'ordre supérieur. Si on suppose que f est de classe $\mathcal{C}^k$ alors l'ordre de dérivation des dérivées partielles d'ordre inférieur à k n'importe pas :

$$\forall r \in \llbracket 2; k \rrbracket, \forall (i_1, \dots, i_r) \in \llbracket 1; n \rrbracket^r, \forall \sigma \in \mathfrak{S}_r, \partial_{i_1, i_2, \dots, i_r} f = \partial_{i_{\sigma(1)}, \dots, i_{\sigma(r)}} f$$

Démonstration : (Non exigible) Pour simplifier les notations on va faire la preuve pour une fonction f définie sur un ouvert U de $\mathbb{R}^2$. Soit $(x_0, y_0) \in U$. Soit h_x, h_y suffisamment petit, on applique le théorème des accroissements finis à $\alpha : x \mapsto f(x, y_0 + h_y) - f(x, y_0)$. Il existe u compris entre 0 et h_x tels que

$$\begin{aligned} f(x_0 + h_x, y_0 + h_y) - f(x_0, y_0 + h_y) - f(x_0 + h_x, y_0) + f(x_0, y_0) &= \alpha(x_0 + h_x) - \alpha(x_0) \\ &= h_x \alpha'(x_0 + u) \\ &= h_x (\partial_x f(x_0 + u, y_0 + h_y) - \partial_x f(x_0 + u, y_0)) \end{aligned}$$

En appliquant maintenant le théorème des accroissements finis à $\beta : y \mapsto \partial_x f(x_0 + u, y)$, il existe v compris entre 0 et h_y tel que

$$\begin{aligned} f(x_0 + h_x, y_0 + h_y) - f(x_0, y_0 + h_y) - f(x_0 + h_x, y_0) + f(x_0, y_0) &= h_x (\beta(y_0 + h_y) - \beta(y_0)) \\ &= h_x h_y \beta'(y_0 + v) \\ &= h_x h_y \partial_{y,x} f(x_0 + u, y_0 + v) \end{aligned}$$

En procédant de même en considérant d'abord $\alpha' : y \mapsto f(x_0 + h_x, y) - f(x_0, y)$ il existe v' compris entre 0 et h_y tel que

$$f(x_0 + h_x, y_0 + h_y) - f(x_0, y_0 + h_y) - f(x_0 + h_x, y_0) + f(x_0, y_0) = h_y (\partial_y f(x_0 + h_x, y_0 + v') - \partial_y f(x_0, y_0 + v'))$$

On considère alors $\beta : x \mapsto \partial_y f(x, y_0 + v')$ et on en déduit qu'il existe u' compris entre 0 et h_x tel que

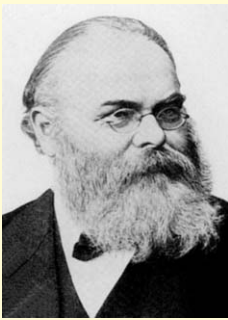
$$f(x_0 + h_x, y_0 + h_y) - f(x_0, y_0 + h_y) - f(x_0 + h_x, y_0) + f(x_0, y_0) = h_x h_y \partial_{x,y} f(x_0 + u', y_0 + v')$$

On en déduit que

$$\partial_{y,x} f(x_0 + u, y_0 + v) = \partial_{x,y} f(x_0 + u', y_0 + v')$$

Il suffit alors de faire tendre (h_x, h_y) vers $(0, 0)$, on obtient que (u, v) et (u', v') tend vers $(0, 0)$. Par continuité de $\partial_{y,x} f$ et de $\partial_{x,y} f$ on obtient le résultat voulu. $\square$

Matheux (Hermann Amandus Schwarz : 1843 - 1921)



Hermann Amandus Schwarz est né le 25 janvier 1843 à Hermisdorf en Silésie et est mort le 30 novembre 1921 à Berlin. Élève de Karl Weierstrass, les notes de cours qu'il prit en 1861 contribuèrent à propager les idées de Weierstrass en Italie et en France. Il produit une étude détaillée du théorème de l'application conforme, énoncé quelques années auparavant par Riemann, et montre son application à la transformation conforme d'un carré en cercle. Ces recherches l'amènent à formuler le principe de symétrie qui porte son nom.

Exemple : Sans l'hypothèse $\mathcal{C}^2$ ce résultat n'est plus vrai. Voici un contre-exemple (du à H. Schwarz)

On pose

$$f : (x, y) \mapsto \begin{cases} x^2 \arctan \frac{y}{x} - y^2 \arctan \frac{x}{y} & \text{si } xy \neq 0 \\ 0 & \text{sinon.} \end{cases}$$

On veut calculer les dérivées partielles « croisées » en $(0, 0)$. Pour cela on doit calculer $\frac{\partial f}{\partial x}(0, y)$ et $\frac{\partial f}{\partial y}(x, 0)$

Soit $y \neq 0$, on pose

$$\phi : x \mapsto f(x, y) = \begin{cases} x^2 \arctan \frac{y}{x} - y^2 \arctan \frac{x}{y} & \text{si } x \neq 0 \\ 0 & \text{sinon.} \end{cases}$$

et on cherche à calculer $\phi'(0)$.

Il suffit de faire un développement limité pour voir que

$$\phi(x) = -yx + o(x)$$

donc ϕ est dérivable en 0 et $\phi'(0) = -y$. On a donc $\frac{\partial f}{\partial x}(0, y) = -y$. De même on trouve que $\frac{\partial f}{\partial y}(x, 0) =$

x . De plus $\frac{\partial f}{\partial x}(0, 0) = \frac{\partial f}{\partial y}(0, 0) = 0$. Finalement,

$$\frac{\partial^2 f}{\partial x \partial y}(0, 0) = 1 \neq -1 = \frac{\partial^2 f}{\partial y \partial x}(0, 0).$$

4.2 Opérations algébriques sur les fonctions de classe $\mathcal{C}^k$

Proposition 18.37 (Opérations algébriques sur les fonctions de classe $\mathcal{C}^k$)

L'ensemble des fonctions de classe $\mathcal{C}^k$ vérifie que

1. Il est stable par combinaisons linéaires
2. Il est stable par composition
3. Si f et g sont des fonctions de classe $\mathcal{C}^k$ et que B est une fonction bilinéaire alors $B(f, g)$ est de classe $\mathcal{C}^k$.
4. Si $f_1, \dots, f_p$ sont des fonctions de classe $\mathcal{C}^k$ et que Φ est une fonction multilinéaire alors $\Phi(f_1, \dots, f_p)$ est de classe $\mathcal{C}^k$.

Exercice : Soit f une fonction de classe $\mathcal{C}^2$ sur U . Soit $\mathcal{B} = (e_1, \dots, e_n)$ On appelle laplacien de f et on note Δf la fonction

$$\begin{aligned} \Delta f : U &\mapsto \mathbf{R} \\ a &\mapsto \sum_{i=1}^n \frac{\partial^2 f}{\partial x_i^2}(a) \end{aligned}$$

On passe en coordonnées polaires, on pose $\phi : (\rho, \theta) = (\rho \cos \theta, \rho \sin \theta)$ et $F = f \circ \phi$, c'est-à-dire $F(\rho, \theta) = f(\rho \cos \theta, \rho \sin \theta)$.

On veut montrer qu'alors

$$\Delta f = \frac{\partial^2 f}{\partial x^2} + \frac{\partial^2 f}{\partial y^2} = \frac{\partial^2 F}{\partial \rho^2} + \frac{1}{\rho} \frac{\partial F}{\partial \rho} + \frac{1}{r^2} \frac{\partial^2 F}{\partial \theta^2}.$$

Correction

On a déjà comment exprimer les dérivées partielles de f (par rapport à x et y) par rapport à celles de F (par rapport à ρ et θ).

On utilise la règle de la chaîne on a $J_F = J_f \times J_\phi$ c'est à-dire :

$$\begin{pmatrix} \frac{\partial F}{\partial \rho} & \frac{\partial F}{\partial \theta} \end{pmatrix} = \begin{pmatrix} \frac{\partial f}{\partial x} & \frac{\partial f}{\partial y} \end{pmatrix} \times \begin{pmatrix} \frac{\partial x}{\partial \rho} & \frac{\partial x}{\partial \theta} \\ \frac{\partial y}{\partial \rho} & \frac{\partial y}{\partial \theta} \end{pmatrix} = \begin{pmatrix} \frac{\partial f}{\partial x} & \frac{\partial f}{\partial y} \end{pmatrix} \times \begin{pmatrix} \cos \theta & -\rho \sin \theta \\ \sin \theta & \rho \cos \theta \end{pmatrix}$$

On peut alors inverser J_ϕ . En effet $\det(J_\phi) = \rho(\cos^2 \theta + \sin^2 \theta) = \rho > 0$. Elle est donc inversible et on a alors $J_f = J_F \times (J_\phi)^{-1}$. C'est-à-dire

$$\begin{pmatrix} \frac{\partial f}{\partial x} & \frac{\partial f}{\partial y} \end{pmatrix} = \begin{pmatrix} \frac{\partial F}{\partial \rho} & \frac{\partial F}{\partial \theta} \end{pmatrix} \times \frac{1}{\rho} \begin{pmatrix} \rho \cos \theta & \rho \sin \theta \\ -\sin \theta & \cos \theta \end{pmatrix}$$

On en déduit les formules

$$\frac{\partial f}{\partial x}(a) = \cos \theta \frac{\partial F}{\partial \rho}(b) - \frac{1}{\rho} \sin \theta \frac{\partial F}{\partial \theta}(b)$$

et

$$\frac{\partial f}{\partial y}(a) = \sin \theta \frac{\partial F}{\partial \rho}(b) + \frac{1}{\rho} \cos \theta \frac{\partial F}{\partial \theta}(b)$$

où $a = \phi(b)$. Ceci peut s'écrire $\frac{\partial f}{\partial x} \circ \phi = G_1$ et $\frac{\partial f}{\partial y} \circ \phi = G_2$ où

$$G_1 : (\rho, \theta) \mapsto \cos \theta \frac{\partial F}{\partial \rho}(b) - \frac{1}{\rho} \sin \theta \frac{\partial F}{\partial \theta}(b) \text{ et } G_2 : (\rho, \theta) \mapsto \sin \theta \frac{\partial F}{\partial \rho}(b) + \frac{1}{\rho} \cos \theta \frac{\partial F}{\partial \theta}(b)$$

On peut donc recommencer :

$$\begin{aligned} \frac{\partial^2 f}{\partial x^2} &= \cos \theta \frac{\partial G_1}{\partial \rho}(b) - \frac{1}{\rho} \sin \theta \frac{\partial G_1}{\partial \theta}(b) \\ &= \cos \theta \left[\cos \theta \frac{\partial^2 F}{\partial \rho^2} + \frac{1}{\rho^2} \sin \theta \frac{\partial F}{\partial \theta} - \frac{1}{\rho} \sin \theta \frac{\partial^2 F}{\partial \rho \partial \theta} \right] \\ &\quad - \frac{1}{\rho} \sin \theta \left[-\sin \theta \frac{\partial F}{\partial \rho} + \cos \theta \frac{\partial^2 F}{\partial \theta \partial \rho} - \frac{1}{\rho} \cos \theta \frac{\partial F}{\partial \theta} - \frac{1}{\rho} \sin \theta \frac{\partial^2 F}{\partial \theta^2} \right] \\ &= \cos^2 \theta \frac{\partial^2 F}{\partial \rho^2} + \frac{2 \cos \theta \sin \theta}{\rho^2} \frac{\partial F}{\partial \theta} - \frac{2 \cos \theta \sin \theta}{\rho} \frac{\partial^2 F}{\partial \rho \partial \theta} + \frac{\sin^2 \theta}{\rho} \frac{\partial F}{\partial \rho} + \frac{1}{\rho^2} \sin^2 \theta \frac{\partial^2 F}{\partial \theta^2} \end{aligned}$$

De même, on trouve

$$\frac{\partial^2 f}{\partial y^2} = \sin^2 \theta \frac{\partial^2 F}{\partial \rho^2} - \frac{2 \cos \theta \sin \theta}{\rho^2} \frac{\partial F}{\partial \theta} + \frac{2 \cos \theta \sin \theta}{\rho} \frac{\partial^2 F}{\partial \rho \partial \theta} + \frac{\cos^2 \theta}{\rho} \frac{\partial F}{\partial \rho} + \frac{1}{\rho^2} \cos^2 \theta \frac{\partial^2 F}{\partial \theta^2}$$

Correction

On obtient bien

$$\Delta f = \frac{\partial^2 f}{\partial x^2} + \frac{\partial^2 f}{\partial y^2} = \frac{\partial^2 F}{\partial \rho^2} + \frac{1}{\rho} \frac{\partial F}{\partial \rho} + \frac{1}{\rho^2} \frac{\partial^2 F}{\partial \theta^2}.$$

5 Exemple d'équations aux dérivées partielles

5.1 Premier ordre

★ Commençons par l'équation sur $\mathbf{R}^2$: $\frac{\partial f}{\partial x} = 0$.

On cherche une fonction de classe $\mathcal{C}^1$ sur $\mathbf{R}^2$. Pour tout $y \in \mathbf{R}$, l'application partielle $x \mapsto f(x, y)$ est constante car de dérivée nulle. On en déduit que pour tout y il existe $\varphi(y)$ tel que $\forall x \in \mathbf{R}, f(x, y) = \varphi(y)$. De plus, φ est de classe $\mathcal{C}^1$ car f l'est.

Réciproquement toutes les fonctions de la forme

$$f : (x, y) \mapsto \varphi(y)$$

pour φ de classe $\mathcal{C}^1$ vérifient l'équation.

★ On pose maintenant

$$\frac{\partial f}{\partial x} - 5 \frac{\partial f}{\partial y} = x. \quad (\text{E})$$

On peut remarquer $(x, y) \mapsto \frac{x^2}{2}$ est une solution. Comme c'est une équation linéaire, on est ramené à résoudre

$$\frac{\partial f}{\partial x} - 5 \frac{\partial f}{\partial y} = 0. \quad (\text{H})$$

On va faire un changement de variables en posant $x = u + \alpha v$ et $y = u + \beta v$ pour $\alpha \neq \beta$.

On voit que

$$\begin{aligned} \varphi : \quad \mathbf{R}^2 &\rightarrow \mathbf{R}^2 \\ (u, v) &\mapsto (x, y) \end{aligned}$$

est une bijection de $\mathbf{R}^2$ dans lui-même (la bijection réciproque est $(x, y) \mapsto (u, v) = \left(\frac{\beta x - \alpha y}{\beta - \alpha}, \frac{y - x}{\beta - \alpha} \right)$).

De plus elle est de classe $\mathcal{C}^1$ car polynomiale. Maintenant si f est une fonction de classe $\mathcal{C}^1$ on pose $F = f \circ \varphi$. On sait alors que

$$\begin{aligned} \frac{\partial F}{\partial u} &= \frac{\partial f}{\partial x} \frac{\partial x}{\partial u} + \frac{\partial f}{\partial y} \frac{\partial y}{\partial u} \\ &= \frac{\partial f}{\partial x} + \frac{\partial f}{\partial y} \end{aligned}$$

et

$$\begin{aligned}\frac{\partial F}{\partial v} &= \frac{\partial f}{\partial x} \frac{\partial x}{\partial v} + \frac{\partial f}{\partial y} \frac{\partial y}{\partial v} \\ &= \alpha \frac{\partial f}{\partial x} + \beta \frac{\partial f}{\partial y}.\end{aligned}$$

On voit donc que l'on prend $\alpha = 1$ et $\beta = -5$. Maintenant si f l'équation (H) alors F vérifie $\frac{\partial F}{\partial v}$. On a donc qu'il existe une fonction φ de classe $\mathcal{C}^1$ telle que

$$\forall (u, v) \in \mathbf{R}^2 F(u, v) = \varphi(u) \text{ et donc } \forall (x, y) \in \mathbf{R}^2, f(x, y) = \varphi\left(\frac{5x + y}{6}\right).$$

Donc les solutions de E sont de la forme :

$$(x, y) \mapsto \frac{x^2}{2} + \varphi\left(\frac{5x + y}{6}\right).$$

Réciproquement toutes les fonctions de cette forme vérifient l'équation (E).

5.2 Second ordre

★ Equation des ondes :

On appelle équation des ondes en dimension 1 l'équation

$$\frac{\partial^2 f}{\partial x^2}(x, t) - \frac{1}{c^2} \frac{\partial^2 f}{\partial t^2}(x, t) = 0. \quad (\text{E})$$

C'est en particulier l'équation qui est vérifiée par une corde vibrante. Par exemple si on considère une corde de longueur ℓ . On note $f(x, t)$ la hauteur du point de la corde à l'abscisse x et à l'instant t . La grandeur c est homogène à une vitesse.

On pose le changement de variables $u = x + ct$ et $v = x - ct$

On voit que

$$\begin{aligned}\psi : \quad \mathbf{R}^2 &\rightarrow \mathbf{R}^2 \\ (x, t) &\mapsto (u, v)\end{aligned}$$

est une bijection de $\mathbf{R}^2$ dans lui même. La bijection réciproque est $\varphi : (u, v) \mapsto (x, y) = \left(\frac{u+v}{2}, \frac{u-v}{2c}\right)$. De plus elle est de classe $\mathcal{C}^2$ car polynomiale. On pose encore $F = f \circ \varphi$ et donc $f = F \circ \psi$.

