

PC*

Lycée Marcelin Berthelot

Cours

Chapitre I

Espaces vectoriels

1. Structures vectorielles

La notion d'*espace vectoriel* naît conceptuellement de la géométrie affine avec l'introduction au XVII^e siècle des coordonnées dans un repère du plan ou de l'espace usuel. Les *vecteurs* sont introduits progressivement au cours de la première moitié du XIX^e siècle, et en 1857, Cayley introduit la notation *matricielle*, qui permet d'harmoniser les notations et de simplifier l'écriture des applications linéaires entre espaces vectoriels.

1.1 Espaces vectoriels

Dans tout le chapitre, $\mathbb{K}$ désigne l'un des deux corps $\mathbb{R}$ ou $\mathbb{C}$.

La structure d'*espace vectoriel* sur $\mathbb{K}$ (ou $\mathbb{K}$ -espace vectoriel) a été décrite en première année : un $\mathbb{K}$ -espace vectoriel E dispose de deux opérations, une addition entre vecteurs et une multiplication entre un scalaire et un vecteur. On conviendra de noter 0_E le vecteur nul de E , pour éviter de le confondre avec le scalaire nul 0.

Lorsqu'on veut illustrer graphiquement un concept lié aux espaces vectoriels, on se réfère à la géométrie : à condition de fixer un point qui représentera par convention le vecteur nul, on peut identifier tout point du plan à un unique vecteur d'un espace vectoriel de dimension 2, ou tout point de l'espace à un vecteur d'un espace vectoriel de dimension 3 (illustration figure 1).

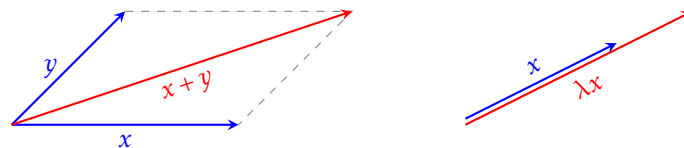


FIGURE 1 – Représentation graphique de l'addition et de la multiplication par un scalaire.

DÉFINITION. (Produit de deux $\mathbb{K}$ -espaces vectoriels) — Si E et F sont deux $\mathbb{K}$ -espaces vectoriels, on munit leur produit cartésien $E \times F$ d'une structure de $\mathbb{K}$ -espace vectoriel en définissant somme et produit externe de la façon suivante :

- (i) pour tout (x, y) et (x', y') dans $E \times F$, $(x, y) + (x', y') = (x + x', y + y')$;
- (ii) pour tout $(x, y) \in E \times F$ et $\lambda \in \mathbb{K}$, $\lambda(x, y) = (\lambda x, \lambda y)$.

Cette définition s'étend naturellement au produit d'un nombre fini quelconque de $\mathbb{K}$ -espaces vectoriels.

■ Quelques exemples de référence

- Le $\mathbb{K}$ -espace vectoriel $\mathbb{K}^p$ est l'espace vectoriel obtenu sur le produit cartésien $\mathbb{K} \times \dots \times \mathbb{K}$, autrement dit sur l'ensemble des p -uplets $(x_1, x_2, \dots, x_p)$ où $x_1, x_2, \dots, x_p$ sont éléments de $\mathbb{K}$;
- le $\mathbb{K}$ -espace vectoriel $\mathbb{K}[X]$ est l'ensemble des polynômes à coefficients dans $\mathbb{K}$;
- le $\mathbb{K}$ -espace vectoriel $\mathcal{M}_{n,p}(\mathbb{K})$ est l'ensemble des matrices n lignes et p colonnes à coefficients dans $\mathbb{K}$.

On conviendra d'identifier les espaces vectoriels $\mathbb{K}^p$ et $\mathcal{M}_{p,1}(\mathbb{K})$, autrement dit de confondre le vecteur $(x_1, \dots, x_p)$

de $\mathbb{K}^p$ avec la matrice colonne $\begin{pmatrix} x_1 \\ \vdots \\ x_p \end{pmatrix}$.

1.2 Sous-espaces vectoriels

Si E est un $\mathbb{K}$ -espace vectoriel, un *sous-espace vectoriel* de E est une partie H non vide et stable par combinaison linéaire. H est alors lui aussi muni d'une structure de $\mathbb{K}$ -espace vectoriel, ce qui justifie sa dénomination.

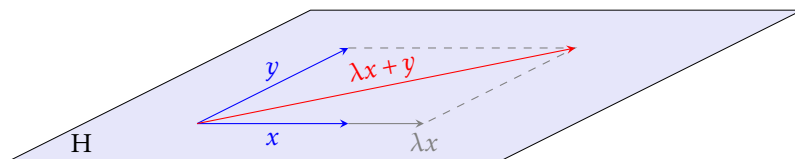


FIGURE 2 – Une représentation graphique en perspective d'un sous-espace vectoriel.

Pour prouver qu'une partie H est un sous-espace vectoriel de E , on utilise le plus souvent le résultat suivant :

PROPOSITION 1.1 — H est un sous-espace vectoriel de E si et seulement si :

- (i) $0_E \in H$ (ou $H \neq \emptyset$);
- (ii) $\forall (x, y) \in H^2, \forall \lambda \in \mathbb{K}, \lambda x + y \in H$.

Exercice 1

Soit E le $\mathbb{R}$ -espace vectoriel des applications de $\mathbb{R}$ dans $\mathbb{R}$. parmi les sous-ensembles suivants, indiquez ceux qui sont des sous-espaces vectoriels de E :

- a. L'ensemble des fonctions 1-périodiques;
- b. l'ensemble des fonctions croissantes;
- c. l'ensemble des fonctions monotones;
- d. l'ensemble des fonctions majorées;
- e. l'ensemble des fonctions bornées;
- f. l'ensemble des fonctions lipschitziennes.

PROPOSITION 1.2 — Soit E un $\mathbb{K}$ -espace vectoriel, et $(H_i)_{i \in I}$ une famille de sous-espaces vectoriels. Alors $\bigcap_{i \in I} H_i$ est un sous-espace vectoriel de E .

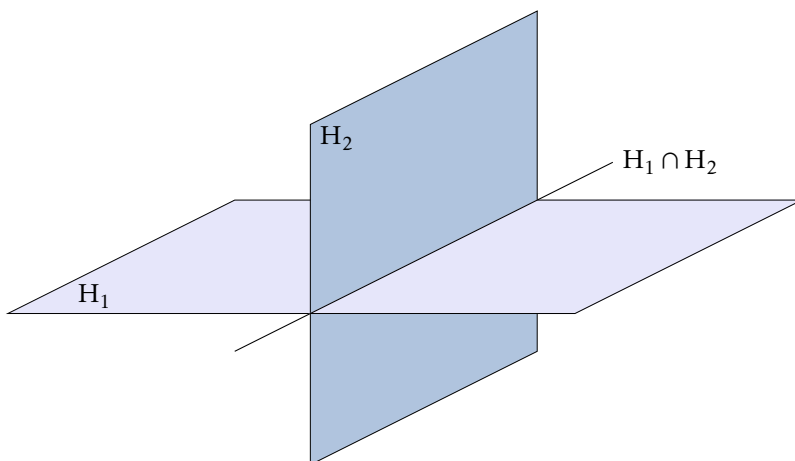


FIGURE 3 – L'intersection de deux sous-espaces vectoriels est un sous-espace vectoriel.

Attention. En revanche, la réunion de deux sous-espaces vectoriels n'est pas, sauf dans le cas trivial où l'un est inclus dans l'autre, un sous-espace vectoriel.

■ Familles génératrices d'un sous-espace vectoriel

DÉFINITION. — Soit E un $\mathbb{K}$ -espace vectoriel, et $\mathcal{A} = \{a_1, \dots, a_n\}$ une famille finie de vecteurs de E . On appelle combinaison linéaire des vecteurs de $\mathcal{A}$ tout vecteur x pouvant s'écrire sous la forme :

$$x = \sum_{i=1}^n \lambda_i a_i \quad \text{avec } (\lambda_1, \dots, \lambda_n) \in \mathbb{K}^n.$$

THÉORÈME 1.3 — L'ensemble des combinaisons linéaires des vecteurs de $\mathcal{A}$ forme un sous-espace vectoriel de E , que l'on note $\text{Vect}(\mathcal{A})$ ou $\text{Vect}(a_1, \dots, a_n)$. C'est le sous-espace vectoriel engendré par la famille $\mathcal{A}$.

À l'inverse, on dira que la famille $\mathcal{A}$ est une famille génératrice du sous-espace vectoriel $\text{Vect}(\mathcal{A})$. Lorsqu'on parle de famille génératrice sans préciser le sous-espace vectoriel dont il est question, c'est qu'il s'agit d'une famille génératrice de l'espace E tout entier.

Remarque. $\text{Vect}(\mathcal{A})$ est le plus petit (au sens de l'inclusion) des sous-espaces vectoriels contenant $\mathcal{A}$.

Remarque. Lorsque $\mathcal{A} = \{a\}$ est composé d'un seul vecteur, on peut écrire $\text{Vect}(a) = \{\lambda a \mid \lambda \in \mathbb{K}\}$ sous la forme plus concise : $\text{Vect}(a) = \mathbb{K}a$.

1.3 Somme de sous-espaces vectoriels

Lorsque H_1 et H_2 sont deux sous-espaces vectoriels d'un même $\mathbb{K}$ -espace vectoriel E , on note

$$H_1 + H_2 = \{x_1 + x_2 \mid x_1 \in H_1 \text{ et } x_2 \in H_2\}.$$

PROPOSITION 1.4 — $H_1 + H_2$ est un sous-espace vectoriel. En outre, si $\mathcal{A}_1$ et $\mathcal{A}_2$ sont des parties génératrices respectivement de H_1 et H_2 , $\mathcal{A}_1 \cup \mathcal{A}_2$ est une partie génératrice de $H_1 + H_2$.

En d'autres termes, $H_1 + H_2$ est le plus petit sous-espace vectoriel (au sens de l'inclusion) contenant H_1 et H_2 .

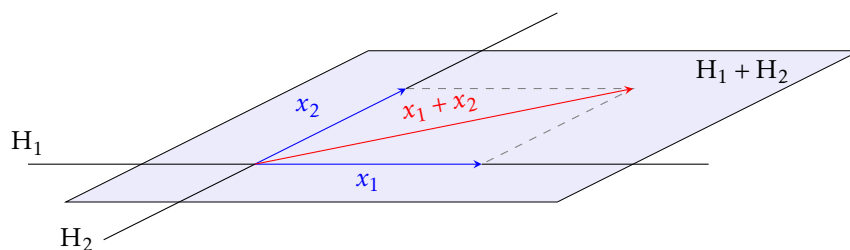


FIGURE 4 – La somme de deux droites vectorielles est en général un plan.

Tout vecteur x de $H_1 + H_2$ peut donc se décomposer sous la forme $x = x_1 + x_2$ avec $x_1 \in H_1$ et $x_2 \in H_2$, mais cette décomposition est-elle unique ? Le résultat suivant a pour objet de répondre à cette question.

PROPOSITION 1.5 — Soient H_1 et H_2 deux sous-espaces vectoriels de E . Il y a équivalence entre :

- (i) $\forall x \in H_1 + H_2, \quad \exists!(x_1, x_2) \in H_1 \times H_2 \mid x = x_1 + x_2$
- (ii) $H_1 \cap H_2 = \{0_E\}$

Autrement dit, pour qu'il y ait unicité de la décomposition, il faut et il suffit que $H_1 \cap H_2 = \{0_E\}$. On dit dans ce cas que la somme $H_1 + H_2$ est directe, et on la note : $H_1 \oplus H_2$.

Pour finir, notons que de cette notion de somme de deux sous-espaces vectoriels découle la notion de sous-espaces supplémentaires :

DÉFINITION. — Lorsque H_1 et H_2 vérifient : $E = H_1 \oplus H_2$, on dit que ces deux sous-espaces sont supplémentaires.

Exercice 2

On considère l'espace vectoriel $E = \mathcal{C}^0([0, 1], \mathbb{R})$ des fonctions continues de $[0, 1]$ dans $\mathbb{R}$. On note H_1 l'ensemble des fonctions constantes et H_2 l'ensemble des fonctions $f \in E$ telles que $\int_0^1 f(t) dt = 0$. Montrer que H_1 et H_2 sont deux sous-espaces vectoriels supplémentaires de E .

L'exemple de la division euclidienne

Considérons l'espace vectoriel $E = \mathbb{K}[X]$ des polynômes à coefficients dans $\mathbb{K}$; il s'agit d'un $\mathbb{K}$ -espace vectoriel. Si M est un polynôme non nul, l'ensemble des multiples de M , noté : $M.\mathbb{K}[X] = \{MQ \mid Q \in \mathbb{K}[X]\}$, est un sous-espace vectoriel de $\mathbb{K}[X]$. En posant $n = \deg M$, l'identité de la division euclidienne affirme pour tout $P \in \mathbb{K}[X]$ l'existence d'un *unique* couple $(Q, R) \in \mathbb{K}[X]^2$ tel que :

$$P = MQ + R \quad \text{et} \quad \deg R \leq n - 1.$$

Autrement dit, tout polynôme P se décompose de manière unique comme somme d'un polynôme $MQ \in M.\mathbb{K}[X]$ et d'un polynôme $R \in \mathbb{K}_{n-1}[X]$. Ainsi, les sous-espaces vectoriels $M.\mathbb{K}[X]$ et $\mathbb{K}_{n-1}[X]$ sont des sous-espaces vectoriels supplémentaires de $\mathbb{K}[X]$. On peut donc écrire $\mathbb{K}[X] = M.\mathbb{K}[X] \oplus \mathbb{K}_{n-1}[X]$ lorsque $n = \deg M$.

■ Projections vectorielles

Considérons deux sous-espaces vectoriels supplémentaires H_1 et H_2 de E : $E = H_1 \oplus H_2$. Pour tout $x \in E$, il existe un unique couple $(x_1, x_2) \in H_1 \times H_2$ tel que $x = x_1 + x_2$. On définit l'application $p : E \rightarrow E$ qui à tout $x \in E$ associe $p(x) = x_1$; il s'agit de la *projection vectorielle* sur H_1 *parallèlement à* H_2 .

On a $H_1 = \text{Im } p = \text{Ker}(p - \text{Id}_E)$ et $H_2 = \text{Ker } p$ donc on peut écrire : $E = \text{Ker } p \oplus \text{Ker}(p - \text{Id}_E)$.

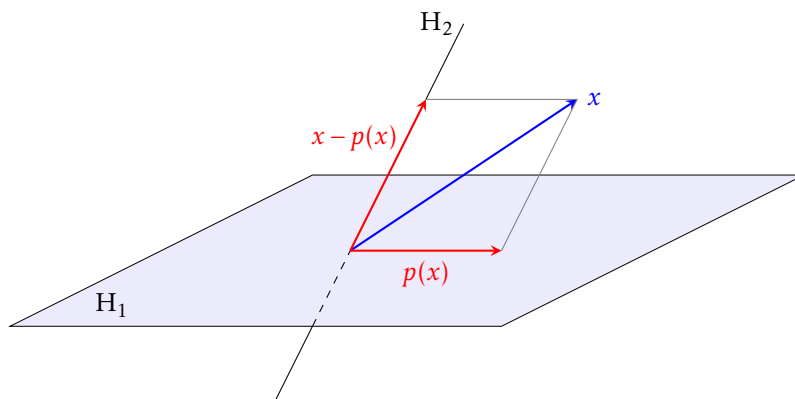


FIGURE 5 – La projection sur H_1 parallèlement à H_2 .

Remarque. Si p est la projection vectorielle sur H_1 parallèlement à H_2 , alors $\text{Id}_E - p$ est la projection sur H_2 parallèlement à H_1 .

THÉORÈME 1.6 — Un endomorphisme $p \in \mathcal{L}(E)$ est une projection vectorielle si et seulement si $p \circ p = p$. Dans ce cas, p est la projection sur $\text{Im } p = \text{Ker}(p - \text{Id}_E)$ parallèlement à $\text{Ker } p$.

Exercice 3

On considère deux projections p et q d'un même espace vectoriel E . Montrer que $\text{Im } p = \text{Im } q$ si et seulement si $p \circ q = q$ et $q \circ p = p$.

Donner une condition analogue pour caractériser l'égalité $\text{Ker } p = \text{Ker } q$.

■ Somme de plusieurs sous-espaces vectoriels

Si $H_1, \dots, H_p$ sont des sous-espaces vectoriels de E , on peut définir de manière analogue leur somme :

$$H_1 + H_2 + \dots + H_p = \{x_1 + x_2 + \dots + x_p \mid x_i \in H_i, 1 \leq i \leq p\}.$$

Lorsque la décomposition d'un vecteur $x \in H_1 + H_2 + \dots + H_p$ est unique, on dira que cette somme est *directe*, et on la notera $H_1 \oplus H_2 \oplus \dots \oplus H_p$.

Comment caractériser une somme directe ? Pour répondre à cette question, on peut adopter une démarche récursive en écrivant : $x = \underbrace{(x_1 + x_2 + \dots + x_{p-1})}_{\in H_1 + H_2 + \dots + H_{p-1}} + \underbrace{x_p}_{\in H_p}$

Ainsi, la somme est directe si et seulement si les sommes $H = H_1 + H_2 + \dots + H_{p-1}$ et $H + H_p$ sont directes. Cela conduit au résultat suivant :

THÉORÈME 1.7 — La somme $H_1 + H_2 + \dots + H_p$ est directe si et seulement si :

- (i) la somme $H_1 \oplus H_2 \oplus \dots \oplus H_{p-1}$ est directe ;
- (ii) $(H_1 \oplus H_2 \oplus \dots \oplus H_{p-1}) \cap H_p = \{0_E\}$.

Attention. Il n'existe pas de critère simple pour vérifier qu'une somme de $n \geq 3$ sous-espaces vectoriels est directe. Ou bien on justifie l'unicité de la décomposition directement, ou bien on procède récursivement à l'aide du résultat précédent. Par exemple, pour prouver qu'une somme $H_1 + H_2 + H_3$ est directe il faut prouver successivement les deux égalités : $H_1 \cap H_2 = \{0_E\}$ puis $(H_1 \oplus H_2) \cap H_3 = \{0_E\}$.

■ Famille de projecteurs associée à une somme directe

Considérons maintenant une famille $(H_1, \dots, H_n)$ de sous-espaces vectoriels vérifiant : $E = \bigoplus_{k=1}^n H_k$. Tout vecteur $x \in E$ s'écrit de manière unique : $x = \sum_{k=1}^n x_k$, avec $x_k \in H_k$. On peut donc définir les endomorphismes $p_k : x \mapsto x_k$ pour $1 \leq k \leq p$. Ainsi, p_k est la projection vectorielle sur H_k parallèlement à $\bigoplus_{i \neq k} H_i$.

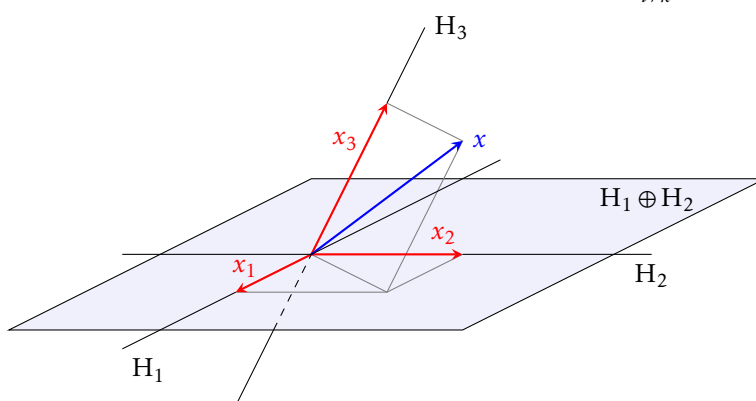


FIGURE 6 – x_3 est la projection sur H_3 parallèlement à $H_1 \oplus H_2$.

■ Familles libres

DÉFINITION. — Une famille finie $(a_1, \dots, a_n)$ de vecteurs non nuls de E est dite libre lorsque la somme $\mathbb{K}a_1 + \dots + \mathbb{K}a_n$ est directe, c'est à dire lorsque tout vecteur x appartenant à cette somme se décompose de manière unique sous la forme :

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

On dit encore que les vecteurs $a_1, \dots, a_n$ sont linéairement indépendants. Une famille qui n'est pas libre est dite liée.

Il existe essentiellement trois manières de prouver la liberté d'une famille de vecteurs : on peut bien entendu recourir à la définition en justifiant l'unicité de la décomposition, ou utiliser l'un des deux résultats suivants.

PROPOSITION 1.8 — La famille $(a_1, \dots, a_n)$ est libre si et seulement si elle vérifie la propriété :

$$(i) \quad \forall (\lambda_1, \dots, \lambda_n) \in \mathbb{K}^n, \quad \sum_{i=1}^n \lambda_i a_i = 0_E \implies \lambda_1 = \dots = \lambda_n = 0.$$

Exercice 4

Soit E un espace vectoriel, et $f \in \mathcal{L}(E)$. On suppose l'existence d'un vecteur $x \in E$ et d'un entier n tel que $f^{n-1}(x) \neq 0_E$ et $f^n(x) = 0_E$. Montrer que la famille $(x, f(x), f^2(x), \dots, f^{n-1}(x))$ est libre.

Le second résultat adopte une approche récursive :

PROPOSITION 1.9 — Soit $(a_1, \dots, a_n)$ une famille libre, et $a_{n+1} \in E$. Alors $(a_1, \dots, a_{n+1})$ est libre si et seulement si $a_{n+1} \notin \text{Vect}(a_1, \dots, a_n)$.

Autrement dit, pour prouver que la famille $(a_1, \dots, a_n)$ est libre il suffit de prouver que $(a_1, \dots, a_{n-1})$ est libre puis que a_n n'est pas combinaison linéaire des vecteurs $a_1, \dots, a_{n-1}$.

Exercice 5

On considère n réels ordonnés $\alpha_1 < \alpha_2 < \dots < \alpha_n$ ainsi que les fonctions $f_i = x \mapsto e^{\alpha_i x}$ de $\mathcal{F}(\mathbb{R}, \mathbb{R})$. Prouver par récurrence que les fonctions $(f_1, f_2, \dots, f_n)$ forment une famille libre.

1.4 Bases d'un espace vectoriel

DÉFINITION. — Une base $(e_1, \dots, e_p)$ est une famille libre et génératrice de E , c'est à dire lorsque tout vecteur x de E se décompose de manière unique sous la forme :

$$x = \sum_{i=1}^p x_i e_i \quad \text{avec} \quad (x_1, \dots, x_p) \in \mathbb{K}^p.$$

On a donc dans ce cas : $E = \mathbb{K}e_1 \oplus \dots \oplus \mathbb{K}e_p$.

Ainsi, le caractère *générateur* de la famille (e) traduit l'existence de la décomposition de tout vecteur de E , le caractère *libre*, l'unicité de cette décomposition.

Remarque. Les liens entre base et décomposition de l'espace en somme directe sont profonds : si on dispose d'une décomposition de E en somme directe $E = H_1 \oplus H_2 \oplus \dots \oplus H_p$, on obtient une base de (e) en réunissant des bases de chacun des sous-espaces vectoriels $H_1, H_2, \dots, H_p$.

Plus formellement, si $(e_1, \dots, e_{i_1})$ est une base de H_1 , $(e_{i_1+1}, \dots, e_{i_2})$ une base de H_2 , $\dots$, $(e_{i_{p-1}-1}, \dots, e_p)$ une base de H_p , alors $(e_1, \dots, e_p)$ est une base de E . Une telle base sera dite *adaptée* à la décomposition en somme directe $E = H_1 \oplus \dots \oplus H_p$.

À l'inverse, à partir d'une base $(e_1, \dots, e_p)$ de E on peut obtenir une décomposition en somme directe de E en fractionnant cette base. Si on considère par exemple un entier $k \in \llbracket 1, p-1 \rrbracket$ et si on pose $H_1 = \text{Vect}(e_1, \dots, e_k)$ et $H_2 = \text{Vect}(e_{k+1}, \dots, e_p)$ on obtient une décomposition de E en somme directe de deux sous-espaces supplémentaires $E = H_1 \oplus H_2$.

■ Dimension d'un espace vectoriel

Dans le cours de première année a été prouvé un résultat important : si un espace vectoriel contient une base de cardinal fini, toutes les autres bases ont même cardinal, appelé dimension de l'espace vectoriel.

Les conséquences de ce résultat sont nombreuses, et en particulier :

PROPOSITION 1.10 — Si E est un $\mathbb{K}$ -espace vectoriel de dimension p , toute famille libre (respectivement génératrice) de cardinal p est une base.

En outre, toute famille génératrice contient au moins p éléments, et toute famille libre contient au plus p éléments.

THÉORÈME 1.11 (de la base incomplète) — Soit (e) une famille libre et (g) une famille génératrice d'un espace vectoriel E . Alors il existe une base (b) telle que $(e) \subset (b) \subset (e \cup g)$. Autrement dit, on peut « compléter » une famille libre par certains éléments d'une famille génératrice pour former une base.

Cet énoncé possède une version simplifiée (en prenant pour (g) l'ensemble des vecteurs de E , puis en prenant pour (e) l'ensemble vide) :

COROLLAIRE — Toute famille libre peut être complétée pour former une base de E (théorème de la base incomplète); de toute famille génératrice on peut extraire une base de E (théorème de la base extraite).

Une application fréquente du théorème de la base incomplète consiste, à partir d'une base $(e_1, \dots, e_k)$ d'un sous-espace vectoriel H de E , à compléter celle-ci pour obtenir une base $(e_1, \dots, e_k, e_{k+1}, \dots, e_p)$ de E . Une telle base est dite *adaptée* à H .

PROPOSITION 1.12 — Si E et F sont deux $\mathbb{K}$ -espaces vectoriels de dimensions finies, il en est de même de $E \times F$, et $\dim(E \times F) = \dim(E) + \dim(F)$.

COROLLAIRE — On en déduit par une récurrence immédiate que si $E_1, E_2, \dots, E_k$ sont des $\mathbb{K}$ -espaces vectoriels de dimensions finies, il en est de même de $E_1 \times \dots \times E_k$, et $\dim(E_1 \times \dots \times E_k) = \sum_{i=1}^k \dim E_i$.

PROPOSITION 1.13 (Formule de Grassmann) — Si H_1 et H_2 sont deux sous-espaces vectoriels d'un $\mathbb{K}$ -espace vectoriel de dimension finie, alors $\dim(H_1 + H_2) = \dim H_1 + \dim H_2 - \dim(H_1 \cap H_2)$.

Il existe une formule qui généralise la formule de Grassmann au cas d'une somme de k sous-espaces vectoriels, mais elle est trop compliquée pour être utilisable en pratique. On se contentera donc du résultat suivant :

PROPOSITION 1.14 — Si $H_1, \dots, H_k$ sont des sous-espaces vectoriels de dimensions finies, il en est de même de leur somme, et $\dim\left(\sum_{i=1}^k H_i\right) \leq \sum_{i=1}^k \dim H_i$, avec égalité si et seulement si la somme est directe.

Remarque. Ceci donne un moyen alternatif pour prouver qu'une somme est directe, pour peu qu'on sache calculer la dimension de la somme.

■ Représentation matricielle des vecteurs en dimension finie

Matrice associée à un vecteur

Étant donnée une base $(e_1, \dots, e_p)$ de E , l'application : $\phi : \mathcal{M}_{p,1}(\mathbb{K}) \rightarrow E$ qui à une matrice colonne $X = \begin{pmatrix} x_1 \\ \vdots \\ x_p \end{pmatrix}$ associe le vecteur $\sum_{k=1}^p x_k e_k$ est un isomorphisme. Pour tout $x \in E$, $X = \phi^{-1}(x)$ est la *matrice des composantes* de x dans la base (e) , et sera notée : $X = \text{Mat}_e(x)$.

Matrice associée à une famille de vecteurs

Si $(x_1, \dots, x_k)$ est une famille de vecteurs de E et $X_1, \dots, X_k$ les matrices colonnes associées à ces vecteurs dans la base (e) , on appelle *matrice associée* à la famille $(x_1, \dots, x_k)$ dans la base (e) la matrice $A \in \mathcal{M}_{p,k}(\mathbb{K})$ formée des colonnes $X_1, \dots, X_k$:

$$A = \begin{pmatrix} \uparrow & & \uparrow \\ X_1 & \dots & X_k \\ \downarrow & & \downarrow \end{pmatrix} = \begin{pmatrix} e_1 \\ \vdots \\ e_i \\ \vdots \\ e_p \end{pmatrix} \begin{pmatrix} x_1 & \dots & x_j & \dots & x_k \end{pmatrix}$$

$x_1 \dots x_j \dots x_k$

→ i^{e} coefficient de x_j dans la base (e)

Le *rang* de la famille $(x_1, \dots, x_k)$ est la dimension de l'espace vectoriel qu'ils engendrent ; on a donc $\text{rg}(x_1, \dots, x_k) = \text{rg}(X_1, \dots, X_k) = \text{rg } A$.

Exemple. Considérons $E = \mathbb{R}^4$ et notons (e) la base canonique. Définissons les quatre vecteurs :

$$a = (1, 2, 3, 4), \quad b = (1, 1, 1, 3), \quad c = (2, 1, 1, 1), \quad d = (3, 1, 0, 3)$$

et posons $H = \text{Vect}(a, b, c, d)$. Quelle est la dimension de H ? Pour répondre à cette question, posons $A = \text{Mat}_{(e)}(a, b, c, d)$ et calculons $\text{rg}(A)$ en appliquant la méthode de Gauss-Jordan sur les colonnes de A :

$$A = \begin{pmatrix} 1 & 1 & 2 & 3 \\ 2 & 1 & 1 & 1 \\ 3 & 1 & 1 & 0 \\ 4 & 3 & 1 & 3 \end{pmatrix}$$

On réalise les opérations $C_2 \leftarrow C_2 - C_1$, $C_3 \leftarrow C_3 - 2C_1$, $C_4 \leftarrow C_4 - 3C_1$:

$$\text{rg } A = \text{rg} \begin{pmatrix} 1 & 0 & 0 & 0 \\ 2 & -1 & -3 & -5 \\ 3 & -2 & -5 & -9 \\ 4 & -1 & -7 & -9 \end{pmatrix}$$

Réalisons maintenant les opérations $C_3 \leftarrow C_3 - 3C_2$, $C_4 \leftarrow C_4 - 5C_2$:

$$\text{rg } A = \text{rg} \begin{pmatrix} 1 & 0 & 0 & 0 \\ 2 & -1 & 0 & 0 \\ 3 & -2 & 1 & 1 \\ 4 & -1 & -4 & -4 \end{pmatrix}$$

Et enfin l'opération $C_4 \leftarrow C_4 - C_3$:

$$\text{rg } A = \text{rg} \begin{pmatrix} 1 & 0 & 0 & 0 \\ 2 & -1 & 0 & 0 \\ 3 & -2 & 1 & 0 \\ 4 & -1 & -4 & 0 \end{pmatrix}$$

Ainsi, H est un sous-espace vectoriel de dimension $\text{rg } A = 3$, et la famille (a, b, c, d) est une famille génératrice qui n'est pas libre : il ne s'agit pas d'une base.

Pourquoi avoir agi sur les colonnes plutôt que sur les lignes¹ ? La matrice A est la matrice de quatre vecteurs de H ; toute combinaison linéaire de ces vecteurs donne de nouveaux vecteurs de H . Ainsi, *réaliser des opérations élémentaires sur les colonnes de A crée de nouvelles familles de vecteurs de H sans en modifier le rang*. Les vecteurs qui apparaissent dans la matrice finale sont donc toujours des vecteurs générateurs de H , mais cette fois les trois premiers forment une famille libre, et donc une base de H : la famille (a, b', c') avec $b' = (0, -1, -2, -1)$ et $c' = (0, 0, 1, -4)$ est une base de H .

1. Rappelons que les opérations élémentaires sur les lignes comme sur les colonnes ne modifient pas le rang.

Exercice 6

On considère l'espace vectoriel $E = \mathbb{K}_n[X]$ des polynômes de degré inférieur ou égal à n , ainsi que la famille de vecteurs $(P_0, \dots, P_n)$ définie par : $P_k = X^k(1 - X)^{n-k}$. Quelle forme particulière prend la matrice associée à la famille (P) dans la base canonique ? En déduire que (P) est une base de E .

Par un raisonnement analogue, prouver que toute famille de polynômes $(Q_0, \dots, Q_n)$ vérifiant $\deg Q_k = k$, $0 \leq k \leq n$, est une base de E .

Matrice de passage entre deux bases

Considérons un $\mathbb{K}$ -espace vectoriel E de dimension finie p , et (e) et (e') deux bases. Nous qualifierons la base (e) d'*ancienne base*, et (e') de *nouvelle base*.

Étant donné un vecteur $x \in E$, on souhaite exprimer ses nouvelles coordonnées $X' = \text{Mat}_{(e')}(x)$ en fonction de ses anciennes coordonnées $X = \text{Mat}_{(e)}(x)$.

On suppose connaître l'expression des vecteurs de la nouvelle base (e') dans l'ancienne base (e) :

$$\forall j \in \llbracket 1, p \rrbracket, \quad e'_j = \sum_{i=1}^p \lambda_{ij} e_i$$

ce qui revient à considérer la matrice $P = \text{Mat}_e(e'_1, \dots, e'_p) = (\lambda_{ij}) \in \mathcal{M}_p(\mathbb{K})$. On dit que P est la *matrice de passage* de (e) vers (e') .

THÉORÈME 1.15 (formule de changement de base) — La matrice $P = \text{Mat}_e(e')$ est une matrice inversible, et la formule de changement de base s'exprime sous la forme : $X' = P^{-1}X$.

Remarque. De l'égalité $X = PX' = (P^{-1})^{-1}X'$ il résulte que P^{-1} est la matrice de passage de (e') vers (e) .

2. Applications linéaires

2.1 Rappels

Une application linéaire est une application entre deux espaces vectoriels qui respecte l'addition des vecteurs et la multiplication scalaire, ou, en d'autres termes, qui préserve les combinaisons linéaires. On adoptera donc la définition suivante :

DÉFINITION. — Soit E et F deux $\mathbb{K}$ -espaces vectoriels, et $u : E \rightarrow F$ une application. On dit que u est linéaire lorsque : $\forall (x, y) \in E^2, \forall \lambda \in \mathbb{K}, u(\lambda x + y) = \lambda u(x) + u(y)$.

On note $\mathcal{L}(E, F)$ le $\mathbb{K}$ -espace vectoriel des applications linéaires de E vers F ; si E et F sont de dimensions finies, la dimension de cet espace vectoriel est égal à $\dim E \times \dim F$.

Enfin, lorsque $F = E$ on notera $\mathcal{L}(E) = \mathcal{L}(E, E)$, et les éléments de $\mathcal{L}(E)$ seront appelés des *endomorphismes*.

■ Matrice associée à une application linéaire

Soit E un $\mathbb{K}$ -espace vectoriel de dimension p , et F un $\mathbb{K}$ -espace vectoriel de dimension n . On note $(e_1, \dots, e_p)$ une base de E , et $(f_1, \dots, f_n)$ une base de F .

À une application linéaire $u \in \mathcal{L}(E, F)$ on associe la matrice $A = \text{Mat}_f(u(e_1), \dots, u(e_p))$ (la matrice des composantes des vecteurs $u(e_1), \dots, u(e_p)$ dans la base (f)), matrice que l'on note $\text{Mat}_{e,f}(u)$. On a donc :

$$A = \text{Mat}_{e,f}(u) = \begin{matrix} & \begin{matrix} f_1 & & & & f_n \end{matrix} \\ \begin{matrix} \vdots \\ \vdots \\ \vdots \end{matrix} & \begin{pmatrix} a_{11} & \dots & a_{1j} & \dots & a_{1p} \\ \vdots & & \vdots & & \vdots \\ a_{n1} & \dots & a_{nj} & \dots & a_{np} \end{pmatrix} \end{matrix}$$

→ coordonnées de $u(e_j)$ dans la base (f)

$u(e_1) \quad \dots \quad u(e_j) \quad \dots \quad u(e_p)$

Remarque. L'application $\phi : \mathcal{L}(E, F) \rightarrow \mathcal{M}_{np}(\mathbb{K})$ définie par $\phi(u) = \text{Mat}_{e,f}(u)$ établit un isomorphisme entre $\mathcal{L}(E, F)$ et $\mathcal{M}_{n,p}(\mathbb{K})$; c'est ce résultat qui permet de justifier sans peine que $\dim \mathcal{L}(E, F) = np = \dim E \times \dim F$.

Exercice 7

Soient E et F deux espaces vectoriels de dimensions finies, et $u \in \mathcal{L}(E, F)$.

On pose $\mathcal{H} = \{v \in \mathcal{L}(F, E) \mid v \circ u = 0\}$.

Soit $v \in \mathcal{H}$. Quelle particularité possède la matrice associée à v dans une base adaptée à $\text{Im } u$? En déduire l'expression de $\dim \mathcal{H}$ en fonction des dimensions de E et de F et du rang de u .

L'application d'une application linéaire à un vecteur est lié au produit matriciel par le résultat suivant :

THÉORÈME 2.1 — Si $x \in E$, on pose $X = \text{Mat}_e(x)$ et $Y = \text{Mat}_f(u(x))$. Alors $Y = AX$.

Formule de changement de base pour les applications linéaires

Soient (e') et (f') deux nouvelles bases, respectivement de E et F . On note $P \in \mathcal{GL}_p(\mathbb{K})$ la matrice de passage de (e) vers (e') et $Q \in \mathcal{GL}_n(\mathbb{K})$ la matrice de passage de (f) vers (f') .

On note $A' = (a'_{ij}) = \text{Mat}_{e',f'}(u)$ la matrice associée à l'application linéaire u dans les nouvelles bases (e') et (f') .

On souhaite exprimer A' en fonction de A , matrice associée à u dans les anciennes bases (e) et (f) .

La matrice associée à $u(x)$ est égale à AX dans la base (f) , et à $A'X'$ dans la base (f') . Des formules de changement de base pour les vecteurs on déduit que : $A'X' = Q^{-1}AX$. Or $X' = P^{-1}X$, donc : $A'P^{-1}X = Q^{-1}AX$. Ceci étant vrai pour tout $X \in \mathcal{M}_{p,1}(\mathbb{K})$, on en déduit que $A'P^{-1} = Q^{-1}A$, soit : $A' = Q^{-1}AP$.

Exemple. On pose $E = \mathbb{R}^4$, $F = \mathbb{R}^3$, on note (e) et (f) les bases canoniques respectivement de E et F , et on considère l'application linéaire $u \in \mathcal{L}(E, F)$ définie par $\text{Mat}_{e,f}(u) = \begin{pmatrix} 4 & 5 & -7 & 7 \\ 2 & 1 & -1 & 3 \\ 1 & -1 & 2 & 1 \end{pmatrix} = A$. On souhaite obtenir la matrice $A' = \text{Mat}_{e',f'}(u)$ relative aux changements de bases définis par :

$$\begin{cases} e'_1 = e_1 \\ e'_2 = e_2 \\ e'_3 = 4e_1 + e_2 - 3e_4 \\ e'_4 = -7e_1 + e_3 + 5e_4 \end{cases} \quad \text{et} \quad \begin{cases} f'_1 = 4f_1 + 2f_2 + f_3 \\ f'_2 = 5f_1 + f_2 - f_3 \\ f'_3 = f_3 \end{cases}$$

Pour obtenir A' nous avons deux possibilités :

(i) définir les matrices $P = \text{Mat}_{(e)}(e') = \begin{pmatrix} 1 & 0 & 4 & -7 \\ 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & -3 & 5 \end{pmatrix}$ et $Q = \text{Mat}_{(f)}(f') = \begin{pmatrix} 4 & 5 & 0 \\ 2 & 1 & 0 \\ 1 & -1 & 1 \end{pmatrix}$ et calculer $Q^{-1}AP$;

(ii) exprimer directement les vecteurs $u(e'_i)$ dans la base (f') .

La première méthode s'avère longue en calcul ; elle ne sera quasiment jamais employée.

La seconde méthode peut s'avérer elle aussi fastidieuse, sauf si les vecteurs (e') et (f') ont été judicieusement choisis, ce qui s'avérera en général le cas.

Et en effet :

$$\begin{aligned} u(e'_1) &= u(e_1) = f'_1 & u(e'_3) &= 4u(e_1) + u(e_2) - 3u(e_4) = 0_E \\ u(e'_2) &= u(e_2) = f'_2 & u(e'_4) &= -7u(e_1) + u(e_3) + 5u(e_4) = 0_E \end{aligned}$$

donc $A' = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix} = \text{Mat}_{e',f'}(u)$ et on peut affirmer que $A = QA'P^{-1}$ sans avoir besoin de réaliser le calcul.

Le théorème 2.4 permettra d'expliquer la façon dont ont été choisies les bases (e') et (f') .

Formule de changement de base pour les endomorphismes

Il s'agit d'un cas particulier du précédent, avec : $F = E$, $(f) = (e)$, $(f') = (e')$. On obtient : $A' = P^{-1}AP$.

Deux matrices A et A' liées par une relation de ce type sont dites *semblables*. Garder toujours à l'esprit que deux matrices semblables sont deux matrices qui peuvent être associées au même endomorphisme, mais exprimées dans des bases différentes.

Exercice 8

Soit $A \in \mathcal{M}_n(\mathbb{K})$ une matrice vérifiant $A^n = 0$ et $A^{n-1} \neq 0$. Montrer que la matrice A est semblable à la matrice

$$A' = \begin{pmatrix} 0 & 1 & 0 & \cdots & 0 \\ & \ddots & \ddots & \ddots & \\ & & 0 & 1 & 0 \\ & & & \ddots & \ddots \\ 0 & \cdots & & & 0 \end{pmatrix}$$

En déduire que les matrices $A = \begin{pmatrix} 2 & 2 & -3 \\ 5 & 1 & -5 \\ -3 & 4 & 0 \end{pmatrix}$ et $T = \begin{pmatrix} 1 & 1 & 0 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{pmatrix}$ sont semblables, puis calculer explicitement une matrice P vérifiant $A = PTP^{-1}$.

■ Trace d'un endomorphisme

DÉFINITION. — On appelle trace d'une matrice carrée $A = (a_{ij}) \in \mathcal{M}_p(\mathbb{K})$ le scalaire : $\text{tr } A = \sum_{i=1}^p a_{ii}$, c'est à dire la somme des éléments diagonaux de cette matrice.

On définit ainsi une forme linéaire sur l'espace $\mathcal{M}_p(\mathbb{K})$ des matrices carrées d'ordre p , autrement dit une application linéaire de $\mathcal{M}_p(\mathbb{K})$ dans $\mathbb{K}$. Cette forme linéaire va pouvoir à son tour être définie sur l'espace $\mathcal{L}(E)$ des endomorphismes d'un espace vectoriel E de dimension finie grâce au résultat suivant, et surtout son corollaire :

PROPOSITION 2.2 — Si $(A, B) \in \mathcal{M}_p(\mathbb{K})^2$ on a $\text{tr}(AB) = \text{tr}(BA)$.

COROLLAIRE — Si $A \in \mathcal{M}_p(\mathbb{K})$ et $P \in \mathcal{GL}_p(\mathbb{K})$ alors $\text{tr}(P^{-1}AP) = \text{tr } A$.

Du corollaire précédent on déduit que si $u \in \mathcal{L}(E)$ et $A = \text{Mat}_e(u)$, alors $\text{tr } A$ ne dépend pas du choix de la base (e) . On peut donc définir la trace de u par l'intermédiaire de la trace d'une matrice associée à A dans une base quelconque :

DÉFINITION. — Si E est un $\mathbb{K}$ -espace vectoriel de dimension finie et $u \in \mathcal{L}(E)$ un endomorphisme de E , on appelle trace de u la trace de la matrice $\text{Mat}_e(u)$, où (e) est une base quelconque de E .

L'application $u \mapsto \text{tr } u$ est une forme linéaire sur $\mathcal{L}(E)$, autrement dit une application linéaire de $\mathcal{L}(E)$ dans $\mathbb{K}$. De la proposition 2.2 il résulte :

COROLLAIRE — Si u et v sont deux endomorphismes d'un même $\mathbb{K}$ -espace vectoriel E , alors $\text{tr}(u \circ v) = \text{tr}(v \circ u)$.

Base canonique de $\mathcal{M}_{np}(\mathbb{K})$

Il s'agit bien entendu de la base $(E_{ij})_{\substack{1 \leq i \leq n \\ 1 \leq j \leq p}}$ formée des matrices dont tous les coefficients sont nuls sauf un, égal à 1 :

$$E_{ij} = \begin{pmatrix} 0 & \cdots & 0 \\ \vdots & & \vdots \\ 0 & \cdots & 0 \end{pmatrix} \quad \begin{matrix} \text{1} \text{ à la position } (i,j) \end{matrix}$$

Il est bon de connaître la formule donnant le produit de deux matrices de cette forme ; c'est le résultat suivant :

$$E_{ij}E_{kl} = \delta_{j,k}E_{il}.$$

où $\delta_{j,k}$ désigne le *symbole de Kronecker* : $\delta_{j,k} = \begin{cases} 1 & \text{si } j = k \\ 0 & \text{si } j \neq k \end{cases}$.

Exercice 9

En utilisant la base canonique de $\mathcal{M}_p(\mathbb{K})$, prouver que toute forme linéaire $\phi : \mathcal{M}_p(\mathbb{K}) \rightarrow \mathbb{K}$ vérifiant : $\forall (A, B) \in \mathcal{M}_p(\mathbb{K})^2, \phi(AB) = \phi(BA)$ est proportionnelle à la trace.

2.2 Image et noyau d'une application linéaire

Nous allons maintenant nous intéresser aux liens qui existent entre sous-espaces vectoriels et applications linéaires.

PROPOSITION 2.3 — Soit $u : E \rightarrow F$ une application linéaire, H_1 et H_2 des sous-espaces vectoriels respectivement de E et F . Alors $u(H_1)$ et $u^{-1}(H_2)$ sont respectivement des sous-espaces vectoriels de F et de E .

Attention. Attention à la notation $u^{-1}(H_2)$, qui pourrait faire croire à tort que u est supposée bijective. Il n'en est rien, il s'agit de la notion d'*image réciproque* définie par :

$$u^{-1}(H_2) = \{x \in E \mid u(x) \in H_2\}.$$

Exemples. En appliquant cette propriété aux sous-espaces vectoriels $H_1 = E$ et $H_2 = \{0_F\}$, on définit *image* et *noyau* d'une application linéaire :

$\text{Im } u = u(E) = \{y \in F \mid \exists x \in E \text{ tel que } u(x) = y\}$ est un sous-espace vectoriel de F (l'*image* de u) ;

$\text{Ker } u = u^{-1}(\{0_F\}) = \{x \in E \mid u(x) = 0_F\}$ est un sous-espace vectoriel de E (le *noyau* de u).

Rappelons que ces deux sous-espaces vectoriels permettent de caractériser l'injectivité et la surjectivité d'une application linéaire :

$$u \text{ est injective si et seulement si } \text{Ker } u = \{0_E\}, \text{ et } u \text{ est surjective si et seulement si } \text{Im } u = F.$$

Remarque. Ces notions de noyau et d'image interviennent dans la résolution d'un système linéaire du type : $u(x) = y$, d'inconnue $x \in E$:

cette équation possède une solution si et seulement si $y \in \text{Im } u$, et dans ce cas, l'ensemble des solutions prend la forme $\{x_0 + h \mid h \in \text{Ker } u\}$, où x_0 est une solution particulière quelconque.

DÉFINITION. — Lorsque u est bijective, l'application u^{-1} est aussi linéaire. On dit alors que u est un isomorphisme, et que E et F sont des espaces vectoriels isomorphes.

Lorsqu'ils sont de dimensions finies, deux espaces isomorphes sont de même dimension.

Nous allons maintenant aborder un théorème très important, qui lie image et supplémentaire du noyau. Il s'agit du résultat suivant :

THÉORÈME 2.4 (Théorème du rang - forme géométrique) — Soit $u \in \mathcal{L}(E, F)$ une application linéaire, et H un supplémentaire de $\text{Ker } u$ dans E . Alors la restriction de u à H réalise un isomorphisme entre H et $\text{Im } u$.

En d'autres termes, l'application $u_H : \begin{pmatrix} H & \longrightarrow & \text{Im } u \\ x & \longmapsto & u(x) \end{pmatrix}$ est un isomorphisme.

Remarque. Lorsque E et F sont de dimensions finies, considérons une base $(e_1, \dots, e_r)$ de H et une base $(e_{r+1}, \dots, e_p)$ de $\text{Ker } u$. On obtient ainsi une base $(e_1, \dots, e_r, e_{r+1}, \dots, e_p)$ de E . Le théorème précédent nous permet d'affirmer que $(f_1 = u(e_1), \dots, f_r = u(e_r))$ est une base de $\text{Im } u$, que l'on peut compléter pour former une base $(f_1, \dots, f_r, f_{r+1}, \dots, f_n)$ de F . La matrice associée à u pour les bases (e) et (f) est alors la matrice suivante :

$$\begin{pmatrix} 1 & & 0 & \cdots & 0 \\ & \ddots & & & \\ & & 1 & & \\ 0 & \cdots & & 0 & \\ \vdots & & & & \\ 0 & \cdots & & & 0 \end{pmatrix} = \left(\begin{array}{c|c} I_r & O \\ \hline O & O \end{array} \right)$$

Notons que l'exemple donné en page 10 illustre ce résultat.

COROLLAIRE (Théorème du rang) — Soit E un $\mathbb{K}$ -espace vectoriel de dimension finie, F un $\mathbb{K}$ -espace vectoriel, et $u \in \mathcal{L}(E, F)$ une application linéaire. Alors $\text{Ker } u$ et $\text{Im } u$ sont de dimension finie, et :

$$\dim E = \dim(\text{Ker } u) + \dim(\text{Im } u).$$

COROLLAIRE — Si F est de dimension finie et si $\dim E = \dim F$, alors :

$$u \text{ injective} \iff u \text{ surjective} \iff u \text{ bijective}.$$

En particulier, pour les endomorphismes en dimension finie, injectivité, surjectivité et bijectivité sont des notions équivalentes.

Exercice 10

Soit E un $\mathbb{K}$ -espace vectoriel de dimension finie, et $(u, v) \in \mathcal{L}(E)^2$. Montrer, en appliquant le théorème du rang à la restriction de u à $\text{Im } v$, que : $\text{rg}(u \circ v) \geq \text{rg } u + \text{rg } v - \dim E$.

En déduire que $\dim(\text{Ker } u^2) \leq 2 \dim(\text{Ker } u)$.

■ Application à l'interpolation de Lagrange

En analyse numérique, l'interpolation est une opération mathématique consistant à déterminer une fonction à partir de la donnée d'un nombre fini de valeurs, et vérifiant éventuellement certaines propriétés supplémentaires.

Dans le cas particulier de l'interpolation de Lagrange on considère un entier $n \in \mathbb{N}$, $x_0, \dots, x_n$ des scalaires deux à deux distincts, et $y_0, \dots, y_n$ des scalaires quelconques. Le problème consiste à déterminer le ou les polynômes $P \in \mathbb{K}[X]$ (s'ils existent) vérifiant : $\forall k \in \llbracket 0, n \rrbracket, P(x_k) = y_k$, et si possible de degré minimal.

Considérons l'application linéaire $u : \mathbb{K}[X] \rightarrow \mathbb{K}^{n+1}$ définie par :

$$\forall P \in \mathbb{K}[X], \quad u(P) = (P(x_0), \dots, P(x_n)).$$

Si on note $y = (y_0, \dots, y_n)$, il s'agit de résoudre le système linéaire : $u(P) = y$, d'inconnue $P \in \mathbb{K}[X]$.

LEMME — Le noyau de u est constitué des multiples du polynôme $N = \prod_{i=0}^n (X - x_i)$.

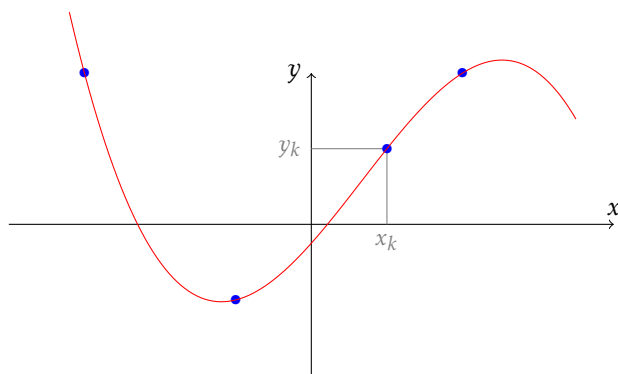


FIGURE 7 – Un polynôme de degré trois passant par quatre points d'interpolation.

Sachant que $\mathbb{K}_n[X]$ est un supplémentaire de $N \cdot \mathbb{K}[X]$ (principe de la division euclidienne par N), on en déduit que u réalise un isomorphisme entre $\mathbb{K}_n[X]$ et l'image de u . Mais alors $\dim(\text{Im } u) = n + 1$, et puisque $\text{Im } u \subset \mathbb{K}^{n+1}$ on a $\text{Im } u = \mathbb{K}^{n+1}$. Autrement dit, u est un endomorphisme surjectif, et :

THÉORÈME 2.5 — Il existe un unique polynôme P de $\mathbb{K}_n[X]$ tel que : $\forall k \in \llbracket 0, n \rrbracket, P(x_k) = y_k$.

Nous venons donc de démontrer que le problème de l'interpolation de Lagrange possède une unique solution P_L de degré inférieur ou égal à N ; les autres solutions s'écrivent : $P = P_L + N \cdot Q$, où Q est un polynôme quelconque. Mais tout ceci ne nous dit pas comment calculer P_L . Pour ce faire, nous allons introduire une nouvelle base de $\mathbb{K}_n[X]$, la base des *polynômes d'interpolation de Lagrange*, dans laquelle l'expression de P_L sera très simple.

THÉORÈME 2.6 — Posons pour tout entier $k \in \llbracket 0, n \rrbracket$, $L_k = \prod_{i \neq k} \frac{X - x_i}{x_k - x_i}$. Ces polynômes forment une base de $\mathbb{K}_n[X]$

pour laquelle : $\forall P \in \mathbb{K}_n[X], P = \sum_{k=0}^n P(x_k) L_k$.

Les polynômes L_k sont les *polynômes d'interpolation de Lagrange* aux points $x_0, \dots, x_n$.

Il devient alors évident que le polynôme P_L s'écrit : $P_L = \sum_{k=0}^n y_k L_k$.

Exemple. Déterminons le polynôme d'interpolation de degré minimal répondant aux conditions d'interpolation : $P(-3) = 2, P(-1) = -1, P(1) = 1, P(2) = 2$ (c'est celui représenté figure 7).

On commence par calculer les quatre polynômes de Lagrange associés aux réels $-3, -1, 1, 2$:

$$\begin{aligned} L_0 &= \frac{(X+1)(X-1)(X-2)}{(-3+1)(-3-1)(-3-2)} = -\frac{1}{40}(X^3 - 2X^2 - X + 2) \\ L_1 &= \frac{(X+3)(X-1)(X-2)}{(-1+3)(-1-1)(-1-2)} = \frac{1}{12}(X^3 - 7X + 6) \\ L_2 &= \frac{(X+3)(X+1)(X-2)}{(1+3)(1+1)(1-2)} = -\frac{1}{8}(X^3 + 2X^2 - 5X - 6) \\ L_3 &= \frac{(X+3)(X+1)(X-1)}{(2+3)(2+1)(2-1)} = \frac{1}{15}(X^3 + 3X^2 - X - 3) \end{aligned}$$

Le polynôme d'interpolation recherché est donc :

$$P = 2L_0 - L_1 + L_2 + 2L_3 = -\frac{1}{8}X^3 + \frac{1}{4}X^2 + \frac{9}{8}X - \frac{1}{4}$$

Adoptons maintenant une démarche naïve pour résoudre le problème de l'interpolation de Lagrange : posons

[illegible]
$$\begin{pmatrix} 1 & x_1 & x_1^2 & \cdots & x_1^{n-1} \\ 1 & x_2 & x_2^2 & \cdots & x_2^{n-1} \\ \vdots & \vdots & \vdots & & \vdots \\ 1 & x_n & x_n^2 & \cdots & x_n^{n-1} \end{pmatrix} \begin{pmatrix} a_0 \\ a_1 \\ \vdots \\ a_{n-1} \end{pmatrix} = \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix}$$
$$V(x_1, x_2, \dots, x_n) = \begin{vmatrix} 1 & x_1 & x_1^2 & \cdots & x_1^{n-1} \\ 1 & x_2 & x_2^2 & \cdots & x_2^{n-1} \\ \vdots & \vdots & \vdots & & \vdots \\ 1 & x_n & x_n^2 & \cdots & x_n^{n-1} \end{vmatrix}$$

THÉORÈME 2.7 — $V(x_1, x_2, \dots, x_n) = \prod_{i=2}^n \prod_{j=1}^{j-1} (x_j - x_i)$, formule qu'on retiendra sous la forme plus concise :

$$V(x_1, x_2, \dots, x_n) = \prod_{i < j} (x_j - x_i).$$

Si $P = \sum_{k=0}^n a_k X^k$, on définit l'endomorphisme $P(u) = \sum_{k=0}^n a_k u^k$. En bref :

$$\begin{aligned} P(X) &= a_n X^n + a_{n-1} X^{n-1} + \cdots + a_1 X + a_0 \\ P(u) &= a_n u^n + a_{n-1} u^{n-1} + \cdots + a_1 u + a_0 \text{Id} \end{aligned}$$

$$\forall (P, Q) \in \mathbb{K}[X]^2, \quad (PQ)(u) = P(u) \circ Q(u).$$

PC* – Lycée Marcelin Berthelot

Attention. Si P et Q sont deux polynômes vérifiant $PQ = 0$, on sait que l'on peut en déduire que $P = 0$ ou $Q = 0$. **Ce n'est pas le cas des polynômes d'un endomorphisme** : on peut avoir $(PQ)(u) = 0$ sans pour autant en déduire que $P(u) = 0$ ou $Q(u) = 0$.

