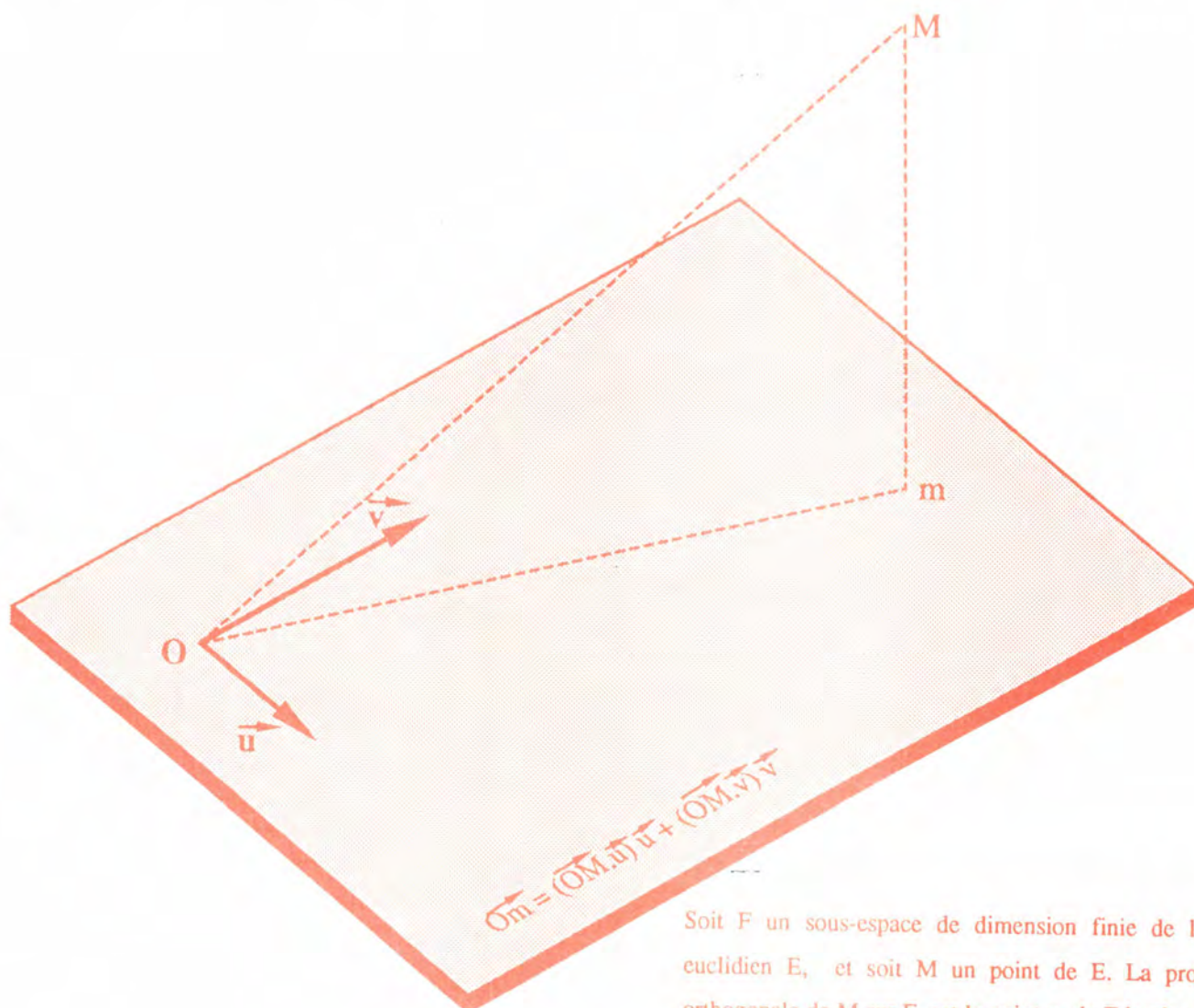




Soit F un sous-espace de dimension finie de l'espace euclidien E , et soit M un point de E . La projection orthogonale de M sur F , est le point m de F le plus proche de M . Si (e_i) est une base orthonormée de F , alors m est donné par la formule: $m = \sum_i \langle M, e_i \rangle e_i$

S.M.2

M.P.I.

Cours de Mathématiques

Année 1989-90

Fascicule 1

Tirage 1992



Avertissement

Ce cours a été conçu pour une section de seconde année de DEUG SM, à l'université Nancy I, en 1989-90. Il s'agit d'une section (dénommée Math-Physique-Informatique) à horaire de mathématiques renforcé (180 à 200h, options exclues). La majorité des étudiants concernés se destinent à une licence de mathématiques, quelques uns veulent entreprendre des études de physique, d'informatique, de mécanique, ou entrer dans une école d'ingénieurs. Il s'agit donc d'un public qui, n'ayant pas été habitué à l'abstraction au niveau du lycée, sera appelé à manipuler dès l'année suivante des notions purement théoriques. Notons que beaucoup d'étudiants suivent en outre une option probabilité, et une option algèbre et géométrie; ceci explique que certaines structures fondamentales - comme les groupes par exemple - n'apparaissent pas ici.

BIBLIOTHEQUE IREM



106492

La règle du jeu

L'enseignement de mathématiques de SM2 MPI ne comporte aucun cours magistral; les étudiants sont répartis en groupes, chacun de ceux-ci est confié à un unique enseignant qui assure la totalité de l'horaire. Ce cours écrit a d'abord pour fonction d'harmoniser le travail des groupes; il contient ce que vous devrez savoir pour l'examen. Il devrait donc être l'un de vos principaux outils de travail.

Tout cours de mathématiques contient une partie de rappels (j'entends par là tout ce qu'on vous a enseigné l'année précédente et que vous vous êtes empressés d'oublier); ceux-ci sont disséminés dans (à peu près) tous les chapitres, mais essentiellement les chapitres 1, 2, 3, 5 et 6. Ils sont constitués de fascicules de résultats, résumant sans démonstrations ce que vous êtes sensés connaître, et d'exercices. Nous n'aurons pas le temps d'y consacrer beaucoup de temps; mais comme il s'agit de sujets sur les quels vous avez déjà beaucoup travaillé vous devez être capables de faire ces révisions seuls. N'oubliez pas d'y consacrer le temps nécessaire ! Vous venez d'horizons très divers, il est donc possible que les paragraphes à classer dans les rappels ne soient pas exactement les mêmes pour tous.

Le cours est composé de 20 chapitres (en deux fascicules), et de cinq chapitres un peu spéciaux intitulés "outils fondamentaux du raisonnement". Il est complété par de très nombreux exercices.

La présentation du cours: Vous constaterez que les mots "définition", "théorème",... ont disparu. Ils ont été remplacés par l'utilisation de trois graphismes. Les passages les plus importants sont écrits en italique (*abcdef...*) il est souhaitable qu'ils soient mémorisés. Les parties écrites en caractères droits (*abcdef...*) sont pour l'essentiel des démonstrations; elles doivent être comprises, mais il serait illusoire de les apprendre par cœur. Les commentaires sont en petits caractères (*abcdef...*).

Le cours est parsemé d'exercices (marqués de lettres). Certains d'entre eux vous demandent d'écrire vous mêmes des démonstrations qui n'ont pas été écrites (par exemple si deux démonstrations sont à peu près identiques, la première est écrite, et on vous demande d'écrire la seconde).

Très souvent il s'agit de vous faire observer des exemples ou des contre-exemples essentiels à la compréhension du texte. Dans tous les cas il est souhaitable qu'en apprenant votre cours vous traitiez ces exercices.

Il va de soi que ce texte a été conçu pour être lu par un étudiant de DEUG, et qu'il remplace le cours oral que vous avez l'habitude de noter. Vous ne devez attendre des enseignants que l'explication des parties difficiles, et quelques commentaires. Cette année, vous devrez donc (entre autres) apprendre à lire un texte mathématique, ce que vous avez probablement peu l'habitude de faire.

Les outils fondamentaux du raisonnement: Il s'agit de questions qui étaient enseignées dès le collège dans les années 70-80, et qui ont disparu des programmes, ou plutôt qui n'y figurent plus que par quelques abréviations usuelles. Il n'y a donc là rien de bien compliqué, c'est seulement un peu plus abstrait que les mathématiques que vous avez pratiquées. Vous devez considérer que ce ne sont que quelques mises au point sur des sujets que vous avez déjà rencontrés; elles sont nécessaires pour toutes les études de mathématiques (ou d'informatique) au delà du DEUG. Ces chapitres comportent un minimum de théorie, vous y trouverez surtout des précisions de vocabulaire, des règles simples, et des exercices pour vous entraîner; ils sont numérotés 1b, 6b,... et se trouvent ainsi associés aux parties du cours dans les quelles ils sont d'usage constant.

Les exercices de fin de chapitre: Chaque chapitre est accompagné d'exercices (en général plus de 20). Ceux-ci sont marqués d'étoiles, qui constituent un indice de difficulté. Une étoile signifie qu'il s'agit d'une application directe du cours. Deux étoiles signifient qu'il est nécessaire de faire preuve d'un peu d'imagination. Les quelques exercices qui sont marqués de trois étoiles peuvent être considérés comme un peu difficiles. Cette classification fort imparfaite doit faciliter votre travail personnel, et éviter que celui-ci se réduise à recopier proprement les solutions des exercices traités en TD; elle devrait vous permettre de sélectionner les exercices faciles que vous êtes capables de traiter - comme le jour de l'examen - sans l'aide d'un enseignant.

Les compléments: Il s'agit de paragraphes qui ne font pas vraiment partie du cours; il n'entrent donc pas dans le contrat minimum exigé pour être reçu à l'examen.

Algèbre linéaire

Avertissement: Ce premier chapitre contient essentiellement des rappels du cours d'algèbre linéaire de première année. C'est pourquoi la plupart des résultats sont donnés sans démonstration.

§ 1 : Espaces vectoriels - Sous-espaces vectoriels

On dit qu'un ensemble est un espace vectoriel (ou un $\mathbb{R}$ -espace vectoriel) si :

D'une part on sait ajouter deux éléments de E .

D'autre part on sait multiplier les éléments de E par les scalaires (i.e.: par les éléments de $\mathbb{R}$).

De telle façon que:

Propriétés de la somme : $(E, +)$ est une groupe commutatif, autrement dit:

$$(\forall a)(\forall b) : a+b = b+a$$

$$(\forall a)(\forall b)(\forall c) : (a+b)+c = a+(b+c)$$

$$(\exists 0) \text{ tel que } (\forall a) : a+0 = 0+a = a$$

$$(\forall a)(\exists b) \text{ tel que } a+b = b+a = 0. \text{ Cet élément } b \text{ est noté } -a.$$

Propriétés du produit :

$$(\forall a)(\forall b)(\forall \lambda) : \lambda(a+b) = \lambda a + \lambda b$$

$$(\forall a)(\forall \lambda)(\forall \mu) : (\lambda+\mu)a = \lambda a + \mu a$$

$$(\forall a)(\forall \lambda)(\forall \mu) : (\lambda\mu)a = \lambda(\mu a)$$

$$(\forall a) : 1a = a \text{ et } (-1)a = -a$$

En fait ces 8 propriétés expriment que, dans E , 'on peut faire les calculs, et les abus de langage, que l'on a pris l'habitude de faire dans le calcul des vecteurs de la géométrie'.

Sous-espaces vectoriels: Un sous-ensemble non vide F de E est un sous-espace vectoriel, si il est stable par combinaisons linéaires, autrement dit si:

D'une part la somme de deux vecteurs de F est un vecteur de F .

D'autre part le produit d'un vecteur de F par un scalaire est un vecteur de F (donc $0 \in F$)

On se persuade facilement que F est alors un espace vectoriel.

Exercice a: 1) Démontrer que l'intersection de deux sous-espaces vectoriels de E est un sous-espace vectoriel de E .

2) Démontrer que l'intersection d'une famille quelconque de sous-espaces vectoriels de E , est un sous-espace vectoriel de E .

Exercice b: Dans l'espace vectoriel $\mathbb{R}^3$ (ou dans l'espace des vecteurs de la géométrie) exhiber deux sous-espaces vectoriels dont la réunion n'est pas un sous-espace vectoriel.

§ 2 : Familles libres - Familles génératrices - Bases

Familles libres: Une famille $\{e_1 \dots e_n\}$ de vecteurs de E , est dite libre (on dit aussi que ces vecteurs sont indépendants) si la seule famille de scalaires $\lambda_1 \dots \lambda_n$ telle que $\lambda_1 e_1 + \dots + \lambda_n e_n = 0$ est la famille nulle ($\forall i : \lambda_i = 0$).

Autrement dit si je sais que les e_i sont indépendants, et que $\lambda_1 e_1 + \dots + \lambda_n e_n = 0$, je peux conclure que tous les λ_i sont nuls.

Exercice c: Montrer que toute partie d'une famille libre est libre.

Sous-espace engendré: Soit X une partie quelconque de E , on appelle sous-espace vectoriel engendré par X , le plus petit sous-espace vectoriel qui contient X . Nous le noterons $[X]$. C'est l'ensemble des combinaisons linéaires de la forme $\lambda_1 x_1 + \dots + \lambda_p x_p$, où les x_i sont dans X .

En particulier, si $X = \{x_1, \dots, x_n\}$ est fini, $[X]$ est l'ensemble des combinaisons linéaires de la forme $\lambda_1 x_1 + \dots + \lambda_n x_n$.

Si le sous-espace $[X]$ est E tout entier, on dit que X est une partie génératrice de E .

Exercice d: Montrer que toute partie de E qui contient une partie génératrice, est elle-même génératrice.

Exercice e: On donne deux parties X et Y de E . Montrer que :

- a) $X \subset Y$ implique $[X] \subset [Y]$
- b) $[X] \cap [Y]$ n'est pas toujours égal à $[X \cap Y]$
- c) $[X \cup Y] = [X] \cup [Y]$

Les bases: Soit $\{e_1 \dots e_n\}$ une partie finie de E ; les conditions suivantes sont équivalentes :

- a) $\{e_1 \dots e_n\}$ est libre et génératrice.
 - b) Pour tout vecteur e de E , il existe une unique famille de scalaires $\{\lambda_1 \dots \lambda_n\}$ telle que $e = \lambda_1 e_1 + \dots + \lambda_n e_n$.
 - c) $\{e_1 \dots e_n\}$ est libre, et toute famille de vecteurs qui contient strictement $\{e_1 \dots e_n\}$ n'est pas libre (autrement dit $\{e_1 \dots e_n\}$ est maximale parmi les familles libres).
 - d) $\{e_1 \dots e_n\}$ est génératrice, et toute famille de vecteurs qui est strictement plus petite que $\{e_1 \dots e_n\}$ n'est pas génératrice (autrement dit $\{e_1 \dots e_n\}$ est minimale parmi les familles génératrices).
- Si ces conditions sont vérifiées, on dit que $\{e_1 \dots e_n\}$ est une base de l'espace vectoriel E , et les scalaires $\{\lambda_1 \dots \lambda_n\}$ sont appelés les coordonnées de e dans cette base.

Si E est de dimension finie (c'est à dire s'il existe une famille finie qui est génératrice), toute partie libre de E est contenue dans une base. Toute partie génératrice de E contient une base.

Théorème de la dimension : Si un espace vectoriel E possède une base ayant n éléments, toute autre base de E a aussi n éléments. On dit que c'est un espace vectoriel de dimension n .

Dans un espace de dimension n :

- Toute partie génératrice a au moins n éléments; si elle a exactement n éléments, c'est une base.
- Toute partie libre a au plus n éléments; si elle a exactement n éléments, c'est une base.

Théorème de la base incomplète: Si $\{e_1 \dots e_p, e_{p+1} \dots e_{p+q}\}$ est une partie génératrice de E , telle que $\{e_1 \dots e_p\}$ soit une partie libre, alors il existe une base de E qui contient $\{e_1 \dots e_p\}$ et est contenue dans $\{e_1 \dots e_p, e_{p+1} \dots e_{p+q}\}$.

Les espaces de dimension infinie: Il existe des espaces ayant, quel que soit p , des systèmes libres à p éléments. Par exemple l'espace des polynômes: quelque soit p , $\{1, x, \dots, x^{p-1}\}$ est libre. De tels espaces ne peuvent posséder une base (au sens de la définition ci-dessus); on dit qu'ils sont de dimension infinie. On en verra de nombreux exemples au chapitre 3 ci-dessous. Dans de tels espaces, toute partie génératrice a une infinité d'éléments.

§ 3 : Applications linéaires - Equations linéaires

Soit E et F deux espaces vectoriels; une application $f: E \rightarrow F$ est dite linéaire si :

- * D'une part elle commute à la somme ; c'est à dire

$$(\forall e \in E) (\forall e' \in E) \quad f(e+e') = f(e) + f(e')$$

- * D'autre part elle commute au produit par les scalaires; c'est à dire

$$(\forall e \in E) (\forall \lambda \text{ scalaire}) \quad f(\lambda e) = \lambda f(e)$$

Noyau et image: A toute application linéaire $f: E \rightarrow F$ on associe:

D'une part le noyau de f (noté $\text{Ker } f$) qui est un sous-espace de E . C'est l'ensemble des solutions de l'équation $f(x) = 0$. Autrement dit $\text{Ker } f = \{x \in E \text{ tels que } f(x) = 0\}$.

D'autre part l'image de f (notée $\text{Im } f$) qui est un sous-espace de F . C'est l'ensemble des vecteurs a de F , tels que l'équation $f(x) = a$ ait (au moins) une solution. Ce que l'on peut encore écrire: $\text{Im } f = \{a \in F \text{ tels qu'il existe } x \in E \text{ vérifiant } f(x) = a\}$. Noter que $\dim \text{Im } f$ est appelé le rang de f .

On retiendra :

$$\dim \text{Ker } f + \dim \text{Im } f = \dim E \quad (E \text{ espace de départ})$$

$$"\text{Ker } f = \{0\}" \Leftrightarrow "f \text{ est injective}"$$

$$"\text{Im } f = F" \Leftrightarrow "f \text{ est surjective}"$$

$$\text{Si } f \text{ est injective: } \dim E = \dim \text{Im } f \leq \dim F$$

$$\text{Si } f \text{ est surjective: } \dim E \geq \dim \text{Im } f = \dim F$$

$$\text{Si } E \text{ et } F \text{ sont de même dimension finie, alors : } "f \text{ injective}" \Leftrightarrow "f \text{ surjective}" \Leftrightarrow "f \text{ bijective}"$$

Equations linéaires: Soit une application linéaire $f: E \rightarrow F$, et soit a dans F . Considérons l'équation $f(x) = a$.

Pour qu'elle ait au moins une solution, il faut et il suffit que a soit dans $\text{Im } f$.

- * Si x_1 est une solution et y un élément du noyau de f , alors $f(x_1+y) = a + 0 = a$; donc x_1+y est une autre solution.

- * Inversement si x_1 et x_2 sont deux solutions, alors $f(x_1) - f(x_2) = a - a = 0$; donc $x_1 - x_2$ est un élément du noyau de f .

Autrement dit: si nous connaissons une solution de l'équation $f(x) = a$, en lui ajoutant les éléments du noyau, nous obtenons d'autres solutions de l'équation; et nous les obtenons toutes ainsi. L'équation $f(x) = 0$, est souvent appelée "équation homogène" associée à "l'équation non homogène" $f(x) = a$.

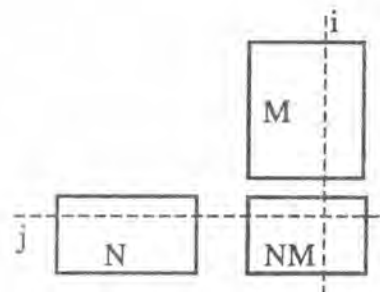
Lorsque l'on est en dimension finie, la dimension de $\text{Im } f$ est appelée le rang de l'équation linéaire. Si ce rang est égal à $\dim F$, l'équation a (quel que soit a) au moins une solution. Si ce rang est égal à $\dim E$, f est injective, et (quel que soit a) l'équation a au plus une solution.

§ 4 : Le calcul matriciel

Lorsque l'on s'est donné d'une part une application linéaire $f: E \rightarrow F$, d'autre part des bases $\{e_i\}$ de E et $\{f_j\}$ de F , on construit un tableau de scalaires à $n=\dim E$ colonnes et $p=\dim F$ lignes. La i -ème colonne de ce tableau est formée de coordonnées de $f(e_i)$ dans la base $\{f_j\}$. Ce tableau est la "matrice-de- f -dans-les-bases- $\{e_i\}$ -et- $\{f_j\}$ ". Attention: cette matrice dépend de f , mais aussi des deux bases; ne jamais parler de "la matrice de f ", cela n'a pas de sens.

Composées d'applications linéaires - Produits de matrices: Soit $\phi: E \rightarrow F$ et $\psi: F \rightarrow G$ deux applications linéaires, alors la composée $\psi \circ \phi: E \rightarrow G$ est une application linéaire. Soit $\{e_i\}$, $\{f_j\}$ et $\{g_k\}$ des bases de E , F et G . Soit M la matrice de ϕ dans les bases $\{e_i\}$ et $\{f_j\}$. Soit N la matrice de ψ dans les bases $\{f_j\}$ et $\{g_k\}$. Alors, par définition, le produit NM est la matrice de $\psi \circ \phi$ dans les bases $\{e_i\}$ et $\{g_k\}$.

On le calcule grâce à la règle suivante: Le terme qui est à l'intersection de la i -ème colonne et de la j -ème ligne dans NM est la somme des produits terme à terme de la j -ème ligne de N , et de la i -ème colonne de M .



Une notation commode : Pour écrire en abrégé que M est la matrice de ϕ dans les bases $\{e_i\}$ et $\{f_j\}$, nous écrirons $(E, e_i) \xrightarrow{M} (F, f_j)$. De même pour écrire que N est la matrice de $\psi: F \rightarrow G$ dans les bases $\{f_j\}$ et $\{g_k\}$ nous écrirons $(F, f_j) \xrightarrow{N} (G, g_k)$. Alors pour composer ϕ et ψ , nous accolons ces deux écritures:

$$(E, e_i) \xrightarrow{M} (F, f_j) \xrightarrow{N} (G, g_k)$$

Puis nous oublions l'espace intermédiaire, en remarquant que la composée s'écrit $\psi \circ \phi$ (et non $\phi \circ \psi$!!), et que les matrices se composent dans le même sens; il reste:

$$(E, e_i) \xrightarrow{NM} (G, g_k)$$

Application aux matrices de changement de base: Soit $\{e_i\}$ et $\{e_j\}$ deux bases de E , et soit A la matrice formée des coordonnées des $\{e_i\}$ dans la base $\{e_j\}$, elle correspond à $(E, e_i) \xrightarrow{A} (E, e_j)$. Soit B la matrice des coordonnées des $\{e_j\}$ dans la base $\{e_i\}$, elle correspond à $(E, e_j) \xrightarrow{B} (E, e_i)$.

Nous aurons :

$$(E, e_i) \xrightarrow{A} (E, e_j) \xrightarrow{B} (E, e_i)$$

Donc:

$$(E, e_i) \xrightarrow{BA} (E, e_i)$$

Donc $AB = I_n$, soit $B = A^{-1}$.

Application aux changements de bases: Soit M la matrice d'une application $\phi: E \rightarrow E$ dans la base $\{e_i\}$ (au départ et à l'arrivée), nous écrirons:

$$\begin{array}{ccc} (E, e_i) & \xrightarrow{M} & (E, e_i) \\ \text{Id} \uparrow B=A^{-1} & & \text{Id} \downarrow A \end{array}$$

Donc:

$$(E, e_j) \xrightarrow{AMA^{-1}} (E, e_j)$$

Espaces vectoriels de matrices et d'applications linéaires: Nous savons ajouter deux applications linéaires de E dans F ($(f+g)(x) = f(x) + g(x)$). Nous savons multiplier une application

linéaire par un scalaire $((\lambda f)(x) = \lambda f(x))$. Les nouvelles applications ainsi obtenues sont linéaires. Ceci munit l'ensemble $\mathfrak{L}(E, F)$ des applications linéaires de E dans F , d'une somme et d'un produit par les scalaires, et en fait un espace vectoriel.

De la même façon nous savons ajouter deux matrices ayant même nombre de lignes et même nombre de colonnes $((m_i^j) + (n_i^j) = (m_i^j + n_i^j))$, et nous savons multiplier une matrice par un scalaire $(\lambda(m_i^j) = (\lambda m_i^j))$. Ceci munit l'ensemble $\mathfrak{M}(n, p)$ des matrices à n lignes et p colonnes, d'une somme et d'un produit par les scalaires, et en fait un espace vectoriel.

Et lorsque nous donnons des bases $\{e_i\}_{i=1, \dots, n}$ et $\{f_j\}_{j=1, \dots, p}$ de E et F , la correspondance qui à ϕ dans $\mathfrak{L}(E, F)$, associe sa matrice dans ces bases est une application linéaire bijective:

$$\mathfrak{L}(E, F) \rightarrow \mathfrak{M}(n, p)$$

Notons que $\mathfrak{M}(n, p)$ est un espace de dimension np (avec pour base les matrices dont un coefficient est égal à 1, et dont les autres sont nuls).

Donc $\mathfrak{L}(E, F)$ (qui est isomorphe à $\mathfrak{M}(n, p)$) est aussi de dimension $np = \dim E \times \dim F$.

Deux erreurs à ne pas commettre: Si M et N sont deux matrices $n \times n$, on a en général $MN \neq NM$; de même si A et B sont deux endomorphismes de E , on a en général $A \circ B \neq B \circ A$.

Si l'on a deux matrices M et N telles que $MN = 0$, on ne peut pas en conclure que l'une des deux est nulle. De même si $A: E \rightarrow F$ et $B: F \rightarrow G$ sont deux applications linéaires telles que $B \circ A = 0$, on ne peut pas en conclure que l'une des deux est nulle (mais seulement que $\text{Ker } B \supset \text{Im } A$).

§ 5 : Le dual

Par définition le dual d'un espace vectoriel E est l'ensemble des formes linéaires sur E , c'est à dire des applications linéaires de E dans $\mathbb{R}$. On le note $\mathfrak{L}(E, \mathbb{R})$, ou encore E^* .

C'est, d'après le §4 ci-dessus, un espace vectoriel. Et, si E est de dimension finie, il a même dimension que E .

Bases du dual: Supposons que E soit de dimension n , et muni d'une base $\{e_1 \dots e_n\}$. Dans E^* nous disposons de n éléments privilégiés $\{e'_1 \dots e'_n\}$, définis comme suit : e'_i est l'application linéaire dont la matrice, dans les bases $\{e_i\}$ de E , et $\{1\}$ de $\mathbb{R}$, s'écrit:

$$M_i = (0, \dots, 1, \dots, 0) \quad (\text{le } 1 \text{ étant à la place } i)$$

Autrement dit e'_i est la seule application linéaire de E dans $\mathbb{R}$ telle que $e'_i(e_i) = 1$ et $(\forall j \neq i) e'_i(e_j) = 0$.

Ces éléments e'_i forment une base de E^* (c'est un cas particulier de ce que nous avons vu au §4), appelée la base duale de $\{e_i\}$.

Si ϕ s'écrit $\phi = \lambda_1 e'_1 + \dots + \lambda_n e'_n$, alors $(\forall j) \phi(e_j) = \lambda_1 e'_1(e_j) + \dots + \lambda_n e'_n(e_j) = \lambda_j$. Autrement dit $\phi = \lambda_1 e'_1 + \dots + \lambda_n e'_n$, signifie que la matrice de ϕ (dans les bases $\{e_i\}$ de E , et $\{1\}$ de $\mathbb{R}$) s'écrit: $\mu = (\lambda_1, \dots, \lambda_n)$.

Transposée d'une application linéaire: Soit $g: E \rightarrow F$ linéaire, à toute forme linéaire $\phi: F \rightarrow \mathbb{R}$, nous pouvons associer $\phi \circ g: E \rightarrow \mathbb{R}$. Ceci nous donne une application (elle est linéaire, c'est facile à vérifier) de F^* dans E^* . Nous la noterons g^* , et nous l'appellerons la transposée de g .

Donnons nous des bases $\{e_i\}$ et $\{f_j\}$ de E et de F , et notons $M = (m_{ij})$ la matrice de g dans ces bases. Nous allons chercher la matrice de g^* dans les bases duales $\{e'_i\}$ et $\{f'_j\}$. Sa i -ème colonne est

formée des coordonnées de $g^*(f'_j)$ dans la base $\{e'_i\}$. Pour calculer les coordonnées de $g^*(f'_j)$ dans la base $\{e'_i\}$, nous calculons ses valeurs sur les éléments de la base $\{e_i\}$. Nous avons

$$(g^*(f'_j))(e_i) = (f'_j \circ g)(e_i) = f'_j(g(e_i))$$

$$(g^*(f'_j))(e_i) = f'_j(m_1^1 f_1 + \dots + m_i^n f_n) = m_1^1 f'_j(f_1) + \dots + m_i^n f'_j(f_n) = m_i^j f'_j(f_j) = m_i^j$$

Ainsi la matrice N de g^* dans les bases $\{f'_j\}$ et $\{e'_i\}$ a les mêmes coefficients que M ; mais attention ils ne sont pas rangés de la même façon. Si l'une a p lignes et n colonnes, l'autre a p colonnes et n lignes. La r -ième ligne de M , est la r -ième colonne de N . La s -ième colonne de M est la s -ième ligne de N . On dit que N est la transposée de M , que l'on notera M^t .

Exemple: Si $M = \begin{pmatrix} 1 & 2 \\ 3 & 4 \\ 5 & 6 \end{pmatrix}$ alors $M^t = \begin{pmatrix} 1 & 3 & 5 \\ 2 & 4 & 6 \end{pmatrix}$

§ 6 : Sous-espaces supplémentaires et sommes directes

Positions relatives de deux sous-espaces: Lorsque F et G sont deux sous-espaces de E , on note $F+G$ le sous-espace de E engendré par $F \cup G$. Si E est de dimension n , alors F , G , $F+G$ et $F \cap G$ sont de dimension finie, et l'on a :

$$\dim F + \dim G = \dim (F+G) + \dim (F \cap G)$$

Cas de deux sous-espaces supplémentaires: Etant donnés deux sous-espaces F et G d'un espace E , les conditions suivantes sont équivalentes:

- a) $F \cap G = \{O\}$ et $F+G = E$
- b) $\dim F + \dim G = \dim E$ et $F \cap G = \{O\}$
- c) $\dim F + \dim G = \dim E$ et $F+G = E$
- d) tout vecteur de E s'écrit, de façon unique, comme la somme d'un vecteur de F et d'un vecteur de G .
- e) si $\{f_1, \dots, f_p\}$ est une base de F , et $\{g_1, \dots, g_q\}$ une base de G , alors $\{f_1, \dots, f_p, g_1, \dots, g_q\}$ est une base de E .

Si ces conditions sont vérifiées, on dit que F et G sont supplémentaires (dans E); la somme de F et G est alors notée $F \oplus G$, au lieu de $F+G$.

Cas de 3 (ou 4,...) sous-espaces supplémentaires: Etant donnés 3 sous-espaces F, G, H de E , les conditions suivantes sont équivalentes :

- c') $\dim F + \dim G + \dim H = \dim E$ et $F+G+H = E$.
- d') tout vecteur de E s'écrit, de façon unique, comme la somme d'un vecteur de F , d'un vecteur de G et d'un vecteur de H .
- e') si $\{f_1, \dots, f_p\}$ est une base de F , $\{g_1, \dots, g_q\}$ une base de G et $\{h_1, \dots, h_r\}$ une base de H , alors $\{f_1, \dots, f_p, g_1, \dots, g_q, h_1, \dots, h_r\}$ est une base de E .

Lorsqu'elles sont vérifiées, on dit que F, G et H sont supplémentaires (dans E).

On notera que dès qu'on a plus de deux sous-espaces les conditions a et b n'ont pas d'équivalent. Par exemple dans un espace de dimension 3, il ne suffit pas que 3 sous-espaces de dimension 1 aient 2 à 2 une intersection réduite à $\{O\}$, pour qu'ils soient supplémentaires (cf : exercice 14). Voici, en exercice une autre caractérisation:

Exercice f: Soit F, G et H trois sous-espaces d'un espace E .

a) On suppose que F, G et H sont supplémentaires, montrer que $F \cap (G+H) = \{0\}$ et $H \cap G = \{0\}$.

b) On suppose (E étant de dimension finie) que $\dim F + \dim G + \dim H = \dim E$ que $F \cap (G+H) = \{0\}$ et que $H \cap G = \{0\}$, montrer que F, G et H sont supplémentaires.

Produits: Etant donnés deux espaces vectoriels A et B , on peut construire un espace vectoriel, noté $A \times B$ (appelé ^(*)produit de A et B), qui contient A et B comme sous-espaces, et dans lequel ces deux sous-espaces sont supplémentaires.

Construction: Un élément de $A \times B$ est un couple (a, b) , où a est un élément de A , et b un élément de B . La somme de (a, b) et de (a', b') est $(a+a', b+b')$. Le produit de (a, b) par le scalaire λ est $(\lambda a, \lambda b)$.

Laissons au lecteur le soin de vérifier que les 8 propriétés de la définition du § 1, sont vérifiées, c'est à dire que l'on a bien défini ainsi un nouvel espace vectoriel. Dans cet espace $A \times B$, les éléments de la forme $(a, 0)$ forment un sous-espace, qui s'identifie naturellement à A . De même B s'identifie au sous-espace formé des couples $(0, b)$. Et ces deux sous-espaces sont supplémentaires puisque tout couple (a, b) s'écrit de façon unique

$$(a, b) = (a, 0) + (0, b)$$

§ 7 : Algèbre linéaire sur un corps K

Dans tout ce qui précède les scalaires qui intervenaient, étaient des nombres réels. Une étude détaillée des démonstrations - que nous n'avons pas faites - montrerait que les propriétés des nombres réels qui nous ont servi, sont les suivantes:

Premier type de propriétés : Nous pouvons ajouter deux nombres réels, nous pouvons les multiplier; et ces opérations vérifient les propriétés suivantes:^(*)

α) $(\forall a) (\forall b) : a + b = b + a$ (on dit que l'addition est commutative).

β) $(\forall a) (\forall b) (\forall c) : (a + b) + c = a + (b + c)$ (on dit que l'addition est associative).

γ) $(\forall a) : a + 0 = 0 + a = a$ (on dit que 0 est élément neutre de l'addition).

δ) $(\forall a)$, il existe un élément, noté $-a$, tel que : $a + (-a) = (-a) + a = 0$.

ϵ) $(\forall a) (\forall b) : ab = ba$ (on dit que la multiplication est commutative).

ϕ) $(\forall a) (\forall b) (\forall c) : (ab)c = a(bc)$ (on dit que la multiplication est associative).

η) $(\forall a) : a \cdot 1 = 1 \cdot a = a$ (on dit que 1 est élément neutre de la multiplication).

ι) $(\forall a) (\forall b) (\forall c) : a(b+c) = ab + ac$ et $(b+c)a = ba + ca$ (on dit que la multiplication est distributive par rapport à l'addition).

Deuxième type de propriétés :

κ) $(\forall a \neq 0)$, il existe un élément, noté $\frac{1}{a}$ ou a^{-1} , tel que $aa^{-1} = a^{-1}a = 1$.

De nombreux objets mathématiques possèdent une addition et une multiplication vérifiant ces neuf propriétés; citons l'ensemble $\mathbb{C}$ des nombres complexes; et l'ensemble $\mathbb{Q}$ des nombres rationnels. D'autres ne possèdent que les propriétés du premier type; par exemple l'ensemble $\mathbb{Z}$ des nombres entiers, ou l'ensemble $\mathbb{R}[X]$ des polynômes à coefficients réels.

^(*) On l'appelle aussi la somme directe de A et B ; et on le note aussi $A \oplus B$.

^(*) Les propriétés $\alpha, \beta, \gamma, \delta$ signifient que l'on est en présence d'un groupe commutatif.

Vocabulaire: Les ensembles d'objets mathématiques munis d'une addition et d'une multiplication qui possèdent les propriétés du premier type (α à 1) sont appelés des anneaux commutatifs. S'ils possèdent aussi la propriété κ , on les appelle des corps.

Espaces vectoriels sur le corps K : Nous appellerons espace vectoriel sur le corps K (ou K -espace vectoriel), la donnée d'un ensemble E sur lequel sont définies deux opérations:

* une somme.

* un produit par les éléments de K (qui seront appelés des scalaires).

Ces opérations étant astreintes à vérifier les 8 propriétés écrites au §1.

Les espaces vectoriels sur les corps les plus divers, sont des outils souvent utilisés en algèbre. Dans toute la suite du cours nous n'utiliserons que les espaces vectoriels sur $\mathbb{R}$, et les espaces vectoriels sur $\mathbb{C}$. Nous admettrons - sans démonstration - que tous les résultats, toutes les méthodes de calcul, et toutes les constructions connues (c'est à dire rappelées dans ce chapitre et le suivant) pour les $\mathbb{R}$ -espaces vectoriels, sont valables (quel que soit le corps K) pour les K -espaces vectoriels.

Exercices sur le chapitre 1

Exercice 1: Soit $\mathcal{P}$ l'espace vectoriel des vecteurs du plan P . Trouver deux sous-espaces de $\mathcal{P}$, dont la réunion n'est pas un sous-espace vectoriel de $\mathcal{P}$.

Soit E un espace vectoriel et soit (F_n) une suite de sous-espaces de E , tels que, quel que soit n , on ait : $F_{n+1} \supset F_n$; montrer que la réunion de tous les F_n est un sous-espace vectoriel de E .

* **Exercice 2:** Montrer qu'une famille libre ne peut contenir le vecteur nul.

* **Exercice 3:** Soit une partie X d'un espace vectoriel E , et soit un vecteur e de E . Montrer que $[X] = [X \cup \{e\}]$ si et seulement si e est dans $[X]$.

** **Exercice 4:** Soit $\{e_1, \dots, e_p\}$ et $\{f_1, \dots, f_q\}$ deux familles libres dans un espace vectoriel E . Montrer que la famille $\{e_1, \dots, e_p, f_1, \dots, f_q\}$ est libre si et seulement si $\{e_1, \dots, e_p\} \cap \{f_1, \dots, f_q\} = \{0\}$.

* **Exercice 5:** Soit F et G deux sous-espaces de dimension 2 dans $\mathbb{R}^4$; montrer que
ou bien F et G sont supplémentaires
ou bien il existe 3 vecteurs a, b, c tels que $F = [a, b]$ et $G = [b, c]$.
ou bien $F = G$

* **Exercice 6:** On donne deux sous-espaces F (de dimension 1) et G (de dimension 2) de $\mathbb{R}^3$; montrer que :
ou bien $G \supset F$
ou bien F et G sont supplémentaires.