On utilise la règle de la chaîne sur la relation $f = F \circ \psi$ ce qui donne

$$\frac{\partial f}{\partial x}(x, t) = \frac{\partial F}{\partial u}(u, v) + \frac{\partial F}{\partial v}(u, v)$$

et

$$\frac{\partial f}{\partial t}(x, t) = c \frac{\partial F}{\partial u}(u, v) - c \frac{\partial F}{\partial v}(u, v)$$

On peut donc poser

$$G_1 : (u, v) \mapsto \frac{\partial F}{\partial u}(u, v) + \frac{\partial F}{\partial v}(u, v)$$

et

$$G_2 : (u, v) \mapsto c \frac{\partial F}{\partial u}(u, v) - c \frac{\partial F}{\partial v}(u, v)$$

On a donc

$$\frac{\partial f}{\partial x} = G_1 \circ \psi \text{ et } \frac{\partial f}{\partial t} = G_2 \circ \psi$$

On peut donc appliquer la règle de la chaîne qui affirme que

$$\begin{aligned} \frac{\partial^2 f}{\partial x^2}(x, t) &= \frac{\partial}{\partial x} \left(\frac{\partial f}{\partial x} \right) (x, t) \\ &= \frac{\partial G_1}{\partial u}(u, v) \frac{\partial u}{\partial x} + \frac{\partial G_1}{\partial v}(u, v) \frac{\partial v}{\partial x} \\ &= \frac{\partial G_1}{\partial u}(u, v) + \frac{\partial G_1}{\partial v}(u, v) \\ &= \frac{\partial}{\partial u} \left(\frac{\partial F}{\partial u} + \frac{\partial F}{\partial v} \right) + \frac{\partial}{\partial v} \left(\frac{\partial F}{\partial u} + \frac{\partial F}{\partial v} \right) \\ &= \frac{\partial^2 F}{\partial u^2} + 2 \frac{\partial^2 F}{\partial u \partial v} + \frac{\partial^2 F}{\partial v^2} \end{aligned}$$

En utilisant le théorème de Schwarz pour affirmer que $\frac{\partial F}{\partial u \partial v} = \frac{\partial F}{\partial v \partial u}$

De même,

$$\begin{aligned} \frac{\partial^2 f}{\partial t^2}(x, t) &= \frac{\partial}{\partial t} \left(\frac{\partial f}{\partial t} \right) (x, t) \\ &= \frac{\partial G_2}{\partial u}(u, v) \frac{\partial u}{\partial t} + \frac{\partial G_2}{\partial v}(u, v) \frac{\partial v}{\partial t} \\ &= c \frac{\partial G_2}{\partial u}(u, v) - c \frac{\partial G_2}{\partial v}(u, v) \\ &= c \frac{\partial}{\partial u} \left(c \frac{\partial F}{\partial u} - c \frac{\partial F}{\partial v} \right) - c \frac{\partial}{\partial v} \left(c \frac{\partial F}{\partial u} - c \frac{\partial F}{\partial v} \right) \\ &= c^2 \frac{\partial^2 F}{\partial u^2} - 2c^2 \frac{\partial^2 F}{\partial u \partial v} + c^2 \frac{\partial^2 F}{\partial v^2} \end{aligned}$$

On en déduit que si f est une solution alors F vérifie $4 \frac{\partial^2 F}{\partial u \partial v} = 0$.

Il existe alors h de classe $\mathcal{C}^1$ sur $\mathbf{R}$ telle que $\frac{\partial F}{\partial v} = h(v)$ puis $F : (u, v) \mapsto H(v) + G(u)$ où H est une primitive de h .

Finalement,

$$f : (x, y) \mapsto H(x - ct) + G(x + ct).$$

Réciproquement les fonctions de ce type vérifient l'équation.

6 Optimisation

6.1 Point critique et extremum

On veut généraliser le théorème qui dit que si un point intérieur est un extremum d'une fonction f alors la dérivée de f y est nulle.

Définition 18.38

Soit $f : U \rightarrow \mathbf{R}$ une application différentiable. Un point a est dit critique si $df(a) = 0_{\mathcal{L}(E, \mathbf{R})}$ (ce qui revient à $\nabla f(a) = 0$ dans le cas où E est euclidien).

Remarque : Pour montrer qu'un point est critique, il suffit de vérifier que toutes les dérivées partielles sont nulles.

Définition 18.39

Soit $f : U \rightarrow \mathbf{R}$ une application. Un point a est un maximum local (resp. minimum local) s'il existe un voisinage V de a tel que

$$\forall x \in V, f(x) \leq f(a) \quad (\text{resp. } f(x) \geq f(a))$$

Si f est un maximum local ou un minimum local on dit que a est un extremum local.

Proposition 18.40

Soit $f : U \rightarrow \mathbf{R}$ une application différentiable. Si un point a est un extremum local alors a est un point critique.

ATTENTION

- Ce n'est pas une équivalence, par exemple le point $a = (0, 0)$ pour $f : (x, y) \mapsto x^2 - y^2$.
- Cela n'est vrai que si U est un ouvert. Sinon, on peut avoir des problèmes si on est « au bord » de l'ensemble de définition - comme dans le cas des fonctions de $\mathbf{R}$ dans $\mathbf{R}$.

Démonstration : On suppose que a est un extremum local. Soit v un vecteur de E .

L'application $\theta : t \mapsto f(a + tv)$ admet un extremum local en 0. On en déduit que $\theta'(0) = 0$. Or $\theta'(0) = D_v f(a) = df(a).v$.

On en déduit que $df(a)$ est l'application nulle et que a est critique. $\square$

Exemple : On cherche les extremums locaux sur $U = \mathbf{R} \times \mathbf{R}_+^*$ de

$$f : (x, y) \mapsto y(x^2 + (\ln y)^2)$$

On calcule les dérivées partielles :

$$\forall (x, y) \in U, \frac{\partial f}{\partial x} = 2xy \quad \text{et} \quad \frac{\partial f}{\partial y} = x^2 + (\ln y)^2 + 2 \ln y.$$

On en déduit que

$$(x, y) \text{ critique} \iff \begin{cases} 2xy = 0 \\ x^2 + (\ln y)^2 + 2 \ln y = 0 \end{cases} \iff \begin{cases} x = 0 \\ (\ln y)(\ln y + 2) = 0 \end{cases}$$

Les deux points critiques sont $(0, 1)$ et $(0, e^{-2})$.

- On a $f(0, 1) = 0$ et pour tout $(x, y) \in U$, $f(x, y) \geq 0 = f(0, 1)$. On en déduit que $(0, 1)$ est un minimum (local et global)
- On regarde la fonction au voisinage de $(0, e^{-2})$. On étudie

$$\theta(x) = f(x, e^{-2}) = e^{-2}(x^2 + 4)$$

On voit que $\theta'(0) = 0$ et $\theta''(0) > 0$. On voit donc que $(0, e^{-2})$ ne peut être qu'un minimum. On regarde maintenant

$$\psi(y) = f(0, y) = y \ln(y)^2$$

On voit que $\psi'(y) = (\ln y)^2 + 2 \ln y$ et donc $\psi'(y) = 0$. De plus $\psi''(y) = 2 \frac{1 + \ln y}{y}$. En particulier, $\psi''(e^{-2}) < 0$. De ce fait $(0, e^{-2})$ ne peut pas être un minimum.

En conclusion, $(0, e^{-2})$ n'est pas un extremum. Nous verrons plus

6.2 Matrice hessienne

Définition 18.41 (matrice hessienne)

Soit $f : U \rightarrow \mathbf{R}$ une fonction de classe $\mathcal{C}^2$. Soit $a \in U$, on appelle matrice Hessienne de f en a (relativement à une base $\mathcal{B} = (e_1, \dots, e_n)$ de E) et on note $H_f(a) \in \mathcal{M}_n(\mathbf{R})$ la matrice définie par

$$\forall (i, j) \in \llbracket 1 ; n \rrbracket^2, (H_f(a))[i, j] = \partial_{i,j} f(a)$$

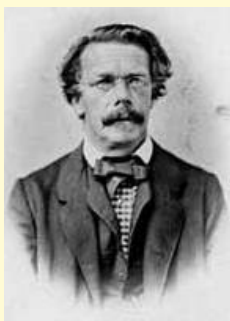
On a donc

$$H_f(a) = \begin{pmatrix} \partial_{1,1} f(a) & \cdots & \partial_{1,n} f(a) \\ \vdots & & \vdots \\ \partial_{n,1} f(a) & \cdots & \partial_{n,n} f(a) \end{pmatrix}$$

Remarques :

1. D'après le théorème de Schwarz, la matrice Hessienne est une matrice symétrique réelle.
2. La majorité du temps, la fonction f sera définie sur un ouvert de $\mathbf{R}^n$. On calculera la matrice hessienne relativement à la base canonique.

Matheux (Ludwig Otto Hesse : 1811 - 1874)



Ludwig Otto Hesse né le 22 avril 1811 à Königsberg et mort le 4 août 1874 à Munich est un mathématicien prussien qui a travaillé sur les invariants algébriques. Il a donné son nom à la matrice hessienne (c'est James Sylvester qui a décidé de nommer la matrice ainsi).

Exemple : On considère la fonction $f : \mathbf{R}^2 \rightarrow \mathbf{R}$ définie par $f : (x, y) \mapsto (y - x)^3 + 6xy$.

On calcule les dérivées partielles :

$$\forall (x, y) \in \mathbf{R}^2, \frac{\partial f}{\partial x} = -3(y - x)^2 + 6y \text{ et } \frac{\partial f}{\partial y} = 3(y - x)^2 + 6x$$

On peut alors calculer les dérivées secondes puis la matrice hessienne.

$$\forall (x, y) \in \mathbf{R}^2, H_f(x) = \begin{pmatrix} 6(y-x) & -6(y-x+1) \\ -6(y-x+1) & 6(y-x) \end{pmatrix}$$

Théorème 18.42 (Formule de Taylor-Young à l'ordre 2)

Soit $f : U \rightarrow \mathbf{R}$ une fonction de classe $\mathcal{C}^2$. Soit $a \in U$ et $h \in E$ tel que $a + h \in U$, on note encore $h \in \mathcal{M}_{n,1}(\mathbf{R})$ la matrice des coordonnées de h . On a alors

$$f(a+h) = f(a) + df(a).h + \frac{1}{2}h^\top H_f(a)h + o(\|h\|^2)$$

Remarques :

1. La majorité du temps la fonction f sera définie sur un ouvert U de $\mathbf{R}^n$. Dans ce cas, les coordonnées de h et la matrice hessienne seront calculées relativement à la base canonique mais on peut tout aussi bien travailler dans une autre base à condition de prendre la même base pour calculer les coordonnées de h et pour calculer la matrice hessienne.
2. Dans le cas d'une fonction définie sur un ouvert de $\mathbf{R}^n$ on peut travailler en coordonnées pour obtenir

$$f(a_1 + h_1, \dots, a_n + h_n) = f(a_1, \dots, a_n) + \sum_{i=1}^n h_i \partial_i f(a) + \frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n h_i h_j \partial_{ij} f(a) + o(\|h\|^2)$$

Démonstration : (Non exigible) Pour alléger les notations on travaille avec une fonction f définie sur un ouvert de $\mathbf{R}^n$.

Soit $h = (h_1, \dots, h_n) \in \mathbf{R}^n$ tel que $a+h \in U$. D'après le théorème fondamental de l'analyse, en posant $\gamma : [0, 1] \rightarrow \mathbf{R}^n$ défini par $\gamma : t \mapsto a + th$,

$$f(a+h) - f(a) = \int_0^1 df(a+th).h dt = \int_0^1 \sum_{i=1}^n \partial_i f(a+th) h_i dt$$

Soit $i \in \llbracket 1; n \rrbracket$, on peut utiliser la formule de Taylor-Young à l'ordre 1 sur la fonction $x \mapsto \partial_i f$ pour évaluer $\partial_i f(a+th)$. On a

$$\partial_i f(a+th) = \partial_i f(a) + d(\partial_i f)(a).th + \|th\| \varepsilon_i(th) = \partial_i f(a) + \sum_{j=1}^n \partial_{ji} f(a) th_j + \|th\| \varepsilon_i(th)$$

ou ε_i est une fonction qui tend vers 0 en $0_{\mathbf{R}^n}$. On en déduit que

$$\begin{aligned} f(a+h) &= f(a) + \sum_{i=1}^n \int_0^1 h_i \left(\partial_i f(a) + \sum_{j=1}^n \partial_{ji} f(a) th_j + \|th\| \varepsilon_i(th) \right) dt \\ &= f(a) + df(a).h + \frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n h_i h_j \partial_{ji} f(a) + \sum_{i=1}^n \sum_{j=1}^n h_i \|h\| \int_0^1 t \varepsilon_i(th) dt \end{aligned}$$

En utilisant la caractérisation séquentielle et le théorème de convergence dominée, on obtient que

$$\lim_{h \rightarrow 0} \int_0^1 t \varepsilon_i(th) dt = \int_0^1 0 dt = 0$$

On a donc que

$$f(x+h) = f(x) + df(x).h + \frac{1}{2}h^\top H_f(x)h + o(\|h\|^2)$$

□

Corollaire 18.43

Soit $f : U \rightarrow \mathbf{R}$ une fonction de classe $\mathcal{C}^2$. Soit $a \in U$. Si f admet un minimum local en a alors a est un point critique et $H_f(a) \in \mathcal{S}_n^+(\mathbf{R})$.

Remarque : On peut adapter cet énoncé au cas des maximums locaux. En effet si a est un maximum local alors $df(a) = 0$ et $-H_f(a) \in \mathcal{S}_n^+(\mathbf{R})$.

Démonstration : On travaille avec une fonction définie sur $\mathbf{R}^n$. On suppose que a est un minimum local. On sait déjà que cela implique que $df(a) = 0$. On en déduit en utilisant la formule de Taylor-Young à l'ordre 2 que

$$f(a+h) = f(a) + \frac{1}{2}h^\top H_f(a)h + o(\|h\|^2)$$

En particulier pour $h \in \mathbf{R}^n \setminus \{0\}$ et t petit,

$$0 \geq f(a+th) - f(a) = t^2 \frac{1}{2}h^\top H_f(a)h + t^2 \|h\| \varepsilon(th)$$

Cela montre que, pour t petit,

$$\frac{1}{2}h^\top H_f(a)h + \|h\| \varepsilon(th) \geq 0$$

Il suffit alors de faire tendre t vers 0 pour obtenir que $h^\top H_f(a)h \geq 0$. □

Remarque : Cela n'est pas une équivalence. Il suffit de prendre $f : (x, y) \mapsto x^3$ pour voir que $a = (0, 0)$ vérifie $df(a) = 0$, $H_f(a) = 0 \in \mathcal{S}_2^+(\mathbf{R})$ et pourtant le point $a = (0, 0)$ n'est pas un minimum.

Corollaire 18.44

Soit $f : U \rightarrow \mathbf{R}$ une fonction de classe $\mathcal{C}^2$. Soit $a \in U$. Si a est un point critique et $H_f(a) \in \mathcal{S}_n^{++}(\mathbf{R})$ alors a est un minimum local (strict) de la fonction.

Démonstration : Il suffit là encore d'utiliser la formule de Taylor-Young. On considère h un vecteur unitaire de E . En reprenant les calculs précédents en supposant que $H_f(a) \in \mathcal{S}_n^{++}(\mathbf{R})$, pour t petit,

$$f(a+th) = f(a) + t^2 (h^\top H_f(a)h + \varepsilon(th))$$

On sait que la fonction $q : h \mapsto h^\top H_f(a)h$ est continue. D'après le théorème des bornes atteintes, elle atteint son minimum sur la boule unité qui est compacte. Comme $H_f(a) \in \mathcal{S}_n^{++}(\mathbf{R})$, ce minimum que l'on peut noter α est strictement positif. Comme la limite de ε en $0_{\mathbf{R}^n}$ est nulle, il existe $\delta > 0$ tel que pour $x \in B(0, \delta)$, $|\varepsilon(x)| \leq \frac{\alpha}{2}$. Cela montre que pour $t \in]-\delta, \delta[$ non nul, $f(a+th) - f(a) > 0$. □

Corollaire 18.45 (Cas des fonctions de deux variables)

Soit $f : U \rightarrow \mathbf{R}$ une fonction de classe $\mathcal{C}^2$ où U est un ouvert de $\mathbf{R}^2$. Soit $a \in U$ un point critique.

1. Si $\det(H_f(a)) > 0$, a est un extremum local. Si $\text{tr}(H_f(a)) < 0$ c'est un maximum et si $\text{tr}(H_f(a)) > 0$ c'est un minimum.
2. Si $\det(H_f(a)) < 0$ la fonction n'admet pas d'extremum local en a . C'est un point col.
3. Si $\det(H_f(a)) = 0$, c'est indéterminé

Démonstration : Notons $H_f(a) = \begin{pmatrix} r & s \\ s & t \end{pmatrix}$. Pour un vecteur $h = (x, y)$, on veut connaître le signe de

$$q(h) = h^\top H_f(a) h = rx^2 + 2sxy + ty^2$$

Rappelons que l'on a vu dans le chapitre sur les endomorphismes autoadjoints que la matrice $H_f(a)$ appartient à $\mathcal{S}_n^{++}(\mathbf{R})$ si et seulement si $r > 0$ et $\det(H_f(a)) = rt - s^2 > 0$.

1. Si $\det(H_f(a)) > 0$, en particulier r et t sont du même signe (et ne sont pas nuls). Ils sont donc du signe de $\text{tr}(H_f(a))$. On en déduit que si $\text{tr}(H_f(a)) > 0$ alors $H_f(a) \in \mathcal{S}_2^{++}(\mathbf{R})$ ce qui montre que a est un minimum.

Si $\text{tr}(H_f(a)) < 0$ alors, de même, $-H_f(a) \in \mathcal{S}_2^{++}(\mathbf{R})$ ce qui montre que a est un maximum.

2. Procédons par contraposée. Si a est un minimum $H_f(a) \in \mathcal{S}_n^+(\mathbf{R})$ et donc $\det(H_f(a)) \geq 0$. De même si a est un maximum $-H_f(a) \in \mathcal{S}_n^+(\mathbf{R})$ et donc $\det(-H_f(a)) = \det(H_f(a)) \geq 0$.
3. Dans le cas où $\det(H_f(a)) = 0$ on peut avoir un minimum par exemple si $f : (x, y) \mapsto x^2$, un maximum par exemple si $f : (x, y) \mapsto -x^2$ ou un point qui n'est pas un extremum par exemple $f : (x, y) \mapsto x^3$.

□

Exercice : On pose $f : (x, y) \mapsto (y - x)^3 + 6xy$.

1. Déterminer les points critiques
2. Déterminer si ce sont des extremums.
3. On pose $K = \{(x, y) \in [-1, 1]^2, y \geq x\}$. Déterminer les extremas de f sur K .