Considérons par exemple une projection vectorielle u : on a $u^2 - u = 0$. Si on pose $P = X$ et $Q = X - 1$ on a $PQ = X^2 - X$ donc $(PQ)(u) = 0$, mais on a pas en général $P(u) = 0$ ou $Q(u) = 0$ (sauf si $u = 0$ ou $u = \text{Id}$).

DÉFINITION. — Si $u \in \mathcal{L}(E)$ et $P \in \mathbb{K}[X]$, on dit que P est un polynôme annulateur de u lorsque $P(u) = 0$.

Par exemple, $X^2 - X$ est un polynôme annulateur de toute projection vectorielle, $X^2 - 1$ un polynôme annulateur de toute symétrie vectorielle.

PROPOSITION 2.9 — Lorsque E est un espace vectoriel de dimension finie, tout endomorphisme $u \in \mathcal{L}(E)$ possède un polynôme annulateur.

Exemple. De façon symétrique, on définit la notion de polynôme annulateur d'une matrice carrée $A \in \mathcal{M}_n(\mathbb{R})$.

Considérons une matrice $A \in \mathcal{M}_2(\mathbb{K})$, et posons $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$.

$$\text{On a } A^2 = \begin{pmatrix} a^2 + bc & (a+d)b \\ (a+d)c & bc + d^2 \end{pmatrix} = \begin{pmatrix} (a+d)a + bc - ad & (a+d)b \\ (a+d)c & (a+d)d + bc - ad \end{pmatrix} = (a+d) \begin{pmatrix} a & b \\ c & d \end{pmatrix} - (ad - bc) \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \\ = (\text{tr } A)A - (\det A)I_2$$

Autrement dit, le polynôme $P = X^2 - (\text{tr } A)X + (\det A)$ est un polynôme annulateur de A .

■ Application au calcul de l'inverse

Considérons un endomorphisme $u \in \mathcal{L}(E)$, et M un polynôme annulateur de u , de degré d . Supposons de plus le

coefficient constant de M non nul : $M = \sum_{k=0}^d a_k X^k$ avec $a_0 \neq 0$.

$$\text{Alors } M(u) = 0 \iff a_0 \text{Id} = - \sum_{k=1}^d a_k u^k = u \circ \left(- \sum_{k=1}^d a_k u^{k-1} \right) \text{ donc } u \text{ est inversible, d'inverse } u^{-1} = -\frac{1}{a_0} \left(- \sum_{k=1}^d a_k u^{k-1} \right).$$

Exemple. Si $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$, on a vu que $A^2 - (\text{tr } A)A + (\det A)I_2 = 0$ donc si A est inversible, $A((\text{tr } A)I_2 - A) = (\det A)I_2$.

$$\text{On a donc } A^{-1} = \frac{1}{\det A} ((\text{tr } A)I_2 - A) = \frac{1}{ad - bc} \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}.$$

■ Application au calcul des puissances de u

Pour calculer u^n , on peut réaliser la division euclidienne de X^n par M : $X^n = MQ + R$, avec $\deg R < d$. Ainsi, $u^n = M(u) \circ Q(u) + R(u) = R(u)$ puisque $M(u) = 0$. Le calcul de u^n se ramène à celui de $R(u)$, ce qui peut être intéressant lorsque le degré d du polynôme annulateur est petit, puisque $\deg R < d$.

Exercice 11

a. Soit $p \in \mathbb{N}^*$. Déterminer le reste de la division euclidienne de $(1 + X)^n$ par $X(X - p)$.

b. On pose $U = \begin{pmatrix} 1 & \cdots & \cdots & 1 \\ \vdots & \ddots & & \vdots \\ \vdots & & \ddots & \vdots \\ 1 & \cdots & \cdots & 1 \end{pmatrix} \in \mathcal{M}_p(\mathbb{R})$, et $A = \begin{pmatrix} 2 & 1 & \cdots & 1 \\ 1 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 1 \\ 1 & \cdots & 1 & 2 \end{pmatrix} \in \mathcal{M}_p(\mathbb{R})$. Déterminer un polynôme annulateur de U , et en déduire A^n pour $n \in \mathbb{N}$, ainsi que l'inverse de A , s'il existe.

■ Polynôme minimal (notion hors programme)

Dans la seconde application, nous avons vu que nous avons intérêt à utiliser un polynôme annulateur de degré d le plus petit possible. Un tel polynôme existe toujours ; son degré est défini par :

$$d = \min\{k \in \mathbb{N} \mid (\text{Id}, u, u^2, \dots, u^k) \text{ est liée}\}$$

et il est caractérisé par :

$$\text{la famille } (\text{Id}, u, u^2, \dots, u^{d-1}) \text{ est libre et } u^d \in \text{Vect}(\text{Id}, u, u^2, \dots, u^{d-1}).$$

De plus, il est unique si on fixe son coefficient dominant :

THÉORÈME 2.10 — Il existe un unique polynôme annulateur et unitaire de degré minimal ; il est appelé le polynôme minimal de u .

THÉORÈME 2.11 — Si M est le polynôme minimal de u , les polynômes annulateurs de u sont les multiples de M .

2.4 Sous-espaces stables

■ Matrices définies par blocs

Considérons une matrice $A \in \mathcal{M}_{n,p}(\mathbb{K})$ ainsi que deux entiers $i \in \llbracket 1, n-1 \rrbracket$ et $j \in \llbracket 1, p-1 \rrbracket$. Divisons les lignes de A en deux ensembles : les lignes dont les indices sont compris entre 1 et i et celles dont les indices sont compris entre $i+1$ et n . Faisons de même avec les colonnes en distinguant celles dont les indices sont compris entre 1 et j de celles dont les indices sont compris entre $j+1$ et p .

En procédant de la sorte, on divise la matrice A en quatre blocs :

$$A = \begin{pmatrix} A_1 & A_2 \\ A_3 & A_4 \end{pmatrix} \quad \begin{matrix} \uparrow i \\ \downarrow n-i \end{matrix} \quad \text{avec} \quad A_1 \in \mathcal{M}_{i,j}(\mathbb{K}), A_2 \in \mathcal{M}_{i,p-j}(\mathbb{K}), A_3 \in \mathcal{M}_{n-i,j}(\mathbb{K}), A_4 \in \mathcal{M}_{n-i,p-j}(\mathbb{K}).$$

$\leftarrow j \quad \leftarrow p-j$

Une telle matrice sera dite *définie par blocs*.

Pour peu que le découpage soit identique, la définition par bloc de deux matrices est évidemment compatible avec l'addition :

$$\text{si } A' = \begin{pmatrix} A'_1 & A'_2 \\ A'_3 & A'_4 \end{pmatrix} \quad \text{alors} \quad \lambda A + A' = \begin{pmatrix} \lambda A_1 + A'_1 & \lambda A_2 + A'_2 \\ \lambda A_3 + A'_3 & \lambda A_4 + A'_4 \end{pmatrix}$$

mais le fait le plus remarquable est que le découpage par blocs est *compatible avec la multiplication*, pour peu que les découpages conduisent à des produits « licites » de matrices :

$$\text{si } B = \begin{pmatrix} B_1 & B_2 \\ B_3 & B_4 \end{pmatrix} \in \mathcal{M}_{p,q}(\mathbb{K}) \quad \begin{matrix} \uparrow j \\ \downarrow p-j \end{matrix} \quad \text{alors} \quad AB = \begin{pmatrix} A_1 B_1 + A_2 B_3 & A_1 B_2 + A_2 B_4 \\ A_3 B_1 + A_4 B_3 & A_3 B_2 + A_4 B_4 \end{pmatrix} \begin{matrix} \uparrow i \\ \downarrow n-i \end{matrix} \in \mathcal{M}_{n,q}(\mathbb{K})$$

$\leftarrow k \quad \leftarrow q-k \quad \leftarrow k \quad \leftarrow q-k$

Autrement dit, les matrices définies par blocs se multiplient entre elles tout comme si les blocs étaient des scalaires, à condition que chaque multiplication corresponde à une multiplication « légale » de matrices (en ce qui concerne les dimensions).

Ces propriétés s'étendent par récurrence au cas d'un découpage des lignes et/ou des colonnes en un nombre arbitraire de subdivisions.

DÉFINITION. — Une matrice carrée $A \in \mathcal{M}_p(\mathbb{K})$ est dite diagonale par bloc lorsqu'il existe une subdivision de $\llbracket 1, p \rrbracket$ telle que :

$$A = \begin{pmatrix} \boxed{A_{11}} & & & \\ & \boxed{A_{22}} & & \\ & & \ddots & \\ & & & \boxed{A_{kk}} \end{pmatrix}$$

$\xleftarrow{i_1} \xleftarrow{i_2} \cdots \xleftarrow{i_k}$
 $\begin{matrix} \uparrow i_1 \\ \uparrow i_2 \\ \vdots \\ \uparrow i_k \end{matrix}$

(Tous les blocs sont nuls hormis les blocs diagonaux, qui sont tous carrés.)

Une matrice carrée $A \in \mathcal{M}_p(\mathbb{K})$ est dite triangulaire par bloc lorsqu'il existe une subdivision de $\llbracket 1, p \rrbracket$ telle que :

$$A = \begin{pmatrix} \boxed{A_{11}} & \boxed{A_{12}} & \cdots & \boxed{A_{1k}} \\ & \boxed{A_{22}} & \cdots & \boxed{A_{2k}} \\ & & \ddots & \vdots \\ & & & \boxed{A_{kk}} \end{pmatrix}$$

$\xleftarrow{i_1} \xleftarrow{i_2} \cdots \xleftarrow{i_k}$
 $\begin{matrix} \uparrow i_1 \\ \uparrow i_2 \\ \vdots \\ \uparrow i_k \end{matrix}$

(Tous les blocs diagonaux sont carrés, et les blocs situés sous la diagonale sont nuls.)

■ Sous-espaces stables

DÉFINITION. — Soit H un sous-espace vectoriel de E , et $u \in \mathcal{L}(E)$ un endomorphisme. On dit que H est stable par u lorsque $u(H) \subset H$.

Considérons une base adaptée à un sous-espace vectoriel H , c'est-à-dire construite à partir d'une base $(e_1, \dots, e_k)$ de H puis complétée pour former une base $(e_1, \dots, e_k, e_{k+1}, \dots, e_p)$ de E . Alors H est stable par u si et seulement si la matrice associée à u dans cette base (e) est de la forme :

$$\begin{matrix} \begin{matrix} \uparrow k \\ \uparrow p-k \end{matrix} \end{matrix} \left(\begin{array}{cc} \boxed{A} & \boxed{C} \\ \boxed{O} & \boxed{D} \end{array} \right) \begin{matrix} \xleftarrow{k} \xleftarrow{p-k} \end{matrix}$$

En effet, nous avons : $\forall j \in \llbracket 1, k \rrbracket, u(e_j) \in H = \text{Vect}(e_1, \dots, e_k)$.

Lorsque H est stable par u , la restriction de u à H définit donc un endomorphisme u_H de H dont la matrice dans la base $(e_1, \dots, e_k)$ est la matrice A . Cet endomorphisme s'appelle l'*induit* de u sur H .

Remarque. Dans une base $(e'_1, \dots, e'_p)$ de E pour laquelle ce sont les vecteurs $(e'_{p-k+1}, \dots, e'_p)$ qui forment une

base de H , la matrice d'un endomorphisme stabilisant H est de la forme :

$$\begin{pmatrix} \boxed{D} & \boxed{O} \\ \boxed{C} & \boxed{A} \end{pmatrix}$$

Exemple. $\text{Ker } u$ et $\text{Im } u$ sont des sous-espaces vectoriels stables de u . En effet, dans une base adaptée à $\text{Ker } u$, la matrice associée à u prend la forme :

$$\begin{pmatrix} \boxed{O} & \boxed{C} \\ \boxed{O} & \boxed{D} \end{pmatrix}$$

et dans une base adaptée à $\text{Im } u$ la matrice associée à u prend la forme :

$$\begin{pmatrix} \boxed{A} & \boxed{C} \\ \boxed{O} & \boxed{O} \end{pmatrix}$$

PROPOSITION 2.12 — Si $P \in \mathbb{K}[X]$ alors $\text{Ker } P(u)$ est un sous-espace stable de u .

Exercice 12

Soit E un $\mathbb{K}$ -espace vectoriel, et $p \in \mathcal{L}(E)$ une projection vectorielle. Montrer que $u \in \mathcal{L}(E)$ commute avec p si et seulement si $\text{Ker } p$ et $\text{Im } p$ sont stables par u .

Décomposition de l'espace en somme de sous-espaces stables

Considérons enfin une famille $(H_1, \dots, H_k)$ de sous-espaces vectoriels telle que : $E = H_1 \oplus H_2 \oplus \dots \oplus H_k$, et une base $(e_1, \dots, e_p)$ adaptée à cette décomposition. Alors un endomorphisme $u \in \mathcal{L}(E)$ stabilise *chacun* de ces sous-espaces vectoriels si et seulement si la matrice associée à u dans cette base est diagonale par bloc :

$$\text{Mat}_{(e)}(u) = \begin{pmatrix} \boxed{A_1} & & & \\ & \boxed{A_2} & & \\ & & \ddots & \\ & & & \boxed{A_k} \end{pmatrix} = A$$

Remarque. Avec les notations ci-dessus, on a : $\text{rg } A = \sum_{j=1}^k \text{rg } A_j$ et $\text{tr } A = \sum_{j=1}^k \text{tr } A_j$.

En outre, si v est un endomorphisme ayant aussi $H_1, H_2, \dots, H_k$ comme sous-espaces stables, et si $B = \text{Mat}_{(e)}(v)$,

alors :

$$B = \begin{pmatrix} B_1 & & \\ & B_2 & \\ & & \ddots \\ & & & B_k \end{pmatrix} \quad \text{et} \quad AB = \begin{pmatrix} A_1 B_1 & & \\ & A_2 B_2 & \\ & & \ddots \\ & & & A_k B_k \end{pmatrix}$$

En particulier, on notera que pour tout entier $n \in \mathbb{N}$,

$$A^n = \begin{pmatrix} A_1^n & & \\ & A_2^n & \\ & & \ddots \\ & & & A_k^n \end{pmatrix}$$

■ Déterminant d'une matrice définie par blocs

Il n'existe pas de formule simple pour calculer le déterminant d'une matrice définie par blocs, à l'exception du cas des matrices triangulaires par blocs. Commençons par le cas d'une matrice définie par quatre blocs :

PROPOSITION 2.13 — Soit $A \in \mathcal{M}_n(\mathbb{K})$, et $k \in \llbracket 1, n-1 \rrbracket$ un entier induisant la même partition des lignes et des colonnes en deux sous-ensembles $\llbracket 1, k \rrbracket$ et $\llbracket k+1, n \rrbracket$. On suppose de plus le bloc correspondant aux indices de lignes $\llbracket k+1, n \rrbracket$ et aux indices de colonnes $\llbracket 1, k \rrbracket$ (autrement dit le bloc en bas à gauche) nul. Alors :

$$A = \begin{pmatrix} A_1 & A_2 \\ O & A_4 \end{pmatrix} \implies \det A = \det(A_1) \times \det(A_4).$$

On en déduit aisément par récurrence le :

COROLLAIRE — Le déterminant d'une matrice triangulaire par bloc est égal au produit des déterminants des blocs diagonaux :

$$\begin{vmatrix} A_{11} & A_{12} & \cdots & A_{1k} \\ & A_{22} & \cdots & A_{2k} \\ & & \ddots & \vdots \\ & & & A_{kk} \end{vmatrix} = \det A_{11} \times \det A_{22} \times \cdots \times \det A_{kk}.$$

Exercice 13

Soient A, B, C, D quatre matrices de $\mathcal{M}_n(\mathbb{K})$. On suppose que C et D commutent et que D est inversible.

Calculer le produit $\begin{pmatrix} A & B \\ C & D \end{pmatrix} \begin{pmatrix} D & O \\ -C & D^{-1} \end{pmatrix}$ et en déduire : $\det \begin{pmatrix} A & B \\ C & D \end{pmatrix} = \det(AD - BC)$.

2.5 Endomorphismes nilpotents (notion hors-programme)

DÉFINITION. — Un endomorphisme $u \in \mathcal{L}(E)$ est dit nilpotent lorsqu'il existe un entier $p \in \mathbb{N}^*$ tel que $u^p = 0$. Le plus petit entier p vérifiant cette condition, autrement dit tel que $u^p = 0$ et $u^{p-1} \neq 0$, est appelé l'indice de nilpotence de u .

THÉORÈME 2.14 — Soit u un endomorphisme nilpotent d'indice p , et $x \in E$ un vecteur vérifiant $u^{p-1}(x) \neq 0_E$. Alors la famille $(x, u(x), \dots, u^{p-1}(x))$ est libre.

COROLLAIRE — Lorsque l'espace vectoriel est de dimension n , l'indice d'un endomorphisme nilpotent est inférieur ou égal à n .

Intéressons nous maintenant au cas où l'indice de nilpotence de u est égal à la dimension n de E . Dans ce cas, quel que soit $x \in E$ vérifiant $u^{n-1}(x) \neq 0_E$, la famille $(x, u(x), \dots, u^{n-1}(x))$ est libre et de cardinal n donc constitue une base de E , base dans laquelle la matrice associée à u est de la forme :

$$J = \begin{pmatrix} 0 & 0 & \cdots & \cdots & 0 \\ 1 & \ddots & \ddots & & \vdots \\ 0 & \ddots & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & \ddots & 0 \\ 0 & \cdots & 0 & 1 & 0 \end{pmatrix}$$

Exercice 14

Montrer que J et J^T sont deux matrices semblables.

Chapitre II

Réduction des endomorphismes

Réduire un endomorphisme, c'est trouver une base dans laquelle la matrice associée à cet endomorphisme est la plus simple possible, de manière à faciliter les calculs que l'on peut être amené à effectuer sur celui-ci. La réduction sans doute la plus utilisée en dimension finie est la réduction de Jordan, qui décompose l'espace en somme directe de sous-espaces stables, l'endomorphisme agissant de manière très simple sur chacun de ces sous-espaces.

1. Introduction

Nous allons commencer par observer l'action de la réduction de Jordan sur un exemple, pour apprécier l'intérêt qu'il y a à réduire un endomorphisme.

Exemple. Considérons l'endomorphisme u de $E = \mathbb{R}^4$ défini par sa matrice sur la base canonique (e) :

$$\text{Mat}_{(e)}(u) = \begin{pmatrix} 5 & 8 & 6 & 5 \\ 0 & 2 & 0 & 0 \\ -1 & -4 & 0 & -3 \\ 1 & 4 & 2 & 5 \end{pmatrix} = A$$

Nous allons effectuer le changement de base sur la base (e') définie par la matrice de passage :

$$\text{Mat}_{(e)}(e') = \begin{pmatrix} -1 & 0 & 1 & -1 \\ 0 & 0 & -1 & -1 \\ 1 & -1 & 0 & 1 \\ 1 & 1 & 1 & 1 \end{pmatrix} = P$$

Bien entendu, nous ne savons pas pour l'instant comment ont été choisis ces vecteurs formant la nouvelle base ; c'est là tout l'enjeu de ce chapitre. Mais observons déjà le résultat de ce changement de base.

Nous l'avons déjà dit au chapitre précédent, calculer $P^{-1}AP$ est la plus-part du temps une mauvaise option ; il est préférable de calculer les vecteurs $u(e'_k)$, $k \in \{1, 2, 3, 4\}$, et chercher à les exprimer dans la base (e') . On calcule donc :

$$\begin{aligned} u(e'_1) &= -u(e_1) + u(e_3) + u(e_4) = -4e_1 + 4e_3 + 4e_4 = 4e'_1 \\ u(e'_2) &= -u(e_3) + u(e_4) = -e_1 - 3e_3 + 3e_4 = e'_1 + 4e'_2 \\ u(e'_3) &= u(e_1) - u(e_2) + u(e_4) = 2e_1 - 2e_2 + 2e_4 = 2e'_3 \\ u(e'_4) &= -u(e_1) - u(e_2) + u(e_3) + u(e_4) = -2e_1 - 2e_2 + 2e_3 + 2e_4 = 2e'_4 \end{aligned}$$

Nous obtenons $\text{Mat}_{(e')}(u) = \begin{pmatrix} \boxed{4} & \boxed{1} & 0 & 0 \\ 0 & 4 & 0 & 0 \\ 0 & 0 & \boxed{2} & 0 \\ 0 & 0 & 0 & \boxed{2} \end{pmatrix} = P^{-1}AP.$

Cette nouvelle matrice est constituée de deux blocs diagonaux, qui correspondent à la décomposition de l'espace en deux sous-espaces stables : $E = H_1 \oplus H_2$ avec $H_1 = \text{Vect}(e'_1, e'_2)$, $H_2 = \text{Vect}(e'_3, e'_4)$.

Sur le plan vectoriel H_2 l'endomorphisme u agit comme une homothétie :

$$\forall x \in H_2, u(x) = 2x.$$

Sur le plan vectoriel H_1 , l'action de u est un peu plus compliquée : c'est l'addition d'une homothétie de rapport 4 et d'un endomorphisme nilpotent v défini par $v(e'_1) = 0_E$ et $v(e'_2) = e'_1$:

$$\forall x \in H_1, u(x) = 4x + v(x) \quad \text{avec } v^2(x) = 0_E.$$

Il est beaucoup plus facile de travailler avec la base (e') qu'avec la base (e) ; par exemple, le calcul de u^n s'obtient très simplement dans la base (e') :

$$\forall x \in H_1, u^n(x) = 4^n x + n4^{n-1}v(x), \quad \forall x \in H_2, u^n(x) = 2^n x$$

égalités qui se traduisent matriciellement par : $A^n = P \begin{pmatrix} \boxed{4^n} & \boxed{n4^{n-1}} & 0 & 0 \\ 0 & 4^n & 0 & 0 \\ 0 & 0 & \boxed{2^n} & 0 \\ 0 & 0 & 0 & \boxed{2^n} \end{pmatrix} P^{-1}$.

Exercice 1

On considère l'endomorphisme $u \in \mathcal{L}(\mathbb{K}^3)$ défini par la matrice $A = \begin{pmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \\ -1 & 1 & 1 \end{pmatrix}$.

Déterminer les droites vectorielles stables par u , et en déduire une base (e) de $\mathbb{K}^3$ pour laquelle $\text{Mat}_{(e)}(u)$ est diagonale.

2. Éléments propres

Dans l'exemple introductif que nous venons de traiter, sur deux des sous-espaces de la décomposition (H_2 et H_3) l'endomorphisme agit comme une homothétie. Ce sont ces sous-espaces particuliers qui vont nous intéresser.

2.1 Valeurs et vecteurs propre

Dans cette section, sauf mention explicite du contraire, E désigne un $\mathbb{K}$ -espace vectoriel, de dimension finie ou non.

DÉFINITION. — On dit qu'un scalaire $\lambda \in \mathbb{K}$ est une valeur propre d'un endomorphisme $u \in \mathcal{L}(E)$ lorsqu'il existe un vecteur non nul $x \in E$ tel que $u(x) = \lambda x$. Dans ce cas, on dit que x est un vecteur propre associé à la valeur propre λ .

On note $\text{Sp}(u)$ l'ensemble des valeurs propres de u ; c'est le spectre de u .

DÉFINITION. — Si λ est une valeur propre de u , on note $E_\lambda(u) = \text{Ker}(u - \lambda \text{Id}_E)$; il s'agit du sous-espace propre associé à la valeur propre λ . C'est un sous-espace vectoriel de E stable par u .

Attention. Le vecteur nul n'est pas un vecteur propre; les vecteurs propres associés à une valeur propre λ sont les éléments *non nuls* du sous-espace propre $E_\lambda(u)$, sous-espace qui est au moins de dimension 1.

Remarque. La restriction de u au sous-espace propre $E_\lambda(u)$ est l'homothétie vectorielle de rapport λ .

Exercice 2

Soit $E = \mathcal{C}^\infty(\mathbb{R}, \mathbb{R})$ le $\mathbb{R}$ -espace vectoriel des applications de classe $\mathcal{C}^\infty$ sur $\mathbb{R}$, et $D : f \mapsto f'$ l'opérateur de dérivation. Déterminer les éléments propres (valeurs et vecteurs propres) de D .

THÉORÈME 2.1 — Si $\lambda_1, \dots, \lambda_k$ sont des valeurs propres deux à deux distinctes de u , la somme $E_{\lambda_1}(u) \oplus \dots \oplus E_{\lambda_k}(u)$ est directe.

Chacun de ces sous-espaces propres étant au minimum de dimension 1, on en déduit :

COROLLAIRE — Si E est un espace vectoriel de dimension finie p , tout endomorphisme a au plus p valeurs propres distinctes.

■ Traduction matricielle en dimension finie

Considérons maintenant un $\mathbb{K}$ -espace vectoriel E de dimension finie, et (e) une base de E . L'égalité $u(x) = \lambda x$ se traduit matriciellement par $AX = \lambda X$, où $A = \text{Mat}_{(e)}(u) \in \mathcal{M}_p(\mathbb{K})$ et $X = \text{Mat}_{(e)}(x) \in \mathcal{M}_{p,1}(\mathbb{K})$, ce qui nous amène aux définitions suivantes (rappelons que l'espace des matrices colonnes $\mathcal{M}_{p,1}(\mathbb{K})$ est identifié à $\mathbb{K}^p$) :

DÉFINITION. — Soit $A \in \mathcal{M}_p(\mathbb{K})$ une matrice carrée. Un scalaire $\lambda \in \mathbb{K}$ est une valeur propre de A lorsqu'il existe un vecteur non nul $x \in \mathbb{K}^p$ tel que $Ax = \lambda x$. Le vecteur x est un vecteur propre associé à la valeur propre λ . En outre, on appelle sous-espace propre associé à la valeur propre λ le sous-espace vectoriel :

$$\text{Ker}(A - \lambda I) = \{x \in \mathbb{K}^p \mid Ax = \lambda x\}.$$

D'après le corollaire du théorème 2.1, une matrice $p \times p$ ne peut avoir plus de p valeurs propres distinctes.

Il y a bien entendu parfaite équivalence entre éléments spectraux d'un endomorphisme et éléments spectraux d'une matrice qui lui est associée par le choix d'une base.

PROPOSITION 2.2 — Un scalaire λ est valeur propre de $A \in \mathcal{M}_p(\mathbb{K})$ si et seulement si $\det(\lambda I - A) = 0$.

Ce dernier résultat nous indique la démarche à suivre pour étudier les éléments propres en dimension finie :

1. déterminer les valeurs propres de A en résolvant l'équation $\det(\lambda I - A) = 0$;
2. pour chaque valeur propre λ , résoudre le système linéaire $(A - \lambda I)X = 0$ pour déterminer une base du sous-espace propre correspondant ;
3. Lorsque cela est possible, construire une base formée de vecteurs propres, et établir la formule de changement de base $A = PDP^{-1}$.

Exercice 3

Déterminer les éléments propres des matrices suivantes et le cas échéant, former une base de vecteurs propres :

$$A_1 = \begin{pmatrix} 5 & 2 & 6 \\ -4 & -1 & -8 \\ 0 & 0 & 2 \end{pmatrix} \quad A_2 = \begin{pmatrix} 5 & -3 & -2 \\ -3 & 5 & 2 \\ 6 & -6 & -2 \end{pmatrix} \quad A_3 = \begin{pmatrix} 0 & 2 & 1 \\ -4 & 6 & 1 \\ 4 & -4 & 2 \end{pmatrix}$$

2.2 Polynôme caractéristique

En dimension finie, nous venons de constater que déterminer les valeurs propres d'un endomorphisme $u \in \mathcal{L}(E)$ revient à résoudre l'équation $\det(u - \lambda \text{Id}_E) = 0$. Nous allons nous intéresser à la nature de cette équation, en démontrant qu'il s'agit d'une équation *polynomiale*.

Considérons une base quelconque (e) de E , et $A = \text{Mat}_e(u) = (a_{ij})$. Alors :

$$\det(x\text{Id}_E - u) = \det(xI - A) = \begin{vmatrix} x - a_{11} & -a_{12} & \dots & -a_{1p} \\ -a_{21} & x - a_{22} & \dots & -a_{2p} \\ \vdots & \vdots & \ddots & \vdots \\ -a_{p1} & -a_{p,p-1} & \dots & x - a_{pp} \end{vmatrix}$$

THÉORÈME 2.3 — L'application $x \mapsto \det(xI - A)$ est une fonction polynomiale ; le polynôme qui lui est associé est un polynôme unitaire de degré p appelé polynôme caractéristique de la matrice A . Il est noté χ_A .

Remarque. Le déterminant d'un endomorphisme ne dépendant pas de la base choisie pour effectuer le calcul, on définit de même le *polynôme caractéristique* d'un endomorphisme : le polynôme canoniquement associé à la fonction polynomiale $x \mapsto \det(x\text{Id}_E - u)$.

Exemple. Lorsque $A = \begin{pmatrix} a & c \\ b & d \end{pmatrix}$, son polynôme caractéristique est :

$$\chi_A(x) = \begin{vmatrix} x-a & -c \\ -b & x-d \end{vmatrix} = (x-a)(x-d) - bc = x^2 - (a+d)x + ad - bc = x^2 - (\text{tr } A)x + \det A.$$

Remarque. Pour une matrice $A \in \mathcal{M}_p(\mathbb{K})$ de taille $p \times p$, le coefficient constant de χ_A est égal à $(-1)^p \det A$ et le coefficient de X^{p-1} égal à $-\text{tr } A$.

Exercice 4

Soit E un espace vectoriel de dimension finie, et $u \in \mathcal{L}(E)$ un endomorphisme de rang 1. Déterminer son polynôme caractéristique.

■ Ordre de multiplicité d'une valeur propre

DÉFINITION. — les valeurs propres d'un endomorphisme u sont les racines de son polynôme caractéristique. On appelle ordre de multiplicité d'une valeur propre son ordre de multiplicité en tant que racine du polynôme caractéristique.

PROPOSITION 2.4 — Lorsque le polynôme caractéristique est scindé, notons $\lambda_1, \dots, \lambda_p$ les valeurs propres de u , en répétant autant de fois que sa multiplicité chacune des valeurs propres. Alors $\det u = \prod_{k=1}^p \lambda_k$ et $\text{tr } u = \sum_{k=1}^p \lambda_k$.

Enfin, on notera qu'il existe un lien entre ordre de multiplicité de la valeur propre et la dimension du sous-espace propre correspondant :

THÉORÈME 2.5 — La dimension d'un sous-espace propre est inférieure ou égale à l'ordre de multiplicité de la valeur propre correspondante.

Ce résultat a plusieurs conséquences intéressantes. Considérons par exemple une valeur propre simple (c'est à dire de multiplicité égale à 1). Le sous-espace propre associé n'étant pas réduit à $\{0_E\}$, on en déduit qu'il est obligatoirement de dimension 1.

Nous verrons d'autres conséquences de ce résultat dans les sections suivantes.

2.3 Diagonalisation en dimension finie

Dans toute cette section on suppose que E est un $\mathbb{K}$ -espace vectoriel de dimension finie p .

DÉFINITION. — Un endomorphisme $u \in \mathcal{L}(E)$ est dit diagonalisable lorsqu'il existe une base (e) dans laquelle la matrice $\text{Mat}_{(e)}(u)$ est diagonale.

Traduction matricielle

Considérons une base *quelconque* (e) , et $A = \text{Mat}_{(e)}(u)$. u est diagonalisable s'il existe une base (e') telle que $D = \text{Mat}_{(e')}(u)$ est diagonale. Si on note $P = \text{Mat}_{(e)}(e')$ la matrice de passage de (e) vers (e') nous disposons de la relation : $D = P^{-1}AP$, qu'on peut écrire $A = PDP^{-1}$. Ceci conduit à la définition :

DÉFINITION. — Une matrice carrée $A \in \mathcal{M}_p(\mathbb{K})$ est dite diagonalisable lorsqu'il existe une matrice inversible $P \in \mathcal{GL}_p(\mathbb{K})$ telle que $A = PDP^{-1}$.

Exemple. Les matrices A_1 et A_2 de l'exercice 3 sont diagonalisables : nous avons dans les deux cas trouvé une base formée de vecteurs propres.

Exercice 5

Soit $A \in \mathcal{M}_p(\mathbb{K})$ une matrice triangulaire supérieure dans laquelle tous les coefficients diagonaux sont égaux. Peut-elle être diagonalisable ?

Remarque. Lorsqu'un endomorphisme u est diagonalisable, la base (e) pour laquelle $\text{Mat}_{(e)}(u)$ est diagonale est constituée de vecteurs propres. Dès lors, on ne s'étonnera pas des nombreuses définitions équivalentes que l'on va obtenir en faisant intervenir la théorie spectrale.

THÉORÈME 2.6 — Soit $u \in \mathcal{L}(E)$ un endomorphisme de E , et $\text{Sp}(u) = \{\lambda_1, \dots, \lambda_k\}$ le spectre de u . Alors u est diagonalisable si et seulement si $E = E_{\lambda_1}(u) \oplus \dots \oplus E_{\lambda_k}(u)$.

COROLLAIRE — Soit $u \in \mathcal{L}(E)$ un endomorphisme de E , $\text{Sp}(u) = \{\lambda_1, \dots, \lambda_k\}$ le spectre de u . Alors u est diagonalisable si et seulement si $\sum_{i=1}^k \dim E_{\lambda_i}(u) = \dim E$.

Exemple. La matrice A_3 de l'exercice 3 n'est pas diagonalisable. Nous n'avons trouvé que deux sous-espaces propres, chacun de dimension 1.

COROLLAIRE — Un endomorphisme u de $\mathcal{L}(E)$ est diagonalisable si et seulement si son polynôme caractéristique est scindé sur le corps de base $\mathbb{K}$, et si pour toute valeur propre la dimension du sous-espace propre associé est égale à sa multiplicité dans le polynôme caractéristique.

Exemple. Reprenons une nouvelle fois les exemples de l'exercice 3 :

- le polynôme caractéristique de A_1 est égal à $(X-1)(X-2)(X-3)$; A_1 possède trois sous-espaces propres de dimension 1 donc A_1 est diagonalisable;
- le polynôme caractéristique de A_2 est égal à $(X-2)^2(X-4)$; le sous-espace propre associé à la valeur propre 2 est de dimension 2, celui associé à la valeur propre 4 de dimension 1, donc A_2 est diagonalisable;
- le polynôme caractéristique de A_3 est égal à $(X-2)^2(X-4)$; le sous-espace propre associé à la valeur propre 2 est de dimension 1 donc A_3 n'est pas diagonalisable.

Un cas particulier

Lorsque E est de dimension p et lorsque u possède p valeurs propres distinctes, chacun des sous-espaces propres est de dimension au moins égale à 1 donc la somme des sous-espaces propres est au moins de dimension p . Ceci prouve que la somme de ces sous-espaces propres est égale à E , donc u est diagonalisable, et indique en plus que chacun de ces sous-espaces propres est de dimension 1. C'est le cas par exemple de la matrice A_1 .

Cette situation n'est pas caractéristique de tous les endomorphismes diagonalisables (comme le montre par exemple la matrice A_2), mais quand elle se produit, nous donne une façon simple de justifier que l'endomorphisme est diagonalisable :

PROPOSITION 2.7 — Si E est de dimension p et si $u \in \mathcal{L}(E)$ possède p valeurs propres distinctes alors u est diagonalisable.

Exercice 6

Soit $z \in \mathbb{C}$. Montrer que la matrice $M = \begin{pmatrix} 0 & 0 & z \\ 1 & 0 & 0 \\ 1 & 1 & 0 \end{pmatrix} \in \mathcal{M}_3(\mathbb{C})$ est diagonalisable, sauf pour deux valeurs de z qu'on précisera.

Attention. Pour finir cette section, observons que la dernière caractérisation de la diagonalisation fait intervenir le corps de base $\mathbb{K}$. Lorsqu'il s'agit de diagonaliser un endomorphisme, le corps de base est imposé par l'espace vectoriel, mais lorsqu'il s'agit de diagonaliser une matrice à coefficients réels, il est possible de la considérer comme un élément de $\mathcal{M}_p(\mathbb{R})$ mais aussi comme un élément de $\mathcal{M}_p(\mathbb{C})$. En d'autres termes, une matrice à coefficients réels peut être diagonalisable dans $\mathcal{M}_p(\mathbb{C})$ sans être diagonalisable dans $\mathcal{M}_p(\mathbb{R})$.

Considérons par exemple la matrice $A = \begin{pmatrix} 1 & -1 \\ 1 & 1 \end{pmatrix}$. On calcule $\chi_A = (X-1)^2 + 1$, donc A n'a pas de valeurs propres réelles : elle n'est pas diagonalisable dans $\mathcal{M}_2(\mathbb{R})$. En revanche, elle dispose de deux valeurs propres complexes distinctes $1-i$ et $1+i$ donc est diagonalisable dans $\mathcal{M}_2(\mathbb{C})$.

2.4 Projecteurs spectraux d'un endomorphisme diagonalisable

Considérons une projection vectorielle $p \in \mathcal{L}(E)$ sur H_1 parallèlement à H_2 : on a $E = H_1 \oplus H_2$, et dans une base (e) adaptée à cette décomposition on a

$$\text{Mat}_{(e)}(p) = \begin{pmatrix} \boxed{\begin{array}{c} 1 \\ \diagdown \\ 1 \end{array}} & \\ & \boxed{\begin{array}{c} 0 \\ \diagdown \\ 0 \end{array}} \end{pmatrix}$$

L'endomorphisme p est diagonalisable, $\text{Sp}(p) = \{0, 1\}$, et $H_1 = \text{Ker}(p - \text{Id}_E)$ et $H_2 = \text{Ker } p$ sont les sous-espaces propres associés.

On peut de même considérer la symétrie vectorielle $s \in \mathcal{L}(E)$ par rapport à H_1 , parallèlement à H_2 : sur la même base (e) on a cette fois

$$\text{Mat}_{(s)}(s) = \begin{pmatrix} \boxed{\begin{array}{c} 1 \\ \diagdown \\ 1 \end{array}} & \\ & \boxed{\begin{array}{c} -1 \\ \diagdown \\ -1 \end{array}} \end{pmatrix}$$

L'endomorphisme s est diagonalisable, $\text{Sp}(s) = \{-1, 1\}$, $H_1 = \text{Ker}(s - \text{Id}_E)$ et $H_2 = \text{Ker}(s + \text{Id}_E)$.

Observons enfin que $s = p - (\text{Id}_E - p) = 1 \times p_1 + (-1) \times p_2$, où $p_1 = p$ est la projection vectorielle sur H_1 parallèlement à H_2 et $p_2 = \text{Id}_E - p$ est la projection vectorielle sur H_2 parallèlement à H_1 .

Considérons enfin un endomorphisme diagonalisable $u \in \mathcal{L}(E)$, et la décomposition de E en somme des sous-espaces propres : $E = E_{\lambda_1}(u) \oplus \cdots \oplus E_{\lambda_k}(u)$. Si (e) est une base adaptée à la décomposition de l'espace, nous avons :

$$\text{Mat}_e(u) = \begin{pmatrix} \boxed{\begin{array}{c} \lambda_1 \\ \diagdown \\ \lambda_1 \end{array}} & & \\ & \boxed{\begin{array}{c} \lambda_2 \\ \diagdown \\ \lambda_2 \end{array}} & \\ & & \ddots \\ & & & \boxed{\begin{array}{c} \lambda_k \\ \diagdown \\ \lambda_k \end{array}} \end{pmatrix}$$

La famille $(p_1, \dots, p_k)$ associée à cette décomposition de l'espace est appelée la famille des *projecteurs spectraux*

de u . Rappelons que pour tout $i \in \llbracket 1, k \rrbracket$, p_i est la projection sur $E_{\lambda_i}(u)$ parallèlement à $\bigoplus_{j \neq i} E_{\lambda_j}(u)$. Ainsi,

$$\text{Mat}_e(p_i) = \begin{pmatrix} \boxed{\begin{array}{c} 0 \\ \diagdown \\ 0 \end{array}} & & \\ & \ddots & \\ & & \boxed{\begin{array}{c} 1 \\ \diagdown \\ 1 \end{array}} & & \\ & & & \ddots & \\ & & & & \boxed{\begin{array}{c} 0 \\ \diagdown \\ 0 \end{array}} \end{pmatrix} \quad \leftarrow i^{\text{e}} \text{ bloc}$$

On dispose alors de manière évidente des égalités $\text{Id}_E = \sum_{j=1}^k p_j$ et $u = \sum_{j=1}^k \lambda_j p_j$, et plus généralement :

PROPOSITION 2.8 — Pour tout entier $n \in \mathbb{N}$ on a
$$u^n = \sum_{j=1}^k \lambda_j^n p_j.$$

Remarque. Lorsque u est inversible (c'est à dire lorsque 0 n'est pas valeur propre de u) cette formule s'étend sur $\mathbb{Z}$.

interprétation matricielle

Considérons la diagonalisation de la matrice A_1 obtenue dans l'exercice 3 : $A_1 = PDP^{-1}$ avec $D = \begin{pmatrix} 1 & & \\ & 2 & \\ & & 3 \end{pmatrix}$.

Dans la base de diagonalisation, les trois projecteurs spectraux sont associés aux matrices

$$\begin{pmatrix} 1 & & \\ & 0 & \\ & & 0 \end{pmatrix}, \quad \begin{pmatrix} 0 & & \\ & 1 & \\ & & 0 \end{pmatrix}, \quad \begin{pmatrix} 0 & & \\ & 0 & \\ & & 1 \end{pmatrix}.$$

Dans la base initiale, les trois projecteurs spectraux sont donc associés aux matrices

$$U = P \begin{pmatrix} 1 & & \\ & 0 & \\ & & 0 \end{pmatrix} P^{-1}, \quad V = P \begin{pmatrix} 0 & & \\ & 1 & \\ & & 0 \end{pmatrix} P^{-1}, \quad W = P \begin{pmatrix} 0 & & \\ & 0 & \\ & & 1 \end{pmatrix} P^{-1}.$$

On a $I = U + V + W$, $A_1 = U + 2V + 3W$, et plus généralement : $\forall n \in \mathbb{N}$, $A_1^n = P \begin{pmatrix} 1 & & \\ & 2^n & \\ & & 3^n \end{pmatrix} P^{-1} = U + 2^n V + 3^n W$;

le calcul des matrices U , V et W permet donc d'exprimer aisément A_1^n .

Exercice 7

On considère la matrice A_2 de l'exercice 3. Justifier l'existence (mais sans les calculer) de deux matrices U et V telles que pour tout $n \in \mathbb{N}$, $A_2^n = 2^n U + 4^n V$.

Montrer que ces matrices U et V peuvent s'exprimer en fonction des matrices I et A , et en déduire une expression de A^n en fonction de I et de A .

■ Application à la recherche du commutant d'un endomorphisme diagonalisable

Considérons un endomorphisme u diagonalisable, et posons $\text{Sp}(u) = \{\lambda_1, \dots, \lambda_k\}$. Nous allons chercher à caractériser le *commutant* de u , c'est-à-dire l'ensemble des endomorphismes $v \in \mathcal{L}(E)$ qui vérifient : $u \circ v = v \circ u$.

Le raisonnement que nous allons tenir tient essentiellement au fait suivant :

$$\text{pour tout } x \in E_{\lambda_i}(u), \quad u(v(x)) = u \circ v(x) = v \circ u(x) = v(u(x)) = v(\lambda_i x) = \lambda_i v(x)$$

égalité qui montre que pour tout $x \in E_{\lambda_i}(u)$, $v(x) \in E_{\lambda_i}(u)$: *le sous-espace propre $E_{\lambda_i}(u)$ est stable par v* .

Ceci montre que dans une base adaptée à la décomposition de l'espace en somme de sous-espaces propres, la matrice associée à v est diagonale par blocs.

Bref, dans une telle base nous avons :

$$\text{Mat}_{(e)}(u) = \begin{pmatrix} \boxed{\lambda_1 I} & & \\ & \boxed{\lambda_2 I} & \\ & & \ddots \\ & & & \boxed{\lambda_k I} \end{pmatrix} \quad \text{et} \quad \text{Mat}_{(e)}(v) = \begin{pmatrix} \boxed{A_1} & & \\ & \boxed{A_2} & \\ & & \ddots \\ & & & \boxed{A_k} \end{pmatrix}$$

Réciproquement, il est évident que ces deux matrices (et donc les endomorphismes u et v) commutent. Nous avons donc prouvé la

PROPOSITION 2.9 — Si u est un endomorphisme diagonalisable, les endomorphismes qui commutent avec u sont ceux qui laissent stables les sous-espaces propres.

COROLLAIRE — Lorsque le polynôme caractéristique de u est scindé à racines simples, le commutant est un espace de dimension $p = \dim E$, et les projecteurs spectraux de u en constituent une base.

Exercice 8

Soit $A \in \mathcal{M}_2(\mathbb{R})$ une matrice admettant -1 et 8 pour valeurs propres. Justifier l'existence d'une unique matrice $B \in \mathcal{M}_2(\mathbb{R})$ vérifiant $B^3 = A$, puis exprimer B en fonction de I et de A .

2.5 Polynômes annulateurs et théorie spectrale

Étant donné un endomorphisme u de E et un polynôme $P \in \mathbb{K}[X]$, nous avons les implications :

- si λ est valeur propre de u , $P(\lambda)$ est valeur propre de $P(u)$;
- si $P(u) = 0$, toute valeur propre λ de u est racine de P .

Ce dernier résultat montre qu'il existe un lien entre les racines d'un polynôme annulateur de u et ses valeurs propres. C'est ce que nous allons étudier dans cette partie.

Considérons maintenant un endomorphisme diagonalisable u ; il existe une base (e) pour laquelle :

$$\text{Mat}_e(u) = \begin{pmatrix} \boxed{\begin{matrix} \lambda_1 & \\ & \lambda_1 \end{matrix}} & & \\ & \boxed{\begin{matrix} \lambda_2 & \\ & \lambda_2 \end{matrix}} & \\ & & \ddots \\ & & & \boxed{\begin{matrix} \lambda_k & \\ & \lambda_k \end{matrix}} \end{pmatrix} \quad \text{avec} \quad \text{Sp}(u) = \{\lambda_1, \dots, \lambda_k\}.$$

On constate que le polynôme $P = \prod_{j=1}^k (X - \lambda_j) = \prod_{\lambda \in \text{Sp}(A)} (X - \lambda)$ annule u (on peut même constater que c'est le polynôme minimal). Notons qu'il s'agit d'un polynôme scindé à racines simples.

Nous venons de constater que lorsque u est diagonalisable, il existe un polynôme scindé à racines simples qui annule u . Le fait remarquable est qu'il s'agit d'une équivalence, comme le prouve le théorème :

THÉORÈME 2.10 — Soit $u \in \mathcal{L}(E)$ un endomorphisme de E . Alors u est diagonalisable si et seulement si u est annulé par un polynôme scindé à racines simples.

Attention. Cette preuve ne permet pas d'affirmer que les λ_j sont les valeurs propres de u , car rien ne dit qu'on a bien $E_{\lambda_j}(u) \neq \{0_E\}$. Tout au plus peut-on affirmer que $\text{Sp}(u) \subset \{\lambda_1, \dots, \lambda_k\}$.

Notons en revanche que lorsqu'on connaît les valeurs propres de u , on peut en déduire le résultat suivant :

COROLLAIRE — u est diagonalisable si et seulement s'il est annulé par le polynôme $\prod_{\lambda \in \text{Sp}(u)} (X - \lambda)$.

Notons pour finir que ce résultat permet de prouver le résultat suivant :

PROPOSITION 2.11 — Si u est diagonalisable et si H est un sous-espace vectoriel stable par u , alors l'endomorphisme induit par u sur H est aussi diagonalisable.

Exercice 9

Soit $u \in \mathcal{L}(E)$ un endomorphisme pour lequel il existe une famille libre (e) vérifiant : $u(e_1) = e_1$ et $u(e_2) = e_1 + e_2$. L'endomorphisme u est-il diagonalisable ?

2.6 Le théorème de Cayley-Hamilton

Nous venons donc d'établir un lien entre les deux chapitres d'algèbre linéaire de ce cours : la notion de polynôme annulateur et la notion d'endomorphisme diagonalisable. Il nous reste à énoncer un dernier résultat.

Lorsque u est diagonalisable, le polynôme caractéristique χ_u de u est un multiple de $\prod_{\lambda \in \text{Sp}(u)} (X - \lambda)$. Ce dernier polynôme étant annulateur, il en est de même de χ_u . Le théorème qui suit affirme que ceci reste vrai même lorsque u n'est pas diagonalisable :

THÉORÈME 2.12 (Cayley-Hamilton) — Le polynôme caractéristique de u est un polynôme annulateur de u .

Ce résultat présente bien évidemment l'intérêt de nous fournir un polynôme annulateur de u , mais ce dernier ne sera pas forcément de degré minimal (on se souvient néanmoins que le polynôme minimal de u se trouve parmi ses diviseurs).

3. Matrices et endomorphismes trigonalisables

Nous n'avons pour l'instant pas abordé le cas des endomorphismes non diagonalisables car ce n'est pas un des objectifs principaux de ce cours, mais nous en avons vu un exemple avec la matrice A_3 de l'exercice 3. Nous allons montrer maintenant qu'à défaut d'être diagonalisable, cette matrice est *trigonalisable*, c'est-à-dire semblable à une matrice triangulaire supérieure.

Exercice 10

Trouver une matrice inversible P telle que $A_3 = P \begin{pmatrix} 4 & 0 & 0 \\ 0 & 2 & 1 \\ 0 & 0 & 2 \end{pmatrix} P^{-1}$.

Ceci nous amène aux définitions suivantes :

DÉFINITION. — Un endomorphisme $u \in \mathcal{L}(E)$ est dit *trigonalisable* s'il existe une base de E dans laquelle la matrice associée à u est triangulaire supérieure.

Une matrice A est dite *trigonalisable* si et seulement si elle est semblable à une matrice triangulaire supérieure, c'est à dire s'il existe une matrice triangulaire supérieure T et une matrice inversible P telles que $A = PTP^{-1}$.

Le résultat majeur dont on dispose est le suivant :

THÉORÈME 3.1 — *Un endomorphisme $u \in \mathcal{L}(E)$ (ou une matrice $A \in \mathcal{M}_p(\mathbb{K})$) est trigonalisable si et seulement si son polynôme caractéristique est scindé.*

Remarque. Puisque tout polynôme complexe est scindé, une conséquence importante de ceci est que toute matrice est trigonalisable dans $\mathcal{M}_p(\mathbb{C})$, mais pas nécessairement dans $\mathcal{M}_p(\mathbb{R})$.

Chapitre III

Espaces euclidiens

Élément important de calcul en géométrie euclidienne, le produit scalaire apparaît cependant assez tard dans l'histoire des mathématiques. On en trouve trace chez Hamilton en 1843 lorsqu'il crée le corps des quaternions ou encore chez Peano (associé à un calcul d'aire), et n'est initialement défini qu'à l'aide du cosinus d'un angle. Sa qualité de forme bilinéaire symétrique ne sera exploitée en algèbre linéaire que plus tard et, de propriété, deviendra définition.

Un espace muni d'un produit scalaire sera dit *préhilbertien*², le terme *euclidien* étant réservé aux espaces de dimensions finies.

1. Espaces préhilbertiens

1.1 Produit scalaire

Dans toute cette section, E désigne un $\mathbb{R}$ -espace vectoriel de dimension quelconque.

DÉFINITION. — Un produit scalaire sur E est une forme bilinéaire $\phi : E \times E \rightarrow \mathbb{R}$ vérifiant :

- $\forall (x, y) \in E^2, \phi(x, y) = \phi(y, x)$ (ϕ est symétrique);
- $\forall x \in E, \phi(x, x) \geq 0$ et $\phi(x, x) = 0 \Rightarrow x = 0_E$ (ϕ est définie positive).

Un produit scalaire est donc une forme bilinéaire symétrique définie positive³.

Remarque. Une forme bilinéaire symétrique qui vérifie seulement la propriété $\forall x \in E, \phi(x, x) \geq 0$ sans être nécessairement définie positive est dite *positive*.

On notera par la suite les notations usuelles : $\phi(x, y) = \langle x | y \rangle$, et $\|x\| = \sqrt{\langle x | x \rangle}$, cette dernière expression désignant la *norme euclidienne* associée au produit scalaire⁴.

Un $\mathbb{R}$ -espace vectoriel muni d'un produit scalaire est appelé un *espace préhilbertien réel*.

PROPOSITION 1.1 — L'application $(A, B) \mapsto \text{tr}(A^T B)$ est un produit scalaire sur $\mathcal{M}_{n,p}(\mathbb{R})$. Il s'agit du produit scalaire canonique de $\mathcal{M}_{n,p}(\mathbb{R})$: la base canonique est orthonormée pour ce produit scalaire.

PROPOSITION 1.2 — Soit $\omega : [a, b] \rightarrow \mathbb{R}_+^*$ une fonction continue à valeurs strictement positives, et E l'ensemble des fonctions continues $f : [a, b] \rightarrow \mathbb{R}$. Alors l'application $(f, g) \mapsto \int_a^b f(t)g(t)\omega(t)dt$ est un produit scalaire sur E .

Exemple. L'application $(P, Q) \mapsto \int_{-1}^1 P(t)Q(t)dt$ est un produit scalaire sur $\mathcal{C}^0([-1, 1], \mathbb{R})$, mais aussi sur $\mathbb{R}[X]$.

Utilisation de la bilinéarité

En utilisant la bilinéarité et la symétrie du produit scalaire, on obtient les deux développements suivants :

$$\begin{aligned}\forall (x, y) \in E^2, \quad \|x + y\|^2 &= \|x\|^2 + 2\langle x | y \rangle + \|y\|^2 \\ \|x - y\|^2 &= \|x\|^2 - 2\langle x | y \rangle + \|y\|^2\end{aligned}$$

Ces développements conduisent à diverses *identités de polarisation*, autrement dit des relations qui définissent le produit scalaire à partir de la norme :

$$\langle x | y \rangle = \frac{1}{2}(\|x + y\|^2 - \|x\|^2 - \|y\|^2) \quad \langle x | y \rangle = \frac{1}{2}(\|x\|^2 + \|y\|^2 - \|x - y\|^2) \quad \langle x | y \rangle = \frac{1}{4}(\|x + y\|^2 - \|x - y\|^2)$$

2. Comme ce terme le laisse entendre, il existe aussi des espaces *hilbertiens*, mais leur étude n'est pas au programme.

3. Ces différents termes proviennent de l'étude générale des formes bilinéaires.

4. Il s'agit en effet d'une norme au sens topologique du terme.

On remarquera que ces identités impliquent qu'à une norme euclidienne donnée ne peut correspondre qu'un seul produit scalaire.

THÉORÈME 1.3 (Inégalité de Cauchy-Schwarz) — $\forall (x, y) \in E^2, |\langle x | y \rangle| \leq \|x\| \times \|y\|$.

COROLLAIRE — Il y a égalité dans l'inégalité de Cauchy-Schwarz (autrement dit, $|\langle x | y \rangle| = \|x\| \times \|y\|$) si et seulement si la famille (x, y) est liée.

Exercice 1

Soit $(x_1, x_2, \dots, x_n) \in \mathbb{R}^n$. Montrer que $\left(\sum_{k=1}^n x_k\right)^2 \leq n \sum_{k=1}^n x_k^2$. Dans quel cas y-a-t-il égalité?

THÉORÈME 1.4 (Inégalité triangulaire) — $\forall (x, y) \in E^2, \|x + y\| \leq \|x\| + \|y\|$.

Remarque. Il y a égalité dans l'inégalité triangulaire lorsque $\langle x | y \rangle = \|x\| \times \|y\|$, c'est à dire lorsqu'il y a égalité dans l'inégalité de Cauchy-Schwarz et qu'en plus $\langle x | y \rangle \geq 0$, ce qui impose $x = 0_E$ ou $y = \lambda x$ avec $\lambda \geq 0$.

1.2 Orthogonalité

DÉFINITION. — Soit E un espace préhilbertien réel.

- (i) On dit que deux vecteurs x et y sont orthogonaux lorsque $\langle x | y \rangle = 0$.
- (ii) On dit qu'un vecteur x est orthogonal à un sous-espace vectoriel H lorsque $\forall y \in H, \langle x | y \rangle = 0$.
- (iii) Enfin, deux sous-espaces vectoriels H_1 et H_2 sont orthogonaux lorsque $\forall (x, y) \in H_1 \times H_2, \langle x | y \rangle = 0$.

Remarque. On peut noter que deux sous-espaces vectoriels orthogonaux sont nécessairement en somme directe. En effet, si $x \in H_1 \cap H_2$ alors $\langle x | x \rangle = 0$, ce qui impose $x = 0_E$. On dit alors que la somme $H_1 \oplus H_2$ est une *somme directe orthogonale*, et on pourra éventuellement la noter $H_1 \overset{\perp}{\oplus} H_2$.

THÉORÈME 1.5 (Pythagore) — Soient H_1 et H_2 deux sous-espaces vectoriels orthogonaux, et un vecteur $x = x_1 + x_2 \in H_1 \overset{\perp}{\oplus} H_2$. Alors $\|x\|^2 = \|x_1\|^2 + \|x_2\|^2$.

Nous pouvons noter que réciproquement, si nous avons $\|x_1 + x_2\|^2 = \|x_1\|^2 + \|x_2\|^2$, alors x_1 et x_2 sont nécessairement orthogonaux.

DÉFINITION. — Soit A une partie quelconque de E . On appelle orthogonal de A l'ensemble

$$A^\perp = \{x \in E \mid \forall a \in A, \langle x | a \rangle = 0\}$$

des vecteurs orthogonaux à tout élément de A . Il s'agit d'un sous-espace vectoriel de E .

PROPOSITION 1.6 — Si H est un sous-espace vectoriel et A une partie génératrice de H , alors $H^\perp = A^\perp$.

Remarque. L'intérêt majeur de ce dernier résultat est qu'en dimension finie, déterminer l'orthogonal d'un sous-espace vectoriel H revient à déterminer l'orthogonal d'une base de H .

Lorsque H est un sous-espace vectoriel, $H^\perp$ est donc le plus grand des sous-espaces vectoriels (au sens de l'inclusion) qui soit en somme directe orthogonale avec H : $H \oplus H^\perp$.