* **Exercice 7:** Soit E un espace de dimension finie, et soit F un sous-espace de E .
a) Montrer que $\dim F \leq \dim E$
b) Montrer que $\dim F = \dim E$ implique $F = E$.

Exercice 8: Soit $\{e_1, \dots, e_n\}$ une famille de vecteurs de E , et soit f une application linéaire de E dans F .

a) Si $\{f(e_1), \dots, f(e_n)\}$ est une famille libre de F , alors $\{e_1, \dots, e_n\}$ est une famille libre de E .

b) Réciproquement si f est injective et si $\{e_1, \dots, e_n\}$ est une famille libre de E , alors $\{f(e_1), \dots, f(e_n)\}$ est une famille libre de F .

** **Exercice 9:** Soit $f : E \rightarrow F$ et $g : F \rightarrow G$ deux applications linéaires,

a) Montrer que $g \circ f : E \rightarrow G$ est linéaire

b) Montrer que $\text{Ker}(g \circ f) \supset \text{Ker } f$

c) Montrer que $\text{Im } g \supset \text{Im}(g \circ f)$

d) Montrer que si $g \circ f$ est surjective, alors g est surjective

e) Montrer que si $g \circ f$ est injective, alors f est injective.

*** Exercice 10: Soit $f: E \rightarrow E$ linéaire (E de dimension finie).

- Montrer que $(\forall n) : \text{Im } f^n \supset \text{Im } f^{n+1}$
- Montrer que $\text{Im } f^n = \text{Im } f^{n+1}$ si et seulement si $\text{Im } f^n \cap \text{Ker } f = \{0\}$.
En déduire que $\text{Im } f^n = \text{Im } f^{n+1}$ entraîne $(\forall p > 0) : \text{Im } f^n = \text{Im } f^{n+p}$.
- Montrer que $\text{Ker } f^{n+1} \supset \text{Ker } f^n$.
- Montrer que $\text{Ker } f^{n+1} = \text{Ker } f^n$ implique $(\forall p > 0) : \text{Ker } f^{n+p} = \text{Ker } f^n$.

* Exercice 11: Soit E un espace vectoriel, et ϕ un endomorphisme de E .

- Montrer que si E est de dimension finie, les assertions " ϕ est injective", " ϕ est surjective" et " ϕ est bijective" sont équivalentes.
- Soit E l'espace des polynômes, supposons que ϕ associe à tout polynôme $P(X)$ le polynôme $XP(X)$; montrer que ϕ est linéaire injective et non surjective.
- Construire de même une application linéaire surjective mais non injective de E dans E .

** Exercice 12: Soit $u: E \rightarrow F$ linéaire et surjective. Soit $\{f_1, \dots, f_k\}$ une base de F .

- Soit $\phi_1, \dots, \phi_k$ des vecteurs tels que $(\forall i) u(\phi_i) = f_i$. Montrer que les vecteurs $\phi_1, \dots, \phi_k$ sont indépendants.
- Soit G le sous-espace engendré par les vecteurs $\phi_1, \dots, \phi_k$. Montrer que $E = G \oplus \text{Ker } u$.
- On choisit une base $\{\alpha_1, \dots, \alpha_m\}$ de $\text{Ker } u$. Ecrire la matrice de u dans les bases $\{\phi_1, \dots, \phi_k, \alpha_1, \dots, \alpha_m\}$ et $\{f_1, \dots, f_k\}$.
- Application: Soit $u: \mathbb{R}^4 \rightarrow \mathbb{R}^2$ dont la matrice, dans les bases naturelles, s'écrit :

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

On prend pour f_1, f_2 la base naturelle de $\mathbb{R}^2$, construire $\phi_1, \phi_2, \alpha_1, \alpha_2$.

* Exercice 13: Soit $f: E \rightarrow E$ linéaire, et telle que f^2 soit l'identité.

Montrer que $\text{Im}(f - \text{Id}) \subset \text{Ker}(f + \text{Id})$ et $\text{Im}(f + \text{Id}) \subset \text{Ker}(f - \text{Id})$.

Montrer que tout vecteur e de E s'écrit $e = (f - \text{Id})(u) + (f + \text{Id})(v)$.

En déduire que $\text{Im}(f - \text{Id})$ et $\text{Im}(f + \text{Id})$ sont supplémentaires.

* Exercice 14: Dans $\mathbb{R}^2$, on considère 3 sous-espaces 2 à 2 distincts de dimension 1 ; leurs intersections 2 à 2 sont réduites à $\{0\}$. Montrer qu'ils ne sont pas supplémentaires.

Dans $\mathbb{R}^3$, on considère 3 sous-espaces 2 à 2 distincts de dimension 1 ; sont-ils supplémentaires ?

** Exercice 15: On considère un endomorphisme f d'un espace E de dimension n , on suppose que $(f - \text{Id}) \circ (f - 3\text{Id}) = 0$

- Démontrer que $\text{Im}(f - 3\text{Id}) \subset \text{Ker}(f - \text{Id})$; en déduire que $\dim(\text{Ker}(f - \text{Id})) + \dim(\text{Ker}(f - 3\text{Id})) \geq \dim E$.
- Démontrer que $\text{Ker}(f - \text{Id}) \cap \text{Ker}(f - 3\text{Id}) = \{0\}$. Qu'en résulte-t-il pour les dimensions de ces deux noyaux ?
- Montrer que $\text{Ker}(f - \text{Id})$ et $\text{Ker}(f - 3\text{Id})$ sont supplémentaires.
- Application: On considère $f: \mathbb{R}^4 \rightarrow \mathbb{R}^4$ dont la matrice dans la base naturelle est:

$$\frac{1}{2} \begin{pmatrix} 5 & 1 & 1 & 1 \\ 1 & 3 & -1 & 1 \\ 1 & -1 & 5 & -1 \\ 1 & 1 & -1 & 3 \end{pmatrix}$$

Vérifier que $(f - \text{Id}) \circ (f - 3\text{Id}) = 0$. Trouver une base de chacun des deux noyaux, puis écrire la matrice de f dans la base de $\mathbb{R}^4$ obtenue en réunissant ces deux bases.

e) Soit maintenant un endomorphisme f d'un espace E de dimension infinie, tel que $(f - \text{Id}) \circ (f - 3\text{Id}) = 0$. Démontrer que $(\forall x) \quad x = \frac{-1}{2} (f - 3\text{Id})(x) + \frac{1}{2} (f - \text{Id})(x)$. En déduire que dans ce cas $\text{Ker}(f - \text{Id})$ et $\text{Ker}(f - 3\text{Id})$ sont encore supplémentaires.

Exercice 16: Soit $f: E \rightarrow E$ une application linéaire.

a) Soit α un nombre et soit E_α l'ensemble des $x \in E$ tels que $f(x) = \alpha x$. Montrer que E_α est un sous-espace vectoriel de E .

b) Soit $\alpha \neq \beta$, montrer que $E_\alpha \cap E_\beta = \{0\}$.

c) On suppose que $f^2 - 3f + 2\text{Id}_E = 0$. Montrer que $(\forall x \in E) f(x) - x$ est dans E_2 , et $f(x) - 2x$ dans E_1 . En déduire que E_1 et E_2 sont supplémentaires.

Exercice 17: Soit M une matrice $n \times n$, on note E l'ensemble des matrices $n \times n$ qui sont des combinaisons linéaires de $I_n, M, M^2, \dots, M^p, \dots$

a) Montrer que E est un sous-espace vectoriel de l'espace des matrices $n \times n$.

b) On suppose que $M^3 + M + I_n = 0$; montrer que M^3 et M^4 sont combinaisons linéaires de I_n, M et M^2 . Montrer, par récurrence sur p , que $(\forall p) M^p$ est combinaison linéaire de I_n, M et M^2 . En déduire que I_n, M et M^2 engendrent l'espace E .

Exercice 18: Effectuer les produits matriciels suivants

$$\begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & -1 \\ 3 & 7 & 7 \\ 9 & 8 & 7 \end{pmatrix} \begin{pmatrix} 3 & 2 & 5 & 7 \\ 8 & 9 & 0 & -7 \\ -6 & -3 & 4 & 5 \end{pmatrix}$$

$$\begin{pmatrix} -1 & 0 & 6 & 7 & 9 \\ 8 & -2 & 10 & 9 & 8 \\ 2 & 3 & -5 & -6 & 7 \\ 8 & 9 & 4 & -2 & -5 \end{pmatrix} \begin{pmatrix} 1 & -2 & 3 \\ 5 & 3 & -6 \\ 4 & 5 & -1 \\ 3 & 7 & -2 \\ -4 & 8 & 7 \end{pmatrix}$$

$$\begin{pmatrix} 6 & 3 & 4 \\ 7 & -1 & -5 \\ 7 & 7 & 9 \\ 1 & 4 & 7 \end{pmatrix} \begin{pmatrix} 7 & 9 & 2 & 5 \\ 7 & 1 & 0 & -7 \\ 3 & -3 & 1 & 5 \end{pmatrix}$$

$$\begin{pmatrix} -1 & 5 & 0 & 6 & 7 \\ 9 & 0 & 8 & -2 & 10 \\ 9 & 8 & 6 & -7 & 2 \\ 3 & -5 & -6 & 7 & 8 \\ -3 & 9 & 4 & -2 & -5 \end{pmatrix} \begin{pmatrix} 1 & 0 & -2 \\ 3 & 7 & 5 \\ 3 & -6 & 4 \\ 6 & 1 & -3 \\ 2 & -4 & 5 \end{pmatrix}$$

Exercice 19: Soit deux applications linéaires et surjectives φ et ψ de E dans F . On suppose que F est de dimension finie.

a) On suppose qu'il existe une application linéaire α de F dans lui-même telle que $\alpha \circ \varphi = \psi$. Démontrer que α est un isomorphisme de F sur lui-même, et démontrer que $\text{Ker } \varphi = \text{Ker } \psi$.

b) On suppose maintenant que $\text{Ker } \varphi = \text{Ker } \psi$. Démontrer que $(\forall y \in F) (\exists y' \in F)$ tel que $(\forall x \in E)$ si $\varphi(x) = y$ alors $\psi(x) = y'$. En déduire qu'il existe une application α de F dans lui-même telle que $\alpha \circ \varphi = \psi$. Puis montrer que α est linéaire.

Exercice 20 Soit deux applications linéaires $\varphi: E \rightarrow F$ et $\psi: F \rightarrow E$ telles que $\psi \circ \varphi = \text{Id}_E$. Montrer que φ est injective. Montrer que $\text{Im } \varphi$ et $\text{Ker } \psi$ sont supplémentaires dans F .

Exercice 21: Soit F et G deux sous-espaces supplémentaires dans E .

a) Soit H un sous-espace de E ; est-ce que $H \cap F$ et $H \cap G$ sont supplémentaires dans H ?

b) Soit $\varphi: K \rightarrow E$ linéaire; est-ce que $\varphi^{-1}(F)$ et $\varphi^{-1}(G)$ sont supplémentaires dans K ?

Exercice 22: Soit $f: E \rightarrow F$ une application linéaire. On note $f^*: F^* \rightarrow E^*$ sa transposée.

a) On suppose que E et F sont de dimension finie. En comparant les matrices de f et de f^* , montrer que f et f^* ont même rang. En déduire que f est injective si et seulement si f^* est surjective; et que f est surjective si et seulement si f^* est injective.

b) On ne suppose plus que E et F sont de dimension finie (donc on ne peut plus raisonner en termes de bases et de matrices). On suppose f surjective, montrer que $f^*(\varphi) = 0$ implique $\varphi = 0$. Qu'en résulte-t-il pour f^* ?

- * Exercice 23: Soit $u : E \rightarrow F$ et $v : F \rightarrow G$ des applications linéaires, montrer que $(v \circ u)^* = u^* \circ v^* : G^* \rightarrow E^*$
 Soit M et N deux matrices telles que le produit MN existe. Montrer que $N^t M^t$ existe et est égal à $(MN)^t$?

- * Exercice 24: Soit $f : E \rightarrow F$ et $g : E \rightarrow G$ deux applications linéaires, on définit une application $f \oplus g$ de E dans $F \oplus G$ en posant $(f \oplus g)(x) = (f(x), g(x))$.
 a) Montrer que $f \oplus g$ est linéaire.
 b) Montrer que $\text{Ker}(f \oplus g) = \text{Ker } f \cap \text{Ker } g$.

Exercices sur le § 7

- * Exercice 25: Les ensembles suivants sont munis d'une addition et d'une multiplication ; les quels sont des anneaux commutatifs ? Les quels sont des corps ?
 $\mathbb{N}$ ensemble des entiers positifs
 $\mathbb{Z}(\sqrt{2})$ ensemble des nombres de la forme $a + b\sqrt{2}$, où a et b sont entiers relatifs .
 $\mathbb{Q}(\sqrt{2})$ ensemble des nombres de la forme $a + b\sqrt{2}$, où a et b sont rationnels .
 $\mathbb{Q}[X]$ ensemble des polynômes à coefficients rationnels .
 $\mathbb{Z}[X]$ ensemble des polynômes à coefficients entiers relatifs .
- * Exercice 26: Soit $\mathbb{Q}(\sqrt{3})$ l'ensemble des nombres de la forme $a + b\sqrt{3}$, où a et b sont rationnels ; montrer que c'est un espace vectoriel de dimension 2 sur $\mathbb{Q}$.
- * Exercice 27: L'ensemble des matrices 2×2 (à coefficients réels) est muni d'une somme et d'un produit. Est-ce un anneau commutatif ? Est-ce un corps ?
- ** Exercice 28: Soit E un espace vectoriel sur $\mathbb{R}$, et soit $f : E \rightarrow E$ linéaire et telle que $f^2 = -\text{Id}$.
 a) Montrer que f est injective et surjective .
 b) Pour tout nombre complexe $z = a + bi$, et pour tout e dans E , on pose $ze = ae + bf(e)$; montrer que l'on obtient ainsi un $\mathbb{C}$ -espace vectoriel .
 c) Montrer que si une famille de vecteurs est indépendante dans E considéré comme un $\mathbb{C}$ -espace vectoriel, elle est indépendante dans E considéré comme un $\mathbb{R}$ -espace vectoriel. L'inverse est-il vrai ?

Outils fondamentaux du raisonnement

Première partie: Les opérations ensemblistes

§ 1 : Rappel des notations classiques de la théorie ensembliste.

Appartenance et inclusion: On se gardera de confondre le signe d'appartenance " $\in$ " et le signe d'inclusion " $\subset$ ". L'écriture $x \in E$ se lit "x appartient à E" (ou: x est un élément de E); tandis que l'écriture $A \subset E$ se lit A est un sous-ensemble de E.

On remarquera que $A=B$ signifie aussi que l'on a à la fois $A \subset B$ et $B \subset A$; cette remarque est souvent un puissant outil de démonstration: pour démontrer que $A=B$, on démontre successivement que $x \in A \Rightarrow x \in B$, puis que $x \in B \Rightarrow x \in A$.

Réunion: $A \cup B$ est la réunion de A et B. Dire que x appartient à $A \cup B$, c'est dire que x appartient à l'un au moins des ensembles A et B. Si $(A_i)_{i \in I}$ est une famille d'ensembles (l'ensemble d'indices I est fini ou infini), $\bigcup_{i \in I} A_i$ est la réunion de tous les ensembles A_i ; dire que x appartient à $\bigcup_{i \in I} A_i$, c'est dire que x appartient à l'un au moins des A_i .

En termes de formules:

$$A \cup B = \{x \text{ tels que ou bien } x \in A, \text{ ou bien } x \in B \text{ (ou bien les deux)}\}$$

$$\bigcup_{i \in I} A_i = \{x \text{ tels que } (\exists i) \text{ tel que } x \in A_i\}$$

Intersection: $A \cap B$ est l'intersection de A et B. Dire que x appartient à $A \cap B$ c'est dire que x appartient à la fois à A et à B. Si $(A_i)_{i \in I}$ est une famille d'ensembles (l'ensemble d'indices I est fini ou infini), $\bigcap_{i \in I} A_i$ est l'intersection de tous les ensembles A_i ; dire que x appartient à $\bigcap_{i \in I} A_i$, c'est dire que x appartient à tous les A_i .

En termes de formules:

$$A \cap B = \{x \text{ tels que } x \in A \text{ et } x \in B\}$$

$$\bigcap_{i \in I} A_i = \{x \text{ tels que } (\forall i \in I) x \in A_i\}$$

Le complémentaire: Si A est un sous-ensemble de E, le complémentaire de A dans E (noté $C_E A$) est l'ensemble des points de E qui n'appartiennent pas à A. Autrement dit on a $A \cup C_E A = E$ et aussi $A \cap C_E A = \emptyset$.

§ 2 : Injections - Surjections - Bijections.

Dire que $f: A \rightarrow B$ est injective c'est dire que $(\forall b \in B)$ l'équation $f(x) = b$ n'a au plus qu'une solution. Dire que $f: A \rightarrow B$ est surjective c'est dire que $(\forall b \in B)$ l'équation $f(x) = b$ a au moins une solution. Dire que $f: A \rightarrow B$ est bijective c'est dire que $(\forall b \in B)$ l'équation $f(x) = b$ a exactement une

solution. Dans ce cas il existe une application de B dans A appelée l'inverse (ou la réciproque) de f , telle que $f \circ g = \text{Id}_B$ et $g \circ f = \text{Id}_A$.

En termes de formules:

$$f \text{ injective} \iff [(\forall x \in E)(\forall x' \in E) f(x) = f(x') \text{ implique } x = x']$$

$$f \text{ surjective} \iff [(\forall y \in F)(\exists x \in E) \text{ tel que } f(x) = y]$$

Image: Si $f: E \rightarrow F$, et si A est un sous-ensemble de E , l'image de A par f (qui est notée $f(A)$) est un sous-ensemble de F ; c'est l'ensemble des points y de F tels que l'équation $f(x) = y$ ait au moins une solution dans A . Notons que $f(E)$ est aussi appelé l'image de l'application f .

Si B est un sous-ensemble de F , l'image réciproque de B par f (notée $f^{-1}(B)$) est un sous-ensemble de E ; c'est l'ensemble des points x de E tels que $f(x)$ soit dans B .

En termes de formules:

$$f(A) = \{y \in F \text{ tels que } (\exists x \in A) \text{ tel que } f(x) = y\}$$

$$f^{-1}(B) = \{x \in E \text{ tels que } f(x) \in B\}$$

Exercices sur le chapitre 1b

**** Exercice 1:** Soit $f: A \rightarrow B$ et $g: B \rightarrow C$ tels que $g \circ f: A \rightarrow C$ soit injective, montrer que f est injective. Montrer qu'il peut arriver que $g \circ f$ soit injective et que g soit non injective.

Soit $f: A \rightarrow B$ injective, montrer qu'il existe $h: B \rightarrow A$ telle que $h \circ f = \text{Id}_A$.

**** Exercice 2:** Soit $f: A \rightarrow B$ et $g: B \rightarrow C$ tels que $g \circ f: A \rightarrow C$ soit surjective, montrer que g est surjective. Montrer qu'il peut arriver que $g \circ f$ soit surjective et que f soit non surjective.

Soit $f: A \rightarrow B$ surjective, montrer qu'il existe $h: B \rightarrow A$ telle que $f \circ h = \text{Id}_B$.

*** Exercice 3:** Montrer qu'il existe des applications de $\mathbb{N}$ dans $\mathbb{N}$ qui sont injectives et non bijectives. Quelle condition faut-il imposer à un ensemble E , pour pouvoir affirmer que toute application injective de E dans E est bijective ?

*** Exercice 4:** Montrer qu'il existe des applications de $\mathbb{N}$ dans $\mathbb{N}$ qui sont surjectives et non bijectives. Quelle condition faut-il imposer à un ensemble E , pour pouvoir affirmer que toute application surjective de E dans E est bijective ?

*** Exercice 5:** On considère l'application φ de $\mathbb{Z}$ dans $\mathbb{N}$ définie par $\varphi(a) = 2a$ si $a \geq 0$ et $\varphi(a) = -1 - 2a$ si $a < 0$. Montrer quelle est bijective.

*** Exercice 6:** Soit A_1, A_2 , et B trois ensembles, a-t-on $(A_1 \cap A_2) \cap B = (A_1 \cap B) \cap (A_2 \cap B)$? Soit $(A_i)_{i \in I}$ une famille d'ensembles (I est fini ou infini); et soit B un ensemble, a-t-on $(\bigcap_{i \in I} A_i) \cap B = \bigcap_{i \in I} (A_i \cap B)$?

*** Exercice 7:** Soit A_1, A_2 , et B trois ensembles, a-t-on $(A_1 \cup A_2) \cap B = (A_1 \cap B) \cup (A_2 \cap B)$? Soit $(A_i)_{i \in I}$ une famille d'ensembles (I est fini ou infini); et soit B un ensemble, a-t-on $(\bigcup_{i \in I} A_i) \cap B = \bigcup_{i \in I} (A_i \cap B)$?

*** Exercice 8:** Soit A_1, A_2 , et B trois ensembles, a-t-on $(A_1 \cap A_2) \cup B = (A_1 \cup B) \cap (A_2 \cup B)$? Soit $(A_i)_{i \in I}$ une famille d'ensembles (I est fini ou infini). Soit B un ensemble, a-t-on $(\bigcap_{i \in I} A_i) \cup B = \bigcap_{i \in I} (A_i \cup B)$?

- ** Exercice 9:** Soit une application $f: E \rightarrow F$, et soit deux sous-ensembles A et B de E .
- Montrer que $f(A) \cup f(B) = f(A \cup B)$; mais qu'il peut arriver que l'on n'ait pas $f(A) \cap f(B) = f(A \cap B)$.
 - Montrer que si f est injective $f(A) \cap f(B) = f(A \cap B)$.
- * Exercice 10:** Soit une application $f: E \rightarrow F$, et soit deux sous-ensembles A et B de F .
- Montrer que $f^{-1}(A) \cup f^{-1}(B) = f^{-1}(A \cup B)$, et que $f^{-1}(A) \cap f^{-1}(B) = f^{-1}(A \cap B)$.
 - Est-il possible que $A \cap B$ soit non vide, et que $f^{-1}(A) \cap f^{-1}(B)$ soit vide ?
- ** Exercice 11:** Soit $f: E \rightarrow F$.
- Soit $B \subset F$, montrer que $\bigcup_E f^{-1}(B) = f^{-1}(\bigcup_F B)$.
 - Soit $A \subset E$, est-ce que $f(\bigcup_E A) = \bigcup_F f(A)$? Que devient cette réponse si f est injective ? surjective ? bijective ?
- * Exercice 12:** Soit A et B deux sous-ensembles de E , montrer que $\bigcup_E A \cup \bigcup_E B = \bigcup_E (A \cup B)$, et que $\bigcup_E A \cap \bigcup_E B = \bigcup_E (A \cap B)$.
Soit $(A_i)_{i \in I}$ une famille de sous-ensembles de E ; montrer que $\bigcup_E (\bigcup_{i \in I} A_i) = \bigcup_{i \in I} (\bigcup_E A_i)$.
- ** Exercice 13:** a) Soit $f: A \rightarrow B$. Quelle condition doit-on imposer à f , pour que $((\forall B' \subset B) \text{ l'on ait } f(f^{-1}(B')) = B' ?$
b) Soit $f: A \rightarrow B$. Quelle condition doit-on imposer à f , pour que $((\forall A' \subset A) \text{ l'on ait } f^{-1}(f(A')) = A' ?$
- ** Exercice 14:** Soit E un ensemble et $\mathcal{P}(E)$ l'ensemble des parties de E . Soit φ une application de E dans $\mathcal{P}(E)$; posons $F = \{x \in E \text{ tels que } x \notin \varphi(x)\}$; montrer qu'il ne peut exister aucun point e de E tel que $\varphi(e) = F$. En déduire qu'il n'existe aucune application surjective de E dans $\mathcal{P}(E)$.

Méthodes de calcul dans les espaces de dimension finie

Avertissement: Ce chapitre contient des rappels du cours d'algèbre linéaire de première année. C'est pourquoi la plupart des résultats sont énoncés sans démonstration.

§ 1 : La méthode du pivot.

Transformer un tableau de scalaires par la méthode du pivot, c'est lui faire subir des transformations d'un des trois types suivants:

Transformations de première espèce : A une ligne on ajoute λ fois (λ scalaire quelconque) une autre ligne.

$$\text{Exemple : } \begin{array}{ccc} 1 & 2 & 3 \\ 1 & 3 & 0 \\ 2 & 1 & 4 \end{array} \Rightarrow \begin{array}{ccc} 1 & 2 & 3 \\ 1 & 3 & 0 \\ 0 & -5 & 4 \end{array} \quad (\ell_3 - 2 \times \ell_2)$$

Transformations de deuxième espèce: on multiplie une ligne par un scalaire non nul.

$$\text{Exemple : } \begin{array}{ccc} 1 & 2 & 3 \\ 1 & 3 & 0 \\ 2 & 1 & 4 \end{array} \Rightarrow \begin{array}{ccc} 1 & 2 & 3 \\ 2 & 6 & 0 \\ 2 & 1 & 4 \end{array} \quad (2 \times \ell_2)$$

Transformations de troisième espèce: on échange deux lignes.

$$\text{Exemple : } \begin{array}{ccc} 1 & 2 & 3 \\ 1 & 3 & 0 \\ 2 & 1 & 4 \end{array} \Rightarrow \begin{array}{ccc} 1 & 2 & 3 \\ 2 & 1 & 4 \\ 1 & 3 & 0 \end{array}$$

Réduire un tableau par la méthode du pivot c'est, au moyen de telles transformations, le transformer en un tableau ayant un maximum de colonnes dont tous les coefficients, à l'exception d'un seul, sont nuls (ce coefficient non nul est appelé un pivot). Dans cette réduction, on est amené à faire de multiples choix; toutefois le nombre des pivots que l'on obtient est indépendant de ces choix; ce nombre est le rang du tableau.

Interprétation vectorielle de la méthode du pivot: Considérons un espace vectoriel E , muni d'une base $\{e_1, \dots, e_n\}$, et p vecteurs $\{V_1, \dots, V_p\}$ de E . Les coordonnées des V_j dans cette base, forment les colonnes d'un tableau à n lignes et p colonnes.

Faire une transformation de 1-ère, 2-ème ou 3-ème espèce sur ce tableau, c'est le remplacer par le tableau des coordonnées des V_j dans une autre base de E .

Ainsi réduire le tableau des coordonnées d'une famille de vecteurs par la méthode du pivot, c'est trouver une base dans laquelle certains de ces vecteurs ont des coordonnées particulièrement simples (ce sont des vecteurs de la base, éventuellement multipliés par un scalaire); on remarquera d'ailleurs que cette nouvelle base n'est pas déterminée explicitement, on sait seulement qu'elle existe.

§ 2 : Utilisation de la méthode du pivot en algèbre linéaire.

Dans ce qui suit E est un espace de dimension n , et $\{e_1, \dots, e_n\}$ une base de E .

Problème 1: Etant donnés p vecteurs $V_1, \dots, V_p$ dont on connaît les coordonnées dans la base $\{e_1, \dots, e_n\}$, trouver une base de $F = [V_1, \dots, V_p]$; et exprimer les vecteurs V_j dans cette base.

Exemple: $n = 4, p = 3$; nous réduisons le tableau des coordonnées de V_i

$$\begin{array}{ccc|ccc|ccc|ccc} 1 & 2 & 3 & & & & 1 & 0 & -1 & & & & 1 & 0 & -1 \\ 2 & 0 & -2 & \Rightarrow & 0 & -4 & -8 & \Rightarrow & 0 & -4 & -8 & \Rightarrow & 0 & 1 & 2 \\ 3 & 1 & -1 & & 0 & -5 & -10 & & 0 & 0 & 0 & & 0 & 0 & 0 \\ 4 & 3 & 2 & & 0 & -5 & -10 & & 0 & 0 & 0 & & 0 & 0 & 0 \end{array}$$

Il existe une nouvelle base $\{\varepsilon_1, \dots, \varepsilon_4\}$ telle que $V_1 = \varepsilon_1$, $V_2 = \varepsilon_2$ et $V_3 = -\varepsilon_1 + 2\varepsilon_2$. Donc les vecteurs V_1 et V_2 (c'est à dire les vecteurs correspondant aux colonnes contenant les pivots) forment une base de F , et $V_3 = -V_1 + 2V_2$ est l'expression de V_3 dans cette base.

Problème 2: Etant donnés q vecteurs $V_1, \dots, V_q$ qui engendrent E , et tels que $V_1, \dots, V_p$ ($p < q$) soient indépendants, trouver une base de E extraite de $\{V_1, \dots, V_q\}$ et contenant $\{V_1, \dots, V_p\}$.

Puisque $V_1, \dots, V_q$ engendrent E , en réduisant le tableau de leurs coordonnées, nous aurons n pivots. Puisque $V_1, \dots, V_p$ sont indépendants, nous pouvons choisir les p premiers pivots dans les p premières colonnes.

Nous obtenons ainsi le tableau des coordonnées des V_j dans une nouvelle base formée de vecteurs qui sont (au produit par un scalaire près) des vecteurs V_j ; et parmi ces V_j il y a les p premiers.

Problème 3: Etant donné un système libre $\{V_1, \dots, V_p\}$, le compléter en une base.

Nous considérons le système générateur $\{V_1, \dots, V_p, e_1, \dots, e_n\}$, et nous sommes ramenés au problème précédent. On remarquera alors que, lorsque l'on a choisi les p premiers pivots dans les p premières colonnes, le tableau est tout réduit (c'est à dire que $n-p$ des n dernières colonnes n'ont pas été modifiées).

Exemple: Dans $\mathbb{R}^4$, $V_1 = (1; 3; 0; 1)$ et $V_2 = (0; 1; 2; 1)$.

$$\begin{array}{cccc|cccc|cccc|cccc} 1 & 0 & 1 & 0 & 0 & 0 & 1 & 0 & 1 & 0 & 0 & 0 & 1 & 0 & 1 & 0 & 0 & 0 \\ 3 & 1 & 0 & 1 & 0 & 0 & 0 & 1 & -3 & 1 & 0 & 0 & 0 & 1 & -3 & 1 & 0 & 0 \\ 0 & 2 & 0 & 0 & 1 & 0 & 0 & 2 & 0 & 0 & 1 & 0 & 0 & 0 & 6 & -2 & 1 & 0 \\ 1 & 1 & 0 & 0 & 0 & 1 & 0 & 1 & -1 & 0 & 0 & 1 & 0 & 0 & 2 & -1 & 0 & 1 \end{array}$$

Les colonnes correspondant à e_3 et e_4 sont inchangées. Les vecteurs V_1, V_2, e_3, e_4 forment une base.

Problème 4: Trouver un supplémentaire d'un sous-espace F dont on connaît une base $\{V_1, \dots, V_p\}$.

On complète $\{V_1, \dots, V_p\}$ en une base (problème précédent); les vecteurs que l'on a ainsi ajoutés forment une base d'un supplémentaire.

Problème 5: On donne un espace vectoriel F muni d'une base $\{f_1, \dots, f_p\}$, et une application linéaire $\phi: F \rightarrow E$ dont on connaît la matrice M dans les bases $\{f_1, \dots, f_p\}$ et $\{e_1, \dots, e_n\}$. Trouver une base du noyau et une base de l'image de ϕ .

L'image est le sous-espace engendré par les vecteurs $\phi(f_i)$; nous sommes donc ramenés au problème 1.

Exemple: $n = 3$ et $p = 4$ $M = \begin{pmatrix} 1 & 1 & 3 & 0 \\ 2 & 1 & 4 & 1 \\ 3 & 2 & 7 & 1 \end{pmatrix} \Rightarrow \begin{pmatrix} 1 & 1 & 3 & 0 \\ 1 & 0 & 1 & 1 \\ 1 & 0 & 1 & 1 \end{pmatrix} \Rightarrow \begin{pmatrix} 1 & 1 & 3 & 0 \\ 1 & 0 & 1 & 1 \\ 0 & 0 & 0 & 0 \end{pmatrix}$

Donc $\phi(f_2)$ et $\phi(f_4)$ forment une base de l'image de ϕ . Mais de plus, dans la nouvelle base $\{\varepsilon_i\}$ nous avons :

$$\phi(f_1) = \varepsilon_1 + \varepsilon_2 ; \phi(f_2) = \varepsilon_1 ; \phi(f_3) = 3\varepsilon_1 + \varepsilon_2 ; \phi(f_4) = \varepsilon_2$$

Donc : $\phi(f_1) = \phi(f_2) + \phi(f_4)$ et $\phi(f_3) = 3\phi(f_2) + \phi(f_4)$

Et ainsi les vecteurs $f_1 - f_2 - f_4$ et $f_3 - 3f_2 - f_4$ sont dans $\text{Ker } \phi$; ils sont indépendants, et au nombre de deux, donc ils forment une base de $\text{Ker } \phi$.

Nous retiendrons que chaque colonne sans pivot donne un vecteur de $\text{Ker } \phi$, et que ces vecteurs forment une base de ce noyau.

Problème 6: Inverser une matrice.

Soit M une matrice $n \times n$ inversible; écrivons côte à côte M et I_n , puis réduisons (le tableau $n \times 2n$ complet) en choisissant les pivots dans M .

Exemple: $n = 3$

$$\begin{array}{ccc|ccc|ccc|ccc} 1 & 0 & 1 & 1 & 0 & 0 & 1 & 0 & 1 & 1 & 0 & 0 & 0 & -1 & 1 & 2 & -1 & 0 & 0 & 0 & 1 & 5/6 & -1/2 & 1/6 \\ 2 & 1 & 1 & 0 & 1 & 0 & 1 & 1 & 0 & -1 & 1 & 0 & 1 & 1 & 0 & -1 & 1 & 0 & 1 & 0 & 0 & 1/6 & 1/2 & -1/6 \\ 1 & 3 & 4 & 0 & 0 & 1 & -3 & 3 & 0 & -4 & 0 & 1 & 0 & 6 & 0 & -7 & 3 & 1 & 0 & 6 & 0 & -7 & 3 & 1 \end{array}$$

Puis, grâce à des transformations de deuxième et troisième espèces, nous ferons de la partie gauche du tableau la matrice I_n :

$$\begin{array}{ccc|ccc} 0 & 0 & 1 & 5/6 & -1/2 & 1/6 & 1 & 0 & 0 & 1/6 & 1/2 & -1/6 \\ 1 & 0 & 0 & 1/6 & 1/2 & -1/6 & 0 & 1 & 0 & -7/6 & 1/2 & 1/6 \\ 0 & 6 & 0 & -7 & 3 & 1 & 0 & 0 & 1 & 5/6 & -1/2 & 1/6 \end{array}$$

Alors la partie droite du tableau final est la matrice inverse de M . En effet le tableau initial était formé des coordonnées des vecteurs $\varepsilon_1, \dots, \varepsilon_n, e_1, \dots, e_n$ exprimés dans la base $\{e_1, \dots, e_n\}$. Le tableau final est formé des coordonnées des mêmes vecteurs dans une autre base, qui n'est autre que $\{\varepsilon_1, \dots, \varepsilon_n\}$ (pour s'en convaincre, regarder la partie gauche du tableau final). Ainsi M est la matrice de changement de base de $\{\varepsilon_1, \dots, \varepsilon_n\}$ à $\{e_1, \dots, e_n\}$; et la partie droite du tableau final la matrice de changement de base de $\{e_1, \dots, e_n\}$ à $\{\varepsilon_1, \dots, \varepsilon_n\}$, c'est à dire la matrice inverse de M .

Problème 7: Résoudre le système linéaire $\phi(x) = b$, où $\phi: E \rightarrow F$, connaissant les coordonnées b_i de b dans une base $\{f_j\}$ de F , et la matrice M de ϕ dans les bases $\{e_i\}$ de E et $\{f_j\}$ de F .

On réduit le tableau formé de M et de la colonne des b_i , en choisissant les pivots dans M ; c'est à dire que l'on remplace $\{f_j\}$ par une autre base de F . Ceci nous donne alors (Pb. 5) une base de $\text{Im } \phi$; si b est combinaison linéaire des éléments de cette base, il y a (au moins) une solution; sinon il n'y a pas de solution.

§ 3 : Le calcul des déterminants

A tout tableau carré ($n \times n$) de scalaires, est associé un scalaire appelé son déterminant. Avec les propriétés suivantes:

Propriété 1: $\det A = \det A^t$ (où A^t est le tableau transposé du tableau A)

Propriété 2: Les transformations de première espèce (sur les lignes, voir § 1) ne changent pas le déterminant. Donc (Prop. 1) les transformations de première espèce sur les colonnes ne changent pas le déterminant.

Propriété 3: Multiplier une ligne (ou une colonne) du tableau par λ multiplie le déterminant par λ .

En particulier: Si on multiplie tous les coefficients du tableau par λ , on multiplie son déterminant par λ^n .

Si M est une matrice carrée d'ordre n , $\det(\lambda M) = \lambda^n \det M$.

Propriété 4: Echanger deux lignes (ou deux colonnes), multiplie le déterminant par -1 .

Propriété 5: Un tableau diagonal (c'est à dire dont les seuls coefficients non nuls, ont mêmes indices de ligne et de colonne) a pour déterminant le produit de ses coefficients diagonaux.

Propriété 6: Si M et N sont deux matrices carrées d'ordre n , $\det(MN) = \det M \times \det N$

Le déterminant comme polynôme de ses coefficients: Soit $A = (a_{ij}^i)$ un tableau $n \times n$; soit $\mathcal{S}_n$ l'ensemble des permutations de $\{1, \dots, n\}$. Alors

$$\det A = \sum_{\sigma \in \mathcal{S}_n} |\sigma| a_{\sigma(1)}^1 \dots a_{\sigma(n)}^n$$

où $|\sigma|$ est la signature de la permutation σ , c'est à dire:

* $+1$, si σ est le produit d'un nombre pair de transpositions (i.e.: d'échanges de deux indices)

* -1 , si σ est le produit d'un nombre impair de transpositions.

Conséquence: Si les a_{ij}^i sont des polynômes en t , alors $\det A$ est aussi un polynôme en t .

Développement d'un déterminant: Soit $A = (a_{ij}^i)$ un tableau $n \times n$; on appelle mineur de la place (i, j) le déterminant du tableau $(n-1) \times (n-1)$ obtenu en rayant la ligne et la colonne qui contiennent la place (i, j) .

Le cofacteur de la place (i, j) , est le produit de ce mineur par $(-1)^{i+j}$.