Correction

1. La fonction est différentiable. On calcule les dérivées partielles :

$$\forall (x, y) \in \mathbb{R}^2, \frac{\partial f}{\partial x} = -3(y-x)^2 + 6y \text{ et } \frac{\partial f}{\partial y} = 3(y-x)^2 + 6x.$$

On en déduit que (x, y) est critique si et seulement si

$$\begin{cases} -3(y-x)^2 + 6y = 0 \\ 3(y-x)^2 + 6x = 0 \end{cases} \iff \begin{cases} -3(y-x)^2 + 6y = 0 \\ x+y = 0 \end{cases} \iff \begin{cases} 4x^2 + 2x = 0 \\ y = -x \end{cases}$$

Les points critiques sont donc $a_1 = (0, 0)$ et $a_2 = (-\frac{1}{2}, \frac{1}{2})$.

2. On étudie le point a_1 .

On a

$$\forall (x, y) \in \mathbb{R}^2, \partial_{x,x}f(x, y) = 6(y-x), \partial_{x,y}f(x, y) = \partial_{y,x}f(x, y) = -6(y-x)+6 \text{ et } \partial_{y,y}f(x, y) = 6(y-x)$$

On en déduit que

$$H_f(a_1) = \begin{pmatrix} 0 & 6 \\ 6 & 0 \end{pmatrix}$$

Comme $\det(H_f(a_1)) = -36 < 0$, a_1 n'est pas un extremum ; c'est un point col.

On étudie alors a_2 . Cette fois

$$H_f(a_2) = \begin{pmatrix} 6 & 0 \\ 0 & 6 \end{pmatrix}$$

On a $\det(H_f(a_2)) > 0$. Le point a_2 est un extremum. De plus $\text{tr}(H_f(a_2)) = 12 > 0$, c'est donc un minimum.

3. L'ensemble K est un compact. Comme f est continue, la fonction f admet un minimum et un maximum global sur K . On a vu que $a_2 \in K$ est un minimum local. Il faut maintenant regarder ce qui se passe sur le bord de K .

On se ramène à trois segments.

On regarde $\theta_1 : t \mapsto f(t, 1) = (1-t)^3 + 6t$ sur $[-1, 1]$, $\theta_2 : t \mapsto f(-1, t) = (t+1)^3 - 6t$ sur $[-1, 1]$ et $\theta_3 : t \mapsto f(t, t) = 6t^2$ sur $[-1, 1]$. En établissant les tableaux de variations, on voit que le maximum est atteint en $(1, 1)$ et $(-1, -1)$ où la fonction vaut 6. Le minimum est a_2 car les valeurs minimales sur les bords sont positives.

on voit que le maximum est atteint en $(1, 1)$ et $(-1, -1)$ où la fonction vaut 6. Le minimum est a_2 car les valeurs minimales sur les bords sont positives.

6.3 Optimisation sous contrainte

Nous allons maintenant nous intéresser aux extremums d'une fonction f en nous limitant à un sous-ensemble fermé.

Théorème 18.46 (Optimisation sous contrainte)

Soit f, g deux fonctions de classe $\mathcal{C}^1$ définies sur un ouvert U de E à valeurs dans $\mathbf{R}$. On note $X = g^{-1}(\{0\})$ l'ensemble des zéros de g .

Si $x \in X$, $dg(x) \neq 0$ et si la restriction de f à X admet un extremum local en x alors $df(x)$ et $dg(x)$ sont colinéaires.

En particulier, dans le cas où E est euclidien, on obtient que $\text{Grad } f(x)$ est colinéaire à $\text{Grad } g(x)$.

Démonstration : On sait avec nos hypothèses que $T_x X = \text{Ker}(dg(x))$.

Soit $u \in T_x X$, il existe un arc γ à valeurs dans X tel que $\gamma(0) = x$ et $\gamma'(0) = u$. La fonction $f \circ \gamma : t \mapsto f(\gamma(t))$ atteint un extremum en 0 donc sa dérivée est nulle. On a

$$0 = (f \circ \gamma)'(0) = df(x) \cdot \gamma'(0) = df(x) \cdot u$$

On en déduit que $u \in \text{Ker}(df(x))$.

On a montré que $df(x)$ et $dg(x)$ sont deux formes linéaires telles que $\text{Ker}(dg(x)) \subset \text{Ker}(df(x))$, cela implique qu'il existe $\lambda \in \mathbf{R}$ telle que $df(x) = \lambda dg(x)$. $\square$

Exemple : Soit $n \geq 3$. On considère un polygone convexe à n sommets inscrit dans le cercle unité. On veut maximiser son périmètre. On pose donc $0 = x_1 < x_2 < \dots < x_n < 2\pi$ et on note pour $k \in \llbracket 1; n \rrbracket$ les points A_k d'affixe e^{ix_k} et on considère le polygone de sommets $A_1, \dots, A_n$.

Pour $k \in \llbracket 1; n-1 \rrbracket$, la longueur du segment $[A_k, A_{k+1}]$ est $|e^{ix_{k+1}} - e^{ix_k}|$. En factorisant par l'angle moitié on obtient que

$$|e^{ix_{k+1}} - e^{ix_k}| = 2 \sin\left(\frac{x_{k+1} - x_k}{2}\right)$$

De même pour $k = n$, la longueur du segment $[A_n, A_1]$ est $2 \sin\left(\frac{2\pi + x_1 - x_n}{2}\right)$

Posons pour simplifier pour $k \in \llbracket 1; n-1 \rrbracket$, $t_k = \frac{x_{k+1} - x_k}{2}$ et $t_n = \frac{2\pi + x_1 - x_n}{2}$ de sorte que le périmètre du polygone est donné par

$$f : (t_1, \dots, t_n) \mapsto 2 \sum_{k=1}^n \sin(t_k)$$

On remarque de plus que

$$\sum_{k=1}^n t_k = \frac{1}{2} \left(\sum_{k=1}^{n-1} (x_{k+1} - x_k) + 2\pi + x_1 - x_n \right) = \pi$$

On essaye donc de maximiser sur $(\mathbf{R}_+^*)^n$ la fonction f sous la contrainte $g(t_1, \dots, t_n) = 0$ où $g : (t_1, \dots, t_n) \mapsto \sum_{k=1}^n t_k - \pi$.

Soit $t = (t_1, \dots, t_n)$ un éventuel maximum local, on calcule $\nabla g(t) = \begin{pmatrix} 1 \\ \vdots \\ 1 \end{pmatrix} \neq 0$ ainsi que $\nabla f(t) =$

$2 \begin{pmatrix} \cos(t_1) \\ \vdots \\ \cos(t_n) \end{pmatrix}$. On en déduit que $\nabla f(t)$ est colinéaire à $\nabla g(t)$ ce qui implique que

$$\cos(t_1) = \dots = \cos(t_n)$$

comme $t_1, \dots, t_n$ appartiennent à $[0, \pi]$ et que $\cos$ est injective sur $[0, \pi]$ on obtient que nécessairement $t_1 = \dots = t_n = \frac{\pi}{n}$ et par conséquence que pour tout $k \in \llbracket 1 ; n \rrbracket$, $x_k = \frac{2k\pi}{n}$.

On a donc montré que s'il existe un maximum local de f sur $(\mathbf{R}_+^*)^n$ sous la contrainte $g(t) = 0$ alors $t_1 = \dots = t_n = \frac{\pi}{n}$.

Il reste à montrer l'existence de ce maximum. On considère donc f sur $(\mathbf{R}_+)^n$. Dans ce cas

$$X = \{(t_1, \dots, t_n) \in [0, \pi]^n \mid g(t_1, \dots, t_n) = 0\}$$

On voit donc que X est un compact ce qui montre que la fonction f qui est continue admet un maximum global sur X .

De plus, on voit que le maximum $(t_1, \dots, t_n)$ appartient à $t_1 \neq 0, \dots, t_n \neq 0$. En effet si on a, par exemple, $t_1 = 0$ et $t_2 > 0$, en remplaçant t_1 et t_2 par $\frac{t_2}{2}$ et $\frac{t_2}{2}$ on augmente strictement le périmètre.

Algèbre générale

1	Groupes	504
1.1	Généralités	505
1.2	Sous-groupe	506
1.3	Morphismes de groupes	509
2	Groupes monogènes	512
2.1	Groupe $(\mathbb{Z}/n\mathbb{Z}, +)$	512
2.2	Générateurs, groupes monogènes et groupes cycliques	515
3	Ordre d'un élément	517
3.1	Définition	517
3.2	Théorème de Lagrange	519
4	Anneaux	522
4.1	Définitions	522
4.2	Morphisme d'anneaux	525
4.3	Anneaux intègres et corps	527
5	L'anneau $(\mathbb{Z}/n\mathbb{Z}, +, \times)$	527
5.1	Définition	527
5.2	Théorème Chinois	532
5.3	Indicatrice d'Euler	536
6	Factorisation des polynômes de $K[X]$	539
6.1	Polynômes irréductibles	539
6.2	Décomposition en produit de facteurs irréductibles	540

Dans ce chapitre nous allons revoir et compléter les notions d'algèbre générale vues en première année.

1 Groupes

1.1 Généralités

Définition 19.1 (Groupe)

Une groupe est un couple formé d'un ensemble non vide et d'une loi de composition interne telle que

- La loi est associative
- Il existe un élément neutre
- Tout élément du groupe est symétrisable

Note

1. De manière générale, on notera multiplicativement les groupes. La loi sera notée $\cdot$ ou $*$. L'élément neutre sera souvent noté e (ou 1). Pour tout g de G , son symétrique sera noté g^{-1} , il vérifie que $gg^{-1} = g^{-1}g = e$. On parle alors d'inverse de g .
2. Si le groupe est abélien, c'est-à-dire commutatif, on notera des fois $+$ la loi. Dans ce cas, l'élément neutre sera noté 0 et le symétrique de g sera noté $-g$. On parle alors d'opposé.

Exemples :

1. $(\mathbb{Z}, +)$, $(\mathbb{Q}, +)$, $(\mathbb{R}, +)$ et $(\mathbb{C}, +)$ sont des groupes
2. $(\mathbb{R}^*, \times)$ est un groupe. De manière générale pour tout anneau A , $(A^\times, \times)$ est un groupe où $A^\times$ est l'ensemble des éléments inversibles.
3. Pour tout entier n , on note $\mathfrak{S}_n$ l'ensemble des bijections de $\{1, 2, \dots, n\}$ dans lui-même. Alors $(\mathfrak{S}_n, \circ)$ est un groupe.

De manière générale si X est un ensemble, l'ensemble des bijections de cet ensemble dans lui-même est un groupe pour la composition.

4. Pour tout entier $n \geq 1$, $\text{GL}_n(\mathbb{K})$ est un groupe (non abélien si $n \geq 2$).
5. Pour tout entier $n \geq 1$, $\mathcal{O}_n(\mathbb{R})$ est un groupe.

Proposition-Définition 19.2 (Produit de groupes)

Soit $G_1, \dots, G_p$ des groupes, le produit $G = G_1 \times \dots \times G_p$ est muni d'une structure de groupe en posant

$$\forall ((g_1, \dots, g_p), (h_1, \dots, h_p)) \in G^2, (g_1, \dots, g_p) * (h_1, \dots, h_p) = (g_1 * h_1, \dots, g_p * h_p)$$

L'élément neutre de G étant

$$e_G = (e_1, \dots, e_p)$$

où pour tout $i \in \llbracket 1; p \rrbracket$ e_i est l'élément neutre de G_i . De même, le symétrique de $(g_1, \dots, g_p)$ est $(g_1^{-1}, \dots, g_p^{-1})$.

On dit que G est le produit direct des groupes $G_1, \dots, G_p$.

Exemple : On en déduit que $(\mathbb{Z}^p, +)$ est un groupe.

1.2 Sous-groupe

Définition 19.3

Soit $(G, *)$ un groupe, une partie H de G est un sous-groupe s'il elle est stable par la loi $*$ et si $(H, *)$ est un groupe.

Théorème 19.4 (Caractérisation des sous-groupes)

Avec les notations précédentes, H est un sous-groupe de G si et seulement si

- i) H n'est pas vide
- ii) Pour tout g, g' dans H , $g(g')^{-1} \in H$

Démonstration :

- $\Rightarrow$ Si H est un sous-groupe alors H n'est pas vide et si g, g' alors $g(g')^{-1} \in H$.
- $\Leftarrow$ On suppose que H n'est pas vide.

Soit $g \in H$, (qui existe car $H \neq \emptyset$) la formule ii) affirme alors que $e = gg^{-1} \in H$.

De plus pour tout $g \in H$ on a donc $g^{-1} = eg^{-1} \in H$.

Finalement, pour tout g, g' dans H , on applique ii) à g et $(g')^{-1}$ (ce dernier appartenant à H d'après le point précédent) pour avoir que

$$gg' = g((g')^{-1})^{-1} \in H$$

Donc H est un sous-groupe.

□

Exercice : Soit $\omega \in \mathbb{C}$. On pose $H = \{a + \omega b \in \mathbb{C} \mid (a, b) \in \mathbb{Z}^2\}$. Montrer que H est un sous-groupe de $(\mathbb{C}, +)$.

Correction

On vérifie les axiomes.

- On voit que H n'est pas vide car $0 = 0 + 0\omega \in H$.
- Soit $u = a + b\omega$ et $v = a' + b'\omega$ deux éléments de H où a, a', b et b' sont des entiers,

$$u - v = (a - a') + (b - b')\omega$$

appartient à H car $a - a' \in \mathbb{Z}$ et $b - b' \in \mathbb{Z}$.

Finalement, H est bien un sous-groupe de $(\mathbb{C}, +)$.

Proposition 19.5

Soit $(G, *)$ un groupe. Si $(H_i)_{i \in I}$ est une famille de sous-groupe de G alors $\bigcap_{i \in I} H_i$ est un sous-groupe.

Démonstration : Notons e l'élément neutre de G .

- Comme pour tout $i \in I$, H_i est un sous-groupe de G , $e \in H_i$. De ce fait, $e \in \bigcap_{i \in I} H_i$ et donc $\bigcap_{i \in I} H_i \neq \emptyset$.
- Soit g, g' dans $\bigcap_{i \in I} H_i$, alors pour tout $i \in I$, $g \in H_i$ et $g' \in H_i$ donc $g(g')^{-1} \in H_i$. Cela implique que $g(g')^{-1} \in \bigcap_{i \in I} H_i$.

On a bien montré que $\bigcap_{i \in I} H_i$ est un sous-groupe de G . □

Proposition-Définition 19.6

Soit $(G, *)$ un groupe et A une partie de G . Il existe un plus petit sous-groupe de G contenant A . On l'appelle le sous-groupe engendré par A et se note $\langle A \rangle$.

Démonstration : On considère l'ensemble X des sous-groupes de G contenant A . Ce n'est pas un ensemble vide car $G \in X$. On considère

$$H = \bigcap_{K \in X} K$$

C'est bien un sous-groupe de G comme intersection de sous-groupes de G .

De plus comme pour tout $K \in X$, $A \subset K$ alors $A \subset H$.

Pour finir, tout sous-groupe K contenant A contient aussi H par construction.

On a bien montré que H était un sous-groupe contenant A et qu'il était le plus petit. □

ATTENTION

Cette construction est agréable d'un point de vue théorique. Cependant elle ne peut pas être utilisée pour déterminer réellement $\langle A \rangle$.

Notation : Soit $g \in G$, on notera $\langle g \rangle$ pour $\langle \{g\} \rangle$ le plus petit sous-groupe de G contenant g .

Exemples :

1. Soit G un groupe $\langle \emptyset \rangle = \{e\}$.
2. Soit G un groupe et $g \in G$, $\langle g \rangle = \{g^n \mid n \in \mathbb{Z}\}$.

En effet, si on pose $H = \{g^n \mid n \in \mathbb{Z}\}$, on peut vérifier que H est un sous-groupe de G et que $g \in H$.

- $H \neq \emptyset$ car $e = g^0 \in H$
- Soit g^p et g^q deux éléments de H , $g^p(g^q)^{-1} = g^{p-q} \in H$.

De ce fait $\langle g \rangle \subset H$

Réciproquement :

- comme $\langle g \rangle$ est un sous-groupe, il contient g^0 .
- Il contient g , il contient donc (par récurrence) les g^n pour $n \in \mathbb{N}^*$.
- Il contient aussi g^{-1} et donc $(g^{-1})^n$ pour $n \geq 0$.

Au final $\boxed{H \subset \langle g \rangle}$ et donc $\boxed{H = \langle g \rangle}$

En particulier pour $G = \mathbb{Z}$ et $g = m$, on a que $\langle m \rangle = m\mathbb{Z}$.

3. On considère G le groupe des bijections du plan dans lui même (ou le groupe des isométries).

On pose r la rotation d'angle $\frac{2\pi}{5}$ et s la symétrie par rapport à (Ox) . On peut aussi imaginer cela comme des bijections de $\mathbb{C}$ dans $\mathbb{C}$. Dans ce cas

$$r : z \mapsto e^{2i\pi/5}z \text{ et } s : z \mapsto \bar{z}.$$

Le groupe $H = \langle r, s \rangle$ est formé de toutes les termes de la formes

$$r^{\alpha_1} s^{\beta_1} r^{\alpha_2} s^{\beta_2} \dots r^{\alpha_p} s^{\beta_p}$$

où les α_i et β_i appartiennent à $\mathbb{Z}$.

On remarque alors que $r^5 = e$ et $s^2 = e$ donc $r^{-1} = r^4$ et $s^{-1} = s$. Soit $\alpha_i \in \mathbb{Z}$, on peut écrire la division euclidienne par 5,

$$\alpha_i = 5q_i + u_i \text{ où } u_i \in \llbracket 0 ; 4 \rrbracket$$

On a $r^{\alpha_i} = r^{u_i}$.

De ce fait, on peut supposer que les α_i sont dans $\llbracket 0 ; 4 \rrbracket$.

On peut faire de même en faisant la division euclidienne de β_i par 2; on peut donc supposer que tous les β_i appartiennent à $\{0, 1\}$.