Attention cependant, cela ne signifie pas pour autant que cette somme soit égale à E . Il faudra en effet supposer en plus que E est de dimension finie pour pouvoir affirmer que H et $H^\perp$ sont des sous-espaces supplémentaires. Si H_1 et H_2 sont deux sous-espaces vectoriels de E , on dispose enfin des équivalences :

$$H_1 \text{ et } H_2 \text{ sont orthogonaux} \iff H_1 \subset H_2^\perp \iff H_2 \subset H_1^\perp.$$

1.3 Espaces euclidiens

DÉFINITION. — Une famille finie $(e_1, \dots, e_p)$ de vecteurs de E est dite orthonormée lorsque :

$$\forall (i, j) \in \llbracket 1, p \rrbracket^2, \quad \langle e_i | e_j \rangle = \delta_{ij} = \begin{cases} 1 & \text{si } i = j \\ 0 & \text{si } i \neq j \end{cases}$$

PROPOSITION 1.7 — Une famille orthonormée est libre. En particulier, lorsque E est de dimension finie n , une famille orthonormée constituée de n vecteurs est une base de E , dite base orthonormée.

On appelle *espace euclidien* tout espace préhilbertien réel de dimension finie. Le résultat précédent définit la notion de base orthonormée, mais ne prouve pas l'existence de celles-ci. C'est l'objet du théorème suivant :

THÉORÈME 1.8 — Tout espace euclidien possède des bases orthonormées.

Nous reviendrons sur cette construction une fois définie la notion de *projection orthogonale* ; elle prendra alors le nom de *procédé d'orthonormalisation de Gram-Schmidt*.

■ Expression du produit scalaire dans une base orthonormée

Une fois acquise l'existence de bases orthonormées dans un espace euclidien, il reste à constater que les calculs relatifs au produit scalaire sont très simples une fois exprimés dans une telle base.

Soit $(e_1, \dots, e_n)$ une base orthonormée d'un espace euclidien E , et $x = \sum_{i=1}^n x_i e_i$, $y = \sum_{j=1}^n y_j e_j$. On pose $X = \text{Mat}_e(x)$

$$\text{et } Y = \text{Mat}_e(y). \text{ Alors : } \langle x | y \rangle = \sum_{i=1}^n x_i y_i = X^T Y \quad \text{et} \quad \|x\| = \sqrt{\sum_{i=1}^n x_i^2} = \sqrt{X^T X}.$$

En outre, on peut noter que $\langle e_k | x \rangle = \sum_{i=1}^n x_i \langle e_k | e_i \rangle = x_k$, donc on dispose dans un espace euclidien d'une expression simple pour caractériser la décomposition dans une base orthonormée : $x = \sum_{k=1}^n \langle e_k | x \rangle e_k$.

PROPOSITION 1.9 — Toute forme linéaire de E s'écrit de manière unique : $x \mapsto \langle a | x \rangle$, où a est un vecteur de E .

1.4 Projection orthogonale

Revenons maintenant à la notion d'orthogonal d'un sous-espace vectoriel H de E . Nous avons vu que $H \oplus H^\perp$ est une somme directe, mais ce n'est que lorsque H est de dimension finie qu'on sera assuré d'être en présence de deux sous-espaces supplémentaires :

THÉORÈME 1.10 — Si E est un espace préhilbertien et H un sous-espace vectoriel de dimension finie ; alors $E = H \oplus H^\perp$.

COROLLAIRE — Lorsque E est un espace euclidien et H un sous-espace vectoriel de E , on a $\dim(H^\perp) = \dim E - \dim H$ et $H^{\perp\perp} = H$.

Remarque. Dans un espace préhilbertien de dimension quelconque, on peut seulement affirmer que $H \subset H^{\perp\perp}$.

DÉFINITION. — On appelle projection orthogonale sur un sous-espace vectoriel H de dimension finie d'un espace préhilbertien E la projection vectorielle sur H parallèlement à $H^\perp$.

Lorsque p est la projection orthogonale sur H et $(e_1, \dots, e_k)$ une base orthonormée de H on a :

$$\forall x \in E, \quad p(x) = \sum_{j=1}^k \langle e_j | x \rangle e_j$$

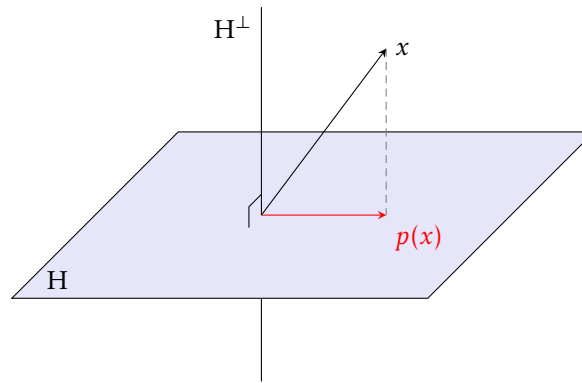


FIGURE 1 – Représentation graphique d'une projection orthogonale.

Remarque. Dans le cas de la projection sur un hyperplan H on a $\dim H = \dim E - 1$ donc $\dim H^\perp = 1$. Si a est un vecteur *unitaire* de $H^\perp$, la projection orthogonale sur $H^\perp$ s'écrit $x \mapsto \langle a | x \rangle a$ et celle sur H s'écrit donc $x \mapsto x - \langle a | x \rangle a$.

Lorsqu'on ne dispose pas d'une base orthonormée de H , on utilise pour caractériser le vecteur $p(x)$ le résultat suivant :

PROPOSITION 1.11 — $p(x)$ est l'unique vecteur de E vérifiant les conditions : $\begin{cases} p(x) \in H \\ x - p(x) \in H^\perp \end{cases}$

Distance à un sous-espace vectoriel

Si $x \in E$ et si H est un sous-espace vectoriel de E , on appelle *distance* de x à H la quantité : $d(x, H) = \inf\{\|x - h\| \mid h \in H\}$. Dans le cas où H est un sous-espace vectoriel de dimension finie d'un espace préhilbertien réel, le résultat suivant nous permet de calculer cette distance :

PROPOSITION 1.12 — L'application $\begin{pmatrix} H & \longrightarrow & \mathbb{R} \\ h & \longmapsto & \|x - h\| \end{pmatrix}$ atteint un minimum en un unique point, à savoir $h = p(x)$. Autrement dit, $d(x, H) = \|x - p(x)\|$.

Exercice 2

Soit E un espace euclidien de dimension 4, (e) une base orthonormée et $u = 3e_1 + 2e_2 - e_3 + e_4$, $v = 2e_1 + 5e_2 - e_4$. On note $H = \text{Vect}(u, v)$. Calculer la distance de H au vecteur $w = e_1 + e_2 + e_3 + e_4$.

Remarque. Nous reviendrons au chapitre IX sur la notion de distance, dans un cadre plus général, celui des espaces vectoriels normés.

■ Orthonormalisation par la méthode de Gram-Schmidt

Considérons une famille libre $(x_1, \dots, x_k)$ de E , et notons p_j la projection orthogonale sur $\text{Vect}(x_1, \dots, x_{j-1})$ (avec la convention $p_1 = 0$). Alors la famille (e) définie par les formules suivantes :

$$\forall j \in \llbracket 1, k \rrbracket, \quad e_j = \frac{x_j - p_j(x_j)}{\|x_j - p_j(x_j)\|}$$

est une famille orthonormée vérifiant : $\forall j \in \llbracket 1, k \rrbracket, \text{Vect}(e_1, \dots, e_j) = \text{Vect}(x_1, \dots, x_j)$. C'est en outre l'unique famille vérifiant en plus les conditions : $\forall j \in \llbracket 1, k \rrbracket, \langle x_j | e_j \rangle > 0$.

Exemple. Considérons l'espace euclidien $\mathbb{R}^3$ muni du produit scalaire usuel, ainsi que la famille de vecteurs

$$x_1 = \begin{pmatrix} 0 \\ 1 \\ 2 \end{pmatrix}, x_2 = \begin{pmatrix} 1 \\ 2 \\ 3 \end{pmatrix} \text{ et } x_3 = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}, \text{ et appliquons lui la méthode de Gram-Schmidt :}$$

$$- e_1 = \frac{x_1}{\|x_1\|} \text{ donc } e_1 = \frac{1}{\sqrt{5}} \begin{pmatrix} 0 \\ 1 \\ 2 \end{pmatrix}.$$

$$- p(x_2) = \langle e_1 | x_2 \rangle e_1 = \frac{8}{5} \begin{pmatrix} 0 \\ 1 \\ 2 \end{pmatrix} \text{ donc } x_2 - p(x_2) = \frac{1}{5} \begin{pmatrix} 5 \\ 2 \\ -1 \end{pmatrix} \text{ et } e_2 = \frac{1}{\sqrt{30}} \begin{pmatrix} 5 \\ 2 \\ -1 \end{pmatrix}.$$

$$- p(x_3) = \langle e_1 | x_3 \rangle e_1 + \langle e_2 | x_3 \rangle e_2 = \frac{2}{3} \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix} \text{ donc } x_3 - p(x_3) = \frac{1}{3} \begin{pmatrix} 1 \\ -2 \\ 1 \end{pmatrix} \text{ et } e_3 = \frac{1}{\sqrt{6}} \begin{pmatrix} 1 \\ -2 \\ 1 \end{pmatrix}.$$

Exercice 3

On muni $\mathbb{R}[X]$ d'un produit scalaire quelconque. À l'aide du procédé de Schmidt appliqué à la base canonique de $\mathbb{R}[X]$, justifier l'existence d'une unique famille $(P_n)_{n \in \mathbb{N}}$ telle que :

- pour tout $n \in \mathbb{N}$, $\deg P_n = n$;
- pour tout $n \in \mathbb{N}$, $\text{cdom}(P_n) = 1$;
- pour tout $i \neq j$, $\langle P_i | P_j \rangle = 0$.

2. Endomorphismes dans un espace euclidien

2.1 Isométries vectorielles

DÉFINITION. — Si E est un espace préhilbertien, on appelle isométrie vectorielle un endomorphisme $u \in \mathcal{L}(E)$ compatible avec le produit scalaire, c'est à dire vérifiant : $\forall (x, y) \in E^2$, $\langle u(x) | u(y) \rangle = \langle x | y \rangle$.

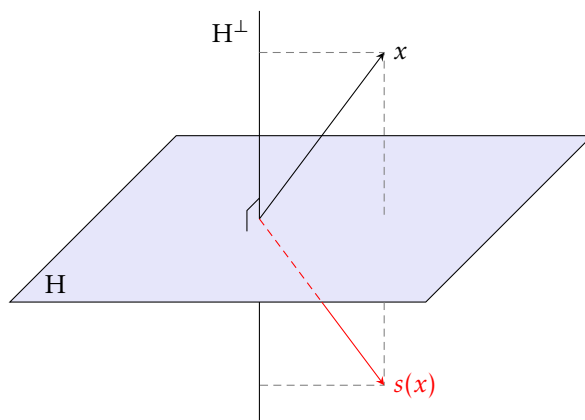
Une telle application est *a fortiori* compatible avec la norme euclidienne : en posant $y = x$ on obtient $\forall x \in E$, $\|u(x)\| = \|x\|$, ce qui explique leur nom. Le fait remarquable est que la réciproque est vraie :

PROPOSITION 2.1 — $u \in \mathcal{L}(E)$ est une isométrie vectorielle si et seulement si $\forall x \in E$, $\|u(x)\| = \|x\|$.

En conséquence de quoi une isométrie vectorielle est injective : en effet, lorsque $u(x) = 0_E$ nous avons $\|x\| = \|u(x)\| = 0$ et donc $x = 0_E$. Et en particulier, lorsque E est de dimension finie, une isométrie vectorielle est nécessairement inversible. Un endomorphisme inversible étant appelé un *automorphisme*, en dimension finie les isométries vectorielles portent aussi le nom d'*automorphisme orthogonal*.

On notera $\mathcal{O}(E)$ l'ensemble des isométries vectorielles de E ; il est appelé le *groupe orthogonal* de E .

Exemple. On appelle *symétrie orthogonale* par rapport à un sous-espace vectoriel H la symétrie par rapport à H , parallèlement à $H^\perp$. Il s'agit d'une isométrie vectorielle.



Posons $x = x_1 + x_2$ avec $x_1 \in H$ et $x_2 \in H^\perp$. Alors $\|s(x)\|^2 = \|x_1\|^2 + \|x_2\|^2 = \|x\|^2$ donc s préserve la norme ; il s'agit bien d'une isométrie vectorielle.

Attention. Une symétrie orthogonale est un automorphisme orthogonal (ie une isométrie vectorielle), mais ce n'est pas le cas d'une projection orthogonale (qui, hormis l'identité, n'est pas inversible).

Remarque. Une symétrie orthogonale par rapport à un hyperplan (un sous-espace vectoriel de dimension $p-1$) est aussi appelée une *réflexion*. En dimension 2, les réflexions sont donc les symétries orthogonales par rapport aux droites, en dimension 3 les symétries orthogonales par rapport aux plans.

Exercice 4

Soient a et b deux vecteurs non nuls distincts d'un espace euclidien E vérifiant $\|a\| = \|b\|$. Montrer qu'il existe une unique réflexion s telle que $s(a) = b$.

PROPOSITION 2.2 — Soit E un espace euclidien, $u \in \mathcal{O}(E)$ une isométrie vectorielle, et H un sous-espace vectoriel de E stable par u . Alors $H^\perp$ est aussi stable par u .

■ Interprétation matricielle d'une isométrie vectorielle

PROPOSITION 2.3 — Soit $(e_1, \dots, e_p)$ une base orthonormée de E , et $u \in \mathcal{L}(E)$. Alors u est une isométrie vectorielle si et seulement si $(u(e_1), \dots, u(e_p))$ est une base orthonormée.

COROLLAIRE — Soit $(e_1, \dots, e_n)$ une base orthonormée, $u \in \mathcal{L}(E)$, et $A = \text{Mat}_e(u)$. Alors u est une isométrie vectorielle si et seulement si $A^T A = I$.

Une matrice $A \in \mathcal{M}_p(\mathbb{R})$ vérifiant l'identité $A^T A = I$ est appelée une matrice *orthogonale*. On note $\mathcal{O}_p(\mathbb{R})$ l'ensemble des matrices orthogonales de $\mathcal{M}_p(\mathbb{R})$; ensemble qu'on appelle le *groupe orthogonal* d'ordre p .

Remarque. Si on observe que $\text{Mat}_{(e)}(u) = \text{Mat}_{(e)}(u(e_1), u(e_2), \dots, u(e_p))$, on peut affirmer qu'une matrice orthogonale est une matrice dont les colonnes forment une famille orthonormée pour le produit scalaire usuel. C'est souvent par l'intermédiaire de cette propriété que l'on reconnaît une matrice orthogonale. Une autre conséquence de cette observation réside dans la :

PROPOSITION 2.4 — La matrice de passage entre deux bases orthonormées est une matrice orthogonale.

Structure de groupe

Le vocable *groupe* a une signification particulière en mathématiques, et ce n'est pas par hasard s'il est employé ici. Sans rentrer dans les détails, l'emploi de ce terme implique les propriétés suivantes :

- (i) la matrice I_p est orthogonale : $I_p \in \mathcal{O}_p(\mathbb{R})$;
- (ii) si A et B sont orthogonales, AB est aussi orthogonale : $(A, B) \in \mathcal{O}_p(\mathbb{R})^2 \implies AB \in \mathcal{O}_p(\mathbb{R})$;
- (iii) si A est orthogonale, A^{-1} aussi : $A \in \mathcal{O}_p(\mathbb{R}) \implies A^{-1} \in \mathcal{O}_p(\mathbb{R})$.

Notons en outre pour ce dernier point que $A^{-1} = A^T$.

PROPOSITION 2.5 — Soit u une isométrie vectorielle (respectivement A une matrice orthogonale). Alors $\det u \in \{-1, 1\}$ ($\det A \in \{-1, 1\}$).

Ce dernier résultat permet de séparer les isométries vectorielles en deux classes : ceux dont le déterminant est égal à 1 (les isométries directes) : $\mathcal{SO}(E) = \{u \in \mathcal{O}(E) \mid \det u = 1\}$, qui forment eux aussi un groupe, appelé le *groupe spécial orthogonal*, et ceux dont le déterminant est égal à -1 (les isométries indirectes, qui n'ont pas de structure algébrique particulière). On notera bien entendu de même : $\mathcal{SO}_p(\mathbb{R}) = \{A \in \mathcal{O}_p(\mathbb{R}) \mid \det A = 1\}$.

Exemple. Si u est une réflexion, alors $\det u = -1$.

Orientation d'un espace euclidien

Considérons l'ensemble $\mathcal{B}$ des bases orthonormées d'un espace euclidien E . Si (e) et (e') sont deux éléments de $\mathcal{B}$, $P = \text{Mat}_{(e)}(e')$ est une matrice orthogonale donc $\det P = \pm 1$. On définit donc une relation $\mathcal{R}$ sur $\mathcal{B}$ en posant : $(e) \mathcal{R} (e') \iff \det \text{Mat}_{(e)}(e') = 1$.

PROPOSITION 2.6 — La relation $\mathcal{R}$ est une relation d'équivalence qui possède deux classes d'équivalence distinctes.

DÉFINITION. — Orienter l'espace E , c'est choisir l'une de ces deux classes d'équivalence; les bases orthonormées de cette classe seront qualifiées de bases directes, les autres de bases indirectes.

Remarque. Pour orienter l'espace, il suffit de choisir une base (e) et la qualifier de directe. Une fois ce choix fait, une base orthonormée (e') sera directe si $\det_{(e)}(e') = 1$, et indirecte si $\det_{(e)}(e') = -1$.

PROPOSITION 2.7 — Si l'espace E est orienté et si $u \in \mathcal{O}(E)$ est une isométrie vectorielle, alors u appartient à $\mathcal{SO}(E)$ si et seulement si l'image par u d'une base orthonormée directe est une base orthonormée directe. Autrement dit, les isométries directes sont celles qui préservent l'orientation de l'espace.

Une dernière conséquence de la notion de base orthonormée directe est le

THÉORÈME 2.8 — Si (e) et (e') sont deux bases orthonormées directes et $(x_1, \dots, x_p)$ une famille de p vecteurs de E alors $\det_{(e)}(x_1, \dots, x_p) = \det_{(e')}(x_1, \dots, x_p)$. Autrement dit, le déterminant d'une famille de vecteurs ne dépend pas du choix de la base orthonormée directe dans laquelle on réalise le calcul.

2.2 Isométries vectorielles d'un plan euclidien

Dans cette section, E désigne un plan euclidien orienté, (e_1, e_2) une base orthonormée directe, et $u \in \mathcal{O}(E)$.

Notons $A = \text{Mat}_{(e)}(u) = \begin{pmatrix} a & c \\ b & d \end{pmatrix}$. A est une matrice orthogonale, ce qui se traduit par

$$\begin{cases} a^2 + b^2 = 1 \\ c^2 + d^2 = 1 \\ ac + bd = 0 \end{cases}$$

La première égalité traduit l'existence d'un réel α (défini de manière unique modulo 2π) tel que $a = \cos \alpha$ et $b = \sin \alpha$. De même, la seconde égalité traduit l'existence d'un réel β pour lequel $d = \cos \beta$ et $c = \sin \beta$.

La troisième égalité s'écrit alors $\cos \alpha \sin \beta + \sin \alpha \cos \beta = 0$, soit $\sin(\alpha + \beta) = 0$. Ainsi, nous avons $\beta \equiv -\alpha \pmod{\pi}$, ce qui laisse deux possibilités (sachant que β est unique modulo 2π) : $\beta = -\alpha$ ou $\beta = \pi - \alpha$.

Les matrices de $\mathcal{O}_2(\mathbb{R})$ sont donc de deux types uniquement :

$$A_1 = \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix} \quad (\text{lorsque } \beta = -\alpha) \quad \text{ou} \quad A_2 = \begin{pmatrix} \cos \alpha & \sin \alpha \\ \sin \alpha & -\cos \alpha \end{pmatrix} \quad (\text{lorsque } \beta = \pi - \alpha).$$

$\det A_1 = 1$ et $\det A_2 = -1$ donc les isométries vectorielles directes sont associées aux matrices de type A_1 , et les isométries indirectes aux matrices de type A_2 .

■ Isométries directes du plan euclidien orienté

Nous avons donc $\mathcal{SO}_2(\mathbb{R}) = \{R(\alpha) \mid \alpha \in [0, 2\pi[\}$, où $R(\alpha)$ désigne la matrice $\begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix}$.

PROPOSITION 2.9 — $\mathcal{SO}_2(\mathbb{R})$ est un groupe commutatif, et $R(\alpha)R(\beta) = R(\alpha + \beta)$.

COROLLAIRE — Si $\text{Mat}_{(e)}(u) = R_\alpha$, la valeur de α est indépendante du choix de la base orthonormée directe (e) ; on dit que u est la rotation d'angle α .

Remarque. Les matrices de $\mathcal{SO}_2(\mathbb{R})$ sont aussi les matrices de passage d'une base orthonormée directe à une autre; ainsi nous venons de prouver que nous ne pouvons passer d'une base orthonormée directe à une autre que par l'action d'une rotation.

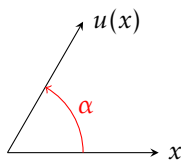
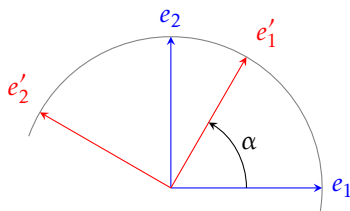
FIGURE 2 – Action d'une rotation vectorielle d'angle α .

FIGURE 3 – Passage d'une base orthonormée directe à une autre.

Comment mesurer l'angle d'une rotation ?

Si x est un vecteur non nul, il est toujours possible de construire une base orthonormée directe (e_1, e_2) telle que $e_1 = \frac{x}{\|x\|}$. On a alors $u(e_1) = \cos \alpha e_1 + \sin \alpha e_2$ et

$$\langle e_1 | u(e_1) \rangle = \cos \alpha \quad \text{et} \quad \det(e_1, u(e_1)) = \sin \alpha$$

On en déduit deux formules qui permettent de calculer $\cos \alpha$ et $\sin \alpha$ et par leur intermédiaire de déterminer l'angle d'une rotation à partir de l'image d'un vecteur non nul quelconque :

$$\cos \alpha = \frac{\langle x | u(x) \rangle}{\|x\|^2} \quad \text{et} \quad \sin \alpha = \frac{\det(x, u(x))}{\|x\|^2}$$

ce dernier déterminant pouvant être calculé dans une base orthonormée directe quelconque.

■ Isométries indirectes du plan euclidien orienté

Revenons à la matrice $A_2 = \begin{pmatrix} \cos \alpha & \sin \alpha \\ \sin \alpha & -\cos \alpha \end{pmatrix}$ et cherchons à la diagonaliser.

On calcule sans peine $\chi_{A_2}(x) = x^2 - 1 = (x-1)(x+1)$ donc $\text{Sp}(A_2) = \{-1, 1\}$; la matrice A_2 est diagonalisable.

On résout $A_2 X = X \iff X \in \text{Vect} \begin{pmatrix} \cos(\alpha/2) \\ \sin(\alpha/2) \end{pmatrix}$ et $A_2 X = -X \iff X \in \text{Vect} \begin{pmatrix} -\sin(\alpha/2) \\ \cos(\alpha/2) \end{pmatrix}$

Posons $P = \begin{pmatrix} \cos(\alpha/2) & -\sin(\alpha/2) \\ \sin(\alpha/2) & \cos(\alpha/2) \end{pmatrix}$; P est la matrice de la rotation d'angle $\alpha/2$ donc la base (e') obtenue par

$\text{Mat}_{(e)}(e')$ est une base orthonormée directe dans laquelle $\text{Mat}_{(e')}(u) = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$. L'isométrie indirecte $u \in \mathcal{O}(E)$

est donc une symétrie orthogonale par rapport à la droite engendrée par le vecteur e'_2 .

Nous avons donc prouvé que les isométries indirectes du plan euclidien orienté sont les réflexions, autrement dit les symétries orthogonales par rapports aux droites.

2.3 Endomorphismes autoadjoints

DÉFINITION. — On dit qu'un endomorphisme $u \in \mathcal{L}(E)$ est autoadjoint lorsqu'il vérifie : $\forall (x, y) \in E^2$,

$$\langle u(x) | y \rangle = \langle x | u(y) \rangle.$$

THÉORÈME 2.10 — Si (e) est une base orthonormée de E et $A = \text{Mat}_e(u)$, alors u est autoadjoint si et seulement si $A^T = A$, c'est à dire si et seulement si A est symétrique.

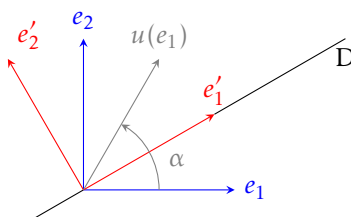


FIGURE 4 – Action d'une réflexion par rapport à la droite D.

Remarque. Pour cette raison, un endomorphisme autoadjoint est aussi appelé un endomorphisme *symétrique*, mais cette appellation peut être trompeuse, car si (e) n'est pas une base orthonormée, la matrice associée dans cette base à un endomorphisme autoadjoint peut ne pas être symétrique.

Adjoint d'un endomorphisme (hors programme)

Lorsque (e) est une base orthonormée de E , $u \in \mathcal{L}(E)$ et $A = \text{Mat}_{(e)}(u)$ on a

$$\langle u(x) | y \rangle = (AX)^T Y = X^T (A^T Y) = \langle x | u^*(y) \rangle$$

où u^* est l'endomorphisme défini par $\text{Mat}_{(e)}(u^*) = A^T$. Cet endomorphisme est appelé l'*adjoint* de l'endomorphisme u (il est facile de montrer que sa définition ne dépend pas du choix de la base orthonormée (e)). On comprend dès lors la dénomination des endomorphismes autoadjoints : les endomorphismes $u \in \mathcal{L}(E)$ qui vérifient $u^* = u$.

PROPOSITION 2.11 — L'ensemble $\mathcal{S}(E)$ des endomorphismes autoadjoints de E est un sous-espace vectoriel de $\mathcal{L}(E)$, de dimension $\frac{p(p+1)}{2}$.

Exercice 5

Soient u et v deux endomorphismes autoadjoints. Montrer que $u \circ v$ est autoadjoint si et seulement si $u \circ v = v \circ u$.

■ Réduction des endomorphismes autoadjoints

De nombreuses applications des endomorphismes autoadjoints résultent du fait que ce sont les seuls endomorphismes diagonalisables dans les bases orthonormées, résultat que nous allons nous attacher à prouver maintenant.

PROPOSITION 2.12 — Si u est un endomorphisme autoadjoint, ses sous-espaces propres sont en somme directe orthogonale.

PROPOSITION 2.13 — Soit H un espace vectoriel stable par un endomorphisme autoadjoint u . Alors $H^\perp$ est aussi stable par u .

LEMME — Un endomorphisme autoadjoint possède au moins une valeur propre réelle.

THÉORÈME 2.14 (théorème spectral) — Un endomorphisme autoadjoint est diagonalisable dans une base orthonormée.

COROLLAIRE — Si A est une matrice symétrique, il existe une matrice diagonale D et une matrice orthogonale P telles que $A = PDP^T$ (rappelons que $P^{-1} = P^T$).

Exercice 6

Diagonaliser sur une base orthonormée la matrice $A = \begin{pmatrix} 2 & -1 & 2 \\ -1 & 2 & 2 \\ 2 & 2 & -1 \end{pmatrix}$.

2.4 Formes bilinéaires symétriques

Nous allons maintenant revenir à la définition d'un produit scalaire : une forme bilinéaire, symétrique, définie positive.

Commençons par considérer une forme bilinéaire $b : E \times E \rightarrow \mathbb{R}$, où E est un espace euclidien.

Si (e) est une base orthonormée de E , posons $x = \sum_{i=1}^p x_i e_i$ et $y = \sum_{j=1}^p y_j e_j$. Alors :

$$b(x, y) = \sum_{i=1}^p \sum_{j=1}^p x_i y_j b(e_i, e_j) = X^T B Y \quad \text{avec} \quad B = (b(e_i, e_j))_{1 \leq i, j \leq p}$$

Si on suppose de plus b symétrique alors $b(e_i, e_j) = b(e_j, e_i)$ donc $B \in \mathcal{S}_p(\mathbb{R})$.

Notons $u \in \mathcal{S}(E)$ l'endomorphisme autoadjoint défini par $\text{Mat}_{(e)}(u) = B$. Nous avons démontré que toute forme bilinéaire symétrique de E s'écrit de manière unique sous la forme $(x, y) \mapsto \langle x | u(y) \rangle$ avec $u \in \mathcal{S}(E)$.

D'après le théorème spectral, il existe une base orthonormée (e') formée de vecteurs propres de u . En notant $P \in \mathcal{O}_p(\mathbb{R})$ la matrice de passage de (e) vers (e') et en posant $X' = P^T X$ et $Y' = P^T Y$ on a $b(x, y) = X'^T D Y'$ avec

$$D = \text{Mat}_{(e')}(u) = \text{diag}(\lambda_1, \dots, \lambda_p) \text{ soit } b(x, y) = \sum_{k=1}^p \lambda_k x'_k y'_k.$$

En particulier, $b(x, x) = \sum_{k=1}^p \lambda_k x'^2_k$ donc :

- b est positive lorsque pour tout $k \in \llbracket 1, p \rrbracket$, $\lambda_k \geq 0$, soit $\text{Sp}(u) \subset \mathbb{R}_+$;
- b est définie positive lorsque pour tout $k \in \llbracket 1, p \rrbracket$, $\lambda_k > 0$, soit $\text{Sp}(u) \subset \mathbb{R}_+^*$.

DÉFINITION. — Une endomorphisme autoadjoint $u \in \mathcal{S}(E)$ est dit positif lorsque pour tout $x \in E$, $\langle x | u(x) \rangle \geq 0$; un endomorphisme autoadjoint $u \in \mathcal{S}(E)$ est dit défini positif lorsque pour tout $x \in E \setminus \{0_E\}$, $\langle x | u(x) \rangle > 0$.

On note $\mathcal{S}^+(E)$ l'ensemble des endomorphismes autoadjoints positifs et $\mathcal{S}^{++}(E)$ celui des endomorphismes autoadjoints définis positifs.

THÉORÈME 2.15 — Soit $u \in \mathcal{S}(E)$ un endomorphisme autoadjoint. Alors :

- u est positif si et seulement si $\text{Sp}(u) \subset \mathbb{R}_+$;
- u est défini positif si et seulement si $\text{Sp}(u) \subset \mathbb{R}_+^*$.

Traduits matriciellement ces résultats donnent :

THÉORÈME 2.16 — Soit $A \in \mathcal{S}_p(\mathbb{R})$ une matrice symétrique. Alors :

- $(\forall X \in \mathbb{R}^p, X^T A X \geq 0)$ si et seulement si $\text{Sp}(A) \subset \mathbb{R}_+$;
- $(\forall X \in \mathbb{R}^p \setminus \{0\}, X^T A X > 0)$ si et seulement si $\text{Sp}(A) \subset \mathbb{R}_+^*$.

On note $\mathcal{S}_p^+(\mathbb{R})$ l'ensemble des matrices symétriques positives, et $\mathcal{S}_p^{++}(\mathbb{R})$ celui des matrices symétriques définies positives.

Application à la réduction d'une forme quadratique

Considérons une forme quadratique $q : \begin{pmatrix} \mathbb{R}^2 & \longrightarrow & \mathbb{R} \\ (x_1, x_2) & \longmapsto & ax_1^2 + 2bx_1x_2 + cx_2^2 \end{pmatrix}$. Nous aurons besoin dans le chapitre consacré aux fonctions à plusieurs variables de déterminer si une telle fonction garde un signe constant ou non.

Observons que $q(x_1, x_2) = X^T A X$ où $X = \begin{pmatrix} x_1 \\ x_2 \end{pmatrix}$ et $A = \begin{pmatrix} a & b \\ b & c \end{pmatrix}$. La matrice A étant symétrique se diagonalise dans

une base orthonormée : $A = P D P^T$ où $P \in \mathcal{O}_2(\mathbb{R})$ et $D = \text{diag}(\lambda_1, \lambda_2)$. En posant $X' = P^T X = \begin{pmatrix} x'_1 \\ x'_2 \end{pmatrix}$ on a alors $q(x_1, x_2) = \lambda_1 (x'_1)^2 + \lambda_2 (x'_2)^2$ donc :

- si $A \in \mathcal{S}_2^{++}(\mathbb{R})$, pour tout $(x_1, x_2) \in \mathbb{R}^2 \setminus \{(0, 0)\}$, $q(x_1, x_2) > 0$;
- si $A \in \mathcal{S}_2^+(\mathbb{R})$, pour tout $(x_1, x_2) \in \mathbb{R}^2$, $q(x_1, x_2) \geq 0$.

Matrices de Gram (hors programme)

On appelle *matrice de Gram* toute matrice $A \in \mathcal{M}_p(\mathbb{R})$ qui s'écrit $A = M^T M$ avec $M \in \mathcal{M}_p(\mathbb{R})$.

Une matrice de Gram est bien évidemment symétrique : $A^T = (M^T M)^T = M^T M = A$; de plus, ses valeurs propres sont positives. En effet, pour tout $x \in \mathbb{R}^n$, $\langle Ax | x \rangle = (Ax)^T x = x^T Ax = x^T M^T M x = \|Mx\|^2 \geq 0$.

Le fait remarquable réside dans la réciproque :

PROPOSITION 2.17 — Une matrice $A \in \mathcal{M}_p(\mathbb{R})$ est symétrique positive si et seulement s'il existe $M \in \mathcal{M}_p(\mathbb{R})$ telle que $A = M^T M$. De plus, A est définie positive si et seulement si $M \in \mathcal{GL}_p(\mathbb{R})$.

Chapitre IV

Suites et séries numériques

Outil de base en analyse, la notion de *suite numérique* apparaît très tôt dans l'histoire des sciences, accompagnée de l'idée intuitive de la convergence. Cependant, il faut attendre le XIX^e siècle et les travaux de Cauchy et de Weierstrass pour substituer aux concepts intuitifs qui avaient prévalu jusque là les définitions que nous connaissons.

1. Suites réelles ou complexes

1.1 Convergence des suites numériques

On qualifiera de *suite numérique* toute suite à valeurs dans $\mathbb{R}$ ou $\mathbb{C}$.

DÉFINITION. — Une suite numérique (u_n) est dite bornée lorsqu'il existe un réel $B \geq 0$ tel que pour tout $n \in \mathbb{N}$, $|u_n| \leq B$.

Dans le cas réel, cette définition est équivalente à dire que la suite est majorée et minorée. Cependant on lui préférera en général la définition ci-dessus, qui présente deux avantages :

- cette définition est valable aussi bien dans $\mathbb{R}$ que dans $\mathbb{C}$ (et, au prix d'une modification mineure, dans le cas des espaces vectoriels);
- elle traduit le concept à l'aide d'une inégalité entre nombres positifs, ce qui évite de nombreuses erreurs de manipulations d'inégalités.

DÉFINITION. — On dit qu'une suite (u_n) numérique converge vers une limite finie ℓ lorsque la distance de u_n à ℓ tend vers 0 : $\lim_{n \rightarrow +\infty} |u_n - \ell| = 0$. Ceci revient donc à écrire :

$$\forall \epsilon > 0, \exists N \in \mathbb{N} \mid n \geq N \Rightarrow |u_n - \ell| \leq \epsilon$$

Autrement dit, il existe un rang à partir duquel tous les termes de la suite (u_n) sont à une distance de ℓ inférieure à une quantité arbitrairement petite ϵ .

Exercice 1

Démontrer les propriétés suivantes :

- toute suite convergente est bornée;
- toute suite convergente possède une *unique* limite;
- toute suite extraite d'une suite convergente converge vers la même limite.

THÉORÈME 1.1 (Cesàro) — Soit (u_n) une suite numérique qui converge vers une limite ℓ . Pour tout $n \in \mathbb{N}$ on pose

$$v_n = \frac{1}{n+1} \sum_{k=0}^n u_k. \text{ Montrer que la suite } (v_n) \text{ converge vers } \ell.$$

Exercice 2

Déduire du théorème de Cesàro le *lemme de l'escalier* : si une suite numérique (u_n) vérifie : $\lim(u_{n+1} - u_n) = \ell$ alors $\lim \frac{u_n}{n} = \ell$.

1.2 Le cas particulier des suites réelles

Contrairement à $\mathbb{C}$, $\mathbb{R}$ est muni d'une *relation d'ordre*. Celle-ci confère aux suites réelles des propriétés uniques qui n'ont pas d'équivalent dans $\mathbb{C}$, ainsi que dans les autres ensembles dans lesquels nous étendrons le concept de limite.

Limites infinies

La première particularité des suites réelles est de caractériser deux cas particuliers de divergence : la divergence vers $-\infty$ et vers $+\infty$:

DÉFINITION. — Une suite réelle (u_n) diverge vers $+\infty$ lorsque : $\forall A \in \mathbb{R}, \exists N \in \mathbb{N} \mid n \geq N \Rightarrow u_n \geq A$.

Une suite réelle (u_n) diverge vers $-\infty$ lorsque : $\forall A \in \mathbb{R}, \exists N \in \mathbb{N} \mid n \geq N \Rightarrow u_n \leq A$.

Autrement dit, une suite (u_n) diverge vers $+\infty$ lorsque u_n est, à partir d'un certain rang, supérieure à une quantité arbitrairement grande A .

PROPOSITION 1.2 — Une suite réelle qui diverge vers $+\infty$ est minorée mais pas majorée.

De même, une suite qui diverge vers $-\infty$ est majorée mais non minorée.

Compatibilité avec la relation d'ordre

PROPOSITION 1.3 (passage à la limite dans une inégalité) — Si (u_n) et (v_n) sont deux suites réelles convergeant respectivement vers α et β et vérifiant : $\forall n \in \mathbb{N}, u_n \leq v_n$, alors $\alpha \leq \beta$.

THÉORÈME 1.4 (encadrement) — Soient (u_n) , (v_n) et (w_n) trois suites réelles telles que pour tout $n \in \mathbb{N}$, $u_n \leq v_n \leq w_n$. On suppose que (u_n) et (w_n) convergent vers la même limite ℓ . Alors (v_n) converge vers ℓ .

THÉORÈME 1.5 (minoration) — Soient (u_n) et (v_n) deux suites réelles telles que pour tout $n \in \mathbb{N}$, $u_n \leq v_n$. On suppose que (u_n) diverge vers $+\infty$. Alors (v_n) diverge vers $+\infty$.

■ Suites monotones

THÉORÈME 1.6 — Une suite croissante et majorée converge ; une suite croissante et non majorée diverge vers $+\infty$.

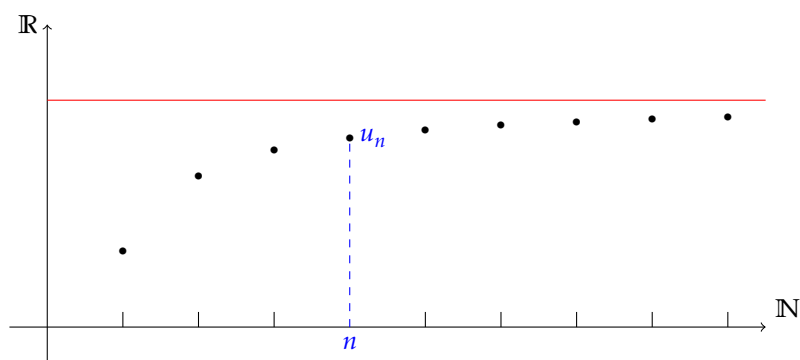


FIGURE 1 – La limite d'une suite croissante et majorée est la borne supérieure de la suite.

Remarque. Bien entendu, une suite décroissante est convergente lorsque elle est minorée, et diverge vers $-\infty$ dans le cas contraire.

Enfin, à la notion de suite monotone est attaché le concept de *suites adjacentes*, utile car fournissant une approximation par défaut et par excès de leur limite commune.

DÉFINITION. — Deux suites (u_n) et (v_n) sont dites adjacentes lorsque (u_n) est croissante, (v_n) décroissante, et $\lim_{+\infty} (v_n - u_n) = 0$.

THÉORÈME 1.7 — Si (u_n) et (v_n) sont adjacentes, alors $\forall n \in \mathbb{N}$, $u_n \leq v_n$, et ces deux suites convergent vers la même limite ℓ ; ℓ est l'unique réel tel que pour tout $n \in \mathbb{N}$, $u_n \leq \ell \leq v_n$.

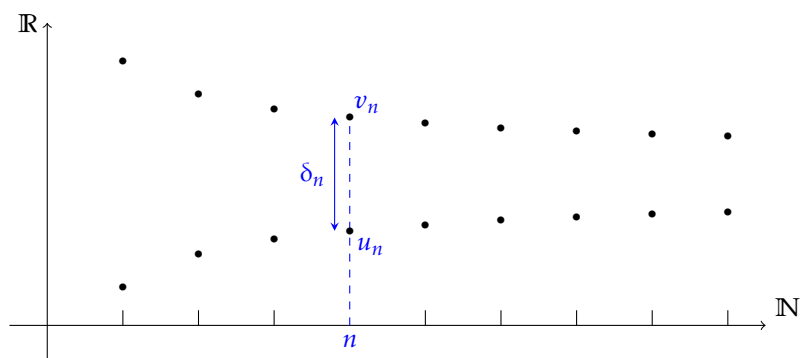


FIGURE 2 – L'écart entre deux suites adjacentes tend vers 0 en décroissant.

Exercice 3

On pose $a_n = \sum_{k=0}^n \frac{1}{k!}$ et $b_n = a_n + \frac{1}{n \cdot n!}$. Montrer que les deux suites (a_n) et (b_n) sont adjacentes, puis montrer que leur limite commune est irrationnelle.

Remarque. Nous aurons l'occasion de prouver plus tard dans l'année que la limite commune aux deux suites de cet exercice est le nombre de Neper e (la base du logarithme naturel). Les deux suites (a_n) et (b_n) permettent donc d'obtenir une approximation par défaut et par excès de cette quantité, en utilisant le script Python suivant.

```
from math import factorial

n = 1
a = 2
while 1 / (n * factorial(n)) > 1e-12:
    n += 1
    a += 1 / factorial(n)
b = a + 1 / (n * factorial(n))
```

```
In [1]: a, b
Out[1]: (2.71828182845823, 2.7182818284590495)
```

Si on fait abstraction des erreurs de calcul inhérentes à la manipulation des flottants en machine, nous pouvons affirmer que $2,718\,281\,828\,458\,23 < e < 2,718\,281\,828\,459\,049\,5$, ce qui fournit les première décimales de $e \approx 2,718\,281\,828\,45\dots$.

1.3 Comparaison asymptotique

Peut-on dire qu'une suite converge lentement ou rapidement? Dans l'absolu, cette question n'a pas de sens, puisque la notion de vitesse est une notion *relative*. Il importe donc d'avoir des éléments de comparaison, composés à la fois de *suites de référence* et d'*outils de comparaison*.

■ Les notations de Landau

Soient (u_n) et (v_n) deux suites numériques telles que : $\forall n \in \mathbb{N}, v_n \neq 0$.

On dit que u_n est *dominée* par v_n lorsque la suite $\left(\frac{u_n}{v_n}\right)$ est bornée. On note dans ce cas $u_n = O(v_n)$.

On dit que u_n est *négligeable* devant v_n lorsque $\lim_{n \rightarrow +\infty} \frac{u_n}{v_n} = 0$. On note dans ce cas $u_n = o(v_n)$.

On dit que (u_n) et (v_n) sont *équivalentes* lorsque $\lim_{n \rightarrow +\infty} \frac{u_n}{v_n} = 1$. On note dans ce cas $u_n \sim v_n$.

On peut noter que les suites (u_n) et (v_n) sont équivalentes si et seulement si $u_n = v_n + o(v_n)$.

Suites de références

Ces notations n'ont d'intérêt que pour comparer des infiniment petits (des suites qui tendent vers 0) ou des infiniment grands (des suites qui tendent vers $+\infty$) entre eux. Si deux suites (u_n) et (v_n) convergent vers une limite commune ℓ , on comparera les suites $(u_n - \ell)$ à $(v_n - \ell)$ entre elles pour mesurer leurs vitesses de convergence relatives.

En outre, *dans la pratique* la suite (v_n) est le plus souvent une suite de référence, c'est-à-dire une suite dont on connaît le comportement. En ce qui nous concerne, les suites de références au voisinage de $+\infty$ seront composées des fonctions $(\ln n)^\alpha$ (avec $\alpha > 0$), n^β (avec $\beta > 0$) et $e^{\gamma n}$ (avec $\gamma > 0$). À ce sujet, rappelons le principe dit des *croissances comparées* :

$$\forall \alpha, \beta, \gamma > 0, \quad \boxed{(\ln n)^\alpha = o(n^\beta) \quad \text{et} \quad n^\beta = o(e^{\gamma n})}$$

Les suites de référence au voisinage de 0 sont les inverses des trois suites précédentes. Ainsi,

$$e^{-\gamma n} = o\left(\frac{1}{n^\beta}\right) \quad \text{et} \quad \frac{1}{n^\beta} = o\left(\frac{1}{(\ln n)^\alpha}\right).$$

Exercice 4

Ordonner les suites ci-dessous à l'aide de la relation « *est négligeable devant* » :

$$n^2 e^n \quad n \ln^2 n + n^2 \quad \frac{n^3}{\ln n} \quad \frac{e^n}{n \ln n} \quad n + \ln \sqrt{n} \quad \frac{n^2}{n + \ln n} \quad n^2 \ln^2 n$$

■ Comparaison logarithmique

La notion que nous allons introduire consiste à comparer deux suites positives (concrètement une suite à étudier et une suite de référence) par le biais du quotient de deux termes consécutifs de ces suites. Cette technique repose sur le résultat suivant :

LEMME — Si (u_n) et (v_n) sont deux suites de réels strictement positifs telles qu'à partir d'un certain rang, $\frac{u_{n+1}}{u_n} \leq \frac{v_{n+1}}{v_n}$, alors $u_n = O(v_n)$.

Dans la pratique, nous nous contenterons de prendre pour l'une de ces deux suites une suite géométrique :

PROPOSITION 1.8 (comparaison à une suite géométrique) — Soit (u_n) une suite de réels strictement positifs. On suppose l'existence d'un réel positif a tel qu'à partir d'un certain rang, $\frac{u_{n+1}}{u_n} \leq a$. Alors $u_n = O(a^n)$.

De même, s'il existe un rang à partir duquel $a \leq \frac{u_{n+1}}{u_n}$ alors $a^n = O(u_n)$.

Exercice 5

Montrer que si $a < e < b$ alors $n^n b^{-n} = O(n!)$ et $n! = O(n^n a^{-n})$.

2. Séries numériques

2.1 Généralités

À une suite réelle ou complexe (u_n) on associe la suite (S_n) dont le terme général est défini par :

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

La suite (S_n) est la suite des *sommes partielles* associée à la série $\sum u_n$ de terme général u_n .

DÉFINITION. — On dit que la série $\sum u_n$ converge lorsque la suite (S_n) converge, et qu'elle diverge dans le cas contraire. En cas de convergence, on pose : $S = \lim_{n \rightarrow +\infty} S_n$, et on écrira : $S = \sum_{k=0}^{+\infty} u_k = \lim_{n \rightarrow +\infty} \sum_{k=0}^n u_k$.

Enfin, lorsque qu'une série converge, on appelle *reste d'ordre n* la quantité : $R_n = S - S_n = \sum_{k=n+1}^{+\infty} u_k$. On définit ainsi une suite (R_n) qui converge vers 0.

Exemple. Séries géométriques.

Soit $a \in \mathbb{C}$, et $u_n = a^n$. Si $a \neq 1$ on a $S_n = \sum_{k=0}^n a^k = \frac{1-a^{n+1}}{1-a}$; si $a = 1$ on a $S_n = n+1$.

- Lorsque $|a| < 1$, $\lim_{n \rightarrow +\infty} a^{n+1} = 0$ donc la série $\sum a^k$ converge, et $\sum_{k=0}^{+\infty} a^k = \frac{1}{1-a}$;
- Lorsque $|a| \geq 1$, la série $\sum a^k$ diverge en vertu de la proposition 2.1.

Attention. Les résultats concernant les opérations sur les limites permettent de prouver que la somme de deux séries convergentes est encore convergente, ou que le produit par un scalaire d'une série convergente est encore convergente. Attention néanmoins à ne pas commettre l'erreur suivante : la série $\sum (u_n + v_n)$ peut être convergente sans que les séries $\sum u_n$ et $\sum v_n$ le soient. Autrement dit, avant d'écrire que :

$$\sum_{n=0}^{+\infty} (u_n + v_n) = \sum_{n=0}^{+\infty} u_n + \sum_{n=0}^{+\infty} v_n$$

il faudra prendre la peine de vérifier que ces séries sont effectivement convergentes.

■ Correspondance fondamentale entre suites et séries

Nous avons associé à une suite (u_n) la suite (S_n) des sommes partielles de la série $\sum u_n$ de terme général u_n . Réciproquement, si (S_n) est une suite *quelconque* on peut, en posant : $u_0 = S_0$ et $\forall n \in \mathbb{N}^*$, $u_n = S_n - S_{n-1}$, la faire apparaître comme la suite des sommes partielles d'une certaine série. Cette égalité permet de déduire des propriétés de la série $\sum (S_n - S_{n-1})$ certains résultats qui concernent la suite (S_n) , à commencer par la

PROPOSITION 2.1 — Si la série $\sum u_n$ converge, la suite (u_n) tend vers 0.

Attention. Ce critère que nous venons d'énoncer n'assure pas à lui seul la convergence de la série ; il existe en effet de nombreuses séries divergentes dont le terme général tend vers 0. Il suffit pour cela que la suite (S_n) diverge et que la suite $(S_n - S_{n-1})$ tende vers 0. C'est le cas par exemple lorsque $S_n = \ln n$.

Mais l'exemple le plus connu est sans conteste la série *harmonique* $\sum \frac{1}{n}$. Les méthodes pour prouver la divergence de cette série sont très nombreuses, et nous verrons plus loin (dans la section « comparaison à une

intégrale ») une méthode plus simple. Dans l'immédiat, nous allons raisonner par l'absurde en supposant la convergence de cette série. Dans ces conditions, la suite $S_{2n} - S_n$ converge vers 0. Mais

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

ce qui est contradictoire.

Égalité télescopique

Lorsqu'on remplace u_k par $S_k - S_{k-1}$ pour $k \geq 1$ dans la relation $S_n = \sum_{k=0}^n u_k$ on obtient :

$$\forall n \in \mathbb{N}, \quad S_n = S_0 + \sum_{k=1}^n (S_k - S_{k-1}).$$

Cette relation, lorsqu'elle est mise en évidence, permet le calcul de certaines sommes, comme par exemple dans l'exercice suivant.

Exercice 6

Prouver la convergence et calculer la somme $\sum_{n=1}^{+\infty} \frac{1}{n(n+1)}$.

2.2 Séries de nombres réels positifs

Considérons maintenant une suite (u_n) de nombres réels *positifs*. Pour tout $n \in \mathbb{N}$ on a : $S_{n+1} - S_n = u_{n+1} \geq 0$, donc la suite des sommes partielles (S_n) est *croissante*. En conséquence :

la série $\sum u_n$ est convergente si et seulement si la suite (S_n) des sommes partielles est majorée.

Nous allons tirer plusieurs conséquences de cette constatation :

THÉORÈME 2.2 (comparaison) — Soient (u_n) et (v_n) deux suites de nombres réels positifs telles que : $\forall n \in \mathbb{N}$, $0 \leq u_n \leq v_n$. Alors la convergence de la série $\sum v_n$ entraîne celle de $\sum u_n$, et la divergence de $\sum u_n$ celle de $\sum v_n$.

Remarque. On peut remplacer l'hypothèse : $u_n \leq v_n$ par l'hypothèse : $u_n = O(v_n)$. En effet, cette nouvelle hypothèse implique l'existence d'un réel $B > 0$ tel que $u_n \leq Bv_n$, et il suffit alors d'appliquer le théorème aux séries $\sum u_n$ et $\sum Bv_n$. On peut donc énoncer le :

COROLLAIRE — Soient (u_n) et (v_n) deux suites de nombres réels positifs telles que : $u_n = O(v_n)$. Alors la convergence de la série $\sum v_n$ entraîne celle de $\sum u_n$, et la divergence de $\sum u_n$ celle de $\sum v_n$.

Exemples. La série $\sum \frac{n}{3^n}$ converge car $\frac{n}{3^n} = O\left(\frac{1}{2^n}\right)$ et $\sum \frac{1}{2^n}$ converge.

La série $\sum \frac{\ln n}{n}$ diverge car $\frac{1}{n} = O\left(\frac{\ln n}{n}\right)$ et $\sum \frac{1}{n}$ diverge.

COROLLAIRE — Deux séries $\sum u_n$ et $\sum v_n$ à terme général positif vérifiant : $u_n \sim v_n$ ont même nature⁵.

Attention. Ce résultat peut être mis en défaut lorsque les suites ne sont pas de signes constants. Cette erreur, très commune, a même été commise par Cauchy dans un article de 1823 consacré aux séries trigonométriques !

5. la *nature* d'une série est le fait pour elle d'être convergente ou divergente

Séries de référence

Ces deux derniers résultats nécessitent de posséder des séries de référence, c'est à dire des séries dont on connaît la nature et à qui on compare les autres séries. En ce qui nous concerne, nos séries de référence seront les séries géométriques (étudiées à la section 2.1) et les séries de Riemann (étudiées à la section 2.3).

Exercice 7

En admettant le résultat du corollaire du théorème 2.4, étudier la nature des séries de terme général :

$$u_n = 3 \ln(n^2 + 1) - 2 \ln(n^3 + 1), \quad v_n = \frac{a^n}{1 + a^{2n}} \quad (a > 0), \quad w_n = \sqrt[n]{n} - 1, \quad x_n = \frac{\ln n}{n^\alpha} \quad (\alpha > 0).$$

Enfin, la comparaison logarithmique à une série géométrique fournit le :

THÉORÈME 2.3 (règle de d'Alembert) — Soit (u_n) une suite de nombres réels strictement positifs, telle que $\lim_{n \rightarrow +\infty} \frac{u_{n+1}}{u_n} = a$. Alors : si $a < 1$, la série $\sum u_n$ converge ; si $a > 1$ elle diverge.

Attention. On ne peut rien conclure lorsque $a = 1$.

Cette règle s'avère particulièrement efficace dans le cas où le quotient $\frac{u_{n+1}}{u_n}$ présente des simplifications notables, comme on pourra l'observer dans l'exercice suivant.

Exercice 8

Déterminer la nature de la série de terme général $u_n = \frac{2 \times 4 \times 6 \times \cdots \times (2n)}{n^n}$.

2.3 Comparaison à une intégrale

Cette section concerne les séries de la forme $\sum f(n)$, où $f : [0, +\infty[\rightarrow \mathbb{R}_+$ est une fonction positive, continue par morceaux, et décroissante.

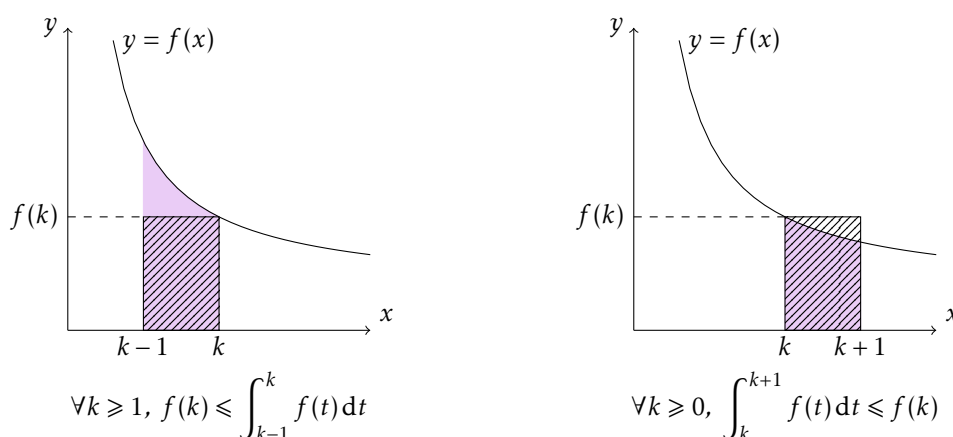


FIGURE 3 – Minoration et majoration de l'intégrale d'une fonction décroissante.

Observons les deux graphes représentés figure 3. Dans les deux cas, on compare l'aire hachurée, égale à $f(k)$, avec l'aire colorée, qui se calcule par l'intermédiaire d'une intégrale.

Pour tout $t \in [k-1, k]$, $f(k) \leq f(t)$ donc $f(k) = \int_{k-1}^k f(k) dt \leq \int_{k-1}^k f(t) dt$.

Pour tout $t \in [k, k+1]$, $f(t) \leq f(k)$ donc $\int_k^{k+1} f(t) dt \leq \int_k^{k+1} f(k) dt = f(k)$.

La première inégalité n'est valable que pour $k \geq 1$; en sommant on obtient :

$$\sum_{k=1}^n f(k) \leq \sum_{k=1}^n \int_{k-1}^k f(t) dt \quad \text{et donc} \quad \boxed{\sum_{k=0}^n f(k) \leq f(0) + \int_0^n f(t) dt} \quad (1)$$

En revanche, la seconde égalité est valable pour $k \geq 0$; en sommant on obtient :

$$\sum_{k=0}^n \int_k^{k+1} f(t) dt \leq \sum_{k=0}^n f(k) \quad \text{et donc} \quad \boxed{\int_0^{n+1} f(t) dt \leq \sum_{k=0}^n f(k)} \quad (2)$$

On en déduit :

THÉORÈME 2.4 — La série $\sum f(n)$ converge si et seulement si la suite $\left(\int_0^n f(t) dt\right)$ converge.

Remarque. Plus tard dans l'année nous dirons que l'intégrale $\int_0^{+\infty} f(t) dt$ converge.