Notons α_j^i ce cofacteur; alors:

$$\det A = \sum_{i=1}^n \alpha_j^i a_{ij}^i \quad \text{et} \quad \det A = \sum_{j=1}^n a_{ij}^i \alpha_j^i$$

Exemple: Si $A = \begin{pmatrix} 2 & 3 & 1 \\ 4 & -2 & -1 \\ 5 & 0 & 7 \end{pmatrix}$, alors $\alpha_1^1 = \det \begin{pmatrix} -2 & -1 \\ 0 & 7 \end{pmatrix}$, $\alpha_3^2 = (-1)^{2+3} \det \begin{pmatrix} 2 & 3 \\ 5 & 0 \end{pmatrix}$, etc

Et nous avons les relations

$$\det A = 2 \alpha_1^1 + 3 \alpha_2^1 + 1 \alpha_3^1 = 4 \alpha_1^2 - 2 \alpha_2^2 - 1 \alpha_3^2 = 5 \alpha_1^3 + 0 \alpha_2^3 + 7 \alpha_3^3$$

$$\det A = 2 \alpha_1^1 + 4 \alpha_1^2 + 5 \alpha_1^3 = 3 \alpha_2^1 - 2 \alpha_2^2 + 0 \alpha_2^3 = 1 \alpha_3^1 - 1 \alpha_3^2 + 7 \alpha_3^3$$

Ces formules sont appelées les développements du déterminant de A par rapport à ses lignes et par rapport à ses colonnes. Elles permettent, comme ici, de calculer le déterminant d'un tableau 3×3 , lorsqu'on sait calculer les déterminants des tableaux 2×2 (mais pour cela on peut aussi utiliser la règle de Sarrus). Plus généralement pour calculer le déterminant d'un tableau $n \times n$, il suffira de calculer les déterminants de n tableaux $(n-1) \times (n-1)$.

Conséquence: Inversion d'une matrice carrée. Soit A une matrice $n \times n$ inversible (c'est à dire que $\det A$ est non nul); nous avons les relations

$$\det A = \sum_{j=1}^n \alpha_j^i a_j^i \quad \text{et} \quad \det A = \sum_{j=1}^n a_j^i \alpha_j^i$$

$$0 = \sum_{i=1}^n \alpha_i^j a_k^i \quad (\text{si } k \neq j) \quad \text{et} \quad 0 = \sum_{j=1}^n a_j^i \alpha_k^j \quad (\text{si } k \neq i)$$

Ces deux dernières formules résultent de la remarque suivante: la somme $\sum_{i=1}^n \alpha_i^j a_k^i$ est le déterminant du tableau obtenu en remplaçant $(\forall i)$ le terme a_j^i par a_k^i , c'est à dire du tableau obtenu en remplaçant la j -ème ligne du tableau initial par sa k -ème ligne; ainsi cette somme est le déterminant d'un tableau qui a deux lignes identiques, elle est nulle. L'autre somme est le déterminant d'un tableau qui a deux colonnes égales.

Ces formules peuvent se lire ainsi: notons B la matrice des cofacteurs du tableau A (le coefficient qui se trouve à la place (i,j) dans B , est le cofacteur de la place (i,j) dans A), alors $A \cdot B^t = B^t \cdot A$ est la matrice diagonale dont les coefficients diagonaux sont égaux à $\det A$. Donc B^t est le produit de A^{-1} par $\det A$.

Exemple: Soit $A = \begin{pmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 0 \end{pmatrix}$, la matrice des cofacteurs s'écrit $B = \begin{pmatrix} -48 & 42 & -3 \\ 24 & -21 & 6 \\ -3 & 6 & -3 \end{pmatrix}$. La transposée de B est $\begin{pmatrix} -48 & 24 & -3 \\ 42 & -21 & 6 \\ -3 & 6 & -3 \end{pmatrix}$. On vérifie que $AB = BA = 27 I_3$; et $\det A = 27$.

Ceci nous donne un algorithme pour inverser les matrices au moyen des déterminants, on écrit la matrice B des cofacteurs, on la transpose, puis on la divise par $\det A$.

Cette méthode d'inversion des matrices est peu efficace en grandes dimensions; en dimension 3, elle est commode; et en dimension 2, elle est particulièrement rapide:

Si $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$, alors $B = \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}$, donc $B^t = \begin{pmatrix} d & -b \\ -c & a \end{pmatrix}$, et pour avoir A^{-1} il suffit de diviser par $ad-bc$.

§ 4 : Utilisation des déterminants en algèbre linéaire

Etude du rang d'un système de vecteurs: Considérons un espace E muni d'une base $\{e_1, \dots, e_n\}$ et des vecteurs $V_1, \dots, V_p$ de cet espace. Les coordonnées de ces vecteurs donnent un tableau à p colonnes et n lignes.

Les vecteurs V_j sont indépendants si et seulement si on peut extraire de ce tableau (au moins) un tableau $p \times p$ (ce qui suppose $p \leq n$) dont le déterminant est non nul.

Plus généralement, si tous les déterminants d'ordre $p \times p$, $(p-1) \times (p-1), \dots, (k+1) \times (k+1)$ que l'on peut extraire du tableau, sont nuls, et s'il existe (au moins) un déterminant extrait $k \times k$ non nul, alors le tableau est de rang k (c'est à dire que le sous-espace engendré par les vecteurs $V_1, \dots, V_p$ est de dimension k). De plus les k vecteurs correspondant aux colonnes ayant servi à fabriquer le déterminant $k \times k$ non nul, sont indépendants et forment une base de ce sous-espace.

Exemple: Dans $\mathbb{R}^4$ on donne les vecteurs $V_1(1;3;4;2)$, $V_2(3;2;5;-1)$, $V_3(5;6;11;1)$ et $V_4(1;0;1;-1)$. Le tableau de ces vecteurs s'écrit:

$$\begin{array}{cccc} 1 & 3 & 5 & 1 \\ 3 & 2 & 6 & 0 \\ 4 & 5 & 11 & 1 \\ 2 & -1 & 1 & -1 \end{array}$$

On vérifie que son déterminant est nul. Ceci prouve que le sous-espace engendré par nos 4 vecteurs n'est pas de dimension 4. On vérifie que tous les tableaux 3×3 que l'on peut extraire (attention, il y en a 16) ont un déterminant nul. Ceci prouve que le sous-espace engendré par nos 4 vecteurs est au plus de dimension 2. Le déterminant formé des colonnes 2 et 3 et des lignes 1 et 2 est non nul, donc cette dimension est exactement 2, et les vecteurs V_2 et V_3 forment une base.

Résolution d'un système linéaire: Soit E un espace muni d'une base $\{e_1, \dots, e_n\}$, et des vecteurs $V_1, \dots, V_p, B$ de E . Nous allons chercher tous les systèmes de scalaires $\{x_1, \dots, x_p\}$ tels que $\sum_{i=1}^p x_i V_i = B$;

ce qui revient à résoudre un système linéaire à n équations en les p inconnues $x_1, \dots, x_p$.

Pour cela nous chercherons d'abord le rang k du tableau des coordonnées des V_i (comme ci-dessus).

Systèmes de Cramer: Si $k = p = n$ le système est dit "de Cramer", il a une solution unique, donnée par les formules:

$$\Delta x_k = \Delta_k$$

où Δ est le déterminant du tableau T (carré puisque $n = p$) des coordonnées des V_j ; et où Δ_k est le déterminant du tableau obtenu en remplaçant la k -ième colonne de T , par les coordonnées de B .

Exemple de système de Cramer:

$$\begin{cases} 2x - y + 3z = 1 \\ x + y - 4z = 5 \\ 3x + 2y + 5z = 1 \end{cases}$$

Son déterminant est $\det \begin{pmatrix} 2 & -1 & 3 \\ 1 & 1 & -4 \\ 3 & 2 & 5 \end{pmatrix} = 40$; c'est un système de Cramer, il a une unique solution qui est:

$$40x = \det \begin{pmatrix} 1 & -1 & 3 \\ 5 & 1 & -4 \\ 1 & 2 & 5 \end{pmatrix} \quad 40y = \det \begin{pmatrix} 2 & 1 & 3 \\ 1 & 5 & -4 \\ 3 & 1 & 5 \end{pmatrix} \quad 40z = \det \begin{pmatrix} 2 & -1 & 1 \\ 1 & 1 & 5 \\ 3 & 2 & 1 \end{pmatrix}$$

Systèmes qui ne sont pas de Cramer: On choisit un déterminant $k \times k$ qui est non nul (on l'appelle alors le déterminant principal); ce déterminant fait intervenir k équations, que l'on appellera équations principales, et k inconnues, que l'on appellera inconnues principales.

Pour que le système ait (au moins) une solution, il est nécessaire et suffisant que le vecteur B soit dans le sous-espace engendré par $V_1, \dots, V_p$; pour le vérifier on écrit le tableau des coordonnées des vecteurs V_j correspondant aux inconnues principales, et on lui adjoint la colonne des coordonnées de B ; on obtient un tableau $(k+1) \times n$, et on vérifie que tous les déterminants $(k+1) \times (k+1)$ que l'on peut en extraire sont nuls (ces déterminants sont quelque fois appelés déterminants caractéristiques).

S'il existe des solutions, ce sont celles du système formé des k équations principales. Pour les calculer, on donne des valeurs arbitraires aux inconnues non principales, et il reste alors un système à k équations et k inconnues (les k inconnues principales) qui est "de Cramer". On notera que, dès qu'il existe des inconnues non principales, il ne peut y avoir unicité de la solution.

Exemple: Soit à résoudre

$$\begin{cases} 2x - y + z + t = 5 \\ 3x + y - 3z + 4t = 1 \\ 5x \quad \quad - 2z + 5t = a \end{cases}$$

Le tableau des premiers membres s'écrit:

$$\begin{array}{cccc} 2 & -1 & 1 & 1 \\ 3 & 1 & -3 & 4 \\ 5 & 0 & -2 & 5 \end{array}$$

Les quatre déterminants 3×3 que l'on peut en extraire sont nuls, donc le rang n'est pas 3; en fait le rang est 2, puisqu'il existe des déterminants 2×2 non nuls. On en choisit un (par exemple $\begin{vmatrix} 1 & -3 \\ 0 & -2 \end{vmatrix}$) formé des colonnes 2 et 3, et des lignes 2 et 3; dès lors les inconnues principales sont y et z, les équations principales sont la deuxième et la troisième.

Il y a des solutions si et seulement si le tableau $\begin{array}{ccc} -1 & 1 & 5 \\ 1 & -3 & 1 \\ 0 & -2 & a \end{array}$ (formé des colonnes 2 et 3 du précédent et de la colonne des seconds membres) a un déterminant nul; c'est à dire si $2a - 12 = 0$.

Si $a = 6$, pour tout choix d'une valeur x_0 de x, et d'une valeur t_0 de t, on a une solution unique, obtenue en résolvant le système de Cramer en y et z:

$$\begin{cases} y - 3z = 1 - 3x_0 - 4t_0 \\ -2z = 6 - 5x_0 - 5t_0 \end{cases}$$

§ 5 : Quand doit-on utiliser l'une ou l'autre méthode ?

Nous disposons de deux méthodes pour calculer dans les espaces de dimension finie: la méthode du pivot, et le calcul des déterminants. Quelle est la plus commode ? On va voir que ceci dépend du type de calculs que l'on veut faire.

La méthode du pivot est peu commode lorsque l'on traite des problèmes comportant des paramètres.

Considérons par exemple le système linéaire, d'inconnues x et y $\begin{cases} (1-t)x - ty = 3 \\ ty + (1+t)y = 1 \end{cases}$

Pour le résoudre par la méthode du pivot, nous devons réduire le tableau $\left\{ \begin{array}{cc|c} 1-t & -t & 3 \\ t & 1+t & 1 \end{array} \right\}$

Les seuls pivots possibles sont $1-t$, $-t$, t , $1+t$. Puisqu'un pivot doit être non nul, il nous faudra donc, selon le choix du pivot, supposer que t est différent de 1, 0 ou -1; et il nous faudra reprendre ensuite le calcul (avec un autre choix de pivot) dans ce cas particulier. Pourtant le déterminant du système vaut $(1+t) \times (1-t) + t^2 = 1$; il n'est nul pour aucune valeur de t. Notre système est 'de Cramer' pour tout t. Le calcul par la méthode du pivot introduit une valeur singulière inutile, il est maladroit. Si l'on traite, par la méthode du pivot, un système à paramètres, on verra apparaître ainsi, à chaque choix de pivot, des 'valeurs singulières inutiles' (donc une discussion inutile). Donc tout problème paramétrique doit, dans la mesure du possible être traité au moyen des déterminants.

Mais la méthode des déterminants donne (lorsque l'on considère de grands tableaux) des calculs plus longs que la méthode du pivot.

Essayons par exemple de déterminer si un tableau $n \times n$ est de rang n. Nous pouvons le réduire au moyen de transformations de première espèce (ce qui ne change pas le déterminant); puis au moyen de transformations de troisième espèce, amener les pivots sur la diagonale (ce qui modifie le signe du déterminant de façon contrôlable); et enfin faire le produit des termes diagonaux du tableau final. Chaque transformation de première espèce nécessite $n(n-1)$ soustractions, autant de multiplications, et autant de divisions; soit en tout

$3n(n-1)$ opérations. Et il y a au plus n pivots. Ainsi le nombre total d'opérations est un polynôme de degré 3 en n . Si nous voulons éviter l'emploi des réductions de première espèce, il nous faut utiliser la formule:

$$\det A = \sum_{\sigma \in \mathcal{A}_n} |\sigma| a_{\sigma(1)}^1 \dots a_{\sigma(n)}^n$$

Nous calculerons alors une somme de $n!$ termes (car $\mathcal{A}_n$ a $n!$ éléments); et chacun de ces éléments est un produit de n coefficients du tableau; il y aura donc $n \times n!$ opérations à effectuer. Il est clair que, pour n grand, $n \times n!$ est beaucoup plus grand qu'un polynôme de degré 3 en n . Lorsque l'on calcule avec un ordinateur la question se pose en termes de temps de calcul; pour de grandes valeurs de n la méthode des déterminants donne des temps de calcul tout à fait inacceptables.

Il ne faudrait toutefois pas croire que la méthode du pivot est la panacée. En effet supposons que notre déterminant soit nul. Dans chaque calcul l'ordinateur va faire apparaître des erreurs d'arrondi, qui vont s'accumuler. Ainsi au lieu de trouver après k réductions ($k =$ le rang du tableau) des colonnes nulles, nous verrons apparaître des colonnes dont les coefficients sont petits (ce sont des sommes d'erreurs d'arrondi). Ainsi aucun tableau n'apparaîtra jamais comme de rang plus petit que n ; on pourra toujours trouver n pivots, mais dans certains cas les derniers d'entre eux seront très petits; et il faudra définir un critère permettant de conclure (sans l'avoir réellement démontré) que le rang est plus petit que n .

Exercices sur le chapitre 2

Exercice 1: Dans $\mathbb{R}^4$, on considère les vecteurs $V_1(1;2;3;1)$, $V_2(2;1;0;3)$, $V_3(2;1;0;1)$, $V_4(2;1;1;1)$ et $V_5(3;0;2;1)$. En utilisant successivement la méthode du pivot, puis les déterminants:

- Montrer que V_1 et V_2 sont indépendants.
- Montrer que ces 5 vecteurs engendrent $\mathbb{R}^4$.
- Fabriquer une base de $\mathbb{R}^4$ formée de vecteurs V_j , et contenant V_1 et V_2 .

Exercice 2: Dans $\mathbb{R}^4$, on considère les vecteurs $V_1(3;0;1;-2)$ et $V_2(2;1;3;0)$. En utilisant successivement la méthode du pivot, puis les déterminants:

- Montrer qu'ils sont indépendants.
- Trouver deux vecteurs de la base naturelle de $\mathbb{R}^4$, qui forment avec V_1 et V_2 une base de $\mathbb{R}^4$.

Exercice 3: Dans $\mathbb{R}^5$, on considère $V_1(1;3;0;2;1)$, $V_2(2;1;1;0;3)$ et $V_3(1;2;0;3;4)$.

Montrer qu'ils sont indépendants et trouver deux supplémentaires (distincts) du sous-espace qu'ils engendrent.

Exercice 4: Déterminer si les matrices suivantes sont inversibles, et calculer leurs éventuelles inverses:

a) par la méthode du pivot.

b) au moyen des déterminants.

$$A = \begin{pmatrix} 1 & 2 & 1 \\ 2 & 0 & 3 \\ 3 & 1 & 2 \end{pmatrix} \quad B = \begin{pmatrix} 2 & 3 & 5 \\ 1 & 2 & 3 \\ 0 & 1 & 1 \end{pmatrix} \quad C = \begin{pmatrix} 2 & 3 & 8 \\ 7 & 6 & 4 \\ 5 & 2 & 1 \end{pmatrix}$$

$$D = \begin{pmatrix} 1 & 0 & 1 & 0 \\ 1 & 1 & 1 & 1 \\ 3 & 2 & 3 & 4 \\ 4 & 1 & 0 & 1 \end{pmatrix} \quad E = \begin{pmatrix} 1 & 0 & 2 & 1 \\ 2 & 0 & 0 & 1 \\ 0 & 1 & 3 & 4 \\ 2 & 1 & 0 & 1 \end{pmatrix}$$

Exercice 5: Dans l'espace E_4 des polynômes de degré au plus 4, on considère:

$$P_1(X) = X^4 + X^2 + 1, \quad P_2(X) = X^3 + X \quad \text{et} \quad P_3(X) = X^4 + 2X.$$

- Montrer qu'ils sont indépendants.
- Déterminer des polynômes P_4 et P_5 , tels que $\{P_1, P_2, P_3, P_4, P_5\}$ soit une base de E_4 .
- Est-il possible de choisir pour P_4 et P_5 des polynômes de degré 4 ?
- Est-il possible de choisir pour P_4 et P_5 des polynômes de degré 1 ? de degré 0 ?

Exercice 6: Soit f l'application linéaire de $\mathbb{R}^4$ dans lui-même dont la matrice dans la base naturelle est

$$M = \begin{pmatrix} 1 & 2 & 3 & 4 \\ 4 & 5 & 6 & 1 \\ 4 & 6 & 3 & 2 \\ 1 & 5 & 9 & 8 \end{pmatrix}$$

- Montrer que les vecteurs $V_1 = (1;2;1;2)$, $V_2 = (3;1;2;5)$, $V_3 = (3;2;3;1)$ et $V_4 = (3;1;2;6)$ forment une base de $\mathbb{R}^4$.
- Quelle est la matrice de f dans la base $\{V_1, V_2, V_3, V_4\}$?

Exercice 7: Soit $f: \mathbb{R}^4 \rightarrow \mathbb{R}^2$ dont la matrice dans les bases naturelles $\{e_1, e_2, e_3, e_4\}$ et $\{f_1, f_2\}$ est

$$M = \begin{pmatrix} 3 & 2 & 1 & 5 \\ 6 & 1 & 2 & 4 \end{pmatrix}$$

Trouver une base $\{\varepsilon_1, \varepsilon_2, \varepsilon_3, \varepsilon_4\}$ de $\mathbb{R}^4$ telle que la matrice de f dans les bases $\{\varepsilon_1, \varepsilon_2, \varepsilon_3, \varepsilon_4\}$ et $\{f_1, f_2\}$ s'écrive

$$M' = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \end{pmatrix}$$

- * **Exercice 8:** Soit f et g les applications linéaires de $\mathbb{R}^3$ dans lui-même dont les matrices dans la base naturelle $\{\varepsilon_1, \varepsilon_2, \varepsilon_3\}$ s'écrivent

$$M = \begin{pmatrix} 1 & 2 & 3 \\ 4 & 3 & 2 \\ 1 & 3 & 2 \end{pmatrix} \quad \text{et} \quad N = \begin{pmatrix} 3 & 2 & 1 \\ 1 & 3 & 2 \\ 1 & 3 & 1 \end{pmatrix}$$

Calculer les matrices des applications $f \circ g$ et $g \circ f$ dans la base naturelle; puis dans la base $\{\varepsilon_2, \varepsilon_3, \varepsilon_1\}$.

- * **Exercice 9:** Dans $\mathbb{R}^4$, on considère les vecteurs $V_1(1;2;3;0)$, $V_2(2;0;0;3)$, $V_3(1;1;0;1)$, $V_4(2;1;1;0)$ et $V_5(4;0;2;1)$. En utilisant successivement la méthode du pivot, puis les déterminants:

- Montrer qu'ils engendrent $\mathbb{R}^4$.
- Déterminer toutes les bases de $\mathbb{R}^4$ que l'on peut extraire de cette famille de vecteurs.

- * **Exercice 10:** Dans $\mathbb{R}^4$, on considère les vecteurs $V_1(1;2;3;3)$, $V_2(2;1;1;0)$, $V_3(2;2;0;1)$, $A(3;0;2;1)$. En utilisant d'abord la méthode du pivot, puis les déterminants déterminer si A est dans le sous-espace engendré par V_1 , V_2 et V_3 .

- * **Exercice 11:** Dans $\mathbb{R}^4$, on considère $V_1(1;3;0;1)$, $V_2(2;1;0;3)$ et $V_3(1;0;3;4)$.

Montrer qu'ils sont indépendants et trouver deux supplémentaires (distincts) du sous-espace qu'ils engendrent.

- * **Exercice 12:** Dans $\mathbb{R}^5$, on considère $V_1(1;3;1;2;3)$, $V_2(2;1;0;0;3)$ et $V_3(1;1;0;2;4)$.

Montrer qu'ils sont indépendants et déterminer pour quelles valeurs de α et β le vecteur $W(\alpha, \beta, 3, \alpha+1, \beta-2)$ est dans le sous-espace qu'ils engendrent.

- * **Exercice 13:** Soit f l'application linéaire de $\mathbb{R}^4$ dans $\mathbb{R}^3$ dont la matrice dans les bases naturelles est

$$M = \begin{pmatrix} 1 & 2 & 3 & 4 \\ 6 & 1 & 4 & 6 \\ 3 & 5 & 9 & 8 \end{pmatrix}$$

- Montrer que les vecteurs $V_1(1;0;1;2)$, $V_2(3;1;3;5)$, $V_3(3;2;0;1)$ et $V_4(3;1;2;0)$ forment une base de $\mathbb{R}^4$.
- Montrer que les vecteurs $W_1(1,3,2)$, $W_2(3,2,1)$ et $W_3(4,3,1)$ forment une base de $\mathbb{R}^3$.
- Quelle est la matrice de f dans ces nouvelles bases ?

- * **Exercice 14:** Dans $\mathbb{R}^5$, on considère les vecteurs $V_1(3;0;3;1;-2)$ et $V_2(4;1;1;3;0)$. En utilisant successivement la méthode du pivot, puis les déterminants:

- Montrer qu'ils sont indépendants.
- Trouver trois vecteurs de la base naturelle de $\mathbb{R}^5$, qui forment avec V_1 et V_2 une base de $\mathbb{R}^5$.
- Trouver tous les systèmes de trois vecteurs de la base naturelle qui forment avec V_1 et V_2 une base de $\mathbb{R}^5$.

- * **Exercice 15** Pour quelles valeurs de λ le système suivant a-t-il des solutions ? Calculer alors ces solutions.

$$\begin{array}{rrrrrrr} x & + & y & + & (\lambda-1)z & + & 2t & = & 1 \\ 2\lambda x & + & y & + & (\lambda+2)z & - & \lambda t & = & \lambda \\ (2+3\lambda)x & + & (1+3\lambda)y & - & 3z & + & t & = & 3 \end{array}$$

- * **Exercice 16:** Résoudre le système linéaire suivant, d'abord par la méthode du pivot, puis au moyen des déterminants.

$$\begin{array}{rrrrrrr} 2x & - & y & + & 2z & + & t & = & 1 \\ x & - & 3y & + & 4z & - & t & = & 0 \\ 3x & + & y & - & z & + & 2t & = & 1 \\ & & - & 5y & + & 7z & - & 2t & = & 0 \end{array}$$

- * **Exercice 17:** Pour quelles valeurs du paramètre λ le système suivant a-t-il des solutions ? Calculer alors ces solutions.

$$\begin{array}{rrrrrrr} \lambda x & + & 2y & + & (\lambda-1)z & = & 1 \\ 2x & + & \lambda y & + & (\lambda+2)z & = & \lambda \\ (2+\lambda)x & + & (2+\lambda)y & - & 3z & = & 3 \end{array}$$

Exercice 18: Résoudre le système linéaire suivant, d'abord par la méthode du pivot, puis au moyen des déterminants.

$$\begin{array}{rrcrcl} 2x & - & 2y & + & 3z & + & 2t & = & \alpha \\ x & + & 3y & + & 5z & - & t & = & 0 \\ x & + & 3y & - & 2z & + & 2t & = & 1 \\ & & 5y & - & 7z & - & 2t & = & 2 \end{array}$$

Exercice 19: Calculer les déterminants suivants:

$$A = \begin{pmatrix} 1 & a & a^2 \\ 1 & b & b^2 \\ 1 & c & c^2 \end{pmatrix} \quad B = \begin{pmatrix} 1 & a & a^2 & a^3 \\ 1 & b & b^2 & b^3 \\ 1 & c & c^2 & c^3 \\ 1 & d & d^2 & d^3 \end{pmatrix} \quad C = \begin{pmatrix} 1 & t & t & t \\ t & 1 & t & t \\ t & t & 1 & t \\ t & t & t & 1 \end{pmatrix}$$

$$D = \begin{pmatrix} 1 & a & a^2 & a^4 \\ 1 & b & b^2 & b^4 \\ 1 & c & c^2 & c^4 \\ 1 & d & d^2 & d^4 \end{pmatrix} \quad E = \begin{pmatrix} 1 & t & t^2 & t^3 \\ t & 1 & t^2 & 1 \\ t^2 & t^3 & 1 & t \\ t^3 & 1 & t & t^2 \end{pmatrix}$$

(Dans C et dans E, on cherchera d'abord à mettre $(1-t)^3$ en facteur)

Exercice 20: Dans $\mathbb{R}^5$, on considère les sous-espaces:

E engendré par $V_1(1;0;0;2;4)$, $V_2(2;1;2;3;4)$, $V_3(3;1;4;1;5)$ et $V_4(0;0;-2;4;3)$

F engendré par $W_1(4;1;2;7;12)$, $W_2(6;2;6;6;13)$, $W_3(4;0;1;2;4)$ et $W_4(6,3,7,11,21)$

Déterminer :

- Une base de E.
- Une base de F.
- Une base de $E+F$ contenant la base de E trouvée en a.
- Une base de $E \cap F$.

Exercice 21: Dans $\mathbb{R}^6$ on donne les vecteurs:

$V_1(0;1;1;2;0;0)$

$V_2(1;4;5;8;5;7)$

$V_3(1;2;3;6;1;1)$

$V_4(2;3;5;10;2;2)$

$V_5(0;1;1;1;2;3)$

$W_1(0;1;2;1;1;1)$

$W_2(1;6;8;11;6;8)$

$W_3(0;3;5;3;4;5)$

On appelle E le sous-espace engendré par V_1, V_2, V_3, V_4, V_5 . On appelle F le sous-espace engendré par W_1, W_2, W_3 .

- Parmi les systèmes de vecteurs (V_1, V_2) , (V_1, V_3, V_5) , (V_2, V_3, V_5) , (V_1, V_2, V_4, V_5) , l'un est une base de E. Trouver lequel.
- Trouver une base de F. Trouver une base de $E \cap F$. Quelle est la dimension de $E+F$?
- On considère le sous-espace G de E formé des vecteurs de E (u, v, w, x, y, z) tels que $u+v+w+x+y+z=0$. Trouver une base de G.
- Trouver une base d'un supplémentaire de E dans $\mathbb{R}^6$.

Extrait d'un partiel du 7.11.81

Exercice 22: On considère la matrice

$$M = \begin{pmatrix} 1 & 1 & 0 & 1 \\ 1 & 3 & 2 & 3 \\ -1 & 5 & 3 & 1 \\ 0 & 4 & 2 & 1 \\ 2 & 2 & 0 & 4 \end{pmatrix}$$

Trouver une matrice N telle que $N.M = I_4$.

Extrait d'une épreuve de septembre 1981

** Exercice 23: Soit $M = \begin{pmatrix} 1 & 2 & 1 \\ 3 & 2 & 1 \\ 1 & 1 & 3 \end{pmatrix}$.

a) Montrer que $M^3 - 6M^2 + 3M + 10I_3 = 0$.

b) En déduire que M est inversible, et qu'il existe des nombres α, β, γ tels que $M^{-1} = \alpha I_3 + \beta M + \gamma M^2$.

c) Calculer explicitement M^{-1} , de trois façons: par la méthode du pivot, par les déterminants, et grâce à la formule ci dessus.

** Exercice 24: Soit $M = \begin{pmatrix} 1 & 2 & 2 \\ 2 & 1 & 0 \\ 0 & 1 & 4 \end{pmatrix}$

a) Montrer que $M^3 - 6M^2 + 5M + 8I_3 = 0$.

b) En déduire que M est inversible, et qu'il existe des nombres α, β, γ tels que $M^{-1} = \alpha I_3 + \beta M + \gamma M^2$.

c) Calculer explicitement M^{-1} , de trois façons: par la méthode du pivot, par les déterminants, et grâce à la formule ci dessus.

*** Exercice 25: Soit M une matrice 3×3 telle que $(M - I_3)(M + I_3)^2 = 0$.

a) Montrer que $M^3, M^4, M^5, \dots$ sont combinaisons linéaires de M^2, M et Id .

b) On suppose que $M^2 \neq I$, et que $(M+I)^2 \neq 0$. Soit E le sous-espace vectoriel de $\mathcal{M}(3,3)$ engendré par I_3, M et les puissances de M . Montrer qu'il est de dimension 3, et que I_3, M et M^2 en forment une base.

c) Montrer que M est inversible, et que $M^{-1}, M^{-2}, \dots$ sont dans E .

d) Application: $M = \begin{pmatrix} 1 & 0 & 0 \\ 3 & -1 & 0 \\ 2 & 4 & -1 \end{pmatrix}$. Calculer M^6 et M^{-3} .

Utilisation de l'algèbre linéaire en analyse

Les premiers exemples d'espaces vectoriels que vous avez rencontrés sont l'espace $\vec{\mathbb{P}}$ des vecteurs du plan $\mathbb{P}$, puis l'espace $\vec{\mathbb{E}}$ des vecteurs de l'espace. L'algèbre linéaire vous apparaît peut être comme un outil vous permettant de mieux présenter les calculs de géométrie analytique, et de les généraliser à l'espace abstrait $\mathbb{R}^n$.

Pourtant dans les chapitres qui suivent les espaces vectoriels que nous rencontrerons seront rarement issus de l'algèbre ou de la géométrie; ce seront des espaces de fonctions, qui nous serviront à traiter des problèmes d'analyse. L'objet de ce chapitre 3 est de présenter les plus simples de ces espaces, et les principales applications linéaires qui les relient.

Ceci signifie que nous nous intéresserons moins aux fonctions prises individuellement, qu'à l'ensemble de celles qui vérifient telle ou telle propriété. Et l'intuition géométrique qui est liée à la notion d'espace vectoriel, nous servira à interpréter les problèmes d'analyse. C'est dans cet esprit que les rappels sur les équations différentielles linéaires ont été reformulés. C'est aussi dans cet esprit qu'a été posé (complément du chapitre 3, et complément 1 du chapitre 4) le problème de l'interpolation (Que sait-on d'une fonction lorsque l'on connaît ses valeurs en n points?).

§1 : Quelques espaces vectoriels de fonctions

Soit X un ensemble, notons $\mathcal{F}(X, \mathbb{R})$ l'ensemble des applications de X dans $\mathbb{R}$. Nous savons ajouter deux applications de X dans $\mathbb{R}$: $(\forall x \in X) \quad (f+g)(x) = f(x) + g(x)$. Nous savons multiplier une application de X dans $\mathbb{R}$ par un scalaire réel: $(\forall x \in X) \quad (\lambda f)(x) = \lambda f(x)$. Ceci fait de $\mathcal{F}(X, \mathbb{R})$ un espace vectoriel sur $\mathbb{R}$.

Dans les problèmes d'analyse que nous rencontrerons l'ensemble X sera un intervalle I (borné ou non, fermé ou non), et nous nous intéresserons rarement à $\mathcal{F}(I, \mathbb{R})$ tout entier, mais plutôt à certains de ses sous-espaces. En voici quelques uns.

L'espace $\mathcal{C}^0(I, \mathbb{R})$ des fonctions continues sur I .

Dire que $\mathcal{C}^0(I, \mathbb{R})$ (que l'on notera aussi $\mathcal{C}(I, \mathbb{R})$) est un sous-espace vectoriel de $\mathcal{F}(I, \mathbb{R})$, c'est à dire que la somme de deux fonctions continues est continue, et que le produit d'une fonction continue par un scalaire, est une fonction continue.

L'espace $\mathcal{C}^1(I, \mathbb{R})$ des fonctions de classe C^1 sur I .

Nous appellerons ainsi les fonctions (continues et) dérivables sur I , dont la dérivée est elle même continue sur I . Si f et g sont de classe C^1 , alors $f+g$ et λf ($\lambda \in \mathbb{R}$) sont (c'est bien connu) de classe C^1 ; ce qu'on traduit par: $\mathcal{C}^1(I, \mathbb{R})$ est un sous-espace vectoriel de $\mathcal{F}(I, \mathbb{R})$ - et d'ailleurs aussi de $\mathcal{C}^0(I, \mathbb{R})$.

L'espace $\mathcal{C}^2(I, \mathbb{R})$ des fonctions de classe C^2 sur I .

Nous appellerons ainsi les fonctions (continues et) 2 fois dérivables sur I , dont la dérivée seconde est une fonction continue sur I . Si f et g sont de classe C^2 , alors $f+g$ et λf le sont aussi. Ce qu'on traduit par: $\mathcal{C}^2(I, \mathbb{R})$ est un sous espace vectoriel de $\mathcal{F}(I, \mathbb{R})$ - et d'ailleurs aussi de $\mathcal{C}^0(I, \mathbb{R})$ et de $\mathcal{C}^1(I, \mathbb{R})$.

L'espace $\mathcal{C}^p(I, \mathbb{R})$ des fonctions de classe C^p sur I .

Plus généralement une fonction est dite de classe C^p , si elle a des dérivées jusqu'à l'ordre p et si sa dérivée d'ordre p est continue. Si f et g sont de classe C^p , alors $(\forall k \leq p) (f+g)^{(k)} = f^{(k)} + g^{(k)}$ et $(\lambda f)^{(k)} = \lambda f^{(k)}$ (où $\lambda \in \mathbb{R}$); donc $f+g$ et λf sont aussi de classe C^p ; ce qu'on exprime en disant que $\mathcal{C}^p(I, \mathbb{R})$ est un sous-espace vectoriel de $\mathcal{F}(I, \mathbb{R})$ - et d'ailleurs aussi de $\mathcal{C}^q(I, \mathbb{R})$, pour $q < p$.

L'espace $\mathcal{C}^\infty(I, \mathbb{R})$ des fonctions de classe C^∞ sur I .

Une fonction est dite indéfiniment dérivable, ou de classe C^∞ , si elle a des dérivées de tous ordres. Ainsi $\mathcal{C}^\infty(I, \mathbb{R}) = \bigcap_{k \in \mathbb{N}} \mathcal{C}^k(I, \mathbb{R})$. Et l'intersection d'une famille quelconque de sous-espaces vectoriels d'un espace E , est un sous-espace vectoriel de E . Donc $\mathcal{C}^\infty(I, \mathbb{R})$ est un sous-espace de $\mathcal{F}(I, \mathbb{R})$ - et aussi de tous les $\mathcal{C}^k(I, \mathbb{R})$.

L'espace des fonctions polynômes sur I .

C'est le sous-espace des précédents engendré par les vecteurs (ie: les applications de I dans $\mathbb{R}$) $x \rightarrow 1$; $x \rightarrow x$; $x \rightarrow x^2$; ...; $x \rightarrow x^n$, etc; Nous le noterons $\text{Polyn}(I, \mathbb{R})$.

Remarque 1 : Le dernier de ces sous-espaces est - c'est bien connu - de dimension infinie (ie: il possède, quelque soit k , des systèmes libres à k éléments). Il en résulte que tous les autres sont aussi de dimension infinie. En fait $\mathcal{F}(X, \mathbb{R})$ est de dimension finie si (et seulement si) X est fini (Les applications qui prennent la valeur 1 en un point, et 0 en tous les autres, en forment alors une base).

Remarque 2: Dans tout ceci nous pouvons remplacer les fonctions à valeurs réelles par les fonctions à valeurs complexes (le domaine de définition restant un intervalle de $\mathbb{R}$). Nous obtiendrons alors des sous-espaces $\mathcal{C}^0(I, \mathbb{C}), \dots, \mathcal{C}^k(I, \mathbb{C}), \dots, \mathcal{C}^\infty(I, \mathbb{C})$ et $\text{Polyn}(I, \mathbb{C})$ (polynômes à coefficients complexes) de $\mathcal{F}(I, \mathbb{C})$.

Notons qu'une fonction à valeurs complexes est dite continue si les fonctions $x \rightarrow \text{Re}(f(x))$ et $x \rightarrow \text{Im}(f(x))$ sont continues. (Pour des précisions sur cette définition, voir chap 10). Et f est dite dérivable si $x \rightarrow \text{Re}(f(x))$ et $x \rightarrow \text{Im}(f(x))$ sont dérivables; la dérivée de f est alors définie par :

$$\frac{d}{dx}(f(x)) = \frac{d}{dx}(\text{Re } f(x)) + i \frac{d}{dx}(\text{Im } f(x))$$

§ 2 : Quelques applications linéaires en analyse

La dérivation: Les formules classiques $(f+g)' = f' + g'$ et $(\lambda f)' = \lambda f'$ ($\lambda \in \mathbb{R}$) traduisent le fait que la dérivation définit, pour tout $k \geq 1$, une application linéaire de $\mathcal{C}^k(I, \mathbb{R})$ dans $\mathcal{C}^{k-1}(I, \mathbb{R})$ (et aussi de $\mathcal{C}^\infty(I, \mathbb{R})$ dans lui-même).

L'intégrale: Les formules $\int_a^b (f(t) + g(t)) dt = \int_a^b f(t) dt + \int_a^b g(t) dt$ et $\int_a^b \lambda f(t) dt = \lambda \int_a^b f(t) dt$

traduisent le fait que l'application qui à f associe son intégrale sur $[a, b]$, est une forme linéaire sur $\mathcal{C}^0([a, b], \mathbb{R})$.