Maintenant, si on regarde sr^k , avec la représentation par les applications de $\mathbb{C}$ dans $\mathbb{C}$ on a pour tout $z \in \mathbb{C}$

$$s \circ r^k(z) = s(r^k(z)) = \overline{e^{2ik\pi/5}z} = e^{-2ik\pi/5}\bar{z} = r^{-k}(s(z)).$$

C'est-à-dire que $sr^k = r^{-k}s$. On peut donc faire passer tous les r^{α_i} à gauche. Par exemple

$$r^3 s^1 r^2 s^0 r^4 = r^3 sr^6 = r^3 sr = r^3 r^{-1} s = r^2 s$$

On en déduit finalement que H n'a que 10 éléments

$$H = \{e, r, r^2, r^3, r^4, s, sr, sr^2, sr^3, sr^4\}$$

Ces 10 éléments sont distincts. C'est de groupe diédral D_5 des isométries qui préservent un pentagone régulier.

Proposition 19.7 (Sous-groupes de $(\mathbb{Z}, +)$)

Les sous-groupes de $(\mathbb{Z}, +)$ sont les parties de la forme $m\mathbb{Z}$.

Démonstration :

- Si $H = \{0\}$ alors $H = 0\mathbb{Z}$.
- Si $H \neq \{0\}$ alors il existe $n \in H$ avec $n \neq 0$, comme $|n| \in H$ (si $n \in H$ alors $-n \in H$) on en déduit que $H \cap \mathbb{N}^*$ n'est pas vide. Comme c'est une partie non vide de $\mathbb{N}$, on pose

$$m = \min(H \cap \mathbb{N}^*) \in \mathbb{N}^*$$

Comme H est un groupe qui contient m , alors $m\mathbb{Z} = \langle m \rangle \subset H$.

Réciproquement supposons par l'absurde que $H \not\subset m\mathbb{Z}$. il existe donc un élément $n \in H$ qui ne soit pas un multiple de m . On fait alors la division euclidienne de n par m :

$$n = mq + r \text{ où } r \in \llbracket 1 ; m - 1 \rrbracket$$

On en déduit que $r = n - mq \in m\mathbb{Z}$ ce qui est absurde car $0 < r < m$.

Finalement, on a bien $H = m\mathbb{Z}$.

□

Exercice : Soit H un sous-groupe de $(\mathbb{R}, +)$. Montrer que H est dense dans $\mathbb{R}$ ou qu'il existe $\alpha \in \mathbb{R}$ tel que $H = \alpha\mathbb{Z}$.

1.3 Morphismes de groupes

Si on considère deux groupes $(G, \top)$ et $(G', *)$. Quand on considère une application f de G dans G' on voudra souvent que cette application conserve la structure. Précisément

Définition 19.8

Soit $(G, \top)$ et $(G', *)$ deux groupes. On appelle *morphisme de groupe* de $(G, \top)$ dans $(G', *)$ toute application f de G dans G' telle que

$$\forall (g_1, g_2) \in G^2, f(g_1 \top g_2) = f(g_1) * f(g_2).$$

Remarque : Dans la plupart des cas, si les structures sur G et G' sont évidentes, on omettra le terme groupe pour parler simplement de morphisme de G dans G' .

Exemples :

1. Soit G un groupe. L'application identité Id_G de G dans lui-même est un morphisme.
2. Soit G et G' deux groupes dont les éléments neutres sont e et e' . L'application constante égale à e' est un morphisme.

C'est la seule application constante qui vérifie cela car nous verrons plus loin que si f est un morphisme alors $f(e) = e'$.

3. De $(\mathbb{Z}, +)$ dans lui-même, les applications $n \mapsto a.n$ sont des morphismes.
4. L'application $\ln$ est un morphisme de $(\mathbb{R}_+^*, \times)$ dans $(\mathbb{R}, +)$ et inversement pour $\exp$.
5. L'application $\theta \mapsto e^{i\theta}$ de $(\mathbb{R}, +)$ dans $(\mathcal{U}, \times)$ où $\mathcal{U} = \{z \in \mathbb{C} \mid |z| = 1\}$ est un morphisme.
6. On sait que toute permutation $\sigma \in \mathfrak{S}_n$ peut s'écrire comme un produit de transpositions. Si σ s'écrit de deux manières comme un produit de transpositions

$$\sigma = \tau_1 \circ \tau_2 \circ \cdots \circ \tau_p \text{ et } \sigma = \gamma_1 \circ \gamma_2 \circ \cdots \circ \gamma_q$$

alors p et q ont la même parité. On appelle alors la signature de la permutation σ et on note $\varepsilon(\sigma)$ pour $(-1)^{I(\sigma)}$ où $I(\sigma)$ est le nombre de transpositions dans une décomposition de σ en un produit de transpositions.

La signature $\varepsilon : \mathfrak{S}_n \rightarrow \{\pm 1\}$ est un morphisme de groupe

7. Le déterminant est un morphisme de groupe de $(\text{GL}_n(\mathbb{K}), \times)$ dans $(\mathbb{K}^*, \times)$.

Terminologie :

1. Un morphisme de groupe bijectif s'appelle un isomorphisme.

2. Un morphisme de groupe bijectif dont l'ensemble de départ et d'arrivé sont les mêmes s'appelle un automorphisme.

Proposition 19.9

Soit $(G, \top)$ et $(G', *)$ deux groupes et f un morphisme de G dans G' . On note e et e' les éléments neutres respectifs.

1. On a $f(e) = e'$.
2. Pour tout g de G , $f(g^{-1}) = (f(g))^{-1}$.
3. Pour tout g de G et tout n de $\mathbb{Z}$, $f(g^n) = (f(g))^n$.

Démonstration :

1. On a $f(e \top e) = f(e)$. Or $f(e \top e) = f(e) * f(e)$. En multipliant par $(f(e))^{-1}$ on obtient

$$(f(e))^{-1} * f(e) = (f(e))^{-1} * f(e) * f(e) \iff e' = f(e).$$

2. Soit $g \in G$, $f(g^{-1}) * f(g) = f(g^{-1}g) = f(e) = e'$ et pareil pour $f(g) * f(g^{-1}) = e'$.

On en déduit que $f(g^{-1}) = (f(g))^{-1}$.

3. Par récurrence on obtient la propriété pour $n \in \mathbb{N}$, il suffit ensuite d'utiliser le résultat précédent pour obtenir la propriété pour $n < 0$.

□

Proposition 19.10

Soit $(G, \top)$ et $(G', *)$ deux groupes et f un morphisme de G dans G' .

1. Si H est un sous-groupe de G alors $f(H) = \{g' \in G' \mid \exists h \in H, g' = f(h)\}$ est un sous-groupe de G' .
2. Si H' est un sous-groupe de G' alors $f^{-1}(H') = \{g \in G \mid f(g) \in H'\}$ est un sous-groupe de G .

Démonstration :

1. — Comme $e \in H$, $e' = f(e) \in f(H)$ et donc $f(H) \neq \emptyset$.
— Soit g_1, g_2 dans $f(H)$. Il existe h_1, h_2 dans H tels que $f(h_1) = g_1$ et $f(h_2) = g_2$.
On a alors $g_1 g_2^{-1} = f(h_1) * (f(h_2))^{-1} = f(h_1) * f(h_2^{-1}) = f(h_1 \top h_2^{-1})$. Comme $h_1 \top h_2^{-1} \in H$ (puisque H est un sous-groupe de G), $g_1 g_2^{-1} \in f(H)$.

On a montré que $f(H)$ est un sous-groupe de G' .

2. — Comme $f(e) = e' \in H'$. Cela implique $e \in f^{-1}(H')$ et donc $f^{-1}(H') \neq \emptyset$.
— Soit g_1, g_2 dans $f^{-1}(H')$.

$$f(g_1 \top g_2^{-1}) = f(g_1) * f(g_2^{-1}) = f(g_1) * f(g_2)^{-1}$$

Comme H' est un sous-groupe et que $f(g_1)$ et $f(g_2)$ sont dans H' alors $f(g_1 \top g_2^{-1}) = f(g_1) * f(g_2)^{-1} \in H'$.

Cela montre que $g_1 \top g_2^{-1} \in f^{-1}(H')$.

On a montré que $f^{-1}(H)$ est un sous-groupe de G .

Exemples :

1. Soit $f : \theta \mapsto e^{i\theta}$ le morphisme de groupe de $(\mathbf{R}, +)$ dans $(\mathcal{U}, \times)$. Soit $\mathbf{Z} \subset \mathbf{R}$ qui est un sous-groupe de $\mathbf{R}$. Alors $f(\mathbf{Z})$ est un sous-groupe de $\mathcal{U}$.

De même si $H' = \{1, j, j^2\}$ qui est un sous-groupe de $\mathcal{U}$, alors $f^{-1}(H') = \frac{2\pi}{3}\mathbf{Z}$ qui est un sous-groupe de $(\mathbf{R}, +)$.

2. Soit $f : n \mapsto a.n$ de $\mathbf{Z}$ dans $\mathbf{Z}$. Alors $f(\mathbf{Z})$ est $a\mathbf{Z}$.

Définition 19.11

Soit $(G, \top)$ et $(G', *)$ deux groupes et f un morphisme de G dans G' .

1. On appelle noyau de f et on note $\text{Ker}(f)$ le sous-groupe $f^{-1}(\{e'\})$.
2. On appelle image de f et on note $\text{Im}(f)$ le sous-groupe $f(G)$.

Exemples :

1. Soit E un espace euclidien. On a vu que $O(E)$ était un groupe pour la multiplication. Maintenant $SO(E)$ est le noyau du déterminant de $O(E)$ dans $(\mathbf{R}^*, \times)$. On en déduit que $SO(E)$ est un sous-groupe de $O(E)$.
2. Le noyau de la signature $\varepsilon : \mathfrak{S}_n \rightarrow \{\pm 1\}$ est un sous-groupe de $\mathfrak{S}_n$. Il s'appelle le groupe alterné $\mathfrak{A}_n$.

Pour $n = 3$, il y a six permutations :

- L'identité e
- Les transpositions $\tau_1 = (1, 2)$, $\tau_2 = (1, 3)$ et $\tau_3 = (2, 3)$
- Les 3-cycles $\gamma_1 = (1, 2, 3)$, $\gamma_2 = (1, 3, 2)$

On a alors $\mathfrak{A}_3 = \{e, \gamma_1, \gamma_2\}$.

3. Soit $(G, *)$ un groupe et $g \in G$. On peut considérer

$$\begin{aligned} \varphi_g : \mathbf{Z} &\rightarrow G \\ n &\mapsto g^n \end{aligned}$$

C'est un morphisme de groupe et

$$\text{Im}(\varphi_g) = \langle g \rangle.$$

Proposition 19.12

Avec les notations précédentes,

1. L'application f est injective si et seulement si $\text{Ker}(f) = \{e\}$.
2. L'application f est surjective si et seulement si $\text{Im}(f) = G'$.

Exemple : L'application $f : \theta \mapsto e^{i\theta}$ est surjective mais pas injective. Son noyau est $2\pi\mathbf{Z}$.

Proposition 19.13

Soit $(G, \top)$ et $(G', *)$ deux groupes et f un isomorphisme de G dans G' . L'application réciproque $f^{-1} : G' \rightarrow G$ est un morphisme de groupe.

Démonstration : Soit (g'_1, g'_2) deux éléments de G' , on veut montrer que $f^{-1}(g'_1 * g'_2) = f^{-1}(g'_1) \top f^{-1}(g'_2)$. Comme f est injective (car bijective), il suffit de montrer que

$$f(f^{-1}(g'_1 * g'_2)) = f(f^{-1}(g'_1) \top f^{-1}(g'_2)).$$

Or

$$f(f^{-1}(g'_1 * g'_2)) = g'_1 * g'_2$$

et

$$f(f^{-1}(g'_1) \top f^{-1}(g'_2)) = f(f^{-1}(g'_1)) * f(f^{-1}(g'_2)) = g'_1 * g'_2.$$

□

Exercice : 5 ; 8 ; 9 p 59

2 Groupes monogènes

2.1 Groupe $(\mathbb{Z}/n\mathbb{Z}, +)$

Proposition-Définition 19.14

Soit n un entier non nul.

1. La relation $m_1 \mathcal{R}_n m_2 \iff n \mid m_1 - m_2$ de congruence modulo n est une relation d'équivalence.
2. On note $\mathbb{Z}/n\mathbb{Z}$ l'ensemble des classes d'équivalence de la relation $\mathcal{R}_n$

Exemple : Pour $n = 5$, il y a 5 classes d'équivalence $C_0 = \{0, 5, -5, 10, -10, \dots\}$, $C_1 = \{1, 6, -4, 11, -9, \dots\}$, $C_2 = \{2, 7, -3, 12, -8, \dots\}$, C_3 et C_4 .

Notation : On notera par la suite C_m ou $\overline{m}$ la classe de l'entier m dans $\mathbb{Z}/n\mathbb{Z}$.

Proposition 19.15

Pour $n \geq 1$, l'ensemble $\mathbb{Z}/n\mathbb{Z}$ est en bijection avec $[[0; n-1]]$.

Démonstration : Soit $\varphi : [[0; n-1]] \rightarrow \mathbb{Z}/n\mathbb{Z}$ qui associe à un entier m sa classe modulo n .

- L'application φ est injective car pour $(m, m') \in [[0; n-1]]^2$, si on suppose que $\varphi(m) = \varphi(m')$ alors $m \equiv m' [n]$ donc $n \mid m - m'$.

Cependant, comme $0 \leq m \leq n-1$ et $0 \leq m' \leq n-1$ on a donc $-(n-1) \leq m - m' \leq n-1$ et donc $m - m' = 0$.

L'application est bien injective.

- Montrons que φ est surjective. Soit C une classe de $\mathbb{Z}/n\mathbb{Z}$. Soit m un élément de la classe. On fait la division euclidienne de m par n ,

$$m = nq + r \text{ où } r \in [[0; n-1]]$$

En particulier $m \equiv r [n]$ et donc C est la classe de r , c'est-à-dire $C = \varphi(r)$.

□

Remarque : Cela signifie que chaque classe C de $\mathbb{Z}/n\mathbb{Z}$ a un et un seul représentant dans $[[0; n-1]]$.

Notation : On note souvent $\mathbb{Z}/n\mathbb{Z} = \{\bar{0}, \bar{1}, \dots, \overline{n-1}\}$.

Théorème 19.16

Soit $n \in \mathbb{N}^*$. Soit C_1 et C_2 deux classes de $\mathbb{Z}/n\mathbb{Z}$ et $m_1 \in C_1$, $m_2 \in C_2$, la classe à laquelle appartient $m_1 + m_2$ ne dépend pas du choix de m_1 et de m_2 .

Démonstration : Soit m_1 et m'_1 deux représentants de C_1 et m_2 et m'_2 deux représentants de C_2 . On a donc $m_1 \equiv m'_1 [n]$ et $m_2 \equiv m'_2 [n]$. On sait alors que

$$(m_1 + m_2) \equiv (m'_1 + m'_2) [n]$$

Cela montre que

$$\overline{m_1 + m_2} = \overline{m'_1 + m'_2}$$

□

Remarque : Cela signifie que l'on peut « ajouter » les classes.

Définition 19.17

On peut donc définir une loi de composition interne $+$ sur $\mathbb{Z}/n\mathbb{Z}$.

$$\begin{aligned} + : \mathbb{Z}/n\mathbb{Z} \times \mathbb{Z}/n\mathbb{Z} &\rightarrow \mathbb{Z}/n\mathbb{Z} \\ (C_1, C_2) &\mapsto C_1 + C_2 \end{aligned}$$

où $C_1 + C_2$ est la classe de la somme d'un représentant de C_1 et d'un représentant de C_2 .

Exemple : Dans $\mathbb{Z}/5\mathbb{Z}$ on a donc $\bar{2} + \bar{2} = \bar{4}$ et $\bar{4} + \bar{3} = \bar{7} = \bar{2}$.

Proposition 19.18

Avec la loi définie ci-dessus, $(\mathbb{Z}/n\mathbb{Z}, +)$ est un groupe abélien.

Démonstration : Montrons les axiomes des groupes

— La loi est associative. Soit trois classes de $\mathbb{Z}/n\mathbb{Z}$ et m_1, m_2, m_3 des représentants,

$$(\overline{m_1 + m_2}) + \overline{m_3} = \overline{m_1 + m_2 + m_3} = \overline{(m_1 + m_2) + m_3} = \overline{m_1 + (m_2 + m_3)} = \overline{m_1} + (\overline{m_2 + m_3})$$

— La loi est commutative. Soit deux classes de $\mathbb{Z}/n\mathbb{Z}$ et m_1, m_2 , des représentants,

$$\overline{m_1} + \overline{m_2} = \overline{m_1 + m_2} = \overline{m_2 + m_1} = \overline{m_2} + \overline{m_1}.$$

— La classe de 0 est un élément neutre car

$$\forall m \in [[0; n-1]], \bar{0} + \bar{m} = \overline{0 + m} = \bar{m} = \overline{m + 0}$$

— Toute classe a un opposé. En effet

$$\forall m \in \llbracket 0; n-1 \rrbracket, \overline{m} + \overline{-m} = \overline{0}$$

On a donc $-\overline{m} = \overline{-m} = \overline{n-m}$

On a bien montré que $(\mathbb{Z}/n\mathbb{Z}, +)$ était un groupe commutatif. $\square$

Exemple : Voici la table de $(\mathbb{Z}/4\mathbb{Z}, +)$.

+	0	1	2	3
0	0	1	2	3
1	1	2	3	0
2	2	3	0	1
3	3	0	1	2

Note (Groupe quotient - Hors programme)

La construction de $\mathbb{Z}/n\mathbb{Z}$ est un cas particulier d'un groupe quotient.

De manière plus générale, si H est un sous-groupe de G , on peut construire une relation d'équivalence sur G par

$$\forall (g, g') \in G^2, g \mathcal{R}_H g' \iff g(g')^{-1} \in H.$$

Dans le cas de $\mathbb{Z}/n\mathbb{Z}$, $H = n\mathbb{Z}$ est un sous-groupe de $(\mathbb{Z}, +)$. On a bien défini

$$m_1 \mathcal{R}_n m_2 \iff n \mid (m_1 - m_2) \iff m_1 - m_2 \in H$$

On vérifie aisément que $\mathcal{R}_H$ est une relation d'équivalence.