Séries de Riemann

L'application de ce théorème aux fonctions $x \mapsto \frac{1}{x^\alpha}$ donne nos principales séries de référence :

COROLLAIRE (Séries de Riemann) — $\boxed{\text{La série } \sum \frac{1}{n^\alpha} \text{ converge si et seulement si } \alpha > 1.}$

2.4 Équivalent des sommes partielles et des restes

Lorsque une série numérique converge, la suite des restes tend vers 0 (c'est un infiniment petit) ; il est donc légitime de chercher un équivalent simple du reste.

Lorsque une série à terme général positif diverge, la suite de ses sommes partielles diverge vers $+\infty$ (c'est un infiniment grand) ; il est donc légitime de chercher un équivalent simple de la somme partielle.

La technique de comparaison à une intégrale permet dans certains cas de répondre à ces questions.

Exemple. Nous avons : $\forall k \geq 1, \int_k^{k+1} \frac{dt}{t^2} \leq \frac{1}{k^2} \leq \int_{k-1}^k \frac{dt}{t^2}$, donc en sommant :

$$\int_{n+1}^{N+1} \frac{dt}{t^2} \leq \sum_{k=n+1}^N \frac{1}{k^2} \leq \int_n^N \frac{dt}{t^2} \quad \text{ce qui donne :} \quad \frac{1}{n+1} - \frac{1}{N+1} \leq \sum_{k=n+1}^N \frac{1}{k^2} \leq \frac{1}{n} - \frac{1}{N}.$$

En faisant tendre N vers $+\infty$ on obtient : $\frac{1}{n+1} \leq \sum_{k=n+1}^{+\infty} \frac{1}{k^2} \leq \frac{1}{n}$ et donc : $\sum_{k=n+1}^{+\infty} \frac{1}{k^2} \sim \frac{1}{n}$. Nous avons obtenu un équivalent du reste de la série convergente $\sum \frac{1}{n^2}$.

Exercice 9

Appliquer de nouveau la technique de comparaison à une intégrale, mais cette fois-ci pour encadrer une somme partielle d'une série divergente : $S_n = \sum_{k=1}^n \frac{1}{k}$.

En déduire un équivalent de cette somme lorsque n tend vers $+\infty$.

Remarque. Il est possible d'obtenir une formule plus précise que dans l'exercice précédent, en prouvant

l'existence d'une constante $\gamma \approx 0,577 \dots$, appelée *constante d'Euler* vérifiant : $\boxed{\sum_{k=1}^n \frac{1}{k} = \ln n + \gamma + o(1).}$

■ Un complément (hors programme)

Le résultat qui suit donne une méthode alternative pour déterminer un équivalent du reste d'une série convergente, ou de la somme partielle d'une série convergente, toujours dans le cadre d'une série à terme général positif. Ce résultat étant hors-programme, il doit être démontré avant d'être utilisé.

PROPOSITION 2.5 — Soient (u_n) et (v_n) deux suites à terme général positif telles que $u_n \sim v_n$. Alors :

- si $\sum u_n$ converge il en est de même de $\sum v_n$, et $\sum_{k=n+1}^{+\infty} v_k \underset{n \rightarrow +\infty}{\sim} \sum_{k=n+1}^{+\infty} u_k$;
- si $\sum u_n$ diverge il en est de même de $\sum v_n$, et $\sum_{k=0}^n v_k \underset{n \rightarrow +\infty}{\sim} \sum_{k=0}^n u_k$.

2.5 Séries alternées

Nous allons maintenant étudier un cas particulier de séries à terme général réel, mais qui ne sont plus de signe constant. Nous adoptons la définition suivante :

DÉFINITION. — Une série alternée est une série de la forme $\sum (-1)^n a_n$, la suite (a_n) étant formée de nombres réels positifs.

L'intérêt de ces séries est que l'on dispose d'un critère très simple assurant leur convergence ; il s'agit du résultat suivant :

THÉORÈME 2.6 (Critère spécial des séries alternées) — Si (a_n) est une suite décroissante qui tend vers 0, la série alternée $\sum (-1)^n a_n$ est une série convergente.

Exemple. La série $\sum \frac{(-1)^{n-1}}{n}$ vérifie les conditions du critère spécial des séries alternées, donc converge. Nous pouvons illustrer cette convergence en observant le comportement des sommes partielles présenté figure 4.

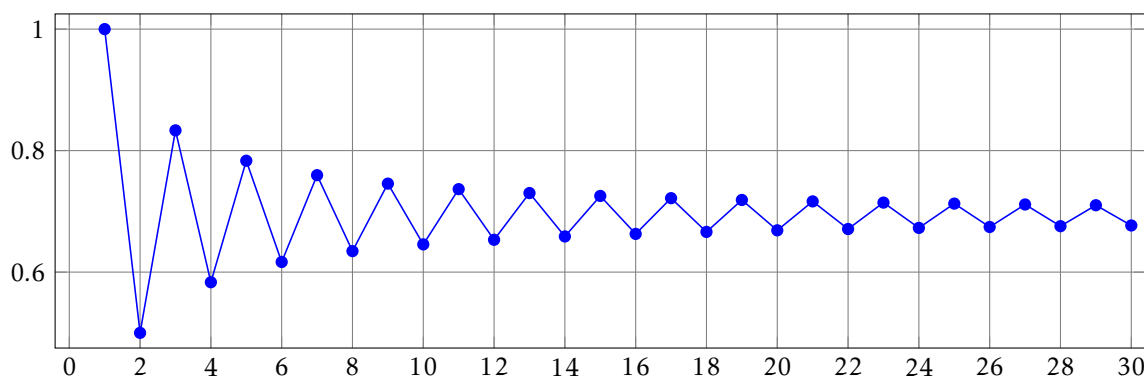


FIGURE 4 – La suite des sommes partielles de la série de terme général $\frac{(-1)^{n-1}}{n}$.

Cette figure indique clairement comment procéder pour prouver ce théorème : montrer que les suites (S_{2n}) et (S_{2n+1}) sont adjacentes. Nous en déduisons aussi le résultat suivant :

COROLLAIRE — Si la série $\sum (-1)^n a_n$ vérifie les hypothèses du critère spécial des séries alternées, le reste $r_n = \sum_{k=n+1}^{+\infty} (-1)^k a_k$ vérifie : $|r_n| \leq a_{n+1}$. De plus, r_n est du signe de son premier terme, à savoir du signe de $(-1)^{n+1} a_{n+1}$.

Remarque. Dans le cas particulier de la série $\sum \frac{(-1)^{n-1}}{n}$ il est possible, en séparant les termes pairs des termes impairs, d'en calculer la somme.

En effet on a : $S_{2p} = \sum_{n=1}^{2p} \frac{(-1)^{n-1}}{n} = \sum_{k=1}^p \frac{1}{2k-1} - \sum_{k=1}^p \frac{1}{2k} = I_p - J_p$, avec $I_p = \sum_{k=1}^p \frac{1}{2k-1}$ et $J_p = \sum_{k=1}^p \frac{1}{2k}$.

Par ailleurs, $I_p + J_p = \sum_{n=1}^{2p} \frac{1}{n} = \ln(2p) + \gamma + o(1)$, et $J_p = \frac{1}{2} \sum_{k=1}^p \frac{1}{k} = \frac{1}{2} \ln p + \frac{\gamma}{2} + o(1)$ donc :

$$I_p - J_p = (I_p + J_p) - 2J_p = \ln(2p) + \gamma - \ln p - \gamma + o(1) = \ln 2 + o(1).$$

En passant à la limite, $\sum_{n=1}^{+\infty} \frac{(-1)^{n-1}}{n} = \ln 2$.

Exercice 10

Soit $\alpha > 0$. Pour tout $n \geq 1$ on pose $u_n = \ln\left(1 + \frac{(-1)^n}{n^\alpha}\right)$. Effectuer un développement asymptotique à deux termes de u_n , puis expliquer comment l'utiliser pour prouver la convergence de la série $\sum u_n$.

2.6 Séries absolument convergentes

Le critère spécial relatif aux séries alternées s'applique dans un cadre relativement étroit : il faut que le terme général soit réel, de signe alterné et décroissante en valeur absolue, même si l'exercice 10 a montré comment il pouvait être utilisé dans un cadre *un peu* plus général.

Pour prouver la convergence d'une série à terme général complexe, ou à terme général réel mais sans alternance de signe, il ne reste alors (dans le cadre de notre programme) qu'une seule possibilité : l'absolue convergence, qui repose sur le théorème suivant :

THÉORÈME 2.7 — Si la série de terme général positif $\sum |u_n|$ converge, il en est de même de la série $\sum u_n$. On dit alors que la série $\sum u_n$ est absolument convergente.

Remarque. Si la série $\sum u_n$ est absolument convergente, l'inégalité triangulaire se généralise en :

$$\left| \sum_{n=0}^{+\infty} u_n \right| \leq \sum_{n=0}^{+\infty} |u_n|.$$

Exemple. La fonction zêta de Riemann et la fonction êta de Dirichlet sont respectivement définies pour une variable complexe z par : $\zeta(z) = \sum_{n=1}^{+\infty} \frac{1}{n^z}$ et $\eta(z) = \sum_{n=1}^{+\infty} \frac{(-1)^{n-1}}{n^z}$.

Si $z = x + iy$ avec $(x, y) \in \mathbb{R}^2$, on a $\frac{1}{n^z} = \frac{e^{-iy \ln n}}{n^x}$ donc $\left| \frac{1}{n^z} \right| = \frac{1}{n^x}$; ainsi les fonction ζ et η sont (au moins) définies sur l'ensemble $\{z \in \mathbb{C} \mid \Re(z) > 1\}$.

Exercice 11

Démontrer que lorsque $\Re(z) > 1$, $\eta(z) = (1 - 2^{1-z})\zeta(z)$.

■ Semi-convergence

Lorsque $x \in \mathbb{R}$, le critère spécial des séries alternées permet de prouver que $\eta(x)$ est définie pour $x > 0$. Plus généralement, une technique hors-programme (la transformation d'Abel) permet de prouver que $\eta(z)$ est définie lorsque $\Re(z) > 0$.

Ainsi, lorsque $0 < x \leq 1$, la série $\sum \frac{(-1)^{n-1}}{n^x}$ est un exemple de série convergente qui n'est pas absolument convergente. On parle alors de série *semi-convergente*.

Comme exemple type de semi-convergence on pourra donc citer $\eta(1) = \sum_{n=1}^{+\infty} \frac{(-1)^{n-1}}{n}$, qui est une série convergente d'après le critère spécial, mais qui n'est pas absolument convergente car la série harmonique $\sum \frac{1}{n}$ diverge.

2.7 Produit de Cauchy de deux séries

DÉFINITION. — Le produit de Cauchy de deux séries $\sum u_n$ et $\sum v_n$ est la série $\sum w_n$ de terme général

$$w_n = \sum_{i+j=n} u_i v_j.$$

Remarque. L'expression de w_n doit être comprise ainsi : on réalise la somme de tous les termes de la forme $u_i v_j$ pour lesquels les entiers i et j vérifient la condition $i + j = n$.

Cette condition est équivalente aux conditions $i \in \llbracket 0, n \rrbracket$ et $j = n - i$ donc on peut aussi écrire $w_n = \sum_{i=0}^n u_i v_{n-i}$.

Si en revanche on observe que $j \in \llbracket 0, n \rrbracket$ et $i = n - j$ on écrira $w_n = \sum_{j=0}^n u_{n-j} v_j$.

Attention. Si la suite (u_n) n'est définie que pour $n \geq 1$, il faut adapter la définition : la suite w_n ne sera définie que pour $n \geq 1$ et la condition $i + j = n$ se traduira par $i \in \llbracket 1, n \rrbracket$ et $j = n - i$ ou par $j \in \llbracket 0, n - 1 \rrbracket$ et $i = n - j$:

$$\forall n \geq 1, \quad w_n = \sum_{i=1}^n u_i v_{n-i} = \sum_{j=0}^{n-1} u_{n-j} v_j.$$

De même, si (u_n) et (v_n) ne sont définies que pour $n \geq 1$ la suite (w_n) ne sera définie que pour $n \geq 2$ par

$$\forall n \geq 2, \quad w_n = \sum_{i=1}^{n-1} u_i v_{n-i} = \sum_{j=1}^{n-1} u_{n-j} v_j.$$

LEMME — Soient $\sum a_n$ et $\sum b_n$ deux séries à terme général positif ($a_n \geq 0$ et $b_n \geq 0$) et convergentes. Alors leur produit de Cauchy $\sum c_n$ converge, et :

$$\sum_{n=0}^{+\infty} c_n = \left(\sum_{n=0}^{+\infty} a_n \right) \left(\sum_{n=0}^{+\infty} b_n \right).$$

THÉORÈME 2.8 — Soient $\sum u_n$ et $\sum v_n$ deux séries absolument convergentes. Alors leur produit de Cauchy $\sum w_n$ converge absolument, et

$$\sum_{n=0}^{+\infty} w_n = \left(\sum_{n=0}^{+\infty} u_n \right) \left(\sum_{n=0}^{+\infty} v_n \right).$$

Exercice 12

Soit (u_n) une suite numérique telle que la série $\sum u_n$ converge absolument. En faisant apparaître un produit de Cauchy, montrer que la série de terme général $w_n = \frac{1}{2^n} \sum_{k=0}^n 2^k u_k$ converge absolument, puis exprimer sa somme.

2.8 La formule de Stirling : un équivalent de $n!$

La formule de Stirling fournit un équivalent de $n!$. Cette formule a été démontrée en deux temps : De Moivre prouve l'existence d'une constante C telle que $n! \sim C \sqrt{n} \left(\frac{n}{e} \right)^n$, puis Stirling trouve la valeur de cette constante, $C = \sqrt{2\pi}$.

Le résultat de Moivre

Exercice 13

Pour tout $n \geq 1$ on pose $u_n = \frac{n!}{n^n e^{-n} \sqrt{n}}$ et $v_n = \ln\left(\frac{u_{n+1}}{u_n}\right)$. Prouver la convergence de la série $\sum v_n$ et en déduire l'existence d'une constante $C > 0$ telle que $n! \sim C n^n e^{-n} \sqrt{n}$.

L'apport de Stirling

Il repose sur les intégrales de Wallis $I_n = \int_0^{\pi/2} (\sin t)^n dt$.

Exercice 14

- À l'aide d'une intégration par parties, prouver que pour tout $n \geq 2$, $nI_n = (n-1)I_{n-2}$ et en déduire que pour tout $p \in \mathbb{N}$, $I_{2p} = \frac{\pi}{2} \frac{(2p)!}{(2^p p!)^2}$ et $I_{2p+1} = \frac{(2^p p!)^2}{(2p+1)!}$.
- Justifier que pour tout $n \geq 2$, $I_n \leq I_{n-1} \leq I_{n-2}$, et en déduire que $\lim_{n \rightarrow \infty} \frac{I_{n-1}}{I_n} = 1$.
- Exprimer $\lim_{p \rightarrow \infty} \frac{I_{2p}}{I_{2p+1}}$ en fonction de la constante C et en déduire que $C = \sqrt{2\pi}$.

Les résultats combinés de ces deux exercices prouvent la *formule de Stirling* : $n! \sim \sqrt{2\pi n} \left(\frac{n}{e}\right)^n$. Cette formule est à connaître mais la preuve n'est pas exigible.

Chapitre V

Suites et séries de fonctions

C'est au XIX^e siècle que les mathématiciens comprennent peu à peu l'importance qu'il y a à distinguer différents types de convergence pour une suite de fonctions. Weierstrass, en 1840, est le premier à utiliser le terme de *convergence uniforme* et à comprendre qu'il s'agit d'une des idées fondamentales de l'analyse.

Dans ce chapitre, nous allons considérer une suite de fonctions $f_n : I \rightarrow \mathbb{K}$, $n \in \mathbb{N}$ et donner un sens à la notion de convergence simple puis de convergence uniforme de la suite de fonctions (f_n) . Ces notions seront ensuite étendues aux séries de fonctions.

1. Convergence d'une suite de fonctions

Dans toute cette partie, on considère un intervalle I de $\mathbb{R}$, et une suite de fonctions (f_n) définies sur I et à valeurs dans $\mathbb{K} = \mathbb{R}$ ou $\mathbb{C}$.

1.1 Convergence simple

La première façon d'étudier la convergence de la suite de fonctions (f_n) consiste, *pour chaque valeur* $x \in I$, à étudier la convergence de la suite numérique $(f_n(x))$. Ceci conduit à la définition :

DÉFINITION. — On dit que la suite de fonctions (f_n) converge simplement vers $f : I \rightarrow \mathbb{K}$ lorsque pour tout $x \in I$, la suite numérique $(f_n(x))_{n \in \mathbb{N}}$ converge vers $f(x)$.

Exemple. Considérons l'intervalle $[0, \pi]$ et la suite de fonctions (f_n) définie par $f_n : x \mapsto (\sin x)^n$.

- Si $x \neq \pi/2$ on a $\sin x \in [0, 1[$ donc $\lim_{n \rightarrow +\infty} (\sin x)^n = 0$;
- Si $x = \pi/2$ on a $\sin x = 1$ donc $\lim_{n \rightarrow +\infty} (\sin x)^n = 1$.

On en déduit que la suite (f_n) converge simplement sur l'intervalle $[0, \pi]$ vers la fonction $f : x \mapsto \begin{cases} 0 & \text{si } x \neq \pi/2 \\ 1 & \text{si } x = \pi/2 \end{cases}$.

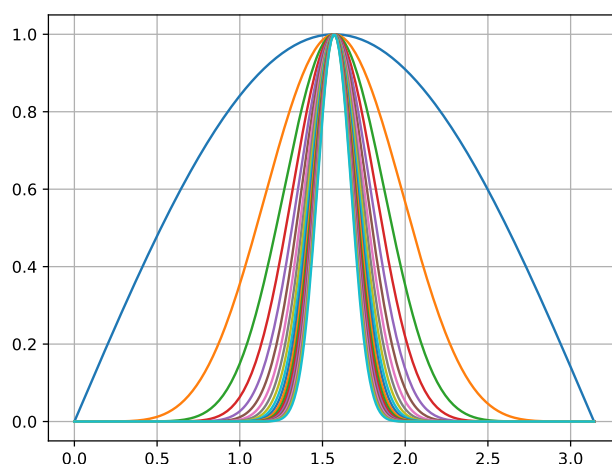


FIGURE 1 – Le graphe des premières fonctions f_n .

On peut déjà faire une première observation : bien que toutes les fonctions f_n soient continues sur $[0, \pi]$, leur limite simple f présente une discontinuité en $\pi/2$. C'est là un des défauts de la convergence simple sur lequel on reviendra : les propriétés locales (continuité, limite, ...) ne sont pas préservées par ce mode de convergence.

Exercice 1

Pour tout $n \in \mathbb{N}^*$, on définit la fonction $f_n : [0, +\infty[\rightarrow \mathbb{R}$ par : $f_n(x) = \begin{cases} \left(1 - \frac{x}{n}\right)^{-n} & \text{si } x < n \\ 0 & \text{si } x \geq n \end{cases}$.
Déterminer sa limite simple f .

Nous l'avons vu sur le premier exemple : la convergence simple ne préserve pas la continuité. Elle ne préserve pas non plus le passage à la limite : en général, sous la seule hypothèse de convergence simple, $\lim_{x \rightarrow a} \lim_{n \rightarrow +\infty} f_n(x) \neq \lim_{n \rightarrow +\infty} \lim_{x \rightarrow a} f_n(x)$, comme le montre l'exercice ci-dessus (avec $a = +\infty$).

En l'absence d'hypothèses supplémentaires, les seules propriétés préservées par la convergence simple sont celles qui ne font pas intervenir le comportement local des fonctions, comme par exemple :

PROPOSITION 1.1 — Soit (f_n) une suite de fonctions croissantes qui converge simplement sur l'intervalle I vers une fonction f . Alors f est aussi croissante sur I .

PROPOSITION 1.2 — Soit (f_n) une suite de fonctions positives qui converge simplement sur l'intervalle I vers une fonction f . Alors f est aussi positive sur I .

Pour obtenir des propriétés plus fortes, il faut adopter une définition de la convergence plus exigeante.

1.2 Convergence uniforme

DÉFINITION. — Soient I et J deux intervalles tels que $J \subset I$, et $f : I \rightarrow \mathbb{K}$ une fonction. Lorsque f est bornée sur J on appelle norme uniforme de f sur J la quantité

$$\|f\|_{\infty, J} = \sup\{|f(x)| \mid x \in J\}$$

Dans le cas particulier où $J = I$ (intervalle de définition de f) on se contentera de noter $\|f\|_{\infty}$ au lieu de $\|f\|_{\infty, I}$.

DÉFINITION. — On dit que la suite de fonctions (f_n) converge uniformément sur J vers une fonction $f : J \rightarrow \mathbb{K}$ lorsque les fonctions $f_n - f$ sont bornées (à partir d'un certain rang) sur J et

$$\lim_{n \rightarrow +\infty} \|f_n - f\|_{\infty, J} = 0$$

La quantité $\|f_n - f\|_{\infty, J}$ doit être interprétée comme la distance (uniforme) entre f_n et f sur l'intervalle J .

PROPOSITION 1.3 — Si (f_n) converge uniformément vers f , elle converge aussi simplement vers f .

Ce résultat est important car il nous indique la démarche à suivre pour étudier la convergence d'une suite de fonctions (f_n) :

- (i) on détermine limite simple f ;
- (ii) sur un intervalle J sur laquelle la fonction f est définie, on calcule (ou éventuellement on encadre) $\|f_n - f\|_{\infty, J}$ pour étudier la convergence uniforme.

Remarque. Si $J_1 \subset J_2$ on a $\|f_n - f\|_{\infty, J_1} \leq \|f_n - f\|_{\infty, J_2}$ donc la convergence uniforme sur J_2 entraîne la convergence uniforme sur J_1 . En particulier, s'il y a convergence uniforme sur I , il y a *a fortiori* convergence uniforme sur tout intervalle inclus dans I .

Exercice 2

Pour tout $n \in \mathbb{N}^*$, on considère la fonction $f_n : x \mapsto \frac{x\sqrt{n}}{1+nx^2}$, définie sur l'intervalle $[0, +\infty[$.

- Déterminer sa limite simple f sur $[0, +\infty[$.
- Former le tableau des variations de $f_n - f$ sur $[0, +\infty[$, et en déduire la valeur de $\|f_n - f\|_\infty$ sur cet intervalle. La convergence est-elle uniforme sur $[0, +\infty[$?
- Considérons maintenant un réel $\alpha > 0$ fixé. Former le tableau des variations de $f_n - f$ sur $[\alpha, +\infty[$ en distinguant les cas $n \leq \frac{1}{\alpha^2}$ et $n \geq \frac{1}{\alpha^2}$, et en déduire la valeur de $\|f_n - f\|_{\infty, [\alpha, +\infty[}$. La convergence est-elle uniforme sur $[\alpha, +\infty[$?

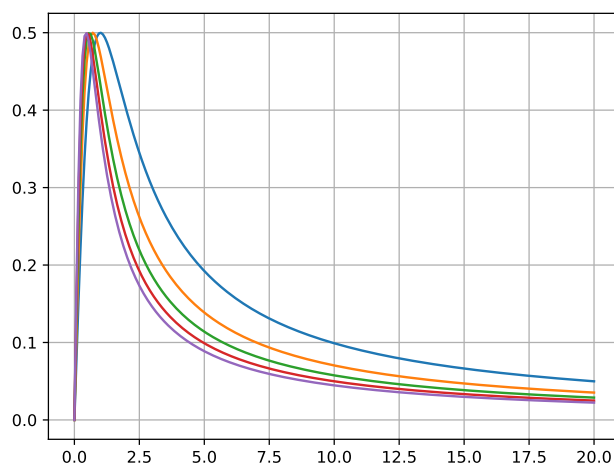


FIGURE 2 – Le graphe des premières fonctions f_n de l'exercice 2.

1.3 Régularité de la limite uniforme

Nous allons maintenant passer en revue les principales propriétés de la convergence uniforme, c'est à dire les propriétés des fonctions f_n qui sont transmises à leur limite uniforme f .

Rappel. Une fonction $f : I \rightarrow \mathbb{K}$ est *continue* en $a \in I$ lorsque pour tout $\epsilon > 0$ il existe un réel $\eta > 0$ tel que pour tout $x \in I$,

$$|x - a| \leq \eta \implies |f(x) - f(a)| \leq \epsilon.$$

THÉORÈME 1.4 — Soit (f_n) une suite de fonctions convergeant uniformément vers une fonction f sur un intervalle I , et $a \in I$ un point en lequel toutes les fonctions f_n sont continues. Alors f est aussi continue en a .

COROLLAIRE — Une limite uniforme de fonctions continues sur I est aussi continue sur I .

Remarque. La continuité étant une notion locale, il n'est pas forcément nécessaire de prouver la convergence uniforme sur I tout entier pour pouvoir justifier de la continuité de la fonction f .

Supposons par exemple $I = [0, +\infty[$. S'il n'y a pas convergence uniforme sur I mais seulement sur tout intervalle $[0, \alpha]$, la limite f sera néanmoins continue sur $[0, +\infty[$. En effet, si on considère un réel $a \geq 0$, il suffit de choisir un réel $\alpha > a$ et d'appliquer le théorème 1.4 sur l'intervalle $[0, \alpha]$: puisqu'il y a convergence uniforme sur $[0, \alpha]$, la fonction f est continue en a . Et puisque a est un réel quelconque de $[0, +\infty[$, f est bien continue sur cet intervalle.

Le même cas se produit lorsque $I =]0, +\infty[$ et lorsqu'il y a convergence uniforme sur tout intervalle de la forme $[\alpha, +\infty[$ avec $\alpha > 0$: tout réel $a > 0$ peut être englobé dans un intervalle de cette forme, et le théorème 1.4 appliqué sur l'intervalle $[\alpha, +\infty[$ permet alors de justifier la continuité de f en a .

Ce type de démarche sera appelée une *preuve par recouvrement* de la continuité de f sur I .

Exercice 3

Soit (f_n) une suite de fonctions continues qui converge uniformément vers f sur l'intervalle I , et (x_n) une suite d'éléments de I qui converge vers $\ell \in I$. Montrer que la suite $(f_n(x_n))_{n \in \mathbb{N}}$ converge vers $f(\ell)$.

Remarque. Ce théorème est un théorème d'interversion de limites : il montre qu'en cas de convergence uniforme et lorsque les fonctions f_n sont continues en a on a $\lim_{x \rightarrow a} \lim_{n \rightarrow +\infty} f_n(x) = \lim_{n \rightarrow +\infty} \lim_{x \rightarrow a} f_n(x)$. Nous admettrons la propriété plus générale suivante :

THÉORÈME 1.5 (théorème de la double limite) — Soit (f_n) une suite de fonctions qui converge uniformément vers une fonction f sur I , et a un point adhérent à I (qui peut éventuellement être égal à $\pm\infty$). On suppose que pour tout $n \in \mathbb{N}$, la fonction f_n possède une limite ℓ_n en a . Alors la suite (ℓ_n) admet elle-même une limite ℓ , et $\lim_{x \rightarrow a} f(x) = \ell$.

Autrement dit, ce théorème étend la relation $\lim_{x \rightarrow a} \lim_{n \rightarrow +\infty} f_n(x) = \lim_{n \rightarrow +\infty} \lim_{x \rightarrow a} f_n(x)$ dans le cas où a est adhérent à I , en garantissant l'existence des limites.

■ Intégration d'une suite de fonctions

THÉORÈME 1.6 — Soit (f_n) une suite de fonctions continues qui converge uniformément vers une fonction f sur un segment $[a, b]$. Alors la suite numérique $\left(\int_a^b f_n(t) dt\right)$ converge vers $\int_a^b f(t) dt$. Autrement dit :

$$\lim_{n \rightarrow +\infty} \int_a^b f_n(t) dt = \int_a^b \lim_{n \rightarrow +\infty} f_n(t) dt.$$

Exercice 4

Étudier la convergence simple sur $\left[0, \frac{\pi}{2}\right]$ de la suite (f_n) définie par $f_n(x) = n(\cos x)^n \sin x$, puis calculer $\int_0^{\frac{\pi}{2}} \lim_{n \rightarrow +\infty} f_n(x) dx$ et $\lim_{n \rightarrow +\infty} \int_0^{\frac{\pi}{2}} f_n(x) dx$. La convergence est-elle uniforme sur $\left[0, \frac{\pi}{2}\right]$?

■ Dérivation d'une suite de fonctions

Observons figure 3 deux fonctions « proches » pour la norme uniforme. On constate que ces deux fonctions délimitent également des aires algébriques proches, ce que traduit le théorème 1.6. En revanche, ces deux mêmes fonctions peuvent avoir des dérivées très éloignées pour la norme uniforme. Il ne faut donc pas s'étonner que le théorème de dérivation d'une suite de fonctions ait une hypothèse de convergence uniforme portant non pas sur la suite (f_n) mais sur la suite des dérivées (f'_n) .

THÉORÈME 1.7 — Soit (f_n) une suite de fonctions de classe $\mathcal{C}^1$ sur I , telle que :

- (i) (f_n) converge simplement vers une fonction f sur I ;
- (ii) (f'_n) converge uniformément vers une fonction g sur I .

Alors f est de classe $\mathcal{C}^1$ sur I , et $f' = g$.

Remarque. À l'instar de la continuité, la dérivabilité est une propriété *locale*, ce qui permet d'effectuer une preuve par recouvrement de la dérivabilité de la fonction f : pour prouver que f est de classe $\mathcal{C}^1$ sur I , il suffit de prouver que f est de classe $\mathcal{C}^1$ sur un ensemble d'intervalles recouvrant I .

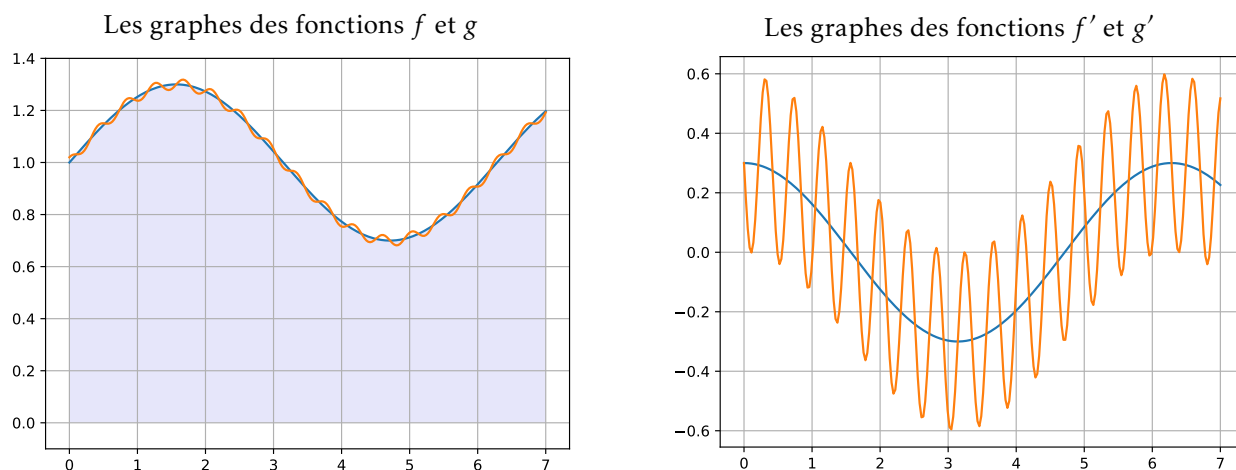


FIGURE 3 – Deux fonctions proches pour la norme uniforme, et leurs dérivées.

Extension aux fonctions de classe $\mathcal{C}^k$

On peut ajouter une conclusion supplémentaire au théorème 1.7 : non seulement f est de classe $\mathcal{C}^1$ sur I , mais en plus, la convergence de (f_n) vers f est uniforme sur tout segment.

Considérons alors une suite de fonctions (f_n) de classes $\mathcal{C}^2$ sur I vérifiant :

- (i) (f_n) converge simplement vers une fonction f sur I ;
- (ii) (f'_n) converge simplement vers une fonction g sur I ;
- (iii) (f''_n) converge uniformément vers une fonction h sur I .

Compte tenu du théorème 1.7 et de son complément, les propriétés (ii) et (iii) prouvent :

- (iv) la fonction g est de classe $\mathcal{C}^1$ sur I , $g' = h$ et la convergence de (f'_n) vers g est uniforme sur tout segment inclus dans I .

Mais alors, le théorème 1.7 associé aux propriétés (i) et (iv) prouve que f est de classe $\mathcal{C}^2$ sur I , que $f' = g$ et donc que $f'' = h$.

En généralisant on obtient :

PROPOSITION 1.8 — Soit (f_n) une suite de fonctions de classes $\mathcal{C}^k$ telle que :

- (i) (f_n) converge simplement vers une fonction f sur I ;
- (ii) pour tout $i \in \llbracket 1, k-1 \rrbracket$, $(f_n^{(i)})$ converge simplement vers une fonction g_i sur I ;
- (iii) $(f_n^{(k)})$ converge uniformément vers une fonction g_k sur I .

Alors f est de classe $\mathcal{C}^k$, et pour tout $i \in \llbracket 1, k \rrbracket$, $f^{(i)} = g_i$.

2. Convergence d'une série de fonctions

Nous allons maintenant adapter les définitions et résultats précédents au cas des séries de fonctions.

2.1 Convergence simple et absolue

Dans cette partie, on considère un intervalle I et une suite de fonctions (f_n) définies sur I .

DÉFINITION. — On dit que la série de fonctions $\sum f_n$ converge simplement lorsque pour tout $x \in I$, la série numérique $\sum f_n(x)$ converge, et qu'elle converge absolument lorsque pour tout $x \in I$, la série numérique $\sum |f_n(x)|$ converge.

Nous avons déjà vu dans le chapitre consacré aux séries numériques que la convergence absolue entraîne la convergence simple.

Dans le cas où la série $\sum f_n$ converge simplement, on définit une fonction $S : I \rightarrow \mathbb{K}$ en posant :

$$\forall x \in I, \quad S(x) = \sum_{n=0}^{+\infty} f_n(x).$$

2.2 Convergence uniforme et normale

DÉFINITION. — On dit que la série de fonction $\sum f_n$ converge uniformément vers une fonction $S : I \rightarrow \mathbb{K}$ lorsque la suite des sommes partielles converge uniformément vers S sur I , c'est à dire :

$$\lim_{n \rightarrow +\infty} \left\| S - \sum_{k=0}^n f_k \right\|_{\infty, I} = 0.$$

Nous savons déjà que la convergence uniforme entraîne la convergence simple. Si cette dernière est déjà acquise, nous pouvons définir la fonction *reste au rang n* en posant :

$$\forall x \in I, \quad R_n(x) = S(x) - \sum_{k=0}^n f_k(x) = \sum_{k=n+1}^{+\infty} f_k(x)$$

ce qui permet de retenir comme définition alternative le résultat suivant :

PROPOSITION 2.1 — La série de fonction $\sum f_n$ converge uniformément sur I lorsque :

- (i) la série de fonctions $\sum f_n$ converge simplement sur I ;
- (ii) $\lim_{n \rightarrow +\infty} \|R_n\|_{\infty, I} = 0$.

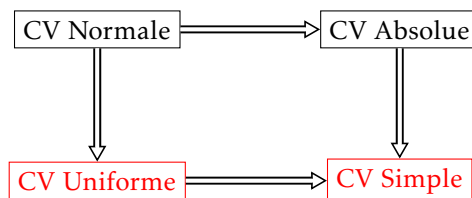
Nous allons maintenant introduire une notion spécifique aux séries de fonctions et qui va constituer un cas particulier de convergence uniforme, en adoptant la définition suivante :

DÉFINITION. — On dit que la série de fonctions $\sum f_n$ converge normalement (= au sens de la norme) sur I lorsque la série numérique $\sum \|f_n\|_{\infty, I}$ converge.

Le fait majeur, qui donne tout son intérêt à la notion de convergence normale, est le

THÉORÈME 2.2 — Toute série normalement convergente est uniformément convergente.

Remarque. On peut résumer les 4 modes de convergence possible d'une série de fonctions par le schéma suivant :



Les modes qui nous sont utiles sont :

- la convergence simple pour définir la fonction S ;
- la convergence uniforme pour utiliser les théorèmes de régularité qui assureront la continuité, la dérivabilité, etc., de la fonction S .

De ce fait, l'étude d'une série de fonctions suivra peu ou prou le modèle suivant :

- (1) établir la convergence simple de $\sum f_n$ sur I ;
- (2) si la convergence est absolue, calculer $\|f_n\|_\infty$ pour prouver la convergence normale sur I ou sur tout segment inclus dans I ;
- (3) si la convergence n'est pas absolue, majorer $\|R_n\|_\infty$ en vue de prouver la convergence uniforme sur I ou sur tout segment inclus dans I .

On notera que le point (3) intervient essentiellement dans le cadre du critère spécial relatif aux séries alternées, critère qui donne une majoration du reste (revoir le cours sur les séries numériques).

Exemple. La fonction zêta de Riemann est définie par : $\zeta(x) = \sum_{n=1}^{+\infty} \frac{1}{n^x}$.

Pour tout $n \geq 1$ et $x > 0$, posons $f_n(x) = \frac{1}{n^x}$.

La technique de comparaison à une intégrale permet de prouver la convergence simple de la série $\sum f_n$ sur $]1, +\infty[$; ainsi la fonction ζ est définie sur $]1, +\infty[$.

Le tableau des variations de la fonction f_n sur $]1, +\infty[$ est le suivant :

x	1	α	$+\infty$
$f_n(x)$	$\frac{1}{n}$		0

Sur l'intervalle $]1, +\infty[$, nous avons $\|f_n\|_\infty = \frac{1}{n}$, mais comme la série $\sum \frac{1}{n}$ diverge, la convergence n'y est pas normale. En revanche, sur l'intervalle $[\alpha, +\infty[$ (avec un réel arbitraire $\alpha > 1$) nous avons $\|f_n\|_{\infty, [\alpha, +\infty[} = \frac{1}{n^\alpha}$, et puisque la série $\sum \frac{1}{n^\alpha}$ converge, la convergence y est normale, donc uniforme. *A fortiori*, la convergence est uniforme sur tout segment inclus dans $]1, +\infty[$.

Exemple. La fonction êta de Dirichlet est définie par : $\eta(x) = \sum_{n=1}^{+\infty} \frac{(-1)^{n-1}}{n^x}$.

Le critère spécial relatif aux séries alternées prouve la convergence simple de la série $\sum (-1)^{n-1} f_n$ sur $]0, +\infty[$; ainsi la fonction η est définie sur $]0, +\infty[$.

De plus, toujours d'après le critère spécial, $|R_n(x)| = \left| \sum_{k=n+1}^{+\infty} (-1)^{k-1} f_k(x) \right| \leq |f_{n+1}(x)| = \frac{1}{n^x}$.

Sur l'intervalle $]0, +\infty[$ on en déduit que $\|R_n\|_\infty \leq 1$, ce qui est insuffisant pour prouver la convergence uniforme.

En revanche, pour tout $\beta > 0$ nous avons sur l'intervalle $[\beta, +\infty[$: $\|R_n\|_{\infty, [\beta, +\infty[} \leq \frac{1}{n^\beta}$ et $\lim_{n \rightarrow +\infty} \frac{1}{n^\beta} = 0$, ce qui prouve la convergence uniforme sur $[\beta, +\infty[$. *A fortiori*, la convergence est uniforme sur tout segment inclus dans $]0, +\infty[$.

Exercice 5

Pour tout $n \in \mathbb{N}$ on définit les fonctions $f_n : x \mapsto \frac{1}{1+n^2x}$ et $g_n : x \mapsto \frac{(-1)^n}{1+nx}$. Montrer que les séries $\sum f_n$ et $\sum g_n$ convergent simplement sur $]0, +\infty[$. Sur quels intervalles peut-on établir la convergence uniforme?

2.3 Régularité de la somme d'une série de fonctions

Nous allons maintenant traduire les théorèmes relatifs aux suites de fonctions dans les cas particulier des séries, en les appliquant à la suite des sommes partielles :

THÉORÈME 2.3 — Si la série de fonctions $\sum f_n$ converge uniformément sur I et si pour tout $n \in \mathbb{N}$, la fonction f_n est continue en un point a de I , alors la somme S est continue en a . En particulier, si toutes les fonctions f_n sont continues sur I , il en est de même de leur somme.

Remarque. À l'instar des suites de fonctions, il est fréquent d'avoir à procéder par recouvrement pour prouver la continuité d'une fonction définie par une série.

Exemple. Compte tenu des deux exemples traités dans la section précédente, on peut affirmer que la fonction zêta de Riemann est continue sur tout intervalle $[\alpha, +\infty[$ avec $\alpha > 1$ donc par recouvrement sur $]1, +\infty[$. De même, la fonction éta de Dirichlet est continue sur tout intervalle $[\alpha, +\infty[$ avec $\alpha > 0$ donc par recouvrement sur $]0, +\infty[$.

THÉORÈME 2.4 (théorème de la double limite) — Soit $\sum f_n$ une série de fonctions qui converge uniformément sur l'intervalle I , et a un point adhérent à I (qui peut éventuellement prendre la valeur $\pm\infty$). On suppose que pour tout $n \in \mathbb{N}$ la fonction f_n possède une limite ℓ_n en a . Alors la série numérique $\sum \ell_n$ converge, et $\lim_{x \rightarrow a} S(x) = \sum_{n=0}^{+\infty} \ell_n$.

Autrement dit, sous réserve de convergence uniforme sur I , $\lim_{x \rightarrow a} \sum_{n=0}^{+\infty} f_n(x) = \sum_{n=0}^{+\infty} \lim_{x \rightarrow a} f_n(x)$.

Exemple. La convergence uniforme sur l'intervalle $[2, +\infty[$ permet grâce à ce théorème de calculer la limite en $+\infty$ de la fonction zêta : $\lim_{x \rightarrow +\infty} \zeta(x) = \lim_{x \rightarrow +\infty} \sum_{n=1}^{+\infty} \frac{1}{n^x} = \sum_{n=1}^{+\infty} \lim_{x \rightarrow +\infty} \frac{1}{n^x} = 1$ car $\lim_{x \rightarrow +\infty} \frac{1}{n^x} = \begin{cases} 1 & \text{si } n = 1 \\ 0 & \text{si } n \geq 2 \end{cases}$.

Il permet aussi de prouver que la convergence de cette même série ne peut être uniforme sur un intervalle de la forme $]1, \alpha]$ avec $\alpha > 1$ puisque pour tout $n \in \mathbb{N}$, $\lim_{x \rightarrow 1} \frac{1}{n^x} = \frac{1}{n}$. S'il y avait convergence uniforme, la série $\sum \frac{1}{n}$ convergerait, ce qui n'est pas.

■ Intégration de la somme d'une série de fonctions

Le théorème d'intégration d'une suite de fonctions appliqué à la suite des sommes partielles d'une série de fonctions fournit le résultat suivant :

THÉORÈME 2.5 — Soit (f_n) une suite d'applications continues sur $[a, b]$, telle que la série $\sum f_n$ converge uniformément sur ce segment. Alors la série $\sum \int_a^b f_n(t) dt$ converge, et :

$$\sum_{n=0}^{+\infty} \int_a^b f_n(t) dt = \int_a^b \left(\sum_{n=0}^{+\infty} f_n(t) \right) dt.$$

Exercice 6

- Montrer que la fonction $S : x \mapsto \sum_{n=1}^{+\infty} n e^{-nx}$ est définie et continue sur $]0, +\infty[$.
- Calculer $\int_1^x S(t) dt$ pour tout $x > 0$ et en déduire une expression de $S(x)$ sans symbole de sommation.

■ Dérivation de la somme d'une série de fonctions

THÉORÈME 2.6 — Soit (f_n) une suite de fonctions de classe $\mathcal{C}^1$ sur un intervalle I , à valeurs réelles ou complexes. On suppose que la série $\sum f_n$ converge simplement sur I , et que la série $\sum f'_n$ converge uniformément sur I . Alors la fonction $x \mapsto \sum_{n=0}^{+\infty} f_n(x)$ est de classe $\mathcal{C}^1$ sur I , et sa dérivée est la fonction $x \mapsto \sum_{n=0}^{+\infty} f'_n(x)$.

Remarque. Comme pour la continuité, il est fréquent de devoir procéder par recouvrement.

Exemple. Pour montrer que la fonction ζ de Riemann est de classe $\mathcal{C}^1$ sur son intervalle de définition $]1, +\infty[$, nous devons considérer la série des dérivées $\sum f'_n(x) = \sum -\frac{\ln n}{n^x}$.

Sur tout intervalle $[\alpha, +\infty[$ nous avons $\|f'_n\|_{\infty, [\alpha, +\infty[} = \frac{\ln n}{n^\alpha}$. Si β désigne un réel vérifiant : $1 < \beta < \alpha$, nous avons $\|f'_n\|_{\infty, [\alpha, +\infty[} = o\left(\frac{1}{n^\beta}\right)$, donc la série $\sum \|f'_n\|_{\infty, [\alpha, +\infty[}$ converge.

La convergence de la série des dérivées $\sum f'_n$ est normale, donc uniforme, sur l'intervalle $[\alpha, +\infty[$; on peut donc affirmer que la fonction ζ est de classe $\mathcal{C}^1$ sur cet tout intervalle de la forme $[\alpha, +\infty[$, puis par recouvrement sur $]1, +\infty[$, et que :

$$\forall x > 1, \quad \zeta'(x) = - \sum_{n=1}^{+\infty} \frac{\ln n}{n^x}.$$

Extension aux fonctions de classe $\mathcal{C}^k$

Enfin, à l'instar des suites de fonctions, on établit par récurrence un résultat permettant de prouver directement que la somme d'une série de fonctions est de classe $\mathcal{C}^k$, $k \geq 2$:

PROPOSITION 2.7 — soit (f_n) une suite de fonctions de classe $\mathcal{C}^k$ sur I , telles que :

- (i) $\sum f_n$ converge simplement sur I ;
- (ii) pour tout $i \in \llbracket 1, k-1 \rrbracket$, $\sum f_n^{(i)}$ converge simplement sur I ;
- (iii) $\sum f_n^{(k)}$ converge uniformément sur I .

Alors la fonction $S : x \mapsto \sum_{n=0}^{+\infty} f_n(x)$ est de classe $\mathcal{C}^k$ sur I , et pour tout $i \in \llbracket 1, k \rrbracket$, $S^{(i)}(x) = \sum_{n=0}^{+\infty} f_n^{(i)}(x)$.

Exercice 7

On considère la fonction $S : x \mapsto \sum_{n=0}^{+\infty} \frac{e^{-nx}}{n^2 + 1}$. Montrer que S est continue sur $[0, +\infty[$ et de classe $\mathcal{C}^2$ sur $]0, +\infty[$, puis établir une équation différentielle d'ordre 2 vérifiée par S sur l'intervalle $]0, +\infty[$.

Chapitre VI

Séries entières

Les séries entières sont des séries numériques de la forme $\sum a_n z^n$, où (a_n) est une suite réelle ou complexe, et z un élément de $\mathbb{R}$ ou $\mathbb{C}$ (la série est dite *entière* du fait qu'elle ne fait intervenir que des puissances entières). Ces séries possèdent des propriétés de convergence remarquables, que nous allons étudier dans la première partie. Dans un second temps, nous étudierons les propriétés de la fonction d'une *variable réelle* :

$$x \mapsto \sum_{n=0}^{+\infty} a_n x^n.$$

L'extension de ces propriétés au cas d'une variable complexe constitue la théorie des fonctions *analytiques*, le pilier central de l'analyse complexe.

1. Rayon de convergence

1.1 Définition d'une série entière

DÉFINITION. — Étant donnée une suite (a_n) de nombres complexes, on appelle série entière la série de fonctions $\sum a_n z^n$ de la variable complexe z . Son domaine de convergence est l'ensemble des nombres $z \in \mathbb{C}$ pour lesquels cette série converge.

Nous allons commencer par étudier le domaine de convergence d'une série entière dans un cas simple, en faisant deux hypothèses supplémentaires :

- (i) il existe un rang N à partir duquel $a_n \neq 0$;
- (ii) la suite $\left| \frac{a_{n+1}}{a_n} \right|$ converge vers une limite $\ell > 0$.

L'objectif de ces deux hypothèses est de permettre l'application du critère de d'Alembert :

$$\text{si } u_n = |a_n z^n| \text{ alors } \frac{u_{n+1}}{u_n} = |z| \left| \frac{a_{n+1}}{a_n} \right| \text{ donc } \lim \frac{u_{n+1}}{u_n} = \ell |z|.$$

Ainsi :

- si $|z| < \frac{1}{\ell}$, la série positive $\sum u_n$ converge donc la série $\sum a_n z^n$ converge absolument ;
- si $|z| > \frac{1}{\ell}$, la suite positive (u_n) diverge vers $+\infty$ donc la série $\sum a_n z^n$ diverge grossièrement.

En posant $R = \frac{1}{\ell}$ nous avons mis en évidence dans le plan complexe l'existence d'un disque $\mathcal{D}$ de centre 0 et de rayon R tel que :

- lorsque z est à l'intérieur du disque $\mathcal{D}$, la convergence de la série entière $\sum a_n z^n$ est absolue ;
- lorsque z est à l'extérieur du disque $\mathcal{D}$, la divergence de la série entière $\sum a_n z^n$ est grossière ;
- lorsque z est sur le bord de $\mathcal{D}$ (c'est-à-dire $|z| = R$) on ne peut pas conclure.

(illustration figure 1.)

Exemple. Considérons les trois séries entières $\sum z^n$, $\sum \frac{z^n}{n^2}$ et $\sum \frac{z^n}{n}$. Elles vérifient toutes trois les hypothèses (i) et (ii) avec $\ell = 1$ donc dans les trois cas le disque $\mathcal{D}$ est de rayon $R = 1$.

- Pour $|z| = 1$, la série $\sum z^n$ diverge (son terme général ne tend pas vers 0). Le domaine de convergence est le disque ouvert.

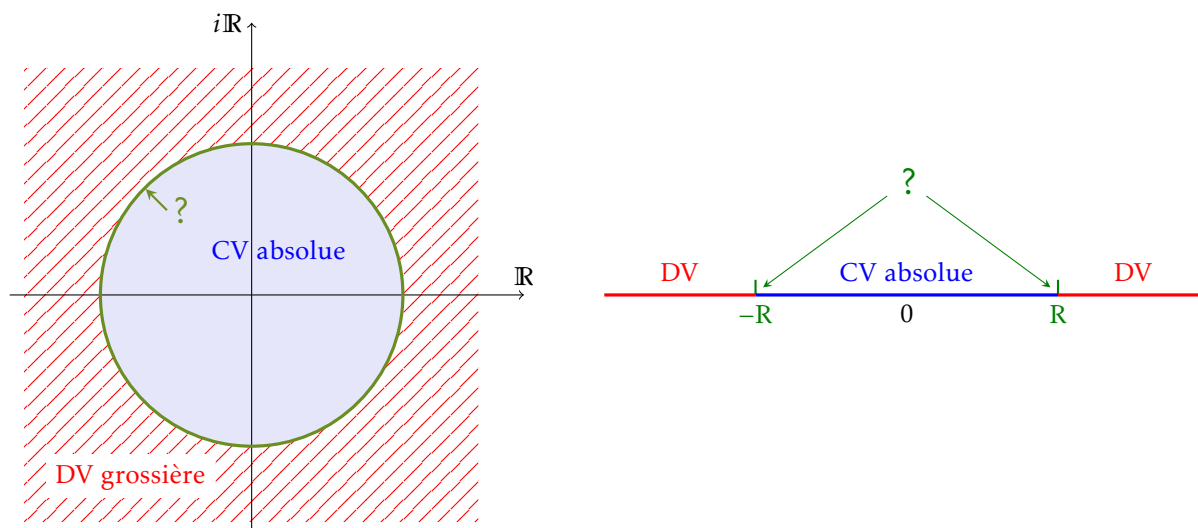


FIGURE 1 – Le disque de convergence et sa restriction au cas réel.

- Pour $|z| = 1$, la série $\sum \frac{z^n}{n^2}$ converge absolument car la série $\sum \frac{1}{n^2}$ converge. Son domaine de convergence est le disque fermé.
- Pour $z = 1$, la série $\sum \frac{z^n}{n}$ diverge, pour $z = -1$ la série $\sum \frac{z^n}{n}$ converge (par application du critère spécial), et pour $|z| = 1$ et $z \neq \pm 1$ on ne sait pas étudier la convergence de la série. Le domaine de convergence est « entre » le disque ouvert et le disque fermé.

Ces trois exemples montrent qu'il n'y a pas à espérer une règle générale concernant les valeurs de z pour lesquelles $|z| = R$.

Pour résumer, qu'avons-nous observé ?

lorsqu'une série entière vérifie les hypothèses (i) et (ii), il existe un disque $\mathcal{D}$ tel que la série converge absolument à l'intérieur de ce disque et diverge grossièrement à l'extérieur de ce disque.

Nous allons maintenant montrer que nous pouvons nous affranchir des hypothèses (i) et (ii), et que cette propriété est une propriété générale des séries entières.

1.2 Le lemme d'Abel

La généralisation du résultat que nous cherchons à obtenir repose sur le lemme suivant :

LEMME (Abel) — Soit $z_0 \in \mathbb{C} \setminus \{0\}$ tel que la suite $(a_n z_0^n)$ soit bornée. Alors pour tout $z \in \mathbb{C}$ tel que $|z| < |z_0|$ la série $\sum a_n z^n$ est absolument convergente.

Considérons alors l'ensemble $\mathcal{A} = \{\rho \in \mathbb{R}_+ \mid \text{la suite } (a_n \rho^n) \text{ est bornée}\}$ et $R = \sup \mathcal{A} \in \mathbb{R}_+ \cup \{+\infty\}$ (autrement dit, on convient que si $\mathcal{A}$ n'est pas majoré alors $R = +\infty$). On dispose du :

THÉORÈME 1.1 — Soit $\sum a_n z^n$ une série entière, et $R = \sup \mathcal{A}$. Alors :

- Si $R = 0$, le domaine de convergence se réduit à $\{0\}$;
- si $R = +\infty$, le domaine de convergence est égal à $\mathbb{C}$ tout entier ;
- si $0 < R < +\infty$, on a :
 - lorsque $|z| < R$, la série $\sum a_n z^n$ converge absolument ;
 - lorsque $|z| > R$, la série $\sum a_n z^n$ diverge.

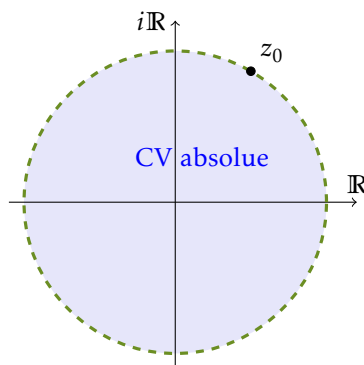


FIGURE 2 – Illustration du lemme d'Abel.

On appelle disque ouvert de convergence le disque de centre 0 et de rayon R . Le réel R est le rayon de convergence de la série entière.

Attention. Rappelons encore une fois qu'on ne peut rien dire *a priori* de la convergence sur le cercle de rayon R .

Exemple. Pour tout $\alpha \in \mathbb{R}$, le rayon de convergence de la série entière $\sum n^\alpha z^n$ est égal à 1.

Exercice 1

On considère deux réels α et β vérifiant : $0 < \alpha < \beta$, ainsi que la suite (a_n) définie par :

$$\forall p \in \mathbb{N}, \quad a_{2p} = \alpha^{2p} \quad \text{et} \quad a_{2p+1} = \beta^{2p+1}.$$

Déterminer l'ensemble $\mathcal{A}$ puis la valeur du rayon de convergence de la série entière $\sum a_n z^n$.

Remarque. Nous avons vu dans la section précédente que lorsque *certaines hypothèses sont réalisées* le critère de d'Alembert permet d'obtenir facilement la valeur du rayon de convergence. C'est le cas par exemple pour l'exercice suivant :

Exercice 2

À l'aide du critère de d'Alembert, déterminer le rayon de convergence de la série entière $\sum \frac{n^n}{n!} z^n$.

Cette démarche est séduisante car facile d'utilisation, mais il ne faut pas en faire une utilisation systématique, car les hypothèses nécessaires peuvent aisément être mises en défaut ! C'est par exemple le cas de l'exercice 1 :

la suite $\left| \frac{a_{n+1}}{a_n} \right|$ vaut $\beta \left(\frac{\beta}{\alpha} \right)^{2p}$ lorsque $n = 2p$, et $\alpha \left(\frac{\alpha}{\beta} \right)^{2p+1}$ lorsque $n = 2p+1$, donc :

$$\lim_{+\infty} \left| \frac{a_{2p+1}}{a_{2p}} \right| = +\infty \quad \text{et} \quad \lim_{+\infty} \left| \frac{a_{2p+2}}{a_{2p+1}} \right| = 0.$$

La suite $\left(\left| \frac{a_{n+1} z^{n+1}}{a_n z^n} \right| \right)$ ne possède donc pas de limite.

On s'autorisera néanmoins à utiliser directement le résultat suivant :

PROPOSITION 1.2 — Soit (a_n) une suite numérique ne s'annulant pas, telle que le quotient $\left| \frac{a_{n+1}}{a_n} \right|$ possède une limite $\ell \in \mathbb{R}_+ \cup \{+\infty\}$. Alors le rayon de convergence de $\sum a_n z^n$ est égal à $\frac{1}{\ell}$ (avec la convention $\frac{1}{0} = +\infty$ et $\frac{1}{+\infty} = 0$).

Exercice 3

Calculer à l'aide du critère de d'Alembert le rayon de convergence de $\sum \frac{(1+i)^n}{n} z^{3n}$.

■ Comparaison du rayon de convergence de deux séries entières

Notons pour finir que la comparaison de l'ordre de grandeur des termes généraux à une conséquence sur les rayons de convergence :

THÉORÈME 1.3 — Soit $\sum a_n z^n$ et $\sum b_n z^n$ deux séries entières de rayons de convergence respectifs R_a et R_b , et telles que $a_n = O(b_n)$. Alors $R_a \geq R_b$.

COROLLAIRE — Soit $\sum a_n z^n$ et $\sum b_n z^n$ deux séries entières de rayons de convergence respectifs R_a et R_b , et telles que $|a_n| \sim |b_n|$. Alors $R_a = R_b$.