La valeur en un point: A toute f associons sa valeur au point x . Evidemment $(f+g)(x) = f(x) + g(x)$ et $(\lambda f)(x) = \lambda f(x)$; autrement dit nous avons défini une forme linéaire sur $\mathcal{C}^0(I, \mathbb{R})$. Nous la noterons souvent ev_x (lire: évaluation en x).

Primitives nulles en a : A toute fonction f dans $\mathcal{C}^0(I, \mathbb{R})$, nous pouvons associer la primitive de f qui est nulle en a (a étant un point choisi une fois pour toutes). C'est la fonction: $x \rightarrow \int_a^x f(t) dt$, qui est de classe C^1 . Nous obtenons ainsi une application linéaire de $\mathcal{C}^0(I, \mathbb{R})$ dans $\mathcal{C}^1(I, \mathbb{R})$; et aussi de $\mathcal{C}^k(I, \mathbb{R})$ dans $\mathcal{C}^{k+1}(I, \mathbb{R})$.

D'autres exemples: En modifiant, composant, ajoutant,... ces quatre exemples, on obtient beaucoup d'autres applications linéaires définies sur des espaces de fonctions. En voici quelques unes (Voir aussi les exercices, en particulier les premiers)

* A toute f on associe $x \rightarrow \int_{x_0}^x f(t) \cos t dt$

C'est une application linéaire $\mathcal{C}^0(I, \mathbb{R}) \rightarrow \mathcal{C}^1(I, \mathbb{R})$.

* A toute f on associe la dérivée k -ième $f^{(k)}$.

C'est une application linéaire de $\mathcal{C}^p(I, \mathbb{R})$ dans $\mathcal{C}^{p-k}(I, \mathbb{R})$.

* A toute f dans $\mathcal{C}^k(I, \mathbb{R})$, on associe son polynôme de Taylor à l'ordre k en x_0

$$f(x_0) + (x-x_0) f'(x_0) + \dots + \frac{(x-x_0)^k}{k!} f^{(k)}(x_0)$$

On a ainsi défini une application linéaire de $\mathcal{C}^k(I, \mathbb{R})$ dans l'espace des polynômes de degré au plus k .

Remarque: Tout ceci reste valable lorsque l'on considère des fonctions à valeurs complexes. Les espaces obtenus sont alors des espaces vectoriels sur $\mathbb{C}$. Notons que par définition, lorsque f est à valeurs complexes:

$$\int_a^b f(t) dt = \int_a^b \operatorname{Re}(f(t)) dt + i \int_a^b \operatorname{Im}(f(t)) dt$$

§ 3 : Equations différentielles linéaires du premier ordre

Soit a une fonction continue sur I . A toute fonction y dans $\mathcal{C}^1(I, \mathbb{R})$ associons la fonction $x \rightarrow y'(x) + a(x) y(x)$; nous obtenons une application Φ linéaire (vérification facile) de $\mathcal{C}^1(I, \mathbb{R})$ dans $\mathcal{C}^0(I, \mathbb{R})$.

Soit b dans $\mathcal{C}^0(I, \mathbb{R})$; rechercher les fonctions y telles que $\Phi(y) = b$, c'est résoudre l'équation différentielle: $y' + ay = b$. Ainsi notre équation différentielle est une équation linéaire (au sens du § 3 du chapitre I). Nous allons donc:

* D'abord rechercher $\text{Ker } \Phi$, autrement dit résoudre l'équation homogène $y' + ay = 0$. On vérifie que $x \rightarrow \lambda e^{-A(x)}$ (où on a posé $A(x) = \int_{x_0}^x a(t) dt$) est ($\forall \lambda \in \mathbb{R}$) une solution, et qu'il n'y en a pas d'autre. Autrement dit $\text{Ker } \Phi$ est un espace de dimension 1, engendré par la fonction $x \rightarrow e^{-A(x)}$.

* Ensuite chercher une solution particulière de l'équation proposée. On vérifie facilement que $x \rightarrow e^{-A(x)} \int_{x_0}^x b(t) e^{A(t)} dt$ en est une. La solution générale de notre équation différentielle est :

$$x \rightarrow \lambda e^{-A(x)} + e^{-A(x)} \int_{x_0}^x b(x) e^{A(x)} dx$$

On voit que la "recette" classique qui permet de résoudre ces équations, est mot pour mot la démarche générale de résolution des équations linéaires.

§ 4 : Equations différentielles linéaires du second ordre à coefficients constants.

Soit a et b dans $\mathcal{C}^0(I, \mathbb{R})$; définissons $\Phi: \mathcal{C}^2(I, \mathbb{R}) \rightarrow \mathcal{C}^0(I, \mathbb{R})$, par $\Phi(y) = y'' + ay' + by$. Cette application Φ est linéaire (vérification facile).

Soit g dans $\mathcal{C}^0(I, \mathbb{R})$; résoudre l'équation linéaire $\Phi(y) = g$, c'est résoudre l'équation différentielle du 2^d ordre :

$$y'' + ay' + by = g$$

Nous allons, selon la méthode exposée au § 3 du chapitre 1, scinder le problème en deux.

* On cherchera d'abord le noyau de Φ , c'est à dire les solutions de l'équation différentielle homogène $y'' + ay' + by = 0$. Ce noyau est de dimension 2 (nous le démontrerons au chap. 20, admettons le pour l'instant). Il nous faut en trouver une base, c'est à dire 2 éléments indépendants.

Si a et b sont des constantes, nous écrirons l'équation $r^2 + ar + b = 0$

Si elle a :

- a) 2 racines réelles distinctes r et r' , on prend pour base $x \rightarrow e^{rx}$ et $x \rightarrow e^{r'x}$.
- b) une racine réelle double r , on prend pour base $x \rightarrow e^{rx}$ et $x \rightarrow x e^{rx}$.
- c) deux racines complexes conjuguées $\alpha \pm i\omega$, on prend pour base $x \rightarrow e^{\alpha x} \sin \omega x$ et $x \rightarrow e^{\alpha x} \cos \omega x$.

Notons que si a et b ne sont pas des constantes nous n'avons actuellement aucun moyen pour trouver une base de $\text{Ker } \Phi$.

On cherchera ensuite une solution particulière de l'équation $\Phi(y) = g$. On peut retenir qu'il y en a toujours une. Se reporter à un cours de 1^{ère} année pour trouver les "recettes" permettant de calculer une telle solution.

Complément : Eléments d'une théorie de l'interpolation.

Le problème de l'interpolation: Etant donnés des nombres $a_0, \dots, a_n$ ($2 \leq n$ distincts) et $b_0, \dots, b_n$, trouver une fonction f dans $\mathcal{C}^k(\mathbb{R}, \mathbb{R})$ telle que $(\forall i) f(a_i) = b_i$.

Considérons ($k \geq 0$ arbitrairement choisi) l'application $\Phi: \mathcal{C}^k(\mathbb{R}, \mathbb{R}) \rightarrow \mathbb{R}^{n+1}$ définie par $\Phi(f) = (f(a_0), \dots, f(a_n))$. On vérifie aisément qu'elle est linéaire.

Montrons que Φ n'est pas injective: l'espace $\mathcal{C}^k(\mathbb{R}, \mathbb{R})$ n'est pas de dimension finie, nous pouvons donc y choisir $n+2$ vecteurs indépendants, qui engendrent un sous-espace F de dimension

$n+2$. Si Φ était injective, sa restriction à F le serait aussi. Ce serait donc une application injective d'un espace de dimension $n+2$ dans un espace de dimension $n+1$. Ce qui est impossible; c'est donc que Φ n'est pas injective.

Montrons que Φ est surjective: Considérons le polynôme

$$P_0(x) = \frac{(x - a_1)(x - a_2) \dots (x - a_n)}{(a_0 - a_1)(a_0 - a_2) \dots (a_0 - a_n)}$$

Il prend la valeur 1 en a_0 , et 0 en $a_1, \dots, a_n$. Donc $\Phi(P_0)$ est le premier vecteur de la base naturelle de $\mathbb{R}^{n+1}$. De même si

$$P_1(x) = \frac{(x - a_0)(x - a_2) \dots (x - a_n)}{(a_1 - a_0)(a_1 - a_2) \dots (a_1 - a_n)}$$

alors $\Phi(P_1)$ est le second vecteur de la base naturelle de $\mathbb{R}^{n+1}$. Et ainsi de suite. Il en résulte que l'image de Φ contient tous les vecteurs de la base naturelle de $\mathbb{R}^{n+1}$; elle contient donc aussi l'ensemble de leurs combinaisons linéaires, c'est à dire $\mathbb{R}^{n+1}$ tout entier.

Il en résulte que, quelles que soient $b_0, \dots, b_n$, il existe plusieurs fonctions de $\mathcal{E}^k(\mathbb{R}, \mathbb{R})$ qui ($\forall i$) prennent la valeur b_i en a_i .

Recherche d'un polynôme solution: Considérons maintenant le sous-espace E_n de $\mathcal{E}^k(\mathbb{R}, \mathbb{R})$ formé des polynômes de degré au plus n ($\dim E_n = n+1$). Alors la restriction de Φ à E_n est une bijection de E_n sur $\mathbb{R}^{n+1}$. En effet les polynômes P_i ci-dessus sont dans E_n ; donc $\Phi(E_n)$ contient les $\Phi(P_i)$ c'est à dire les vecteurs de la base naturelle de $\mathbb{R}^{n+1}$; donc $\Phi(E_n)$ est $\mathbb{R}^{n+1}$ tout entier. Ainsi la restriction de Φ à E_n est une surjection de E_n sur $\mathbb{R}^{n+1}$; c'est aussi une bijection puisque les deux espaces ont même dimension.

Il existe ainsi un polynôme P de degré n et un seul, qui, quel que soit i , prend la valeur b_i en a_i . Et pour calculer P nous disposons de 2 méthodes.

Première méthode de calcul: Le polynôme P s'écrit $p_0 x^n + p_1 x^{n-1} + \dots + p_n$. Ecrire qu'il prend la valeur b_i en a_i , c'est écrire que les coefficients p_i vérifient la relation linéaire

$$p_0 a_i^n + p_1 a_i^{n-1} + \dots + p_n = b_i$$

Et nous obtenons un système linéaire ($n+1$ équations, $n+1$ inconnues $p_0, \dots, p_n$) dont nous savons qu'il a une solution unique (en résolvant on doit donc constater qu'il est "de Cramer").

Seconde méthode de calcul: Le polynôme $b_0 P_0 + b_1 P_1 + \dots + b_n P_n$ (où les P_i sont les polynômes définis ci dessus) prend la valeur b_i en a_i , puisque:

$$(b_0 P_0 + b_1 P_1 + \dots + b_n P_n)(a_i) = b_0 P_0(a_i) + b_1 P_1(a_i) + \dots + b_n P_n(a_i) = b_i P_i(a_i) = b_i$$

(Puisque $P_j(a_i)$ vaut 0 si $i \neq j$, et 1 si $i=j$); c'est le polynôme P . On se gardera de croire que cette seconde méthode, apparemment très simple, est beaucoup plus rapide que la première; car le calcul explicite des polynômes P_i est assez pénible.

Il reste que le problème le plus souvent rencontré est le suivant: étant donnée une fonction dont on connaît seulement les valeurs b_i aux points a_i , essayer de retrouver f . Notre polynôme d'interpolation P et f prennent les mêmes valeurs en $n+1$ points; ils ne sont pas pour autant égaux. Il faut donc chercher à quelles conditions P peut être considéré comme une "bonne approximation" de f .

Exercices sur le chapitre 3

* Exercice 1: Voici 6 applications de $\mathcal{C}^0([0, \infty), \mathbb{R})$ dans $\mathbb{R}$. Quelles sont celles qui sont linéaires ?

$$a) f \rightarrow \int_0^a \cos(f(t)) dt$$

$$b) f \rightarrow \int_0^a f(t) \cos t dt$$

$$c) f \rightarrow \int_0^a f(1) f(t) \cos t dt$$

$$d) f \rightarrow \int_0^a f^2(t) dt$$

$$e) f \rightarrow \int_0^a \left[\int_0^u f(t) dt \right] du$$

$$f) f \rightarrow f(a) + \int_0^a f(t) e^{-t} dt$$

* Exercice 2: Voici 6 applications de $\mathcal{C}^\infty([0, 1], \mathbb{R})$ dans lui-même. Quelles sont celles qui sont linéaires ?

$$a) \text{ à } f \text{ on associe } (x \rightarrow \int_0^x f^3(t) dt)$$

$$b) \text{ à } f \text{ on associe } (x \rightarrow \int_0^{x^2} f(t^3) dt)$$

$$c) \text{ à } f \text{ on associe } (x \rightarrow f'(x^2))$$

$$d) \text{ à } f \text{ on associe } (x \rightarrow f''(x^2 - x + 1))$$

$$e) \text{ à } f \text{ on associe } (x \rightarrow \left(\int_0^1 f(t) \cos t dt \right) \cos x + \left(\int_0^1 f(t) \sin t dt \right) \sin x)$$

$$f) \text{ à } f \text{ on associe } (x \rightarrow \left(\int_0^1 f(t) dt \right) \times \left(\int_0^x f(t) dt \right))$$

* Exercice 3: Soit $\phi: \mathcal{C}^\infty(\mathbb{R}, \mathbb{R}) \rightarrow \mathcal{C}^\infty(\mathbb{R}, \mathbb{R})$ qui à toute fonction associe sa dérivée seconde, montrer que ϕ est linéaire et déterminer son noyau et son image.

** Exercice 4: Soit $\phi: \mathcal{C}^\infty(\mathbb{R}, \mathbb{R}) \rightarrow \mathcal{C}^\infty(\mathbb{R}, \mathbb{R})$ qui à toute fonction f associe la fonction $t \rightarrow \int_3^t f(\tau) e^\tau d\tau$.

Montrer que ϕ est linéaire et déterminer son noyau et son image.

** Exercice 5: A toute fonction f continue sur $\mathbb{R}$, nous associons une fonction L en posant: $L(x) = \int_{-x}^{+x} f(t) dt$

a) Montrer que l'on a ainsi défini une application linéaire de $\mathcal{C}^0(\mathbb{R}, \mathbb{R})$ dans $\mathcal{C}^1(\mathbb{R}, \mathbb{R})$.

b) Montrer que l'image de cette application linéaire, est le sous-espace vectoriel de $\mathcal{C}^1(\mathbb{R}, \mathbb{R})$ formé des fonctions impaires.

c) Quel est le noyau de cette application linéaire ?

* Exercice 6: On définit une application de $\mathcal{C}^k([0, 1], \mathbb{R})$ dans $\mathcal{C}^0([0, 1], \mathbb{R})$, en associant à toute f , sa dérivée $f^{(k)}$ d'ordre k .

Montrer que l'on a ainsi défini une application linéaire de $\mathcal{C}^k([0, 1], \mathbb{R})$ dans $\mathcal{C}^0([0, 1], \mathbb{R})$.

Quel est son noyau ? Montrer qu'elle est surjective (on pourra commencer par $k=1$ ou $k=2$)

* Exercice 7: Des exemples d'applications linéaires d'un espace dans lui-même, qui sont surjectives et non injectives.

a) A toute f dans $\mathcal{C}^\infty([0, 1], \mathbb{R})$, on associe sa dérivée. Montrer qu'on a ainsi défini une application linéaire Φ de $\mathcal{C}^\infty([0, 1], \mathbb{R})$ dans lui-même. Montrer qu'elle est surjective. Quel est son noyau ?

b) Quel est le noyau de $\Phi^k (= \Phi \circ \Phi \circ \dots \circ \Phi, k \text{ fois})$? Quelle est son image ?

* Exercice 8: Soit F l'application linéaire $\mathcal{C}^0(\mathbb{R}, \mathbb{R}) \rightarrow \mathcal{C}^0(\mathbb{R}, \mathbb{R})$ qui à toute f , associe la fonction $x \rightarrow f(-x)$.

- Quel est le noyau de $F + \text{Id}$? Son image ?
- Montrer qu'ils sont supplémentaires.
- Calculer F^2 , et comparer avec l'exercice 13 du chapitre 1.

* Exercice 9: Soit E le sous-espace de $\mathcal{C}^0([0, 2\pi], \mathbb{R})$ engendré par les fonctions

$$f_0: x \rightarrow 1; \quad f_1: x \rightarrow \cos x; \quad f_2: x \rightarrow \cos 2x; \quad \dots; \quad f_p: x \rightarrow \cos p x; \quad \dots$$

- A toute f on associe $\int_0^{2\pi} f_k(t) f(t) dt$. Montrer qu'on a ainsi défini une forme linéaire sur E . Montrer que f_i est dans

son noyau si $i \neq k$. Est-ce que f_k est dans ce noyau ?

- En déduire que, quelque soit p , les vecteurs $f_0, \dots, f_p$ de E sont indépendants.

** Exercice 10: On considère l'équation différentielle $y'' + 3y' + 2y = 0$.

- La résoudre
- A tout couple (α, β) on associe la seule solution Y de cette équation telle que $Y(1) = \alpha$ et $Y'(1) = \beta$. Montrer que l'on définit ainsi une application linéaire de $\mathbb{R}^2$ dans l'espace des solutions. Montrer que cette application est bijective.
- A (α, β) on associe $(Y(0), Y'(0))$. Que peut-on dire de l'application de $\mathbb{R}^2$ dans lui-même ainsi définie ?

** Exercice 11: On définit une application linéaire de $\mathcal{C}^1(\mathbb{R}, \mathbb{R})$ dans lui-même, en associant à toute fonction f , la fonction

$$x \rightarrow \int_0^x t f'(t) dt$$

- Montrer que son image est formée des fonctions g telles que $g(0) = 0$, et que $\frac{g'(x)}{x}$ se prolonge en une fonction continue en 0.
- Quel est son noyau ?

*** Exercice 12: On cherche les solutions de l'équation différentielle $y'' + 2y' + y = \cos t$ qui prennent les valeurs α et β en 0 et en π respectivement.

On note Φ l'application de $\mathcal{C}^2([0, \pi], \mathbb{R})$ dans $\mathcal{C}^0([0, \pi], \mathbb{R})$ qui à y associe $y'' + 2y' + y$.

- Déterminer $\text{Ker } \Phi$.
- Déterminer une solution de l'équation $y'' + 2y' + y = \cos t$.
- On considère l'application $\theta: \text{Ker } \Phi \rightarrow \mathbb{R}^2$ qui à y associe $\theta(y) = (y(0), y(\pi))$. Démontrer qu'elle est surjective. En déduire que le problème a une unique solution quels que soient α et β . Puis calculer cette solution.
- Reprendre tout le problème, mais avec l'équation différentielle $y'' + y = te^t$ (on observera que dans ce cas θ n'est ni injective, ni surjective).

** Exercice 13: On se place dans $\mathcal{C}^\infty(\mathbb{R}, \mathbb{R})$, et on définit les 3 applications

$$\begin{aligned} \Phi_0: \mathcal{C}^\infty(\mathbb{R}, \mathbb{R}) &\rightarrow \mathcal{C}^\infty(\mathbb{R}, \mathbb{R}) & \text{par } \Phi_0(u) &= u'' + 4u' - 5u \\ \Phi_1: \mathcal{C}^\infty(\mathbb{R}, \mathbb{R}) &\rightarrow \mathcal{C}^\infty(\mathbb{R}, \mathbb{R}) & \text{par } \Phi_1(u) &= u'' + 6u' \\ \varphi: \mathcal{C}^\infty(\mathbb{R}, \mathbb{R}) &\rightarrow \mathcal{C}^\infty(\mathbb{R}, \mathbb{R}) & \text{par } (t \rightarrow u(t)) &\rightarrow (t \rightarrow e^{-t} u(t)) \end{aligned}$$

- Montrer quelles sont linéaires. Montrer que $\Phi_1 \circ \varphi = \varphi \circ \Phi_0$. En déduire que φ définit un isomorphisme de $\text{Ker } \Phi_0$ sur $\text{Ker } \Phi_1$.
- Résoudre l'équation $u''(t) + 4u'(t) - 5u(t) = \cos 2t$.

Exercice 14: Soit $\varphi: \mathcal{C}^0(\mathbb{R}, \mathbb{R}) \rightarrow \mathcal{C}^0(\mathbb{R}, \mathbb{R})$ l'application qui à toute fonction f associe $f \circ \cos$.

Montrer qu'elle est linéaire. Quel est son noyau ?

- *** Exercice 15 : Soit $\varphi: \mathcal{C}^0(\mathbb{R}, \mathbb{R}) \rightarrow \mathcal{C}^0(\mathbb{R}, \mathbb{R})$ l'application qui à toute fonction f associe $t \rightarrow f(t+2\pi) - f(t)$.
- Montrer qu'elle est linéaire.
 - Quel est son noyau ?
 - Montrer que les fonctions affines sont dans l'image de φ . Montrer (Plus difficile) que les fonctions nulles en 0 et en 2π sont dans l'image de φ (on montrera que $(\forall g$ nulle en 0 et en 2π) il existe f telle que $\varphi(f) = g$ et qui est nulle entre 0 et en 2π). En déduire que φ est surjective.
- ** Exercice 16: On dit qu'une fonction (définie sur $\mathbb{R}$ ou sur une partie de $\mathbb{R}$) est lipschitzienne, s'il existe λ tel que $(\forall x$ et $y)$ $|f(x) - f(y)| \leq \lambda |x - y|$.
- Soit $\mathcal{L}(a, b)$ l'ensemble des fonctions lipschitziennes de $[a, b]$ dans $\mathbb{R}$. Montrer que $\mathcal{L}(a, b)$ est un sous-espace vectoriel de $\mathcal{C}^0([a, b], \mathbb{R})$. Montrer que $\mathcal{C}^1([a, b], \mathbb{R})$ est un sous-espace vectoriel de $\mathcal{L}(a, b)$ (on pourra utiliser le théorème des accroissements finis).
 - Soit $\mathcal{L}(\mathbb{R})$ l'ensemble des fonctions lipschitziennes de $\mathbb{R}$ dans $\mathbb{R}$. Montrer que $\mathcal{L}(\mathbb{R})$ est un sous-espace vectoriel de $\mathcal{C}^0(\mathbb{R}, \mathbb{R})$. Est-ce que $\mathcal{C}^1(\mathbb{R}, \mathbb{R})$ est un sous-espace vectoriel de $\mathcal{L}(\mathbb{R})$?
- * Exercice 17: Soit $\varphi: \mathcal{C}^0(\mathbb{R}, \mathbb{R}) \rightarrow \mathcal{C}^0(\mathbb{R}, \mathbb{R})$ l'application qui à toute fonction f associe $t \rightarrow f(t) - \int_0^t \tau f(\tau) d\tau$.
- Montrer qu'elle est linéaire.
 - Déterminer $\text{Ker } \varphi$ (on montrera d'abord que toute fonction qui est dans $\text{Ker } \varphi$ est dérivable et est solution de l'équation différentielle $f'(t) = t f(t)$).
- * Exercice 18: Soit E l'ensemble des suites $(u_n)_{n \in \mathbb{N}}$ de nombres réels. Montrer que E est un espace vectoriel.
- Montrer que l'ensemble F des suites convergentes est un sous-espace vectoriel de E .
 - Montrer que l'application qui à tout élément de F associe sa limite, est linéaire.
- ** Exercice 19: Pour tout nombre réel $T (>0)$ on note E_T l'ensemble des fonctions f continues de $\mathbb{R}$ dans $\mathbb{R}$ telles que $(\forall x)$ $f(x+T) = f(x)$.
- E_T est-il un sous espace vectoriel de $\mathcal{C}^0(\mathbb{R}, \mathbb{R})$?
 - Soit E la réunion de tous les E_T , est-ce un sous-espace vectoriel de $\mathcal{C}^0(\mathbb{R}, \mathbb{R})$?
 - Soit $E_{\mathbb{N}}$ la réunion des E_T pour tous les T entiers positifs, est-ce un sous-espace vectoriel de $\mathcal{C}^0(\mathbb{R}, \mathbb{R})$?
 - Soit $E_{\mathbb{Q}}$ la réunion des E_T pour tous les T rationnels positifs, est-ce un sous-espace vectoriel de $\mathcal{C}^0(\mathbb{R}, \mathbb{R})$?
- * Exercice 20: Soit φ l'application de $\mathcal{C}^0([-1, 1], \mathbb{R})$ dans $\mathcal{C}^0([-1, 0], \mathbb{R}) \times \mathcal{C}^0([0, 1], \mathbb{R})$, qui à toute f associe $(f|_{[-1, 0]}, f|_{[0, 1]})$. Elle est linéaire. Est-elle surjective ? Est-elle injective ?
- * Exercice 21: On donne des nombres $(\alpha_1, \dots, \alpha_n)$ deux à deux distincts. Et on considère dans $\mathcal{C}^0(\mathbb{R}, \mathbb{R})$ les fonctions $e^{\alpha_1 t}, \dots, e^{\alpha_n t}$. Montrer qu'elles sont linéairement indépendantes. (On pourra montrer d'abord qu'il existe α tel que, pour toute valeur de i sauf une, la fonction $e^{\alpha} e^{\alpha_i t}$ ait pour limite $+\infty$ en $+\infty$)

Produits scalaires

Les espaces $\mathcal{P}$ et $\mathcal{E}$ de la géométrie possèdent un certain nombre de propriétés qui n'ont pas été généralisées à $\mathbb{R}^n$ (dans les chapitres 1 et 2). Ce sont toutes celles qui font intervenir les notions de norme, de produit scalaire, d'angle, ... Nous allons généraliser ces notions à des espaces absolument quelconques.

Nous généraliserons d'abord la notion de produit scalaire, puis nous en déduirons toutes les autres. Les espaces $\mathcal{P}$ et $\mathcal{E}$ de la géométrie possèdent un produit scalaire naturel. Au contraire, un espace vectoriel $\mathcal{F}$ ne possède pas naturellement un produit scalaire; on peut en définir de nombreux. Et lorsque l'on a choisi l'un d'entre eux (on dit aussi que l'on a muni $\mathcal{F}$ d'une structure euclidienne), il est possible de parler de norme d'un vecteur, de base orthonormée, d'angle, ... Ces notions dépendent du produit scalaire choisi.

Lorsque $\mathcal{F}$ est muni d'une structure euclidienne, parmi les applications linéaires de $\mathcal{F}$ dans $\mathcal{F}$ (on parle plutôt d'endomorphisme de $\mathcal{F}$), certaines jouent un rôle particulier, ce sont celles qui conservent le produit scalaire (i.e.: le produit scalaire de deux vecteurs est égal à celui de leurs images). On les appelle les applications orthogonales de $\mathcal{F}$ dans $\mathcal{F}$, ou les isométries de $\mathcal{F}$ dans $\mathcal{F}$.

Il faut comprendre ces constructions dans le même esprit que le chapitre 3 ci-dessus. Mettre une structure euclidienne sur un espace $\mathcal{F}$, et en particulier sur un espace de fonctions, c'est se donner les moyens d'appliquer à $\mathcal{F}$ un certain nombre de schémas de raisonnements qui nous sont familiers en géométrie. Le complément 1 est une illustration de ce principe dans le cas des espaces de polynômes.

§ 1 : La notion de produit scalaire

Soit E un espace vectoriel réel, un produit scalaire (ou une structure euclidienne) sur E , est une application $E \times E \rightarrow \mathbb{R}$ (notée ici $(x, y) \rightarrow \langle x, y \rangle$), telle que les propriétés suivantes soient vérifiées:

- 1) $(\forall x) (\forall y) \langle x, y \rangle = \langle y, x \rangle$
- 2) $(\forall y)$ l'application $x \rightarrow \langle x, y \rangle$ ($E \rightarrow \mathbb{R}$) est linéaire (donc, d'après 1, $(\forall x)$ l'application $y \rightarrow \langle x, y \rangle$ ($E \rightarrow \mathbb{R}$) est linéaire).
- 3) $(\forall x \text{ non nul}) \langle x, x \rangle$ est un nombre strictement positif.

Nous aurons l'occasion, en cours d'année, d'utiliser des produits scalaires sur divers espaces de dimension finie, et aussi sur des espaces de fonctions. Voici, en exercice, quelques exemples.

Exercice a: Dans $\mathbb{R}^2$, si $u = (x, y)$ et $u' = (x', y')$, on pose :

$$\langle u, u' \rangle = 2x x' + 2y y' + x y' + y x'$$

Montrer que l'on a défini un produit scalaire sur $\mathbb{R}^2$ (pour montrer la condition 3, on écrira $2x^2 + 2y^2 + 2xy = 2(x\frac{y}{2})^2 + \frac{3}{2}y^2$).

Exercice b: Soit E l'espace des fonctions continues de $[0,1]$ dans $\mathbb{R}$. On pose : $\langle f, g \rangle = \int_0^1 f(t) g(t) dt$.

Montrer que l'on a ainsi défini un produit scalaire sur E (pour montrer que $\langle f, f \rangle = 0$ implique $f = 0$, on remarque que $\langle f, f \rangle = 0$, implique $\varphi(x) = \int_0^x f(t) f(t) dt = 0$ quel que soit x , et on calcule $\varphi'(x)$).

Exercice c: Dans l'exercice b, remplaçons $\int_0^1 f(t) g(t) dt$ par $\int_0^{1/2} f(t) g(t) dt$. A-t-on encore un produit scalaire ?

Exercice d: Soit E_n l'espace des polynômes de degré au plus n . Soit $n+1$ nombres $x_0, \dots, x_n$ (2 à 2 distincts). On pose $\langle P, Q \rangle = \sum_{i=0}^n P(x_i) Q(x_i)$

Montrer que l'on a ainsi défini un produit scalaire sur E_n (Rappel : Le seul polynôme de degré $\leq n$ qui est nul en $x_0, \dots, x_n$, est le polynôme nul).

Exercice e: Sur $\mathbb{R}^4$ on pose:

$$\langle (x, y, z, t), (X, Y, Z, T) \rangle = 4xX + 4yY + 3zZ + 8tT + (xY + XY) + 2(xZ + XZ) + (yT + YT) + 3(zT + ZT)$$

Montrer que l'on a ainsi défini un produit scalaire. Pour montrer la condition 3, on écrira, selon la méthode dite de Gauss,

$$\langle (x, y, z, t), (x, y, z, t) \rangle = (\alpha x + \beta y + \gamma z + \delta t)^2 + F(y, z, t); \text{ puis } \langle (x, y, z, t), (x, y, z, t) \rangle = (\alpha x + \beta y + \gamma z + \delta t)^2 + (\varepsilon y + \varphi z + \eta t)^2 + G(z, t); \text{ etc.}$$

§ 2 : Norme euclidienne et angles

Lemme de Schwarz: *Quels que soient x et y , $(\langle x, y \rangle)^2 \leq \langle x, x \rangle \langle y, y \rangle$.*

En effet : $\langle x, y \rangle^2 - \langle x, x \rangle \langle y, y \rangle$ est le Δ réduit du trinôme du second degré en t :

$$\langle x, x \rangle + 2t \langle x, y \rangle + t^2 \langle y, y \rangle = \langle x + ty, x + ty \rangle$$

Ce trinôme est toujours positif ou nul, donc son Δ réduit est négatif ou nul.

Norme euclidienne: La norme euclidienne du vecteur x est le nombre $\|x\|$ défini par

$$\|x\| = \sqrt{\langle x, x \rangle}$$

Ses propriétés fondamentales sont :

$$\alpha) \|x\| = 0 \Leftrightarrow x = 0$$

$$\beta) (\forall x) (\forall \lambda) \|\lambda x\| = |\lambda| \|x\|$$

$$\gamma) (\forall x) (\forall y) \|x + y\| \leq \|x\| + \|y\| \text{ (inégalité triangulaire)}$$

$$\gamma') (\forall x) (\forall y) \|x - y\| \geq |\|x\| - \|y\||$$

Les propriétés α et β résultent simplement des propriétés du produit scalaire. La propriété γ est une conséquence du Lemme de Schwarz ; en effet :

$$\begin{aligned} \|x + y\|^2 &= \langle x + y, x + y \rangle \\ &= \langle x, x \rangle + \langle y, y \rangle + 2 \langle x, y \rangle \\ &\leq \langle x, x \rangle + \langle y, y \rangle + 2 \sqrt{\langle x, x \rangle \langle y, y \rangle} = (\|x\| + \|y\|)^2 \end{aligned}$$

La propriété γ' est l'application de γ à $x = y + (x - y)$ et à $y = x + (y - x)$.

Remarque: Sur un même espace nous pouvons définir de nombreux produits scalaires. Par exemple sur $\mathbb{R}^2$ si $e = (x, y)$ et $e' = (x', y')$, nous pouvons poser:

$$\langle e, e' \rangle = x x' + y y'$$

$$\text{ou bien } \langle e, e' \rangle = x x' + 2y y'$$

$$\text{ou bien } \langle e, e' \rangle = 2x x' + 2y y' + x y' + y x'$$

A ces produits scalaires distincts correspondent des normes distinctes.

Angles: L'angle des deux vecteurs non nuls est défini par $(X, Y) = \arccos \frac{\langle X; Y \rangle}{\|X\| \|Y\|}$

On notera que, d'après le lemme de Schwartz $-1 \leq \frac{\langle X; Y \rangle}{\|X\| \|Y\|} \leq 1$, ce qui donne un sens à notre définition. Nous obtenons ainsi un angle entre 0 et π . Dès que la dimension de l'espace est supérieure à 2, il n'est plus possible de définir des angles orientés comme dans le plan de la géométrie élémentaire.

En particulier, deux vecteurs sont dits orthogonaux si leur produit scalaire est nul (c'est-à-dire, si l'un au moins d'entre eux est nul, ou si leur angle est $\pi/2$)

Exercice f: Soit $X = (1; 3)$ et $Y = (6; -1)$ dans $\mathbb{R}^2$, déterminer leur angle pour chacun des produits scalaires définis dans la remarque ci-dessus. Vérifier ainsi que l'angle de 2 vecteurs dépend du produit scalaire choisi.

Exercice g: Démontrer que, quels que soient x et y , $2\langle x, y \rangle = \|x+y\|^2 - \|x\|^2 - \|y\|^2$. En déduire que l'on a le "théorème de Pythagore" : si x et y sont orthogonaux $\|x+y\|^2 = \|x\|^2 + \|y\|^2$.

§ 3 : Familles orthogonales et orthonormées

Une famille de vecteurs $\{e_i\}$ est dite orthogonale si, $(\forall i, j) \langle e_i, e_j \rangle = 0$. Elle est dite orthonormée si, de plus $(\forall i) \|e_i\| = 1$

Remarque : Une famille orthogonale finie dont les vecteurs sont tous non nuls, est une famille libre.

En effet $\sum_i \lambda_i e_i = 0$ entraîne, pour tout j , $\langle 0, e_j \rangle = \langle \sum_i \lambda_i e_i, e_j \rangle = \sum_i \lambda_i \langle e_i, e_j \rangle = \lambda_j \|e_j\|^2$; et

puisque $\|e_j\| \neq 0$, ceci implique $\lambda_j = 0$.

En particulier dans un espace de dimension n , les familles orthonormées ont au plus n éléments; et toute famille orthonormée ayant n éléments est une base orthonormée.

Le procédé de Schmidt: Pour fabriquer des familles orthonormées, et en particulier des bases orthonormées, on dispose du procédé de Schmidt :

Etant donnée une famille libre $\{e_1, \dots, e_p\}$ de vecteurs de E , il existe une famille $\{\varepsilon_1, \dots, \varepsilon_p\}$ orthonormée telle que $(\forall k)$ les sous espaces vectoriels $[e_1, \dots, e_k]$ et $[\varepsilon_1, \dots, \varepsilon_k]$ coïncident.

Construction des ε_i :

* On pose $\varepsilon_1 = \frac{1}{\|e_1\|} e_1$

* Le vecteur $e_2 + \lambda \varepsilon_1$ est orthogonal à ε_1 , pourvu que $\langle e_2, \varepsilon_1 \rangle + \lambda = 0$. On choisit donc $\lambda = -\langle e_2, \varepsilon_1 \rangle$ et on pose:

$$\varepsilon_2 = \frac{1}{\|e_2 + \lambda \varepsilon_1\|} (e_2 + \lambda \varepsilon_1)$$

* Le vecteur $e_3 + \mu_1 \varepsilon_1 + \mu_2 \varepsilon_2$ est orthogonal à ε_1 et ε_2 pourvu que $\langle e_3, \varepsilon_1 \rangle + \mu_1 = 0$ et $\langle e_3, \varepsilon_2 \rangle + \mu_2 = 0$. On choisit donc $\mu_1 = -\langle e_3, \varepsilon_1 \rangle$ et $\mu_2 = -\langle e_3, \varepsilon_2 \rangle$; et on pose:

$$\varepsilon_3 = \frac{1}{\|e_3 + \mu_1 \varepsilon_1 + \mu_2 \varepsilon_2\|} (e_3 + \mu_1 \varepsilon_1 + \mu_2 \varepsilon_2)$$

* Etc...

Si la famille $\{e_i\}$ a p éléments, après p étapes, on obtient une famille orthonormée à p éléments. Si on part d'une suite infinie $\{e_1, \dots, e_p, \dots\}$ libre (c'est-à-dire telle que $(\forall N) \{e_1, \dots, e_N\}$ soit libre), par récurrence sur N , cette construction nous donne une suite infinie $\{e_1, \dots, e_N, \dots\}$ qui est orthonormée.

Si l'on part d'une base $\{e_1, \dots, e_n\}$ de E , on construit une base orthonormée de E .

Utilisation des bases orthonormées.

Calcul du produit scalaire: En base orthonormée si $X = \sum_i x_i e_i$ et $Y = \sum_i y_i e_i$, alors le produit scalaire est donné par la formule $\langle X, Y \rangle = \sum_i x_i y_i$.

Dans une base quelconque, puisque le produit scalaire est linéaire par rapport à ses deux arguments, on a $\langle X, Y \rangle = \sum_{i,j} x_i y_j \langle e_i, e_j \rangle$. En base est orthonormée, $\langle e_i, e_i \rangle = 1$ et $\langle e_i, e_j \rangle = 0$ pour $i \neq j$, d'où la formule annoncée.

Calcul des coordonnées: En base orthonormée si $X = \sum_i x_i e_i$, alors $(\forall j) x_j = \langle X, e_j \rangle$

En effet $\langle X, e_j \rangle = \sum_i x_i \langle e_i, e_j \rangle$ et si la base est orthonormée ceci se réduit à x_j .

Exercice h: Sur $\mathbb{R}^3$ on définit un produit scalaire par :

$$\langle (x, y, z), (x', y', z') \rangle = 2x x' + x y' + y x' + 2y y' + z z'$$

Fabriquer une base orthonormée en appliquant le procédé de Schmidt à la base naturelle $(1; 0; 0)$, $(0; 1; 0)$ et $(0; 0; 1)$.