La question est alors de savoir si l'opération de G « passe au quotient » c'est-à-dire peut-on étendre la loi de G aux classes.

Cela revient à dire que si $x_1 \mathcal{R}_H x'_1$ et $x_2 \mathcal{R}_H x'_2$ alors a-t-on

$$x_1 x_2 \mathcal{R}_H x'_1 x'_2 \iff x_1 x_2 (x'_2)^{-1} (x'_1)^{-1} \in H$$

Dans le cas où G est un groupe quelconque cela n'est pas toujours vrai; par contre si G est abélien (ce qui est le cas de $(\mathbb{Z}, +)$) cela fonctionne car

$$x_1 x_2 (x'_2)^{-1} (x'_1)^{-1} = x_1 (x'_1)^{-1} x_2 (x'_2)^{-1} \in H$$

car $x_1 (x'_1)^{-1} \in H$ et $x_2 (x'_2)^{-1} \in H$ par hypothèses.

2.2 Générateurs, groupes monogènes et groupes cycliques

Définition 19.19

Soit $(G, *)$ un groupe et $g \in G$. On dit que g est un générateur de G si $\langle g \rangle = G$.
C'est-à-dire si pour tout $h \in G$ il existe $n \in \mathbb{Z}$ tel que $h = g^n$.

Remarque : Si la loi du groupe est notée additivement, g est un générateur si pour tout h , il existe $n \in \mathbb{Z}$, $h = ng$.

Exemples :

1. Le groupe $(\mathbb{Z}, +)$ est engendré par 1 (ou par -1).
2. Le groupe $(\mathbb{R}, +)$ n'a pas de générateur. En effet pour tout $x_0 \in \mathbb{R}$, $\langle x_0 \rangle = x_0\mathbb{Z} \neq \mathbb{R}$.

Proposition 19.20

Soit $n \geq 1$. Soit $m \in \mathbb{Z}$, $\bar{m}$ est un générateur de $\mathbb{Z}/n\mathbb{Z}$ si et seulement si $n \wedge m = 1$

Exemples :

1. Dans $\mathbb{Z}/6\mathbb{Z}$, on voit que $\bar{5}$ est un générateur. En effet,

r	0	1	2	3	4	5
$r\bar{5}$	$\bar{0}$	$\bar{5}$	$\bar{4}$	$\bar{3}$	$\bar{2}$	$\bar{1}$

En fait, on voit que $\bar{5} = -\bar{1}$.

2. Par contre, $\bar{3}$ n'est pas un générateur de $\mathbb{Z}/6\mathbb{Z}$ car $r\bar{3} = \bar{0}$ ou $\bar{3}$.

Démonstration :

- $\Rightarrow$ Si $\bar{m}$ est un générateur, cela signifie qu'il existe $r \in \mathbb{Z}$ tel que $r\bar{m} = \bar{1} \Rightarrow r\bar{m} = \bar{1}$.

Cela implique que $rm - 1$ est divisible par n . On en déduit qu'il existe un entier r' tel que

$$rm + r'n = 1$$

D'après la relation de Bézout on obtient alors que $m \wedge n = 1$.

- $\Leftarrow$ Il suffit de remonter le calcul précédent pour trouver un entier r tel que $r\bar{m} = \bar{1}$. Maintenant pour tout $k \in \llbracket 0 ; n-1 \rrbracket$,

$$kr\bar{m} = \bar{k}.$$

On obtient bien que $\bar{m}$ engendre $\mathbb{Z}/n\mathbb{Z}$.

□

Définition 19.21

Soit $(G, *)$ un groupe.

1. S'il existe g tel que $\langle g \rangle = G$, on dit que G est monogène.
2. Si de plus G est fini, il est dit cyclique.

Exemples :

1. Le groupe $(\mathbb{Z}, +)$ est monogène

2. Le groupe $(\mathbb{Z}/n\mathbb{Z}, +)$ est cyclique car $\langle 1 \rangle = \mathbb{Z}/n\mathbb{Z}$.

Théorème 19.22

Soit $(G, *)$ un groupe monogène.

1. Si G est infini alors G est isomorphe à $(\mathbb{Z}, +)$
2. Si G est fini alors G est isomorphe à $(\mathbb{Z}/n\mathbb{Z}, +)$ où $n = |G|$.

Démonstration : On suppose que G est monogène. Soit g un générateur de G , on considère le morphisme

$$\begin{aligned} \varphi_g : \mathbb{Z} &\rightarrow G \\ n &\mapsto g^n \end{aligned}$$

Par définition, il est surjectif car G est monogène.

Le noyau de φ_g est donc un sous-groupe de $(\mathbb{Z}, +)$. Il y a donc deux possibilités :

- Si $\text{Ker} \varphi_g = \{0\}$ alors φ_g est injectif. C'est donc un isomorphisme de $\mathbb{Z}$ dans G .
- Si $\text{Ker} \varphi_g = n\mathbb{Z}$. Alors φ_g « se factorise » de la manière suivante

$$\begin{array}{ccc} \mathbb{Z} & \xrightarrow{\varphi_g} & G \\ \downarrow \pi & \searrow f & \uparrow \\ \mathbb{Z}/n\mathbb{Z} & & \end{array}$$

où $\pi : \mathbb{Z} \rightarrow \mathbb{Z}/n\mathbb{Z}$ est le morphisme de groupe qui envoie tout entier m sur $\overline{m}$.

Nous devons montrer que pour toute classe C de $\mathbb{Z}/n\mathbb{Z}$, si on choisit un représentant m de la classe alors $\varphi_g(m)$ ne dépend que de la classe mais pas du choix du représentant.

En effet soit m et m' deux entiers tels que $\overline{m} = \overline{m'}$ alors il existe $q \in \mathbb{Z}$ tel que $m = m' + qn$.

Donc

$$\varphi_g(m) = g^m = g^{m'+qn} = g^{m'} * (g^n)^q = g^{m'} * e = g^{m'} = \varphi_g(m')$$

puisque, comme $n \in \text{Ker}(\varphi_g)$, $g^n = e$.

On peut donc bien construire $f : \mathbb{Z}/n\mathbb{Z} \rightarrow G$ tel que $f \circ \pi = \varphi_g$.

Il est clair que f est aussi surjective car φ_g l'était.

De plus c'est encore un morphisme de groupes car si on se donne deux classes représentées par m et m' ,

$$f(\overline{m} + \overline{m'}) = f(\overline{m + m'}) = f(\pi(m + m')) = \varphi_g(m + m') = \varphi_g(m) * \varphi_g(m') = f(\overline{m}) * f(\overline{m'}).$$

Pour finir, f est cette fois injective. Si on considère une classe représentée par m telle que $f(\overline{m}) = e$. On a donc $\varphi_g(m) = e$ ce qui implique $m \in n\mathbb{Z} = \text{Ker} \varphi_g$ et donc $\overline{m} = \overline{0}$.

En conclusion f est bien un isomorphisme de $\mathbb{Z}/n\mathbb{Z}$ dans G .

□

Exemple : On considère $\mathbb{U}_n = \{z \in \mathbb{C} \mid z^n = 1\}$ l'ensemble des racines n -ièmes de 1 dans $\mathbb{C}$. On a déjà vu que c'était un groupe pour la multiplication. Si on pose $\omega = e^{2i\pi/n}$. On a bien $\mathbb{U}_n = \langle \omega \rangle$.

De fait

$$\begin{aligned}\varphi_\omega : \mathbf{Z} &\rightarrow G \\ n &\mapsto \omega^n\end{aligned}$$

se factorise par

$$\begin{aligned}f : \mathbf{Z}/n\mathbf{Z} &\rightarrow G \\ n &\mapsto \omega^n.\end{aligned}$$

Cela permet de déterminer facilement les générateurs de $\mathbb{U}_n$. Pour tout $\overline{m} \in \mathbf{Z}/n\mathbf{Z}$ on a que $\overline{m}$ engendre $\mathbf{Z}/n\mathbf{Z}$ si et seulement si $f(\overline{m}) = e^{2im\pi/n}$ engendre $\mathbb{U}_n$. De ce fait, les générateurs de $\mathbb{U}_n$ sont les $e^{2ik\pi/n}$ où $k \wedge n = 1$.

Par exemple pour $n = 6$, $\mathbb{U}_6$ est isomorphe à $\mathbf{Z}/6\mathbf{Z}$.

Le groupe $(\mathbf{Z}/6\mathbf{Z}, +)$ a deux générateurs, $\overline{1}$ et $\overline{5} = -\overline{1}$. On en déduit que $\mathbb{U}_6$ a deux générateurs, $\omega = e^{2i\pi/6}$ et $\omega^5 = e^{10i\pi/6} = e^{-2i\pi/6}$.

On peut ici vérifier « à la main » que $\omega^2, \omega^3, \omega^4$ n'engendrent pas $\mathbb{U}_6$.

- Le complexe $\omega^2 = e^{4i\pi/6} = e^{2i\pi/3} = j$ engendre $\{1, j, j^2\}$.
- Le complexe $\omega^3 = e^{6i\pi/6} = e^{i\pi} = -1$ engendre $\{1, -1\}$.
- Le complexe $\omega^4 = e^{8i\pi/6} = e^{4i\pi/3} = j^2$ engendre $\{1, j, j^2\}$ car $(j^2)^2 = j^4 = j$.

3 Ordre d'un élément

3.1 Définition

Définition 19.23

Soit $(G, *)$ un groupe et x un élément de G . On dit que x est d'ordre fini s'il existe un entier $n \geq 1$ tel que $x^n = e$.

Dans ce cas, on appelle ordre de x et on note $\text{ord}(x)$ le plus petit entier strictement positif tel que $x^n = e$.

Remarque : On a donc que si x est d'ordre fini,

$$\forall n \in \mathbf{N}^*, n = \text{ord}(x) \iff x^n = e \text{ et } \forall k \in \llbracket 1; n-1 \rrbracket, x^k \neq e.$$

Exemples :

1. Dans $\mathbf{Z}/12\mathbf{Z}$, l'élément $\overline{3}$ est d'ordre 4 car $2.\overline{3} = \overline{6}$, $3.\overline{3} = \overline{9}$ et $4.\overline{3} = \overline{0}$.
2. Dans $(\mathbf{C}^\times, \times)$, l'élément i est d'ordre 4. Par contre 2 n'est pas d'ordre fini.

Proposition 19.24

Soit $(G, \times)$ un groupe et x un élément d'ordre fini noté d

1. Soit $n \in \mathbf{Z}$, $x^n = e \iff d \mid n$.
2. Le groupe $\langle x \rangle$ est de cardinal d . On a

$$\langle x \rangle = \{e = x^0, x, x^2, \dots, x^{d-1}\}.$$

Démonstration :

1. Le sens $\Leftarrow$ est immédiat. Si $d \mid n$, alors en posant $n = kd$, on a

$$x^n = (x^d)^k = e^k = e$$

À l'inverse montrons par contraposée que si d ne divise pas n alors $x^n \neq e$. En effet, on peut écrire la division euclidienne de n par d ,

$$n = qd + r \text{ où } r \in \llbracket 1; d-1 \rrbracket$$

(on a $r > 0$ car d ne divise pas n). On a alors

$$x^n = x^{qd+r} = (x^d)^q x^r = e x^r = x^r.$$

Maintenant, par définition de l'ordre, $x^r \neq e$.

2. On regarde $\langle x \rangle = \{x^n \mid n \in \mathbb{Z}\}$.

— Pour tout entier n , par division euclidienne

$$n = qd + r \text{ où } r \in \llbracket 0; d-1 \rrbracket$$

on a $x^n = x^r$. On en déduit que $\langle x \rangle = \{x^n \mid n \in \llbracket 0; d-1 \rrbracket\}$.

— De plus pour n_1, n_2 dans $\llbracket 0; d-1 \rrbracket^2$, $x^{n_1} = x^{n_2} \Rightarrow n_1 = n_2$ car

$$x^{n_1} = x^{n_2} \Rightarrow x^{n_1} (x^{n_2})^{-1} = e \Rightarrow x^{n_1 - n_2} = e = x^{n_2 - n_1}$$

Or, par symétrie, on peut supposer $n_1 \geq n_2$ et donc $n_1 - n_2 \in \llbracket 0; d-1 \rrbracket$ ce qui implique que $n_1 - n_2 = 0$ par définition de l'ordre de x .

On a bien $\langle x \rangle = \{e = x^0, x, x^2, \dots, x^{d-1}\}$. De ce fait, $|\langle x \rangle| = d$.

□

Remarque : On peut aussi reprendre l'application

$$\begin{aligned} \varphi_x : \mathbb{Z} &\rightarrow G \\ n &\mapsto x^n \end{aligned}$$

Par définition de l'ordre de x , $\text{Ker } \varphi_x = n\mathbb{Z}$ où $n = \text{ord}(x)$. On en déduit que φ_x induit une application de $\mathbb{Z}/n\mathbb{Z}$ dans $\langle x \rangle$ qui est bijective.

Exercice : Soit G un groupe cyclique de cardinal n , montrer que pour tout diviseur d de n , il existe un sous-groupe de G de cardinal d .

Correction

On suppose que G est cyclique. Notons g un générateur qui est d'ordre n . Soit d un diviseur de n et k tel que $kd = n$. On pose $h = g^k$.

L'élément h est d'ordre d puisque $h^d = (g^k)^d = g^{kd} = g^n = e$ et que si $i < d$, $ki < kd = n$ d'où $h^i = g^{ki} \neq e$.

On considère $H = \langle h \rangle$, on a

$$H = \{e = h^0, h, h^2, \dots, h^{d-1}\}$$

On en déduit que H est un sous-groupe de cardinal d de G .

3.2 Théorème de Lagrange

Définition 19.25

Soit G un groupe fini, on appelle ordre de G son cardinal.

Remarque : En particulier, d'après ce qui précède on a que l'ordre de x est l'ordre de $\langle x \rangle$.

Théorème 19.26

Soit $(G, *)$ un groupe fini. Soit $x \in G$, l'ordre de x divise l'ordre de G

Démonstration : Faisons la démonstration avec une « astuce » dans le cas abélien.

On pose $a = \prod_{g \in G} g$.

On considère l'application Φ de G dans G définie par $\Phi : g \mapsto xg$. C'est une bijection car sa réciproque est $g \mapsto x^{-1}g$.

On en déduit $G = \{g, g \in G\} = \{\Phi(h), h \in G\}$

Cela nous donne

$$a = \prod_{g \in G} g = \prod_{h \in G} \Phi(h) = \prod_{h \in G} (xh) = x^{|G|} \prod_{h \in G} h = x^{|G|} a$$

On en déduit que $x^{|G|} = e$. Cela signifie d'après la proposition précédente que $\text{ord}(x) \mid |G|$. $\square$

Corollaire 19.27

Soit $(G, *)$ un groupe fini de cardinal n et $g \in G$, $g^n = e$

Exemple : Ce théorème sert à déterminer tous les sous-groupes d'un groupes. Par exemple dans $\mathfrak{S}_3$ qui est de cardinal 6, tous les éléments sont d'ordre 1, 2, 3 ou 6.

Exercice : Soit $n \geq 1$. Soit G un sous-groupe fini de $\text{GL}_n(\mathbb{C})$. Montrer que toute matrice $A \in G$ est diagonalisable.

Note (Théorème de Lagrange - Hors programme)

Il existe une version plus générale du résultat précédent

Théorème 19.28

Soit $(G, *)$ un groupe fini. Soit H un sous-groupe de G . L'ordre de H divise l'ordre de G .

Remarque : Le résultat précédent s'en déduit en considérant le sous-groupe $\langle x \rangle$.

Démonstration : On considère la relation sur G ,

$$\forall (x, y) \in G^2, x \mathcal{R}_H y \iff x^{-1}y \in H.$$

On montre aisément que c'est une relation d'équivalence.

On sait alors que les classes d'équivalences forment une partition de G . En particulier, il n'y a qu'un nombre fini de classes d'équivalences. Notons $C_1, \dots, C_p$ ces classes.

Soit $i \in \llbracket 1 ; p \rrbracket$ et x_i un élément de C_i , on a pour $y \in G$,

$$\begin{aligned} y \in C_i &\iff x_i^{-1}y \in H \\ &\iff \exists h \in H, x_i^{-1}y = h \\ &\iff \exists h \in H, y = x_i h \end{aligned}$$

On peut alors noter $C_i = x_i H = \{x_i h \mid h \in H\}$. Ce sont les « translations » de H .

Notons maintenant que $|x_i H| = |H|$. En effet l'application

$$\begin{aligned} \psi : H &\rightarrow x_i H \\ h &\mapsto x_i h \end{aligned}$$

est bijective. Elle est surjective par définition et de plus elle est injective car pour h_1 et h_2 on a

$$x_i h_1 = x_i h_2 \implies h_1 = h_2$$

Pour conclure, comme les $x_1 H, \dots, x_p H$ forment une partition de G ,

$$|G| = \sum_{i=1}^p |x_i H| = p|H|$$

On a bien

$$|H| \mid |G|.$$

□

Exercice : Soit G de cardinal n et d un entier divisant n .

On a vu que si G était cyclique alors il existait un sous-groupe H d'ordre d .

Le but de l'exercice est de montrer que cela n'est plus vrai si G n'est plus cyclique.

On considère $G = \mathfrak{A}_4$ le groupe alterné qui est le sous-groupe de $\mathfrak{S}_4$ des permutations paires.

$$\mathfrak{A}_4 = \{\sigma \in \mathfrak{S}_4, \varepsilon(\sigma) = 1\}$$

1. Quel est le cardinal de $\mathfrak{A}_4$? Le montrer rigoureusement.
2. Faire la liste des éléments de $\mathfrak{A}_4$ en les classant selon la forme de leur décomposition en cycles à support disjoint.
3. Supposons par l'absurde qu'il existe un sous-groupe H d'ordre 6 dans $\mathfrak{A}_4$.
 - (a) Montrer que pour tout $g \in \mathfrak{A}_4$, $g^2 \in H$.

On pourra séparer selon que g appartienne à H ou pas.

Dans le cas où $g \notin H$, on pourra considérer $\psi_g : x \mapsto gx$ de $\mathfrak{A}_4$ dans $\mathfrak{A}_4$.