Exercice 4

Soit (a_n) une suite vérifiant : $\lim a_n = 0$. Que dire du rayon de convergence de la série entière $\sum a_n z^n$?

1.3 Opérations algébriques sur les séries entières

■ Somme de deux séries entières

THÉORÈME 1.4 — Soient $\sum a_n z^n$ et $\sum b_n z^n$ deux séries entières, R_a et R_b leurs rayons de convergence respectifs. Alors le rayon de convergence de la série $\sum (a_n + b_n) z^n$ est supérieur ou égal à $\min(R_a, R_b)$, et :

$$\forall z \in \mathbb{C}, \quad |z| < \min(R_a, R_b) \implies \sum_{n=0}^{+\infty} a_n z^n + \sum_{n=0}^{+\infty} b_n z^n = \sum_{n=0}^{+\infty} (a_n + b_n) z^n.$$

Remarque. Supposons $R_a < R_b$ et considérons $z \in \mathbb{C}$ tel que $R_a < |z| < R_b$. Alors $\sum a_n z^n$ diverge et $\sum b_n z^n$ converge donc $\sum (a_n + b_n) z^n$ diverge. Ceci prouve que $R \leq R_a = \min(R_a, R_b)$, et donc que lorsque $R_a \neq R_b$, alors $R = \min(R_a, R_b)$.

■ Produit de deux séries entières

Considérons deux séries entières $\sum a_n z^n$ et $\sum b_n z^n$ de rayons de convergences respectifs R_a et R_b . Le produit de Cauchy de ces deux séries a pour terme général :

$$\sum_{p+q=n} (a_p x^p)(b_q x^q) = \left(\sum_{p+q=n} a_p b_q \right) x^n$$

Il s'agit du terme général d'une série entière. En appliquant le théorème prouvé dans le chapitre consacré aux séries numériques on obtient :

THÉORÈME 1.5 — Soient $\sum a_n z^n$ et $\sum b_n z^n$ deux séries entières de rayons de convergences respectifs R_a et R_b . Alors la série entière $\sum \left(\sum_{p+q=n} a_p b_q \right) z^n$ a un rayon de convergence supérieur ou égal au $\min(R_a, R_b)$. En outre, pour tout

$z \in \mathbb{C}$ tel que $|z| < \min(R_a, R_b)$ on a : $\left(\sum_{n=0}^{+\infty} a_n z^n \right) \left(\sum_{n=0}^{+\infty} b_n z^n \right) = \sum_{n=0}^{+\infty} \left(\sum_{p+q=n} a_p b_q \right) z^n$.

Exercice 5

On considère une série entière $\sum a_n z^n$ de rayon de convergence $R_a > 0$, et on définit la suite (b_n) en posant :

$\forall n \in \mathbb{N}, b_n = \sum_{k=0}^n a_k$. Que dire du rayon de convergence R_b de la série entière $\sum b_n z^n$?

Donner une relation liant les sommes $\sum_{n=0}^{+\infty} a_n z^n$ et $\sum_{n=0}^{+\infty} b_n z^n$ au voisinage de 0.

■ Dérivation formelle

Terminons avec un résultat qui nous sera utile dans la suite de ce chapitre (pour prouver le théorème 2.2) :

PROPOSITION 1.6 — La série entière $\sum na_n z^n$ a même rayon de convergence que la série $\sum a_n z^n$.

2. Séries entières d'une variable réelle

Nous allons désormais considérer une suite réelle ou complexe (a_n) telle que la série entière $\sum a_n z^n$ ait un rayon de convergence $R > 0$; ceci permet de définir une fonction numérique $S :]-R, R[\rightarrow \mathbb{C}$ (éventuellement définie en $-R$ et en R) à l'aide de l'égalité : $\forall x \in]-R, R[, S(x) = \sum_{n=0}^{+\infty} a_n x^n$.

L'intervalle $] -R, R[$ sera appelé *l'intervalle ouvert de convergence* (sachant que S peut en outre être définie en $\pm R$).

2.1 Convergence normale

THÉORÈME 2.1 — Soit $\sum a_n x^n$ une série entière de rayon de convergence R . Alors la convergence est normale sur tout segment $[-r, r]$ avec $r < R$.

Attention. Rappelons que ceci ne signifie pas qu'il y ait convergence normale (ou même uniforme) sur l'intervalle ouvert $] -R, R[$. Il suffit de considérer la série $\sum x^n$ pour s'en convaincre.

COROLLAIRE — La fonction $S : x \mapsto \sum_{n=0}^{+\infty} a_n x^n$ est continue sur l'intervalle ouvert de convergence.

Attention. Même si la fonction S est définie en $\pm R$, cela n'implique pas sa continuité en ces points.

2.2 Dérivation et intégration d'une série entière

THÉORÈME 2.2 — La fonction S est de classe $\mathcal{C}^1$ sur $] -R, R[$ et pour tout $x \in] -R, R[, S'(x) = \sum_{n \geq 1} n a_n x^{n-1}$.

Exercice 6

Calculer sur l'intervalle ouvert de convergence la somme $\sum_{n=1}^{+\infty} n x^n$.

COROLLAIRE — La fonction S est de classe $\mathcal{C}^\infty$ sur l'intervalle $] -R, R[$, et les dérivées successives s'obtiennent par dérivation terme à terme. De plus, pour tout $n \in \mathbb{N}$, $a_n = \frac{S^{(n)}(0)}{n!}$.

Cette dernière formule va avoir une conséquence importante :

PROPOSITION 2.3 — Deux séries entières dont les rayons de convergence sont non nuls ont des sommes égales si et seulement si tous leurs coefficients sont égaux.

En particulier, une série entière de rayon de convergence non nul aura une somme non identiquement nulle dès lors que l'un au moins de ses coefficients sera non nul.

En appliquant le théorème 2.2 à la série primitive, on obtient le résultat suivant :

PROPOSITION 2.4 — La fonction $T : x \mapsto \sum_{n=0}^{+\infty} \frac{a_n}{n+1} x^{n+1}$ définit l'unique primitive s'annulant en zéro de S sur l'intervalle $] -R, R[$.

2.3 Développement en série entière

Considérons une fonction $f :]-r, r[\rightarrow \mathbb{C}$ de classe $\mathcal{C}^\infty$ sur son intervalle de définition.

DÉFINITION. — On dit que f est développable en série entière au voisinage de 0 lorsqu'il existe une série entière $\sum a_n z^n$ de rayon de convergence $R \geq r$ telle que :

$$\forall x \in]-r, r[, \quad f(x) = \sum_{n=0}^{+\infty} a_n x^n.$$

D'après ce qui a été dit à la section précédente, si f est développable en série entière alors pour tout $n \in \mathbb{N}$, $a_n = \frac{f^{(n)}(0)}{n!}$, mais ceci n'est pas suffisant pour assurer l'existence de ce développement. En revanche, cette formule nous permet d'affirmer que si un tel développement existe, *ce dernier est unique* et coïncide avec la série de Taylor de f .

Pour prouver qu'une fonction est développable en série entière, différentes possibilités s'offrent à nous : une méthode fréquemment utilisée consiste à considérer un problème de Cauchy dont la fonction f est l'unique solution, et à déterminer les solutions de ce système qui peuvent s'écrire sous forme d'une somme de série entière. C'est ce que nous ferons pour déterminer le développement en série entière des fonctions exponentielle et $x \mapsto (1+x)^\alpha$. Une autre possibilité consiste à effectuer des manipulations à base de somme ou de produit de séries usuelles (par exemple pour obtenir le développement des fonctions trigonométriques) ou encore en utilisant les propriétés d'intégration et de dérivation des développements usuels ; nous procéderons par exemple ainsi pour obtenir les développements des fonctions $x \mapsto \arctan x$ et $x \mapsto \ln(1+x)$.

■ Développements usuels

La fonction exponentielle

Nous admettons que pour tout $\alpha \in \mathbb{C}$, la fonction $f : x \mapsto e^{\alpha x}$ est l'unique solution sur $\mathbb{R}$ du problème de Cauchy :

$$\begin{cases} y'(x) = \alpha y(x) \\ y(0) = 1 \end{cases}$$

Cherchons une solution de ce problème sous forme d'une série entière $y(x) = \sum_{n=0}^{+\infty} a_n x^n$ de rayon de convergence $R > 0$:

y est solution si et seulement si :

$$\begin{aligned} & \begin{cases} \forall x \in]-R, R[, \sum_{n=1}^{+\infty} n a_n x^{n-1} = \alpha \sum_{n=0}^{+\infty} a_n x^n \\ a_0 = 1 \end{cases} \iff \begin{cases} \forall x \in]-R, R[, \sum_{n=0}^{+\infty} (n+1) a_{n+1} x^n = \sum_{n=0}^{+\infty} \alpha a_n x^n \\ a_0 = 1 \end{cases} \\ & \iff \begin{cases} \forall n \in \mathbb{N}, (n+1) a_{n+1} = \alpha a_n \\ a_0 = 1 \end{cases} \iff \forall n \in \mathbb{N}, a_n = \frac{\alpha^n}{n!}. \end{aligned}$$

Le critère de d'Alembert nous permet de déterminer que cette série entière a un rayon de convergence infini, ce qui permet de conclure, en invoquant l'unicité de la solution d'un problème de Cauchy : $\forall x \in \mathbb{R}, e^{\alpha x} = \sum_{n=0}^{+\infty} \frac{\alpha^n}{n!} x^n$.

En prenant $\alpha = 1$ puis $\alpha = -1$ on obtient en particulier :

$$\boxed{\forall x \in \mathbb{R}, \quad e^x = \sum_{n=0}^{+\infty} \frac{1}{n!} x^n} \quad \text{et} \quad \boxed{\forall x \in \mathbb{R}, \quad e^{-x} = \sum_{n=0}^{+\infty} \frac{(-1)^n}{n!} x^n}.$$

Les fonctions trigonométriques

Sachant que pour tout $x \in \mathbb{R}$, $\operatorname{ch} x = \frac{e^x + e^{-x}}{2}$ et $\operatorname{sh} x = \frac{e^x - e^{-x}}{2}$ on en déduit :

$$\forall x \in \mathbb{R}, \quad \operatorname{ch} x = \sum_{p=0}^{+\infty} \frac{x^{2p}}{(2p)!} \quad \text{et} \quad \forall x \in \mathbb{R}, \quad \operatorname{sh} x = \sum_{p=0}^{+\infty} \frac{x^{2p+1}}{(2p+1)!}.$$

En prenant cette fois $\alpha = i$ puis $\alpha = -i$, on obtient de même :

$$\forall x \in \mathbb{R}, \quad \cos x = \sum_{p=0}^{+\infty} \frac{(-1)^p}{(2p)!} x^{2p} \quad \text{et} \quad \forall x \in \mathbb{R}, \quad \sin x = \sum_{p=0}^{+\infty} \frac{(-1)^p}{(2p+1)!} x^{2p+1}.$$

Les sommes géométriques

La formule déjà connue de sommation des sommes géométriques donne :

$$\forall x \in]-1, 1[, \quad \frac{1}{1-x} = \sum_{n=0}^{+\infty} x^n \quad \text{et} \quad \forall x \in]-1, 1[, \quad \frac{1}{1+x} = \sum_{n=0}^{+\infty} (-1)^n x^n$$

On en déduit en intégrant :

$$\forall x \in]-1, 1[, \quad \ln(1-x) = - \sum_{n=1}^{+\infty} \frac{x^n}{n} \quad \text{et} \quad \forall x \in]-1, 1[, \quad \ln(1+x) = \sum_{n=1}^{+\infty} \frac{(-1)^{n-1}}{n} x^n$$

et enfin :

$$\forall x \in]-1, 1[, \quad \arctan x = \sum_{p=0}^{+\infty} \frac{(-1)^p x^{2p+1}}{2p+1}.$$

La fonction $x \mapsto (1+x)^\alpha$

Enfin, pour obtenir le développement en série entière de la fonction $x \mapsto (1+x)^\alpha$ (avec $\alpha \in \mathbb{C}$) nous allons de nouveau admettre que cette fonction est l'unique solution sur l'intervalle $] -1, 1[$ du problème de Cauchy :

$$\begin{cases} (1+x)y'(x) = \alpha y(x) \\ y(0) = 1 \end{cases}$$

Exercice 7

Chercher l'unique série entière qui soit solution de ce problème de Cauchy sur l'intervalle $] -1, 1[$, calculer son rayon de convergence, et en déduire la formule ci-dessous.

$$\forall x \in]-1, 1[, \quad (1+x)^\alpha = \sum_{n=0}^{+\infty} \binom{\alpha}{n} x^n \quad \text{en ayant noté} \quad \binom{\alpha}{n} = \frac{\alpha(\alpha-1)\cdots(\alpha-n+1)}{n!}.$$

Attention. Malgré les apparences il ne s'agit pas ici d'un coefficient binomial puisqu'en général α n'est pas un nombre entier. Lorsque vous utilisez cette notation, il faut prendre garde à ne pas appliquer la fonction factorielle à des arguments non entiers.

Exercice 8

À l'aide de cette formule, obtenir le développement en série entière de la fonction $x \mapsto \frac{1}{\sqrt{1+x}}$ sur l'intervalle $] -1, 1[$.

Remarque. Les formules que nous venons d'établir ne vous sont pas inconnues : elles coïncident avec les développements limités usuels appris en première année. Ce n'est pas étonnant puisque série de Taylor et polynômes de Taylor partagent les mêmes coefficients. D'ailleurs, les développements limités peuvent être établis à partir du développement en série entière en utilisant le résultat ci-dessous.

PROPOSITION 2.5 — Soit $\sum a_n x^n$ une série entière de rayon de convergence $R > 0$, et S sa fonction somme. Alors S admet pour tout entier $n \in \mathbb{N}$ un développement limité d'ordre n en zéro donné par :

$$S(x) = \sum_{k=0}^n a_k x^k + o(x^n).$$

2.4 Séries géométrique et exponentielle d'une variable complexe

Le cours de première année a défini l'exponentielle d'un nombre complexe : si $z = x + iy \in \mathbb{C}$, alors $\exp(z)$ (ou e^z) = $e^x e^{iy} = e^x (\cos y + i \sin y)$. À la section précédente nous avons obtenu un développement de $e^{\alpha x}$ pour $\alpha \in \mathbb{C}$ et $x \in \mathbb{R}$. En posant $\alpha = z$ et $x = 1$ on obtient :

$$\forall z \in \mathbb{C}, \quad e^z = \sum_{n=0}^{+\infty} \frac{z^n}{n!}.$$

Il est intéressant d'observer que la propriété fondamentale de la fonction exponentielle, à savoir : $\forall (z, z') \in \mathbb{C}^2$, $\exp(z + z') = \exp(z) \times \exp(z')$ peut être prouvée à partir de cette expression. En effet, la convergence absolue de cette série autorise un produit de Cauchy :

$$e^z \times e^{z'} = \left(\sum_{p=0}^{+\infty} \frac{z^p}{p!} \right) \left(\sum_{q=0}^{+\infty} \frac{z'^q}{q!} \right) = \sum_{n=0}^{+\infty} \sum_{p=0}^n \left(\frac{1}{p!} \times \frac{1}{(n-p)!} z^p z'^{n-p} \right) = \sum_{n=0}^{+\infty} \frac{1}{n!} \sum_{p=0}^n \binom{n}{p} z^p z'^{n-p} = \sum_{n=0}^{+\infty} \frac{1}{n!} (z + z')^n = e^{z+z'}.$$

Similairement, les sommes géométriques étudiées en première année fournissent un deuxième exemple de fonction définie sur une partie du plan complexe et développables en série entière :

$$\forall z \in \mathbb{C}, \quad |z| < 1 \implies \frac{1}{1-z} = \sum_{n=0}^{+\infty} z^n.$$

Notons que les différents développements en série entière des fonctions réelles que nous avons obtenus pourraient être utilisés pour prolonger les fonctions correspondantes dans le plan complexe (ou le disque unité suivant les cas), mais nous ne nous aventurerons pas plus loin dans cette direction. Nous nous contenterons d'admettre le résultat suivant :

PROPOSITION 2.6 — Soit $\sum a_n z^n$ une série entière complexe de rayon de convergence $R > 0$ et de somme $S(z)$. Alors la fonction S est continue sur le disque ouvert $\{z \in \mathbb{C} \mid |z| < R\}$.

Chapitre VII

Probabilités

La théorie mathématique des probabilités naît au XVI^e siècle sous l'impulsion de Jérôme Cardan puis de Blaise Pascal qui analysent les jeux de hasard. Des avancées majeures sont ensuite réalisées par Kolmogorov au début du XX^e siècle, qui fait la connexion entre la théorie de la mesure de Borel, l'intégration de Lebesgue et les probabilités, donnant à ces dernières des fondements incontestés.

1. Ensembles dénombrables et familles sommables

1.1 Ensembles dénombrables

Le cours de première année s'est restreint aux variables aléatoires à valeurs dans un ensemble fini ; cette année, nous allons étendre nos connaissances aux variables aléatoires à valeurs dans un ensemble infini, *mais pas à n'importe lesquels* : seuls les plus « simples » des ensembles infinis seront abordés, les ensembles dits *dénombrables* c'est-à-dire ceux qui peuvent être mis en bijection avec $\mathbb{N}$.

DÉFINITION. — On dit d'un ensemble E qu'il est :

- fini lorsqu'il existe un entier $n \in \mathbb{N}$ tel que E est en bijection avec $\llbracket 1, n \rrbracket$;
- dénombrable lorsqu'il est en bijection avec $\mathbb{N}$;
- discret s'il est fini ou dénombrable. On dira aussi que E est au plus dénombrable.

Si E est un ensemble dénombrable, il existe donc une bijection $\phi : \mathbb{N} \rightarrow E$. En posant pour tout $n \in \mathbb{N}$, $x_n = \phi(n)$ il devient possible de définir E en *extension*, c'est-à-dire sous la forme : $E = \{x_n \mid n \in \mathbb{N}\}$.

Exemple. L'ensemble $2\mathbb{N}$ des entiers pairs est dénombrable, puisqu'il peut être défini en extension : $2\mathbb{N} = \{2n \mid n \in \mathbb{N}\}$, ce qui correspond à la bijection $\phi : \mathbb{N} \rightarrow 2\mathbb{N}$, $\phi(n) = 2n$.

Il en est bien entendu de même de l'ensemble $2\mathbb{N} + 1$ des entiers impairs : $2\mathbb{N} + 1 = \{2n + 1 \mid n \in \mathbb{N}\}$.

Plus généralement, on dispose du résultat suivant :

PROPOSITION 1.1 — Toute partie d'un ensemble dénombrable est finie ou dénombrable.

Par exemple, l'ensemble $\mathcal{P}$ des nombres premiers est infini (vous avez du démontrer ceci en première année) donc dénombrable puisqu'inclus dans $\mathbb{N}$. Il existe donc une suite (p_n) telle que $\mathcal{P} = \{p_n \mid n \in \mathbb{N}\}$.

PROPOSITION 1.2 — Soit E un ensemble dénombrable et F un ensemble au plus dénombrable. Alors $E \cup F$ est dénombrable.

COROLLAIRE — $\mathbb{Z}$ est un ensemble dénombrable.

PROPOSITION 1.3 — Soit E un ensemble dénombrable et F un ensemble non vide au plus dénombrable. Alors le produit cartésien $E \times F$ est dénombrable.

■ Réunion et intersection dénombrables

DÉFINITION. — Soit $(A_n)_{n \in \mathbb{N}}$ une suite de parties d'un même ensemble E . On définit les ensembles $\bigcup_{n \in \mathbb{N}} A_n$ et $\bigcap_{n \in \mathbb{N}} A_n$ par :

$$x \in \bigcup_{n \in \mathbb{N}} A_n \iff \exists n \in \mathbb{N} \mid x \in A_n \quad \text{et} \quad x \in \bigcap_{n \in \mathbb{N}} A_n \iff \forall n \in \mathbb{N}, x \in A_n$$

PROPOSITION 1.4 — Une union finie ou dénombrable d'ensembles dénombrables est dénombrable.

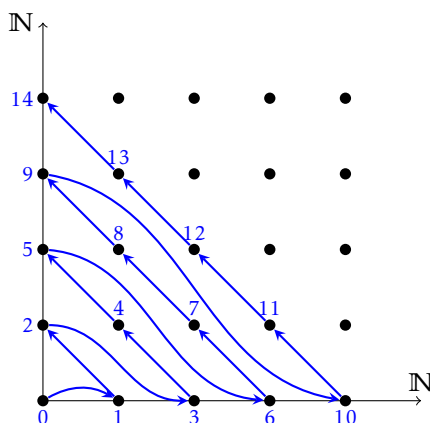


FIGURE 1 – $\mathbb{N} \times \mathbb{N}$ est dénombrable car on peut énumérer ses éléments.

COROLLAIRE — $\mathbb{Q}$ est un ensemble dénombrable.

Jusqu'à présent, nous n'avons vu que des ensembles dénombrables, pour la bonne et simple raison qu'il est plus facile de prouver qu'un ensemble est dénombrable que de prouver qu'il ne l'est pas. C'est Cantor qui le premier a donné des exemples d'ensembles non dénombrables, en utilisant une méthode qui maintenant porte son nom : *l'argument de la diagonale de Cantor*. Nous admettrons le résultat suivant :

THÉORÈME 1.5 — L'ensemble $\mathcal{P}(\mathbb{N})$ des parties de $\mathbb{N}$ n'est pas dénombrable. $\mathbb{R}$ n'est pas dénombrable. L'ensemble $\{0, 1\}^{\mathbb{N}}$ des suites à valeurs dans $\{0, 1\}$ n'est pas dénombrable.

Les ensembles cités sont d'une certaine manière « trop gros » pour être dénombrables.

1.2 Familles sommables de réels positifs

Étant donné une famille de réels positifs $(x_i)_{i \in I}$ indexée par un ensemble I , nous allons chercher à déterminer s'il est possible de donner un sens à la somme $\sum_{i \in I} x_i$ de ces derniers.

Lorsque I est un ensemble fini : $I = \{i_1, \dots, i_n\}$, cette question ne paraît guère intéressante : la commutativité de l'addition implique que la somme $x_{i_1} + \dots + x_{i_n}$ reste inchangée quel que soit la manière de permuer ses éléments. Mais ceci n'a plus rien d'évident lorsque I est un ensemble infini.

Attention. Il ne s'agit pas ici du concept de série numérique : dans le cas d'une série numérique, les éléments de la famille ont été ordonnés au préalable, et on détermine ensuite la convergence ou la divergence de la suite des sommes partielles. Nous cherchons ici à nous affranchir de l'ordre des éléments dans la somme.

DÉFINITION. — Soit $(x_i)_{i \in I}$ une famille de réels positifs. Cette famille est dite sommable lorsque l'ensemble

$$\left\{ \sum_{i \in J} x_i \mid J \subset I \text{ et } J \text{ finie} \right\}$$

est majoré. Dans ce cas, on note $\sum_{i \in I} x_i$ la borne supérieure de cet ensemble.

Remarque. Lorsque la famille de réels positifs $(x_i)_{i \in I}$ n'est pas sommable, on notera par commodité $\sum_{i \in I} x_i = +\infty$.

PROPOSITION 1.6 — Si la famille $(x_i)_{i \in I}$ est sommable, l'ensemble $\{i \in I \mid x_i \neq 0\}$ est au plus dénombrable.

Il résulte de la proposition ci-dessus que dans la pratique, on pourra toujours supposer, lorsque la famille est sommable, que I est un ensemble fini ou dénombrable.

■ Lien avec les séries numériques

THÉORÈME 1.7 — Soit $(x_n)_{n \in \mathbb{N}}$ une suite de nombres réels positifs. Alors la famille $(x_n)_{n \in \mathbb{N}}$ est sommable si et seulement si la série numérique $\sum x_n$ converge, et dans ce cas, $\sum_{n \in \mathbb{N}} x_n = \sum_{n=0}^{+\infty} x_n$.

Ce résultat nous donne une manière très simple d'étudier la sommabilité d'une famille dénombrable de réels positifs : il suffit de les ordonner d'une manière arbitraire puis d'étudier la convergence de la série numérique afférente.

Maintenant que nous avons trouvé un moyen de nous affranchir de l'ordre de sommation d'une famille de réels positifs, il nous reste à énoncer deux formules couramment utilisées dans les calculs :

THÉORÈME 1.8 (sommation par paquets) — Soit $(I_n)_{n \in \mathbb{N}}$ une partition dénombrable de I , et $(x_i)_{i \in I}$ une famille sommable de réels positifs. Alors $\sum_{i \in I} x_i = \sum_{n=0}^{+\infty} \left(\sum_{i \in I_n} x_i \right)$.

THÉORÈME 1.9 (Fubini) — Soit $(x_{i,j})_{(i,j) \in I \times J}$ une famille sommable de réels positifs. Alors

$$\sum_{(i,j) \in I \times J} x_{i,j} = \sum_{i \in I} \sum_{j \in J} x_{i,j} = \sum_{j \in J} \sum_{i \in I} x_{i,j}$$

Remarque. Nous l'avons dit, lorsqu'une famille $(x_i)_{i \in I}$ de réels positifs n'est pas sommable, on s'autorisera à écrire $\sum_{i \in I} x_i = +\infty$. Ceci a pour conséquence que les deux théorèmes ci-dessus peuvent s'appliquer *sans justification préalable de la sommabilité*. Obtenir à la fin des calculs une somme finie justifiera *a posteriori* la sommabilité de la famille, et au contraire obtenir une somme divergente prouvera la non sommabilité de cette famille.

1.3 Familles sommables de réels quelconques

DÉFINITION. — Une famille $(x_i)_{i \in I}$ est dite sommable lorsque la famille de réels positifs $(|x_i|)_{i \in I}$ est sommable.

Une fois cette définition posée, on définit la somme à l'aide du résultat suivant :

THÉORÈME 1.10 — Soit $(x_i)_{i \in I}$ une famille sommable de réels quelconques. Pour tout $i \in I$ on pose $x_i^+ = \max(x_i, 0)$ et $x_i^- = \max(0, -x_i)$. Alors les familles de réels positifs $(x_i^+)_{i \in I}$ et $(x_i^-)_{i \in I}$ sont sommables, et on pose par définition :

$$\sum_{i \in I} x_i = \sum_{i \in I} x_i^+ - \sum_{i \in I} x_i^-$$

On dispose alors des résultats suivants, que nous admettrons. Notez cependant que contrairement aux familles sommables de réels positifs pour lesquels l'obtention d'un résultat fini à la fin des calculs justifie *a posteriori* la sommabilité de la famille, il est indispensable, dans le cas d'une famille de réels quelconques, de justifier la sommabilité *en préalable à tout calcul*.

PROPOSITION 1.11 — Soit $(x_i)_{i \in I}$ une famille sommable. Alors $\left| \sum_{i \in I} x_i \right| \leq \sum_{i \in I} |x_i|$.

PROPOSITION 1.12 — soient $(x_i)_{i \in I}$ et $(y_i)_{i \in I}$ deux familles de nombres réels telles que pour tout $i \in I$, $|x_i| \leq y_i$. Alors la sommabilité de $(y_i)_{i \in I}$ entraîne celle de $(x_i)_{i \in I}$.

PROPOSITION 1.13 — Soit $(x_i)_{i \in I}$ et $(y_i)_{i \in I}$ deux familles sommables, et $\lambda \in \mathbb{R}$. Alors la famille $(\lambda x_i + y_i)_{i \in I}$ est sommable, et $\sum_{i \in I} (\lambda x_i + y_i) = \lambda \sum_{i \in I} x_i + \sum_{i \in I} y_i$.

THÉORÈME 1.14 (somme par paquets) — Soit $(I_n)_{n \in \mathbb{N}}$ une partition dénombrable de I , et $(x_i)_{i \in I}$ une famille sommable. Alors $\sum_{i \in I} x_i = \sum_{n=0}^{+\infty} \left(\sum_{i \in I_n} x_i \right)$.

THÉORÈME 1.15 (Fubini) — Soit $(x_{i,j})_{(i,j) \in I \times J}$ une famille sommable. Alors $\sum_{(i,j) \in I \times J} x_{i,j} = \sum_{i \in I} \sum_{j \in J} x_{i,j} = \sum_{j \in J} \sum_{i \in I} x_{i,j}$.

2. Espaces probabilisés

2.1 Expérience aléatoire et univers

DÉFINITION. — On appelle expérience aléatoire une expérience qui, reproduite dans des conditions identiques, peut conduire à des résultats différents non prévisibles à l'avance. L'ensemble des résultats possibles de cette expérience est appelé univers et est classiquement noté Ω .

Exemples. Examinons tout d'abord quelques expériences aléatoires et l'univers qui leur est associé :

- on lance trois dés à 6 faces. Dans ce cas, on choisira pour univers $\Omega = \llbracket 1, 6 \rrbracket^3$.
- on lance une pièce de monnaie jusqu'à obtenir Face. Ici pourra choisir $\Omega = \mathbb{N} \cup \{+\infty\}$ si on choisit de représenter une expérience par le nombre d'essais infructueux.
- on casse une baguette de bois en trois et on mesure les longueurs des trois morceaux. En fixant à 1 la longueur de la baguette, l'univers peut être représenté par $\Omega = \{(x, y, z) \in]0, 1[^3 \mid x + y + z = 1\}$.

Le premier exemple correspond à un univers fini, le second à un univers dénombrable, le troisième à un univers non dénombrable.

On observera que la description de l'univers ne nous indique pas la façon dont l'expérience est réalisée : les dés, la pièce, sont-ils pipés ou non ? Suivant quel protocole la baguette est-elle brisée ? Ce sont ces informations qui vont conditionner le choix de la probabilité que nous allons associer à cet univers.

2.2 Tribu et événements

Considérons un univers Ω . Lorsque ce dernier est fini, on appelle *événement* toute partie de Ω . En conjonction avec le vocabulaire de la théorie des ensembles, ont été définies les notions suivantes :

- un événement est dit *élémentaire* si c'est un singleton ;
- l'événement *certain* est l'événement Ω ;
- l'événement *impossible* est l'événement $\emptyset$;
- l'événement *non A*, *contraire* de l'événement A , est l'événement $\bar{A} = \Omega \setminus A$ (le résultat de l'expérience n'appartient pas à A) ;
- si A et B sont des événements, l'événement $A \text{ et } B$ est l'événement $A \cap B$ (le résultat de l'expérience se trouve dans A et dans B) ;
- si A et B sont des événements, l'événement $A \text{ ou } B$ est l'événement $A \cup B$ (le résultat de l'expérience se trouve dans A ou dans B) ;
- les événements A et B sont dits *incompatibles* lorsque $A \cap B = \emptyset$ (le résultat de l'expérience ne peut se trouver à la fois dans A et dans B) ;
- Si A et B sont deux événements, on dit que A *entraîne* B lorsque $A \subset B$ (si le résultat de l'expérience se trouve dans A , il se trouve aussi dans B).

Exemple. Considérons le lancer de trois dés associé à l'univers $\Omega = \llbracket 1, 6 \rrbracket^3$.

$A = \{(x, y, z) \in \Omega \mid x + y + z \leq 10\}$ est l'événement : « la somme des trois dés est inférieure ou égale à 10 ».

$B = \{(x, y, z) \in \Omega \mid x + y + z \geq 10\}$ est l'événement : « la somme des trois dés est supérieure ou égale à 10 ».

$A \text{ ou } B$ est l'événement certain ; $A \text{ et } B$ est l'événement : « la somme des trois dés est égale à 10 ». Enfin, $\text{non } A$ est l'événement « la somme des trois dés est strictement supérieure à 10 » donc $\text{non } A$ entraîne B .

Une fois la notion d'événement définie, l'étape suivante dans la construction d'un espace probabilisé consiste à définir une probabilité $\mathbb{P}(A)$ mesurant la chance de réalisation d'un événement A . Or lorsque l'univers Ω est infini, il n'est en général pas possible de définir cette probabilité pour toutes les parties de Ω ; il faut se restreindre à un sous-ensemble $\mathcal{A}$ de $\mathcal{P}(\Omega)$ qu'on appelle une *tribu*, et qui en quelque sorte contient les événements dont on pourra mesurer la probabilité de réussite.

Plus formellement nous adopterons la définition suivante :

DÉFINITION. — Si Ω est un ensemble, on appelle tribu sur Ω une partie $\mathcal{A}$ de $\mathcal{P}(\Omega)$ vérifiant :

- $\Omega \in \mathcal{A}$ (l'événement certain appartient à la tribu);
- pour tout $A \in \mathcal{A}$, l'événement contraire $\bar{A}$ appartient à $\mathcal{A}$;
- $\mathcal{A}$ est stable par réunion dénombrable, c'est-à-dire que si $(A_n)_{n \in \mathbb{N}}$ est une suite d'éléments de $\mathcal{A}$ alors $\bigcup_{n \in \mathbb{N}} A_n$ appartient à $\mathcal{A}$.

Désormais, le terme d'événement désignera un élément d'une tribu $\mathcal{A}$, supposée définie précédemment.

PROPOSITION 2.1 — Si $\mathcal{A}$ est une tribu sur l'univers Ω , alors :

- $\emptyset \in \mathcal{A}$ (l'événement impossible appartient à la tribu);
- si A et B sont deux événements de la tribu $\mathcal{A}$, il en est de même de $A \cup B$ et de $A \cap B$;
- $\mathcal{A}$ est stable par intersection dénombrable, c'est-à-dire que si $(A_n)_{n \in \mathbb{N}}$ est une suite d'éléments de $\mathcal{A}$ alors $\bigcap_{n \in \mathbb{N}} A_n$ appartient à $\mathcal{A}$.

Exemple. $\{\emptyset, \Omega\}$ est une tribu, appelée *tribu triviale* puisqu'elle ne mesure que deux événements : l'événement certain et l'événement impossible.

Exemple. À l'inverse, $\mathcal{P}(\Omega)$ est la tribu la plus fine qui soit. Cependant, à l'exception des univers finis ou dénombrables, cette tribu ne peut engendrer que des espaces probabilisés sans intérêt.

Exemple. Considérons de nouveau l'expérience consistant à jeter une pièce jusqu'à obtenir Face, mais choisissons cette fois l'univers $\Omega = \{0, 1\}^{\mathbb{N}}$ (autrement dit, dans l'univers des possibles on joue à Pile ou Face indéfiniment). Cet univers n'est pas dénombrable, il est donc nécessaire de définir une tribu sur laquelle on pourra ensuite définir une probabilité. Compte tenu du problème qui nous intéresse on admet l'existence d'une tribu $\mathcal{A}$ dans laquelle « Face apparaît pour la première fois au n^{e} tirage » est un événement noté A_n .

Compte tenu des propriétés des tribus, l'événement $A = \bigcup_{n \in \mathbb{N}^*} A_n$ appartient à $\mathcal{A}$ (il s'agit de l'événement « Face apparaît au moins une fois ») ainsi que l'événement contraire $\bar{A}$ (« la pièce tombe indéfiniment sur Pile »). Tous les événements nécessaires à l'étude de l'expérience sont bien présents dans la tribu.

Exercice 1

Soit $\mathcal{A}$ une tribu de $\mathbb{R}$ contenant toutes les demi-droites $[a, +\infty[$, $a \in \mathbb{R}$. Montrer que cette tribu contient tous les intervalles de $\mathbb{R}$.

2.3 Définition d'une probabilité

Nous sommes maintenant en mesure de donner la définition générale d'une probabilité.

DÉFINITION. — Soit Ω un univers et $\mathcal{A}$ une tribu sur Ω . On appelle probabilité sur $(\Omega, \mathcal{A})$ une application $\mathbb{P} : \mathcal{A} \rightarrow [0, 1]$ vérifiant :

- $\mathbb{P}(\Omega) = 1$;
- pour toute suite dénombrable $(A_n)_{n \in \mathbb{N}}$ d'événements de $\mathcal{A}$ deux-à-deux incompatibles la série $\sum \mathbb{P}(A_n)$ converge, et $\mathbb{P}\left(\bigcup_{n \in \mathbb{N}} A_n\right) = \sum_{n=0}^{+\infty} \mathbb{P}(A_n)$.

On appelle espace probabilisé le triplet $(\Omega, \mathcal{A}, \mathbb{P})$ constitué d'un univers, d'une tribu sur Ω et d'une probabilité sur $(\Omega, \mathcal{A})$.

Commençons par observer que les propriétés sur les univers finis qui ont été établies dans le cours de première année restent vérifiées :

PROPOSITION 2.2 — Une probabilité vérifie les propriétés suivantes :

- $\mathbb{P}(\emptyset) = 0$;
- si A et B sont deux événements incompatibles, $\mathbb{P}(A \cup B) = \mathbb{P}(A) + \mathbb{P}(B)$;
- si A est un événement, $\mathbb{P}(\bar{A}) = 1 - \mathbb{P}(A)$;
- si A et B sont deux événements, $\mathbb{P}(A \cap B) + \mathbb{P}(A \cup B) = \mathbb{P}(A) + \mathbb{P}(B)$;
- si $A \subset B$ sont deux événements, alors $\mathbb{P}(A) \leq \mathbb{P}(B)$.

Exemple. Lorsque l'univers Ω est fini et $\mathcal{A} = \mathcal{P}(\Omega)$ il existe une unique probabilité, appelée *probabilité uniforme* telle que pour tout $\omega \in \Omega$, $\mathbb{P}(\{\omega\}) = \frac{1}{\text{card } \Omega}$. Dans ce cas, pour tout événement A , $\mathbb{P}(A) = \frac{\text{card } A}{\text{card } \Omega}$.

Exemple. Revenons maintenant sur l'expérience consistant à jeter une pièce de monnaie jusqu'à obtenir Face. Nous avons admis qu'on pouvait définir une tribu $\mathcal{A}$ sur l'univers $\Omega = \{0, 1\}^{\mathbb{N}}$ qui contient tous les événements A_n : « Face apparaît pour la première fois au n^{e} tirage ».

Si on note $p \in]0, 1[$ la probabilité pour la pièce de tomber sur Face, les éléments de l'univers sont des suites d'épreuves de Bernoulli indépendantes de paramètre p , et les éléments de A_n les suites qui débutent par $n-1$ échecs suivis d'une réussite donc $\mathbb{P}(A_n) = (1-p)^{n-1}p$.

Les événements A_n étant deux à deux incompatibles ($i \neq j \implies A_i \cap A_j = \emptyset$) on a $\mathbb{P}\left(\bigcup_{n \in \mathbb{N}} A_n\right) = \sum_{n=1}^{+\infty} \mathbb{P}(A_n) =$

$\sum_{n=1}^{+\infty} (1-p)^{n-1}p = p \times \frac{1}{1-(1-p)} = 1$. L'événement $A = \bigcup_{n \in \mathbb{N}^*} A_n$ (« Face apparaît au moins une fois ») vérifie $\mathbb{P}(A) = 1$,

l'événement $\bar{A}$ (« la pièce tombe indéfiniment sur Pile ») vérifie $\mathbb{P}(\bar{A}) = 1 - \mathbb{P}(A) = 0$.

L'événement $\bar{A}$ est dit « quasi-impossible », ou « négligeable » : bien qu'il soit un événement envisageable (il n'est pas égal à l'événement impossible $\emptyset$) sa probabilité est nulle. À l'inverse, l'événement A est dit « quasi-certain », ou « presque sûr ».

Voyons maintenant quelques résultats propres aux univers infinis :

THÉORÈME 2.3 (limite monotone) — Soit $(\Omega, \mathcal{A}, \mathbb{P})$ un espace probabilisé. Alors :

- pour toute suite d'événements (A_n) croissante au sens de l'inclusion ($A_n \subset A_{n+1}$), la suite $(\mathbb{P}(A_n))$ converge, et

$$\mathbb{P}\left(\bigcup_{n \in \mathbb{N}} A_n\right) = \lim \mathbb{P}(A_n) ;$$

- pour toute suite d'événements (A_n) décroissante au sens de l'inclusion ($A_{n+1} \subset A_n$), la suite $(\mathbb{P}(A_n))$ converge, et

$$\mathbb{P}\left(\bigcap_{n \in \mathbb{N}} A_n\right) = \lim \mathbb{P}(A_n).$$

Remarque. Lorsque la suite (A_n) n'est pas monotone au sens de l'inclusion, on peut néanmoins appliquer le théorème de la limite monotone à la suite des « union partielles » ou la suite des « intersections partielles ».

En effet, la suite $B_n = \bigcup_{k=0}^n A_k$ est croissante donc $\mathbb{P}\left(\bigcup_{n \in \mathbb{N}} A_n\right) = \mathbb{P}\left(\bigcup_{n \in \mathbb{N}} B_n\right) = \lim \mathbb{P}(B_n)$.

De même la suite $C_n = \bigcap_{k=0}^n A_k$ est décroissante donc $\mathbb{P}\left(\bigcap_{n \in \mathbb{N}} A_n\right) = \mathbb{P}\left(\bigcap_{n \in \mathbb{N}} C_n\right) = \lim \mathbb{P}(C_n)$.

PROPOSITION 2.4 (sous-additivité) — Soit $(\Omega, \mathcal{A}, \mathbb{P})$ un espace probabilisé. Pour toute suite d'événements (A_n) ,

$\mathbb{P}\left(\bigcup_{n \in \mathbb{N}} A_n\right) \leq \sum_{n=0}^{+\infty} \mathbb{P}(A_n)$ (cette somme peut éventuellement être égale à $+\infty$).

Exercice 2 [Lemme de Borel-Cantelli]

Soit (A_n) une suite d'événements; pour tout $p \in \mathbb{N}$ on pose $B_p = \bigcup_{n \geq p} A_n$ puis $A^* = \bigcap_{p \in \mathbb{N}} B_p$.

On suppose que la série $\sum \mathbb{P}(A_n)$ converge. Montrer que $\mathbb{P}(A^*) = 0$.

■ Probabilité sur un univers dénombrable

Considérons maintenant un univers dénombrable Ω , que l'on peut donc décrire par extension : $\Omega = \{\omega_n \mid n \in \mathbb{N}\}$. Nous allons prouver le résultat suivant, qui montre qu'il est toujours possible de définir une probabilité sur $(\Omega, \mathcal{P}(\Omega))$ à partir de la valeur de $\mathbb{P}$ sur les singletons :

THÉORÈME 2.5 — Soit (p_n) une suite de réels positifs telle que la série $\sum p_n$ converge et $\sum_{n=0}^{+\infty} p_n = 1$. Alors il existe une unique probabilité $\mathbb{P}$ sur $(\Omega, \mathcal{P}(\Omega))$ telle que pour tout $n \in \mathbb{N}$, $\mathbb{P}(\{\omega_n\}) = p_n$.

Exemple. Soit $\theta > 0$ et $p_n = e^{-\theta} \frac{\theta^n}{n!}$. Il est facile de vérifier que $0 \leq p_n \leq 1$ et que $\sum_{n=0}^{+\infty} p_n = e^{-\theta} \sum_{n=0}^{+\infty} \frac{\theta^n}{n!} = 1$. La suite (p_n) définit donc une probabilité sur $(\mathbb{N}, \mathcal{P}(\mathbb{N}))$ en posant $\mathbb{P}(\{n\}) = e^{-\theta} \frac{\theta^n}{n!}$, appelée *loi de Poisson de paramètre θ* . Nous aurons l'occasion d'y revenir.

Exercice 3

- Soit $\mathbb{P}$ une probabilité sur $(\mathbb{N}, \mathcal{P}(\mathbb{N}))$. Montrer que $\lim \mathbb{P}(\{n\}) = 0$.
- Soit (a_n) une suite strictement décroissante de réels positifs de limite nulle. Déterminer une constante $\lambda > 0$ pour qu'il existe une probabilité $\mathbb{P}$ sur $(\mathbb{N}, \mathcal{P}(\mathbb{N}))$ vérifiant : $\mathbb{P}([n, +\infty[)) = \lambda a_n$.

2.4 Conditionnement et indépendance

Dans toute la suite du cours, $(\Omega, \mathcal{A}, \mathbb{P})$ désigne un espace probabilisé.

■ Probabilité conditionnelle

DÉFINITION. — Si A et B sont deux événements tels que $\mathbb{P}(B) > 0$, on appelle probabilité conditionnelle de A sachant B le réel $\mathbb{P}_B(A) = \frac{\mathbb{P}(A \cap B)}{\mathbb{P}(B)}$, réel qu'on pourra aussi noter $\mathbb{P}(A \mid B)$.

THÉORÈME 2.6 — $\mathbb{P}_B$ est une probabilité sur $(\Omega, \mathcal{A})$.

Remarque. On dispose donc de l'égalité $\mathbb{P}(A \cap B) = \mathbb{P}(B)\mathbb{P}(A \mid B)$ lorsque $\mathbb{P}(B) \neq 0$. Lorsque $\mathbb{P}(B) = 0$, on peut observer que cette égalité garde un sens (celui de « $0 = 0$ ») même si $\mathbb{P}(A \mid B)$ n'est pas formellement défini puisque $A \cap B \subset B \Rightarrow 0 \leq \mathbb{P}(A \cap B) \leq \mathbb{P}(B) = 0$.

Si A et B sont deux événements quelconques, la formule $\mathbb{P}(A \cap B) = \mathbb{P}(B)\mathbb{P}(A \mid B)$ est appelée *formule des probabilités composées*.

DÉFINITION. — On appelle système complet d'événements toute famille $(B_i)_{i \in I}$ finie ou dénombrable d'événements deux-à-deux incompatibles ($i \neq j \Rightarrow B_i \cap B_j = \emptyset$) et telle que $\bigcup_{i \in I} B_i = \Omega$.

En d'autres termes, la famille $(B_i)_{i \in I}$ constitue une partition finie ou dénombrable de Ω .

THÉORÈME 2.7 (formule des probabilités totales) — Soit A un événement et $(B_i)_{i \in I}$ un système complet d'événements.

Alors $\mathbb{P}(A) = \sum_{i \in I} \mathbb{P}(B_i)\mathbb{P}(A \mid B_i)$.

Remarque. La formule reste valable lorsque $\sum_{i \in I} \mathbb{P}(B_i) = 1$, autrement dit lorsque l'événement $\bigcup_{i \in I} B_i$ est presque sûr. On parle alors de système *quasi-complet* d'événements.

PROPOSITION 2.8 (Formule de Bayes) — Soit $(B_i)_{i \in I}$ un système complet d'événements tel que pour tout $i \in I$,

$$\mathbb{P}(B_i) > 0, \text{ et } A \text{ un événement tel que } \mathbb{P}(A) > 0. \text{ Alors } \mathbb{P}(B_i | A) = \frac{\mathbb{P}(B_i)\mathbb{P}(A | B_i)}{\sum_{j \in I} \mathbb{P}(B_j)\mathbb{P}(A | B_j)}.$$

Remarque. Cette formule est souvent utilisée lorsque le système complet est constitué des deux seuls événements B et $\bar{B}$. Dans ce cas, la formule devient : $\mathbb{P}(B | A) = \frac{\mathbb{P}(B)\mathbb{P}(A | B)}{\mathbb{P}(B)\mathbb{P}(A | B) + \mathbb{P}(\bar{B})\mathbb{P}(A | \bar{B})}$.

Exercice 4

Un QCM propose 4 réponses pour chaque question. Soit p la probabilité qu'un étudiant connaisse la bonne réponse à une question donnée. S'il ignore la réponse, il choisit au hasard l'une des réponses proposées. Quel est la probabilité qu'un étudiant connaisse vraiment la bonne réponse lorsqu'il a correctement répondu à une question ?

Remarque. La formule de Bayes a longtemps été appelée formule de probabilité des causes. Elle permet en effet de calculer la probabilité d'une cause (ici le fait d'avoir pris le dé pipé) connaissant celle de sa conséquence (le nombre de 6 obtenus).

Exercice 5

On dépose dans une urne vide une boule blanche puis on joue à Pile ou Face avec une pièce non pipée. Tant que la pièce retombe sur Pile, on ajoute une boule noire dans l'urne. Lorsqu'on obtient Face pour la première fois on tire au hasard une boule de l'urne. Celle-ci est blanche. Quelle est la probabilité qu'il n'y ait aucune boule noire dans l'urne ?

■ Indépendance

De manière informelle, deux événements A et B sont indépendants lorsque le fait de savoir que A est réalisé ne donne aucune information sur la réalisation de B , et réciproquement. Ainsi, lorsque $\mathbb{P}(A) > 0$ et $\mathbb{P}(B) > 0$ on souhaite que $\mathbb{P}(A | B) = \mathbb{P}(A)$ et $\mathbb{P}(B | A) = \mathbb{P}(B)$, ce qui se traduit par $\frac{\mathbb{P}(A \cap B)}{\mathbb{P}(B)} = \mathbb{P}(A)$ et $\frac{\mathbb{P}(B \cap A)}{\mathbb{P}(A)} = \mathbb{P}(B)$. Ces deux égalités sont identiques, et pour pouvoir s'abstraire des hypothèses $\mathbb{P}(A) > 0$ et $\mathbb{P}(B) > 0$ on adoptera la définition suivante :

DÉFINITION. — Deux événements A et B sont dits indépendants lorsque $\mathbb{P}(A \cap B) = \mathbb{P}(A)\mathbb{P}(B)$.

PROPOSITION 2.9 — Si A et B sont indépendants, il en est de même de $\bar{A}$ et B , de A et $\bar{B}$, de $\bar{A}$ et $\bar{B}$.

La notion d'indépendance se généralise à une suite finie ou infinie d'événements de la manière suivante :

DÉFINITION. — Une famille finie ou dénombrable $(A_i)_{i \in I}$ d'événements est dite indépendante lorsque pour tout entier $p \leq \text{card } I$, pour toute p -liste $(i_1, \dots, i_p) \in I^p$ d'indices deux-à-deux distincts, $\mathbb{P}(A_{i_1} \cap \dots \cap A_{i_p}) = \mathbb{P}(A_{i_1}) \cdots \mathbb{P}(A_{i_p})$ (on dit aussi que les événements A_i sont mutuellement indépendants).

On observera que cette définition est très délicate à mettre en œuvre. Ne serait-ce que pour trois événements A , B et C il faut vérifier *chacune* des égalités :

$$\mathbb{P}(A \cap B \cap C) = \mathbb{P}(A)\mathbb{P}(B)\mathbb{P}(C), \quad \mathbb{P}(A \cap B) = \mathbb{P}(A)\mathbb{P}(B), \quad \mathbb{P}(B \cap C) = \mathbb{P}(B)\mathbb{P}(C), \quad \mathbb{P}(C \cap A) = \mathbb{P}(C)\mathbb{P}(A).$$

En particulier, les trois dernières égalités, qui traduisent le fait que ces trois événements sont deux-à-deux indépendants, ne sont pas suffisantes pour s'assurer que les trois événements sont indépendants.

Exercice 6

Soit (A_n) une suite d'événements indépendants ; pour tout $p \in \mathbb{N}$ on pose $B_p = \bigcup_{n \geq p} A_n$ puis $A^* = \bigcap_{p \in \mathbb{N}} B_p$.

a. Justifier que $\mathbb{P}(A^*) = \lim_{p \rightarrow +\infty} \mathbb{P}(B_p)$ et que $\mathbb{P}(B_p) = 1 - \lim_{n \rightarrow +\infty} \prod_{k=p}^n (1 - \mathbb{P}(A_k))$.

b. Montrer que $\prod_{k=p}^n (1 - \mathbb{P}(A_k)) \leq \exp\left(-\sum_{k=p}^n \mathbb{P}(A_k)\right)$ et en déduire que si la série $\sum \mathbb{P}(A_n)$ diverge, $\mathbb{P}(A^*) = 1$.

Remarque. Le résultat ci-dessus, associé au lemme de Borel-Cantelli (voir page 7) constitue la *loi du zéro-un de Borel* : si (A_n) est une suite d'événements indépendants, la probabilité qu'une infinité d'entre eux se réalise est :

- égale à 0 si la série $\sum \mathbb{P}(A_n)$ converge ;
- égale à 1 si la série $\sum \mathbb{P}(A_n)$ diverge.

3. Variables aléatoires

3.1 Définition d'une variable aléatoire

Jusqu'à présent, nous avons beaucoup parlé des événements, autrement dit adopté un point de vue *ensembliste* sur les probabilités. Nous allons maintenant changer de point de vue en choisissant un point de vue *fonctionnel* à l'aide de la notion de *variable aléatoire* qui, contrairement à ce que pourrait laisser supposer son nom, n'est pas une variable mais une fonction. De manière informelle, une variable aléatoire est une grandeur qui dépend du résultat de l'expérience ; ce peut être par exemple :

- le nombre de 6 obtenus dans un lancé de trois dés ;
- le temps d'attente avant d'obtenir Face dans un lancer de pièce ;
- la longueur du plus grand des deux morceaux lorsqu'on brise une baguette de bois en deux.

DÉFINITION. — Si $(\Omega, \mathcal{A})$ est un espace probabilisable et E un ensemble, on appelle variable aléatoire toute fonction $X : \Omega \rightarrow E$ telle que pour tout $e \in E$, $X^{-1}(\{e\}) \in \mathcal{A}$ (autrement dit, $X^{-1}(\{e\})$ est un événement).

Lorsque $E = \mathbb{R}$, la variable aléatoire X sera dite réelle.

Lorsque $X(\Omega)$ (l'ensemble des valeurs que peut prendre X) est fini ou dénombrable, la variable aléatoire X sera dite discrète.

Rappel. La notation $X^{-1}(\{e\})$ désigne l'image réciproque de e , c'est-à-dire l'ensemble des antécédents de e : $X^{-1}(\{e\}) = \{\omega \in \Omega \mid X(\omega) = e\}$.

Exemples.

– Pour l'expérience consistant à lancer trois dés et à compter le nombre de 6, nous pouvons choisir $\Omega = \llbracket 1, 6 \rrbracket^3$, $E = \mathbb{N}$ et $X : \Omega \rightarrow \mathbb{N}$ définie par $X(e_1, e_2, e_3) = \text{card}\{i \in \llbracket 1, 3 \rrbracket \mid e_i = 6\}$.

$X(\Omega) = \{0, 1, 2, 3\}$ donc la variable aléatoire X est discrète (finie).

– Pour l'expérience consistant à lancer une pièce jusqu'à obtenir Face, nous pouvons choisir $\Omega = \{0, 1\}^{\mathbb{N}}$, $E = \mathbb{N} \cup \{+\infty\}$ et $X : \Omega \rightarrow E$ définie par $X((u_n)) = \min\{n \in \mathbb{N}^* \mid u_n = 0\}$. Ici X est une variable aléatoire discrète (dénombrable).

– Pour l'expérience consistant à casser une baguette de deux pour mesurer le plus grand des deux morceaux, nous avons $\Omega =]0, 1[$, $E =]0, 1[$ et $X(x) = \max(x, 1-x)$. Dans cet exemple, X n'est pas une variable aléatoire discrète car $X(\Omega) =]1/2, 1[$ n'est pas dénombrable.

Dans la suite de ce cours nous ne prendrons en considération que des variables aléatoires discrètes.

PROPOSITION 3.1 — Lorsque X est une variable aléatoire discrète, pour tout $U \subset X(\Omega)$, $X^{-1}(U) \in \mathcal{A}$ (autrement dit, $X^{-1}(U)$ est un événement).

Remarque. On introduit la notion de variable aléatoire pour s'intéresser aux chances de réalisation des valeurs de X plutôt qu'aux chances de réalisation des résultats de l'expérience. Autrement dit, cette notion permet d'une certaine façon d'« oublier » l'espace probabilisable $(\Omega, \mathcal{A})$ (qui reste présent, mais dont on se contentera le plus souvent d'admettre son existence) au profit des valeurs prises par X .

Par la suite, l'événement $X^{-1}(U)$ sera noté plus simplement $[X \in U]$.

Par exemple, pour le jeté de trois dés, $[X = 2]$ désigne l'événement « deux des trois dés ont donné un 6 ». Pour le lancer d'une pièce jusqu'à obtenir Face, $[X \geq 3]$ désigne l'événement « il a fallu au moins trois lancers avant d'obtenir un Face ».

L'intérêt du résultat précédent est que puisque $[X \in U]$ est un événement, il est possible de lui associer une probabilité. Il s'agit du résultat suivant :

THÉORÈME 3.2 — Soit $(\Omega, \mathcal{A}, \mathbb{P})$ un espace probabilisé et $X : \Omega \rightarrow E$ une variable aléatoire discrète. Alors l'application $\mathbb{P}_X : \mathcal{P}(X(\Omega)) \rightarrow [0, 1]$ définie par $\mathbb{P}_X(U) = \mathbb{P}(X^{-1}(U)) = \mathbb{P}(X \in U)$ est une probabilité sur $(X(\Omega), \mathcal{P}(X(\Omega)))$, appelée loi de la variable X , ou encore distribution de X .

Exercice 7

Une urne contient initialement une boule blanche et une boule noire. On tire au hasard une de ces boules, on note sa couleur et on la replace dans l'urne accompagnée d'une seconde boule de la même couleur. On réalise ce processus n fois, et on note X_n la variable aléatoire égale au nombre de boules blanches tirées durant ce processus. Déterminer la loi de X_n .

Ce résultat définit la loi de la variable aléatoire discrète X à partir de la loi de probabilité sur Ω . Il existe une réciproque de ce résultat : il est possible de choisir a priori la loi de X et d'en déduire une probabilité sur $X(\Omega)$. De manière plus formelle :

THÉORÈME 3.3 — Soit $(\Omega, \mathcal{A})$ un espace probabilisable et $X : \Omega \rightarrow E$ une variable aléatoire discrète. On note $X(\Omega) = \{x_i \mid i \in I\}$ et on considère une famille discrète $(p_i)_{i \in I}$ de réels positifs telle que $\sum_{i \in I} p_i = 1$. Alors il existe une probabilité $\mathbb{P}_X$ sur $(X(\Omega), \mathcal{P}(X(\Omega)))$ telle que pour tout $i \in I$, $\mathbb{P}_X(X = x_i) = p_i$.

L'intérêt de ce résultat est qu'il sera souvent suffisant de raisonner directement à partir de $\mathbb{P}_X$ sans véritablement avoir besoin d'explicitier formellement l'espace probabilisé $(\Omega, \mathcal{A}, \mathbb{P})$.