Exercice i: On reprend le produit scalaire de l'exercice b. Fabriquer un système orthonormé de 4 vecteurs, en appliquant le procédé de Schmidt aux 4 vecteurs

$$t \rightarrow 1, t \rightarrow t, t \rightarrow t^2 \text{ et } t \rightarrow t^3$$

Remarque : Dans l'espace E_3 des vecteurs de la géométrie, une base $\{e_1, e_2, e_3\}$ définit des sous-espaces $[e_1] \subset [e_1, e_2] \subset [e_1, e_2, e_3] = E_3$. Le procédé de Schmidt consiste à choisir un vecteur unitaire de $[e_1]$; puis un vecteur unitaire orthogonal à e_1 situé dans le plan $[e_1, e_2]$; enfin un vecteur unitaire orthogonal à ce plan.

§ 4 : Sous-espaces orthogonaux

Soit F un sous-espace de E , notons $F^\perp$ l'ensemble des vecteurs e de E qui sont orthogonaux à tous les vecteurs de F , autrement dit $F^\perp = \{e \in E \text{ tels que } (\forall v \in F) \langle v, e \rangle = 0\}$. Il est clair que $F^\perp$ est un sous-espace vectoriel de E . Ce sous-espace est appelé le "sous-espace orthogonal à F ".

Supposons que E soit de dimension finie, et prenons une base orthonormée $\{e_i\}_{i \leq n}$ de E , telle que $F = [e_1, \dots, e_p]$. Alors un vecteur $\sum_{i \leq n} \lambda_i e_i$ est dans $F^\perp$ si et seulement si $\lambda_i = 0$ pour $i \leq p$ (en

effet : d'une part les vecteurs e_j sont dans $F^\perp$ si $j > p$, et d'autre part $\sum_{i \leq n} \lambda_i e_i \in F^\perp$ implique

l'égalité $\langle \sum_{i \leq n} \lambda_i e_i, e_j \rangle = 0$, c'est à dire $\lambda_j = 0$, pour $j \leq p$). Il en résulte que $F^\perp = [e_{p+1}, \dots, e_n]$.

Nous retiendons: si E est de dimension finie, F et $F^\perp$ sont supplémentaires.

Projections orthogonales: Soit E un espace euclidien et soit F un sous-espace de dimension finie de E . Considérons une base orthonormée $[f_i]$ de F . Pour tout vecteur x dans E , posons $\langle x, f_i \rangle = x_i$, et $x' = \sum_i x_i f_i$. Alors, pour tout i , $\langle x', f_i \rangle = \langle x, f_i \rangle$; ce qui prouve que le vecteur $x - x'$ est orthogonal à tous les vecteurs f_i , et par conséquent à tous les vecteurs de F .

Le point x' est appelé la projection orthogonale de x sur F . C'est le point de F qui est le plus proche de x (au sens de la norme euclidienne).

En effet: si $y \in F$ ($y \neq x$), alors

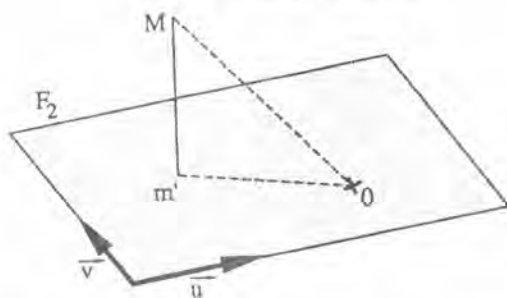
$$\|y - x\|^2 = \|x' - x\|^2 + \|y - x'\|^2 + 2\langle x' - x, y - x' \rangle = \|x' - x\|^2 + \|y - x'\|^2 \geq \|x' - x\|^2$$

Il en résulte que le point x' ainsi construit ne dépend pas du choix de la base $\{f_i\}$.

Exemple: Choisissons une origine O dans l'espace de la géométrie; nous en faisons un espace vectoriel de dimension 3.

Soit F_1 un sous-espace de dimension 1, c'est à dire une droite passant par O . C'est un résultat classique de géométrie élémentaire que la projection orthogonale d'un point M sur F_1 , est le point m de F_1 le plus proche de M . Nous retiendrons que si $\vec{u}$ est un vecteur unitaire de F_1 , alors

$$Om = (OM, \vec{u}) \vec{u}.$$



Soit F_2 un sous-espace de dimension 2, c'est à dire un plan passant par O . C'est un résultat classique que la projection orthogonale m' de M sur F_2 , est le point de F_2 le plus proche de M . Retenons que si $(\vec{u}, \vec{v})$ est une base orthonormée de F_2 , alors:

$$Om' = (OM, \vec{u}) \vec{u} + (OM, \vec{v}) \vec{v}$$

Exercice j: Soit $(O, \vec{i}, \vec{j}, \vec{k})$ un repère orthonormé de l'espace. Soit H le plan d'équation $2x - y + z = 1$. Soit $M(x_0, y_0, z_0)$. Quelles sont les coordonnées de la projection orthogonale de M sur H ? Quelles sont les coordonnées du symétrique de M par rapport à H ?

Exercice k: On reprend le produit scalaire de l'exercice b. Soit F le sous-espace formé des polynômes de degré au plus 1. Quelle est la projection orthogonale de la fonction $x \rightarrow \cos \pi x$ sur F .

§ 5 : Endomorphismes orthogonaux et matrices orthogonales

Soit $f: E^n \rightarrow E^n$ un endomorphisme, les conditions suivantes sont équivalentes

$$1) (\forall X, Y) \langle X, Y \rangle = \langle f(X), f(Y) \rangle$$

$$2) (\forall X) \|X\| = \|f(X)\|$$

3) L'image par f d'une base orthonormée est une base orthonormée.

Un tel endomorphisme est dit orthogonal (on dit aussi que c'est une isométrie vectorielle).

Démontrons l'équivalence de 1) 2) et 3):

1) $\Rightarrow$ 2) Puisque 2) est un cas particulier de 1)

2) $\Rightarrow$ 3) Car si $\{e_i\}$ est orthonormée, 2) entraîne que $\|f(e_i)\| = \|e_i\| = 1$ et aussi que (si $i \neq j$)
 $2\langle f(e_i), f(e_j) \rangle = \|f(e_i) + f(e_j)\|^2 - \|f(e_i)\|^2 - \|f(e_j)\|^2 = \|e_i + e_j\|^2 - \|e_i\|^2 - \|e_j\|^2 = 2 - 1 - 1 = 0.$

3) $\Rightarrow$ 1) Parceque, si $\{e_i\}$ est orthonormée et si $x = \sum_i \lambda_i e_i$ et $y = \sum_i \mu_i e_i$, alors
 $\langle x, y \rangle = \sum_i \lambda_i \mu_i$. Tandis que $\langle f(x), f(y) \rangle = \sum_{i,j} \lambda_i \mu_j \langle f(e_i), f(e_j) \rangle = \sum_i \lambda_i \mu_i$ (car $\{f(e_i)\}$ est orthonormée).

Exercice I: Montrer que tout endomorphisme orthogonal est injectif (donc bijectif si E est de dimension finie).

Soit H une matrice $n \times n$, on dit qu'elle est orthogonale si $H^t H = I_n$. Les matrices orthogonales sont donc inversibles, et leur inverse est égale à leur transposée. En particulier on a aussi $H H^t = I_n$.

Soit $f: E \rightarrow E$ un endomorphisme. Soient $\{e_i\}$ et $\{e_j\}$ deux bases orthonormées, et soit M la matrice de $f: (E, \{e_i\}) \rightarrow (E, \{e_j\})$. Alors M est orthogonale si et seulement si f est orthogonal. En particulier:

Si deux bases sont orthonormées la matrice de changement de base est orthogonale.

Si f est orthogonale, sa matrice dans une base orthonormée $\{e_i\}$ est orthogonale.

En effet les coefficients de la $i^{\text{ème}}$ colonne de M (et de la $i^{\text{ème}}$ ligne de M^t) sont les coefficients de $f(e_i)$ dans la base $\{e_j\}$. Il résulte donc de la définition du produit de matrices, que $M^t M$ est la matrice $(\langle f(e_i), f(e_j) \rangle)$. Donc $M^t M = I_n$ équivaut à :

$$\langle f(e_i), f(e_j) \rangle = \delta_{ij} \quad (= 0 \text{ si } i \neq j, \text{ et } = 1 \text{ si } i = j)$$

Ce qui exprime que les $f(e_i)$ forment une famille orthonormée.

Exemples :

1) En dimension 2 les matrices orthogonales sont les matrices de la forme $\begin{pmatrix} a & c \\ b & d \end{pmatrix}$ telles que $a^2 + b^2 = c^2 + d^2 = 1$ et $ac + bd = 0$. En posant $a = \cos \theta$, $b = \sin \theta$, $c = \cos \varphi$ et $d = \sin \varphi$, on voit que $ac + bd = \cos(\theta - \varphi) = 0$. Donc $\theta - \varphi = \pm \frac{\pi}{2} \pmod{2\pi}$.

* Si $\varphi = \theta + \frac{\pi}{2}$, nous obtenons la matrice $R_\theta = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}$ (dét $R_\theta = 1$) L'application linéaire qui dans une base orthonormée directe admet R_θ pour matrice, est la rotation d'angle θ .

* Si $\varphi = \theta - \frac{\pi}{2}$, nous obtenons la matrice $S_\theta = \begin{pmatrix} \cos \theta & \sin \theta \\ \sin \theta & -\cos \theta \end{pmatrix}$ (dét $S_\theta = -1$) L'application linéaire qui dans la base orthonormée directe $(\vec{i}, \vec{j})$ admet S_θ pour matrice, est la symétrie orthogonale par rapport à la droite D_θ telle que $(\vec{i}, D_\theta) = \frac{\theta}{2} \pmod{\pi}$.

Donc en dimension 2, les endomorphismes orthogonaux sont les rotations et les symétries orthogonales.

2) En dimension 3 les symétries orthogonales sont des endomorphismes orthogonaux. Les rotations aussi. Si l'on compose une rotation avec la symétrie par rapport au plan orthogonal à son axe, on obtient encore un endomorphisme orthogonal. Nous verrons au chapitre 5 (§5) qu'on a ainsi obtenu tous les endomorphismes orthogonaux en dimension 3.

3) En dimension 4: Ecrivons $E = F \oplus F^\perp$, où F (et donc aussi $F^\perp$) est de dimension 2. Soit α et β des endomorphismes orthogonaux de F et $F^\perp$ respectivement. Définissons l'application

$\varphi: (E=F) \oplus F^\perp \rightarrow (E=F) \oplus F^\perp$ par $\varphi(x+y) = \alpha(x) + \beta(y)$. Alors φ est orthogonal. Nous verrons (chap 5 §5) que tous les endomorphismes orthogonaux de E peuvent être obtenus de cette façon.

§ 6 : Extension au cas des espaces vectoriels sur $\mathbb{C}$

Considérons maintenant un espace vectoriel complexe E , et essayons de refaire la théorie ci-dessus. Si $\varphi: E \times E \rightarrow \mathbb{C}$ vérifie 1) et 2), elle ne peut pas vérifier 3); en effet si $\varphi(X, X) > 0$, alors $\varphi(iX, iX) = i^2 \varphi(X, X) < 0$.

C'est pourquoi dans le cas des espaces vectoriels sur $\mathbb{C}$ nous changerons quelque peu la définition:

Nous appellerons *produit hermitien* sur E , une application $E \times E \rightarrow \mathbb{C}$ (on notera $(x|y)$ l'image du couple (x, y)), telle que

$$1) (\forall X, Y) \quad (X|Y) = \overline{(Y|X)}$$

2) $(\forall Y)$ l'application $X \rightarrow (X|Y)$ ($E \rightarrow \mathbb{C}$) est linéaire (donc $(\forall X)$ l'application $Y \rightarrow (X|Y)$ n'est pas linéaire mais vérifie $(X|Y_1 + \lambda Y_2) = (X|Y_1) + \bar{\lambda}(X|Y_2)$).

$$3) \forall X \text{ non nul} : (X|X) > 0.$$

On peut alors reprendre tout ce qui a été dit dans le cas réel; les seules modifications sont:

* Le lemme de Schwarz s'écrit :

$$((X|Y) + (Y|X))^2 \leq 4 (X|X) (Y|Y) \quad (\text{Notons que } (X|Y) + (Y|X) = 2 \operatorname{Re} (X|Y))$$

* Il n'est pas possible de définir l'angle de 2 vecteurs, puisque $\frac{(X|Y)}{\|X\| \|Y\|}$ n'est, en général, pas réel. Mais on a une notion d'orthogonalité définie par $(X|Y) = 0$.

* Au § 4 on parlera d'endomorphismes unitaires au lieu d'endomorphismes orthogonaux. De plus les matrices orthogonales seront remplacées par les matrices unitaires, c'est-à-dire les matrices qui vérifient $\bar{U}^t \cdot U = I_n$.

Complément 1: La méthode des moindres carrés.

Nous avons (complément au chap.3) montré qu'étant donnés $n+1$ nombres ($2 \leq n$ distincts) $a_0, \dots, a_n$, et $n+1$ nombres $b_0, \dots, b_n$, il existe un unique polynôme P , de degré au plus n , tel que $(\forall i) P(a_i) = b_i$.

Par exemple si :

$a_0 = -2$	$b_0 = 0,9968$	$a_3 = 0,5$	$b_3 = 3,502503125$
$a_1 = -1$	$b_1 = -0,9991$	$a_4 = 1$	$b_4 = 7,0031$
$a_2 = 0$	$b_2 = 1,002$	$a_5 = 2$	$b_5 = 17,0072$

Nous obtiendrons $P(x) = 0,0001x^5 + 2x^2 + 4,001x + 1,002$. Une telle solution est en fait peu satisfaisante puisqu'il existe un polynôme Q de degré 2 ($Q(x) = 2x^2 + 4x + 1$), dont la valeur en a_i diffère de b_i de moins de $1/100$ (quelque soit i). Compte tenu des erreurs d'arrondi, ou des inexactitudes expérimentales qui interviennent le plus souvent dans les données d'un problème numérique, il est dommage dans un tel cas d'ignorer cette solution approchée. La méthode des moindres carrés va nous permettre de repérer de tels phénomènes de façon systématique, puisqu'elle va nous donner, pour tout $q < n$, le polynôme de degré q qui - en un certain sens à préciser - approche le mieux le polynôme P .

Un produit scalaire sur l'espace des polynômes de degré au plus n : Notons E_n l'espace (de dim. $n+1$) des polynômes de degré au plus n . Définissons, pour P et Q dans E_n :

$$\langle P, Q \rangle = \sum_{i=0}^n P(a_i) Q(a_i)$$

Il est clair que:

$$a) \langle P, Q \rangle = \langle Q, P \rangle$$

$$b) \langle \lambda P_1 + P_2, Q \rangle = \lambda \langle P_1, Q \rangle + \langle P_2, Q \rangle$$

Et par ailleurs $\langle P, P \rangle = 0$ entraîne que $[P(a_0)]^2 + \dots + [P(a_n)]^2 = 0$, c'est à dire $(\forall i) P(a_i) = 0$. Ce qui entraîne que P est nul (car un polynôme de degré au plus n , ne peut avoir $n+1$ racines). Nous avons donc défini un produit scalaire sur E_n .

Une base orthonormée (pour ce produit scalaire): Nous savons que les polynômes $p_0: x \rightarrow 1$, $p_1: x \rightarrow x$, ..., $p_n: x \rightarrow x^n$ forment une base de E_n . Cette base n'est pas orthonormée, mais nous pouvons lui appliquer le procédé de Schmidt. Nous obtiendrons ainsi une base orthonormée $\varphi_0, \dots, \varphi_n$.

Nous savons que φ_0 et p_0 sont colinéaires. Autrement dit φ_0 est un polynôme de degré 0 (vérifier que $\varphi_0 = \frac{1}{\sqrt{n+1}}$). Nous savons que $\varphi_1 = \alpha p_0 + \beta p_1$ et que φ_1 n'est pas colinéaire à p_0 ; donc φ_1 est un polynôme de degré exactement 1. De même φ_k est dans $[p_0, \dots, p_k]$, mais n'est pas dans $[p_0, \dots, p_{k-1}]$; donc φ_k est de degré exactement k .

Une nouvelle méthode de calcul du polynôme d'interpolation: Grâce à cette base $\varphi_0, \dots, \varphi_n$ nous pouvons donner une nouvelle méthode de calcul du polynôme d'interpolation P .

Nous savons que

$$P = \sum_{j=0}^n \langle P, \varphi_j \rangle \varphi_j$$

Soit

$$P = \sum_{j=0}^n \left[\sum_{i=0}^n P(a_i) \varphi_j(a_i) \right] \varphi_j = \sum_{j=0}^n \left[\sum_{i=0}^n b_i \varphi_j(a_i) \right] \varphi_j$$

Cette formule nous donne des calculs plus compliqués que celles que nous avons vues dans le complément au chap. 3. Mais elle a le mérite de faire apparaître les polynômes:

$$[P]_k = \sum_{j=0}^k \left[\sum_{i=0}^n b_i \varphi_j(a_i) \right] \varphi_j$$

qui sont les projections orthogonales de P sur les espaces E_k (cf: § 4). Ainsi $[P]_k$ est le polynôme de degré au plus k , qui - au sens de notre produit scalaire - est le plus proche de P . Cette distance de P à $[P]_k$ peut se calculer de la façon suivante:

$$\begin{aligned} \|P - [P]_k\|^2 &= \|P - \sum_{j=0}^k \langle P, \varphi_j \rangle \varphi_j\|^2 &= \left\| \sum_{j=k+1}^n \langle P, \varphi_j \rangle \varphi_j \right\|^2 \\ &= \sum_{j=k+1}^n (\langle P, \varphi_j \rangle)^2 &= \sum_{j=0}^n (\langle P, \varphi_j \rangle)^2 - \sum_{j=0}^k (\langle P, \varphi_j \rangle)^2 \\ &= \|P\|^2 - \|[P]_k\|^2 &= \sum_{i=0}^n (b_i)^2 - \sum_{j=0}^k \left(\sum_{i=0}^n b_i \varphi_j(a_i) \right)^2 \end{aligned}$$

Chercher une approximation de P par la méthode des moindres carrés c'est - sans jamais calculer P complètement - déterminer les polynômes $[P]_0, \dots, [P]_k$ jusqu'à ce que la quantité $\|P - [P]_k\|^2$ donnée par cette dernière formule soit suffisamment petite. On obtient ainsi une solution approchée de notre problème d'interpolation, qui a le mérite d'être d'un degré (en général) beaucoup plus petit que la solution exacte.

Les droites de régression: Considérons 10 individus dont on connaît le poids et la taille, par exemple

E_1 (1,75 m , 72 kg)	E_2 (1,80 m , 82 kg)	E_3 (1,85 m , 88 kg)
E_4 (1,70 m ; 74 kg)	E_5 (1,65 m ; 57 kg)	E_6 (1,62 m , 56 kg)
E_7 (1,71 m , 67 kg)	E_8 (1,81 m , 63 kg)	E_9 (1,67 m , 60 kg)
E_{10} (1,77 m , 71kg)		

L'examen rapide d'un tel tableau nous indique que les plus grands sont "généralement" plus lourds que les plus petits. Mais il serait vain d'essayer de construire une fonction donnant le poids quand on connaît la taille; parce qu'il existe des "grands maigres" et des "petits gros". Une telle situation est habituelle en statistiques: les grandeurs étudiées ont un caractère aléatoire, qui rendrait parfaitement illusoire tout polynôme (de degré 10, et même de degré moindre) qu'on pourrait inventer pour rendre compte de la correspondance entre poids et taille. Dans ce cas on va chercher le polynôme de degré 1, qui rend le mieux compte des données, au sens des moindres carrés. C'est

$$[P]_1 = \langle P, \varphi_0 \rangle \varphi_0 + \langle P, \varphi_1 \rangle \varphi_1.$$

Exercice 1 : Avec les données ci-dessus, calculer explicitement φ_0, φ_1 puis $[P]_1$.

Représenter graphiquement les données en portant les tailles en abscisse, et les poids en ordonnées. Représenter la fonction $y = [P]_1(x)$.

Le graphe de cette fonction est appelé la "droite de régression" des poids en fonction des tailles. C'est, pour un statisticien, la fonction qui traduit par une formule mathématique le fait que "les plus grands sont en général les plus lourds".

Complément 2 : Orientation et produit vectoriel

Il reste deux notions, connues en géométrie, et que nous n'avons pas généralisées; ce sont celles de base directe, et de produit vectoriel.

Bases directes: Soit E un $\mathbb{R}$ -espace vectoriel de dimension finie n . Quand va-t-on dire qu'une base de E est directe ?

Nous dirons que deux bases $\{e_1, \dots, e_n\}$ et $\{f_1, \dots, f_n\}$ de E "sont de même orientation", si $\det_{\{e_i\}}(f_1, \dots, f_n) > 0$ (il reviendrait au même de dire que $\det_{\{f_i\}}(e_1, \dots, e_n) > 0$, puisque le produit $\det_{\{e_i\}}(f_1, \dots, f_n) \times \det_{\{f_i\}}(e_1, \dots, e_n)$ vaut 1). Nous partageons alors l'ensemble de toutes les bases en deux classes: ayant choisi une base $B = \{b_1, \dots, b_n\}$, nous mettons dans l'une de ces classes toutes les bases ayant même orientation que B ; dans l'autre, celles qui n'ont pas même orientation que B . Les deux classes de bases ainsi définies ne dépendent pas du choix initial de B . Deux bases d'une même classe sont de même orientation; deux bases qui ne sont pas dans la même classe, ne sont pas de même

orientation. Orienter l'espace E c'est faire jouer un rôle privilégié à l'une de ces deux classes de bases ; ces bases privilégiées sont alors dites directes.

Dans l'espace de la géométrie, on choisit une fois pour toutes celles qui obéissent à la règle des trois doigts (ou du bonhomme d'Ampère). Dans le plan, on choisit les bases $(\vec{i}, \vec{j})$ telles que $\vec{i}$ soit amené sur $\vec{j}$ en tournant d'un angle compris entre 0 et π (dans le sens trigonométrique). On prendra garde qu'un tel choix universel est quelquefois difficile: tracer une base directe sur une feuille de papier calque, retourner la feuille, la base vous paraîtra alors indirecte; autrement dit le sens trigonométrique dépend de la position de l'observateur par rapport au plan considéré.

Le produit mixte: Supposons maintenant que E est orienté et muni d'un produit scalaire. Considérons n vecteurs $\{V_1, \dots, V_n\}$ de E , et une base $\{e_1, \dots, e_n\}$. La quantité $\det_{\{e_i\}}(V_1, \dots, V_n)$ dépend des V_i mais aussi de la base $\{e_1, \dots, e_n\}$ puisque si $\{f_1, \dots, f_n\}$ est une autre base, nous avons:

$$\det_{\{e_i\}}(V_1, \dots, V_n) = \det_{\{e_i\}}(f_1, \dots, f_n) \times \det_{\{f_i\}}(V_1, \dots, V_n)$$

Cependant si l'on décide de n'utiliser que des bases orthonormées directes, cette quantité ne dépend plus que des V_i (puisque, si $\{e_1, \dots, e_n\}$ et $\{f_1, \dots, f_n\}$ sont orthonormées directes, $\det_{\{e_i\}}(f_1, \dots, f_n) = 1$).

Elle est appelée le *produit mixte des vecteurs* $\{V_1, \dots, V_n\}$. Ce produit mixte a évidemment toutes les propriétés du déterminant.

Le produit vectoriel: Ici E est toujours un espace orienté et muni d'un produit scalaire.

Considérons $n-1$ vecteurs $\{V_1, \dots, V_{n-1}\}$ de E . Nous pouvons considérer l'application de E dans $\mathbb{R}$, qui à tout vecteur Y associe $\det_{\{e_i\}}(V_1, \dots, V_{n-1}, Y)$ (où $\{e_1, \dots, e_n\}$ est une base orthonormée directe). D'après les propriétés du déterminant, c'est une application linéaire.

Il existe un unique vecteur X tel que $(\forall Y) \langle X, Y \rangle = \det_{\{e_i\}}(V_1, \dots, V_{n-1}, Y)$. Par définition on l'appelle le *produit vectoriel des vecteurs* $\{V_1, \dots, V_{n-1}\}$.

Démontrons d'abord l'unicité de X ; s'il en existait deux, soit X_1 et X_2 , on aurait $(\forall Y) \langle X_1 - X_2, Y \rangle = 0$. En particulier $\langle X_1 - X_2, X_1 - X_2 \rangle = 0$; donc $X_1 - X_2 = 0$. L'existence de X résulte de la théorie des déterminants: ses coordonnées sont les cofacteurs de la dernière colonne dans le tableau des coordonnées des vecteurs $(V_1, \dots, V_{n-1}, Y)$.

Dans le cas de la géométrie dans l'espace, on retrouve le produit vectoriel classique. Retenons que sur un espace E de dimension n , il n'existe de produit vectoriel que si l'on a choisi une orientation et un produit scalaire; et ce produit vectoriel est alors une fonction de $n-1$ arguments.

Exercices sur le chapitre 4

Exercice 1 : Soit E_3 , l'espace des polynômes de degré au plus 3. On pose

$$\langle f, g \rangle = \int_0^1 f(x) g(x) dx$$

- Montrer que c'est un produit scalaire sur E_3 .
- Trouver une base de E_3 orthonormée pour ce produit scalaire.
- Quel est le sous espace orthogonal au sous espace F engendré par $1-x$ et $2-3x$.

* **Exercice 2 :** Soit E_4 l'espace des polynômes de degré au plus 4. On le munit du produit scalaire

$$\langle f, g \rangle = \sum_{i=0}^4 f(i) g(i)$$

- Trouver la projection du polynôme x^4 , sur l'espace des polynômes de degré 1, puis sur l'espace des polynômes de degré 2, puis sur l'espace des polynômes de degré 3.
- Trouver $f \in E_4$, tel que $f(0) = 0$, $f(1) = 0$, $f(2) = 3$, $f(3) = 6$, $f(4) = 8$.
- Trouver, parmi les polynômes de degré au plus 2, celui qui est le plus proche de f (au sens de la norme associée au produit scalaire ci-dessus).

* **Exercice 3 :** Soit $\{e_1, e_2, e_3\}$ une base orthonormée de l'espace $\vec{E}$ de la géométrie.

- Soit P le plan qui contient les vecteurs $e_1 + 4e_2$ et $2e_1 - e_3$. Trouver un vecteur orthogonal à P .
- Quelle est la projection orthogonale H de $M(X, Y, Z)$ sur P .
- Ecrire la matrice de la symétrie orthogonale par rapport à P .

* **Exercice 4 :** On considère l'espace de la géométrie muni d'une base orthonormée $\{e_1, e_2, e_3\}$.

Montrer que les endomorphismes ϕ qui, dans la base $\{e_1, e_2, e_3\}$ ont les matrices suivantes, sont des symétries. Préciser les éléments caractéristiques de ces symétries. On cherchera les vecteurs X tels que $\phi(X) = X$ et les vecteurs X tels que $\phi(X) = -X$.

$$M_1 = \begin{pmatrix} 0 & 1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 1 \end{pmatrix} \quad M_2 = \begin{pmatrix} \frac{1}{2} & \frac{1}{2} & \frac{1}{\sqrt{2}} \\ \frac{1}{2} & \frac{1}{2} & -\frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} & -\frac{1}{\sqrt{2}} & 0 \end{pmatrix} \quad M_3 = \frac{1}{3} \begin{pmatrix} -1 & 2 & -2 \\ 2 & -1 & -2 \\ -2 & -2 & -1 \end{pmatrix}$$

** **Exercice 5 :** On considère $\mathbb{R}^3$ muni de son "produit scalaire naturel" : $\langle (x, y, z), (x', y', z') \rangle = xx' + yy' + zz'$

Soit ϕ l'endomorphisme de $\mathbb{R}^3$, dont la matrice, dans la base naturelle, est

$$M = \begin{pmatrix} \frac{\sqrt{3}}{2} & \frac{1}{2} & 0 \\ 0 & 0 & 1 \\ \frac{1}{2} & -\frac{\sqrt{3}}{2} & 0 \end{pmatrix}$$

- Vérifier que ϕ est orthogonal. Montrer qu'il existe un vecteur X (non nul) tel que $\phi(X) = X$.
- Soit σ la symétrie par rapport au plan qui contient X et le vecteur $(1, 0, 0)$. Ecrire la matrice de σ , dans la base naturelle.
- Montrer que $\phi \circ \sigma$ et $\sigma \circ \phi$ sont des symétries par rapport à des plans.

Exercice 6: Soit E un espace de dimension 4 muni d'un produit scalaire, et d'une base orthonormée (e_1, e_2, e_3, e_4) . Soit ϕ l'endomorphisme de E dont la matrice, dans la base naturelle, s'écrit

$$M = \frac{1}{2} \begin{pmatrix} 1 & 1 & 1 & 1 \\ 1 & 1 & -1 & -1 \\ 1 & -1 & -1 & 1 \\ 1 & -1 & 1 & -1 \end{pmatrix}$$

a) Vérifier que ϕ est orthogonale.

b) Montrer que ϕ est une symétrie, et préciser les éléments caractéristiques de cette symétrie.

Exercice 7: Même question qu'à l'exercice précédent avec

$$M = \frac{1}{2} \begin{pmatrix} 1 & 1 & 1 & 1 \\ 1 & -1 & -1 & 1 \\ 1 & -1 & 1 & -1 \\ 1 & 1 & -1 & -1 \end{pmatrix}$$

Exercice 8: On considère $\mathbb{R}^3$ muni de son produit scalaire naturel. Soit f l'endomorphisme de $\mathbb{R}^3$ dont la matrice, dans la base naturelle, s'écrit

$$M = \begin{pmatrix} \frac{\sqrt{2}}{2} & \frac{\sqrt{2}}{2} & 0 \\ 0 & 0 & 1 \\ \frac{\sqrt{2}}{2} & -\frac{\sqrt{2}}{2} & 0 \end{pmatrix}$$

a) Montrer qu'il existe un vecteur X (non nul) tel que $f(X) = X$

b) Soit Y un vecteur orthogonal à X et au premier vecteur de la base naturelle, écrire la matrice de la symétrie orthogonale, par rapport à la droite qui porte Y .

c) Ecrire les matrices de $\sigma \circ f$ et $f \circ \sigma$. Montrer que ce sont des rotations.

Exercice 9: On considère dans l'espace un repère orthonormé $(O, \vec{i}, \vec{j}, \vec{k})$. Le plan P a pour équation $2x + 3y - z = 4$.

a) Déterminer la projection $H(\xi, \eta, \zeta)$ du point $M(X, Y, Z)$ sur P .

b) Déterminer le symétrique $M'(X', Y', Z')$ de M par rapport à P .

Exercice 10: Parmi les formules suivantes, quelles sont celles qui définissent un produit scalaire sur $\mathbb{R}^4$?

a) $\langle (x, y, z, t), (x', y', z', t') \rangle = xx' + yy' + zz' + tt' + xt' + x't$.

b) $\langle (x, y, z, t), (x', y', z', t') \rangle = 4xx' + yy' + 2zz' + tt' + xt' + x't + 2yt' + 2y't$.

c) $\langle (x, y, z, t), (x', y', z', t') \rangle = xx' + 3yy' + 2zz' + tt' - xt' - x't + 4yx' + 4y'x$.

Exercice 11: L'espace $\mathbb{R}^4$ est muni du produit scalaire défini par:

$$\langle (x_1, x_2, x_3, x_4), (y_1, y_2, y_3, y_4) \rangle = x_1y_1 + 5x_2y_2 + x_3y_3 + 4x_4y_4 + 2(x_2y_1 + x_1y_2) + (x_4y_3 + x_3y_4).$$

a) Vérifier que ceci définit un produit scalaire.

b) Soit F le sous-espace engendré par $(1; 3; -2; 4)$ et $(3; 1; 2; -1)$. Déterminer une base orthonormée de l'orthogonal de F .

Exercice 12: Soit E l'espace des polynômes. On pose $\langle P, Q \rangle = \int_{-1}^1 P(t) Q(t) dt$.

a) Vérifier qu'on obtient ainsi un produit scalaire sur E

b) Soit $P_n(X) = \frac{d^n}{dx^n} ((1-x^2)^n)$. Vérifier que les polynômes P_n forment une famille orthogonale (pour calculer

$\langle P_n, P_m \rangle$ on pourra faire des intégrations par parties).

c) Déterminer (pour $n = 0, 1, 2, 3, 4$) des α_n tels que la famille $\{\alpha_n P_n\}$ soit orthonormée.

** Exercice 13: Soit E un espace muni d'un produit scalaire. A tout vecteur x on associe l'application φ^x de E dans $\mathbb{R}$, définie par $\varphi^x(y) = \langle x, y \rangle$.

- Montrer que φ^x est un élément du dual E' de E .
- Montrer que l'application Φ de E dans E' qui à x associe φ^x est linéaire.
- Montrer que l'application Φ est injective.
- On suppose E de dimension finie, montrer que l'application Φ est surjective.
- On suppose que $E = \mathcal{C}^0([0,1], \mathbb{R})$, et que le produit scalaire est $\langle f, g \rangle = \int_0^1 f(t)g(t) dt$. Construire une forme

linéaire sur E , qui n'est pas dans l'image de Φ .

* Exercice 14: Dans tout ce qui suit M désigne la matrice:

$$\begin{pmatrix} 1 & 0 & 2 & -3 & -4 \\ -1 & -2 & -4 & 7 & 12 \\ 1 & 1 & 2 & -4 & -6 \\ -1 & 0 & 0 & 1 & 0 \\ 1 & 2 & 2 & -5 & -8 \end{pmatrix}$$

a) Soit A l'endomorphisme de $\mathbb{R}^5$ dont la matrice dans la base naturelle $\{e_1, \dots, e_5\}$ est M . Trouver une base de $\text{Ker } A$, et une base de $\text{Im } A$.

Montrer que la réunion d'une base de $\text{Ker } A$ et d'une base de $\text{Im } A$ est une base de $\mathbb{R}^5$. Comparer $\text{Im } A$ et $\text{Im } A^2$. Comparer $\text{Ker } A$ et $\text{Ker } A^2$.

b) On munit $\mathbb{R}^5$ du produit scalaire (noté $\langle, \rangle$) qui fait de $\{e_1, \dots, e_5\}$ une base orthonormée. Trouver des bases orthonormées B_1 de $\text{Im } A$, et B_2 de $\text{Ker } A$.

Extrait d'une épreuve de juin 1980

** Exercice 15: On se place dans l'espace $\mathcal{C}^0([0,1], \mathbb{R})$, on pose:

$$\begin{aligned} \langle f, g \rangle_1 &= \int_0^1 (1 + \sin \pi t) f(t) g(t) dt & \langle f, g \rangle_2 &= \int_0^1 \sin \pi t f(t) g(t) dt. \\ \langle f, g \rangle_3 &= \int_0^1 \cos \pi t f(t) g(t) dt. & \langle f, g \rangle_4 &= \int_0^{1/2} (1 + \sin \pi t) f(t) g(t) dt. \end{aligned}$$

Parmi ces formules, quelles sont celles qui définissent un produit scalaire ?

* Exercice 16: Soit deux sous-espaces F et G d'un sous espace E muni d'un produit scalaire. Montrer que $(F+G)^\perp = F^\perp \cap G^\perp$. Montrer que, si E est de dimension finie, $(F \cap G)^\perp = F^\perp + G^\perp$.

* Exercice 17: On dit qu'une matrice M à coefficients complexes est unitaire si $\overline{M}^t M = \text{Id}$ (où $\overline{M}$ est la matrice obtenue en remplaçant les coefficients de M par leurs conjugués). Montrer que le produit de deux matrices unitaires est unitaire. Montrer que l'inverse d'une matrice unitaire est unitaire.

** Exercice 18: Soit F et G deux sous-espaces vectoriels d'un espace E . On suppose que $G \cap F = \{0\}$. Montrer qu'il existe un produit scalaire sur E tel que tout vecteur de F soit orthogonal à tout vecteur de G .

Soit maintenant trois sous-espaces F , G et H de E . A quelle condition existe-t-il un produit scalaire sur E tel que trois vecteurs quelconques appartenant respectivement à F à G et à H soient 2 à 2 orthogonaux.

Exercice 19: Dans $\mathbb{R}^5$ (muni du produit scalaire qui fait de la base naturelle une base orthonormée) on donne le sous-espace F engendré par:

$$V_1 = (1; 4; 2; 5; 7)$$

$$V_2 = (2; 1; -1; 3; -4)$$

$$V_3 = (3; 2; 1; 3; 5)$$

$$V_4 = (0; 3; 0; 5; -2)$$

Soit $W = (2; 5; 9; 0; 1)$, déterminer la projection orthogonale de W sur F et sur $F^\perp$.

Exercice 20: Dans $\mathbb{R}^5$ (muni du produit scalaire qui fait de la base naturelle une base orthonormée) on donne le sous-espace F engendré par:

$$V_1 = (1; 4; 5; 5; 7)$$

$$V_2 = (2; 1; 1; 3; -4)$$

$$V_3 = (5; 2; 1; 3; 1)$$

$$V_4 = (-1; 7; 10; 10; 9)$$

Déterminer une base orthonormée de $F^\perp$.

Exercice 21: Soit E de dimension n muni d'un produit scalaire et d'une base orthonormée $(e_1, \dots, e_n)$. Montrer que quels que soient les vecteurs $(V_1, \dots, V_n)$:

$$|\det_{(e_1, \dots, e_n)}(V_1, \dots, V_n)| \leq \prod_{i=1}^n \|V_i\|$$

On commencera par montrer la formule lorsque les vecteurs V_i sont deux à deux orthogonaux, puis on regardera comment les deux termes évoluent lorsqu'on applique à $(V_1, \dots, V_n)$ le procédé de Schmidt.

Dans quels cas a-t-on égalité ?

Exercice 22: (méthode des moindres carrés) On cherche une fonction f telle que :

$$f(0) = 1,5$$

$$f(1) = 2$$

$$f(2) = 2,5$$

$$f(3) = 2$$

$$f(4) = 1$$

$$f(5) = 3$$

1) Trouver - par la méthode des moindres carrés - les polynômes de degré 1, 2 et 3 qui approchent le mieux ces données.

2) L'un de ces polynômes vous semble-t-il approcher valablement les données ? Si non, chercher le polynôme de degré 4 qui approche le mieux les données

Faire une représentation graphique.