- (b) En déduire que tous les 3-cycles appartiennent à H .
- (c) Conclure

Correction

1. On sait que $|\mathfrak{S}_4| = 24$. Vérifions que $|\mathfrak{A}_4| = 12$. Soit $\tau = (1, 2)$ une transposition, l'application

$$\begin{aligned} \Phi : \mathfrak{S}_4 &\rightarrow \mathfrak{S}_4 \\ \sigma &\mapsto \tau \circ \sigma \end{aligned}$$

est une bijection car sa réciproque est $\sigma \mapsto \tau^{-1} \circ \sigma$ (qui est aussi Φ car $\sigma^{-1} = \sigma$).

De plus si on note $K = \mathfrak{S}_4 \setminus \mathfrak{A}_4 = \{\sigma \in \mathfrak{S}_4, \varepsilon(\sigma) = -1\}$ (qui n'est pas un sous-groupe), on voit que Φ envoie les éléments de $\mathfrak{A}_4$ dans K , cela implique que $\Phi(\mathfrak{A}_4) \subset K$ et donc

$$|\mathfrak{A}_4| = |\Phi(\mathfrak{A}_4)| \leq |K|$$

Maintenant, Φ envoie aussi K dans $\mathfrak{A}_4$ donc, par le même argument, $|K| \leq |\mathfrak{A}_4|$.

Finalement $|K| = |\mathfrak{A}_4|$, en utilisant finalement $K \sqcup \mathfrak{A}_4 = \mathfrak{S}_4$ on obtient

$$|\mathfrak{A}_4| = \frac{1}{2}|\mathfrak{S}_4| = 12$$

2. Les éléments de $\mathfrak{A}_4$ sont les suivants :
 - L'identité
 - Il n'y a pas de transpositions car ce sont des permutations impaires.
 - Les 3-cycles ; il y en a 8 puisque pour tout élément i de $[[1; 4]]$ il y a deux trois cycles qui n'ont pas i dans son support. Par exemple pour $i = 1$ il y a les trois cycles $(2, 3, 4)$ et $(2, 4, 3)$.
 - Les compositions de deux transpositions à support disjoint : $(1, 2)(3, 4); (1, 3)(2, 4); (1, 4)(2, 3)$

Correction

3. (a) Soit $g \in H$, il est évidemment que $g^2 \in H$.

Soit $g \in \mathfrak{A}_4$, on suppose que $g \notin H$. On considère $\psi_g : x \mapsto gx$ de $\mathfrak{A}_4$ dans $\mathfrak{A}_4$. Comme $g \notin H$, $\psi_g(H) \subset \mathfrak{A}_4 \setminus H$. Or $|\psi_g(H)| = |H| = |\mathfrak{A}_4 \setminus H|$ donc $\psi_g(H) = \mathfrak{A}_4 \setminus H$.

On en déduit que les éléments de H sont envoyés sur ceux de $\mathfrak{A}_4 \setminus H$ et que réciproquement ceux de $\mathfrak{A}_4 \setminus H$ sont envoyés sur des éléments de H . En particulier $g^2 = \psi_g(\psi_g(e)) \in H$.

- (b) Soit σ un 3-cycle, il est d'ordre 3 donc

$$\sigma = \sigma^4 = (\sigma^2)^2$$

Cela montre que tous les 3-cycles appartiennent à H .

- (c) Il y a 8 permutations qui sont des 3-cycles et $|H| = 6$. C'est absurde.

4 Anneaux

4.1 Définitions

Définition 19.29

Soit A un ensemble muni de deux lois de compositions internes notée $+$ et $\times$. On dit que $(A, +, \times)$ est un anneau si

1. $(A, +)$ est un groupe commutatif (d'où la notation $+$)
2. la loi $\times$ est associative
3. il existe un élément neutre pour $\times$
4. la loi $\times$ est distributive par rapport à $+$

Si la loi $\times$ est aussi commutative, on dit que A est un anneau commutatif.

Remarques :

1. Cette définition correspondent aux règles de calculs classiques dans $\mathbb{Z}, \mathbb{Q}, \mathbb{R}$.
2. L'ensemble $(A, \times)$ n'est pas forcément un groupe car des éléments peuvent pas avoir de symétrique.
3. On peut voir que $A = \{0\}$ est un anneau que l'on appelle l'anneau nul. Ce n'est pas le plus intéressant et on essaye souvent d'exclure ce cas.

Notations :

1. On note 0 ou 0_A l'élément neutre pour $+$.
2. On note 1 ou 1_A l'élément neutre pour $\times$.
3. On abrège $x \times y$ en xy .

Définition 19.30

Soit $(A, +, \times)$ un anneau. Un élément x de A est dit inversible s'il est inversible pour la multiplication. C'est-à-dire s'il existe y dans A tel que $xy = yx = 1$.

Dans ce cas y est unique et on le note y^{-1} . On dit aussi que x est une unité

Remarques :

1. La notion d'inversible réfère à la loi $\times$ et pas la loi $+$. Tous les éléments sont « inversibles » par $+$.
2. A priori, il y a des éléments qui ne sont pas inversibles dans un anneau.
3. On ne note pas $\frac{1}{x}$ l'inverse pour éviter d'écrire par la suite $\frac{y'}{y}$ qui pourrait être $y' \times \frac{1}{y}$ ou $\frac{1}{y} \times y'$ qui sont à priori différents.
4. On note la plupart du temps $A^\times$ l'ensemble des éléments inversibles d'un anneau.
5. Quand un élément x est inversible, on étend les notations x^n au cas où n est négatif.

Proposition 19.31

Soit $(A, +, \times)$ un anneau. L'ensemble des éléments inversibles forme un groupe multiplicatif noté $(A^\times, \times)$.

Démonstration : On ne peut pas nécessairement voir $A^\times$ comme un sous-groupe d'un groupe particulier. Il faut revenir à la définition.

- La loi $\times$ est une loi interne à $A^\times$ car si x, y appartiennent à $A^\times$ alors

$$xyy^{-1}x^{-1} = xx^{-1} = 1_A$$

Cela montre que $xy \in A^\times$ (et que $(xy)^{-1} = y^{-1}x^{-1}$).

- La loi $\times$ est associative (par définition d'un anneau)
- Il y a un élément neutre 1_A pour $\times$ dans A donc à fortiori dans $A^\times$.
- Pour tout élément $x \in A^\times$, il existe, par définition, un élément $y = x^{-1} \in A^\times$ tel que $xy = yx = 1_A$.

□

Exemples :

1. Dans $\mathbb{Z}$, 1 et -1 sont les seuls inversibles. On a donc $\mathbb{Z}^\times = \{\pm 1\}$.
2. Les éléments inversibles de $\mathbb{R}$ sont tous les réels non nuls.
3. Les éléments inversibles de $\mathbb{R}^{\mathbb{R}}$ sont les fonctions qui ne s'annulent pas.

ATTENTION

Dans les anneaux $\mathbb{R}$ et $\mathbb{C}$ on a l'habitude de noter $\mathbb{R}^* = \mathbb{R} \setminus \{0\}$ et $\mathbb{C}^* = \mathbb{C} \setminus \{0\}$. Dans ces anneaux (qui sont des corps - voir plus loin), on a $\mathbb{R}^\times = \mathbb{R}^*$ et $\mathbb{C}^\times = \mathbb{C}^*$. Ce n'est pas le cas en général. Par exemple $\mathbb{Z}^\times = \{\pm 1\} \neq \mathbb{Z} \setminus \{0\}$.

Définition 19.32 (Sous-anneau)

Soit $(A, +, \times)$ un anneau. On appelle, sous-anneau de A tout sous groupe de $(A, +)$ contenant 1 et qui est stable par multiplication.

Exemples :

1. On voit $\mathbb{Z}$ est un sous-anneau de $\mathbb{R}$.
2. $\{0\}$ n'est pas un sous-anneau de $\mathbb{R}$ car il ne contient pas 1.

Exercice : Montrer que $\mathbb{Z}[\sqrt{5}] = \{a + b\sqrt{5} \mid (a, b) \in \mathbb{Z}^2\}$ est un sous-anneau de $\mathbb{R}$.

L'élément $1 + \sqrt{5}$ est-il inversible ? Et $2 + \sqrt{5}$?

Correction

- On vérifie aisément que $\mathbb{Z}[\sqrt{5}]$ est un sous-groupe de $(\mathbb{R}, +)$ car il est non vide et que si $u = a + b\sqrt{5}$ et $v = a' + b'\sqrt{5}$ appartiennent à $\mathbb{Z}[\sqrt{5}]$ alors

$$u - v = (a + b\sqrt{5}) - (a' + b'\sqrt{5}) = (a - a') + (b - b')\sqrt{5} \in \mathbb{Z}[\sqrt{5}]$$

- On vérifie aussi que $1 = 1 + 0\sqrt{5} \in \mathbb{Z}[\sqrt{5}]$
- Pour finir, si $u = a + b\sqrt{5}$ et $v = a' + b'\sqrt{5}$ appartiennent à $\mathbb{Z}[\sqrt{5}]$ alors

$$u \times v = (a + b\sqrt{5}) \times (a' + b'\sqrt{5}) = (aa' + 5bb') + (ab' + a'b)\sqrt{5} \in \mathbb{Z}[\sqrt{5}]$$

Cherchons un inverse $u = a + b\sqrt{5}$ à $1 + \sqrt{5}$. On doit donc avoir, en utilisant ce qui précède

$$\begin{cases} a + 5b = 1 \\ b + a = 0 \end{cases} \iff \begin{cases} 4b = 1 \\ a = -b \end{cases}$$

Ce système n'a pas de solutions (dans $\mathbb{Z}$) donc $1 + \sqrt{5}$ n'est pas inversible dans $\mathbb{Z}[\sqrt{5}]$.

Cherchons un inverse $u = a + b\sqrt{5}$ à $2 + \sqrt{5}$. On doit donc avoir, en utilisant ce qui précède

$$\begin{cases} 2a + 5b = 1 \\ a + 2b = 0 \end{cases} \iff \begin{cases} b = 1 \\ a = -2b \end{cases}$$

On en déduit que $-2 + \sqrt{5}$ est l'inverse de $2 + \sqrt{5}$.

Proposition 19.33

Soit $A_1, \dots, A_p$ des anneaux. Le produit $A = \prod_{i=1}^p A_i$ est muni d'une structure d'anneau par les opérations

$$+ : \begin{array}{ccc} A \times A & \rightarrow & A \\ ((x_1, \dots, x_p), (y_1, \dots, y_p)) & \mapsto & (x_1 + y_1, \dots, x_p + y_p) \end{array}$$

et

$$\times : \begin{array}{ccc} A \times A & \rightarrow & A \\ ((x_1, \dots, x_p), (y_1, \dots, y_p)) & \mapsto & (x_1 y_1, \dots, x_p y_p) \end{array} .$$

On a alors $0_A = (0, \dots, 0)$ et $1_A = (1, \dots, 1)$.

Exercice : Montrer que $A^\times = \prod_{i=1}^p A_i^\times$.

Correction

— Soit $u = (x_1, \dots, x_p) \in A^\times$ et $v = (y_1, \dots, y_p)$ son inverse. Par définition,

$$u \times v = (x_1 y_1, \dots, x_p y_p) = (1, \dots, 1) = 1_A$$

Cela implique que pour tout $i \in \llbracket 1; p \rrbracket$, $x_i y_i = 1$ et donc que $x_i \in A_i^\times$. Donc $u \in \prod_{i=1}^p A_i^\times$.

— Réciproquement, soit $u = (x_1, \dots, x_p) \in \prod_{i=1}^p A_i^\times$. Pour tout $i \in \llbracket 1; p \rrbracket$, $x_i \in A_i^\times$ et donc il existe $y_i \in A_i$ tel que $x_i y_i = 1$ (dans A_i).

On pose $v = (y_1, \dots, y_p)$ et

$$uv = (x_1 y_1, \dots, x_p y_p) = (1, \dots, 1) = 1_A$$

Donc u est inversible.

ATTENTION

Ne pas confondre la notion d'idéal et la notion de sous-anneaux.

4.2 Morphisme d'anneaux**Définition 19.34**

Soit $(A, +, \times)$ et $(B, \oplus, \otimes)$ deux anneaux. Une application f de A dans B est un morphisme d'anneau si

- i) $\forall (x, x') \in A^2, f(a + a') = f(a) \oplus f(a')$
- ii) $\forall (x, x') \in A^2, f(a \times a') = f(a) \otimes f(a')$
- iii) $f(1_A) = 1_B$.

Remarque : D'après l'hypothèse $i)$ on a

$$f(0_A) = f(0_A + 0_A) = f(0_A) \oplus f(0_A)$$

ce qui implique que $f(0_A) = 0_A$.

Par contre,

$$f(1_A) = f(1_A \times 1_A) = f(1_A) \otimes f(1_A)$$

mais cela n'implique pas que $f(1_A) = 1_B$ sauf si on sait que $f(1_A)$ est inversible. On doit en effet ajouter la condition $iii)$

Exemple : Pour $x_0 \in \mathbf{R}$, l'application

$$\begin{aligned} \delta_{x_0} : \mathbf{R}[X] &\rightarrow \mathbf{R} \\ P &\mapsto P(x_0) \end{aligned}$$

est un morphisme d'anneau.

Proposition-Définition 19.35

Soit f un morphisme d'anneau de $(A, +, \times)$ dans $(B, \oplus, \otimes)$

1. On appelle image de f et on note $\text{Im} f$ l'ensemble $f(A)$. C'est un sous-anneau.
2. On appelle noyau de f et on note $\text{Ker} f$ l'ensemble $f^{-1}(\{0_B\}) = \{x \in A \mid f(x) = 0_B\}$. C'est un idéal (bilatère).

Démonstration :

1. On veut montrer que $f(A)$ est un sous-anneau.
 - On sait que $f(A)$ est un sous-groupe car un morphisme d'anneau induit un morphisme de groupe si on « oublie » la multiplication.
 - Montrons que $f(A)$ est stable par $\otimes$. Soit $(y, y') \in f(A)^2$, il existe $(x, x') \in A^2$ tels que $f(x) = y$ et $f(x') = y'$. On a alors

$$f(x \times x') = f(x) \otimes f(x') = y \otimes y'.$$

Donc $y \otimes y' \in f(A)$.

- Montrons que $1_B \in f(A)$. On a $1_B = f(1_A) \in f(A)$.
2. Comme pour $\text{Im} f$, $\text{Ker} f$ est un sous-groupe car c'est le noyau du morphisme de groupe $f : (A, +) \rightarrow (B, \oplus)$. De plus pour $x \in A$ et $y \in \text{Ker} f$,

$$f(x \times y) = f(x) \otimes f(y) = f(x) \otimes 0_B = 0_B$$

et

$$f(y \times x) = f(y) \otimes f(x) = 0_B \otimes f(x) = 0_B.$$

On a bien $x \times y \in \text{Ker} f$ et $y \times x \in \text{Ker} f$.

□

Remarque : Le noyau n'est pas un sous-anneau car $1_A \notin \text{Ker} f$.

Exemple : Si on reprend l'exemple précédent, on a $\text{Ker} \delta_{x_0} = \{P \in \mathbf{R}[X] \mid P(x_0) = 0\}$ qui est un idéal de $\mathbf{R}[X]$. C'est $(X - x_0)\mathbf{R}[X]$. On a $\text{Im} \delta_{x_0} = \mathbf{R}$ qui est un sous-anneau de $\mathbf{R}$.

Définition 19.36

Soit f un morphisme d'anneau de $(A, +, \times)$ dans $(B, \oplus, \otimes)$. Si f est bijective, on dit que f est un isomorphisme d'anneaux.

Proposition 19.37

La réciproque d'un isomorphisme d'anneaux est encore un isomorphisme d'anneaux.

Démonstration : Exercice.

4.3 Anneaux intègres et corps

Définition 19.38

Un anneau (non nul) A est dit intègre si

- i) Il est commutatif
- ii) Pour tout x et y dans A , $x \times y = 0 \Rightarrow x = 0$ ou $y = 0$.

Exemple : L'anneau $K[X]$ est intègre.

Définition 19.39

Un ensemble $(K, +, \times)$ est un corps si c'est un anneau commutatif dont tous les éléments non nuls sont inversibles.

Exemples :

1. Les ensembles $\mathbb{Q}, \mathbb{R}, \mathbb{C}$ sont des corps.
2. L'ensemble $\mathbb{Z}$ n'est pas un corps.

Remarque : Un corps est intègre.

Définition 19.40

Soit K un corps et L un sous-anneau de K . On dit que L est un sous-corps de K ou que K est un sur-corps de L .

Exemple : $\mathbb{Q}, \mathbb{R}$ et $\mathbb{C}$.

5 L'anneau $(\mathbb{Z}/n\mathbb{Z}, +, \times)$

5.1 Définition

On va montrer que la multiplication de $\mathbb{Z}$ « passe aussi au quotient » c'est-à-dire qu'elle engendre une multiplication sur $\mathbb{Z}/n\mathbb{Z}$ et que cela en fait un anneau.

Théorème 19.41

Soit $n \in \mathbb{N}^*$. Soit C_1 et C_2 deux classes de $\mathbb{Z}/n\mathbb{Z}$ et $m_1 \in C_1, m_2 \in C_2$.
La classe à laquelle appartient $m_1 \times m_2$ ne dépend pas du choix de m_1 et de m_2 .

Démonstration : Soit m_1 et m'_1 deux représentants de C_1 et m_2 et m'_2 deux représentants de C_2 .

On a donc $m_1 \equiv m'_1 [n]$ et $m_2 \equiv m'_2 [n]$. On en déduit que

$$(m_1 \times m_2) \equiv (m'_1 \times m'_2) [n].$$

En effet, si on pose $m'_1 = m_1 + k_1 n$ et $m'_2 = m_2 + k_2 n$ on a

$$m_1 \times m'_2 = (m_1 + k_1 n) \times (m_2 + k_2 n) = m_1 \times m_2 + n(k_1 m_2 + k_2 m_1 + n k_1 k_2).$$

□

Remarque : Cela signifie que l'on peut « multiplier » les classes.