3.2 Lois discrètes classiques

Certaines lois interviennent régulièrement dans les problèmes de probabilités ; il est donc intéressant de les connaître afin de ne pas refaire à chaque fois les mêmes calculs. Nous allons maintenant passer en revue celles que vous devez connaître.

Remarque. Deux variables aléatoires X et Y qui suivent la même loi seront notées $X \sim Y$. On notera que si $X \sim Y$ alors pour toute fonction f on a $f(X) \sim f(Y)$.

■ Loi uniforme

L'expérience type consiste à considérer une urne contenant n boules numérotées de 1 à n et à effectuer un tirage équiprobable. La variable aléatoire X est le numéro de la boule obtenue.

DÉFINITION. — Soit $n \in \mathbb{N}^*$. On dit qu'une variable aléatoire réelle X suit une loi uniforme de paramètre n lorsque

$X(\Omega) = \llbracket 1, n \rrbracket$ et si pour tout $k \in \llbracket 1, n \rrbracket$, $\mathbb{P}(X = k) = \frac{1}{n}$. On note dans ce cas $X \sim \mathcal{U}(n)$.

■ Loi de Bernoulli

L'expérience type consiste à tirer dans une urne contenant une proportion p de boules blanches. On note X la variable aléatoire égale à 1 si on tire une boule blanche, et 0 sinon. On peut aussi tirer à pile ou face avec une pièce truquée ayant la probabilité p de tomber sur Face et poser $X = 0$ lorsque la pièce tombe sur Pile, et $X = 1$ lorsque la pièce tombe sur Face.

DÉFINITION. — Soit $p \in]0, 1[$. On dit qu'une variable aléatoire réelle X suit une loi de Bernoulli de paramètre p lorsque $X(\Omega) = \{0, 1\}$ et $\mathbb{P}(X = 0) = 1 - p$, $\mathbb{P}(X = 1) = p$. On note dans ce cas $X \sim \mathcal{B}(p)$.

Remarque. Pour des raisons de symétrie il est fréquent d'introduire la quantité $q = 1 - p$.

Variable indicatrice associée à un événement

À tout événement $A \in \mathcal{A}$ on peut associer la variable aléatoire $\mathbb{1}_A$ définie par :

$$\forall \omega \in \Omega, \quad \mathbb{1}_A(\omega) = \begin{cases} 1 & \text{si } \omega \in A \\ 0 & \text{sinon} \end{cases}$$

Cette variable aléatoire $\mathbb{1}_A$ est appelée l'indicatrice de A ; elle suit une loi de Bernoulli de paramètre $p = \mathbb{P}(A)$.

■ Loi géométrique

L'expérience type consiste en une succession infinie d'expériences de Bernoulli indépendantes de paramètre p . On note X le rang du premier succès.

DÉFINITION. — Soit $p \in]0, 1[$. On dit qu'une variable aléatoire réelle X suit une loi géométrique de paramètre p lorsque $X(\Omega) = \mathbb{N}^*$ et pour tout $k \in \mathbb{N}^*$, $\mathbb{P}(X = k) = pq^{k-1}$ avec $q = 1 - p$. On note dans ce cas $X \sim \mathcal{G}(p)$.

PROPOSITION 3.4 — Soit $X \sim \mathcal{G}(p)$. Alors pour tout $(m, n) \in (\mathbb{N}^*)^2$, $\mathbb{P}(X > m + n \mid X > n) = \mathbb{P}(X > m)$.

Ce résultat traduit le fait qu'une loi géométrique est *sans mémoire* : après n expériences les variables $X - n$ et X suivent la même loi : les expériences passées n'influencent pas sur les succès futurs. C'est la raison pour laquelle le fait qu'un nombre ne soit pas sorti depuis longtemps au loto n'augmente pas la probabilité qu'il sorte au tirage suivant.

Exercice 8

Soit X une variable aléatoire sans mémoire à valeurs dans $\mathbb{N}^*$. Montrer que X suit une loi géométrique.

■ Loi binomiale

L'expérience type consiste à effectuer n fois une expérience de Bernoulli et à noter X le nombre de succès.

DÉFINITION. — Soit $n \in \mathbb{N}^*$ et $p \in]0, 1[$. On dit qu'une variable aléatoire réelle X suit une loi binomiale de paramètres (n, p) lorsque $X(\Omega) = \llbracket 0, n \rrbracket$ et pour tout $k \in \llbracket 0, n \rrbracket$, $\mathbb{P}(X = k) = \binom{n}{k} p^k q^{n-k}$ avec $q = 1 - p$. On note dans ce cas $X \sim \mathcal{B}(n, p)$.

■ Loi de Poisson

La dernière loi que nous allons définir est un peu différente des précédentes, dans le sens où elle ne correspond pas à la modélisation d'une expérience précise mais apparaît (dans un certain sens) comme limite des lois binomiales.

THÉORÈME 3.5 (loi des événements rares) — Soit (X_n) une suite de variables aléatoires réelles telle que pour tout $n \in \mathbb{N}$, $X_n \sim \mathcal{B}(n, p_n)$. On suppose $p_n \sim \frac{\lambda}{n}$ avec $\lambda > 0$. Alors pour tout $k \in \mathbb{N}$, $\lim_{n \rightarrow +\infty} \mathbb{P}(X_n = k) = \frac{\lambda^k}{k!} e^{-\lambda}$.

DÉFINITION. — Soit $\lambda > 0$. On dit qu'une variable aléatoire réelle X suit une loi de Poisson de paramètre λ lorsque

$$X(\Omega) = \mathbb{N} \text{ et pour tout } k \in \mathbb{N}, \mathbb{P}(X = k) = \frac{\lambda^k}{k!} e^{-\lambda}. \text{ On note dans ce cas } X \sim \mathcal{P}(\lambda).$$

Remarque. Concrètement, ce résultat affirme que si des événements indépendants ont une très faible probabilité d'apparition, leur distribution, qui suit en principe une loi binomiale, est dans la pratique très voisine d'une loi de Poisson. On estime souvent qu'on peut utiliser l'approximation de $\mathcal{B}(n, p)$ par $\mathcal{P}(\lambda)$ (avec $\lambda = np$) dès lors que $n \geq 50$ et $np < 10$. Dans le cadre de cette approximation les calculs numériques s'en trouvent grandement simplifiés.

Exemple. Un central téléphonique possède 5 lignes. On estime à $n = 1\,200$ le nombre de personnes susceptibles d'appeler le standard sur une journée de huit heures, les appels étant répartis uniformément durant la journée et d'une durée de deux minutes en moyenne.

On souhaite calculer la probabilité que le standard soit saturé à un instant donné. Pour cela, on note X la variable aléatoire égale au nombre de personnes en train de téléphoner à un instant donné et on cherche à

$$\text{calculer } \mathbb{P}(X > 5) = 1 - \sum_{k=0}^5 \mathbb{P}(X = k).$$

Un appel au standard à un instant donné est une éventualité de probabilité $p = \frac{1}{8 \times 30} = \frac{1}{240}$. La variable aléatoire X suit donc une loi binomiale de paramètres (n, p) , et on est dans le cadre de l'approximation par une loi de Poisson de paramètre $\lambda = np = 5$. Effectuons le calcul avec ces deux lois :

```
from scipy.stats import binom, poisson

print(1-sum([binom.pmf(k, 1200, 1/240) for k in range(6)]))
print(1-sum([poisson.pmf(k, 5) for k in range(6)]))
```

```
0.384039090245462
0.38403934516693705
```

Les deux formules donnent effectivement des réponses très proches : de l'ordre de 38,4%.

Exercice 9

Soit X une variable aléatoire suivant une loi de Poisson de paramètre $\lambda > 0$. Est-il plus probable que la valeur de X soit paire ou impaire ?

3.3 Couple de variables aléatoires

DÉFINITION. — Si X et Y sont deux variables aléatoires sur un même espace probabilisable $(\Omega, \mathcal{A})$, on note (X, Y) la variable aléatoire $\omega \mapsto (X(\omega), Y(\omega))$. On appelle loi conjointe de X et de Y la loi de (X, Y) , autrement dit la loi $\mathbb{P}_{(X,Y)}$ définie par :

$$\forall (x, y) \in X(\Omega) \times Y(\Omega), \quad \mathbb{P}_{(X,Y)}(x, y) = \mathbb{P}(X = x \text{ et } Y = y).$$

À l'inverse, si (X, Y) est un couple de variables aléatoires, on appelle lois marginales de (X, Y) les lois de X et de Y .

Connaissant la loi conjointe de X et Y il est facile de retrouver les lois marginales :

$$\mathbb{P}(X = x) = \sum_{y \in Y(\Omega)} \mathbb{P}(X = x \text{ et } Y = y) \quad \text{et} \quad \mathbb{P}(Y = y) = \sum_{x \in X(\Omega)} \mathbb{P}(X = x \text{ et } Y = y).$$

À l'inverse, la connaissance des lois marginales ne permet pas en général de déterminer la loi conjointe, car en général les événements $\{X = x\}$ et $\{Y = y\}$ n'ont aucune raison d'être indépendants. C'est la raison pour laquelle on adopte la définition suivante :

DÉFINITION. — Deux variables aléatoires X et Y sur un même espace probabilisable $(\Omega, \mathcal{A})$ sont dites indépendantes lorsque pour tout $x \in X(\Omega)$ et tout $y \in Y(\Omega)$ les événements $\{X = x\}$ et $\{Y = y\}$ sont indépendants. On a dans ce cas : $\mathbb{P}(X = x \text{ et } Y = y) = \mathbb{P}(X = x) \cdot \mathbb{P}(Y = y)$. L'indépendance des deux variables aléatoires X et Y sera notée $X \perp\!\!\!\perp Y$.

PROPOSITION 3.6 — Soient X et Y deux variables aléatoires indépendantes d'un même espace probabilisable $(\Omega, \mathcal{A})$. Alors pour toutes parties A dans $X(\Omega)$ et B dans $Y(\Omega)$ on a : $\mathbb{P}(X \in A \text{ et } Y \in B) = \mathbb{P}(X \in A) \cdot \mathbb{P}(Y \in B)$.

PROPOSITION 3.7 — Soient X et Y deux variables aléatoires indépendantes d'un même espace probabilisable $(\Omega, \mathcal{A})$, et f et g deux fonctions de $\mathbb{R}$ dans $\mathbb{R}$. Alors les variables aléatoires $f(X)$ et $g(Y)$ sont indépendantes. Autrement dit,

$$X \perp\!\!\!\perp Y \implies f(X) \perp\!\!\!\perp g(Y)$$

Exercice 10

Soient $X \sim \mathcal{P}(\lambda)$ et $Y \sim \mathcal{P}(\mu)$ deux variables aléatoires indépendantes. Montrer que $X + Y \sim \mathcal{P}(\lambda + \mu)$.

Remarque. Lorsque deux variables X et Y ne sont pas indépendantes, on utilise une probabilité conditionnelle pour calculer la probabilité de l'événement $\{X = x \text{ et } Y = y\}$: $\mathbb{P}(X = x \text{ et } Y = y) = \mathbb{P}(X = x \mid Y = y) \cdot \mathbb{P}(Y = y)$.

Exercice 11

Soit X une variable aléatoire qui suit une loi de Poisson de paramètre λ , Y une variable aléatoire qui, lorsque $X = n$, suit une loi binomiale $\mathcal{B}(n, p)$. On pose enfin $Z = X - Y$.

Déterminer les lois de Y et de Z . Les variables Y et Z sont-elles indépendantes ?

Indépendance mutuelle

DÉFINITION. — Soit $(X_i)_{i \in I}$ une famille finie ou dénombrable de variables aléatoires d'un même espace probabilisable $(\Omega, \mathcal{A})$. On dit que ces variables sont mutuellement indépendantes si et seulement si pour toute p -liste $(i_1, \dots, i_p) \in I^p$ d'indices deux-à-deux distincts, et toute p -liste $(x_{i_1}, \dots, x_{i_p}) \in X_{i_1}(\Omega) \times \dots \times X_{i_p}(\Omega)$, les événements $\{X_{i_k} = x_{i_k}\}$ sont indépendants.

À l'instar de l'indépendance d'une famille finie ou dénombrable d'événements, cette définition est particulièrement malcommode à vérifier. En particulier, on notera qu'il n'est pas équivalent de se contenter de vérifier que les variables sont deux-à-deux indépendantes.

Exemple. Si $(X_k)_{1 \leq k \leq n}$ est une famille finie de variables aléatoires mutuellement indépendantes suivant toutes la même loi de Bernoulli de paramètre p , alors $S_n = \sum_{k=1}^n X_k$ suit une loi binomiale de paramètre (n, p) .

Plus généralement, nous rencontrerons fréquemment des familles de variables aléatoires $(X_n)_{n \in \mathbb{N}}$ indépendantes suivant toutes la même loi. Une telle suite de variables aléatoires sera dite *identiquement distribuée*.

Pour finir, nous admettrons le résultat suivant :

THÉORÈME 3.8 (lemme des coalitions) — Soient $X_1, \dots, X_n$ une famille de n variables aléatoires mutuellement indépendantes, et $p \in \llbracket 1, n \rrbracket$. Soit $f : \mathbb{R}^p \rightarrow \mathbb{R}$ et $g : \mathbb{R}^{n-p} \rightarrow \mathbb{R}$ deux fonctions. Alors les variables $f(X_1, \dots, X_p)$ et $g(X_{p+1}, \dots, X_n)$ sont indépendantes.

On peut bien entendu étendre ce résultat à plus de deux coalitions.

3.4 Espérance

Lorsque $X(\Omega)$ est fini, l'espérance d'une variable aléatoire réelle est la moyenne des valeurs qu'elle est susceptible de prendre pondérées par la probabilité d'apparition de ces valeurs : $\mathbb{E}(X) = \sum_{x \in X(\Omega)} x \mathbb{P}(X = x)$.

Lorsque $X(\Omega)$ est infini, nous avons vu dans la première partie de ce chapitre que pour pouvoir donner un sens à cette expression, il fallait pouvoir s'assurer que cette expression ne dépend pas de l'ordre d'indexation choisi pour $X(\Omega)$. Ceci nous conduit à la :

DÉFINITION. — On dit qu'une variable aléatoire réelle et discrète X est d'espérance finie lorsque la famille de nombres réels $(x \mathbb{P}(X = x))_{x \in X(\Omega)}$ est sommable, et on appelle dans ce cas on appelle espérance de X la quantité

$$\mathbb{E}(X) = \sum_{x \in X(\Omega)} x \mathbb{P}(X = x).$$

Remarque. Lorsqu'on décrit par compréhension l'ensemble $X(\Omega) = \{x_n \mid n \in \mathbb{N}\}$, X admet une espérance si et seulement si la série $\sum x_n \mathbb{P}(X = x_n)$ est absolument convergente.

Remarque. Lorsque la variable aléatoire X est à valeurs positives, nous avons vu que l'on pouvait noter $\sum_{x \in X(\Omega)} x \mathbb{P}(X = x) = +\infty$ lorsque cette famille n'est pas sommable. On notera alors $\mathbb{E}(X) = +\infty$ dans ce cas de figure (attention, ceci n'est pas valable lorsque X n'est pas à valeurs positives).

Exemple. Considérons une fois de plus le problème du lancer de pièce jusqu'à obtenir Face. Nous avons montré que si X désigne la variable aléatoire qui compte le nombre de lancers nécessaires nous avons $\mathbb{P}(X = n) = p(1-p)^{n-1}$ où p désigne la probabilité d'obtenir un Face lors d'un lancer.

Puisque $1-p \in]0, 1[$ la série $\sum_{n \geq 1} n(1-p)^{n-1}$ converge donc X est d'espérance finie, et

$$\mathbb{E}(X) = \sum_{n=1}^{+\infty} np(1-p)^{n-1} = \frac{p}{(1-(1-p))^2} = \frac{1}{p}$$

PROPOSITION 3.9 — Soit X une variable aléatoire presque sûrement bornée (autrement dit, il existe $M > 0$ et que $\mathbb{P}(|X| \leq M) = 1$). Alors X est d'espérance finie.

Notons qu'il existe une formule équivalente pour l'espérance d'une variable aléatoire à valeurs dans $\mathbb{N} \cup \{+\infty\}$:

PROPOSITION 3.10 — Si X est une variable aléatoire à valeurs dans $\mathbb{N} \cup \{+\infty\}$ alors $\mathbb{E}(X) = \sum_{n=1}^{+\infty} \mathbb{P}(X \geq n)$.

Les principaux résultats de l'espérance sont les suivants :

THÉORÈME 3.11 (de transfert) — Si f une application à valeurs réelles définie sur $X(\Omega)$, alors $f(X)$ est d'espérance finie si et seulement si la famille $(f(x)\mathbb{P}(X=x))_{x \in X(\Omega)}$ est sommable, et dans ce cas, $\mathbb{E}(f(X)) = \sum_{x \in X(\Omega)} f(x)\mathbb{P}(X=x)$.

COROLLAIRE — X est d'espérance finie il en est de même de $|X|$, et $|\mathbb{E}(X)| \leq \mathbb{E}(|X|)$.

PROPOSITION 3.12 — Soient X et Y deux variables aléatoires d'espérances finies. Alors :

- (i) pour tout $\lambda \in \mathbb{R}$, $\lambda X + Y$ est d'espérance finie, et $\mathbb{E}(\lambda X + Y) = \lambda \mathbb{E}(X) + \mathbb{E}(Y)$ (linéarité de l'espérance);
- (ii) et si, de plus, X et Y sont indépendantes, alors XY est d'espérance finie et $\mathbb{E}(XY) = \mathbb{E}(X)\mathbb{E}(Y)$.

COROLLAIRE — Soient X et Y deux variables aléatoires d'espérances finies. Alors :

- (i) si X est à valeurs positives, $\mathbb{E}(X) \geq 0$ (positivité de l'espérance);
- (ii) si $X \leq Y$ alors $\mathbb{E}(X) \leq \mathbb{E}(Y)$ (croissance de l'espérance).

■ Espérances des lois usuelles

Toutes les lois que nous avons étudiées à la section 3.2 sont d'espérance finie, et la valeur de leur espérance doit être connue.

Loi uniforme Si $X \sim \mathcal{U}(n)$, alors $\mathbb{E}(X) = \frac{n+1}{2}$.

Loi de Bernoulli Si $X \sim \mathcal{B}(p)$, alors $\mathbb{E}(X) = p$.

Loi géométrique Si $X \sim \mathcal{G}(p)$ alors $\mathbb{E}(X) = \frac{1}{p}$.

Loi binomiale Si $X \sim \mathcal{B}(n, p)$ alors $\mathbb{E}(X) = np$.

Loi de Poisson Si $X \sim \mathcal{P}(\lambda)$ alors $\mathbb{E}(X) = \lambda$.

Remarque. La loi binomiale étant la somme de n loi de Bernoulli (indépendantes) on a bien $\mathbb{E}(\mathcal{B}(n, p)) = n \times \mathbb{E}(\mathcal{B}(p))$.

3.5 Variance et écart type

DÉFINITION. — On dit qu'une variable aléatoire réelle X possède un moment d'ordre 1 lorsque X est d'espérance finie, et un moment d'ordre 2 lorsque X^2 est d'espérance finie.

THÉORÈME 3.13 — Si la variable aléatoire X possède un moment d'ordre 2 alors X possède un moment d'ordre 1, et dans ce cas, on appelle variance de X la quantité $\mathbb{V}(X) = \mathbb{E}((X - \mathbb{E}(X))^2)$, et écart type la quantité $\sigma(X) = \sqrt{\mathbb{V}(X)}$.

PROPOSITION 3.14 (Formule de Koenig-Huyghens) — Lorsque X possède un moment d'ordre 2,

$$\mathbb{V}(X) = \mathbb{E}(X^2) - \mathbb{E}(X)^2.$$

PROPOSITION 3.15 — Si a et b sont deux réels et X une variable aléatoire réelle possédant un moment d'ordre 2, alors $\mathbb{V}(aX + b) = a^2 \mathbb{V}(X)$ et $\sigma(aX + b) = |a| \sigma(X)$.

THÉORÈME 3.16 — Si X et Y sont deux variables aléatoires indépendantes admettant un moment d'ordre 2, il en est de même de $X + Y$, et $\mathbb{V}(X + Y) = \mathbb{V}(X) + \mathbb{V}(Y)$.

Exercice 12

Soit X une variable aléatoire possédant un moment d'ordre 2. Quelle est la valeur minimale de la fonction $t \mapsto \mathbb{E}((X - t)^2)$?

■ Variance des lois usuelles

Toutes les lois que nous avons étudiées à la section 3.2 possèdent une variance (à connaître) :

Loi uniforme Si $X \sim \mathcal{U}(n)$, alors $\mathbb{V}(X) = \frac{n^2 - 1}{12}$.

Loi de Bernoulli Si $X \sim \mathcal{B}(p)$, alors $\mathbb{V}(X) = pq = p(1 - p)$.

Loi géométrique Si $X \sim \mathcal{G}(p)$ alors $\mathbb{V}(X) = \frac{q}{p^2} = \frac{1 - p}{p^2}$.

Loi binomiale Si $X \sim \mathcal{B}(n, p)$ alors $\mathbb{V}(X) = npq$.

Loi de Poisson Si $X \sim \mathcal{P}(\lambda)$ alors $\mathbb{V}(X) = \lambda$.

■ Moment d'une variable aléatoire

Espérance et variance se généralisent avec la notion de *moment* : étant donné un entier $r \in \mathbb{N}$, on dit qu'une variable aléatoire réelle X possède un moment d'ordre r lorsque $\mathbb{E}(X^r)$ existe, et un moment centré d'ordre r lorsque $\mathbb{E}((X - \mathbb{E}(X))^r)$ existe. Ainsi, l'espérance est un moment d'ordre 1 et la variance un moment centré d'ordre 2.

Remarque (Variable centrée réduite). Si X est une variable aléatoire possédant un moment d'ordre 2, la variable aléatoire $Y = \frac{X - \mathbb{E}(X)}{\sigma(X)}$ possède une espérance nulle (on dit qu'elle est *centrée*) et un écart type égal à 1 (on dit qu'elle est *réduite*).

L'application $(X, Y) \mapsto \mathbb{E}(XY)$ est une application bilinéaire, symétrique et positive. En conséquence de quoi il est possible d'établir le résultat suivant :

THÉORÈME 3.17 (Inégalité de Cauchy-Schwarz) — Si X et Y possèdent des moments d'ordre 2, alors XY possède un moment d'ordre 1, et $\mathbb{E}(XY)^2 \leq \mathbb{E}(X^2)\mathbb{E}(Y^2)$.

Remarque. Il y a égalité dans l'inégalité de Cauchy-Schwarz si et seulement s'il existe $(\lambda, \mu) \in \mathbb{R}^2$ tel que $\lambda X + \mu Y$ est quasi-sûrement nul, autrement dit lorsque $\mathbb{P}(\lambda X + \mu Y = 0) = 1$.

3.6 Covariance

Nous avons démontré à la proposition 3.12 que lorsque X et Y sont deux variables aléatoires indépendantes, $\mathbb{E}(XY) = \mathbb{E}(X)\mathbb{E}(Y)$. Lorsque X et Y ne sont pas indépendantes, on peut considérer que la quantité $\mathbb{E}(XY) - \mathbb{E}(X)\mathbb{E}(Y)$ mesure le « défaut d'indépendance » de ces deux variables. Pour des raisons pratiques (liées à l'inégalité de Cauchy-Schwarz, voir plus loin), nous allons introduire cette quantité sous une forme légèrement différente. En effet,

$$\begin{aligned}\mathbb{E}((X - \mathbb{E}(X))(Y - \mathbb{E}(Y))) &= \mathbb{E}(XY - \mathbb{E}(X)Y - X\mathbb{E}(Y) + \mathbb{E}(X)\mathbb{E}(Y)) = \mathbb{E}(XY) - \mathbb{E}(X)\mathbb{E}(Y) - \mathbb{E}(X)\mathbb{E}(Y) + \mathbb{E}(X)\mathbb{E}(Y) \\ &= \mathbb{E}(XY) - \mathbb{E}(X)\mathbb{E}(Y).\end{aligned}$$

Ceci conduit à la définition suivante :

DÉFINITION. — Soient X et Y deux variables aléatoires réelles. Sous réserve d'existence on appelle covariance de X et de Y la quantité $\boxed{\text{cov}(X, Y) = \mathbb{E}((X - \mathbb{E}(X))(Y - \mathbb{E}(Y)))}$.

Le théorème 3.17 nous permet d'énoncer :

PROPOSITION 3.18 — Si X et Y sont deux variables aléatoires réelles possédant un moment d'ordre 2 alors $\text{cov}(X, Y)$ existe.

et le calcul réalisé ci-dessus nous permet d'affirmer :

PROPOSITION 3.19 — Lorsque X et Y sont deux variables aléatoires réelles indépendantes possédant un moment d'ordre 2, alors $\text{cov}(X, Y) = 0$. On dira que X et Y ne sont pas corrélées.

Attention. La réciproque de ce résultat est fautive : deux variables aléatoires peuvent ne pas être corrélées sans pour autant être indépendantes.

PROPOSITION 3.20 (propriétés de la covariance) — Soient X , Y et Z trois variables aléatoires réelles possédant des moments d'ordre 2. Alors :

- $\text{cov}(X, Y) = \mathbb{E}(XY) - \mathbb{E}(X)\mathbb{E}(Y)$;
- $\text{cov}(X, Y) = \text{cov}(Y, X)$;
- $\text{cov}(X, 1) = 0$;
- $\forall (a, b) \in \mathbb{R}^2, \text{cov}(X, aY + bZ) = a \text{cov}(X, Y) + b \text{cov}(X, Z)$.

THÉORÈME 3.21 — Soient X et Y deux variables aléatoires réelles possédant des moments d'ordre 2. Alors

$$\mathbb{V}(X + Y) = \mathbb{V}(X) + 2 \text{cov}(X, Y) + \mathbb{V}(Y).$$

En particulier, lorsque ces deux variables aléatoires sont indépendantes, on retrouve le fait que $\mathbb{V}(X + Y) = \mathbb{V}(X) + \mathbb{V}(Y)$.

Remarque. Cette formule se généralise au cas de n variables aléatoires : $\mathbb{V}\left(\sum_{k=1}^n X_k\right) = \sum_{k=1}^n \mathbb{V}(X_k) + 2 \sum_{i < j} \text{cov}(X_i, X_j)$.

En particulier on retiendra le :

COROLLAIRE — Lorsque $X_1, \dots, X_n$ sont des variables aléatoires possédant des moments d'ordre 2 et deux-à-deux indépendantes, $\mathbb{V}(X_1 + \dots + X_n) = \mathbb{V}(X_1) + \dots + \mathbb{V}(X_n)$.

3.7 Inégalités de concentration

Dans la théorie des probabilités, les inégalités de concentration fournissent des bornes sur la probabilité qu'une variable aléatoire dévie d'une certaine valeur (généralement l'espérance de cette variable aléatoire).

THÉORÈME 3.22 (inégalité de Markov) — Soit X une variable aléatoire positive possédant un moment d'ordre 1, et

$$a > 0. \text{ Alors } \mathbb{P}(X \geq a) \leq \frac{\mathbb{E}(X)}{a}.$$

THÉORÈME 3.23 (inégalité de Bienaymé-Tchebychev) — Soit X une variable aléatoire réelle admettant un moment

$$\text{d'ordre 2, et } \alpha > 0. \text{ Alors } \mathbb{P}(|X - \mathbb{E}(X)| \geq \alpha) \leq \frac{\mathbb{V}(X)}{\alpha^2} = \frac{\sigma(X)^2}{\alpha^2}.$$

Que signifie cette inégalité? La probabilité calculée mesure le risque de s'écarter de l'espérance d'une quantité supérieure à α . La majorant obtenu montre que plus l'écart type est faible, plus ce risque est négligeable. Ainsi, un écart type faible caractérise une faible dispersion autour de l'espérance. À l'inverse, un écart type important dénote une grande dispersion des valeurs.

■ Loi faible des grands nombres

Le théorème que nous allons énoncer ensuite va justifier la démarche expérimentale qu'on utilise pour estimer empiriquement une espérance : on réalise un grand nombre d'expérience (en général par le biais d'une simulation numérique) puis on calcule la moyenne des valeurs qu'a prise la variable aléatoire X . Par exemple, pour estimer l'espérance d'une loi géométrique de paramètre $1/10$ on réalise le script suivant :

```
def experience():
    s = 1
    while True:
        if rd.random() < .1:
            return s
        s += 1

n = 100000 # nombre d'expériences
v = 0      # somme des variables aléatoires

for _ in range(n):
    v += experience()
print(v / n)
```

10.02955

Nous obtenons effectivement une valeur proche de l'espérance théorique égale à 10. Le théorème qui suit donne une justification à cet état de fait :

THÉORÈME 3.24 (loi faible des grands nombres) — Soit $(X_n)_{n \geq 1}$ une suites de variables aléatoires deux-à-deux indépendantes et de même loi admettant un moment d'ordre 2. On pose $S_n = X_1 + X_2 + \dots + X_n$. Alors pour tout $\epsilon > 0$,

$$\mathbb{P}\left(\left|\frac{1}{n}S_n - \mathbb{E}(X)\right| \geq \epsilon\right) \leq \frac{\mathbb{V}(X)}{n\epsilon^2}.$$

Remarque. Avec les mêmes hypothèses on en déduit que : $\lim_{n \rightarrow +\infty} \mathbb{P}\left(\left|\frac{1}{n}S_n - \mathbb{E}(X)\right| \geq \epsilon\right) = 0$.

En d'autres termes, plus on réalise un grand nombre d'expériences, plus le risque que la moyenne s'écarte de l'espérance de plus de ϵ est faible.

Exemple. Dans l'exemple numérique ci-dessus nous avons pris $n = 100\,000$ et nous avons $\mathbb{E}(X) = 10$ et $\mathbb{V}(X) = 90$. Pour $\epsilon = 0,1$ nous avons $\frac{\mathbb{V}(X)}{n\epsilon^2} = 0,09$ donc il y a plus de 91% de chance que le résultat obtenu diffère de l'espérance théorique de moins de 0,1.

3.8 Séries génératrices

DÉFINITION. — Soit X une variable aléatoire à valeurs dans $\mathbb{N}$. On appelle série génératrice de X la série entière $G_X(t) = \mathbb{E}(t^X) = \sum_{n \in \mathbb{N}} \mathbb{P}(X = n)t^n$.

Pourquoi s'intéresser à cette série entière? Nous savons que les coefficients d'une série entière de rayon de convergence $R > 0$ sont définis de manière unique, aussi pouvons-nous affirmer que si $R > 0$ la série génératrice d'une variable aléatoire caractérise cette dernière. On peut donc espérer utiliser la souplesse d'utilisation des séries entières pour calculer plus facilement certaines caractéristiques de X , telles l'espérance ou la variance.

LEMME — La série génératrice G_X de X est au moins définie sur $[-1, 1]$.

COROLLAIRE — Si deux variables aléatoires ont même série génératrice sur $] -1, 1[$ alors ces deux variables aléatoires suivent la même loi.

Supposons maintenant $R > 1$. G_X est de classe $\mathcal{C}^\infty$ sur $] -R, R[$ et $G'_X(t) = \sum_{n=1}^{+\infty} n\mathbb{P}(X = n)t^{n-1}$. On voit immédiatement qu'en posant $t = 1$ on obtient $G'_X(1) = \mathbb{E}(X)$. Nous admettrons que ce résultat reste vrai lorsque $R = 1$ à condition que G_X soit dérivable en 1, ce qui nous permet d'énoncer le

THÉORÈME 3.25 — X admet un moment d'ordre 1 (une espérance) si et seulement si G_X est dérivable en 1, et dans ce cas, $\mathbb{E}(X) = G'_X(1)$.

Voyons comment obtenir la variance. Si on suppose toujours $R > 1$ nous avons $G''_X(1) = \sum_{n=1}^{+\infty} n(n-1)\mathbb{P}(X = n)$.

$$\begin{aligned} \text{Par ailleurs, } \mathbb{V}(X) &= \mathbb{E}(X^2) - \mathbb{E}(X)^2 = \sum_{n=1}^{+\infty} n^2\mathbb{P}(X = n) - \mathbb{E}(X)^2 = \sum_{n=1}^{+\infty} n(n-1)\mathbb{P}(X = n) + \sum_{n=1}^{+\infty} n\mathbb{P}(X = n) - \mathbb{E}(X)^2 \\ &= G''_X(1) + G'_X(1) - G'_X(1)^2. \end{aligned}$$

Nous admettrons que sous réserve d'existence cette formule reste vraie lorsque $R = 1$, ce qui donne le

THÉORÈME 3.26 — X admet un moment d'ordre 2 si et seulement si G_X est deux fois dérivable en 1, et dans ce cas, $\mathbb{V}(X) = G''_X(1) + G'_X(1) - G'_X(1)^2$.

■ Séries génératrices les lois usuelles

Loi uniforme Si $X \sim \mathcal{U}(n)$, alors $G_X(t) = \frac{1}{n}(t + t^2 + \dots + t^n)$.

Loi de Bernoulli Si $X \sim \mathcal{B}(p)$, alors $G_X(t) = pt + q$ avec $q = 1 - p$.

Loi géométrique Si $X \sim \mathcal{G}(p)$ alors $G_X(t) = \frac{pt}{1 - qt}$ avec $q = 1 - p$.

Loi binomiale Si $X \sim \mathcal{B}(n, p)$ alors $G_X(t) = (pt + q)^n$ avec $q = 1 - p$.

Loi de Poisson Si $X \sim \mathcal{P}(\lambda)$ alors $G_X(t) = e^{\lambda t - \lambda}$.

Exercice 13

À l'aide des séries génératrices ci-dessus, retrouver l'espérance et la variance des lois usuelles.

■ Série génératrice d'une somme de deux variables aléatoires indépendantes

Considérons pour finir deux variables aléatoires indépendantes X et Y à valeurs dans $\mathbb{N}$.

Nous avons $G_X(t) = \sum_{n=0}^{+\infty} \mathbb{P}(X = n)t^n$ et $G_Y(t) = \sum_{n=0}^{+\infty} \mathbb{P}(Y = n)t^n$ et

$$G_{X+Y}(t) = \sum_{n=0}^{+\infty} \mathbb{P}(X + Y = n)t^n = \sum_{n=0}^{+\infty} \left(\sum_{k=0}^n \mathbb{P}(X = k \text{ et } Y = n - k) \right) t^n = \sum_{n=0}^{+\infty} \left(\sum_{k=0}^n \mathbb{P}(X = k)\mathbb{P}(Y = n - k) \right) t^n.$$

On reconnaît un produit de Cauchy donc

$$G_{X+Y}(t) = \left(\sum_{n=0}^{+\infty} \mathbb{P}(X = n)t^n \right) \cdot \left(\sum_{n=0}^{+\infty} \mathbb{P}(Y = n)t^n \right) = G_X(t)G_Y(t).$$

Nous avons prouvé le :

THÉORÈME 3.27 — Si X et Y sont deux variables aléatoires indépendantes à valeur entières, alors

$$\forall t \in]-1, 1[, \quad \boxed{G_{X+Y}(t) = G_X(t)G_Y(t).}$$

Exemple.

- Si $X_i \sim \mathcal{B}(p)$ ($1 \leq i \leq n$) sont des variables aléatoires indépendantes et S leur somme, alors $S \sim \mathcal{B}(n, p)$.
- Si $X \sim \mathcal{B}(m, p)$ et $Y \sim \mathcal{B}(n, p)$ sont deux variables aléatoires indépendantes, alors $X + Y \sim \mathcal{B}(m + n, p)$.
En effet, $(pt + q)^m \cdot (pt + q)^n = (pt + q)^{m+n}$.
- Si $X \sim \mathcal{P}(\lambda)$ et $Y \sim \mathcal{P}(\mu)$ sont deux variables aléatoires indépendantes, alors $X + Y \sim \mathcal{P}(\lambda + \mu)$.
En effet, $e^{\lambda t - \lambda} \cdot e^{\mu t - \mu} = e^{(\lambda + \mu)t - (\lambda + \mu)}$.

Chapitre VIII

Intégration

L'intégration est un concept fondamental en mathématiques, issu du calcul des aires. À ce titre, on peut considérer que ses racines se trouvent parmi les premiers calculs d'aires et de volumes de l'antiquité. Mais c'est à Leibniz, au XVII^e siècle qu'on doit les fondements de la théorie de l'intégration, en particulier par l'introduction d'un symbolisme reliant intégration et dérivation. C'est d'ailleurs lui qui est à l'origine du symbole $\int$. Il faut néanmoins attendre Riemann (en 1854) pour avoir une première théorie de l'intégration complète, c'est à dire une définition précise de ce qu'est une fonction *intégrable*. Par la suite, d'autres théories, plus élaborées, ont vu le jour, telles l'intégrale de Lebesgue (1902), ou encore l'intégrale de Kurzweil-Henstock (1950).

1. Intégration des fonctions continues par morceaux

1.1 Fonctions continues par morceaux

■ Subdivision d'un intervalle

DÉFINITION. — Une subdivision d'un segment $[a, b]$ est une suite finie $\sigma = (t_0, t_1, \dots, t_n)$ vérifiant :

$$a = t_0 < t_1 < \dots < t_{n-1} < t_n = b.$$

Le pas de la subdivision est le réel $p(\sigma) = \max\{t_{i+1} - t_i \mid i \in \llbracket 0, n-1 \rrbracket\}$. La subdivision est dite régulière lorsque

$$\forall i \in \llbracket 0, n-1 \rrbracket, t_{i+1} - t_i = p(\sigma), \text{ soit encore : } \forall i \in \llbracket 0, n \rrbracket, \boxed{t_i = a + i \frac{b-a}{n}}.$$

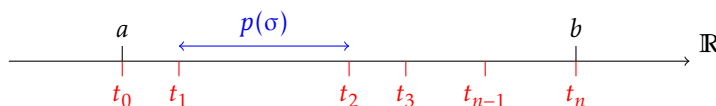


FIGURE 1 – Une subdivision du segment $[a, b]$.

DÉFINITION. — Une fonction numérique $f : [a, b] \rightarrow \mathbb{K}$ est dite continue par morceaux s'il existe une subdivision $\sigma = (t_0, \dots, t_n)$ de $[a, b]$ telle que f soit sur tous les intervalles $]t_k, t_{k+1}[$ la restriction d'une fonction continue sur $[t_k, t_{k+1}]$. Une telle subdivision sera dite adaptée à f .

Remarque. Concrètement, ceci signifie que pour tout $i \in \llbracket 0, n-1 \rrbracket$, f possède une limite (finie) à droite en t_i et à gauche en t_{i+1} . Notons en outre que la fonction f peut être continue en t_i , mais qu'à l'inverse toutes les discontinuités de f (qui doivent être en nombre fini) font partie des points de la subdivision σ (illustration figure 2).

Remarque. On dit qu'une subdivision σ' est *plus fine* qu'une subdivision σ lorsque σ est une sous-suite de σ' , en conséquence de quoi toute subdivision plus fine qu'une subdivision adaptée à f est encore adaptée à f .

Il est alors intéressant d'observer que si σ et σ' sont deux subdivisions quelconques de $[a, b]$, alors $\sigma \cup \sigma'$ est une subdivision à la fois plus fine que σ et que σ' .

Qu'est ce qui peut empêcher une fonction définie sur un segment d'être continue par morceaux ?

- Cette fonction peut présenter un point en lequel il n'y a pas de limite à gauche ou à droite ; c'est par exemple le cas de la fonction $f : x \mapsto \sin(1/x)$ sur le segment $[-1, 1]$, quelle que soit la valeur de $f(0)$;

- cette fonction peut présenter une limite infinie en un point ; c'est par exemple le cas de la fonction $f : x \mapsto \frac{1}{x}$ sur le segment $[-1, 1]$, quelle que soit la valeur de $f(0)$;

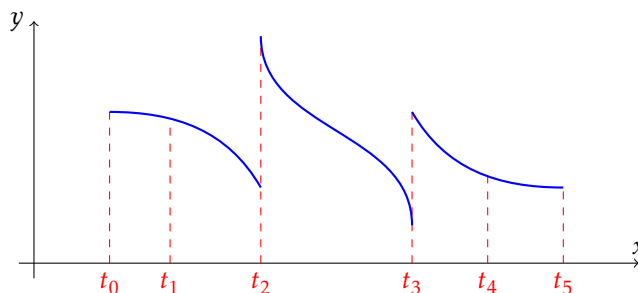


FIGURE 2 – Un exemple de fonction continue par morceaux et d’une subdivision (non minimale) qui lui est adaptée.

– cette fonction peut présenter un nombre *infini* de discontinuités ; c’est par exemple le cas de la fonction $f : x \mapsto x \lfloor 1/x \rfloor$ prolongée par $f(0) = 1$, bien qu’elle possède en tout point une limite à gauche et à droite.

PROPOSITION 1.1 — Toute fonction continue par morceaux sur un segment est bornée.

THÉORÈME 1.2 — L’ensemble $\mathcal{C}_m^0([a, b], \mathbb{K})$ des fonctions continues par morceaux est un sous-espace vectoriel de l’espace $\mathcal{B}([a, b], \mathbb{K})$ des fonctions bornées sur $[a, b]$. De plus, si f et g sont continues par morceaux sur $[a, b]$, il en est de même de leur produit fg .

Remarque. Rappelons qu’une fonction $\phi : [a, b] \rightarrow \mathbb{K}$ est dite *en escalier* lorsqu’il existe une subdivision $\sigma = (t_0, t_1, \dots, t_n)$ telle que f soit constante sur chaque intervalle $]t_i, t_{i+1}[$, $0 \leq i \leq n-1$. Bien entendu, toute fonction en escalier sur $[a, b]$ est continue par morceaux sur cet intervalle, et par une preuve analogue à celle du théorème précédent on montre que les fonctions en escalier constituent un sous-espace vectoriel du $\mathbb{K}$ -espace vectoriel des fonctions continues par morceaux sur $[a, b]$.

Fonctions continues par morceaux sur un intervalle quelconque

DÉFINITION. — Soit I un intervalle quelconque. Une fonction $f : I \rightarrow \mathbb{K}$ est dite continue par morceaux lorsque pour tout segment $[a, b]$ inclus dans I , la restriction de f à $[a, b]$ est continue par morceaux sur $[a, b]$.

Exemple. La fonction $x \mapsto \lfloor x \rfloor$ est continue par morceaux sur $\mathbb{R}$: elle possède en tout point une limite finie à gauche et à droite et, bien que ses discontinuités soient en nombre infini, ne possède qu’un nombre *fini* de discontinuité sur tout segment. Pour ces mêmes raisons, la fonction $f : x \mapsto x \lfloor 1/x \rfloor$ est continue par morceaux sur $]0, 1]$. Elle n’est en revanche pas continue par morceaux sur $[0, 1]$, bien qu’elle soit prolongeable par continuité en 0 !

1.2 Intégrale sur un segment d’une fonction continue par morceaux

Ceci étant posé, définir la valeur de l’intégrale d’une fonction continue par morceaux sur un segment ne pose pas de problème :

DÉFINITION. — Soit $f : [a, b] \rightarrow \mathbb{K}$ une fonction continue par morceaux, et $\sigma = (t_0 = a, t_1, \dots, t_n = b)$ une subdivision adaptée à f . L’intégrale de f sur $[a, b]$ est alors la quantité :

$$\int_{[a, b]} f = \int_a^b f(t) dt = \sum_{k=0}^{n-1} \int_{t_k}^{t_{k+1}} f(t) dt.$$

Remarque. Pour valider cette définition, il faut montrer que cette valeur ne dépend pas du choix de la subdivision subordonnée à f , mais ceci ne présente pas de difficulté.

Cas d'une fonction en escalier

Si $\phi : [a, b] \rightarrow \mathbb{K}$ est une fonction en escalier et σ une subdivision qui lui est adaptée, si $v_k \in \mathbb{K}$ désigne la valeur que prend ϕ sur l'intervalle $]t_k, t_{k+1}[$, alors :
$$\int_{[a,b]} \phi = \sum_{k=0}^{n-1} (t_{k+1} - t_k) v_k.$$

Lorsque $v_k \in \mathbb{R}_+$, cette quantité peut être interprétée comme l'aire délimitée par l'axe des abscisses et la fonction ϕ (illustration figure 3).

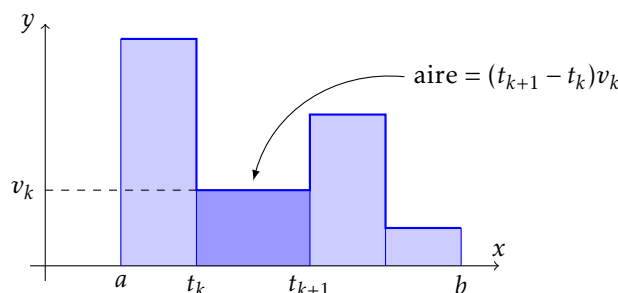


FIGURE 3 – Interprétation graphique de l'intégrale d'une fonction en escalier.

Nous admettrons que cette interprétation graphique reste pertinente pour une fonction continue par morceaux. Rappelons aussi que cette interprétation est à la base d'un résultat du cours de première année : la notion de somme de Riemann.

THÉORÈME 1.3 (Sommes de Riemann) — Si $f : [a, b] \rightarrow \mathbb{K}$ est une fonction continue sur le segment $[a, b]$ alors

$$\lim_{n \rightarrow +\infty} \frac{b-a}{n} \sum_{k=0}^{n-1} f\left(a + k \frac{b-a}{n}\right) = \int_a^b f(t) dt.$$

Exercice 1

Déterminer la limite suivante :
$$\lim_{n \rightarrow +\infty} \sum_{k=1}^n \frac{1}{\sqrt{n^2 + 2kn}}.$$

1.3 Propriétés de l'intégrale

Les propriétés suivantes, établies en première année dans le cas des fonctions continues, s'étendent au cas des fonctions continues par morceaux.

PROPOSITION 1.4 — L'application qui à une fonction continue par morceaux associe son intégrale sur $[a, b]$ est linéaire : si f et g sont continues par morceaux sur $[a, b]$, alors
$$\int_{[a,b]} (\lambda f + g) = \lambda \int_{[a,b]} f + \int_{[a,b]} g.$$

PROPOSITION 1.5 — Soit $f : [a, b] \rightarrow \mathbb{K}$ une fonction continue par morceaux. Alors :
$$\left| \int_{[a,b]} f \right| \leq \int_{[a,b]} |f|.$$

Remarque. La proposition 1.5 appliquée à une fonction à valeurs réelles positives implique le résultat suivant, dite propriété de positivité de l'intégrale :

$$\text{si } f : [a, b] \rightarrow \mathbb{R}_+ \text{ est continue par morceaux, alors } \int_{[a,b]} f \geq 0.$$

Une conséquence immédiate de ce résultat est la propriété dite de croissance de l'intégrale :

si f et g sont deux fonctions réelles continues par morceaux sur $[a, b]$ et telles que $f \leq g$, alors
$$\int_{[a,b]} f \leq \int_{[a,b]} g.$$

COROLLAIRE — Si $f : [a, b] \rightarrow \mathbb{K}$ est continue par morceaux, alors : $\left| \int_{[a,b]} f \right| \leq |b - a| \cdot \|f\|_{\infty, [a,b]}$.

Enfin, sur le même sujet on rappellera un résultat important du cours de première année, mais qui ne s'applique pas aux fonctions continues par morceaux :

PROPOSITION 1.6 — Une fonction continue et à valeurs positives sur $[a, b]$ est nulle si et seulement si son intégrale est nulle.

1.4 Dérivation et intégration

Nous allons maintenant rappeler un résultat vu en première année, souvent connu comme le *théorème fondamental de l'analyse* (le théorème 1.7 de ce document). Ce dernier établit un lien entre intégration et dérivation (un résultat en général attribué à Isaac Newton). Plus précisément, il affirme que le calcul d'une intégrale d'une fonction continue se ramène à la recherche d'une primitive de cette fonction.

■ Primitives et intégrale d'une fonction continue

THÉORÈME 1.7 — Soit I un intervalle, $f : I \rightarrow \mathbb{K}$ une fonction continue, et $a \in I$. Pour tout réel $x \in I$, on note $F(x) = \int_a^x f(t) dt$. Alors F est de classe $\mathcal{C}^1$ sur I , et $F' = f$.

Ce théorème ramène le calcul d'une intégrale à la recherche d'une primitive. Commençons par rappeler la définition suivante :

DÉFINITION. — Soit f une fonction continue sur I . On dit qu'une application $g : I \rightarrow \mathbb{K}$ est une primitive de f lorsque g est de classe $\mathcal{C}^1$, et lorsqu'en tout point de I , $g'(x) = f(x)$.

Les primitives sur un intervalle sont définies « à une constante près » :

PROPOSITION 1.8 — Si g_1 et g_2 sont deux primitives de f , il existe une constante λ telle que $g_2 = g_1 + \lambda$.

Nous pouvons donc préciser le résultat du théorème 1.7 en énonçant le résultat suivant :

PROPOSITION 1.9 — Soit f une fonction continue sur I , et $a \in I$. On définit une fonction F sur I en posant : $\forall x \in I$, $F(x) = \int_a^x f(t) dt$. Alors F est l'unique primitive de f qui s'annule en a .

Voici enfin le résultat qui permet de calculer une intégrale en recherchant une primitive :

COROLLAIRE — Si f est continue et g une primitive quelconque de f , alors : $\forall (a, b) \in I^2$, $\int_a^b f(t) dt = g(b) - g(a)$.

Exercice 2

Soit $f : [0, +\infty[\rightarrow [0, +\infty[$ une application strictement croissante de classe $\mathcal{C}^1$, telle que $f(0) = 0$. Montrer, en appliquant le théorème 1.7, que :

$$\forall x \in [0, +\infty[, \quad xf(x) = \int_0^x f(t) dt + \int_0^{f(x)} f^{-1}(t) dt.$$

Donner une interprétation graphique de ce résultat.

1.5 Changement de variable et intégration par parties

THÉORÈME 1.10 (changement de variable) — Soit $f : I \rightarrow \mathbb{K}$ une fonction continue, $[\alpha, \beta]$ un segment et $\phi : [\alpha, \beta] \rightarrow I$ une fonction de classe $\mathcal{C}^1$. Alors :

$$\int_{\phi(\alpha)}^{\phi(\beta)} f(t) dt = \int_{\alpha}^{\beta} f(\phi(u)) \phi'(u) du.$$

Mise en œuvre pratique

Il y a deux façons d'utiliser cette formule : partir de l'expression de droite pour obtenir celle de gauche, ou procéder dans le sens contraire.

Utiliser la formule de la droite vers la gauche nécessite de reconnaître dans l'expression de l'intégrale à calculer $\int_{\alpha}^{\beta} g(u) du$ un terme de la forme $g(u) = f(\phi(u)) \phi'(u)$ (il faut donc « deviner » f et ϕ). Dans ce cas, il faut donc identifier les expressions $t = \phi(u)$ et $dt = \phi'(u) du$.

Exemple. Soit $I = \int_0^{\pi} \frac{\sin u}{3 + \cos^2 u} du$. On pose $t = \cos u$ de manière à avoir $dt = -\sin u du$. Ainsi,

$$I = \int_{\cos 0}^{\cos \pi} \frac{-dt}{3 + t^2} = \int_{-1}^1 \frac{dt}{3 + t^2} = \frac{1}{3} \int_{-1}^1 \frac{dt}{1 + (t/\sqrt{3})^2} = \frac{1}{\sqrt{3}} \left[\arctan\left(\frac{t}{\sqrt{3}}\right) \right]_{-1}^1 = \frac{2}{\sqrt{3}} \arctan \frac{1}{\sqrt{3}} = \frac{\pi}{3\sqrt{3}}.$$

Utiliser la formule de la gauche vers la droite est souvent utilisé pour faire apparaître une simplification dans l'expression initiale. Une fois posés $t = \phi(u)$ et $dt = \phi'(u) du$ il faut trouver des antécédents par ϕ des bornes de l'intégrale initiale.

Exemple. Soit $I = \int_0^1 \sqrt{1 - t^2} dt$. On pose $t = \sin u$ et $dt = \cos u du$. On choisit $\alpha = 0$ et $\beta = \frac{\pi}{2}$ pour avoir $\sin \alpha = 0$ et $\sin \beta = 1$. Ainsi :

$$I = \int_0^{\frac{\pi}{2}} \cos u \sqrt{1 - \sin^2 u} du = \int_0^{\frac{\pi}{2}} (\cos u)^2 du = \int_0^{\frac{\pi}{2}} \frac{1 - \cos(2u)}{2} du = \left[\frac{u}{2} - \frac{1}{4} \sin(2u) \right]_0^{\frac{\pi}{2}} = \frac{\pi}{4}.$$

THÉORÈME 1.11 (intégration par parties) — Soient f et g deux fonctions de classe $\mathcal{C}^1$ sur $[a, b]$. Alors :

$$\int_a^b f(t) g'(t) dt = \left[f(t) g(t) \right]_a^b - \int_a^b f'(t) g(t) dt.$$

Cette formule peut être représentée par le schéma suivant :

$$\begin{array}{r} + \left| \begin{array}{cc} f(t) & g'(t) \\ f'(t) & g(t) \end{array} \right| \\ - \end{array}$$

Le terme de gauche de la formule se retrouve sur la première ligne, le terme entre crochets sur la diagonale, et le reste intégral sur la dernière ligne :

$$\begin{array}{ccc} + \left| \begin{array}{cc} f(t) & g'(t) \\ f'(t) & g(t) \end{array} \right| & + \left| \begin{array}{cc} f(t) & g'(t) \\ f'(t) & \cancel{g(t)} \end{array} \right| & + \left| \begin{array}{cc} f(t) & g'(t) \\ \cancel{f'(t)} & g(t) \end{array} \right| \\ \int_a^b f(t) g'(t) dt & = & \left[f(t) g(t) \right]_a^b + \left(- \int_a^b f'(t) g(t) dt \right) \end{array}$$

L'intérêt de ce schéma est de permettre d'effectuer plusieurs intégrations par parties successives en une seule étape ; voici par exemple les schémas pour effectuer deux puis trois intégrations par parties successives, et les formules correspondantes :

$$\begin{array}{l}
+ \begin{vmatrix} f(t) & g''(t) \\ f'(t) & g'(t) \\ f''(t) & g(t) \end{vmatrix} \\
- \\
+
\end{array}
\int_a^b f(t)g''(t)dt = \left[f(t)g'(t) - f'(t)g(t) \right]_a^b + \int_a^b f''(t)g(t)dt$$

$$\begin{array}{l}
+ \begin{vmatrix} f(t) & g^{(3)}(t) \\ f'(t) & g''(t) \\ f''(t) & g'(t) \\ f^{(3)}(t) & g(t) \end{vmatrix} \\
- \\
+ \\
-
\end{array}
\int_a^b f(t)g^{(3)}(t)dt = \left[f(t)g''(t) - f'(t)g'(t) + f''(t)g(t) \right]_a^b - \int_a^b f^{(3)}(t)g(t)dt$$

Exercice 3

En effectuant autant d'intégrations par parties que nécessaire, calculer l'intégrale $\int_0^{\frac{\pi}{2}} t^3 \sin t \, dt$.

1.6 Formules de Taylor

Lorsque $f : [a, b] \rightarrow \mathbb{K}$ est de classe $\mathcal{C}^1$ sur $[a, b]$, on peut écrire : $f(b) - f(a) = \int_a^b f'(t) \, dt$.

Cette relation va avoir plusieurs conséquences, à commencer par le résultat suivant :

THÉORÈME 1.12 (Inégalité des accroissements finis) — Soit $f : [a, b] \rightarrow \mathbb{K}$ une fonction numérique de classe $\mathcal{C}^1$ sur $[a, b]$. On suppose l'existence d'un réel k tel que : $\forall t \in [a, b], |f'(t)| \leq k$. Alors :

$$|f(b) - f(a)| \leq k|b - a|$$

La généralisation de cette majoration va passer par plusieurs intégrations par parties successives. En effet, si on adopte le schéma suivant on obtient :

$$\begin{array}{l}
+ \begin{vmatrix} f'(t) & 1 \\ f''(t) & t - b \end{vmatrix} \\
- \\
+
\end{array}
f(b) = f(a) + \int_a^b f'(t) \, dt = f(a) + (b - a)f'(a) + \int_a^b (b - t)f''(t) \, dt$$

En réitérant ce procédé, ceci nous amène au théorème suivant :

PROPOSITION 1.13 (Formule de Taylor avec reste intégral) — Soit $f : I \rightarrow \mathbb{K}$ une fonction de classe $\mathcal{C}^{n+1}$, et $a \in I$.

Alors : $\forall x \in I, f(x) = \sum_{k=0}^n \frac{(x-a)^k}{k!} f^{(k)}(a) + \int_a^x \frac{(x-t)^n}{n!} f^{(n+1)}(t) \, dt$.

Remarque. $T_n : x \mapsto \sum_{k=0}^n \frac{(x-a)^k}{k!} f^{(k)}(a)$ est une fonction polynomiale, appelée *polynôme de Taylor* d'ordre n de f en a . C'est un polynôme dont les dérivées successives jusqu'au rang n coïncident au point a avec celles de f . La quantité $R_n(x) = \int_a^x \frac{(x-t)^n}{n!} f^{(n+1)}(t) \, dt$ est l'expression intégrale de l'erreur qu'on commet en approchant $f(x)$ par $T_n(x)$.

Pour majorer cette erreur, on utilise le résultat suivant :

THÉORÈME 1.14 (inégalité de Taylor-Lagrange) — Soit $f : I \rightarrow \mathbb{K}$ une fonction de classe $\mathcal{C}^{n+1}$, et $a \in I$. On suppose l'existence d'un réel M vérifiant : $\forall t \in I, |f^{(n+1)}(t)| \leq M$. Alors :

$$|f(x) - T_n(x)| \leq M \frac{|x-a|^{n+1}}{(n+1)!}$$

Exercice 4

Appliquer l'inégalité de Taylor-Lagrange entre 0 et 1 à la fonction $x \mapsto \ln(1+x)$ et en déduire : $\sum_{k=0}^{+\infty} \frac{(-1)^k}{k+1} = \ln 2$.