Vecteurs propres et valeurs propres

Les difficultés que l'on rencontre dans la résolution d'un problème d'algèbre linéaire, dépendent souvent de la base dans laquelle on a choisi de travailler. Le choix d'une base bien adaptée permet souvent de simplifier le problème posé. Ainsi (cf: §1 du chapitre 2) la méthode du pivot consiste essentiellement, étant donné p vecteurs, à trouver une base dans laquelle leurs coordonnées sont 'simples'. Ici nous allons essayer de trouver une base 'bien adaptée' à un endomorphisme donné.

Dans ce chapitre, étant donnée une application linéaire f de F dans lui-même (F est de dimension finie, et f est donnée par sa matrice dans une certaine base), nous allons chercher une nouvelle base de F , dans laquelle la matrice de f est 'la plus simple possible'. Une telle base est indispensable si l'on veut établir une formule générale donnant la matrice de f^n en fonction de n (problème dont on verra l'importance aux chapitres 19 et 20). Les bases idéales sont celles où la matrice de f est diagonale. Il n'en existe pas toujours; nous énoncerons quelques conditions (nécessaires ou suffisantes) pour qu'il en existe. Mais nous ne traiterons le problème en entier qu'au chapitre 19.

La recherche de telles bases commence par la recherche des vecteurs V tels que $f(V)$ soit colinéaire à V ; ce sont les vecteurs propres. Nous allons décrire un algorithme permettant de les déterminer. Au §5, on trouvera une étude un peu plus approfondie de quelques cas particuliers.

§ 1 : Définitions générales

Soit f un endomorphisme d'un espace vectoriel E (c'est à dire une application linéaire de E dans E), on dit qu'un vecteur V de E est propre pour f , si V est non nul, et si $f(V)$ est colinéaire à V , c'est-à-dire s'il existe λ tel que $f(V) = \lambda V$. Le scalaire λ est appelé la valeur propre associée au vecteur propre V .

Exercice a: Soit E le plan muni d'une origine O . Soit f la symétrie orthogonale par rapport à une droite D passant par O . C'est une application linéaire de E dans E . Quels sont ses vecteurs propres ?

Exercice b: Dans le même espace E , on considère la rotation de centre O et d'angle $\pi/4$. C'est une application linéaire de E dans E . Quels sont ses vecteurs propres ?

Exercice c: Soit $F: \mathcal{C}^\infty(\mathbb{R}, \mathbb{R}) \rightarrow \mathcal{C}^\infty(\mathbb{R}, \mathbb{R})$ qui à f associe sa dérivée f' . Quels sont les vecteurs propres de F ?

§2 : Recherche des vecteurs propres et valeurs propres en dimension finie

Supposons maintenant E de dimension finie et soit M la matrice de f dans une base $\{e_i\}$ de E . Alors les conditions suivantes sont équivalentes:

$\alpha)$ $f - \lambda \text{id}_E$ n'est pas inversible.

$\beta)$ $M - \lambda I_n$ n'est pas inversible.

$\gamma)$ $\text{Dét}(M - \lambda I_n) = 0$.

$\delta)$ Il existe un vecteur propre V de f , de valeur propre λ .

Pour rechercher les valeurs propres et les vecteurs propres de f , nous allons donc:

* D'abord rechercher les scalaires λ tels que $\det(M - \lambda I_n)$ soit nul. Ce sont les valeurs propres de f .

* Ensuite, pour chacun de ces λ , rechercher les vecteurs V (non nuls) de $\text{Ker}(f - \lambda \text{Id})$. Ce sont les vecteurs propres de f associés à la valeur propre λ .

Exercice d: Soit $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$ dont la matrice, dans la base naturelle est $M = \begin{pmatrix} 1 & 2 \\ 2 & 1 \end{pmatrix}$

Calculer $\det(M - \lambda I_2)$. Résoudre l'équation $\det(M - \lambda I_2) = 0$. Puis déterminer tous les vecteurs propres de f .

Exercice e: Reprendre l'exercice précédent, en remplaçant M par $M_1 = \begin{pmatrix} a & -b \\ b & a \end{pmatrix}$ ($b \neq 0$).

§3 : Le polynôme caractéristique

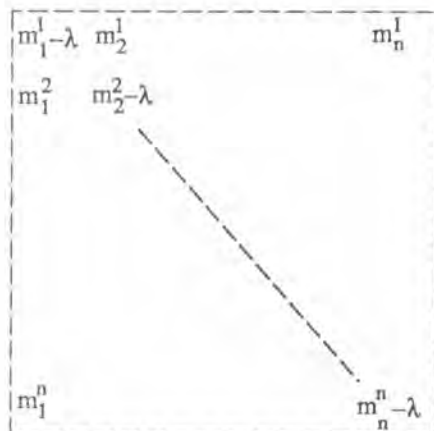
La donnée est toujours $f: E \rightarrow E$, où E est de dimension n , et la matrice M de f dans une base $\{e_1, \dots, e_n\}$ de E .

$\det(M - \lambda I_n)$ est un polynôme de degré n en λ .

La matrice $M - \lambda I_n$ est obtenue en ôtant λ à tous les termes de la diagonale de M . Son déterminant est une expression polynomiale des coefficients (cf: chapitre 2, § 3). Chaque terme du déterminant est le produit de n coefficients de la matrice $M - \lambda I_n$; c'est à dire le produit de n polynômes de degré 0 ou 1 en λ . D'où le résultat.

Nous pouvons observer que le seul terme du déterminant qui est produit de n polynômes de degré 1 en λ , est le produit des coefficients diagonaux. Les autres contiennent au plus $n-2$ polynômes de degré 1 en λ . Il en résulte que $(m_1^1 - \lambda)(m_2^2 - \lambda) \dots (m_n^n - \lambda)$ et $\det(M - \lambda I_n)$ ont mêmes termes en λ^n et en λ^{n-1} . Donc

$$\det(M - \lambda I_n) = (-\lambda)^n + \left(\sum_i m_i^i \right) (-\lambda)^{n-1} + \dots$$



La somme $\sum_i m_i^i$, c'est à dire la somme des coefficients diagonaux de la matrice, s'appelle la

trace de la matrice M . Le terme constant est $\det M (= \det(M - 0I_n))$ (Retenir en particulier que $\det M = 0$ équivaut à 0 est valeur propre). Donc

$$\det(M - \lambda I_n) = (-\lambda)^n + \left(\sum_i m_i^i \right) (-\lambda)^{n-1} + \dots + \det M$$

Si la dimension est 2, cette formule se réduit à: $\det(M_2 - \lambda \text{Id}) = \lambda^2 - \lambda \text{Tr } M_2 + \det M_2$.

Faisons un changement de base, nous aurons une matrice $M_1 = A M A^{-1}$ (ou A est la matrice inversible qui exprime le changement de base). Et

$$\begin{aligned} \det(M_1 - \lambda I_n) &= \det(A M A^{-1} - \lambda A I_n A^{-1}) \\ &= \det[A(M - \lambda I_n) A^{-1}] \\ &= \det A \det(M - \lambda I_n) \det A^{-1} \\ &= \det(M - \lambda I_n) \end{aligned}$$

Nous avons ainsi démontré que le polynôme $\det(M - \lambda I_n)$ ne dépend que de f . Il est le même quelle que soit la base dans laquelle on fait les calculs. On l'appelle le polynôme caractéristique de f .

Exercice f: Soit $f: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ dont la matrice, dans la base naturelle est $M = \begin{pmatrix} 1 & 3 & 1 \\ 2 & 2 & 1 \\ 3 & 1 & 1 \end{pmatrix}$

- 1) Calculer le polynôme caractéristique de f . Montrer que ses racines sont 0, -1 et 5.
- 2) Déterminer les vecteurs propres de f .

Exercice g: Soit $f: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ dont la matrice dans la base naturelle est $M = \begin{pmatrix} 1 & 3 & 1 \\ 0 & 4 & 1 \\ 3 & 0 & 2 \end{pmatrix}$

- 1) Calculer le polynôme caractéristique de f . Montrer que ses racines sont 5 et 1 (double)
- 2) Déterminer $\text{Ker}(f - 5 \text{Id})$ et $\text{Ker}(f - \text{Id})$ (qui sont tous deux de dimension 1)

Exercice h: Soit $f: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ dont la matrice, dans la base naturelle est $M = \begin{pmatrix} 1 & 3 & 1 \\ 2 & 1 & 2 \\ 6 & -2 & 1 \end{pmatrix}$

- 1) Calculer le polynôme caractéristique de f . Montrer que sa seule racine réelle est 5.
- 2) Déterminer tous les vecteurs propres de f .

Exercice i: On donne $f: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ dont la matrice, dans la base naturelle est $M = \begin{pmatrix} 2 & 1 & -1 \\ 1 & 2 & -1 \\ 1 & 1 & 0 \end{pmatrix}$

- a) Calculer le polynôme caractéristique de f . Montrer que ses racines sont 2 et 1 (double).
- b) Montrer qu'à la valeur propre 2 est associée une droite de vecteurs propres, et à la valeur propre 1 un plan de vecteurs propres.

§ 4 : Diagonalisation

Supposons qu'on ait pu trouver n ($= \dim E$) vecteurs propres $\{e_1, \dots, e_n\}$ linéairement indépendants. Ils forment une base. Quel que soit i , $f(e_i) = \lambda_i e_i$; donc la matrice de f dans cette base s'écrit

$$M = \begin{pmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_n \end{pmatrix}$$

Elle est diagonale, et on trouve sur la diagonale les valeurs propres.

Diagonaliser un endomorphisme $f: E \rightarrow E$, c'est chercher une base formée de vecteurs propres. Il n'en existe pas toujours (cf: exercices g et h ci-dessus); nous allons chercher des conditions pour qu'il en existe.

Sous-espaces propres:

A toute valeur propre t associons $\text{Ker}(f - t \text{Id})$ que nous appellerons sous-espace propre associé à t . Alors, en notant r l'ordre de multiplicité de t , comme racine du polynôme caractéristique, nous avons les inégalités:

$$1 \leq \dim(\text{Ker}(f - t \text{Id})) \leq r$$

En effet: d'une part $\text{Ker}(f - t \text{Id})$ n'est pas réduit à $\{0\}$ (puisque t est une valeur propre), donc il est de dimension au moins 1. D'autre part si $\dim(f - t \text{Id})$ est de dimension k , nous pouvons en choisir une base $\{e_1, \dots, e_k\}$, et la compléter en une base $\{e_1, \dots, e_n\}$ de E . Dans cette base la matrice de f s'écrit:

$$M_1 = \left(\begin{array}{c|c} t \times \text{Id}_k & a \\ \hline 0 & b \end{array} \right)$$

Et $\det(M_1 - \lambda \text{Id}) = (t - \lambda)^k Q(\lambda)$; où $Q(\lambda) = \det(b - \lambda \text{Id}_{n-k})$. Ce qui prouve que l'ordre de multiplicité de t comme racine du polynôme caractéristique est au moins égal à k .

L'espace engendré par les vecteurs propres.

Soit $\lambda_1, \dots, \lambda_p$ les valeurs propres de f , et soit $F_1 = \text{Ker}(f - \lambda_1 \text{Id})$, $F_2, \dots, F_p$ les sous-espaces propres. On notera F le sous-espace de E engendré par $F_1 \cup F_2 \dots \cup F_p$.

* Alors tout vecteur e de F s'écrit $e = \sum_i e_i$, où $(\forall i) e_i$ est dans F_i .

* Et cette décomposition est unique. En effet si $e = \sum_i e_i = \sum_i e'_i$ appliquons aux deux membres l'endomorphisme $g_1 = (f - \lambda_2 \text{Id})(f - \lambda_3 \text{Id}) \dots (f - \lambda_n \text{Id})$; nous avons:

$$g_1(e_i) = (\lambda_i - \lambda_2)(\lambda_i - \lambda_3) \dots (\lambda_i - \lambda_n) e_i$$

et $g_1(e'_i) = (\lambda_i - \lambda_2)(\lambda_i - \lambda_3) \dots (\lambda_i - \lambda_n) e'_i$

donc : $g_1(e) = (\lambda_1 - \lambda_2)(\lambda_1 - \lambda_3) \dots (\lambda_1 - \lambda_n) e_1 = (\lambda_1 - \lambda_2)(\lambda_1 - \lambda_3) \dots (\lambda_1 - \lambda_n) e'_1$

D'où $e_1 = e'_1$. Et de même en appliquant $g_2 = (f - \lambda_1 \text{Id})(f - \lambda_3 \text{Id}) \dots (f - \lambda_n \text{Id})$, nous aurons $e_2 = e'_2$. Etc.

Nous avons ainsi montré (cf: chap 1, §6) que F est la somme directe des F_i . Donc en mettant cote à cote une base de F_1 , une base de F_2 , ... et une base de F_p , nous obtenons une base de F formée de vecteurs propres. Et F est de dimension $\sum_i \dim F_i$.

A quelles conditions $f: E \rightarrow E$ est elle diagonalisable ?

Il est clair que F est le plus grand sous-espace de E possédant une base formée de vecteurs propres. Donc f est diagonalisable si et seulement si $F = E$; c'est à dire si $\sum_i \dim F_i = \dim E (= n)$.

Et pour cela il est nécessaire et suffisant que les deux conditions suivantes soient vérifiées:

a) La somme des multiplicités des valeurs propres de f est n . Ceci signifie que le polynôme caractéristique s'écrit comme un produit de facteurs de degré 1.

Si le corps de base est $\mathbb{C}$ cette condition est toujours vérifiée. Si le corps de base est $\mathbb{R}$, il se peut qu'elle ne le soit pas (cf: Exercice h).

b) Pour chaque valeur propre λ_i de multiplicité r : $\dim \text{Ker}(f - \lambda_i \text{Id}) = r$.

Cette condition peut être vérifiée (cf : Exercice i) ; elle peut ne pas l'être (cf : Exercice g).

En particulier :

Si toutes les valeurs propres de f sont (des racines) simples (du polynôme caractéristique), alors f est diagonalisable.

Si l'on travaille sur le corps des réels et si le polynôme caractéristique de f a des racines non réelles, alors f n'est pas diagonalisable.

§ 5: Diagonalisation des matrices orthogonales et des matrices symétriques.

Valeurs propres complexes des endomorphismes réels: Considérons un endomorphisme f d'un espace vectoriel réel E de dimension finie n , et supposons que son polynôme caractéristique ait une racine λ non réelle ($\bar{\lambda}$ est alors aussi une racine). Si M est la matrice de f dans une base $\{e_i\}$, la matrice $M - \lambda I$ est une matrice complexe non inversible. Donc $(M - \bar{\lambda} I)(M - \lambda I)$ est non inversible. Mais

$$(M - \bar{\lambda}I)(M - \lambda I) = M^2 - (\lambda + \bar{\lambda})M + \lambda\bar{\lambda}I$$

est une matrice à coefficients réels; c'est, dans la base (e_i) , la matrice de

$$g = f \circ f - (\lambda + \bar{\lambda})f + \lambda\bar{\lambda}\text{Id}_E : E \rightarrow E$$

Cet endomorphisme g est donc non inversible.

Considérons un vecteur (non nul) V de $\text{Ker } g$, et notons F le sous-espace vectoriel engendré par V et $f(V)$.

1) $f(F) \subset F$ car $f(aV + bf(V)) = af(V) + bf(f(V)) = af(V) + b[(\lambda + \bar{\lambda})f(V) - \lambda\bar{\lambda}V]$. On dit que F est stable par f .

2) F est de dimension 2. En effet si F était de dimension 1, V en serait une base; et il existerait α (réel), tel que $f(V) = \alpha V$; on aurait donc $g(V) = [\alpha^2 - (\lambda + \bar{\lambda})\alpha + \lambda\bar{\lambda}]V = 0$. Donc $\alpha^2 - (\lambda + \bar{\lambda})\alpha + \lambda\bar{\lambda} = 0$, ce qui est impossible car α est réel, λ est non réel et $\alpha^2 - (\lambda + \bar{\lambda})\alpha + \lambda\bar{\lambda} = (\alpha - \lambda)(\alpha - \bar{\lambda})$.

Étudions maintenant l'endomorphisme $f|_F : F \rightarrow F$. Dans la base $(V, f(V))$ sa matrice s'écrit :

$$M_1 = \begin{pmatrix} 0 & -\lambda\bar{\lambda} \\ 1 & \lambda + \bar{\lambda} \end{pmatrix}$$

Mais si l'on pose $\lambda = x + iy$ ($y \neq 0$), les vecteurs V et $\frac{1}{y}(-xV + f(V))$ forment aussi une base de F , et dans cette base $f|_F$ a pour matrice (vérification facile): $M_2 = \begin{pmatrix} x & -y \\ y & x \end{pmatrix}$

On vérifie facilement (sur l'une ou l'autre de ces 2 matrices) que :

$$\det(f|_F - z \text{Id}_F) = z^2 - 2(\text{Re } \lambda)z + |\lambda|^2$$

Endomorphismes semi-simples: Nous dirons qu'un endomorphisme $f : E \rightarrow E$ est semi-simple si pour tout sous-espace stable F (ie : $f(F) \subset F$), il existe un supplémentaire G de F qui est stable (ie : $f(G) \subset G$).

Nous verrons ci-dessous des exemples d'endomorphismes semi-simples: les endomorphismes orthogonaux ou symétriques.

Notons que si $E' \subset E$ est stable par f , et si $f : E \rightarrow E$ est semi-simple, alors $f|_{E'} : E' \rightarrow E'$ est semi-simple. En effet si $F \subset E'$ est stable, prenons un supplémentaire G de F dans E , qui est stable; alors $G \cap E'$ est stable et est un supplémentaire de F dans E' .

Si un endomorphisme f d'un espace E de dimension finie est semi-simple, et si les racines de son polynôme caractéristique sont dans le corps de base (en particulier si E est un espace complexe), alors f est diagonalisable.

Nous raisonnerons par récurrence sur la dimension n de E . Pour $n = 1$, il n'y a rien à démontrer. Supposons que notre énoncé soit vrai pour les espaces de dimension $n-1$, et considérons E de dimension n . Choisissons un vecteur propre V de valeur propre λ_0 (il en existe), il engendre un sous-espace stable F . Soit G un supplémentaire stable.

Formons une base de E en mettant côte à côte une base de F et une base de G . Dans cette base, F a une matrice de la forme:

$$M = \left(\begin{array}{c|cc} \lambda_0 & 0 & 0 \\ \hline 0 & & \\ 0 & N & \end{array} \right)$$

Le polynôme caractéristique de f s'écrit donc $\det(N - \lambda I_{n-1}) \times (\lambda_0 - \lambda)$. Il a toutes ses racines dans le corps de base, donc $\det(N - \lambda I_{n-1})$ (qui est le polynôme caractéristique de $f|_G$) a aussi toutes ses racines dans le corps de base. Par ailleurs $f|_G$ est semi-simple (puisque f l'est); donc on peut appliquer l'hypothèse de récurrence: il existe une base $\{e_1, \dots, e_{n-1}\}$ de G formée de vecteurs propres de $f|_G$. Alors $\{v, e_1, \dots, e_{n-1}\}$ est une base de E formée de vecteurs propres de f .

Si un endomorphisme f d'un espace réel E est semi-simple, alors il existe une décomposition de E en une somme directe $E_1 \oplus \dots \oplus E_k$, où les E_i sont de dimension 1 ou 2, et stables par f .

La démonstration se fait par récurrence sur la dimension de E . Elle est analogue à la précédente. La seule différence est qu'on n'est pas certain de pouvoir construire un sous-espace stable F de dimension 1. Ceci n'est possible que si le polynôme caractéristique a au moins une racine réelle. Si toutes ses racines sont non réelles, on peut (cf: ci-dessus) construire un sous-espace stable de dimension 2.

Endomorphismes orthogonaux et unitaires: Soit E un espace réel euclidien de dimension finie, et soit f un endomorphisme orthogonal de E , alors f est semi-simple (car si F est stable, $F^\perp$ l'est aussi); mais il n'y a aucune raison pour que les racines du polynôme caractéristique soient réelles.

On a donc une décomposition de E en $E_1 \oplus \dots \oplus E_k$, tous de dimension 1 ou 2 et stables par f . Notons que $(\forall i) (\forall j) (\forall x \text{ dans } E_i) \text{ et } (\forall y \text{ dans } E_j) \text{ on a } \langle x, y \rangle = 0$. Choisissons des bases orthonormées de ces E_i , la réunion de ces bases est une base orthonormée de E ; et dans cette base la matrice de f est formée (de coefficients nuls et) de blocs 1×1 (i.e. de coefficients diagonaux) et de blocs 2×2 qui sont des matrices de rotation.

Notons que les racines du polynôme caractéristique de f , sont les racines des polynômes caractéristiques des endomorphismes $f|_{E_i} : E_i \rightarrow E_i$. Elles sont donc de module 1. (Laissons au lecteur le soin de calculer le polynôme caractéristique d'une matrice orthogonale de dimension 1 ou 2).

Dans le cas d'un espace complexe hermitien, considérons un endomorphisme unitaire $f: E \rightarrow E$. Il est semi-simple (si F est stable, $F^\perp$ l'est aussi). Il est donc diagonalisable. Notons que les valeurs propres sont de module 1 (car $\|x\| = \|f(x)\| = |\lambda| \|x\|$ si x est un vecteur propre, de valeur propre λ).

Notons aussi que si x et y sont des vecteurs propres correspondant à des valeurs propres distinctes λ et μ , alors x et y sont des vecteurs orthogonaux; comme dans le cas des endomorphismes orthogonaux. On peut donc choisir une base orthonormée dans laquelle la matrice de f est diagonale. Autrement dit si U est une matrice unitaire, il existe V unitaire telle que $V^{-1}UV$ soit diagonale.

Exercice j: Justifier les affirmations concernant les endomorphismes orthogonaux de dimension 3 et 4, qui ont été faites au chapitre 4 (fin du §5).

Diagonalisation des matrices symétriques réelles: Soit M une matrice réelle $n \times n$ qui est symétrique (ie: $M^t = M$). Soit E un espace euclidien de dimension n , muni d'une base $\{e_1, \dots, e_n\}$ orthonormée. On note f l'endomorphisme dont la matrice, dans cette base, est M .

Montrons que, dans toute base orthonormée de E , la matrice de f est symétrique. Elle s'écrit $M_1 = A M A^{-1}$, où A est orthogonale; donc:

$$(M_1)^t = (A M A^{-1})^t = (A M^t A^t)^t = A M A^{-1} = M_1$$

Montrons que f est semi-simple. Si F est stable choisissons une base orthonormée $\{e_1, \dots, e_n\}$ dont les k premiers vecteurs forment une base de F . La matrice de f dans cette base est de la forme

$$M' = \left(\begin{array}{c|c} m_1 & m_3=0 \\ \hline m_2=0 & m_4 \end{array} \right) \quad (\text{Nota : } m_2 = 0 \text{ car } F \text{ est stable ; } m_3 = 0 \text{ car } M' \text{ est symétrique})$$

Et par conséquent $F^\perp$ qui est engendré par $\{\varepsilon_{k+1}, \dots, \varepsilon_n\}$ est stable.

Enfin $\det(M - \lambda I)$ a toutes ses racines réelles. Si on avait une racine non réelle α , on en déduirait un sous-espace stable F de dimension 2, tel que le polynôme caractéristique de $f|_F$ soit $\lambda^2 - (\alpha + \bar{\alpha})\lambda + \alpha\bar{\alpha}$. En choisissant une base orthonormée $\{\varphi_1, \dots, \varphi_n\}$ dont les 2 premiers vecteurs sont dans F , on obtiendrait une matrice symétrique de la forme:

$$\begin{pmatrix} a & b & 0 & \dots & 0 \\ b & c & 0 & \dots & 0 \\ 0 & 0 & & & \\ & & N & & \\ 0 & 0 & & & \end{pmatrix}$$

Et $\det \begin{pmatrix} a-\lambda & b \\ b & c-\lambda \end{pmatrix} = (a-\lambda)(c-\lambda) - b^2 = \det(f|_F - \lambda \text{Id}_F) = \lambda^2 - (\alpha + \bar{\alpha})\lambda + \alpha\bar{\alpha}$. Ce qui est absurde car $(a-\lambda)(c-\lambda) - b^2 = 0$ a 2 racines réelles (Vérification facile).

Dans ces conditions on peut affirmer que f est diagonalisable. Les vecteurs propres correspondant à des valeurs propres distinctes sont orthogonaux. On peut donc trouver une base orthonormée qui diagonalise f . Autrement dit, *pour toute matrice symétrique réelle M , il existe une matrice orthogonale H telle que HMH^{-1} soit diagonale.*

Cas complexe : dans le cas des *matrices complexes* on remplace symétrique par hermitienne (ie : $M^t = \bar{M}$). On montre alors, de façon analogue, qu'il existe U unitaire telle que UMU^{-1} soit diagonale à coefficients réels.

Exercices sur le chapitre 5

- * Exercice 1: Trouver les vecteurs propres et les valeurs propres des endomorphismes de $\mathbb{R}^3$ dont les matrices, dans la base naturelle s'écrivent :

$$M_a = \begin{pmatrix} 1 & 0 & \sqrt{3} \\ 0 & 1 & 1 \\ \sqrt{3} & 1 & 1 \end{pmatrix}$$

$$M_b = \begin{pmatrix} 1 & 1 & 0 \\ 0 & 1 & 0 \\ 0 & 1 & 1 \end{pmatrix}$$

$$M_c = \begin{pmatrix} 1 & 1 & 0 \\ 1 & 1 & -1 \\ 0 & 1 & 1 \end{pmatrix}$$

$$M_d = \begin{pmatrix} 1 & 1 & 1 \\ 1 & 3 & 1 \\ 1 & 1 & 1 \end{pmatrix}$$

- ** Exercice 2: Trouver une matrice A telle que AMA^{-1} soit diagonale.

$$M_1 = \begin{pmatrix} 2 & 5 & -2 \\ 2 & 11 & -4 \\ 6 & 30 & -11 \end{pmatrix}$$

$$M_2 = \begin{pmatrix} 2 & 1 & 2 & 1 \\ 1 & 1 & 3 & 1 \\ 2 & 3 & 0 & 1 \\ 1 & 1 & 1 & 3 \end{pmatrix}$$

- * Exercice 3: Les matrices suivantes sont-elles diagonalisables ?

$$M = \begin{pmatrix} 3 & 1 & -1 \\ 1 & 5 & -1 \\ 0 & 2 & 2 \end{pmatrix}$$

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

Extrait d'un partiel de 1980

- * Exercice 4: On considère l'endomorphisme φ de $\mathbb{R}^3$ dont la matrice, dans la base naturelle, est

$$M = \begin{pmatrix} -1 & 2 & 2 \\ -1 & 2 & 1 \\ -1 & 1 & 2 \end{pmatrix}$$

- a) Existe-t-il une base dans laquelle la matrice de φ est diagonale ?
b) Existe-t-il une base dans laquelle la matrice de $\varphi \circ \varphi$ est diagonale ?

- * Exercice 5: Calculer M^n ($n \in \mathbb{Z}$), où $M = \begin{pmatrix} 1 & 0 & 2 \\ 1 & 2 & 0 \\ 0 & 1 & 2 \end{pmatrix}$

- * Exercice 6: Les matrices suivantes sont-elles diagonalisables (sur $\mathbb{R}$)

$$\begin{pmatrix} 1 & 1 & 1 \\ 0 & 2 & 1 \\ 1 & 0 & 2 \end{pmatrix} \quad \begin{pmatrix} 0 & 1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 1 \end{pmatrix} \quad \begin{pmatrix} 0 & 0 & 1 \\ 1 & 0 & 0 \\ 0 & 1 & 0 \end{pmatrix}$$

- * Exercice 7:

Existe-t-il une matrice inversible A telle que $A^{-1} \begin{pmatrix} 3 & -1 & -1 \\ 1 & 1 & -1 \\ 2 & -1 & 0 \end{pmatrix} A$ soit diagonale ?

- ** Exercice 8: Existe-t-il une matrice inversible A telle que $AMA^{-1} = N$, sachant que

$$M = \begin{pmatrix} 0 & -1 & -1 \\ -1 & 0 & -1 \\ -1 & -1 & 0 \end{pmatrix}$$

$$N = \frac{1}{2} \begin{pmatrix} -1 & -2 & 1 \\ -3 & -2 & -1 \\ 3 & 4 & 3 \end{pmatrix}$$

- * Exercice 9: Ecrire une matrice non diagonalisable dont la trace est 7, et le déterminant 12.

- * **Exercice 10:** Soit E l'espace des polynômes, et soit ϕ l'endomorphisme de E qui à tout polynôme associe sa dérivée. Quels sont les valeurs propres et les vecteurs propres de ϕ ?
- * **Exercice 11:** Soit E l'espace des polynômes, et soit ϕ l'endomorphisme de E qui à tout polynôme $P(X)$ associe $P(2X)$. Quels sont les valeurs propres et les vecteurs propres de ϕ ?
- ** **Exercice 12:** a) Ecrire une matrice 3×3 dont le polynôme caractéristique est $(1-\lambda)^3$ et qui ne soit pas diagonalisable
b) Ecrire une matrice 3×3 dont le polynôme caractéristique est $(1-\lambda)^2(2-\lambda)$ et qui ne soit pas diagonalisable
- * **Exercice 13:** Trouver les valeurs propres et les vecteurs propres de l'endomorphisme de $\mathbb{R}^4$ dont la matrice, dans la base naturelle, est:

$$M = \begin{pmatrix} 1 & 2 & 1 & 3 \\ 0 & 2 & 4 & -1 \\ 0 & 0 & 2 & 5 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

- ** **Exercice 14:** Soit A un endomorphisme de $\mathbb{R}^3$ dont les valeurs propres sont $1, -1, 3$. Quel est le rang de $\text{Id} - A^2$? Quel est le rang de $\text{Id} - 2A + A^2$? Quel est le rang de $3 \text{Id} - A - 3A^2 + A^3$?
- ** **Exercice 15:** On considère un endomorphisme f de $\mathbb{R}^3$ tel que $f^2 - 2\text{Id} = 0$. Montrer que f est diagonalisable.
- ** **Exercice 16:** Soit f l'endomorphisme de $\mathbb{C}^3$, dont la matrice, dans la base naturelle, s'écrit : $M = \begin{pmatrix} 0 & 2 & 0 \\ 0 & 1 & 1 \\ 1 & 0 & 1 \end{pmatrix}$
Trouver une base de $\mathbb{C}^3$ dans laquelle la matrice de f s'écrit : $P = \begin{pmatrix} 2 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & -1 & 0 \end{pmatrix}$
- * **Exercice 17:** Démontrer que la transposée M^t d'une matrice M est diagonalisable, si et seulement si M est elle-même diagonalisable.
- * **Exercice 18:** On suppose que λ est valeur propre de M , montrer que λ^2 est valeur propre de M^2 .
Montrer que $\lambda^3 + 2\lambda - 1$ est valeur propre de $M^3 + 2M - \text{Id}$.
- * **Exercice 19:** Un endomorphisme A de $\mathbb{R}^n$ vérifie $A^2 + A + \text{Id} = 0$, démontrer que son polynôme caractéristique n'a aucune racine réelle.
- * **Exercice 20:** Soit $\vec{E}$ l'espace des vecteurs de la géométrie, muni d'une base orthonormée $(\vec{i}, \vec{j}, \vec{k})$.
a) Soit A un endomorphisme de E , dont la matrice M , dans la base $(\vec{i}, \vec{j}, \vec{k})$ est symétrique et orthogonale. En utilisant le fait qu'elle est diagonalisable, montrer que A est soit l'identité, soit moins l'identité, soit une symétrie orthogonale par rapport à un plan, soit une symétrie orthogonale par rapport à une droite.
b) Réciproquement soit B une symétrie orthogonale par rapport à un plan P (ou à une droite D), écrire sa matrice dans une base orthonormée dont les deux premiers vecteurs sont dans P . En déduire que sa matrice dans la base $(\vec{i}, \vec{j}, \vec{k})$ est orthogonale et symétrique.

c) Application: identifier les endomorphismes dont les matrices s'écrivent:

$$M_1 = \frac{1}{3} \begin{pmatrix} 1 & 2 & 2 \\ 2 & 1 & -2 \\ 2 & -2 & 1 \end{pmatrix}$$

$$M_2 = \frac{1}{3} \begin{pmatrix} 1 & 2 & 2 \\ 2 & -2 & 1 \\ 2 & 1 & -2 \end{pmatrix}$$

** Exercice 21: Trouver une matrice orthogonale H telle que $H \cdot \begin{pmatrix} 1 & 0 & 1 & 1 \\ 0 & 0 & 1 & 2 \\ 1 & 1 & 0 & 1 \\ 1 & 2 & 1 & -1 \end{pmatrix} \cdot H^{-1}$ soit diagonale.

* Exercice 22: Trouver une matrice orthogonale H telle que $H \cdot \begin{pmatrix} 1 & 0 & 1 & 0 \\ 0 & 1 & 0 & 1 \\ 1 & 0 & 1 & 0 \\ 0 & 1 & 0 & 1 \end{pmatrix} \cdot H^{-1}$ soit diagonale.

** Exercice 23: On considère $\mathbb{R}^4$ muni de son produit scalaire naturel. Soit ϕ l'endomorphisme dont la matrice, dans la base naturelle s'écrit :

$$M = \begin{pmatrix} 1/2 & 1/2 & \frac{1}{\sqrt{2}} & 0 \\ 1/2 & -1/2 & 0 & \frac{1}{\sqrt{2}} \\ 1/2 & 1/2 & -\frac{1}{\sqrt{2}} & 0 \\ 1/2 & -1/2 & 0 & -\frac{1}{\sqrt{2}} \end{pmatrix}$$

Calculer $\det(M - I_4)$ et $\det(M + I_4)$; puis trouver 3 sous espaces F_1, F_2, F_3 de $\mathbb{R}^4$, de dimension 1 ou 2 et tels que

$$\alpha) \mathbb{R}^4 = F_1 \oplus F_2 \oplus F_3$$

$$\beta) \forall i \quad \phi(F_i) \subset F_i$$

** Exercice 24: Même question qu'à l'exercice 23, avec:

$$M = \begin{pmatrix} 1/2 & 1/2 & -1/\sqrt{2} & 0 \\ 1/2 & -1/2 & 0 & 1/\sqrt{2} \\ 1/2 & -1/2 & 0 & -1/\sqrt{2} \\ 1/2 & 1/2 & 1/\sqrt{2} & 0 \end{pmatrix}$$

*** Exercice 25: Soit f l'endomorphisme de $\mathbb{R}^3$ dont la matrice, dans la base naturelle, s'écrit $M = \begin{pmatrix} 1 & 1 & 0 \\ 0 & 1 & 1 \\ 1 & 0 & 1 \end{pmatrix}$

Quels sont les sous espaces F de $\mathbb{R}^3$, qui sont stables par f ?

*** Exercice 26: Même question qu'à l'exercice 25 avec $M = \begin{pmatrix} 1 & 0 & 1 \\ 2 & 0 & 0 \\ 0 & 1 & 1 \end{pmatrix}$

** Exercice 27: Trouver les valeurs propres et les vecteurs propres de l'endomorphisme ϕ de $\mathbb{R}^4$ dont la matrice, dans la base naturelle, s'écrit:

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

Déterminer deux sous-espaces F et G de $\mathbb{R}^4$ de dimension 2, supplémentaires dans $\mathbb{R}^4$, et stables par ϕ .

Extrait d'une épreuve de septembre 1980

Etude des comportements asymptotiques

§ 1 : Les suites de nombres réels.

Suites convergentes:

Rappelons qu'une suite de nombres réels $(u_n)_{n \in \mathbb{N}}$ converge vers u si :

$$(\forall \varepsilon) (\exists N_0) \text{ tel que } (\forall n) \ n \geq N_0 \text{ implique } |u_n - u| \leq \varepsilon$$

Rappelons que si $(u_n)_{n \in \mathbb{N}}$ converge vers u et $(v_n)_{n \in \mathbb{N}}$ vers v , alors :

$$a) (u_n + \lambda v_n)_{n \in \mathbb{N}} \text{ converge vers } u + \lambda v.$$

Ceci signifie que les suites convergentes forment un sous-espace vectoriel $\mathcal{S}$ de l'espace vectoriel $\mathcal{S}$ de toutes les suites, et que l'application qui à toute suite convergente associe sa limite est une forme linéaire sur $\mathcal{S}$.

$$b) (u_n v_n)_{n \in \mathbb{N}} \text{ converge vers } uv.$$

c) Si $u \neq 0$, le nombre u_n n'est nul que pour (au plus) un nombre fini de valeurs de n . Et la suite $(1/u_n)_{n \in \mathbb{N}}$ (qui est donc définie sauf peut être pour un nombre fini de valeurs de n) converge vers $1/u$.

d) Si f est une fonction définie au voisinage de u , et continue en u , alors $(f(u_n))_{n \in \mathbb{N}}$ converge vers $f(u)$.

Exercice a: Ecrire la définition de: $(u_n)_{n \in \mathbb{N}}$ converge vers $+\infty$. Ecrire la définition de $(u_n)_{n \in \mathbb{N}}$ est une suite bornée.

Comparaison de deux suites: Nous aurons souvent à "comparer les vitesses de convergence" de deux suites qui convergent vers 0 , ou vers $+\infty$. Pour cela nous allons introduire cinq relations de comparaison entre les suites:

La majoration stricte :

Nous dirons que $(u_n)_{n \in \mathbb{N}}$ majore $(v_n)_{n \in \mathbb{N}}$ si $(\forall n \in \mathbb{N})$ on a $u_n \geq v_n$.

La majoration à partir d'un certain rang:

Nous dirons que $(u_n)_{n \in \mathbb{N}}$ majore $(v_n)_{n \in \mathbb{N}}$ à partir d'un certain rang si:

$$(\exists N_0 \in \mathbb{N}) \text{ tel que } (\forall n) \ n \geq N_0 \text{ implique } u_n \geq v_n.$$

La relation $\mathcal{O}$ (lire: petit "o"):

Nous dirons que $(u_n)_{n \in \mathbb{N}} = \mathcal{O}((v_n)_{n \in \mathbb{N}})$ (ou encore que $(u_n)_{n \in \mathbb{N}}$ "est" $\mathcal{O}((v_n)_{n \in \mathbb{N}})$) s'il existe une suite $(\varepsilon_n)_{n \in \mathbb{N}}$ qui converge vers 0 , telle que, à partir d'un certain rang, $u_n = \varepsilon_n \cdot v_n$. C'est à dire si:

$$(\forall \alpha > 0) (\exists N_1) \text{ tel que } (\forall n) \ n \geq N_1 \text{ implique } |u_n| \leq \alpha |v_n|$$

Exercice b: Soit $(v_n)_{n \in \mathbb{N}}$ une suite strictement positive. Montrer que l'ensemble des suites $(u_n)_{n \in \mathbb{N}}$ qui sont $\mathcal{O}((v_n)_{n \in \mathbb{N}})$ est un sous-espace vectoriel $\mathcal{V}$ de l'espace vectoriel $\mathcal{S}$ de toutes les suites. On considère une suite $(w_n)_{n \in \mathbb{N}}$, et à toute suite $(u_n)_{n \in \mathbb{N}}$ on associe la suite $(w_n u_n)_{n \in \mathbb{N}}$, montrer que si $(w_n)_{n \in \mathbb{N}}$ est bornée on obtient ainsi une application linéaire de $\mathcal{V}$ dans lui même. A quelles conditions cette application est elle bijective ?