Définition 19.42

On peut donc définir une loi de composition interne $\times$ sur $\mathbb{Z}/n\mathbb{Z}$.

$$\begin{aligned} \times : \mathbb{Z}/n\mathbb{Z} \times \mathbb{Z}/n\mathbb{Z} &\rightarrow \mathbb{Z}/n\mathbb{Z} \\ (C_1, C_2) &\mapsto C_1 \times C_2 \end{aligned}$$

où $C_1 \times C_2$ est la classe du produit d'un représentant de C_1 et d'un représentant de C_2 .

Exemple : Voici les tables de multiplication de $\mathbb{Z}/4\mathbb{Z}$ et $\mathbb{Z}/5\mathbb{Z}$

$\mathbb{Z}/4\mathbb{Z} :$

$\times$	0	1	2	3
0	0	0	0	0
1	0	1	2	3
2	0	2	0	2
3	0	3	2	1

$\mathbb{Z}/5\mathbb{Z} :$

$\times$	0	1	2	3	4
0	0	0	0	0	0
1	0	1	2	3	4
2	0	2	4	1	3
3	0	3	1	4	2
4	0	4	3	2	1

Théorème 19.43

Avec les opérations ci-dessus $(\mathbb{Z}/n\mathbb{Z}, +, \times)$ est un anneau commutatif.

Démonstration : Il faut vérifier tous les axiomes des anneaux.

1. $(\mathbb{Z}/n\mathbb{Z}, +)$ est un groupe commutatif : déjà fait
2. La loi $\times$ est associative : cela se montre comme pour l'associativité de $+$ en remontant au cas de $\times$ dans $\mathbb{Z}$. Précisément, soit C_1, C_2 et C_3 trois classes de $\mathbb{Z}/n\mathbb{Z}$. On considère m_1, m_2 et m_3 dans $\mathbb{Z}$ tel que $C_i = \overline{m_i}$. On a alors

$$\begin{aligned} (\overline{m_1} \times \overline{m_2}) \times \overline{m_3} &= \overline{m_1 \times m_2 \times m_3} \\ &= \overline{(m_1 \times m_2) \times m_3} \\ &= \overline{m_1 \times (m_2 \times m_3)} \\ &= \overline{m_1} \times \overline{m_2 \times m_3} \\ &= \overline{m_1} \times (\overline{m_2} \times \overline{m_3}) \end{aligned}$$

3. La loi $\times$ est commutative : cela se montre comme pour la commutativité de $+$ en remontant au cas de $\times$ dans $\mathbb{Z}$.
4. La loi $\times$ est distributive par rapport à $+$: cela se montre encore en remontant au cas de $\mathbb{Z}$.
5. Il existe un élément neutre pour $\times$: c'est $\bar{1}$ la classe de 1.

En effet pour tout $m \in \mathbb{Z}$,

$$\bar{1} \times \bar{m} = \overline{1 \times m} = \bar{m}$$

□

Proposition 19.44

Soit $n \geq 1$. Soit $m \in \mathbb{Z}$, la classe de m est inversible dans $\mathbb{Z}/n\mathbb{Z}$ si et seulement si n et m sont premiers entre eux.

Remarque : On dit alors que m est inversible modulo n .

Démonstration : On procède par équivalence

$$\begin{aligned}
 \overline{m} \text{ inversible dans } \mathbb{Z}/n\mathbb{Z} &\iff \exists k \in \mathbb{Z}, \overline{k} \times \overline{m} = 1 \\
 &\iff \exists k \in \mathbb{Z}, \overline{k \times m} = 1 \\
 &\iff \exists (k, d) \in \mathbb{Z}^2, km = 1 + dn \\
 &\iff m \wedge n = 1 \text{ d'après la formule de Bézout.}
 \end{aligned}$$

□

Exemple : En reprenant les tables ci-dessus, on voit que 3 est inversible modulo 5 car $3 \times 2 \equiv 1[5]$ mais 2 n'est pas inversible modulo 4.

Corollaire 19.45

Soit $n \geq 1$. L'anneau $\mathbb{Z}/n\mathbb{Z}$ est un corps si et seulement si n est un nombre premier.

Remarque : Pour $n = 1$, on a $\mathbb{Z}/1\mathbb{Z}$ qui est l'anneau nul. Ce n'est pas un corps par définition.

Démonstration :

- $\Rightarrow$: Par contraposée, on montre que si n n'est pas premier alors $\mathbb{Z}/n\mathbb{Z}$ n'est pas un corps. Pour cela on montre que $\mathbb{Z}/n\mathbb{Z}$ n'est pas un anneau intègre. En effet, si n n'est pas premier, on considère d un diviseur de n différent de 1 et de n . On pose $d' = \frac{n}{d}$. Alors

$$\overline{d'}\overline{d} = \overline{n} = 0$$

mais $\overline{d'} \neq 0$ et $\overline{d} \neq 0$.

Cela montre aussi que $\overline{d}$ n'est pas un élément inversible dans $\mathbb{Z}/n\mathbb{Z}$. En effet, s'il l'était et si $\overline{k}$ était l'inverse de $\overline{d}$ on aurait

$$\overline{d'} = \overline{d'} \times 1 = \overline{d'} \times \overline{d} \times \overline{k} = 0$$

ce qui est absurde.

- $\Leftarrow$: On suppose que n est un nombre premier. Soit $x \in \mathbb{Z}/n\mathbb{Z} \setminus \{0\}$. Il a un représentant $m \in \llbracket 1; n-1 \rrbracket$. Maintenant, comme n est premier, $m \wedge n = 1$ et donc x est inversible.

□

ATTENTION

Ce résultat est important. Les corps $F_p = \mathbb{Z}/p\mathbb{Z}$ où p est premier sont des objets importants en mathématiques. Ce sont des corps, on peut donc travailler avec comme dans $\mathbb{Q}, \mathbb{R}$ ou $\mathbb{C}$. En particulier F_2 est le corps naturel pour le calcul en binaire.

La plupart des résultats vus cette année sur les corps se généralisent à F_p . Cependant, ils restent un peu spéciaux car on a $\underbrace{1 + 1 + \dots + 1}_{p \text{ fois}} = 0$.

p fois

Exercice : Résoudre $x^2 - 3x + 4 = 0$ dans $\mathbb{Z}/11\mathbb{Z}$. Résoudre $x^2 + 5x + 11 = 0$ dans $\mathbb{Z}/13\mathbb{Z}$. On utilisera une mise sous-forme canonique.

Correction

Pour la première on réalise une mise sous forme canonique.

Dans $\mathbb{R}$ (ou $\mathbb{C}$), on sait que

$$x^2 + ax + b = \left(x + \frac{a}{2}\right)^2 + b - \frac{a^2}{4}$$

On veut réaliser la même chose dans $\mathbb{Z}/11\mathbb{Z}$. Il faut faire attention à qui est alors $\frac{1}{2}$. C'est l'inverse de 2 c'est-à-dire 6 car $6 \times 2 = 1$ dans $\mathbb{Z}/11\mathbb{Z}$.

On peut donc écrire, dans $\mathbb{Z}/11\mathbb{Z}$,

$$(x - 6 \times 3)^2 = x^2 - 2 \times 6 \times 3 \times x + (6 \times 3)^2$$

C'est-à-dire

$$(x - 7)^2 = x^2 - 3x + (-4)^2 = x^2 - 3x + 5$$

On en déduit

$$x^2 - 3x + 6 = 0 \iff (x - 7)^2 - 1 = 0 \iff (x - 7)^2 = 1 \iff (x - 7) = \pm 1$$

Finalement $S = \{6, 8\}$.

Pour le deuxième, dans $\mathbb{Z}/13\mathbb{Z}$, l'inverse de 2 est 7 car $7 \times 2 = 1$. On en déduit que

$$(x + 7 \times 5)^2 = x^2 + 5x + (7 \times 5)^2$$

C'est-à-dire

$$(x + 9)^2 = x^2 + 5x + (-4)^2 = x^2 + 5x + 3$$

Donc

$$x^2 + 5x + 11 = 0 \iff (x + 9)^2 + 8 = 0 \iff (x + 9)^2 = 5$$

Il reste à savoir s'il existe $y \in \mathbb{Z}/13\mathbb{Z}$ tel que $y^2 = 5$. Pour cela, la plus simple (pour nous) est de calculer les carrés :

y	0	± 1	± 2	± 3	± 4	± 5	± 6
y^2	0	1	4	9	3	12	10

Cela montre que 5 n'est pas un carré modulo 13. L'équation n'a pas de solutions.

Exercice : Les groupes $(\mathbb{Z}/6\mathbb{Z})^\times$, $(\mathbb{Z}/8\mathbb{Z})^\times$ et $(\mathbb{Z}/11\mathbb{Z})^\times$ sont-ils cycliques ?

Correction

- Les éléments inversibles de $\mathbb{Z}/6\mathbb{Z}$ sont $\{1, 5\} = \{1, -1\}$. Le groupe $(\mathbb{Z}/6\mathbb{Z})^\times$ est donc le groupe $(\{\pm 1\}, \times)$ qui est cyclique (il est engendré par -1).
- Les éléments inversibles de $\mathbb{Z}/8\mathbb{Z}$ sont $\{1, 3, 5, 7\} = \{1, 3, -3, -1\}$. Si on regarde alors

$$\langle 1 \rangle = \{1\}; \langle -1 \rangle = \{\pm 1\}; \langle 3 \rangle = \{1, 3\}; \langle -3 \rangle = \{1, -3\}$$

Cela montre que $(\mathbb{Z}/8\mathbb{Z})^\times$ n'est pas cyclique.

- Les éléments inversibles de $\mathbb{Z}/11\mathbb{Z}$ sont $\{1, 2, 3, 4, 5, 6, 7, 8, 9, 10\}$. Si on regarde alors

$$\langle 1 \rangle = \{1\}; \langle -1 \rangle = \{\pm 1\}; \langle 2 \rangle = \{1, 2, 4, 8, 5, 10, 9, 7, 3, 6\} = (\mathbb{Z}/11\mathbb{Z})^\times$$

Cela montre que $(\mathbb{Z}/11\mathbb{Z})^\times$ est cyclique.

5.2 Théorème Chinois

On veut résoudre ce problème :

Un capitaine a une troupe de soldats. Il les regroupe un matin et les compte rapidement. Il décide des ranger en 3 colonnes ayant chacune le même nombre de soldats. A ce moment arrive deux soldats retardataires. Le capitaine est de bonne humeur, il leur reproche gentiment leur retard et ordonne à la troupe de se ranger en 5 colonnes pour que toutes les colonnes aient le même nombre de soldats. C'est alors qu'un général arrive et demande à 5 soldats de sortir des rangs. Le capitaine demande alors à ses hommes de se ranger en 4 colonnes. Combien il y avait-t-il de soldats au premier rassemblement sachant qu'il y en avait entre 50 et 100 ?

Cela revient à trouver les nombres x tels que

$$\begin{cases} x \equiv 0 & [3] \\ x + 2 \equiv 0 & [5] \\ x - 3 \equiv 0 & [4] \end{cases} \iff \begin{cases} x \equiv 0 & [3] \\ x \equiv -2 & [5] \\ x \equiv 3 & [4] \end{cases}$$

Pour cela on va s'intéresser à l'ensemble $\mathbb{Z}/3\mathbb{Z} \times \mathbb{Z}/5\mathbb{Z} \times \mathbb{Z}/4\mathbb{Z}$.

Lemme 19.46

Soit n et m deux entiers, si n divise m l'application $\mathbb{Z} \rightarrow \mathbb{Z}/n\mathbb{Z}$ qui associe à un entier sa classe se factorise par $\mathbb{Z}/m\mathbb{Z}$. De plus l'application $\psi : \mathbb{Z}/m\mathbb{Z} \rightarrow \mathbb{Z}/n\mathbb{Z}$ obtenue est un morphisme d'anneaux.

Démonstration : Pour construire une application de $\mathbb{Z}/m\mathbb{Z}$ dans un ensemble X , on procède comme précédemment, il suffit de construire une application f de $\mathbb{Z}$ dans X tel que si x_1, x_2 sont deux entiers tels que $x_1 \equiv x_2 [m]$ alors $f(x_1) = f(x_2)$.

Dans la suite de la démonstration, pour x un entier relatif, on notera $\bar{x}$ sa classe modulo m et $\tilde{x}$ sa classe modulo n .

Pour tout classe C de $\mathbb{Z}/m\mathbb{Z}$, si on se donne deux représentants x_1 et x_2 de C . Alors $\tilde{x}_1 = \tilde{x}_2$ car comme m divise $x_2 - x_1$, n divise $x_2 - x_1$ puisque n divise m .

On peut donc définir,

$$\begin{aligned} \psi : \mathbb{Z}/m\mathbb{Z} &\rightarrow \mathbb{Z}/n\mathbb{Z} \\ \bar{x} &\mapsto \tilde{x} \end{aligned}$$

Il reste à montrer que ψ est un morphisme d'anneaux.

— Soit $\bar{x}_1, \bar{x}_2$ dans $\mathbb{Z}/m\mathbb{Z}$ on a

$$\psi(\bar{x}_1 + \bar{x}_2) = \psi(\overline{x_1 + x_2}) = \widetilde{x_1 + x_2} = \tilde{x}_1 + \tilde{x}_2 = \psi(\bar{x}_1) + \psi(\bar{x}_2)$$

— De même on a $\psi(\bar{x}_1 \times \bar{x}_2) = \psi(\overline{x_1 \times x_2}) = \widetilde{x_1 \times x_2} = \tilde{x}_1 \times \tilde{x}_2 = \psi(\bar{x}_1) \times \psi(\bar{x}_2)$.

— Pour $x = 1$, on a $\psi(\bar{1}) = \tilde{1}$.

□

ATTENTION

Si n ne divise pas m , cela ne fonctionne plus. Si on essaye par exemple de construire une application de $\mathbb{Z}/6\mathbb{Z}$ dans $\mathbb{Z}/4\mathbb{Z}$ par exemple. Si on regarde la classe $\bar{1}$ de 1 modulo 6, on a $\bar{1} = \bar{7}$ par exemple, mais si on regarde les classes modulo 4, $\tilde{1} \neq \tilde{3} = \tilde{7}$.

Théorème 19.47 (Théorème Chinois)

Soit $(n, m) \in \mathbb{N}^*$. On suppose que n et m sont premiers entre eux le morphisme

$$\pi : \mathbb{Z}/(nm)\mathbb{Z} \rightarrow \mathbb{Z}/n\mathbb{Z} \times \mathbb{Z}/m\mathbb{Z}$$

est un isomorphisme.

Remarque : Ce théorème apparaît dans un ouvrage du mathématicien chinois Qin Jiushao (XIIIe siècle) où il reprend un problème de Sun Zi (IIIe siècle).

Exemple : On peut illustrer cela en regardant pour $n = 3$ et $m = 4$. On considère tous les entiers entre 0 et $11 = 12 - 1$. On peut les répartir selon leur congruence modulo 3 et modulo 4

	0	1	2	3
0	0	9	6	3
1	4	1	10	7
2	8	5	2	11

Démonstration :

- On a déjà vu que l'application π est bien définie. À chaque classe de $\mathbb{Z}/(nm)\mathbb{Z}$, on peut choisir un représentant dans $\mathbb{Z}$. On lui associe alors ses classes dans $\mathbb{Z}/n\mathbb{Z}$ et dans $\mathbb{Z}/m\mathbb{Z}$. Cela ne dépend pas du choix du représentant.
- On a vu que $\mathbb{Z}/(mn)\mathbb{Z} \rightarrow \mathbb{Z}/n\mathbb{Z}$ et $\mathbb{Z}/(mn)\mathbb{Z} \rightarrow \mathbb{Z}/m\mathbb{Z}$ sont des morphismes d'anneaux donc le produit aussi d'après la structure d'anneau produit.
- Il reste à montrer que π est bijectif. Comme

$$|\mathbb{Z}/(nm)\mathbb{Z}| = nm = |\mathbb{Z}/n\mathbb{Z} \times \mathbb{Z}/m\mathbb{Z}|$$

on peut se contenter de montrer l'injectivité. Pour cela on regarde $\text{Ker}\pi$. Soit $C \in \mathbb{Z}/(nm)\mathbb{Z}$ tel que $\pi(C) = 0$. On considère $x \in \mathbb{Z}$ un représentant de C .

On a alors que la classe de x vaut 0 dans $\mathbb{Z}/n\mathbb{Z}$ et dans $\mathbb{Z}/m\mathbb{Z}$ donc $n \mid x$ et $m \mid x$. Comme n et m sont premiers entre eux, $nm \mid x$ ce qui signifie que $C = 0$.

Corollaire 19.48

Soit $n_1, \dots, n_p$ des entiers deux à deux premiers entre eux, il y a un isomorphisme :

$$\mathbb{Z}/(n_1 n_2 \cdots n_p)\mathbb{Z} \rightarrow \prod_{i=1}^p (\mathbb{Z}/n_i \mathbb{Z}).$$

Démonstration : Il suffit de réappliquer le théorème précédent.

ATTENTION

Si n et m ne sont pas premiers entre eux, on peut encore construire

$$\pi : \mathbb{Z}/(nm)\mathbb{Z} \rightarrow \mathbb{Z}/n\mathbb{Z} \times \mathbb{Z}/m\mathbb{Z}$$

car $n \mid nm$ et $m \mid nm$ mais elle n'est plus bijective.

Par exemple pour $n = 4$ et $m = 6$. L'application

$$\pi : \mathbb{Z}/24\mathbb{Z} \rightarrow \mathbb{Z}/4\mathbb{Z} \times \mathbb{Z}/6\mathbb{Z}$$

n'est pas bijective car $\overline{12} \neq 0$ mais $\pi(\overline{1}) = (\overline{0}, \overline{0})$.

Cela permet de revenir à notre problème. En effet 3, 4 et 5 sont deux à deux premiers entre eux. On en déduit que la solution à notre problème qui est l'élément $(\overline{0}, \overline{-2}, \overline{3}) \in \mathbb{Z}/3\mathbb{Z} \times \mathbb{Z}/4\mathbb{Z} \times \mathbb{Z}/5\mathbb{Z}$ est en fait une classe de $\mathbb{Z}/60\mathbb{Z}$. Il reste encore à la trouver. Pour cela il faut déterminer

$$\pi^{-1} : \mathbb{Z}/n\mathbb{Z} \times \mathbb{Z}/m\mathbb{Z} \rightarrow \mathbb{Z}/(nm)\mathbb{Z}.$$

★ **Méthode** : On commence par chercher les antécédents par π de $(0, 1)$ et de $(1, 0)$.