Notons pour finir que la fonction $f^{(n+1)}$ (étant supposée continue) est bornée au voisinage de a , ce qui nous permet de déduire de l'inégalité de Taylor-Lagrange le résultat suivant :

COROLLAIRE (Formule de Taylor-Young) — Soit I un intervalle, $a \in I$ et $f : I \rightarrow \mathbb{K}$ une fonction de classe $\mathcal{C}^{n+1}$. Alors f admet au voisinage de a le développement limité suivant :

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

2. Intégration sur un intervalle

La notion d'intégrale que nous avons définie présente des limitations : l'intervalle d'intégration doit être un segment, et la fonction, continue par morceaux. Dans le cadre des fonctions à valeurs positives, ceci permet d'interpréter l'intégrale comme étant l'aire délimitée par le graphe de la fonction.

Lorsque l'intervalle d'intégration n'est plus un segment, le domaine délimité par la fonction peut ne plus être borné. Nous allons voir cependant que dans certains cas il reste possible de donner un sens à l'aire de ce domaine, par le biais d'un passage à la limite dans des intégrales : c'est la notion d'intégrale généralisée.

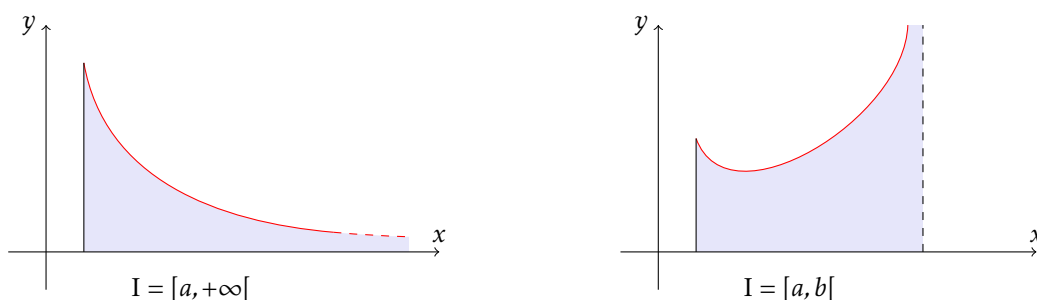


FIGURE 4 – Deux exemples de domaines non bornés, soit parce que l'une des deux bornes est infinie, soit parce que f n'admet pas de limite finie en une des deux bornes.

Nous verrons que cette intégrale généralisée partage un certain nombre de propriétés avec l'intégrale définie, avant d'étudier un théorème d'inversion limite-intégrale adapté aux intégrales généralisées : le théorème de convergence dominée.

Dans toute cette partie, I désigne un intervalle quelconque de $\mathbb{R}$, et $f : I \rightarrow \mathbb{K}$ une fonction continue par morceaux. L'objectif est de donner un sens, lorsque c'est possible, à l'intégrale $\int_I f$; on parlera alors d'intégrale généralisée, ou encore d'intégrale impropre.

2.1 Intégrales convergentes

Si I n'est pas un segment, I ne peut prendre qu'une des trois formes suivantes : $[a, b[$, $]a, b]$, $]a, b[$, avec $a \in \mathbb{R} \cup \{-\infty\}$ et $b \in \mathbb{R} \cup \{+\infty\}$. Nous allons traiter séparément chacun des trois cas.

■ Le cas où $I = [a, b[$, $b \in \mathbb{R} \cup \{+\infty\}$

Pour tout $x \in [a, b[$, f est continue par morceaux sur le segment $[a, x]$, donc $\int_a^x f(t) dt$ a bien un sens.

DÉFINITION. — On dira que l'intégrale $\int_a^b f(t)dt$ est convergente lorsque $\int_a^x f(t)dt$ possède une limite finie quand x tend vers b en restant dans $[a, b[$. On notera alors :

$$\int_a^b f(t)dt = \lim_{x \rightarrow b} \int_a^x f(t)dt.$$

On notera que lorsque f est prolongeable par continuité en b , cette définition est en cohérence avec la notion d'intégrale sur le segment $[a, b]$ puisque dans ce cas, la fonction $x \mapsto \int_a^x f(t)dt$ est une application définie et continue sur $[a, b]$, et en particulier en b . Dans ces conditions, on dira que l'intégrale est *faussement impropre*, puisqu'elle ne correspond pas à l'aire d'un domaine non borné : en prolongeant par continuité la fonction f en b on retrouve l'intégrale d'une fonction continue par morceaux sur un segment.

Remarque. On ne manquera pas de noter la similitude de la démarche avec celle utilisée pour définir la somme d'une série : à la fonction $x \mapsto \int_a^x f(t)dt$ correspondent les sommes partielles $n \mapsto \sum_{k=0}^n u_k$, et il s'agit dans les deux cas de déterminer si ces expressions possèdent une limite (l'une en b , l'autre en $+\infty$).

Exemples.

- L'intégrale de Rieman $\int_1^{+\infty} \frac{dt}{t^\alpha}$ est convergente si et seulement si $\alpha > 1$.
- Pour tout $\alpha \in \mathbb{R}$, l'intégrale $\int_0^{+\infty} e^{-\alpha t} dt$ est convergente si et seulement si $\alpha > 0$.

Pour étudier la convergence des deux exemples précédents, il a été nécessaire de calculer les « intégrales partielles » puis de passer à la limite. En revanche, il n'est pas nécessaire de procéder à ce calcul dans le cas de l'intégrale suivante : $\int_0^1 (t-1)\ln(1-t)dt$ puisqu'il s'agit d'une intégrale faussement impropre : en effet, $\lim_{t \rightarrow 1} (t-1)\ln(1-t) = 0$. Comme on peut le constater sur le graphe ci-dessous, le domaine délimité par le graphe de la fonction est borné :

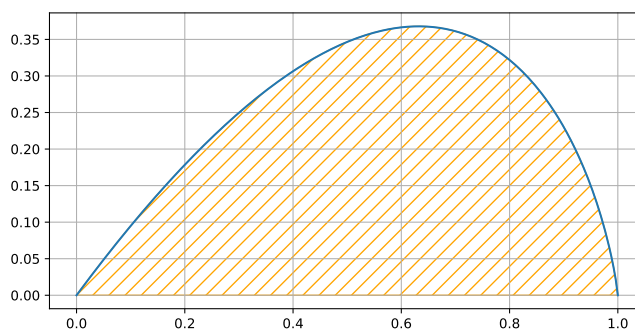


FIGURE 5 – Le graphe de la fonction $t \mapsto (t-1)\ln(1-t)$.

Exercice 5

Discuter en fonction de la valeur de $\beta > 0$ la convergence de l'intégrale $\int_2^{+\infty} \frac{dt}{t(\ln t)^\beta}$.

■ Le cas où $I =]a, b]$, $a \in \{-\infty\} \cup \mathbb{R}$

DÉFINITION. — De manière symétrique, on dira que l'intégrale $\int_a^b f(t)dt$ est convergente lorsque $\int_x^b f(t)dt$ possède une limite finie quand x tend vers a en restant dans $]a, b]$. On notera alors :

$$\int_a^b f(t) dt = \lim_{x \rightarrow a} \int_x^b f(t) dt.$$

Exemples.

- L'intégrale de Rieman $\int_0^1 \frac{dt}{t^\alpha}$ est convergente si et seulement si $\alpha < 1$.
- L'intégrale $\int_0^1 \ln(t) dt$ est convergente (et est égale à -1).

En revanche, l'intégrale $\int_0^{2\pi} \frac{\sin(t)}{t} dt$ est faussement impropre, car $\lim_{t \rightarrow 0} \frac{\sin(t)}{t} = 1$.

Exercice 6

Étudier la convergence de l'intégrale $\int_1^2 \frac{dt}{\sqrt{t^2 - 1}}$.

■ Le cas où $I =]a, b[$

Dans ce dernier cas, nous allons utiliser la relation de Chasles pour nous ramener aux deux cas précédents. Considérons en effet un point c de l'intervalle $]a, b[$.

DÉFINITION. — On dira que l'intégrale $\int_a^b f(t) dt$ est convergente lorsque les deux intégrales $\int_a^c f(t) dt$ et $\int_c^b f(t) dt$ sont convergentes, et on posera alors :

$$\int_a^b f(t) dt = \int_a^c f(t) dt + \int_c^b f(t) dt.$$

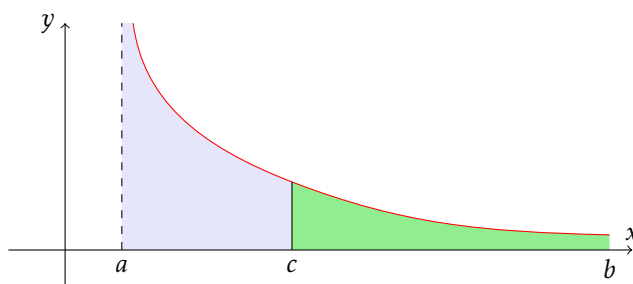


FIGURE 6 – Étude de la convergence d'une intégrale lorsque $I =]a, b[$

Il est aisé de vérifier que cette définition ne dépend pas du choix du point $c \in]a, b[$ (illustration figure 6).

Remarque. Pour chacun des trois types d'intervalles, une intégrale qui n'est pas convergente sera bien entendu dite *divergente*.

Par exemple, l'intégrale $\int_0^{+\infty} \frac{dt}{t^\alpha}$ est toujours divergente, puisque $\int_0^1 \frac{dt}{t^\alpha}$ ne converge que si $\alpha < 1$, alors que $\int_1^{+\infty} \frac{dt}{t^\alpha}$ ne converge que si $\alpha > 1$.

Remarque. Cette définition permet d'étendre sans peine au cas des intégrales convergentes la relation de Chasles, la propriété de linéarité des intégrales, ainsi que les propriétés de croissance et de positivité.

Exercice 7

Soit $f : [0, +\infty[\rightarrow \mathbb{R}$ une fonction continue.

- Montrer que si l'intégrale $\int_0^{+\infty} f(t) dt$ converge, alors $\lim_{x \rightarrow +\infty} \int_x^{x+1} f(t) dt = 0$.
- Si on suppose de plus f décroissante, en déduire que $\lim_{x \rightarrow +\infty} f(x) = 0$.
- Toujours en supposant f décroissante, montrer que $f(x) = o\left(\frac{1}{x}\right)$.

■ Intégration par parties et changement de variable

Une intégration par parties sur un intervalle qui n'est pas un segment peut conduire à une erreur : l'intégrale $\int_I u'v$ peut être convergente sans que $\int_I uv'$ le soit. Il convient donc de procéder avec prudence, en utilisant le résultat suivant :

THÉORÈME 2.1 — Soient u et v deux fonctions de classe $\mathcal{C}^1$ sur un intervalle I , telles que les limites aux bornes de I du produit uv existent. Alors les intégrales $\int_I u'v$ et $\int_I uv'$ ont même nature, et en cas de convergence,

$$\int_a^b u'(t)v(t) dt = \left[u(t)v(t) \right]_a^b - \int_a^b u(t)v'(t) dt$$

où a et b désignent les bornes de I .

Remarque. Une autre possibilité consiste à effectuer l'intégration par parties sur un segment (par exemple sur $[a, x]$ lorsque $I = [a, b[$), puis, une fois tous les calculs effectués, passer à la limite (ici en faisant tendre x vers b).

Exemple. Pour tout $n \in \mathbb{N}$, l'intégrale : $\int_0^{+\infty} t^n e^{-t} dt$ est convergente et vaut $n!$.

Exercice 8

Justifier la convergence et calculer la valeur de l'intégrale $\int_0^{+\infty} \ln\left(1 + \frac{1}{t^2}\right) dt$.

En ce qui concerne le changement de variable, on possède le résultat suivant :

THÉORÈME 2.2 — Soient I et J deux intervalles, $\phi : J \rightarrow I$ une bijection de classe $\mathcal{C}^1$, et $f : I \rightarrow \mathbb{C}$ une fonction continue par morceaux. Alors les intégrales $\int_I f$ et $\int_J (f \circ \phi) \times \phi'$ ont même nature, et en cas de convergence,

$$\int_{\phi(a)}^{\phi(b)} f(u) du = \int_a^b f \circ \phi(t) \phi'(t) dt$$

lorsque a et b désignent les extrémités de J et $\phi(a)$ et $\phi(b)$ les limites respectives de ϕ en a et en b .

Exemple. Pour tout $n \in \mathbb{N}$, l'intégrale $\int_0^1 (\ln u)^n du$ est convergente et vaut $(-1)^n n!$.

Exercice 9

Justifier la convergence et calculer la valeur de l'intégrale $\int_0^{+\infty} e^{-\sqrt{t}} dt$.

2.2 Fonctions à valeurs positives

L'inconvénient de la démarche que nous avons suivi jusqu'à présent est de nécessiter le calcul d'une primitive de f pour déterminer la nature de l'intégrale $\int_I f$; or cela n'est pas toujours possible. Nous allons donc tenter de nous affranchir de cette contrainte, en commençant par nous intéresser aux fonctions à valeurs positives.

En effet, lorsque $f : [a, b[\rightarrow \mathbb{R}_+$ est une fonction continue par morceaux et à valeurs *positives*, l'application $F : x \mapsto \int_a^x f(t) dt$ est une fonction *croissante*. Elle possède donc une limite lorsque x tend vers b si et seulement si elle est *majorée*⁶.

De la même façon, lorsque $f :]a, b] \rightarrow \mathbb{R}_+$ est continue par morceaux et à valeurs *positives*, l'application $F : x \mapsto \int_x^b f(t) dt$ est *décroissante* et possède donc une limite lorsque x tend vers a si et seulement si elle est *majorée*.

Comme pour les séries, ces résultats vont engendrer plusieurs théorèmes de comparaison, qui reposent tous sur le résultat suivant :

THÉORÈME 2.3 (comparaison) — Soient f et g deux fonction continues par morceaux sur l'intervalle I , et à valeurs positives. On suppose que pour tout $t \in I$, $0 \leq f(t) \leq g(t)$. Alors la convergence de l'intégrale $\int_I g$ entraîne celle de $\int_I f$.

De ce théorème vont résulter deux corollaires, qui vont permettre de comparer la nature des intégrales que nous allons rencontrer par la suite à la nature d'intégrales de référence. Ces deux corollaires seront énoncés dans le cas où $I = [a, b]$, mais leur énoncé s'adapte sans problème au cas symétrique où $I =]a, b]$.

COROLLAIRE (domination) — Soit $f : [a, b] \rightarrow \mathbb{R}_+$ et $g : [a, b] \rightarrow \mathbb{R}_+$ deux fonctions continues par morceaux, à valeurs positives, telles que $f(t) = O_b(g(t))$. Alors la convergence de $\int_a^b g(t) dt$ entraîne celle de $\int_a^b f(t) dt$.

COROLLAIRE (équivalence) — Soit $f : [a, b] \rightarrow \mathbb{R}_+$ et $g : [a, b] \rightarrow \mathbb{R}_+$ deux fonctions continues par morceaux, à valeurs positives, telles que $f(t) \sim_b g(t)$. Alors les intégrales $\int_a^b g(t) dt$ et $\int_a^b f(t) dt$ ont même nature.

Remarque. Les intégrales de référence que nous utiliserons sont les suivantes :

$$\int_1^{+\infty} \frac{dt}{t^\alpha} \text{ converge ssi } \alpha > 1, \quad \int_0^1 \frac{dt}{t^\alpha} \text{ converge ssi } \alpha < 1, \quad \int_0^{+\infty} e^{-at} dt \text{ converge ssi } a > 0.$$

On pourra rajouter à cette liste le résultat suivant : $\int_0^1 (\ln t) dt$ converge.

Exercice 10

À l'aide du théorème d'équivalence, prouver que l'intégrale $\int_0^{+\infty} \frac{dt}{(t^2 + 1)\sqrt{t^2 + 4}}$ converge.

À l'aide du théorème de domination, prouver que l'intégrale $\int_1^{+\infty} \frac{\ln t}{t^2} dt$ converge.

À l'aide d'un changement de variable, prouver que l'intégrale $\int_0^1 \frac{dt}{\sqrt{1-t^2}}$ converge.

■ L'exemple de la fonction Γ

La fonction Γ est une fonction mathématique définie sur une partie de $\mathbb{R}$ par la formule : $\Gamma(x) = \int_0^{+\infty} e^{-t} t^{x-1} dt$.

- Au voisinage de 0, $e^{-t} t^{x-1} \sim_0 \frac{1}{t^{1-x}}$, donc (théorème d'équivalence) $\int_0^1 e^{-t} t^{x-1} dt$ converge si et seulement si $x > 0$.
- Au voisinage de $+\infty$, $e^{-t} t^{x-1} =_{+\infty} O(e^{-t/2})$, donc (théorème de domination) $\int_1^{+\infty} e^{-t} t^{x-1} dt$ converge pour tout $x \in \mathbb{R}$.

6. On ne manquera pas de faire l'analogie avec les séries à terme général positif : lorsque (u_n) est une suite positive, la suite des sommes partielles de la série $\sum u_n$ est croissante, et la série converge si et seulement si la suite des sommes partielles est majorée.

On en déduit que la fonction Γ est définie sur l'intervalle $]0, +\infty[$.

La propriété la plus simple de la fonction Γ est de vérifier la relation : $\Gamma(x+1) = x\Gamma(x)$, qui résulte d'une intégration par parties :

$$\int_u^v e^{-t} t^x dt = \left[-e^{-t} t^x \right]_u^v + \int_u^v e^{-t} x t^{x-1} dt = e^{-u} u^x - e^{-v} v^x + x \int_u^v e^{-t} t^{x-1} dt$$

qui conduit en faisant tendre u vers 0 et v vers $+\infty$ à : $\Gamma(x+1) = 0 - 0 + x\Gamma(x)$.

Sachant que $\Gamma(1) = 1$ on en déduit aisément la relation : $\forall n \in \mathbb{N}^*, \Gamma(n) = (n-1)!$; d'une certaine façon, la fonction Γ prolonge donc la fonction factorielle à une partie de $\mathbb{R}$.

2.3 Absolue convergence et fonctions intégrables

La section précédente nous a donné des outils pour prouver la convergence des intégrales des fonctions à valeurs positives. Dans le cas général, nous allons nous y ramener à l'aide de la notion d'absolue convergence, notion semblable à celle que l'on connaît déjà dans le cadre des séries numériques.

THÉORÈME 2.4 — Soit $f : I \rightarrow \mathbb{K}$ une fonction continue par morceaux, telle que l'intégrale : $\int_I |f|$ soit convergente. Alors il en est de même de l'intégrale $\int_I f$. On dira que cette dernière intégrale est absolument convergente, et que la fonction f est intégrable sur I .

Remarque. Lorsque $I = [a, b[$ une fonction intégrable sur I sera aussi dite *intégrable en b* . De même, une fonction intégrable sur $]a, b]$ sera aussi dite *intégrable en a* .

THÉORÈME 2.5 — L'espace $L^1(I, \mathbb{K})$ des fonctions intégrables de I vers $\mathbb{K}$ est un $\mathbb{K}$ -espace vectoriel.

Nous pouvons traduire le théorème 2.3 et ses corollaires en termes d'intégrabilité :

THÉORÈME 2.6 — Soit $f : I \rightarrow \mathbb{C}$ et $g : I \rightarrow \mathbb{R}_+$ deux fonctions continues par morceaux, telles que $0 \leq |f| \leq g$. Si g est intégrable sur I , il en est de même de f , et $\left| \int_I f \right| \leq \int_I g$.

Exemple. La fonction $t \mapsto \frac{\cos t}{t^2}$ est intégrable sur $[1, +\infty[$ (ou encore intégrable en $+\infty$) car $\left| \frac{\cos t}{t^2} \right| \leq \frac{1}{t^2}$.

COROLLAIRE — Soit $f : [a, b[\rightarrow \mathbb{C}$ et $g : [a, b[\rightarrow \mathbb{R}_+$ deux fonctions continues par morceaux telles que $|f(t)| = O(g(t))$. Alors si g est intégrable sur $[a, b[$, il en est de même de f .

COROLLAIRE — Soit $f : [a, b[\rightarrow \mathbb{C}$ et $g : [a, b[\rightarrow \mathbb{R}_+$ deux fonctions continues par morceaux telles que $|f(t)| \sim_b g(t)$. Alors si g est intégrable sur $[a, b[$, il en est de même de f .

Les deux derniers résultats s'étendent bien entendu au cas de l'intervalle $]a, b]$.

■ Fonctions de carré intégrable

Une fonction $f : I \rightarrow \mathbb{K}$ est dite de *carré intégrable* lorsque la fonction f^2 est intégrable sur I . Ces fonctions constituent elles aussi un espace vectoriel, mais ce résultat passe par l'obtention du résultat suivant :

LEMME — Si f et g sont deux fonctions de carré intégrable, alors fg est intégrable.

THÉORÈME 2.7 — L'espace $L^2(I, \mathbb{K})$ des fonctions de carré intégrable de I vers $\mathbb{K}$ est un $\mathbb{K}$ -espace vectoriel.

Remarque. Le résultat du lemme permet en outre d'observer que l'application $(f, g) \mapsto \int_I fg$ est une application bilinéaire, symétrique et positive définie sur $L^2(I, \mathbb{R})$. Ceci conduit naturellement à :

THÉORÈME 2.8 (Inégalité de Cauchy-Schwarz) — Si f et g sont deux fonctions de carré intégrables, alors

$$\left| \int_I fg \right| \leq \left(\int_I |f|^2 \int_I |g|^2 \right)^{1/2}$$

■ Un exemple de semi-convergence

La notion d'intégrabilité que nous venons de définir est la seule qui généralise de manière pertinente la notion d'intégration sur un segment ; en effet nous verrons dans la section suivante que les différents théorèmes relatifs aux intégrales à paramètre exigent une hypothèse d'intégrabilité.

Cependant, il existe des intégrales qui sont convergentes sans être absolument convergentes : l'intégrale $\int_I f$ converge mais l'intégrale $\int_I |f|$ diverge. On dit dans ce cas que l'intégrale $\int_I f$ est *semi-convergente*. Attention, dans ce cas la fonction f **n'est pas** intégrable sur I (tout en possédant une intégrale, ce qui peut paraître paradoxal).

L'étude de la semi-convergence n'est pas un objectif du programme, aussi nous nous contenterons de voir un seul exemple :

PROPOSITION 2.9 — L'intégrale de Dirichlet $\int_0^{+\infty} \frac{\sin t}{t} dt$ est semi-convergente.

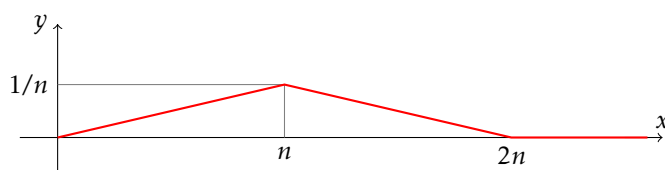
3. Le théorème de convergence dominée

Nous avons vu dans le chapitre sur les suites de fonctions que sous réserve d'une hypothèse de convergence uniforme sur le segment $[a, b]$, on pouvait intervertir passage à la limite et intégration :

$$\lim_{n \rightarrow +\infty} \int_a^b f_n(t) dt = \int_a^b \lim_{n \rightarrow +\infty} f_n(t) dt$$

Malheureusement, ce théorème *ne s'étend pas* au cas de l'intégration sur un intervalle quelconque, comme le montre l'exemple suivant :

Pour tout $n \in \mathbb{N}^*$, $f_n : [0, +\infty[\rightarrow \mathbb{R}$ est la fonction continue et affine par morceaux dont le graphe est donné ci-dessous :



Puisque $\|f_n\|_\infty = \frac{1}{n}$, la suite (f_n) converge uniformément vers la fonction nulle sur $[0, +\infty[$. Pourtant, pour tout $n \in \mathbb{N}$ on a $\int_0^{+\infty} f_n = 1$, donc $\lim_{n \rightarrow +\infty} \int_0^{+\infty} f_n \neq \int_0^{+\infty} \lim_{n \rightarrow +\infty} f_n$.

Nous allons maintenant étudier un théorème permettant de faire une telle interversion dans le cadre d'un intervalle *quelconque*, segment ou pas. Cependant, la preuve de ce résultat sera admise, car inaccessible à ce niveau.

3.1 Le théorème de convergence dominée

Ce théorème s'applique aux fonctions à valeurs réelles ou complexes.

THÉORÈME 3.1 (Théorème de convergence dominée) — Soit (f_n) une suite de fonctions à valeurs réelles ou complexes, continues par morceaux sur I . On suppose que :

- (i) (f_n) converge simplement sur I vers une fonction f continue par morceaux sur I ;
- (ii) il existe une fonction ϕ intégrable sur I , telle que :

$$\forall n \in \mathbb{N}, \quad |f_n| \leq \phi \quad (\text{hypothèse de domination}).$$

Alors les fonctions f_n et f sont intégrables sur I , et :

$$\int_I f = \lim_{n \rightarrow +\infty} \int_I f_n.$$

Remarque. IL n'est pas difficile de justifier l'intégrabilité des fonctions f_n : il s'agit d'une application directe du théorème 2.6. De même, l'hypothèse de convergence simple permet le passage à la limite dans l'inégalité : $\forall t \in I, |f_n(t)| \leq \phi(t)$ pour obtenir : $\forall t \in I, |f(t)| \leq \phi(t)$, ce qui permet d'appliquer de nouveau le théorème 2.6 pour justifier l'intégrabilité de f . En revanche, nous admettrons l'égalité encadrée.

Exemple. Considérons les intégrales de Wallis $\int_0^{\frac{\pi}{2}} (\cos t)^n dt$. La suite de fonctions $f_n : t \mapsto (\cos t)^n$ converge simplement vers la fonction $f : t \mapsto 0$ sur l'intervalle $]0, \frac{\pi}{2}[$. Les fonctions f_n sont dominées par la fonction intégrable $\phi : t \mapsto 1$, donc le théorème de convergence dominée s'applique : $\lim_{n \rightarrow +\infty} \int_0^{\frac{\pi}{2}} (\cos t)^n dt = \int_0^{\frac{\pi}{2}} f(t) dt = 0$.

Exemple. Considérons pour $n \geq 1$ les intégrales $\int_0^{+\infty} e^{-t^n} dt$. La suite de fonctions $f_n : t \mapsto e^{-t^n}$ converge simplement vers la fonction $f : t \mapsto \begin{cases} 1 & \text{si } t \in [0, 1[\\ e^{-1} & \text{si } t = 1 \\ 0 & \text{si } t > 1 \end{cases}$. Les fonctions f_n sont dominées par la fonction intégrable $\phi : t \mapsto \begin{cases} 1 & \text{si } t \in [0, 1[\\ e^{-t} & \text{si } t > 1 \end{cases}$, donc le théorème de convergence dominée s'applique : $\lim_{n \rightarrow +\infty} \int_0^{+\infty} e^{-t^n} dt = \int_0^{+\infty} f(t) dt = 1$.

Exercice 11

Soit $h : [0, +\infty[\rightarrow \mathbb{R}$ une fonction continue et bornée, et, pour $n \in \mathbb{N}^*$, $u_n = \int_0^{+\infty} \frac{h(t)}{1 + n^2 t^2} dt$.

Déterminer la limite de la suite (u_n) .

Lorsque $h(0) \neq 0$ et l'aide du changement de variable $x = nt$, déterminer un équivalent de u_n .

3.2 Intégration terme à terme d'une série de fonctions

Pour inverser série et intégrale, nous allons là encore être obligés d'admettre le résultat suivant :

THÉORÈME 3.2 — Soit (f_n) une suite de fonctions intégrables sur I . On suppose que :

- (i) la série de fonctions $\sum f_n$ converge simplement et la somme $S = \sum_{n=0}^{+\infty} f_n$ est continue par morceaux;
- (ii) la série $\sum \int_I |f_n|$ est convergente.

Alors la fonction S est intégrable sur I , et $\int_I S = \sum_{n=0}^{+\infty} \int_I f_n$.

Remarque. On peut observer que la série $\sum \int_I f_n$ converge absolument puisque $|\int_I f_n| \leq \int_I |f_n|$.

Exemple. Nous allons calculer l'intégrale $\int_0^1 \frac{\ln(1+x)}{x} dx$ à l'aide d'un développement en série.

$$\forall x \in]0, 1[, \ln(1+x) = \sum_{n=1}^{+\infty} (-1)^{n-1} \frac{x^n}{n} \text{ donc } \frac{\ln(1+x)}{x} = \sum_{n=1}^{+\infty} (-1)^{n-1} \frac{x^{n-1}}{n}.$$

Notons $f_n : x \mapsto (-1)^{n-1} \frac{x^{n-1}}{n}$. La série de fonctions $\sum f_n$ converge simplement, et la somme $S : x \mapsto \frac{\ln(1+x)}{x}$ est continue.

On calcule $\int_0^1 \left| (-1)^{n-1} \frac{x^{n-1}}{n} \right| dx = \frac{1}{n^2}$; sachant que la série $\sum \frac{1}{n^2}$ converge, le théorème d'intégration terme à terme s'applique : $\int_0^1 \frac{\ln(1+x)}{x} dx = \sum_{n=1}^{+\infty} \frac{(-1)^{n-1}}{n^2} = \frac{\pi^2}{12}$.

Exemple. Considérons un réel $\alpha > 0$ et, pour $n \in \mathbb{N}^*$, les fonctions $f_n : x \mapsto x^{\alpha-1} e^{-nx}$. Il s'agit de fonctions continues et intégrables sur $]0, +\infty[$. Pour tout $x > 0$, $\sum_{n=1}^{+\infty} x^{\alpha-1} e^{-nx} = \frac{x^{\alpha-1}}{e^x - 1}$ et (en posant $y = nx$) :

$$\int_0^{+\infty} |x^{\alpha-1} e^{-nx}| dx = \frac{1}{n^\alpha} \int_0^{+\infty} y^{\alpha-1} e^{-y} dy = \frac{\Gamma(\alpha)}{n^\alpha}.$$

La série $\sum \frac{1}{n^\alpha}$ converge dès lors que $\alpha > 1$. On en déduit que la fonction $x \mapsto \frac{x^{\alpha-1}}{e^x - 1}$ est intégrable sur $]0, +\infty[$ lorsque $\alpha > 1$ (résultat qu'on pouvait obtenir directement), et dans ce cas :

$$\int_0^{+\infty} \frac{x^{\alpha-1}}{e^x - 1} dx = \Gamma(\alpha) \sum_{n=1}^{+\infty} \frac{1}{n^\alpha} = \Gamma(\alpha) \zeta(\alpha).$$

Par exemple, pour $\alpha = 2$ on obtient : $\int_0^{+\infty} \frac{x}{e^x - 1} dx = \sum_{n=1}^{+\infty} \frac{1}{n^2} = \frac{\pi^2}{6}$.

Exercice 12

La constante de Catalan est le réel : $K = -\int_0^1 \frac{\ln t}{1+t^2} dt$. Établir l'égalité : $K = \sum_{n=0}^{+\infty} \frac{(-1)^n}{(2n+1)^2}$.

■ Une démarche alternative

Considérons l'exemple suivant : on cherche à calculer $\sum_{n=0}^{+\infty} \frac{(-1)^n}{n+1}$ en utilisant la suite de calculs suivante :

$$\sum_{n=0}^{+\infty} \frac{(-1)^n}{n+1} = \sum_{n=0}^{+\infty} (-1)^n \int_0^1 t^n dt \stackrel{?}{=} \int_0^1 \sum_{n=0}^{+\infty} (-1)^n t^n dt = \int_0^1 \frac{dt}{1+t} = \ln 2.$$

Il nous faut justifier la deuxième égalité de ce calcul.

Essayons d'utiliser le théorème 3.2 en notant $f_n : t \mapsto (-1)^n t^n$ sur $I = [0, 1[$. L'hypothèse (i) est bien vérifiée, mais pas l'hypothèse (ii) car la série $\sum \frac{1}{n+1}$ diverge. Dans ce cas, la solution consiste à utiliser le théorème 3.1 à la suite des restes. On écrit :

$$\int_0^1 \frac{dt}{1+t} = \sum_{n=0}^N \int_0^1 (-1)^n t^n dt + \int_0^1 R_N(t) dt = \sum_{n=0}^N \frac{(-1)^n}{n+1} + \int_0^1 R_N(t) dt \quad \text{avec } R_N(t) = \sum_{n=N+1}^{+\infty} (-1)^n t^n.$$

D'après le critère spécial relatif aux séries alternées, $|R_N(t)| \leq t^{N+1} \leq 1$, ce qui permet d'appliquer le théorème de convergence dominée : $\lim_{N \rightarrow +\infty} \int_0^1 R_N(t) dt = \int_0^1 \lim_{N \rightarrow +\infty} R_N(t) dt = 0$ et ainsi conclure : $\int_0^1 \frac{dt}{1+t} = \sum_{n=0}^{+\infty} \frac{(-1)^n}{n+1}$.

3.3 Intégrales dépendant d'un paramètre

Considérons maintenant une fonction à deux variables $f : A \times I \rightarrow \mathbb{K}$ telle que pour tout $x \in A$, la fonction $t \mapsto f(x, t)$ soit continue par morceaux et intégrable sur I . On peut alors définir une application $g : A \rightarrow \mathbb{K}$ en posant :

$$\forall x \in A, \quad g(x) = \int_I f(x, t) dt.$$

La continuité de f vis-à-vis de sa variable x permet-elle d'en déduire celle de g ? La réponse est malheureusement négative. Considérons à cet effet l'application :

$$f : (x, t) \mapsto x e^{-tx} \quad \text{et} \quad g : x \mapsto \int_0^{+\infty} x e^{-xt} dt.$$

Pour tout $x \geq 0$, l'application $f(x, \cdot)$ est intégrable sur $[0, +\infty[$, donc g est bien définie.

Il est aisé de calculer : $g(0) = 0$ et $\forall x > 0$, $g(x) = 1$, aussi g est discontinue en 0 bien que f soit continue vis-à-vis de x .

Pour en déduire la continuité de la fonction g , il va donc être nécessaire d'avoir une hypothèse supplémentaire : ce sera une hypothèse de *domination*.

■ Continuité sous le signe intégral

THÉORÈME 3.3 — Soit $f : A \times I \rightarrow \mathbb{K}$ une fonction telle que :

- (i) pour tout $x \in A$, la fonction $t \mapsto f(x, t)$ est continue par morceaux ;
- (ii) pour tout $t \in I$, la fonction $x \mapsto f(x, t)$ est continue sur A ;
- (iii) il existe une application ϕ intégrable sur I telle que :

$$\forall (x, t) \in A \times I, \quad |f(x, t)| \leq \phi(t) \quad (\text{hypothèse de domination}).$$

Alors l'application $g : A \rightarrow \mathbb{K}$ définie par : $\forall x \in A, g(x) = \int_I f(x, t) dt$ est définie et continue en tout point de A .

Exemple. Rappelons que la fonction $\Gamma : x \mapsto \int_0^{+\infty} t^{x-1} e^{-t} dt$ est définie sur $]0, +\infty[$. Nous allons prouver qu'elle y est continue.

Notons $f(x, t) = t^{x-1} e^{-t}$ et considérons deux réels $0 < a < b$. Pour tout $x \in [a, b]$ on a : $\forall t \in]0, 1]$, $0 \leq t^{x-1} e^{-t} \leq t^{a-1} e^{-t}$ et pour tout $t \in [1, +\infty[$, $0 \leq t^{x-1} e^{-t} \leq t^{b-1} e^{-t}$. Ainsi, pour tout $t > 0$ on a $|f(x, t)| \leq \phi(t)$ avec

$$\phi(t) = \begin{cases} t^{a-1} e^{-t} & \text{si } t \leq 1 \\ t^{b-1} e^{-t} & \text{si } t \geq 1 \end{cases}$$

L'application ϕ est intégrable sur $]0, +\infty[$ et domine f , donc Γ est continue sur $[a, b]$, puis par recouvrement sur $]0, +\infty[$.

Remarque. Comme nous venons de le voir sur cet exemple, il est possible de procéder par recouvrement, en prouvant par exemple la continuité sur tout segment inclus dans I .

■ Limites aux bornes de l'intervalle de définition

La notion de caractérisation séquentielle permet d'obtenir une version continue du théorème de convergence dominée, sous la forme :

THÉORÈME 3.4 — Soit $f : A \times I \rightarrow \mathbb{K}$ une fonction telle que :

- (i) pour tout $x \in A$, la fonction $t \mapsto f(x, t)$ est continue par morceaux ;
- (ii) pour tout $t \in I$, $f(x, t) \xrightarrow{x \rightarrow a} \ell(t)$, la fonction ℓ étant continue par morceaux sur I ;
- (iii) il existe une application ϕ intégrable sur I telle que :

$$\forall (x, t) \in A \times I, \quad |f(x, t)| \leq \phi(t) \quad (\text{hypothèse de domination}).$$

Alors ℓ est intégrable sur I , et $\int_I f(x, t) dt \xrightarrow{x \rightarrow a} \int_I \ell(t) dt$.

■ Dérivation sous le signe intégral

THÉORÈME 3.5 — Soit $f : A \times I \rightarrow \mathbb{K}$ une application vérifiant les hypothèses suivantes :

- (i) pour tout $x \in A$, la fonction $t \mapsto f(x, t)$ est continue par morceaux et intégrable sur I ;
- (ii) pour tout $t \in I$, $x \mapsto f(x, t)$ est de classe $\mathcal{C}^1$ sur A ;
- (iii) pour tout $x \in A$ la fonction $t \mapsto \frac{\partial f}{\partial x}(x, t)$ est continue par morceaux sur I ;
- (iv) il existe une application ϕ continue par morceaux, positive et intégrable sur I telle que :

$$\forall (x, t) \in A \times I, \quad \left| \frac{\partial f}{\partial x}(x, t) \right| \leq \phi(t) \quad (\text{hypothèse de domination}).$$

Alors la fonction $g : x \mapsto \int_I f(x, t) dt$ est de classe $\mathcal{C}^1$ sur A , et : $\forall x \in A, g'(x) = \int_I \frac{\partial f}{\partial x}(x, t) dt$.

Remarque. À l'instar de la continuité, il est fréquent d'avoir à procéder par recouvrement, par exemple en prouvant à l'aide de ce théorème que g est de classe $\mathcal{C}^1$ sur tout segment inclus dans J .

Exercice 13

On considère la fonction $g : x \mapsto \int_0^{+\infty} \frac{e^{-xt^2}}{1+t^2} dt$.

- a. Montrer que g est définie et continue sur $[0, +\infty[$.
- b. Montrer que g est de classe $\mathcal{C}^1$ sur $]0, +\infty[$ et solution sur cet intervalle de l'équation différentielle :

$$y - y' = \sqrt{\frac{\pi}{4x}}.$$

Extension au cas des fonctions de classe $\mathcal{C}^k$

Enfin, nous admettrons l'extension de ce théorème, à l'instar du théorème équivalent pour les séries de fonctions :

PROPOSITION 3.6 — Soit $f : A \times I \rightarrow \mathbb{K}$ une application vérifiant les hypothèses suivantes :

- (i) pour tout $x \in A$, la fonction $t \mapsto f(x, t)$ est continue par morceaux et intégrable sur I ;
- (ii) pour tout $t \in I$, $x \mapsto f(x, t)$ est de classe $\mathcal{C}^k$ sur A ;
- (iii) pour tout $x \in A$ et tout $i \in \llbracket 1, k-1 \rrbracket$, la fonction $t \mapsto \frac{\partial^i f}{\partial x^i}(x, t)$ est continue par morceaux et intégrable sur I ;
- (iv) pour tout $x \in A$ la fonction $t \mapsto \frac{\partial^k f}{\partial x^k}(x, t)$ est continue par morceaux sur I ;
- (v) il existe une application ϕ continue par morceaux, positive et intégrable sur I telle que :

$$\forall (x, t) \in A \times I, \quad \left| \frac{\partial^k f}{\partial x^k}(x, t) \right| \leq \phi(t) \quad (\text{hypothèse de domination}).$$

Alors la fonction $g : x \mapsto \int_I f(x, t) dt$ est de classe $\mathcal{C}^k$ sur A , et : $\forall x \in A, \forall i \in \llbracket 1, k \rrbracket, g^{(i)}(x) = \int_I \frac{\partial^i f}{\partial x^i}(x, t) dt$.

Chapitre IX

Espaces vectoriels normés

Au tournant du XX^e siècle, les travaux de Hilbert et de Banach relatifs aux espaces de fonctions étendent la notion de limite, initialement restreinte aux suites et fonctions réelles, à un cadre plus vaste : les *espaces vectoriels normés*, ouvrant ainsi la porte à l'*analyse fonctionnelle*, notion maintenant omniprésente dans toutes les branches des mathématiques.

1. Espaces vectoriels normés

1.1 Normes et distances

En géométrie, la *norme* est une extension de la valeur absolue des nombres aux vecteurs. Elle permet de mesurer la *longueur* d'un vecteur, mais définit aussi, nous le verrons, une *distance* entre deux vecteurs. Cette distance nous permettra de définir la notion de suite convergente, puis de limite d'une fonction à valeurs vectorielles.

Dans tout le chapitre, E désigne un $\mathbb{K}$ -espace vectoriel, avec $\mathbb{K} = \mathbb{R}$ ou $\mathbb{C}$.

DÉFINITION. — On appelle norme sur E toute application $N : E \rightarrow \mathbb{R}_+$ vérifiant :

- (i) $N(x) = 0 \implies x = 0_E$;
- (ii) pour tout $x \in E$, pour tout $\lambda \in \mathbb{K}$, $N(\lambda x) = |\lambda|N(x)$;
- (iii) pour tout $(x, y) \in E^2$, $N(x + y) \leq N(x) + N(y)$ (inégalité triangulaire).

On appelle espace vectoriel normé tout espace vectoriel muni d'une norme.

En général, on conviendra d'utiliser la notation usuelle $N(x) = \|x\|$.

PROPOSITION 1.1 (seconde inégalité triangulaire) — Pour tout $(x, y) \in E^2$, $|||x| - |y||| \leq \|x - y\|$.

Remarque. Le terme de *norme* ne vous est pas inconnu : à tout produit scalaire défini sur un espace vectoriel réel est associée une norme définie par : $\|x\| = \sqrt{\langle x | x \rangle}$. Ce type de norme respecte la définition que nous venons de donner, et de telles normes seront qualifiées de *norme euclidienne*.

Mais attention : dans le cas général une norme n'est pas forcément issue d'un produit scalaire.

Exemples. Lorsque $E = \mathbb{R}^p$ on utilise souvent l'une des normes suivantes : si $x = (x_1, \dots, x_p) \in \mathbb{R}^p$,

$$\|x\|_\infty = \max(|x_1|, |x_2|, \dots, |x_p|) \quad \|x\|_1 = \sum_{k=1}^p |x_k| \quad \|x\|_2 = \left(\sum_{k=1}^p x_k^2 \right)^{1/2}$$

Remarquons que les deux premières normes peuvent être aussi définies sur $\mathbb{C}^p$.

Exemples.

- Sur l'espace $\mathcal{B}(I, \mathbb{R})$ des fonctions bornées de I dans $\mathbb{R}$, la norme infinie $\|f\|_{\infty, I} = \sup_{x \in I} |f(x)|$ est une norme ;
- sur l'espace $L^1(I, \mathbb{R})$ des fonctions continues et intégrables sur I , $\|f\|_1 = \int_I |f(t)| dt$ est une norme ;
- sur l'espace $L^2(I, \mathbb{R})$ des fonctions continues et de carré intégrable sur I , $\|f\|_2 = \sqrt{\int_I f(t)^2 dt}$ est une norme euclidienne.

■ Distance entre deux vecteurs

DÉFINITION. — Lorsque E est un espace vectoriel normé et $(x, y) \in E^2$, on appelle distance entre x et y le réel $d(x, y) = \|y - x\|$.

Cette notion de distance est importante ; c'est elle qui nous permettra de généraliser en dimension supérieure les notions d'analyse que sont la convergence des suites, la continuité des fonctions, ...

À chaque norme est associé une distance différente, mais toutes les distances ont en commun les propriétés suivantes :

PROPOSITION 1.2 — Si d est une distance de E , alors :

- $\forall (x, y) \in E^2, \quad d(x, y) = 0 \iff x = y \quad (\text{séparation}) ;$
- $\forall (x, y) \in E^2, \quad d(x, y) = d(y, x) \quad (\text{symétrie}) ;$
- $\forall (x, y, z) \in E^3, \quad d(x, z) \leq d(x, y) + d(y, z) \quad (\text{inégalité triangulaire}).$

DÉFINITION. — On appelle sphère de centre $a \in E$ et de rayon $r > 0$ l'ensemble des vecteurs $x \in E$ vérifiant : $d(a, x) = r$. Autrement dit,

$$S(a, r) = \{x \in E \mid \|x - a\| = r\}.$$

Exercice 1

Dessiner la sphère unité (c'est-à-dire la sphère de centre 0_E de rayon 1) pour chacune des trois normes $\|\cdot\|_\infty$, $\|\cdot\|_1$ et $\|\cdot\|_2$ dans $\mathbb{R}^2$.

Par analogie aux intervalles ouverts et fermés de $\mathbb{R}$, on adopte en outre les définitions suivantes :

DÉFINITION. — On appelle boule ouverte de centre $a \in E$ et de rayon $r > 0$ l'ensemble des vecteurs $x \in E$ vérifiant : $d(a, x) < r$. Autrement dit,

$$\mathring{B}(a, r) = \{x \in E \mid \|x - a\| < r\}.$$

On appelle boule fermée de centre $a \in E$ et de rayon $r > 0$ l'ensemble des vecteurs $x \in E$ vérifiant : $d(a, x) \leq r$. Autrement dit,

$$\overline{B}(a, r) = \{x \in E \mid \|x - a\| \leq r\}.$$

Remarque. Les intervalles sont les seules parties convexes de $\mathbb{R}$, c'est-à-dire vérifiant la propriété :

$$\forall (a, b) \in A^2, \quad [a, b] \subset A.$$

Dans le cas d'un espace vectoriel, on définit la notion de segment en posant :

$$\forall (a, b) \in E^2, \quad [a, b] = \{(1-t)a + tb \mid t \in [0, 1]\}$$

On peut dès lors définir la notion de convexité dans un espace vectoriel :

DÉFINITION. — Une partie A d'un espace vectoriel (normé) E est dite convexe lorsque :

$$\forall (a, b) \in A^2, \quad [a, b] \subset A.$$

PROPOSITION 1.3 — Les boules ouvertes et les boules fermées sont des parties convexes d'un espace vectoriel normé.

Cette notion de boule permet d'étendre certaines propriétés topologiques de $\mathbb{R}$ au cas d'un espace vectoriel normé. Prenons par exemple la notion de partie bornée. Dans le cas réel, une partie bornée est définie ainsi : « une partie A de $\mathbb{R}$ est dite bornée lorsqu'il existe un réel $M > 0$ tel que $A \subset [-M, M]$ ». Dans le cadre des espaces vectoriels normés, cette définition devient :

DÉFINITION. — Soit E un $\mathbb{K}$ -espace vectoriel de dimension finie muni d'une norme $\|\cdot\|$. Une partie A de E est dite bornée lorsqu'il existe $M > 0$ tel que $A \subset \overline{B}(0_E, M)$, autrement dit tel que : $\forall x \in A, \|x\| \leq M$.

Normes équivalentes

Nous avons vu dans l'exercice 1 que la forme des boules dépend de la norme choisie, en conséquence de quoi la notion de partie bornée dépend *a priori* du choix de la norme. Cependant, on peut constater que dans $\mathbb{R}^2$ et pour les trois normes que nous avons pris en exemple, ce n'est pas le cas : si une partie A est bornée pour une de ces trois norme, elle le sera pour les deux autres (illustration figure 1).

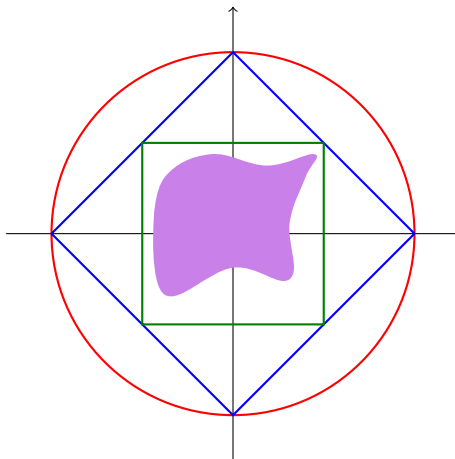


FIGURE 1 – Une partie bornée pour la norme $\|\cdot\|_\infty$ l'est aussi pour les normes $\|\cdot\|_1$ et $\|\cdot\|_2$.

Cette propriété résulte des inégalités suivantes. Pour tout $x \in \mathbb{R}^2$,

- $\|x\|_1 \leq 2\|x\|_\infty$ donc toute partie bornée pour la norme $\|\cdot\|_\infty$ l'est aussi pour la norme $\|\cdot\|_1$;
- $\|x\|_\infty \leq \|x\|_1$ donc toute partie bornée pour la norme $\|\cdot\|_1$ l'est aussi pour la norme $\|\cdot\|_\infty$;

Et de même,

- $\|x\|_2 \leq \sqrt{2}\|x\|_\infty$ donc toute partie bornée pour la norme $\|\cdot\|_\infty$ l'est aussi pour la norme $\|\cdot\|_2$;
- $\|x\|_\infty \leq \|x\|_2$ donc toute partie bornée pour la norme $\|\cdot\|_2$ l'est aussi pour la norme $\|\cdot\|_\infty$;

DÉFINITION. — Deux normes N_1 et N_2 sont dites équivalentes lorsqu'il existe deux réels $\alpha > 0$ et $\beta > 0$ tels que pour tout $x \in E$, $N_1(x) \leq \alpha N_2(x)$ et $N_2(x) \leq \beta N_1(x)$.

PROPOSITION 1.4 — Si deux normes N_1 et N_2 sont équivalentes, toute partie bornée pour l'une de ces deux normes l'est aussi pour l'autre.

Nous admettrons le résultat notable suivant :

THÉORÈME 1.5 — Dans un $\mathbb{K}$ -espace vectoriel de dimension finie, toutes les normes sont équivalentes.

avec pour conséquence immédiate :

COROLLAIRE — Dans un espace vectoriel de dimension finie, la notion de partie bornée est indépendante du choix de la norme.

Attention. On prendra bien garde au fait que l'équivalence des normes n'est valable qu'en dimension finie. Cette hypothèse est primordiale, et a pour conséquence que les différentes notions d'analyse réelle qu'on prolonge au cas d'un espace vectoriel de dimension finie (à commencer par la convergence des suites au paragraphe suivant) ne dépendent pas du choix de la norme utilisée. En revanche, ce théorème est mis en défaut en dimension infinie, avec pour conséquence que dans ces espaces une partie peut être bornée pour une certaine norme, et pas pour d'autres.

1.2 Normes d'opérateurs (notion hors programme)

On appelle *norme matricielle* toute norme sur $\mathcal{M}_n(\mathbb{K})$ qui vérifie en plus la propriété :

$$\forall (A, B) \in \mathcal{M}_n(\mathbb{K})^2, \quad \|AB\| \leq \|A\| \cdot \|B\|.$$

la manière usuelle de définir une norme matricielle consiste à interpréter $A \in \mathcal{M}_n(\mathbb{K})$ comme un endomorphisme de $\mathbb{K}^n$: à partir d'une norme $\|\cdot\|$ sur $\mathbb{K}^n$ on définit sur $\mathcal{M}_n(\mathbb{K})$ la norme :

$$\|A\| = \sup_{x \in \mathbb{K}^n \setminus \{0\}} \frac{\|Ax\|}{\|x\|} = \sup_{\|x\|=1} \|Ax\|.$$

Une telle norme est appelée une *norme d'opérateur*, puisqu'on interprète A comme un opérateur linéaire de $\mathbb{K}^n$ dans lui-même. Certains auteurs parlent de *norme subordonnée* (au choix de la norme sur $\mathbb{K}^n$).

PROPOSITION 1.6 — L'application $A \mapsto \|A\|$ définit une norme matricielle sur $\mathcal{M}_n(\mathbb{K})$.

Exemple. La norme d'opérateur associée à la norme $\|\cdot\|_\infty$ de $\mathbb{K}^n$ est définie par : $\|A\| = \max_{1 \leq i \leq n} \sum_{j=1}^n |a_{ij}|$.

1.3 Suites dans un espace vectoriel normé

DÉFINITION. — On dit qu'une suite (u_n) d'éléments d'un espace vectoriel normé E converge vers $\ell \in E$ lorsque la distance de u_n à ℓ tend vers 0 : $\lim_{n \rightarrow +\infty} \|u_n - \ell\| = 0$.

Géométriquement, cette dernière propriété se traduit ainsi : pour tout $\epsilon > 0$ il existe un rang à partir duquel tous les termes de la suite (u_n) sont dans la boule fermée de centre ℓ de rayon ϵ .

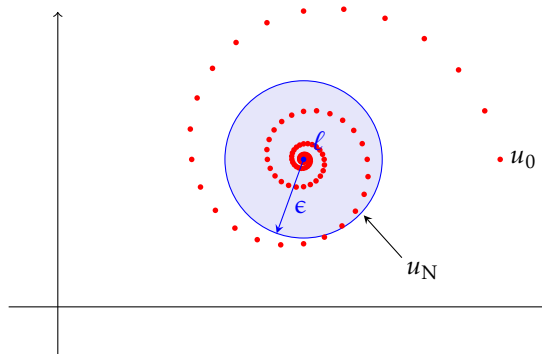


FIGURE 2 – À partir d'un certain rang, tous les termes de la suite (u_n) sont dans $B(\ell, \epsilon)$.

Exercice 2

Soit E un espace vectoriel normé, et (u_n) une suite de vecteurs de E qui converge vers $\ell \in E$. Prouver les propriétés suivantes :

- la suite réelle $(\|u_n\|)$ converge vers $\|\ell\|$;
- la suite (u_n) est bornée.

Remarque. la définition de la convergence dépend *a priori* du choix de la norme utilisée. Cependant, si deux normes sont équivalentes, la convergence pour l'une est équivalente à la convergence pour l'autre. Compte tenu du théorème 1.5, on en déduit :

THÉORÈME 1.7 — Dans un $\mathbb{K}$ -espace vectoriel de dimension finie, la convergence d'une suite et la valeur de la limite ne dépendent pas du choix de la norme.

Remarque. Ce théorème le suggère en creux : en dimension infinie, la notion de convergence dépend du choix de la norme. Et en effet, il est possible de donner des exemples de suites en dimension infinie qui vont converger pour une norme et diverger pour l'autre, voire des exemples de suites qui convergent vers des limites différentes suivant le choix de la norme !

PROPOSITION 1.8 — Soit E un $\mathbb{K}$ -espace vectoriel de dimension finie, $(e_1, \dots, e_p)$ une base de E , (u_n) une suite de vecteurs et $\ell \in E$ un vecteur. On pose :

$$\forall n \in \mathbb{N}, \quad u_n = \sum_{k=1}^p u_{n,k} e_k \quad \text{et} \quad \ell = \sum_{k=1}^p \ell_k e_k.$$

Alors (u_n) converge vers ℓ si et seulement si pour tout $k \in \llbracket 1, p \rrbracket$, la suite $(u_{n,k})$ converge vers ℓ_k .

Autrement dit, en dimension finie l'étude de la convergence d'une suite se ramène à celle de ses coordonnées dans une base.

Exercice 3

Soit $A \in \mathcal{M}_p(\mathbb{K})$ telle que la suite (A^n) converge vers une matrice L . Montrer que L est une matrice de projection.

2. Topologie d'un espace vectoriel normé

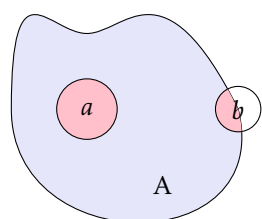
Dans toute cette partie, nous considérons un espace vectoriel E muni d'une norme notée $\|\cdot\|$.

2.1 Ouverts et fermés

Un ensemble *ouvert*, aussi appelé une partie ouverte ou, plus fréquemment, un ouvert, est, de manière informelle, une partie $\mathcal{O}$ de E qui possède la propriété suivante : si a appartient à cet ensemble, $\mathcal{O}$ contiendra aussi tous les points suffisamment proches de a . Quant aux fermés, même si nous ne les définirons pas ainsi, nous verrons qu'il s'agit des parties complémentaires des ouverts.

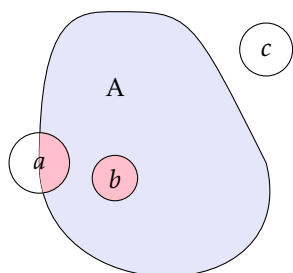
C'est la notion de boule, que nous avons déjà rencontrée, qui permet de définir la notion de *proximité*, ou encore de *voisinage*, dont nous avons besoin pour définir les ouverts.

DÉFINITION. — Soit A une partie de E , et a un point de A . Lorsque A contient une boule (ouverte ou fermée) centrée en a , on dit que a est intérieur à A .



Le point a est intérieur à A , mais pas le point b : quel que soit le rayon de la boule centrée en b , celle-ci ne sera pas incluse dans A .

DÉFINITION. — Un élément $a \in E$ est dit adhérent à une partie A de E lorsque toute boule (ouverte ou fermée) centrée en a contient au moins un point de A : $\forall r > 0, B(a, r) \cap A \neq \emptyset$.



Les points a et b sont adhérents à A , mais pas le point c .

Remarque. Tout point de A est bien entendu adhérent à A .

Exercice 4

Soit A une partie de E et $x \in E$. On pose $d(x, A) = \inf\{\|x - a\| \mid a \in A\}$. Montrer que x est adhérent à A si et seulement si $d(x, A) = 0$.

THÉORÈME 2.1 (caractérisation séquentielle) — Un point $a \in E$ est adhérent à A si et seulement s'il existe une suite (u_n) d'éléments de A qui converge vers a .

■ Intérieur, adhérence et frontière

DÉFINITION. — Lorsque A est une partie quelconque de E , on appelle

- intérieur de A l'ensemble $\mathring{A}$ des points intérieurs à A ;
- adhérence de A l'ensemble $\overline{A}$ des points adhérents à A ;
- frontière de A l'ensemble $\text{Fr}(A) = \overline{A} \setminus \mathring{A}$.