La relation O (lire: grand "O") :

Nous dirons que $(u_n)_{n \in \mathbb{N}} = O((v_n)_{n \in \mathbb{N}})$ (ou encore que $(u_n)_{n \in \mathbb{N}}$ "est" $O((v_n)_{n \in \mathbb{N}})$) s'il existe une suite $(\epsilon_n)_{n \in \mathbb{N}}$ bornée, telle que, à partir d'un certain rang, $u_n = \epsilon_n v_n$. C'est à dire si :

$$(\exists N_0)(\exists M > 0) \text{ tel que } (\forall n) \ n \geq N_0 \text{ implique } |u_n| \leq M |v_n| \quad (\text{ou } \left| \frac{u_n}{v_n} \right| \leq M \text{ si } v_n \neq 0)$$

L'équivalence asymptotique:

Nous dirons que $(u_n)_{n \in \mathbb{N}}$ est équivalente à $(v_n)_{n \in \mathbb{N}}$ s'il existe une suite $(\epsilon_n)_{n \in \mathbb{N}}$ qui converge vers 1, telle que, à partir d'un certain rang, $u_n = \epsilon_n v_n$.

Nous noterons alors $(u_n)_{n \in \mathbb{N}} \sim (v_n)_{n \in \mathbb{N}}$.

Exercice c: Démontrer que si $(u_n)_{n \in \mathbb{N}} - (v_n)_{n \in \mathbb{N}} = o((v_n)_{n \in \mathbb{N}})$, alors $(v_n)_{n \in \mathbb{N}} - (u_n)_{n \in \mathbb{N}} = o((u_n)_{n \in \mathbb{N}})$. Que peut-on alors dire de $(u_n)_{n \in \mathbb{N}}$ et $(v_n)_{n \in \mathbb{N}}$?

Exercice d: 1) On considère des suites $(u_n)_{n \in \mathbb{N}}$, $(u'_n)_{n \in \mathbb{N}}$, $(v_n)_{n \in \mathbb{N}}$ et $(v'_n)_{n \in \mathbb{N}}$ telles que $(u_n)_{n \in \mathbb{N}}$ soit équivalente à $(u'_n)_{n \in \mathbb{N}}$ et que $(v_n)_{n \in \mathbb{N}}$ soit équivalente à $(v'_n)_{n \in \mathbb{N}}$, démontrer que $(u_n v_n)_{n \in \mathbb{N}}$ est équivalente à $(u'_n v'_n)_{n \in \mathbb{N}}$ et que (si v_n n'a qu'un nombre fini de termes nuls) $(u_n/v_n)_{n \in \mathbb{N}}$ est équivalente à $(u'_n/v'_n)_{n \in \mathbb{N}}$.

2) Montrer que $(\frac{1}{n})_{n \geq 1}$ est équivalente à $(\frac{1}{n+1})_{n \in \mathbb{N}}$. Montrer que $(\frac{1}{n+2})_{n \in \mathbb{N}}$ est équivalente à $(\frac{n}{1+n^2})_{n \in \mathbb{N}}$. Montrer que $(\frac{1}{n} - \frac{1}{n+2})_{n \geq 1}$ n'est pas équivalente à $(\frac{1}{n+1} - \frac{n}{1+n^2})_{n \in \mathbb{N}}$.

Trouver d'autres exemples de suites $(u_n)_{n \in \mathbb{N}}$, $(u'_n)_{n \in \mathbb{N}}$, $(v_n)_{n \in \mathbb{N}}$, $(v'_n)_{n \in \mathbb{N}}$ telles que $(u_n)_{n \in \mathbb{N}}$ soit équivalente à $(u'_n)_{n \in \mathbb{N}}$, que $(v_n)_{n \in \mathbb{N}}$ soit équivalente à $(v'_n)_{n \in \mathbb{N}}$, et que $(u_n - v_n)_{n \in \mathbb{N}}$ ne soit pas équivalente à $(u'_n - v'_n)_{n \in \mathbb{N}}$.

Les suites de référence

La plupart du temps ces relations serviront à comparer une suite que l'on doit étudier, à l'une des suites de la liste suivante (qu'il est donc indispensable de connaître ainsi que leurs relations mutuelles).

$$\begin{array}{cccc} (\frac{1}{n})_{n \geq 1} & (\frac{1}{n^2})_{n \geq 1} & (\frac{1}{n^3})_{n \geq 1} & (\frac{1}{n^\alpha})_{n \geq 1} \quad (\alpha \text{ fixé positif}) \\ (n)_{n \in \mathbb{N}} & (n^2)_{n \in \mathbb{N}} & (n^3)_{n \in \mathbb{N}} & (n^\alpha)_{n \in \mathbb{N}} \quad (\alpha \text{ fixé positif}) \\ (\text{Log } n)_{n \geq 1} & (\frac{1}{\text{Log } n})_{n \geq 2} & (\sqrt{n})_{n \in \mathbb{N}} & (\frac{1}{\sqrt{n}})_{n \geq 1} \\ (e^n)_{n \in \mathbb{N}} & (2^n)_{n \in \mathbb{N}} & (10^n)_{n \in \mathbb{N}} & (\alpha^n)_{n \in \mathbb{N}} \quad (\alpha \text{ fixé } \neq 0) \quad (\frac{1}{e^n})_{n \in \mathbb{N}} \quad \text{Etc.} \end{array}$$

Exercice e : Comparer ces suites deux à deux selon chacune des 5 relations de comparaison ci-dessus (il s'agit de résultats que vous devriez connaître).

§ 2 : Etude des fonctions au voisinage de $+\infty$.

Toutes les fonctions que nous considérons ici, sont définies sur un même intervalle $[a, +\infty[$. Nous allons reprendre à leur sujet tout ce qui a été écrit pour les suites au §1(*).

Exercice f : Ecrire, en utilisant les quantificateurs les définitions de :

- a) $f: [a, +\infty[\rightarrow \mathbb{R}$ a pour limite L en $+\infty$
- b) $f: [a, +\infty[\rightarrow \mathbb{R}$ a pour limite $-\infty$ en $+\infty$
- a) $f: [a, +\infty[\rightarrow \mathbb{R}$ a pour limite $+\infty$ en $+\infty$

(*) Notons qu'on pourrait étudier de même les comportements au voisinage de $-\infty$. Le lecteur fera lui-même la traduction.

Comparaison des fonctions au voisinage de $+\infty$

Nous utiliserons pour les fonctions au voisinage de $+\infty$, quatre relations de comparaison analogues à celles que nous avons définies pour les suites.

La majoration au voisinage de $+\infty$: Soit f et g deux fonctions de $[a, +\infty[$ dans $\mathbb{R}$. Nous dirons que f majore g au voisinage de $+\infty$, si :

$$(\exists A) \text{ tel que } (\forall t \geq A) \quad f(t) \geq g(t)$$

La relation $\mathcal{O}$ (Lire: petit "o"): Soit $f: [a, +\infty[\rightarrow \mathbb{R}$ et $g: [a, +\infty[\rightarrow \mathbb{R}$. Nous dirons que $f = \mathcal{O}(g)$ au voisinage de $+\infty$, s'il existe une fonction γ qui a pour limite 0 en $+\infty$, telle que (pour t supérieur à un certain nombre A) $f(t) = g(t)\gamma(t)$.

Exemples: (résultats classiques à connaître)

$$(\forall k > 0) \quad e^{-t} = \mathcal{O}(t^{-k}) \text{ au voisinage de } +\infty$$

$$(\forall k > 0) \quad t^k = \mathcal{O}(e^t) \text{ au voisinage de } +\infty$$

$$(\forall k > 0) \quad \text{Log } t = \mathcal{O}(t^k) \text{ au voisinage de } +\infty$$

La relation $\mathcal{O}$ (Lire: grand "O"): Soit f et g deux fonctions définies sur $]a, +\infty[$. Nous dirons que $f = \mathcal{O}(g)$ au voisinage de $+\infty$, s'il existe un intervalle $[A, +\infty[$ et un nombre M tels que $(\forall t \in [A, +\infty[)$ on ait: $|f(t)| \leq M |g(t)|$.

L'équivalence asymptotique: Soit f et g deux fonctions définies sur $]a, +\infty[$. Nous dirons que f et g sont équivalentes au voisinage de $+\infty$, s'il existe une fonction γ qui a pour limite 1 en $+\infty$, telle que (pour t supérieur à un certain nombre A) $f(t) = g(t)\gamma(t)$. On le notera souvent $f \sim g$.

Exercice g : Trouver des fonctions f, g, ϕ et γ telles que l'on ait $f \sim g$ et $\phi \sim \gamma$ mais que l'on n'ait pas $f + \phi \sim g + \gamma$.

§ 3 : Etude des fonctions au voisinage de a .

Tout ce qui a été fait au voisinage de $+\infty$, peut être repris au voisinage d'un point a . Notons que trois situations (distinctes mais très analogues) peuvent alors apparaître:

On peut regarder le comportement "au voisinage de a , à droite de a ", c'est à dire sur un intervalle de la forme $]a, a+h[$ ($h > 0$).

On peut regarder le comportement "au voisinage de a , à gauche de a ", c'est à dire sur un intervalle de la forme $]a-h, a[$ ($h > 0$).

On peut regarder le comportement "au voisinage de a ", c'est à dire sur $]a-h, a[\cup]a, a+h[$; les fonctions considérées n'étant alors souvent pas définies en a .

Exercice h : Ecrire les définitions de :

1) f tend vers $+\infty$ en a à droite.

2) f a pour limite à gauche en a , le nombre λ .

Exercice i : En s'inspirant des définitions des § 1 et 2, écrire les définitions de :

1) $f = \mathcal{O}(g)$ au voisinage de a , à droite.

2) f est équivalente à g au voisinage de a .

3) $|f|$ majore g au voisinage de a .

4) $f = \mathcal{O}(g)$ au voisinage de a .

§ 4 : Développement de Taylor

Soit $f:]a, b[\rightarrow \mathbb{R}$ de classe C^r . Soit x_0 un point de $]a, b[$. Nous appellerons k -ième polynôme de Taylor de f en x_0 , le polynôme de degré $k(\leq r)$ (où $f^{(p)}$ est la dérivée p -ième de f):

$$P_k(x) = f(x_0) + (x-x_0)f'(x_0) + \frac{(x-x_0)^2}{2!}f''(x_0) + \dots + \frac{(x-x_0)^k}{k!}f^{(k)}(x_0)$$

C'est le seul polynôme de degré k dont les dérivées en x_0 sont, jusqu'à l'ordre k , égales à celles de f . C'est évidemment une fonction qui "ressemble à f au voisinage de x_0 "; mesurer cette ressemblance c'est évaluer l'erreur faite en remplaçant f par P_k , c'est à dire évaluer $f - P_k$.

Première évaluation (Reste de Young) :

$(f - P_k)(x) = \mathcal{O}((x-x_0)^k)$ au voisinage de x_0 , sous la seule condition que f est de classe C^k .

Deuxième évaluation (Reste de Lagrange): Si sur un intervalle $]x_0 - \alpha, x_0 + \alpha[$, f est de classe C^{k+1} , et si $|f^{(k+1)}(t)| \leq M$ pour tout t dans $]x_0 - \alpha, x_0 + \alpha[$, alors:

$$|f(x) - P_k(x)| \leq \frac{M|x-x_0|^{k+1}}{(k+1)!} \text{ pour tout } x \text{ dans }]x_0 - \alpha, x_0 + \alpha[$$

Troisième évaluation (reste intégral): Si f est de classe C^{k+1}

$$f(x) - P_k(x) = \int_{x_0}^x \frac{(x-\theta)^k}{k!} f^{(k+1)}(\theta) d\theta$$

Nous considérons que ces résultats sont connus; signalons toutefois que la démonstration du troisième fait l'objet de l'exercice 21.

§ 5 : Développement limités et développements de Laurent.

Soit $f:]x_0 - \eta, x_0 + \eta[\rightarrow \mathbb{R}$, on dit que f a un développement limité à l'ordre n en x_0 , s'il existe un polynôme $P(x) = a_0 + a_1(x-x_0) + \dots + a_n(x-x_0)^n$, tel que $f(x) - P(x) = \mathcal{O}((x-x_0)^n)$ au voisinage de x_0 .

Soit f une fonction définie sur $]x_0 - \eta, x_0 + \eta[$, sauf peut être en x_0 , et à valeurs dans $\mathbb{R}$. Nous dirons que f a un développement de Laurent à l'ordre n en x_0 , s'il existe un polynôme de Laurent $P(x) = a_{-k}(x-x_0)^{-k} + \dots + a_0 + \dots + a_n(x-x_0)^n$, tel que $f(x) - P(x) = \mathcal{O}((x-x_0)^n)$ au voisinage de x_0 .

Exemples: La formule de Taylor (avec reste de Young) nous dit que toute fonction de classe C^r au voisinage de x_0 , a un développement limité à l'ordre r en x_0 .

Grâce aux développements de $\cos$ et $\sin$ en 0 , on peut établir le développement de Laurent de $\cotg$ à l'ordre 3 en 0 ; il s'écrit:

$$\cotg x = \frac{1}{x} - \frac{x}{3} - \frac{x^3}{45} + \mathcal{O}(x^3)$$

Sommes de développements : Si $f(x) = P(x) + \mathcal{O}((x-x_0)^n)$ et $g(x) = Q(x) + \mathcal{O}((x-x_0)^n)$ au voisinage de x_0 , alors $f(x) + g(x) = P(x) + Q(x) + \mathcal{O}((x-x_0)^n)$. Autrement dit pour avoir un développement de $f+g$ à l'ordre n en x_0 , on ajoute les développements à l'ordre n en x_0 de f et de g .

Produits de développements : Si f a un développement qui s'écrit

$$f(x) = a_p(x-x_0)^p + \dots + a_n(x-x_0)^n + \mathcal{O}((x-x_0)^n) \quad (p \in \mathbb{Z} \text{ et } a_p \neq 0)$$

nous dirons que ce développement est de longueur effective $L = n-p$. Il s'écrit aussi:

$$f(x) = (x-x_0)^p [a_p + \dots + a_n(x-x_0)^{L} + \mathcal{O}((x-x_0)^L)]$$

Supposons connu le développement de longueur effective L de g en x_0 :

$$g(x) = (x-x_0)^q [b_q + \dots + b_n(x-x_0)^L + \mathcal{O}((x-x_0)^L)] \quad (q \in \mathbb{Z} \text{ et } b_q \neq 0)$$

Alors nous obtiendrons le développement du produit fg de longueur effective L , en faisant le produit de ces deux expressions, et en oubliant dans le produit des deux crochets les termes en $(x-x_0)^m$ pour $m > L$.

Exemple: Les développements de longueur effective 3 en 0 de $\sin x$ et de $1 - \cos x$ s'écrivent:

$$\begin{aligned}\sin x &= x - \frac{x^3}{6} + o(x^4) = x \left[1 - \frac{x^2}{6} + o(x^3) \right] \\ 1 - \cos x &= \frac{x^2}{2} - \frac{x^4}{24} + o(x^5) = x^2 \left[\frac{1}{2} - \frac{x^2}{24} + o(x^3) \right]\end{aligned}$$

Donc :

$$(1 - \cos x) \sin x = x^3 \left\{ \left[1 - \frac{x^2}{6} \right] \left[\frac{1}{2} - \frac{x^2}{24} \right] + o(x^3) \right\} = x^3 \left[\frac{1}{2} - \frac{x^2}{8} + o(x^3) \right]$$

Développement de $1/f$: Si nous connaissons le développement de longueur effective L de f en x_0 , nous pouvons en déduire le développement de longueur effective L en x_0 de $1/f$.

$$f(x) = (x-x_0)^p \left[a_p + \dots + a_n(x-x_0)^L + o((x-x_0)^L) \right] = a_p(x-x_0)^p \left[1 + \dots + \frac{a_n}{a_p}(x-x_0)^L + o((x-x_0)^L) \right]$$

Soit $f(x) = a_p(x-x_0)^p \left[1 - Q(x-x_0) + o((x-x_0)^L) \right]$

Alors $\frac{1}{f(x)} = \frac{1}{a_p}(x-x_0)^{-p} \left[1 + Q(x-x_0) + \dots + [Q(x-x_0)]^L + o((x-x_0)^L) \right]$

Et dans $1 + Q(x-x_0) + \dots + [Q(x-x_0)]^L$, on oubliera les termes en $(x-x_0)^m$ pour $m > L$.

Exemple : Le développement de longueur 4 de $1 - \cos x$ en 0 s'écrit :

$$1 - \cos x = \frac{x^2}{2} - \frac{x^4}{24} + \frac{x^6}{720} + o(x^6) = \frac{1}{2}x^2 \left[1 - \frac{x^2}{12} + \frac{x^4}{360} + o(x^4) \right]$$

Donc :

$$\frac{1}{1 - \cos x} = 2x^{-2} \left[1 + \left(\frac{x^2}{12} - \frac{x^4}{360} \right) + \left(\frac{x^2}{12} - \frac{x^4}{360} \right)^2 + \left(\frac{x^2}{12} - \frac{x^4}{360} \right)^3 + \left(\frac{x^2}{12} - \frac{x^4}{360} \right)^4 + o(x^4) \right]$$

Et, dans le crochet, on oublie les termes en x^k pour $k > 4$.

Développement d'une fonction composée: Soit f et g deux fonctions. On suppose que f a un développement limité à l'ordre n en x_0 , qui s'écrit:

$$f(x) = y_0 + a_1(x-x_0) + \dots + a_n(x-x_0)^n + o((x-x_0)^n) = y_0 + Q(x-x_0) + o((x-x_0)^n)$$

On suppose que g a un développement limité à l'ordre n en y_0 , qui s'écrit:

$$g(y) = b_0 + b_1(y-y_0) + \dots + b_n(y-y_0)^n + o((y-y_0)^n)$$

Alors la fonction composée admet un développement limité à l'ordre n en x_0 , il est obtenu en écrivant:

$$g(f(x)) = b_0 + b_1Q(x-x_0) + \dots + b_n[Q(x-x_0)]^n + o((x-x_0)^n)$$

Puis en oubliant les termes en $(x-x_0)^p$ pour $p > n$.

Développements à l'infini: Soit à étudier une fonction f au voisinage de $\pm\infty$; posons $g(t) = f(\frac{1}{t})$;

nous sommes ramenés à l'étude de g au voisinage de 0. Supposons que g ait un développement au voisinage de 0 qui s'écrit :

$$g(t) = a_pt^p + \dots + a_nt^n + o(t^n) \quad (p \in \mathbb{Z})$$

Nous aurons:

$$f(x) = a_px^{-p} + \dots + a_nx^{-n} + o(x^{-n}) \quad (p \in \mathbb{Z})$$

Un tel développement est appelé un développement à l'infini de f .

Exercice j : Trouver un développement (de longueur effective 3) au voisinage de l'infini de la fonction:

$$x \rightarrow \frac{1 + 2x + 3x^2 + 4x^4 + 5x^5}{2 + 4x^2}$$

Exercices sur le chapitre 6

** Exercice 1: Soit $(u_n)_{n \in \mathbb{N}}$ une suite telle que :

$$(\forall k \text{ entier positif}) (\exists N_0 \in \mathbb{N}) \text{ tel que } (\forall n) n \geq N_0 \text{ implique } |u_n| \leq n^{-k}$$

Une telle suite est dite à décroissance rapide.

a) Pour quelles valeurs de p , a-t-on :

1) $(u_n)_{n \in \mathbb{N}} = O(n^p)$

2) $(u_n)_{n \in \mathbb{N}} = o(n^p)$

3) $(n^p)_{n \in \mathbb{N}}$ majore $(u_n)_{n \in \mathbb{N}}$

4) $(n^p)_{n \in \mathbb{N}}$ majore $(u_n)_{n \in \mathbb{N}}$ à partir d'un certain rang.

b) L'assertion $(u_n)_{n \in \mathbb{N}} = O((e^{-n})_{n \in \mathbb{N}})$ est elle certainement vraie ? Certainement fausse ? Peut être vraie, mais peut être fausse ?

c) L'assertion $(u_n)_{n \in \mathbb{N}} = o((e^{-n})_{n \in \mathbb{N}})$ est elle certainement vraie ? Certainement fausse ? Peut être vraie, mais peut être fausse ?

* Exercice 2: Soit trois suites $(u_n)_{n \in \mathbb{N}}$, $(v_n)_{n \in \mathbb{N}}$ et $(a_n)_{n \in \mathbb{N}}$ telles que :

$$(u_n)_{n \in \mathbb{N}} = O((a_n)_{n \in \mathbb{N}}) \text{ et } (v_n)_{n \in \mathbb{N}} = o((a_n)_{n \in \mathbb{N}})$$

a) Que peut on dire de $(u_n + v_n)_{n \in \mathbb{N}}$ et de $(a_n)_{n \in \mathbb{N}}$

b) Quelles hypothèses doit vérifier $(a_n)_{n \in \mathbb{N}}$ pour qu'on puisse affirmer que $(u_n v_n)_{n \in \mathbb{N}} = o((a_n)_{n \in \mathbb{N}})$?

** Exercice 3: On suppose que $(u_n)_{n \in \mathbb{N}} = o((\frac{1}{n})_{n \geq 1})$. Chacune des assertions suivantes est elle certainement fausse ? Certainement vraie ? Indécidable ?

a) $(\log u_n)_{n \in \mathbb{N}} = o((\frac{1}{n+1})_{n \in \mathbb{N}})$

b) $(\log u_n)_{n \in \mathbb{N}} = o((\log(2+n))_{n \in \mathbb{N}})$

c) $(\log(1+u_n))_{n \in \mathbb{N}} = o((\frac{1}{n+1})_{n \in \mathbb{N}})$

d) $(n)_{n \in \mathbb{N}} = o((\frac{1}{u_n})_{n \in \mathbb{N}})$ (si $(\forall n) : u_n \neq 0$)

e) $(\sqrt{u_n})_{n \in \mathbb{N}} = o((\frac{1}{\sqrt{n+1}})_{n \in \mathbb{N}})$

* Exercice 4: Soit $(u_n)_{n \in \mathbb{N}}$ définie par $u_0 = a, u_1 = b, u_2 = c$ et $u_{n+3} = -2u_{n+2} + u_{n+1} + 2u_n$.

a) Calculer u_n en fonction de a, b, c et n ; en remarquant que l'on a l'égalité matricielle

$$\begin{pmatrix} u_{n+3} \\ u_{n+2} \\ u_{n+1} \end{pmatrix} = \begin{pmatrix} -2 & 1 & 2 \\ 1 & 0 & 0 \\ 0 & 1 & 0 \end{pmatrix} \begin{pmatrix} u_{n+2} \\ u_{n+1} \\ u_n \end{pmatrix}$$

b) Pour quelles valeurs de a, b, c , a-t-on

1) $(u_n)_{n \in \mathbb{N}}$ converge vers 0 ?

2) $(u_n)_{n \in \mathbb{N}}$ converge vers $+\infty$?

c) Soit $a=b=c=1$, pour quels k a-t-on

2) $(u_n)_{n \in \mathbb{N}} = o((k^{-n})_{n \in \mathbb{N}})$?

3) $(u_n)_{n \in \mathbb{N}} = O((\frac{k}{1+n^2})_{n \in \mathbb{N}})$?

* Exercice 5: Soit $(u_n)_{n \in \mathbb{N}}$ une suite. Voici 5 affirmations

(α) $(u_n)_{n \in \mathbb{N}} = o(1)$

(β) $(u_n)_{n \in \mathbb{N}}$ converge vers 0

(γ) $(|u_n|)_{n \in \mathbb{N}} = o(1)$

(δ) $(n \log n)_{n \geq 2}$ majore $(u_n)_{n \in \mathbb{N}}$

(ϵ) $(u_n)_{n \in \mathbb{N}}$ est équivalente à $(\frac{1}{1+n^2})_{n \in \mathbb{N}}$

a) Existe-t-il plusieurs de ces affirmations qui sont logiquement équivalentes ?

b) Il en existe une qui implique toutes les autres. La quelle ?

* **Exercice 6:** Soit $(u_n)_{n \in \mathbb{N}}$ et $(v_n)_{n \in \mathbb{N}}$. Voici 4 affirmations

$$(\alpha) \quad (u_n)_{n \in \mathbb{N}} \text{ est équivalente à } \left(\frac{v_n}{1+n}\right)_{n \in \mathbb{N}}$$

$$(\beta) \quad (u_n)_{n \in \mathbb{N}} = o((v_n)_{n \in \mathbb{N}})$$

$$(\gamma) \quad (u_n)_{n \in \mathbb{N}} = O((v_n)_{n \in \mathbb{N}})$$

$$(\delta) \quad (v_n)_{n \in \mathbb{N}} \text{ majore } (u_n)_{n \in \mathbb{N}}$$

a) Parmi ces affirmations, en existe-t-il deux qui sont contradictoires ?

b) En existe-t-il deux qui sont équivalentes ?

c) En existe-t-il une qui est conséquence des 3 autres ?

* **Exercice 7:** Ecrire au moyen des quantificateurs

a) Que la suite $(u_n)_{n \in \mathbb{N}}$ ne converge pas vers u

b) Que la suite $(u_n)_{n \in \mathbb{N}}$ n'a pas de limite.

c) Que la suite $(u_n)_{n \in \mathbb{N}}$ n'est pas bornée.

d) Que la suite $(u_n)_{n \in \mathbb{N}}$ ne converge pas vers $+\infty$

** **Exercice 8:** Que signifient les écritures suivantes ? En existe-t-il qui sont absurdes ?

$(U_n)_{n \in \mathbb{N}}$ est une suite de nombres réels

$$(\alpha) \quad (\forall \varepsilon > 0) \quad (\forall N_0 \in \mathbb{N}) \quad (\forall n) \quad n \geq N_0 \text{ implique } |U_n| \geq \varepsilon$$

$$(\beta) \quad (\exists N_0 \in \mathbb{N}) \quad (\forall \varepsilon > 0) \quad (\forall n) \quad n \geq N_0 \text{ implique } |U_n| \leq \varepsilon$$

$$(\gamma) \quad (\forall N_0 \in \mathbb{N}), (\exists \varepsilon > 0) \quad (\forall n) \quad n \geq N_0 \text{ implique } |U_n| \leq \varepsilon$$

*** **Exercice 9:** Existe-t-il une suite $(U_n)_{n \in \mathbb{N}}$ qui vérifie (α) et ne vérifie pas (β) ?

$$(\alpha) \quad (\forall \varepsilon > 0) \quad (\exists N_0 \in \mathbb{N}) \quad (\exists p \in \mathbb{N}; p \neq 0) \quad (\forall n) \quad n \geq N_0 \text{ implique } |U_{n+p} - U_n| \geq \varepsilon.$$

$$(\beta) \quad (\exists N_0 \in \mathbb{N}) \quad (\forall \varepsilon > 0) \quad (\exists p \in \mathbb{N}; p \neq 0) \quad (\forall n) \quad n \geq N_0 \text{ implique } |U_{n+p} - U_n| \geq \varepsilon.$$

** **Exercice 10:** Soit $(u_n)_{n \in \mathbb{N}}$ et $(v_n)_{n \in \mathbb{N}}$ deux suites. On suppose que $(u_n)_{n \in \mathbb{N}}$ est majorée à partir d'un certain rang par $(\frac{1}{1+n^2})_{n \in \mathbb{N}}$, et que $(v_n)_{n \in \mathbb{N}} \sim (\frac{1}{2+n})_{n \in \mathbb{N}}$.

a) Est ce que $(u_n + v_n)_{n \in \mathbb{N}}$ est majorée à partir d'un certain rang par $(\frac{1}{1+n})_{n \in \mathbb{N}}$ (certainement, peut être, certainement pas ?)

b) Est ce que $(u_n + v_n)_{n \in \mathbb{N}} = O((\frac{1}{1+n})_{n \in \mathbb{N}})$ (certainement, peut être, certainement pas)

c) Pour quels nombres (réels positifs) k , a-t-on $(u_n + v_n)_{n \in \mathbb{N}} = o((\frac{1}{1+n^k})_{n \in \mathbb{N}})$

d) Pour quels nombres (réels positifs) k , a-t-on $(u_n v_n)_{n \in \mathbb{N}} = o((\frac{1}{1+n^k})_{n \in \mathbb{N}})$

* **Exercice 11:**

On considère les suites $(\frac{1+n}{1+n^2})_{n \in \mathbb{N}}$, $(\frac{1}{1+n})_{n \in \mathbb{N}}$, $(\frac{1}{\sqrt{n^2+1}})_{n \in \mathbb{N}}$, $(e^{-n})_{n \in \mathbb{N}}$, et $(\text{Log}(n^2+1))_{n \in \mathbb{N}}$.

a) Les classer pour la relation "majorer à partir d'un certain rang".

b) Les classer pour la relation o .

c) Les classer pour la relation O .

* **Exercice 12:** On considère les fonctions suivantes, définies sur $]1, \infty[$:

$$f(x) = \text{Log}(1+e^x)$$

$$g(x) = \text{Log}(\text{ch } x)$$

$$h(x) = \text{Log}(\text{sh } x)$$

$$i(x) = \text{Argch}(e^x)$$

$$j(x) = \text{Argsh}(\text{ch } x)$$

$$k(x) = \text{Log}(\text{ch } 2x)$$

a) les classer pour la relation "majorer au voisinage de $+\infty$ ".

b) les classer pour la relation "O au voisinage de $+\infty$ ".

c) les classer pour la relation " o au voisinage de $+\infty$ ".

d) en existe-t-il qui sont équivalentes au voisinage de $+\infty$?

* **Exercice 13:** Classer les fonctions suivantes pour la relation Θ au voisinage de 0 à droite:

$$f(x) = \text{Log}(1+x)$$

$$g(x) = x \text{ Log } x$$

$$h(x) = \frac{-1}{\text{Log } x}$$

$$i(x) = e^{\frac{-1}{x}}$$

$$j(x) = e^{\frac{-1}{x}} \text{ Log } x$$

$$k(x) = \text{Log}(1 + e^{\frac{-1}{x}})$$

** **Exercice 14:** On considère les assertions suivantes:

$$a) (\forall n \in \mathbb{N}) f(x) = \Theta(x^n) \text{ au voisinage de } 0$$

$$b) (\forall n \in \mathbb{N}) f(x) = O(x^n) \text{ au voisinage de } 0$$

$$c) (\forall n \in \mathbb{N}) f(x) \text{ est majorée par } x^n \text{ au voisinage de } 0$$

En existe-t-il qui entraînent toutes les autres? En existe-t-il qui sont conséquences de toutes les autres?

* **Exercice 15 :** La fonction f est définie par $f(x) = 0$ si $x \leq 0$ et $f(x) = e^{\frac{-1}{x}}$ pour $x > 0$. Ecrire le développement limité à l'ordre 342 de f au voisinage de 0.

* **Exercice 16 :** Trouver les développements de Laurent, de longueur effective 4 des fonctions suivantes:

$$f(x) = \frac{1}{\sin x - \text{sh } x} \quad (\text{en } 0)$$

$$g(x) = \frac{1}{\sin x - \cos x} \quad (\text{en } \frac{\pi}{4})$$

$$h(x) = \frac{\text{Log}(\text{ch } x)}{\text{Log}(\cos x)} \quad (\text{en } 0)$$

$$i(x) = \text{tg } x \quad (\text{en } \frac{\pi}{2})$$

$$j(x) = \frac{1}{\text{Log } x} \quad (\text{en } 1)$$

$$k(x) = \frac{\sin x}{\text{Log}(\cos 2x)} \quad (\text{en } 0)$$

* **Exercice 17:** Ecrire les développements limités à l'ordre 4 des fonctions suivantes au voisinage de 1 :

$$f(x) = (x-4)(x+5)(x-2)(x^2+1)$$

$$g(x) = \cos\left(\frac{\pi x}{3}\right)$$

Ecrire ensuite leurs développements à l'ordre 29.

* **Exercice 18:** Soit $f(x) = \frac{2x^5 + x^4 + x^3 + 3x^2 + x + 1}{x(x^3+1)}$ et $g(x) = \frac{2x^5 + x^4 + x^3 + 4x^2 + 3x + 2}{x(x^3+2)}$

a) Trouver des développements de longueur effective 4 de f et de g au voisinage de $+\infty$. En déduire que les courbes représentatives de f et g ont même asymptote au voisinage de $+\infty$. Placer ces courbes par rapport à cette asymptote; et placer ces courbes l'une par rapport à l'autre au voisinage de $+\infty$.

b) Mêmes questions au voisinage de $-\infty$.

* **Exercice 19 :** Soit $f(x) = \frac{x^4 + 1}{x^2 - x + 2}$. Trouver un développement de longueur effective 4 de f au voisinage de $+\infty$. En déduire que la courbe représentative de f est asymptote à une parabole; placer cette courbe et cette parabole l'une par rapport à l'autre.

** **Exercice 20:** La formule de Taylor nous dit que toute fonction de classe C^r a, en tout point x_0 , un développement limité à l'ordre r . Il ne faudrait pas croire que, réciproquement, si f a un développement limité à l'ordre r en x_0 , elle est r fois dérivable en x_0 . Soit f définie par $f(0) = 0$ et $f(x) = x^3 \sin \frac{1}{x}$ pour $x \neq 0$.

a) Calculer f' (pour $x = 0$ et pour $x \neq 0$), puis montrer que f' est continue, mais non dérivable en 0.

b) Montrer que f a un développement limité à l'ordre 2 en 0, et écrire ce développement. Est-ce que f a un développement à l'ordre 3 en 0?

** **Exercice 21:** Soit $f :]-\eta, \eta[\rightarrow \mathbb{R}$ une fonction C^∞ .

a) Ecrire les formules de Taylor avec reste intégral à l'ordre n et à l'ordre $n+1$ de f au voisinage de 0. Montrer que la seconde s'obtient à partir de la première par une intégration par parties.

b) Démontrer, pour toute valeur de n , la formule de Taylor avec reste intégral à l'ordre n .

Outils fondamentaux du raisonnement

Deuxième partie: Quantificateurs et expressions formalisées

§1 : Pourquoi formalise-t-on l'écriture des mathématiques ?

On appelle *quantificateurs* les deux signes $\forall$ (lire "quel que soit") et $\exists$ (lire "il existe"). A quoi servent-ils ?

Considérons la phrase "entre deux nombres distincts, il existe deux rationnels distincts". Elle exprime une propriété de $\mathbb{R}$ et de son sous-ensemble $\mathbb{Q}$. Nous pouvons la traduire par:

$$(\forall a \in \mathbb{R}) (\forall b \in \mathbb{R}, a < b) (\exists x \in \mathbb{Q}) (\exists y \in \mathbb{Q}) \text{ tels que } a < x < y < b$$

Si l'on introduit de telles expressions formalisées, ce n'est pas pour le plaisir absurde de raconter les mathématiques dans un langage inaccessible aux non initiés. Vous avez l'habitude d'écrire les mathématiques dans votre langue maternelle, mais, dès que les situations se compliquent, celle-ci devient insuffisante; car elle manque de précision. Par exemple dans l'expression "entre deux nombres distincts", il fallait comprendre "entre deux nombres distincts quelconques". Qu'apporte le mot que l'on vient d'ajouter ? Peut être rien, puisque vous aviez compris que les deux nombres étaient quelconques; dans la phrase initiale il était sous-entendu. Mais aussi beaucoup car il est essentiel de faire la différence entre les deux assertions:

"entre deux nombres distincts quelconques, il existe deux rationnels distincts"

et "entre deux nombres distincts bien choisis, il existe deux rationnels distincts"

La première est beaucoup plus forte que la seconde. Ainsi l'essentiel de la signification de notre phrase initiale, ce qui sur le plan de la logique a le plus d'importance, était sous-entendu. *Les expressions quantifiées sont d'abord un moyen d'écrire les mathématiques avec un maximum de précision.*

§ 2 : Les règles d'écriture des expressions formalisées

Dans l'expression ci-dessus interviennent plusieurs types de signes mathématiques. Les signes $\mathbb{R}$, $\mathbb{Q}$, $\in$, $<$ représentent des objets mathématiques connus ou des relations connues entre des objets mathématiques. Ils n'ont pas à être définis, tout le monde est supposé les connaître. Toute la signification de l'expression est dans les conditions qui permettent d'affirmer que les objets représentés par les lettres a, b, x, y vérifient $a < x < y < b$. La première règle est la suivante:

Toutes les lettres d'une expression formalisée doivent être définies par un quantificateur.

Chaque lettre u est introduite par l'une des deux formules $(\forall u \in \dots)$ et $(\exists u \in \dots)$; autrement dit on doit dire à quel ensemble d'objets il appartient: et on doit dire si la propriété que l'on va énoncer est vraie pour un élément quelconque de cet ensemble, ou si elle ne l'est que pour un choix convenable.

Exemple: Soit $f: I \rightarrow \mathbb{R}$, l'expression:

$$(\forall \varepsilon > 0) (\exists \alpha > 0) \text{ tel que } |x - y| \leq \alpha \text{ implique } |f(x) - f(y)| \leq \varepsilon$$

n'est pas correcte, car on ne sait pas ce que sont x et y . Si l'on ajoute $(\forall x \in I)$ avant $(\forall \varepsilon)$, et $(\forall y \in I)$ après "tel que", on obtient

$$(\forall x \in I) (\forall \varepsilon > 0) (\exists \alpha > 0) \text{ tel que } (\forall y) |x - y| \leq \alpha \text{ implique } |f(x) - f(y)| \leq \varepsilon$$

On a écrit que f est continue en tout point x de I .

Tandis que si l'on ajoute $(\exists x \in I)$ avant $(\forall \epsilon)$, et $(\forall y \in I)$ après "tel que", on obtient

$$(\exists x \in I) (\forall \epsilon > 0) (\exists \alpha > 0) \text{ tel que } (\forall y) |x - y| \leq \alpha \text{ implique } |f(x) - f(y)| \leq \epsilon$$

On a écrit qu'il existe dans I un point où f est continue.