— On cherche donc x tel que $\pi(\overline{x}) = (0, 1)$. C'est-à-dire un entier x tel que

$$\begin{cases} x \equiv 0 & [n] \\ x \equiv 1 & [m] \end{cases}$$

On sait que n et m sont premiers entre eux donc, d'après la relation de Bézout, il existe $(u, v) \in \mathbb{Z}^2$ tels que

$$un + vm = 1$$

En particulier, $x = un$ répond à notre question. De ce fait $\pi^{-1}(0, 1) = un$.

— On cherche donc y tel que $\pi(\overline{y}) = (1, 0)$. De la même manière $\pi^{-1}(1, 0) = vm$.

— Maintenant soit $\overline{r_1}$ une classe dans $\mathbb{Z}/n\mathbb{Z}$ et $\overline{r_2}$ une classe dans $\mathbb{Z}/m\mathbb{Z}$. On a

$$\pi^{-1}(\overline{r_1}, \overline{r_2}) = \overline{r_1 vm + r_2 un}$$

Dans la formule ci-dessus, si on change le représentant r_1 , cela revient à prendre $r'_1 = r_1 + kn$ et dans ce cas,

$$\overline{r'_1 vm + r_2 un} = \overline{r_1 vm + r_2 un + k v n m} = \overline{r_1 vm + r_2 un}$$

De même si on change le représentant r_2 .

Exemple : Revenons à notre problème initial.

On voit que 3 et $20 = 4 \times 5$ sont premiers entre eux et on a que $7 \times 3 - 1 \times 20 = 1$. On en déduit que $r_1 = -20$ vérifie

$$\begin{cases} r_1 \equiv 1 & [3] \\ r_1 \equiv 0 & [4] \\ r_1 \equiv 0 & [5] \end{cases}$$

De même, 4 et $15 = 3 \times 5$ sont premiers entre eux et on a que $4 \times 4 - 1 \times 15 = 1$. On en déduit que $r_2 = -15$ vérifie

$$\begin{cases} r_2 \equiv 0 & [3] \\ r_2 \equiv 1 & [4] \\ r_2 \equiv 0 & [5] \end{cases}$$

Pour finir, 5 et $12 = 3 \times 4$ sont premiers entre eux et on a que $5 \times 5 - 2 \times 12 = 1$. On en déduit que $r_3 = -24$ vérifie

$$\begin{cases} r_3 \equiv 0 & [3] \\ r_3 \equiv 0 & [4] \\ r_3 \equiv 1 & [5] \end{cases}$$

Finalement la classe des solutions à notre problème est la classe de congruence modulo 60 de $0 \times r_1 + 3 \times r_2 - 2 \times r_3 = -45 + 48 = 3$. Le nombre cherché est donc 63.

Exercice : 66

5.3 Indicatrice d'Euler

Définition 19.49

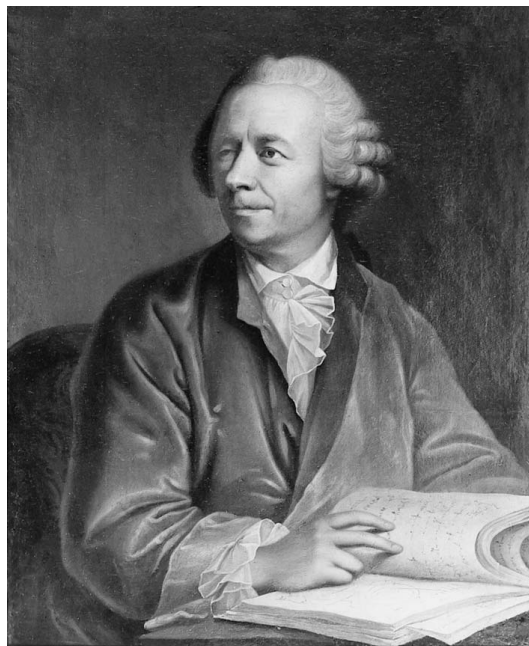
Soit $n \in \mathbb{N}^*$ on note $\varphi(n)$ le nombre d'éléments de $\mathbb{Z}/n\mathbb{Z}$ qui sont inversibles. On a donc

$$\varphi(n) = |(\mathbb{Z}/n\mathbb{Z})^\times| = |\{k \in \llbracket 0; n-1 \rrbracket \mid k \wedge n = 1\}|.$$

La fonction

$$\begin{aligned} \varphi : \mathbb{N}^* &\rightarrow \mathbb{N} \\ n &\mapsto \varphi(n) \end{aligned}$$

s'appelle l'indicatrice d'Euler.



Leonhard Euler
(1707 - 1783)

Exemples :

1. On a $\varphi(1) = 1$ car 0 est premier avec 1 (on a $0 \wedge 1 = 1$).
2. Si n est premier, $\varphi(n) = n - 1$ car tout $k \in \llbracket 1; n-1 \rrbracket$ est premier avec n .
3. Pour les premiers entiers, on a :

n	1	2	3	4	5	6	7	8	9	10	11	12
$\varphi(n)$	1	1	2									

n	1	2	3	4	5	6	7	8	9	10	11	12
$\varphi(n)$	1	1	2	2	4	2	6	4	6	4	10	4

Lemme 19.50

Soit $\pi : (A, +, \times) \rightarrow (B, \oplus, \otimes)$ un isomorphisme d'anneaux.

$$\forall x \in A, x \in A^\times \iff \pi(x) \in B^\times$$

Démonstration :

- $\Rightarrow$ Si x est inversible et soit $y = x^{-1}$ alors $\pi(x) \otimes \pi(y) = \pi(x \times y) = \pi(1) = 1$ donc $\pi(x)$ est inversible dans B .
- $\Leftarrow$ Si $y = \pi(x)$ est inversible dans B , l'élément y^{-1} a un antécédent par π car π est surjective. Notons α cet antécédent, on a $\pi(x \times \alpha) = \pi(x) \otimes \pi(\alpha) = 1_B = \pi(1_A)$. Comme de plus π est injective, on obtient

$$x \times \alpha = 1_A$$

ce qui montre bien que x est inversible

□

Proposition 19.51

Soit n et m deux entiers non nuls premiers entre eux :

$$\varphi(nm) = \varphi(n)\varphi(m)$$

Remarque : On dit que φ est une fonction multiplicative.

Démonstration :

- On utilise le théorème chinois. On sait que $\mathbb{Z}/(nm)\mathbb{Z}$ est isomorphe à $\mathbb{Z}/n\mathbb{Z} \times \mathbb{Z}/m\mathbb{Z}$. Du fait de l'étude des éléments inversibles d'un anneau produit,

$$(\mathbb{Z}/n\mathbb{Z} \times \mathbb{Z}/m\mathbb{Z})^\times = (\mathbb{Z}/n\mathbb{Z})^\times \times (\mathbb{Z}/m\mathbb{Z})^\times.$$

Maintenant, d'après le lemme,

$$x \in (\mathbb{Z}/(nm)\mathbb{Z})^\times \iff \pi(x) \in (\mathbb{Z}/n\mathbb{Z} \times \mathbb{Z}/m\mathbb{Z})^\times.$$

Finalement,

$$\varphi(nm) = |(\mathbb{Z}/(nm)\mathbb{Z})^\times| = |(\mathbb{Z}/n\mathbb{Z} \times \mathbb{Z}/m\mathbb{Z})^\times| = |(\mathbb{Z}/n\mathbb{Z})^\times \times (\mathbb{Z}/m\mathbb{Z})^\times| = \varphi(n)\varphi(m).$$

—

□

Remarque : Cela permet, via la décomposition en facteurs premiers de ramener le calcul de $\varphi(n)$ au calcul de $\varphi(p^s)$ où p est un nombre premier.

Proposition 19.52

Soit p un nombre premier et $s \in \mathbb{N}^*$,

$$\varphi(p^s) = p^s - p^{s-1} = p^s \left(1 - \frac{1}{p}\right).$$

Démonstration : On regarde les p^s nombres compris entre 0 et $p^s - 1$. On cherche ceux qui sont premiers avec p^s . Cela revient à chercher ceux qui sont premiers avec p . Or il n'y a que les p^{s-1} multiples des p qui ne sont pas premiers avec p . Donc

$$\varphi(p^s) = p^s - p^{s-1}.$$

□

Exemple : On veut calculer $\varphi(2020)$. On a $2020 = 2^2 \times 5 \times 101$. De ce fait

$$\varphi(2020) = \varphi(2^2) \times \varphi(5) \times \varphi(101) = 2 \times 4 \times 100 = 800.$$

Théorème 19.53 (Théorème d'Euler)

Soit $n \in \mathbb{N}^*$ et a premier avec n alors

$$a^{\varphi(n)} \equiv 1[n]$$

Remarque : Quand p est premier, $\varphi(p) = p - 1$, on retrouve le petit théorème de Fermat.

Démonstration : Il suffit de considérer le groupe multiplicatif $(\mathbb{Z}/n\mathbb{Z})^\times$. C'est un groupe ayant $\varphi(n)$ élément. Si a est premier avec n , alors $\bar{a} \in (\mathbb{Z}/n\mathbb{Z})^\times$. On sait que l'ordre de $\bar{a}$ divise l'ordre du groupe (théorème de Lagrange) donc $\bar{a}^{\varphi(n)} = \bar{1}$. C'est-à-dire :

$$a^{\varphi(n)} \equiv 1[n].$$

□

Exercice : (Théorème de Wilson)

Montrer que si p est un nombre premier, $(p - 1)! \equiv -1[p]$.

On pourra travailler dans le corps $\mathbb{Z}/p\mathbb{Z}$ où tout élément non nul est inversible

Correction

On veut montrer que dans $F_p = \mathbb{Z}/p\mathbb{Z}$,

$$1 \times 2 \times \cdots \times (p-1) = -1$$

On remarque que pour tout $x \in \llbracket 1; p-1 \rrbracket$ il existe $y \in \llbracket 1; p-1 \rrbracket$ tel que $x \times y = 1$ (modulo p) car x est inversible dans F_p . Dans le produit,

$$1 \times 2 \times \cdots \times (p-1)$$

on peut donc « grouper » tout élément avec son inverse. Il faut juste se méfier des éléments qui sont leur propre inverse ; c'est-à-dire les éléments x tels que $x^2 = 1$. Il n'y en a que deux qui sont 1 et -1 . On en déduit que

$$\prod_{x \in F_p \setminus \{0, 1, -1\}} x = 1$$

Finalement,

$$1 \times 2 \times \cdots \times (p-1) = 1 \times (p-1) = 1 \times (-1) = -1$$

6 Factorisation des polynômes de $K[X]$

C'est un résultat connu que si n est un entier, il peut s'écrire de manière unique comme un produit de nombres premiers. Par exemple $12936 = 2^3 \times 3 \times 7^2 \times 11$. Nous allons essayer de faire de même pour les polynômes de $K[X]$.

6.1 Polynômes irréductibles

Définition 19.54

Un polynôme **non constant** P est dit **irréductible** s'il n'a pas de diviseurs triviaux, c'est-à-dire que P n'est divisible que par les multiples de P et les polynômes constants. Cela s'écrit :

$$\forall Q \in K[X], Q|P \Rightarrow (\exists \lambda \in K, Q = \lambda.P \text{ ou } Q = \lambda).$$

Remarques :

1. C'est un analogue des nombres premiers. Un nombre entier est dit premier s'il n'est divisible que par lui-même et par 1. Avec toujours le fait que les inversibles de $K[X]$ ne sont pas que 1 et -1 mais tous les polynômes constants.
2. On considère qu'un polynôme constant n'est pas irréductible, de la même manière que 1 n'est pas un nombre premier.

Exemples :

1. Les polynômes de degré 1 sont irréductibles.
2. Dans $\mathbb{C}[X]$ il n'y a pas d'autres polynômes irréductibles. C'est le théorème de D'Alembert-Gauss.
3. Dans $\mathbb{R}[X]$ les polynômes irréductibles sont les polynômes de degré 1 ainsi que les polynômes de degré 2, $\alpha X^2 + \beta X + \gamma$ de discriminant strictement négatif.

4. Sur d'autres corps c'est plus compliqué. Par exemple $X^3 - 2$ est un polynôme irréductible dans $\mathbb{Q}[X]$. En effet s'il ne l'était pas il serait divisible par un polynôme de degré 1 et donc il aurait une racine dans $\mathbb{Q}$.
5. Si on regarde $X^5 - 1$ dans $\mathbb{Q}[X]$, on a $(X^5 - 1) = (X - 1)(X^4 + X^3 + X^2 + X + 1)$. Maintenant dans $\mathbb{R}[X]$,

$$X^4 + X^3 + X^2 + X + 1 = (X^2 - 2 \cos(\frac{2\pi}{5})X + 1)(X^2 - 2 \cos(\frac{42\pi}{5})X + 1)$$

Ce qui montre que $X^4 + X^3 + X^2 + X + 1$ est irréductible sur $\mathbb{Q}$.

6. Dans $\mathbb{F}_p = \mathbb{Z}/p\mathbb{Z}$, le polynôme $X^{p-1} - 1$ a $p - 1$ racines. En effet, d'après le petit théorème de Fermat,

$$\forall x \in \llbracket 1 ; p - 1 \rrbracket, x^{p-1} \equiv 1[p]$$

Donc

$$X^{p-1} - 1 = \prod_{i=1}^{p-1} (X - i).$$

Remarque : Attention, contrairement à ce que pourrait faire croire ci qui précède, il n'y a pas équivalence entre être irréductible et ne pas avoir de racines. On a bien

$$(P \text{ est irréductible}) \Rightarrow P \text{ n'a pas de racines.}$$

Cela se démontre bien par contraposée. Mais la réciproque est fausse. Par exemple le polynôme $P = 2X^4 + 6X^2 + 4$ n'a pas de racines réelles de manière évidente. Cependant $P = 2(X^2 + 1)(X^2 + 2)$. Donc P n'est pas irréductible.

6.2 Décomposition en produit de facteurs irréductibles

C'est l'analogue de la décomposition en facteurs premiers dans $\mathbb{Z}$.

Lemme 19.55

Soit P un polynôme irréductible. Il est premier avec tous les polynômes qu'il ne divise pas.

Démonstration : Soit Q un autre polynôme. On suppose que P et Q ne sont pas premiers entre eux. Dans ce cas il existe R non constant qui divise P et Q . Comme R est non constant et qu'il divise P qui est irréductible, il existe $\lambda \in \mathbf{K}^\star$ tel que $P = R\lambda$. On a donc $P|Q$. $\square$

Théorème 19.56

Soit P un polynôme non constant. Il peut s'écrire sous la forme

$$P = \lambda \cdot \prod_{k=1}^p Q_k,$$

où $\lambda \in \mathbf{K}$ et les Q_k sont des polynômes unitaires irréductibles sur $\mathbf{K}$. De plus cette décomposition est unique à l'ordre des facteurs.

Remarque : On peut inclure les polynômes constants en écrivant alors un produit vide.

Démonstration : La démonstration est analogue à celle dans $\mathbb{Z}$.

- Existence : On procède par récurrence forte sur le degré de P . Précisément on pose pour $d \in \mathbf{N}^*$, $\mathcal{P}(d) = \ll \text{tout polynôme de degré } d \text{ admet une décomposition en polynômes irréductibles} \gg$.
 - Soit $d = 1$. La proposition $\mathcal{P}(1)$ est vraie car si P est de degré 1, il est irréductible. Il s'écrit donc λQ où λ est son coefficient dominant.
 - Soit $d \geq 2$. On suppose que pour tout $k \in \llbracket 1; d-1 \rrbracket$ $\mathcal{P}(k)$ est vraie. Soit P un polynôme de degré d . Si P est irréductible, on fait comme ci-dessus en factorisant par son coefficient dominant. S'il n'est pas irréductible, il existe Q_1 un diviseur strict. On écrit alors $P = Q_1 Q_2$. Comme $\deg(Q_1) < \deg(P) = d$ et $\deg(Q_2) < \deg(P) = d$ on peut appliquer l'hypothèse de récurrence à Q_1 et Q_2 puis on fusionne les deux décompositions.
- Par récurrence, on a bien montré que tout polynôme de degré au moins 1 admettait une décomposition en produit de facteurs irréductibles.
- Unicité : Comme dans le cas des entiers, on peut définir pour tout polynôme irréductible Q la valuation Q -adique :

$$\forall P \in \mathbf{K}[X] \setminus \{0\}, v_Q(P) = \max\{k \in \mathbf{N} \mid Q^k \mid P\}$$

Par l'absurde on suppose qu'un polynôme P admet deux décompositions différentes, quitte à les « compléter » par des termes de la forme Q^0 on peut les écrire

$$P = \lambda \prod_{i=1}^p Q_i^{\alpha_i} = \mu \prod_{i=1}^p Q_i^{\beta_i}.$$

Comme λ et μ sont le coefficient dominant de P ils sont égaux. On en déduit que, comme les décompositions sont différentes (à part l'ordre), il existe $i \in \llbracket 1; p \rrbracket$ tel que $\alpha_i \neq \beta_i$. Par symétrie, on peut supposer $i = p$ et que $\alpha_p > \beta_p$. En divisant alors par $Q_p^{\beta_p}$ et par λ on obtient que

$$Q_p^{\alpha_p - \beta_p} \prod_{i=1}^{p-1} Q_i^{\alpha_i} = \prod_{i=1}^{p-1} Q_i^{\beta_i}.$$

Cela implique que Q_p divise $\prod_{i=1}^{p-1} Q_i^{\beta_i}$ ce qui est absurde car pour tout $i \in \llbracket 1; p-1 \rrbracket$, $Q_p \wedge Q_i =$

1 (d'après le lemme) donc Q_p est premier avec $\prod_{i=1}^{p-1} Q_i^{\beta_i}$. La décomposition est bien unique.

□