Exemple. Si A est une boule (ouverte ou fermée) de centre a de rayon r , $\mathring{A}$ est la boule ouverte $\mathring{B}(a, r)$, $\overline{A}$ est la boule fermée $\overline{B}(a, r)$, $\text{Fr}(A)$ est la sphère $S(a, r)$.

Exemples. L'adhérence de $\mathbb{Q}$ dans $\mathbb{R}$ est égal à $\mathbb{R}$ car tout nombre réel est limite d'une suite de nombres rationnels. L'intérieur de $\mathbb{Q}$ dans $\mathbb{R}$ est égal à l'ensemble vide car toute boule de rayon $r > 0$ contient des irrationnels.

Remarque. A l'instar de $\mathbb{Q}$ dans $\mathbb{R}$, une partie A d'un espace vectoriel normé E sera dite *dense* dans E lorsque $\overline{A} = E$.

■ Ouverts et fermés

DÉFINITION. — Une partie $\mathcal{O}$ de E est dite *ouverte* lorsque tous ses points sont intérieurs, c'est à dire :

$$\forall x \in \mathcal{O}, \quad \exists r > 0 \mid B(x, r) \subset \mathcal{O}.$$

Autrement dit, un ouvert est une partie égale à son intérieur.

Exemples.

- Les intervalles ouverts sont des ouverts de $\mathbb{R}$;
- toute boule ouverte est un ouvert;
- $\emptyset$ et E sont des ouverts;
- l'intersection ou la réunion de deux ouverts est un ouvert.

DÉFINITION. — Une partie $\mathcal{F}$ de E est dite *fermée* lorsque tout point adhérent à $\mathcal{F}$ appartient à $\mathcal{F}$, soit encore lorsque toute suite d'éléments de $\mathcal{F}$ convergeant dans E a sa limite dans $\mathcal{F}$.

Autrement dit, un fermé est une partie égale à son adhérence.

Exemple.

- Les intervalles fermés de $\mathbb{R}$ sont des fermés;
- toute boule fermée est un fermé;
- toute sphère $S(a, r) = \{x \in E \mid \|x - a\| = r\}$ est un fermé;
- $\emptyset$ et E sont des fermés;
- la réunion ou l'intersection de deux fermés est un fermé.

PROPOSITION 2.2 — Dans un espace vectoriel de dimension finie, les sous-espaces vectoriels de E sont des fermés.

Enfin, le résultat qui suit établit que les notions d'ouvert et de fermé sont indissociables :

THÉORÈME 2.3 — Une partie $\mathcal{F}$ de E est fermée si et seulement si la partie complémentaire $\mathcal{O} = E \setminus \mathcal{F}$ est ouverte.

Exercice 5

Soit A une partie d'un espace vectoriel normé E .

- Montrer que A est ouvert si et seulement si $A \cap \text{Fr}(A) = \emptyset$;
- Montrer que A est fermé si et seulement si $\text{Fr}(A) \subset A$.

2.2 Limite et continuité

Dans cette section, E et F désigneront deux $\mathbb{R}$ -espaces vectoriels normés de dimensions finies, la norme étant notée $\|\cdot\|$ indépendamment de l'espace.

$\mathcal{U}$ désignera une partie de E , et $f : \mathcal{U} \rightarrow F$ une fonction.

■ Étude locale d'une application

DÉFINITION. — Si a désigne un point de E adhérent à $\mathcal{U}$, on dit que $f(x)$ admet $\ell \in F$ pour limite lorsque x tend vers a lorsque :

$$\forall \epsilon > 0, \exists \eta > 0 \mid \forall x \in \mathcal{U}, \quad \|x - a\| \leq \eta \Rightarrow \|f(x) - \ell\| \leq \epsilon.$$

On notera dans ce cas : $\lim_{x \rightarrow a} f(x) = \ell$.

L'existence d'une limite, et la valeur de cette limite, sont des notions qui ne dépendent pas des normes utilisées si on remplace une norme par une norme équivalente, ce qui est toujours le cas en dimension finie.

Les théorèmes généraux relatifs aux opérations algébriques sur les limites se généralisent sans peine, ainsi que celui relatif à la limite d'une application composée.

Enfin, on peut faire le lien avec les suites de vecteurs :

THÉORÈME 2.4 (caractérisation séquentielle) — $f(x)$ admet ℓ pour limite lorsque x tend vers a si et seulement si pour toute suite (a_n) d'éléments de $\mathcal{U}$ qui converge vers a , la suite $(f(a_n))$ converge vers ℓ .

Ce résultat, associé à la proposition 1.8, permet d'en déduire le

COROLLAIRE — $f(x)$ admet ℓ pour limite si et seulement si chacune des composantes de $f(x)$ dans une base arbitraire de E admet pour limite la composante de ℓ dans cette même base.

■ Relations de comparaison

Soit a un point adhérent à $\mathcal{U}$, et $\phi : \mathcal{U} \rightarrow \mathbb{R}$ une fonction à valeurs réelles ne s'annulant pas en dehors de a .

On dit que f est dominée par ϕ au voisinage de a lorsque f/ϕ est bornée au voisinage de a ; on note alors $f(x) = O_a(\phi(x))$.

On dit que f est négligeable devant ϕ au voisinage de a lorsque $\lim_{x \rightarrow a} \frac{f(x)}{\phi(x)} = 0$; on note alors : $f(x) = o_a(\phi(x))$.

Exemples. $f(x) = O_a(1)$ traduit le fait que f est bornée au voisinage de a .

$f(x) = \ell + o_a(1)$ traduit le fait que $\lim_{x \rightarrow a} f(x) = \ell$.

■ Continuité

DÉFINITION. — f est dite continue en $a \in \mathcal{U}$ lorsque $\lim_{x \rightarrow a} f(x) = f(a)$, autrement dit lorsque :

$$\forall \epsilon > 0, \exists \eta > 0 \mid \forall x \in \mathcal{U}, \quad \|x - a\| \leq \eta \Rightarrow \|f(x) - f(a)\| \leq \epsilon.$$

On notera que la décomposition dans une base de F permet de ramener l'étude de la continuité à des fonctions à valeurs réelles : si $(e_1, \dots, e_n)$ est une base de F et $f = f_1 e_1 + \dots + f_n e_n$, f est continue en a si et seulement si les fonctions à valeurs réelles $f_1, \dots, f_n$ sont continues en a .

Continuité sur une partie

Une fonction $f : \mathcal{U} \subset E \rightarrow F$ est dite continue sur $\mathcal{U}$ lorsque f est continue en tout point de $\mathcal{U}$. Bien entendu les propriétés usuelles de la continuité (opérations algébriques, composition, ...) se prolongent sans réelle modification au cas des espaces vectoriels normés.

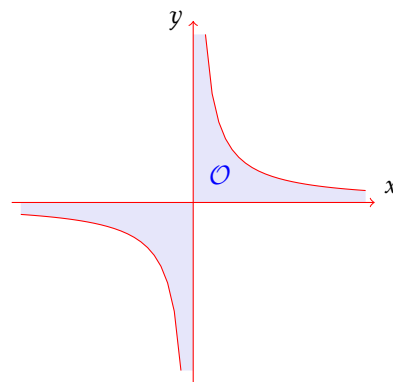
THÉORÈME 2.5 — On considère une fonction $f : E \rightarrow \mathbb{R}$ une fonction continue à valeurs réelles ou complexes.

- Si $\mathcal{O}$ désigne une partie ouverte de $\mathbb{R}$, alors : $f^{-1}(\mathcal{O}) = \{x \in E \mid f(x) \in \mathcal{O}\}$ est un ouvert de E .
- Si $\mathcal{F}$ désigne une partie fermée de $\mathbb{R}$, alors : $f^{-1}(\mathcal{F}) = \{x \in E \mid f(x) \in \mathcal{F}\}$ est un fermé de E .

Ce résultat permet en particulier de donner des exemples simples de parties ouvertes ou fermées. Par exemple, si $f : E \rightarrow \mathbb{R}$ est une fonction continue, et $\alpha \in \mathbb{R}$, la partie $\{x \in E \mid f(x) > \alpha\}$ est ouverte et les parties $\{x \in E \mid f(x) \geq \alpha\}$ et $\{x \in E \mid f(x) = \alpha\}$ fermées.

Exemple. L'ensemble $\mathcal{O} = \{(x, y) \in \mathbb{R}^2 \mid 0 < xy < 1\}$ est un ouvert de $\mathbb{R}^2$.

En effet, l'application $\begin{pmatrix} \mathbb{R}^2 & \longrightarrow & \mathbb{R} \\ (x, y) & \longmapsto & xy \end{pmatrix}$ est continue et $]0, 1[$ est un ouvert de $\mathbb{R}$.



Exercice 6

Soit $A = \{(x_1, \dots, x_n) \in \mathbb{R}^n \mid i \neq j \implies x_i \neq x_j\}$. Montrer que A est un ouvert de $\mathbb{R}^n$.

Nous avons vu que les parties ouvertes possédaient un certain nombre de propriétés communes avec les intervalles ouverts ; il en est de même des parties fermées et des intervalles fermés.

De la même façon, les parties fermées et bornées possèdent des propriétés communes avec les segments, notamment le résultat suivant, que nous admettrons :

THÉORÈME 2.6 (Théorème de la borne atteinte) — Soit $\mathcal{K}$ une partie fermée et bornée d'un $\mathbb{K}$ -espace vectoriel normé E de dimension finie, et $f : \mathcal{K} \rightarrow \mathbb{R}$ une fonction continue. Alors f est bornée et atteint ses bornes.

■ Fonctions lipschitziennes

Une application lipschitzienne est une application possédant une propriété de régularité plus forte que la continuité.

DÉFINITION. — Soient E et F deux $\mathbb{K}$ -espaces vectoriels normés, $\mathcal{U}$ une partie de E , et $k > 0$. Une application $f : \mathcal{U} \rightarrow F$ est dite k -lipschitzienne lorsque :

$$\forall (x, y) \in \mathcal{U}^2, \quad \|f(y) - f(x)\| \leq k \|y - x\|.$$

THÉORÈME 2.7 — Toute application lipschitzienne est continue sur son ensemble de définition.

Exemple. La seconde inégalité triangulaire : $\| \|y\| - \|x\| \| \leq \|y - x\|$ traduit le fait que l'application $x \mapsto \|x\|$ est une application 1-lipschitzienne de E ; il s'agit donc d'une application continue.

2.3 Le cas des applications linéaires

Parmi les applications d'un espace vectoriel vers un autre se trouve en particulier les applications linéaires. Il est légitime de se poser la question de leur continuité. Cette section y répond, en montrant mieux : toute application linéaire est, en dimension finie, lipschitzienne.

Mais tout d'abord, constatons que la définition de cette notion se simplifie dans le cas d'une application linéaire $u \in \mathcal{L}(E, F)$: en effet, u est k -lipschitzienne si et seulement si :

$$\begin{aligned} & \forall (x, y) \in E^2, \quad \|u(y) - u(x)\| \leq k\|y - x\| \\ \iff & \forall (x, y) \in E^2, \quad \|u(y - x)\| \leq k\|y - x\| \\ \iff & \forall z \in E, \quad \|u(z)\| \leq k\|z\| \end{aligned}$$

Nous pouvons maintenant démontrer le :

THÉORÈME 2.8 — Soient E et F deux espaces vectoriels normés de dimensions finies, et $u \in \mathcal{L}(E, F)$. Alors u est lipschitzienne, et donc continue.

■ Applications bilinéaires

Pour finir, un bref mot sur les applications bilinéaires, qui, de manière analogue aux applications linéaires, sont des applications continues en dimension finie.

LEMME — Soient E, F, G trois espaces vectoriels normés de dimensions finies, et $B : E \times F \rightarrow G$ une application bilinéaire. Alors il existe une constante k telle que :

$$\forall (x, y) \in E \times F, \quad \|B(x, y)\| \leq k\|x\| \cdot \|y\|.$$

PROPOSITION 2.9 — Soient E, F, G trois espaces vectoriels normés de dimensions finies, et $B : E \times F \rightarrow G$ une forme bilinéaire. Alors B est continue.

Exemple. Si E est un espace euclidien, l'application $(x, y) \rightarrow \langle x | y \rangle$ est une application continue.

Remarque. Ce résultat s'étend aux fonctions n -linéaires, et en particulier, le déterminant est une application continue de $\mathcal{M}_n(\mathbb{K})$ vers $\mathbb{K}$.

3. Fonctions vectorielles

Dans cette dernière partie, nous allons restreindre l'espace de départ à un intervalle I de $\mathbb{R}$. Ainsi, nous allons nous intéresser plus spécifiquement aux fonctions $f : I \subset \mathbb{R} \rightarrow E$, où E est un espace vectoriel normé de dimension finie. De telles fonctions sont appelées des *fonctions vectorielles*, et cette restriction va nous permettre d'étendre le concept de dérivabilité à de telles fonctions.

3.1 Dérivation des fonctions à valeurs vectorielles

DÉFINITION. — On dit qu'une fonction vectorielle $f : I \rightarrow E$ admet une dérivée en $t_0 \in I$ lorsqu'il existe un vecteur $\ell \in E$ tel que : $\ell = \lim_{t \rightarrow t_0} \frac{f(t) - f(t_0)}{t - t_0}$. On pose alors $f'(t_0) = \ell$.

On a bien entendu de manière équivalente : $f'(t_0) = \lim_{h \rightarrow 0} \frac{f(t_0 + h) - f(t_0)}{h}$.

Les notations de Landau permettent enfin d'exprimer cette définition de la manière suivante : f admet ℓ pour dérivée en $t_0 \in I$ lorsque : $f(t) = f(t_0) + (t - t_0)f'(t_0) + o(t - t_0)$.

PROPOSITION 3.1 — Si f est dérivable en t_0 , alors f est continue en t_0 .

Tout comme la continuité, le recours aux fonctions composantes permet de ramener la dérivation d'une fonction vectorielle à la dérivation des fonctions à valeurs numériques :

PROPOSITION 3.2 — Soit $(e_1, \dots, e_p)$ une base de E , et $f_1, \dots, f_p$ les fonctions coordonnées de f dans cette base. Alors f est dérivable en t_0 si et seulement si les fonctions $f_1, \dots, f_p$ le sont, et dans ce cas :

$$f'(t_0) = f'_1(t_0)e_1 + \dots + f'_p(t_0)e_p.$$

Exemples.

- Une fonction à valeurs complexes $f : I \rightarrow \mathbb{C}$ est dérivable en t_0 si et seulement si les fonctions $\Re f$ et $\Im f$ le sont.
- En cinématique, on obtient les composantes dans une base quelconque du vecteur accélération en dérivant les composantes du vecteur vitesse⁷.

PROPOSITION 3.3 — Soit $f : I \rightarrow E$ une fonction dérivable en t_0 , et $u \in \mathcal{L}(E, F)$ une application linéaire. Alors $u \circ f$ est dérivable en t_0 , et $(u \circ f)'(t_0) = u(f'(t_0))$.

De manière analogue, on peut démontrer le résultat suivant :

PROPOSITION 3.4 — Soient $f : I \rightarrow E$ et $g : I \rightarrow F$ deux applications vectorielles dérivables en t_0 , et $B : E \times F \rightarrow G$ une application bilinéaire. Alors $B(f, g)$ est dérivable en t_0 , et :

$$B(f, g)'(t_0) = B(f'(t_0), g(t_0)) + B(f(t_0), g'(t_0))$$

Exemple. Cette formule généralise bien entendu la formule de dérivation d'un produit fg de deux fonctions à valeurs numériques, mais s'utilise aussi pour dériver une expression faisant intervenir un produit scalaire :

lorsque B est un produit scalaire, on a : $\langle f | g \rangle'(t_0) = \langle f'(t_0) | g(t_0) \rangle + \langle f(t_0) | g'(t_0) \rangle$.

Exercice 7

Soit E un espace euclidien et $f : I \rightarrow E$ une fonction vectorielle dérivable en tout point de I , et telle que $\forall t \in I$, $\|f(t)\| = 1$. Montrer que pour tout $t \in I$, les vecteurs $f(t)$ et $f'(t)$ sont orthogonaux.

Remarque. On peut encore généraliser cette formule au cas d'une application n -linéaire, ce qui est le cas en particulier du déterminant. Ainsi, si $f_1, \dots, f_p$ sont des fonctions définies de I dans E et dérivables en t_0 et (e) une base de E , l'application $\phi : t \mapsto \det_e(f_1(t), \dots, f_p(t))$ est dérivable en t_0 , et :

$$\phi'(t_0) = \sum_{k=1}^p \det_e(f_1(t_0), \dots, f_{k-1}(t_0), f'_k(t_0), f_{k+1}(t_0), \dots, f_p(t_0)).$$

Exercice 8

Soit $A \in \mathcal{M}_n(\mathbb{R})$ une matrice, et $f : \mathbb{R} \rightarrow \mathbb{R}$ la fonction définie par $f(t) = \det(I_n + tA)$. Justifier que f est dérivable en 0, et calculer $f'(0)$.

Enfin, concernant la composée, nous avons :

PROPOSITION 3.5 — Soit I et J deux intervalles, $t_0 \in I$, $\phi : I \rightarrow \mathbb{R}$ une fonction dérivable en t_0 tel que $\phi(I) \subset J$, et $f : J \rightarrow E$ une fonction vectorielle dérivable en $\phi(t_0)$. Alors $f \circ \phi$ est dérivable en t_0 , et :

$$(f \circ \phi)'(t_0) = \phi'(t_0) \times f'(\phi(t_0)).$$

■ Fonction dérivée

DÉFINITION. — Lorsque f est dérivable en tout point de I , on définit une fonction $f' : I \rightarrow E$, appelée fonction dérivée de f . Si f' est à son tour dérivable, on note f'' (ou $f^{(2)}$) sa dérivée, et plus généralement : on note $f^{(0)} = f$, et si $f^{(n)}$ est dérivable, on note $f^{(n+1)}$ sa dérivée.

7. Il faut bien entendu que la base ne soit pas mobile, c'est-à-dire que les vecteurs qui la composent soient indépendants du temps.

On pourra aussi noter $D(f)$ ou $\frac{df}{dt}$ en lieu et en place de f' , et $D^k(f)$ ou $\frac{d^k f}{dt^k}$ pour $f^{(k)}$.

Pour tout entier $n \in \mathbb{N}^*$, on note $\mathcal{C}^n(I, E)$ l'ensemble des fonctions f n fois dérivables sur I , telles que $f^{(n)}$ soit continue. Pour tout entier $n \in \mathbb{N}$, on a : $\mathcal{C}^{n+1}(I, E) \subset \mathcal{C}^n(I, E)$. On pose $\mathcal{C}^\infty(I, E) = \bigcap_{n \in \mathbb{N}} \mathcal{C}^n(I, E)$.

THÉORÈME 3.6 — Si f et g sont des fonctions vectorielles de classe $\mathcal{C}^n$ sur I et B une forme bilinéaire, $B(f, g)$ est aussi de classe $\mathcal{C}^n$, et :

$$B(f, g)^{(n)} = \sum_{k=0}^n \binom{n}{k} B(f^{(k)}, g^{(n-k)}).$$

PROPOSITION 3.7 — Soit I et J deux intervalles, $\phi : I \rightarrow \mathbb{R}$ une fonction numérique de classe $\mathcal{C}^n$ telle que $\phi(I) \subset J$, et $f : J \rightarrow E$ une fonction vectorielle de classe $\mathcal{C}^n$. Alors la fonction vectorielle $f \circ \phi : I \rightarrow E$ est aussi de classe $\mathcal{C}^n$.

Chapitre X

Calcul différentiel

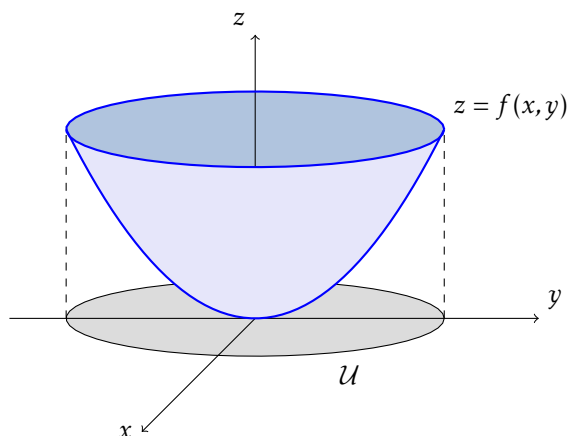
Jusqu'à présent, nous nous sommes cantonnés à l'étude de fonctions d'une *variable*, d'abord à valeurs réelles ou complexes puis à valeurs vectorielles (dans $\mathbb{R}^n$), ces fonctions étaient systématiquement définies sur un intervalle I de $\mathbb{R}$. Nous allons maintenant nous intéresser aux fonctions de *plusieurs variables*, c'est à dire définies sur une partie $\mathcal{U}$ de $\mathbb{R}^p$, à valeurs dans $\mathbb{R}^n$:

$$f : \left(\begin{array}{ccc} \mathcal{U} \subset \mathbb{R}^p & \longrightarrow & \mathbb{R}^n \\ x = (x_1, \dots, x_p) & \longmapsto & f(x) = (f_1(x_1, \dots, x_p), \dots, f_n(x_1, \dots, x_p)) \end{array} \right)$$

Dans ce cours, nous aurons l'occasion, comme pour les fonctions vectorielles, de montrer que l'étude d'une telle fonction se ramène à celle de ses fonctions coordonnées $f_1, \dots, f_n$ et ainsi nous ramener à l'étude des fonctions à valeurs réelles (autrement dit prendre $n = 1$). Pour des raisons pratiques, nos exemples se cantonneront le plus souvent à des fonctions à deux ou trois variables ($p = 2$ ou 3).

Ainsi, lorsque l'on a $p = 2$ et $n = 1$, le graphe $z = f(x, y)$ d'une telle fonction est une nappe paramétrée que l'on peut visualiser et ainsi fournir un support à une interprétation géométrique :

Exemple. $\mathcal{U} = \{(x, y) \in \mathbb{R}^2 \mid x^2 + y^2 < 1\}$, et $f : \left(\begin{array}{ccc} \mathcal{U} & \longrightarrow & \mathbb{R} \\ (x, y) & \longmapsto & x^2 + y^2 \end{array} \right)$



D'un point de vue historique, on peut noter que la notion de fonction à plusieurs variables apparaît très tôt en physique, où l'on étudie souvent des quantités dépendants de plusieurs paramètres. Citons par exemple :

- en mécanique des fluides, la *pression* p est un champ⁸ scalaire qui associe à un point du fluide la pression en ce point ; mathématiquement, cela correspond à une application d'une partie $\mathcal{U}$ de $\mathbb{R}^3$ (ou de $\mathbb{R}^4$, si on tient compte du temps) dans $\mathbb{R}$:

$$p : \left(\begin{array}{ccc} \mathcal{U} & \longrightarrow & \mathbb{R} \\ M = (x, y, z) & \longmapsto & p(M) \end{array} \right)$$

- en électromagnétisme, la *densité de courant* $\vec{j}$ est un champ vectoriel qui associe à tout point de l'espace considéré un vecteur qui décrit le courant électrique qui circule à l'échelle locale ; mathématiquement, cela correspond à une application d'une partie $\mathcal{U}$ de $\mathbb{R}^3$ dans $\mathbb{R}^3$:

$$\vec{j} : \left(\begin{array}{ccc} \mathcal{U} & \longrightarrow & \mathbb{R}^3 \\ M = (x, y, z) & \longmapsto & \vec{j}(M) \end{array} \right)$$

Mais avant de débiter l'étude du concept de différentiabilité, nous allons revenir un instant sur les notions de limite et de continuité, déjà abordées dans le chapitre consacré aux espaces vectoriels normés.

8. en mathématiques, un champ est une application qui associe aux points de l'espace une valeur, scalaire ou vectorielle.

1. Calcul différentiel

Dans cette section, E et F désigneront deux $\mathbb{R}$ -espaces vectoriels normés de dimensions finies, la norme étant notée $\|\cdot\|$ indépendamment de l'espace, mais le plus souvent nous auront $E = \mathbb{R}^2$ ou $\mathbb{R}^3$ et $F = \mathbb{R}$.

$\mathcal{U}$ désignera une partie de E (le plus souvent un ouvert), et $f : \mathcal{U} \rightarrow F$ une fonction à plusieurs variables.

1.1 Étude locale d'une application

Dans le chapitre consacré aux espaces vectoriels normés nous avons donné la définition de la limite de f en un point adhérent à $\mathcal{U}$:

$f(x)$ admet $\ell \in F$ pour limite lorsque x tend vers a lorsque :

$$\forall \epsilon > 0, \exists \eta > 0 \mid \forall x \in \mathcal{U}, \quad \|x - a\| \leq \eta \Rightarrow \|f(x) - \ell\| \leq \epsilon$$

En outre, lorsque $a \in \mathcal{U}$ et $\ell = f(a)$, f est dite continue en a .

Observons sur deux exemples comment se traduit cette définition dans le cadre des fonctions à plusieurs variables.

Exemple. Considérons la fonction $f_1 : \mathbb{R}^2 \setminus \{(0,0)\} \rightarrow \mathbb{R}$ définie par $f_1(x,y) = \frac{x^2 y}{x^2 + y^2}$, et utilisons la norme euclidienne canonique sur $\mathbb{R}^2 : \|(x,y)\|_2 = \sqrt{x^2 + y^2}$.

Sachant que $|x| \leq \|(x,y)\|$ et $|y| \leq \|(x,y)\|$, nous pouvons affirmer que $|f_1(x,y)| \leq \|(x,y)\|$, ce qui implique : $\lim_{(x,y) \rightarrow (0,0)} f_1(x,y) = 0$.

Exemple. Considérons la fonction $f_2 : \mathbb{R}^2 \setminus \{(0,0)\} \rightarrow \mathbb{R}$ définie par $f_2(x,y) = \frac{x^2 y}{x^4 + y^2}$.

Supposons que cette fonction possède une limite ℓ en $(0,0)$. Les théorèmes de composition des limites impliquent que pour toute fonction vectorielle $\phi : I \subset \mathbb{R} \rightarrow \mathcal{U}$ pour laquelle $\lim_{t \rightarrow 0} \phi(t) = (0,0)$ nous avons : $\lim_{t \rightarrow 0} f_2 \circ \phi(t) = \ell$.

Or $\lim_{t \rightarrow 0} f_2(t,t) = \lim_{t \rightarrow 0} \frac{t}{t^2 + 1} = 0$ et $\lim_{t \rightarrow 0} f_2(t,t^2) = \frac{1}{2}$. Ainsi f_2 ne possède pas de limite en $(0,0)$.

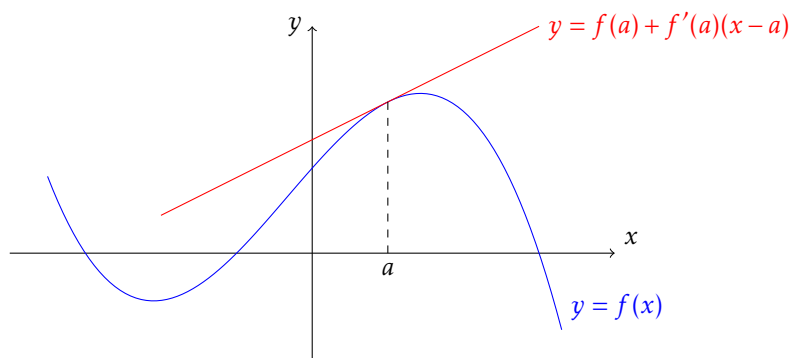
Exercice 1

Déterminer si les fonctions suivantes, définies sur $\mathbb{R}^2 \setminus \{(0,0)\}$, ont une limite finie en $(0,0)$:

$$f(x,y) = \frac{x^2 - y^2}{x^2 + y^2}, \quad g(x,y) = \frac{|x+y|}{x^2 + y^2}, \quad h(x,y) = (x+y) \sin\left(\frac{1}{x^2 + y^2}\right).$$

1.2 Applications différentiables

Pour comprendre comment nous allons généraliser la notion de dérivée aux fonctions à plusieurs variables, observons l'interprétation géométrique qu'on peut faire de la dérivée en $a \in I$ d'une fonction à une variable $f : I \rightarrow \mathbb{R}$.



Au voisinage de a , le graphe de f est approché par une droite, sa tangente. Autrement dit, f est localement approchée par la fonction affine $x \mapsto f(a) + f'(a)(x - a)$, ce qui se traduit par le développement limité suivant :

$$f(x) \underset{a}{=} f(a) + f'(a)(x - a) + o(x - a).$$

qu'on peut écrire de façon équivalente :

$$f(a + h) \underset{0}{=} f(a) + f'(a)h + o(h).$$

Cette approximation affine est formée d'une constante $f(a)$ et d'une application linéaire $h \mapsto f'(a)h$. Ceci nous conduit à adopter la définition suivante :

DÉFINITION. — Soient E et F deux $\mathbb{R}$ -espaces vectoriels normés de dimensions finies, $\mathcal{U}$ un ouvert de E et $f : \mathcal{U} \rightarrow F$ une application.

On dira que f est différentiable en $a \in \mathcal{U}$ lorsqu'il existe une application linéaire $u \in \mathcal{L}(E, F)$ telle que :

$$f(a + h) \underset{0_E}{=} f(a) + u(h) + o(\|h\|).$$

Dans ce cas, l'application linéaire u est appelée la *différentielle* de f en a , et sera notée $df(a)$. Ainsi, on écrira :

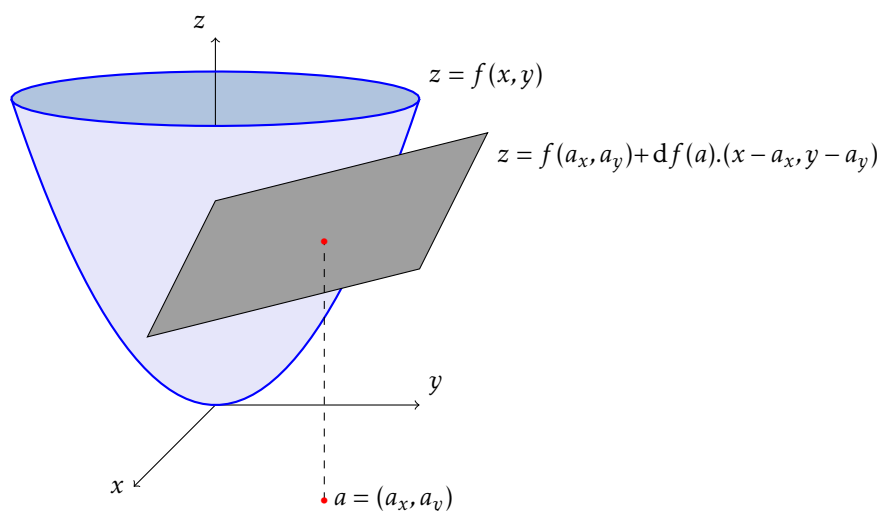
$$\boxed{f(a + h) \underset{0_E}{=} f(a) + df(a).h + o(\|h\|)} \quad \text{ou encore :} \quad \boxed{f(x) \underset{a}{=} f(a) + df(a).(x - a) + o(\|x - a\|)}$$

Notons que la différentiabilité de f en a entraîne *a fortiori* la continuité de f en a , puisqu'une application linéaire est continue (en dimension finie).

Exemple. Dans le cas où $E = \mathbb{R}^2$ et $F = \mathbb{R}$, les fonctions affines prennent la forme suivante :

$$\tilde{u} : \begin{pmatrix} \mathbb{R}^2 & \mapsto & \mathbb{R} \\ (x, y) & \mapsto & \alpha + \beta x + \gamma y \end{pmatrix}$$

et le graphe de cette fonction affine est le plan affine d'équation $z = \alpha + \beta x + \gamma y$. Autrement dit, la nappe d'équation $z = f(x, y)$ est localement approchée par un plan.



Exercice 2

Soit $f : \mathcal{M}_n(\mathbb{R}) \rightarrow \mathcal{M}_n(\mathbb{R})$ l'application définie par $f(M) = M^2$. Montrer que f est différentiable en tout point $A \in \mathcal{M}_n(\mathbb{R})$, et déterminer l'application linéaire $df(A)$.

Remarque. Lorsque f est une application linéaire, l'égalité $f(x) = f(a) + f(x - a)$ montre que sa différentielle en a est égale à elle-même : pour tout $a \in E$, $df(a) = f$.

différence entre dérivée et différentielle

Dans le cas des fonctions numériques, la dérivée de f en a est le réel $f'(a)$, alors que la différentielle de f en a est l'application linéaire $x \mapsto f'(a)x$. En effet, les applications linéaires de $\mathbb{R}$ dans $\mathbb{R}$ s'écrivent de manière unique sous la forme : $x \mapsto \lambda x$, avec $\lambda \in \mathbb{R}$.

On peut se rapprocher encore de la notion de dérivée lorsque $F = \mathbb{R}$: dans ce cas, la différentielle $df(a)$ de f en a est une application linéaire de E dans $\mathbb{R}$, c'est à dire une *forme linéaire*. Or lorsque E est un espace euclidien, nous avons vu que les formes linéaires sur E s'écrivent de manière unique sous la forme $x \mapsto \langle \ell | x \rangle$, avec $\ell \in E$. Cela conduit à la définition :

DÉFINITION. — Lorsque $f : \mathcal{U} \subset E \rightarrow \mathbb{R}$ est différentiable en $a \in E$, il existe un unique vecteur de E , noté $\nabla f(a)$ tel que :

$$\forall h \in E, \quad df(a).h = \langle \nabla f(a) | h \rangle.$$

Le vecteur $\nabla f(a)$ est appelé le gradient de f en a .

Exercice 3

Soit E un espace euclidien, et $f : E \rightarrow \mathbb{R}$ défini par $f(x) = \langle x | x \rangle$. Montrer que f est différentiable en tout $a \in E$, et déterminer le vecteur $\nabla f(a)$.

1.3 Dérivée directionnelle et dérivées partielles

Nous allons maintenant nous attacher à voir comment calculer une différentielle (ou un gradient) dans le cas où $E = \mathbb{R}^p$ et $F = \mathbb{R}$. L'idée que nous allons suivre est d'essayer, autant que faire se peut, de se ramener à des calculs de dérivées de fonctions d'une seule variable.

Sachant que f est définie sur un ouvert $\mathcal{U}$ de $\mathbb{R}^p$, l'application partielle $t \mapsto f(a + tv)$ est, quel que soit le vecteur $v \in \mathbb{R}^p$, $v \neq 0$, une fonction vectorielle définie au voisinage de 0. Lorsque cette fonction est dérivable en 0, on note $D_v f(a)$ cette quantité, qu'on appelle la *dérivée en a selon le vecteur v* .

En particulier, lorsque $v = e_k$ est le k^e vecteur de la base canonique de $\mathbb{R}^p$, cette application prend la forme suivante :

$$t \mapsto f(a_1, a_2, \dots, a_{k-1}, a_k + t, a_{k+1}, \dots, a_p)$$

et sa dérivée en 0 est notée $\partial_k f(a)$ ou $\frac{\partial f}{\partial x_k}(a)$, quantité qu'on appelle la k^e *dérivée partielle d'ordre 1* de f en a . Autrement dit :

$$\forall k \in \llbracket 1, p \rrbracket, \quad \partial_k f(a) = \lim_{t \rightarrow 0} \frac{f(a_1, \dots, a_k + t, \dots, a_p) - f(a_1, \dots, a_k, \dots, a_p)}{t}.$$

PROPOSITION 1.1 — Lorsque f est différentiable en a , f admet en a des dérivées partielles d'ordre 1 et pour tout

$$h = (h_1, \dots, h_p) \in \mathbb{R}^p, \quad df(a).h = \sum_{k=1}^p \partial_k f(a) h_k. \quad \text{Autrement dit,} \quad \nabla f(a) = (\partial_1 f(a), \dots, \partial_p f(a)).$$

La réciproque de ce résultat est fautive : une fonction peut posséder des dérivées partielles en a sans être différentiable en ce point. Cependant, en renforçant un peu les hypothèses on dispose du résultat suivant, avec lequel nous allons maintenant justifier l'existence de la différentielle :

THÉORÈME 1.2 — Soit $\mathcal{U}$ un ouvert de $\mathbb{R}^p$, et $f : \mathcal{U} \rightarrow \mathbb{R}$ une fonction. On suppose que :

- (i) f possède pour tout $k \in \llbracket 1, p \rrbracket$ et en tout point a de $\mathcal{U}$ une dérivée partielle $\partial_k f(a)$;
- (ii) pour tout $k \in \llbracket 1, p \rrbracket$, l'application $\partial_k f : \mathcal{U} \rightarrow \mathbb{R}$ est continue sur $\mathcal{U}$.

Alors f est différentiable en tout point a de $\mathcal{U}$.

Une application f vérifiant ces hypothèses sera dorénavant dite de classe $\mathcal{C}^1$ sur $\mathcal{U}$.

Exemple. Soit $f : \mathbb{R}^2 \rightarrow \mathbb{R}$ défini par $f(x, y) = x^2 y$. f admet en tout point (x, y) des dérivées partielles $\frac{\partial f}{\partial x}(x, y) = 2xy$ et $\frac{\partial f}{\partial y}(x, y) = x^2$ à l'évidence continues, donc f est de classe $\mathcal{C}^1$ et $\nabla f(x, y) = 2xy\vec{e}_1 + x^2\vec{e}_2$.

La différentielle s'écrit donc $df(x, y) : (h_1, h_2) \mapsto 2xyh_1 + x^2h_2$.

Exercice 4

On considère la fonction $f : \mathbb{R}^2 \rightarrow \mathbb{R}$ définie par $f(x, y) = \frac{x^2 y}{x^2 + y^2}$ si $(x, y) \neq (0, 0)$ et $f(0, 0) = 0$. La fonction f est-elle de classe $\mathcal{C}^1$?

Expression de la dérivée directionnelle pour une fonction de classe $\mathcal{C}^1$

Lorsque f est de classe $\mathcal{C}^1$, on a donc $f(a + tv) \underset{t \rightarrow 0}{=} f(a) + df(a).(tv) + o(\|tv\|) \underset{t \rightarrow 0}{=} f(a) + t df(a).(v) + o(t)$ donc

$$D_v f(a) = \lim_{t \rightarrow 0} \frac{f(a + tv) - f(a)}{t} = df(a).(v)$$

1.4 Règle de la chaîne

Considérons une fonction $f : \mathcal{U} \subset \mathbb{R}^p \rightarrow \mathbb{R}$ de classe $\mathcal{C}^1$, I un intervalle et $x_k : I \rightarrow \mathbb{R}$, $1 \leq k \leq p$ des fonctions elles aussi de classe $\mathcal{C}^1$ telles que pour tout $t \in I$, $(x_1(t), \dots, x_p(t)) \in \mathcal{U}$. On peut dès lors définir la fonction $\phi : I \rightarrow \mathbb{R}$ par $\phi(t) = f(x_1(t), \dots, x_p(t))$.

Le résultat suivant, appelé *règle de la dérivation en chaîne*, ou plus simplement *règle de la chaîne*, donne la règle de dérivation d'une telle fonction.

THÉORÈME 1.3 (règle de la chaîne) — Si f et toutes les fonctions $x_1, \dots, x_p$ sont de classe $\mathcal{C}^1$, il en est de même de la fonction ϕ , et

$$\forall t \in I, \quad \phi'(t) = \sum_{k=1}^p x'_k(t) \partial_k f(x_1(t), \dots, x_p(t)) = \sum_{k=1}^p x'_k(t) \frac{\partial f}{\partial x_k}(x_1(t), \dots, x_p(t)).$$

Exercice 5

Soit $f : \mathbb{R}^2 \rightarrow \mathbb{R}$ une fonction de classe $\mathcal{C}^1$, et $g : \mathbb{R}^2 \rightarrow \mathbb{R}$ définie par : $\forall (r, \theta) \in \mathbb{R}^2$, $g(r, \theta) = f(r \cos \theta, r \sin \theta)$. Calculer les dérivées partielles de g en fonction de celles de f , et en déduire l'expression du gradient en coordonnées polaires.

PROPOSITION 1.4 — Soit $\mathcal{U}$ un ouvert convexe de $\mathbb{R}^p$, et $f : \mathcal{U} \rightarrow \mathbb{R}$ une fonction de classe $\mathcal{C}^1$. Alors f est constante sur $\mathcal{U}$ si et seulement si pour tout $a \in \mathcal{U}$, $df(a) = 0$, autrement dit si et seulement si les fonctions $\partial_1 f, \dots, \partial_p f$ sont nulles sur $\mathcal{U}$.

1.5 Dérivées partielles d'ordre deux

Soit $\mathcal{U}$ un ouvert de $\mathbb{R}^p$, et $f : \mathcal{U} \rightarrow \mathbb{R}$ une fonction de classe $\mathcal{C}^1$. Les applications $\partial_i f = \frac{\partial f}{\partial x_i}$ sont des applications définies et continues de $\mathcal{U}$ dans $\mathbb{R}$ et à ce titre peuvent être elles-mêmes de classe $\mathcal{C}^1$. Lorsque c'est le cas, on note $\partial_j \partial_i f = \frac{\partial^2 f}{\partial x_j \partial x_i}$ la dérivée partielle par rapport à la j^{e} variable de $\partial_i f$.

Remarque. Dans le cas particulier où $i = j$ on utilisera plutôt la notation $\partial_i^2 f$ ou $\frac{\partial^2 f}{\partial x_i^2}$.

DÉFINITION. — Une application $f : \mathcal{U} \rightarrow \mathbb{R}$ est dite de classe $\mathcal{C}^2$ lorsqu'elle est de classe $\mathcal{C}^1$ et lorsque pour tout $i \in \llbracket 1, p \rrbracket$, $\partial_i f$ est de classe $\mathcal{C}^1$ sur $\mathcal{U}$.

A priori, l'expression $\partial_j \partial_i f$ signifie que l'on dérive d'abord par rapport à x_i , puis par rapport à x_j . cependant, le théorème suivant, que nous admettrons, montre qu'il n'en est rien dans le cas d'une fonction de classe $\mathcal{C}^2$.

THÉORÈME 1.5 (Théorème de Schwarz) — Soit $f : \mathcal{U} \subset \mathbb{R}^p \rightarrow \mathbb{R}$ une fonction de classe $\mathcal{C}^2$. Alors pour $(i, j) \in \llbracket 1, p \rrbracket^2$, $\partial_j \partial_i f = \partial_i \partial_j f$.

Exercice 6

Soit $f : \mathbb{R}^2 \rightarrow \mathbb{R}$ une fonction de classe $\mathcal{C}^2$, et $g : \mathbb{R}^2 \rightarrow \mathbb{R}$ définie par : $\forall (r, \theta) \in \mathbb{R}^2, g(r, \theta) = f(r \cos \theta, r \sin \theta)$.

a. Calculer les dérivées partielles secondes de g en fonction de celles de f .

On appelle *laplacien* de f la quantité : $\Delta f = \frac{\partial^2 f}{\partial x^2} + \frac{\partial^2 f}{\partial y^2}$.

b. Dédurre des calculs précédents l'expression du laplacien en coordonnées polaires (c'est à dire en fonction des dérivées de g).

■ Équations aux dérivées partielles

En sciences physiques il est fréquent d'avoir à résoudre une équation mêlant les dérivées partielles d'ordre 1 ou 2 d'une même fonction. Le phénomène de propagation des ondes peut par exemple être modélisé par l'équation de d'Alembert $\Delta f - \frac{1}{c^2} \frac{\partial^2 f}{\partial t^2} = 0$. En dimension 1, cette équation s'écrit $\frac{\partial^2 f}{\partial x^2}(x, t) - \frac{1}{c^2} \frac{\partial^2 f}{\partial t^2}(x, t) = 0$.

Il n'y a pas de méthode générale de résolution, mais celle-ci passe le plus souvent par l'utilisation d'un changement de variable.

On appelle *changement de variable de classe $\mathcal{C}^k$* une application bijective $\phi : \mathcal{U} \rightarrow \mathcal{V}$ entre deux ouverts $\mathcal{U}$ et $\mathcal{V}$ de $\mathbb{R}^p$ telle que ϕ et ϕ^{-1} soient toutes deux de classe $\mathcal{C}^k$. Nous nous restreindrons à deux types de changement de variable : les changements de variables affines et le changement de variable en coordonnées polaires.

Un changement de variable *affine* consiste simplement à poser $\begin{cases} u = ax + by + e \\ v = cx + dy + f \end{cases}$ qui se traduit matriciellement par $Y = AX + B$ avec $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$ et $B = \begin{pmatrix} e \\ f \end{pmatrix}$. Pour qu'il soit bijectif, il faut et il suffit que $\det A \neq 0$, et il s'agit alors d'un changement de variable de classe $\mathcal{C}^\infty$.

Par exemple, pour résoudre l'équation de propagation en dimension 1, on posera $\begin{cases} u = x + ct \\ v = x - ct \end{cases}$ et $g(u, v) = f(x, t)$.

Cette dernière équation introduit une nouvelle fonction inconnue g , et peut être comprise de deux façons :

- $g(u, v) = f\left(\frac{u+v}{2}, \frac{u-v}{2c}\right)$ (pour définir g à partir de f);
- $f(x, t) = g(x+ct, x-ct)$ (pour retrouver f une fois déterminé g).

On calcule successivement à l'aide de la règle de la chaîne :

$$\begin{aligned} \frac{\partial f}{\partial x}(x, t) &= \frac{\partial g}{\partial u}(x+ct, x-ct) + \frac{\partial g}{\partial v}(x+ct, x-ct) & \frac{\partial f}{\partial t}(x, t) &= c \frac{\partial g}{\partial u}(x+ct, x-ct) - c \frac{\partial g}{\partial v}(x+ct, x-ct) \\ \frac{\partial^2 f}{\partial x^2}(x, t) &= \frac{\partial^2 g}{\partial u^2}(u, v) + 2 \frac{\partial^2 g}{\partial u \partial v}(u, v) + \frac{\partial^2 g}{\partial v^2}(u, v) & \frac{\partial^2 f}{\partial t^2}(x, t) &= c^2 \frac{\partial^2 g}{\partial u^2}(u, v) - 2c^2 \frac{\partial^2 g}{\partial u \partial v}(u, v) + c^2 \frac{\partial^2 g}{\partial v^2}(u, v) \end{aligned}$$

et alors : $\frac{\partial^2 f}{\partial x^2}(x, t) - \frac{1}{c^2} \frac{\partial^2 f}{\partial t^2}(x, t) = 0 \iff 4 \frac{\partial^2 g}{\partial u \partial v}(u, v) = 0$.

Cette équation est maintenant aisément résoluble : la fonction $\frac{\partial g}{\partial v}$ est indépendante de u donc ne dépend que de v . La fonction g s'écrit donc sous la forme $g(u, v) = \phi(u) + \psi(v)$, où ϕ et ψ sont deux fonctions quelconques de classe $\mathcal{C}^2$. Il reste à revenir à f en concluant que $f(x, t) = \phi(x+ct) + \psi(x-ct)$.

Le changement de variables en coordonnées polaires

En toute rigueur, pour que le changement de variable $\begin{cases} x = r \cos \theta \\ y = r \sin \theta \end{cases}$ soit bijectif et de classe $\mathcal{C}^\infty$ ainsi que sa réciproque, il est nécessaire d'imposer $r \in]0, +\infty[$ et $\theta \in]\alpha, \alpha + 2\pi[$, ce qui restreint (x, y) à $\mathbb{R}^2 \setminus \mathcal{D}_\alpha$, où $\mathcal{D}_\alpha$ est la demi-droite fermée issue de l'origine et d'angle α par rapport à l'axe Ox .

Une fois ceci précisé, l'équation $f(x, y) = g(r, \theta)$ peut s'écrire $g(r, \theta) = f(r \cos \theta, r \sin \theta)$. En revanche, on observera que faute d'une expression simple de θ en fonction de x et de y , il sera plus difficile d'exprimer $f(x, y)$ en fonction de g , x et y , aussi va-t-on dans ce cas calculer les dérivées partielles de g en fonction de celles de f :

$$\frac{\partial g}{\partial r}(r, \theta) = \cos \theta \frac{\partial f}{\partial x}(x, y) + \sin \theta \frac{\partial f}{\partial y}(x, y) \quad \text{et} \quad \frac{\partial g}{\partial \theta}(r, \theta) = -r \sin \theta \frac{\partial f}{\partial x}(x, y) + r \cos \theta \frac{\partial f}{\partial y}(x, y)$$

Par exemple, l'équation aux dérivées partielles $y \frac{\partial f}{\partial x}(x, y) - x \frac{\partial f}{\partial y}(x, y) = 0$ s'écrit $\frac{\partial g}{\partial \theta}(r, \theta) = 0$ en coordonnées polaires, soit $g(r, \theta) = \phi(r)$ où ϕ est une fonction de classe $\mathcal{C}^1$, et $f(x, y) = \phi(\sqrt{x^2 + y^2})$.

De même, l'équation $x \frac{\partial f}{\partial x}(x, y) + y \frac{\partial f}{\partial y}(x, y) = 0$ s'écrit $r \frac{\partial g}{\partial r}(r, \theta) = 0$ en coordonnées polaires, soit $g(r, \theta) = \psi(\theta)$ où ψ est une fonction de classe $\mathcal{C}^1$. Ici, faute d'une expression convenable pour la fonction θ , on se contentera de conclure que $f(x, y) = \psi(\theta(x, y))$.

■ Matrice Hessienne

Nous avons montré que la formule $f(a+h) \underset{h \rightarrow 0}{=} f(a) + hf'(a) + o(h)$ valable pour une fonction numérique de classe $\mathcal{C}^1$ se généralise au cas d'une fonction $f: \mathcal{U} \subset \mathbb{R}^p \rightarrow \mathbb{R}$ de classe $\mathcal{C}^1$ par la formule :

$$f(a+h) \underset{h \rightarrow 0}{=} f(a) + \langle \nabla f(a) | h \rangle + o(\|h\|) \underset{h \rightarrow 0}{=} f(a) + \nabla f(a)^T h + o(\|h\|)$$

en identifiant les vecteurs de $\mathbb{R}^p$ et les matrices colonnes de $\mathcal{M}_{p,1}(\mathbb{R})$.

Nous admettrons que la formule $f(a+h) \underset{h \rightarrow 0}{=} f(a) + hf'(a) + \frac{h^2}{2} f''(a) + o(h^2)$ se généralise pour une fonction $f: \mathcal{U} \subset \mathbb{R}^p \rightarrow \mathbb{R}$ de classe $\mathcal{C}^2$ par la formule :

$$f(a+h) \underset{h \rightarrow 0}{=} f(a) + \nabla f(a)^T h + \frac{1}{2} h^T H_f(a) h + o(\|h\|^2)$$

où $H_f(a) \in \mathcal{M}_p(\mathbb{R})$ est la matrice des dérivées partielles secondes de f en a , autrement dit :

$$[H_f(a)]_{i,j} = \partial_i \partial_j f(a)$$

On notera que cette matrice, appelée *matrice hessienne* de f en a , est une matrice symétrique (d'après le théorème de Schwarz).

1.6 Extremums locaux

Dans cette section, nous considérons un domaine (non nécessairement ouvert) $\mathcal{U}$ de $\mathbb{R}^p$ ainsi qu'une fonction $f: \mathcal{U} \subset \mathbb{R}^p \rightarrow \mathbb{R}$.

DÉFINITION. — On dit que f présente en un point $a \in \mathcal{U}$ un maximum local lorsqu'il existe un réel $r > 0$ tel que pour tout $x \in B(a, r) \cap \mathcal{U}$, $f(x) \leq f(a)$.

On dit que f présente en $a \in \mathcal{U}$ un maximum global lorsque pour tout $x \in \mathcal{U}$, $f(x) \leq f(a)$.

On définit de la même façon les notions de minimum local et de minimum global.

Remarque. De cette définition il résulte immédiatement que tout extremum global est un extremum local, la réciproque n'étant bien évidemment pas vraie.

THÉORÈME 1.6 — Soit $\mathcal{U}$ un ouvert de $\mathbb{R}^p$, $f : \mathcal{U} \rightarrow \mathbb{R}$ une fonction de classe $\mathcal{C}^1$, et $a \in \mathcal{U}$ un point en lequel f présente un extremum local. Alors : $df(a) = 0$ (la forme linéaire nulle).

Remarque. La condition $df(a) = 0$, qui peut s'écrire $\nabla f(a) = 0_E$, est donc une condition *nécessaire* mais non *suffisante* pour que f présente un extremum local en a dans l'ouvert $\mathcal{U}$.

Un point a en lequel $\nabla f(a) = 0_E$ est appelé un *point critique* de f .

Évidemment, la question de la réciproque se pose : un point critique est-il nécessairement un extremum local ? La réponse est négative : ne serait-ce qu'en dimension 1, la fonction $t \mapsto t^3$ présente en 0 un point critique qui n'est pas un extremum local.

Il va néanmoins être possible par l'affirmative dans certains cas, grâce à la formule de Taylor à l'ordre 2. Ainsi, supposons f de classe $\mathcal{C}^2$, et considérons un point critique $a \in \mathcal{U}$ de f . Nous avons alors :

$$f(a+h) = f(a) + \frac{1}{2} h^T H_f(a) h + o(\|h\|^2)$$

Nous savons par ailleurs que la matrice hessienne $H_f(a)$ est symétrique ; elle est donc ortho-diagonalisable. Soit donc (e) une base orthonormée formée de vecteurs propres de $H_f(a)$. Si on décompose le vecteur h dans cette

$$\text{base : } h = \sum_{k=1}^p h_k e_k \text{ alors } h^T H_f(a) h = \sum_{k=1}^p \lambda_k h_k^2.$$

THÉORÈME 1.7 — Si f est de classe $\mathcal{C}^2$ et a un point critique de f , alors :

- si $\text{Sp}(H_f(a)) \subset \mathbb{R}_+^*$ (autrement dit si $H_f(a) \in \mathcal{S}_p^{++}(\mathbb{R})$), f présente en a un minimum local strict ;
- si $\text{Sp}(H_f(a)) \subset \mathbb{R}_-^*$, f présente en a un maximum local strict ;
- si $H_f(a)$ possède deux valeurs propres de signes différents alors f ne présente pas d'extremum en a .

Remarque. Notons que cette étude n'est pas exhaustive ; en particulier lorsque la matrice hessienne n'est pas inversible, $H_f(a)$ admet 0 pour valeur propre et on ne peut conclure.

Le cas de la dimension 2

Lorsque $p = 2$, posons $H_f(a) = \begin{pmatrix} r & s \\ s & t \end{pmatrix}$; autrement dit $r = \partial_1^2 f(a)$, $s = \partial_1 \partial_2 f(a)$ et $t = \partial_2^2 f(a)$.

Les deux valeurs propres λ et μ de $H_f(a)$ vérifient $\lambda + \mu = \text{tr } H_f(a) = r + t$ et $\lambda\mu = \det H_f(a) = rt - s^2$ donc ces deux valeurs propres sont non nulles et de même signe si et seulement si $\det H_f(a) > 0$. Ainsi, f présente en a un extremum local strict si et seulement si $\det H_f(a) > 0$, et cet extremum est :

- un minimum si $\text{tr } H_f(a) > 0$;
- un maximum si $\text{tr } H_f(a) < 0$.

Exemple. Considérons la fonction $f : \mathbb{R}^2 \rightarrow \mathbb{R}$ définie par $f(x, y) = x^3 + y^3 - 3xy$. Ses points critiques vérifient :

$$\begin{cases} \partial_1 f(x, y) = 0 \\ \partial_2 f(x, y) = 0 \end{cases} \iff \begin{cases} 3x^2 - 3y = 0 \\ 3y^2 - 3x = 0 \end{cases}$$

et la résolution de ce système donne deux points critiques $a = (0, 0)$ et $b = (1, 1)$.

On calcule $\partial_1^2 f(x, y) = 6x$, $\partial_1 \partial_2 f(x, y) = 6y$ et $\partial_2^2 f(x, y) = -3$ donc $H_f(a) = \begin{pmatrix} 0 & -3 \\ -3 & 0 \end{pmatrix}$ et $H_f(b) = \begin{pmatrix} 6 & -3 \\ -3 & 6 \end{pmatrix}$.

$\text{Sp } H_f(a) = \{-3, 3\}$ donc f ne présente pas d'extremum local en a (il s'agit d'un point selle).

$\text{Sp } H_f(b) = \{3, 9\}$ donc f présente en b un minimum local.

Exercice 7

Déterminer les points critiques de la fonction $f : (x, y) \mapsto y(x^2 + (\ln y)^2)$ sur $\mathbb{R} \times \mathbb{R}_+^*$, puis déterminer s'il s'agit d'extremums locaux ou pas.

■ Recherche d'extremums globaux sur une partie fermée bornée de $\mathbb{R}^p$

Considérons maintenant une partie $\mathcal{K}$ fermée et bornée de $\mathbb{R}^p$ et une application $f : \mathcal{K} \rightarrow \mathbb{R}$ continue, sur $\mathcal{K}$, de classe $\mathcal{C}^1$ sur $\overset{\circ}{\mathcal{K}}$.

Dans le chapitre consacré aux espaces vectoriels normés, nous avons admis le résultat suivant, qui va nous être de nouveau utile :

Rappel. Si $\mathcal{K}$ est une partie fermée et bornée de $\mathbb{R}^p$ et $f : \mathbb{R}^p \rightarrow \mathbb{R}$ une fonction continue, alors f est bornée et atteint ses bornes sur $\mathcal{K}$.

Ce résultat assure l'existence d'un minimum et d'un maximum global sur $\mathcal{K}$. Ces deux extremums se trouvent ou bien sur la frontière $\text{Fr}(\mathcal{K})$ de $\mathcal{K}$, ou bien dans l'intérieur $\overset{\circ}{\mathcal{K}} = \mathcal{K} \setminus \text{Fr}(\mathcal{K})$ de $\mathcal{K}$. En d'autres termes, les extremums globaux sont à chercher :

- sur la frontière de $\mathcal{K}$;
- et parmi les points critiques de l'intérieur de $\mathcal{K}$.

Exercice 8

Soit $\mathcal{K} = \{(x, y) \in \mathbb{R}^2 \mid x \geq 0, y \geq 0, x + y \leq 1\}$, et $f : \mathcal{K} \rightarrow \mathbb{R}$ définie par $f(x, y) = xy(1 - x - y)$. Déterminer la valeur maximale prise par la fonction f .