On voit donc qu'en faisant varier les quantificateurs qui définissent les deux lettres x et y , on donne des significations très différentes à l'expression.

Remarquons que toutes les formalisations que nous venons d'écrire sont incomplètes. Au lieu d'écrire $(\forall \epsilon > 0)$ nous aurions dû écrire $(\forall \epsilon \in \mathbb{R}, \epsilon > 0)$ ou $(\forall \epsilon \in [0, \infty[)$; le fait que ϵ soit un nombre réel a été sous-entendu. Pour ne pas alourdir exagérément les écritures, nous utiliserons souvent ce type de sous-entendu.

Le rôle de l'emplacement des définitions quantifiées.

L'emplacement des $(\forall x \dots)$ et $(\exists y \dots)$ joue un rôle important. C'est que lorsque nous écrivons $(\forall x \dots) (\exists y \dots)$, on construit y lorsque l'on connaît x , donc y dépend de x . Tandis que si nous écrivons $(\exists y \dots) (\forall x \dots)$, la construction de y se fait avant qu'on connaisse x , donc y ne dépend pas de x ; il est au contraire valable "quelque soit x ".

Exemple; Considérons une fonction $f: I \rightarrow \mathbb{R}$, et les deux assertions

$$(\forall \epsilon > 0) (\forall x \in I) (\exists \alpha > 0) \text{ tel que } (\forall y \in I) |x - y| \leq \alpha \text{ implique } |f(x) - f(y)| \leq \epsilon$$

$$(\forall \epsilon > 0) (\exists \alpha > 0) \text{ tel que } (\forall x \in I) (\forall y \in I) |x - y| \leq \alpha \text{ implique } |f(x) - f(y)| \leq \epsilon$$

Dans la première on affirme que f est continue en tout point x de I . Lorsque l'on a choisi (arbitrairement) ϵ et x , on peut trouver α qui dépend donc de ϵ et de x . Tandis que dans la seconde lorsque l'on a choisi ϵ (arbitrairement), on peut trouver α , qui ne dépend donc que de ϵ (et est valable pour tout x). Cette seconde assertion signifie que f est uniformément continue sur I .

Dans une expression formalisée:

* Echanger deux $(\exists u \dots) (\exists v \dots)$ qui se suivent ne change pas le sens de l'expression.

* Echanger deux $(\forall u \dots) (\forall v \dots)$ qui se suivent ne change pas le sens de l'expression.

* Changer $(\exists u \dots) (\forall v \dots)$ en $(\forall v \dots) (\exists u \dots)$ donne une assertion plus faible (par exemple pour la fonction f de l'exemple ci-dessus, être continue est une condition plus faible qu'être uniformément continue).

Comment nier une assertion ?

Ecrivons que $f: I \rightarrow \mathbb{R}$ n'est pas continue sur I :

$$(\exists \epsilon > 0) (\exists x \in I) (\forall \alpha > 0) (\exists y \in I) \text{ tel que } |x - y| \leq \alpha \text{ et } |f(x) - f(y)| \geq \epsilon$$

La règle est très simple:

Nous avons remplacé tous les $\exists$ par des $\forall$, et tous les $\forall$ par des $\exists$.

Puis nous avons remplacé " $|x - y| \leq \alpha$ implique $|f(x) - f(y)| \leq \epsilon$ " par " $|x - y| \leq \alpha$ et $|f(x) - f(y)| \geq \epsilon$ " (en français on dirait plutôt " $|x - y| \leq \alpha$ mais $|f(x) - f(y)| \geq \epsilon$ ")

Exercice: a) Ecrire que $f: I \rightarrow \mathbb{R}$ n'est pas uniformément continue sur I .

b) Ecrire que f est continue en (au moins) un point de I . Puis écrire que f n'est continue en aucun point de I .

Exercices sur le chapitre 6b.

Exercice 1: Ecrire au moyen des quantificateurs ($\forall$ et $\exists$)

- a) Que la suite $(u_n)_{n \in \mathbb{N}}$ ne converge pas vers u
- b) Que la suite $(u_n)_{n \in \mathbb{N}}$ n'a pas de limite.
- c) Que la suite $(u_n)_{n \in \mathbb{N}}$ n'est pas bornée.
- d) Que la suite $(u_n)_{n \in \mathbb{N}}$ ne converge pas vers $+\infty$

**** Exercice 2:** Que signifient les écritures suivantes ? En existe-t-il qui sont absurdes ?

$(U_n)_{n \in \mathbb{N}}$ est une suite de nombres réels

- $(\alpha) (\forall \varepsilon > 0) (\forall N_0 \in \mathbb{N}) (\forall n) n \geq N_0 \text{ implique } |U_n| \leq \varepsilon$
- $(\beta) (\exists N_0 \in \mathbb{N}) (\forall \varepsilon > 0) (\forall n) n \geq N_0 \text{ implique } |U_n| \geq \varepsilon$
- $(\gamma) (\forall N_0 \in \mathbb{N}), (\exists \varepsilon > 0) (\forall n) n \geq N_0 \text{ implique } |U_n| \leq \varepsilon$

**** Exercice 3:** Existe-t-il une suite $(U_n)_{n \in \mathbb{N}}$ qui vérifie (α) et ne vérifie pas (β) ?

- $(\alpha) (\forall \varepsilon > 0) (\exists N_0 \in \mathbb{N}) (\exists p \in \mathbb{N}; p \neq 0) (\forall n) n \geq N_0 \text{ implique } |U_{n+p} - U_n| \leq \varepsilon$
- $(\beta) (\exists N_0 \in \mathbb{N}) (\forall \varepsilon > 0) (\exists p \in \mathbb{N}; p \neq 0) (\forall n) n \geq N_0 \text{ implique } |U_{n+p} - U_n| \leq \varepsilon$

**** Exercice 4:** Soit une suite $(u_n)_{n \in \mathbb{N}}$, on considère l'expression formalisée

$$\underset{1 \uparrow}{(\forall \varepsilon > 0)} \underset{2 \uparrow}{(\forall N)} \underset{3 \uparrow}{(\exists n \geq N)} \underset{4 \uparrow}{\text{tels que}} |u_n - \lambda| \leq \varepsilon$$

Cette expression n'a pas de sens car λ n'y est pas défini.

- a) On ajoute $(\forall \lambda \in [0, 1])$ à l'un des 4 endroits marqués d'un $\uparrow$. On obtient 4 formules à priori différentes. Trouver, s'il en existe, des exemples de suites vérifiant chacune de ces 4 formules.
- b) On ajoute $(\exists \lambda \in [0, 1])$ à l'un des 4 endroits marqués d'un $\uparrow$. On obtient 4 formules à priori différentes. Trouver des exemples de suites vérifiant chacune de ces 4 formules (et distincts des précédents).

*** Exercice 5:** Ecrire au moyen de quantificateurs les assertions suivantes:

- a) La suite $(u_n)_{n \in \mathbb{N}}$ converge soit vers α soit vers $+\infty$.
- b) La suite $(u_n)_{n \in \mathbb{N}}$ ne converge ni vers α ni vers $+\infty$.
- c) La suite $(u_n)_{n \in \mathbb{N}}$ n'a pas de limite finie.
- d) La suite $(u_n)_{n \in \mathbb{N}}$ n'est pas bornée.

***** Exercice 6:** On considère une suite $(u_n)_{n \in \mathbb{N}}$ qui vérifie:

$$(\forall \varepsilon > 0) (\exists (v_n)_{n \in \mathbb{N}}) (\exists (w_n)_{n \in \mathbb{N}}) \text{ telles que } \begin{cases} (\forall n) u_n = v_n + w_n \\ (\forall \eta > 0) (\exists N) \text{ tel que } (\forall n \geq N) |v_n| \leq \eta \\ (\exists P) \text{ tel que } (\forall n \geq P) |w_n| \leq \varepsilon \end{cases}$$

Montrer que la suite $(u_n)_{n \in \mathbb{N}}$ converge vers 0.

**** Exercice 7:** On considère une suite $(u_n)_{n \in \mathbb{N}}$. Montrer que

$$(\forall \varepsilon > 0) (\exists N) (\forall n) ((n \geq N) \iff (0 \leq u_n \leq \varepsilon))$$

signifie que $(u_n)_{n \in \mathbb{N}}$ est décroissante et tend vers 0.

**** Exercice 8:** On considère une suite $(u_n)_{n \in \mathbb{N}}$. Montrer que

$$(\forall \varepsilon (0 < \varepsilon < 1)) (\exists N) \text{ tel que } (\forall n \geq N) \varepsilon \leq u_n \leq \frac{1}{\varepsilon}$$

signifie que $(u_n)_{n \in \mathbb{N}}$ converge vers 1.

* **Exercice 9:** On considère une suite $(u_n)_{n \in \mathbb{N}}$. Montrer que

$(\forall A > 0)$ l'ensemble $\{n \text{ tels que } |u_n| > A\}$ est fini

signifie que la suite $(u_n)_{n \in \mathbb{N}}$ converge vers 0.

Que signifie $(\exists A > 0)$ tel que l'ensemble $\{n \text{ tels que } |u_n| > A\}$ est fini ?

* **Exercice 10:** On considère une suite $(u_n)_{n \in \mathbb{N}}$ et un nombre λ . Montrer que

$(\forall A > 0)$ l'ensemble $\{n \text{ tels que } |u_n - \lambda| > A\}$ est fini

signifie que la suite $(u_n)_{n \in \mathbb{N}}$ converge vers λ .

* **Exercice 11:** Ecrire en termes de quantificateurs, que la fonction f est dérivable au point a . Ecrire en termes de quantificateurs, que la fonction f est dérivable en tout point de l'intervalle I .

L'intégrale définie

Historiquement la notion d'intégrale remonte au 17^{ème} siècle; elle fut introduite pour calculer des aires et des volumes. Les premiers calculs d'aires de domaines non polygonaux, et de volumes non polyédraux furent faits par Archimède (-287,-212): formules donnant l'aire d'une sphère, le volume d'une boule, calcul d'une valeur approchée de π ,... Le principe de ces calculs était le suivant, pour obtenir un encadrement de l'aire d'un domaine D , il suffit de calculer les aires d'un domaine polygonal D^+ qui contient D , et celle d'un domaine polygonal D^- contenu dans D . Dans les bons cas on peut trouver une suite de domaines D_n^+ (contenant D), et une suite de domaines D_n^- (contenus dans D), dont les aires forment deux suites adjacentes dont la limite commune est l'aire de D . Au 16^{ème} siècle et au début du 17^{ème} siècle, ces calculs d'aires, de volume,... avaient beaucoup occupé les mathématiciens. Chacun s'acharnait à appliquer l'antique méthode à des domaines de plus en plus compliqués, souvent sans aucun souci de rigueur. Les travaux de Leibniz sur les dérivées et les primitives (parus vers 1665^(*)) permettaient de retrouver par des calculs très simples, tous les résultats connus (obtenus avec beaucoup de peine), et d'en découvrir bien d'autres; c'est ce qui explique qu'ils eurent un succès immédiat. En fait la théorie de Leibniz se résumait à peu près aux deux résultats suivants, que vous avez appris dès la classe de terminale:

Pour toute fonction continue f ,
$$\int_a^b f(t) dt = F(b) - F(a), \text{ où } F \text{ est une primitive de } f.$$

Si f est positive,
$$\int_a^b f(t) dt \text{ est l'aire du domaine défini par les inégalités } a \leq x \leq b \text{ et } 0 \leq y \leq f(x).$$

On admettait alors que toute fonction f avait une primitive; on ne faisait même pas la restriction que f doit être continue; ce que vous saviez dès la classe de terminale. Et pour cause: on ne savait pas ce que c'est qu'une fonction continue! On se contentait de traiter un maximum d'exemples, en calculant des primitives de plus en plus savantes. La démarche historique ressemble décidément fort à l'apprentissage qui fut le vôtre.

Ce n'est qu'au début du 19^{ème} siècle qu'on se pencha sur les problèmes liés à l'existence de l'intégrale. Il serait trop long d'expliquer ici les raisons pour les quelles on se mit tout à coup à douter de la validité d'énoncés que l'on considérait jusqu'alors comme évidents. Disons simplement qu'au cours du 18^{ème} siècle on se mit à définir des fonctions de plus en plus compliquées (en prenant des limites de suites de fonctions, ce que nous apprendrons à faire un peu plus tard) dont il était difficile d'imaginer (et encore plus de tracer) le graphique. Pour de telles fonctions il était hasardeux de parler du domaine défini par les inégalités $a \leq x \leq b$ et $0 \leq y \leq f(x)$. La première définition sérieuse de l'intégrale est due à Riemann (1826-1866). L'un des buts du présent chapitre est de décrire l'outil essentiel de cette définition, il s'agit des sommes de Riemann.

(*) Newton publiait presque simultanément un mémoire sur la mécanique, où il introduisait lui aussi la notion de dérivée (c'est à dire la notion de vitesse instantanée). Il s'ensuivit l'une des plus ridicules querelles d'antériorité que la science ait connue.

Remarquons encore que la simili-définition que l'on vous a donnée en terminale est un leurre. Pour démontrer l'existence des primitives, on utilise la notion d'intégrale, et non le contraire comme on a voulu vous le faire croire. Cette inversion des rôles est beaucoup plus qu'un artifice de procédure; c'est elle qui nous permettra d'élargir le champ d'application de la théorie, et de parler de l'intégrale de (certaines) fonctions non continues, et qui n'ont pas de primitives.

§ 1 : Les propriétés fondamentales d'une théorie de l'intégrale.

Qu'est ce qu'une théorie de l'intégrale ?

La théorie de l'intégrale est probablement, à vos yeux, la manipulation du symbole $\int_a^b f(t) dt$; où a et b sont deux nombres ($a < b$ ou $a > b$), et f une fonction continue sur le segment $[a, b]$. Nous allons adopter un point de vue un peu différent. Nous fixerons un intervalle $I =]a, b[$ (ou $[a, b]$, ou $[a, b[$, ou $]a, b]$); il est donc ouvert, ou fermé ou semi ouvert; a et b sont finis ou infinis. Nous nous intéresserons aux fonctions (quelconques a priori) de I dans $\mathbb{R}$. Nous voulons associer à une telle fonction f , un nombre que nous noterons $\int_I f(t) dt$ (ou encore $\int_{[a,b]} f(t) dt$). Ce nombre est, pourvu que $a < b$, ce que vous notez $\int_a^b f(t) dt$ (voir le §4). Et nous voulons que les deux propriétés (bien connues !) suivantes soient vérifiées.

$$(\alpha) \begin{cases} \int_I [f(t) + g(t)] dt = \int_I f(t) dt + \int_I g(t) dt \\ \int_I \lambda f(t) dt = \lambda \int_I f(t) dt \end{cases}$$

$$(\beta) \text{ Si } f \leq g, \text{ alors } \int_I f(t) dt \leq \int_I g(t) dt.$$

Remarquons tout de suite qu'il serait déraisonnable de penser que toutes les fonctions sont intégrables; en effet si $I = [0, \infty[$, et si f est la constante 1, il n'est pas sérieux d'espérer que l'aire comprise entre la courbe $y = f(x)$ et l'axe Ox soit finie. L'intégrale sera donc définie seulement pour les fonctions d'un sous-ensemble de $\mathcal{F}(I, \mathbb{R})$, et pour que la propriété (α) ci-dessus ait une signification, ce sous-ensemble doit être un sous-espace vectoriel. Ce qui nous amène à poser la définition suivante:

Une théorie de l'intégrale sur l'intervalle I est la donnée :

1) d'un sous-espace vectoriel $\mathcal{F}$ de $\mathcal{F}(I, \mathbb{R})$, dont les éléments sont appelés les "fonctions intégrables".

2) d'une forme linéaire $\mathcal{F} \rightarrow \mathbb{R}$ (on note $\int f(t) dt$ l'image de f par cette forme).

Ces données devant vérifier :

$$a) \text{ Si } f \leq g \text{ (i.e.: si } (\forall t) f(t) \leq g(t)), \text{ alors } \int f(t) dt \leq \int g(t) dt.$$

b) Quels que soient u et v (où $u \leq v$) la fonction caractéristique^(*) de l'intervalle (ouvert, fermé, ou semi-ouvert) d'extrémités u et v est intégrable, d'intégrale $v - u$.

(*) Rappelons que la fonction caractéristique du sous-ensemble X de $\mathbb{R}$, est la fonction qui vaut 1 en tous les points de X et 0 en tous les points qui ne sont pas dans X .

Nous esquisserons au §2 la construction d'une théorie de l'intégrale: c'est l'intégrale de Riemann. Les fonctions continues sur $[a, b]$ sont intégrables au sens de Riemann. Les fonctions monotones le sont aussi.

Comment majorer les intégrales? La condition (a) permet de majorer des intégrales. Voici les deux principales règles de majoration que nous emploierons:

* Si I est borné de longueur $\mathcal{L}(I)$, si f est intégrable sur I et si $(\forall t \in I) \ m \leq f(t) \leq M$, alors :

$$m \mathcal{L}(I) \leq \int_I f(t) dt \leq M \mathcal{L}(I)$$

* On a $(\forall t) \ -|f(t)| \leq f(t) \leq |f(t)|$, donc si f est intégrable ainsi que $|f|$:

$$-\int_a^b |f(t)| dt \leq \int_a^b f(t) dt \leq \int_a^b |f(t)| dt$$

Soit encore :

$$\left| \int_a^b f(t) dt \right| \leq \int_a^b |f(t)| dt$$

Nous retiendrons : La valeur absolue de l'intégrale est au plus égale à l'intégrale de la valeur absolue.

Intégrales à valeurs dans un espace vectoriel de dimension finie

Soit E un espace vectoriel de dimension finie, et $\{e_1, \dots, e_n\}$ une base de E . Soit f une fonction de $[a, b]$ dans E ; notons $f_1, \dots, f_n$ ses fonctions coordonnées dans cette base, c'est à dire les fonctions

numériques telles que $(\forall t) f(t) = \sum_{i=1}^n f_i(t) e_i$.

Si les f_i sont intégrables, nous dirons que f est intégrable, et nous poserons :

$$\sum_{i=1}^n \left[\int_I f_i(t) dt \right] e_i = \int_I f(t) dt$$

Notons que cette définition utilise une base particulière de E , mais que le résultat obtenu ne dépend pas de la base utilisée. En effet si $\{\epsilon_1, \dots, \epsilon_n\}$ est une autre base, et si $e_i = \sum_j a_j^i \epsilon_j$, alors :

$$f(t) = \sum_i f_i(t) e_i = \sum_i f_i(t) \left[\sum_j a_j^i \epsilon_j \right] = \sum_j \left[\sum_i a_j^i f_i(t) \right] \epsilon_j$$

Les fonctions coordonnées de f dans la base $\{\epsilon_1, \dots, \epsilon_n\}$, sont les $\phi_j = \sum_i a_j^i f_i(t)$; ce sont des

combinaisons linéaires des f_i ; donc elles sont intégrables dès que les f_i le sont. De plus :

$$\sum_j \left[\int_I \phi_j(t) dt \right] \epsilon_j = \sum_j \left[\int_I \sum_i a_j^i f_i(t) dt \right] \epsilon_j = \sum_j \left[\sum_i a_j^i \left[\int_I f_i(t) dt \right] \right] \epsilon_j$$

$$\sum_j \left[\int_I \phi_j(t) dt \right] \epsilon_j = \sum_i \left[\int_I f_i(t) dt \right] \sum_j a_j^i \epsilon_j = \sum_i \left[\int_I f_i(t) dt \right] e_i$$

Ce qui prouve que la formule qui définit l'intégrale de f , a la même valeur dans la base $\{e_i\}$ et dans la base $\{\epsilon_j\}$.

Remarque: Si E est euclidien, et si f et $\|f\|$ sont intégrables, nous aurons la formule de majoration:

$$\left\| \int_I f(t) dt \right\| \leq \int_I \|f(t)\| dt$$

En effet: Si $\int_I f(t) dt = 0$, il n'y a rien à démontrer. Sinon choisissons une base orthonormée

$\{\varepsilon_1, \dots, \varepsilon_n\}$ telle que ε_1 soit colinéaire à $\int_I f(t) dt$. Alors:

$$\left\| \int_I f(t) dt \right\| = \left| \int_a^b f_1(t) dt \right| \leq \int_a^b |f_1(t)| dt \leq \int_a^b \|f(t)\| dt$$

On retiendra : La norme euclidienne de l'intégrale, est au plus égale à l'intégrale de la norme euclidienne.

§ 2 : L'intégrale de Riemann

Dans ce qui suit nous considérons des fonctions de I dans $\mathbb{R}$ qui sont bornées (f est bornée s'il existe m et M tels que $(\forall t \in I): m \leq f(t) \leq M$). Si I n'est pas un intervalle borné, nous ne considérerons que les fonctions à support borné (f est à support borné s'il existe deux nombres α et β tels que f soit nulle en dehors de l'intervalle $[\alpha, \beta]$).

Si l'on note c la fonction caractéristique de cet intervalle, f est minorée par mc et majorée par Mc ; ainsi ces deux conditions servent à assurer que f est minorée et majorée par des fonctions intégrables simples.

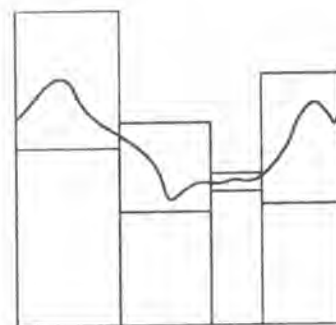
La définition de Riemann est un curieux retour aux sources. Archimède avait calculé des aires en encadrant les domaines par des domaines polygonaux; Leibniz s'était affranchi des lourdeurs de cette méthode, en utilisant l'intégrale (qu'il ne définissait pas !); Riemann définit l'intégrale d'une fonction positive f sur $[a, b]$ (c'est à dire l'aire du domaine $\mathcal{D}(a \leq x \leq b \text{ et } 0 \leq y \leq f(x))$, en introduisant:

d'une part des domaines qui sont réunions de rectangles et qui contiennent $\mathcal{D}$.

d'autre part des domaines qui sont réunions de rectangles et qui sont contenus dans $\mathcal{D}$.

Bien entendu nous ne voulons pas nous restreindre aux fonctions positives; il ne faut donc pas que nous parlions en termes d'aires; et dans tout ce qui suit il n'en sera pas question. Toutefois le lecteur qui voudra comprendre, aura intérêt à faire des dessins dans les quels f est positive, et à traduire les inégalités en termes d'aires de domaines contenus les uns dans les autres.

Les fonctions en escalier: Une fonction f définie sur I est appelée une fonction en escalier, si l'on peut partager f en (un nombre fini d') intervalles J_α (fermés, ouverts, ou semi-ouverts), sur chacun desquels f est constante. Nous ne nous intéresserons qu'aux fonctions en escaliers à support borné. Une telle fonction est somme finie de fonctions de la forme $\lambda_\alpha \chi_\alpha$, où les χ_α sont les fonctions caractéristiques des intervalles J_α ; elle est donc nécessairement dans l'espace $\mathcal{F}$ des fonctions intégrables (que nous voulons définir), et son intégrale est $\sum_\alpha \lambda_\alpha \mathcal{L}(J_\alpha)$.



Sommes de Riemann: Soit maintenant une fonction f bornée et à support borné. Notons $\mathcal{Z}(f, +)$ l'ensemble des fonctions en escalier à support borné, qui majorent f ; notons $\mathcal{Z}(f, -)$ l'ensemble des fonctions en escalier à support borné qui minorent f . Ces deux ensembles sont non vides; en effet si $(\forall t) m \leq f(t) \leq M$ et si f est nulle en dehors de l'intervalle borné J , la fonction $m\chi_J$ est dans $\mathcal{Z}(f, -)$, et la

fonction $M_X J$ est dans $\mathcal{Z}(f, +)$. Notons que c'est ici, et ici seulement, que nous utilisons le fait que f est bornée et à support borné. Les intégrales des éléments de $\mathcal{Z}(f, +)$ sont appelées les sommes de Riemann supérieures de f , les intégrales des éléments de $\mathcal{Z}(f, -)$ sont les sommes de Riemann inférieures de f .

La condition (a) (cf: §1) implique clairement que si f est intégrable, son intégrale est au plus égale aux sommes de Riemann supérieures, et au moins égale aux sommes de Riemann inférieures.

La définition de Riemann:

Soit φ un élément de $\mathcal{Z}(f, -)$, toute fonction de $\mathcal{Z}(f, +)$ est supérieure à φ ; donc l'intégrale de φ est un minorant de l'ensemble X^+ des sommes de Riemann supérieures. Il en résulte que X^+ est un ensemble de nombres réels minoré, nous savons que ceci implique que X^+ a une borne inférieure*; celle-ci est appelée l'intégrale supérieure de f , on la notera $I(f, +)$. De même l'ensemble X^- des sommes de Riemann inférieures est majoré (par une quelconque des sommes de Riemann supérieures), donc il a une borne supérieure $I(f, -)$ appelée l'intégrale inférieure de f .

Puisque tout élément de X^+ est supérieur ou égal à tout élément de X^- , nous avons $I(f, +) \geq I(f, -)$.

Par définition, la fonction f sera dite intégrable au sens de Riemann, si $I(f, +) = I(f, -)$. Et ce nombre est l'intégrale de f , nous le noterons $\int_I f(t) dt$.

On démontre:

Que, si I est fermé et borné, toute fonction continue sur I est intégrable sur I . Si I n'est pas borné, toute fonction continue sur I et qui est nulle en dehors d'un intervalle borné, est intégrable sur I (voir §3).

Que si I est borné, toute fonction bornée monotone sur I , est intégrable sur I (cf: Ex 2).

Que les fonctions intégrables sur I forment un espace vectoriel $\mathcal{F}$, et que l'intégrale est une forme linéaire.

Que l'intégrale est croissante.

Ainsi le programme que nous nous étions fixé est rempli; nous avons défini une 'théorie de l'intégrale'.

Remarque: Pour que f soit intégrable, il suffit que

$$(\forall \varepsilon > 0) (\exists f_+ \text{ et } f_- \text{ en escalier}) \text{ telles que } f_- \leq f \leq f_+ \text{ et } \int (f_+ - f_-)(t) dt \leq \varepsilon.$$

(Démonstration facile).

§ 3 : L'intégrabilité des fonctions continues.

Montrons d'abord que si I est fermé et borné les fonctions C^1 sont intégrables: Soit donc f de classe C^1 sur $I = [\alpha, \beta]$; $|f'|$ est bornée sur $[\alpha, \beta]$ (car continue), c'est à dire qu'il existe K tel que $(\forall t \in [\alpha, \beta]) |f'(t)| \leq K$.

Soit $\varepsilon > 0$; nous allons construire deux fonctions en escalier, f_+ et f_- , telles que $f_- \leq f \leq f_+$ et que $\int [f_+ - f_-](t) dt \leq \varepsilon$ (ceci démontrera que f est intégrable; cf: fin du §2).

Pour cela choisissons une partition de I en intervalles de longueur au plus $\eta = \frac{\varepsilon}{2K|\beta - \alpha|}$. Soit J l'un de ces intervalles et x un point de J . D'après le théorème des accroissements finis, quelque soit t dans J , $|f(x) - f(t)| \leq K \ell(J) \leq K\eta \leq \frac{\varepsilon}{2|\beta - \alpha|}$. Il en résulte que

* Nous utilisons ici le fait que tout ensemble de nombres réels minoré, a une borne inférieure (voir chapitre 8)

$$(\forall t \in J) \quad f(x) - \frac{\varepsilon}{2|\beta-\alpha|} \leq f(t) \leq f(x) + \frac{\varepsilon}{2|\beta-\alpha|}.$$

La fonction f_+ vaudra $f(x) + \frac{\varepsilon}{2|\beta-\alpha|}$ sur J , tandis que la fonction f_- vaudra $f(x) - \frac{\varepsilon}{2|\beta-\alpha|}$ sur J . Il est clair que $f_- \leq f \leq f_+$, et que $f_+ - f_-$ est la fonction constante $2 \frac{\varepsilon}{2|\beta-\alpha|}$; donc

$$\int_1 [f_+ - f_-](t) dt \leq 2 \frac{\varepsilon}{2|\beta-\alpha|} |\beta-\alpha| = \varepsilon.$$

Extension aux fonctions continues:

L'hypothèse que f est C^1 est trop forte. Nous l'avons utilisée pour construire (via le théorème des accroissements finis) un nombre $\eta(>0)$ tel que

$$|x-t| \leq \eta(J) \leq \eta \text{ implique } |f(x) - f(t)| \leq \frac{\varepsilon}{2|\beta-\alpha|}.$$

Nous aurions pu obtenir ceci par un argument de continuité uniforme. Rappelons qu'une fonction f est uniformément continue si:

$$(\forall \mu > 0) (\exists \eta > 0) \text{ tel que } (\forall t) (\forall x) \quad |t-x| \leq \eta \text{ implique } |f(t) - f(x)| \leq \mu.$$

Rappelons aussi que toute fonction continue sur un intervalle fermé et borné est uniformément continue. Nous reviendrons sur ces questions au chapitre 10.

Ainsi nous pouvons reprendre la démonstration ci-dessus. La fonction f est uniformément continue sur $[\alpha, \beta]$, puisqu'elle est continue. Nous appliquerons cette uniforme continuité en prenant $\mu = \frac{\varepsilon}{2|\beta-\alpha|}$, ce qui nous donnera η , et tout le reste est inchangé.

Remarquons que l'argument de continuité uniforme est, dans le cas des fonctions C^1 , remplacé par une utilisation du théorème des accroissements finis. Ou plutôt, le fait que la dérivée soit majorée, donne - via les accroissements finis - une majoration beaucoup plus précise de $|f(t) - f(x)|$.

§ 4 : Quelques fonctions intégrables non continues.

L'espace vectoriel $\mathcal{F}$ des fonctions intégrables contient outre les fonctions continues que nous avons envisagées au §3, bon nombre de fonctions non continues. Certaines sont fort compliquées. Celles que nous rencontrerons sont fort simples; elles sont de deux types:

Les fonctions "continues-avec-sauts": Ce sont les fonctions (bornées et nulles en dehors d'un intervalle $[\alpha, \beta]$ contenu dans I) qui sont continues sauf en un point (ou plus généralement en un nombre fini de points) de I , et qui en ce point (ou en ces points) ont une limite à droite et une limite à gauche. Ces fonctions sont intégrables (on ne le démontrera pas, mais on va voir sur un exemple, que c'est bien clair).

Exemple: La fonction f définie (sur $[-1, +1]$) par $f(x) = x^2 + 2x$ pour $x > 0$, et $f(x) = 1$ pour $x \leq 0$. Elle est la somme de la fonction g définie par $g(x) = f(x)$ si $x > 0$ et $g(x) = 0$ si $x \leq 0$ (qui est continue donc intégrable), et de la fonction en escalier, qui vaut 0 si $x > 0$, et 1 si $x \leq 0$. Donc elle est intégrable.

Les fonctions négligeables: On nomme ainsi les fonctions f dont la valeur absolue $|f|$ est intégrable, d'intégrale nulle. Les fonctions négligeables sont intégrables d'intégrale nulle (mais bien sûr il existe des fonctions intégrables d'intégrale nulle et qui ne sont pas négligeables).

Exercice a: 1) Soit f une fonction négligeable positive, et soit g telle que $0 \leq g \leq f$, montrer que g est négligeable (pour cela on montrera que $0 \leq g \leq f$ implique $0 \leq I(g, -) \leq I(g, +) \leq I(f, +)$).

2) En déduire que si f est négligeable, f est intégrable d'intégrale nulle (on appliquera (1) à $\frac{f+|f|}{2}$ et à $\frac{f-|f|}{2}$).

La plus simple des fonctions négligeables est la fonction χ_x qui vaut 1 en un point x et qui est nulle partout ailleurs.

Exercice b: La fonction (en escalier) partout nulle minore χ_X ; qu'en déduire pour $I(\chi_X, -)$? Montrer que, quelque soit $\varepsilon (>0)$ il existe une fonction en escalier f qui majore χ_X et dont l'intégrale est ε . Qu'en déduire pour $I(\chi_X, +)$? Montrer que χ_X est négligeable.

Les fonctions négligeables forment un sous-espace vectoriel de $\mathcal{F}$. Cet espace vectoriel contient toutes les fonctions nulles sauf en un nombre fini de points.

Exercice c: Comparer $|f+g|$ à $|f|+|g|$, et (en utilisant la première question de l'exercice a) montrer que la somme de deux fonctions négligeables, est négligeable. En déduire que les fonctions négligeables forment un espace vectoriel.

Ensembles quarrables au sens de Riemann, ensembles négligeables: Nous dirons qu'un sous-ensemble X de $\mathbb{R}$ est quarrable au sens de Riemann, si sa fonction caractéristique χ_X est intégrable. L'intégrale de χ_X est appelée la mesure de X . Les intervalles bornés (ouverts, fermés ou semi ouverts) sont quarrables, puisque leurs fonctions caractéristiques sont des fonctions en escalier. Leur mesure est leur longueur. Pour un exemple d'ensemble non quarrable voir l'exercice 1.

Exercice d: Soit A et B deux ensembles quarrables. Montrer que $A \cap B$ est quarrable si et seulement si $A \cup B$ est quarrable; quelle relation y a-t-il entre les mesures de ces 4 ensembles?

Un ensemble X est dit négligeable, s'il est quarrable et de mesure nulle; autrement dit si sa fonction caractéristique est négligeable. Les ensembles finis sont négligeables. Ce ne sont pas les seuls (cf: exercice 10). Toute fonction bornée qui est nulle en dehors d'un ensemble négligeable, est une fonction négligeable.

Exercice e: Soit X un ensemble négligeable, et f une fonction bornée (on supposera que $(\forall t) |f(t)| \leq M$) qui est nulle en tout point qui n'appartient pas à X . En utilisant la première question de l'exercice a, montrer que f est négligeable.

Si f est intégrable et si g est négligeable, alors $f+g$ est intégrable (comme somme de deux fonctions intégrables), et a même intégrale que f . Si f est intégrable, et si une fonction bornée h coïncide avec f en dehors d'un ensemble négligeable X , alors h est intégrable, et a même intégrale que f (car $h = f+g$, où g est bornée et nulle en dehors de X). Ainsi, par exemple, modifions une fonction intégrable f en un nombre fini de points, nous obtenons une nouvelle fonction intégrable, qui a même intégrale que f . En particulier si f est une fonction "continue-avec-sauts", vis à vis de l'intégrale, nous pouvons oublier ses valeurs en les points de discontinuité; en ces points nous pouvons la rendre continue à droite, ou continue à gauche, ou..., nous aurons toujours une fonction intégrable, et toujours la même intégrale.

§ 5 : Retour aux primitives.

Restriction de l'intervalle d'intégration: Soit une fonction f intégrable sur I , et un intervalle J contenu dans I , alors la restriction $f|_J$ de f à J est une fonction intégrable sur J . D'après la remarque qui termine de §2, pour nous en convaincre, il nous faut, quel que soit ε , construire deux fonctions (sur J) en escalier g_+ et g_- telles que $g_- \leq f|_J \leq g_+$, et que $\int_J [g_+ - g_-](t) dt \leq \varepsilon$. Or puisque f est intégrable (sur I)

il existe deux fonctions (sur I) en escalier f_+ et f_- telles que $f_- \leq f \leq f_+$, et que $\int_I [f_+ - f_-](t) dt \leq \varepsilon$. Nous

prendrons donc pour g_+ et g_- les restrictions à J de f_+ et f_- .

Retour à des notations plus classiques: Soit f une fonction définie sur un intervalle I , et intégrable sur I . Soit a et b deux points de I , tels que $a < b$. La fonction $f|_{[a,b]}$ est intégrable, son intégrale est - classiquement - notée de deux façons:

$$\int_{[a,b]} f(t) dt = \int_a^b f(t) dt = - \int_b^a f(t) dt$$

Avec ces notations nous avons la classique relation de Chasles:

$$(\forall a) (\forall b) (\forall c) \quad \int_a^b f(t) dt + \int_b^c f(t) dt = \int_a^c f(t) dt$$

Primitives d'une fonction continue: Fixons a , pour tout x dans I , nous pouvons poser

$\varphi(x) = \int_a^x f(t) dt$, nous obtenons une application de I dans $\mathbb{R}$. Si f est une fonction continue, φ est une fonction dérivable, et $(\forall t) \varphi'(t) = f(t)$.

Pour démontrer ce résultat, écrivons:

$$\begin{aligned} \varphi(t+h) - \varphi(t) - h f(t) &= \int_a^{t+h} f(\tau) d\tau - \int_a^t f(\tau) d\tau - h f(t) \\ &= \int_t^{t+h} f(\tau) d\tau - \int_t^{t+h} f(t) d\tau = \int_t^{t+h} [f(\tau) - f(t)] d\tau \end{aligned}$$

Quelque soit $\varepsilon (> 0)$ nous pouvons trouver η tel que $|\tau - t| \leq \eta$ implique $|f(\tau) - f(t)| \leq \varepsilon$ (c'est la continuité de f). Si $|h| \leq \eta$, pour tout τ dans l'intervalle $[t, t+h]$ nous aurons $|f(\tau) - f(t)| \leq \varepsilon$; et par conséquent:

$$|\varphi(t+h) - \varphi(t) - h f(t)| = \left| \int_t^{t+h} [f(\tau) - f(t)] d\tau \right| \leq \int_t^{t+h} |f(\tau) - f(t)| d\tau \leq \int_t^{t+h} \varepsilon d\tau = \varepsilon |h|.$$

Ce qui démontre que φ est dérivable en t , de dérivée $f(t)$.

C'est le lien fondamental entre primitivation et intégration; puisque deux primitives d'une même fonction diffèrent d'une constante, on en déduit que, pour toute primitive F de f , on a: $\int_a^b f(t) dt = F(b) - F(a)$.

Et nous avons ainsi fait le lien avec l'idée originale de Leibniz.

Dés lors pour calculer l'intégrale d'une fonction simple f deux stratégies sont possibles. L'une consiste à calculer une primitive de f . L'autre consiste à calculer des sommes de Riemann de f , pour obtenir un encadrement du résultat; ou plus généralement à calculer l'intégrale d'une fonction f^+ qui majore f , et l'intégrale d'une fonction f^- qui minore f , en s'arrangeant pour que l'intégrale de $f^+ - f^-$ soit suffisamment petite (cf: méthode des trapèzes).

§ 6 : Les principales méthodes du calcul des primitives

La formule d'intégration par parties:

$$\text{Elle s'écrit : } \int u(t)v'(t) dt = u(t)v(t) - \int u'(t)v(t) dt.$$

Exercice f: Calculer $\int t^n e^{\alpha t} dt$. Calculer $\int t^n \cos \beta t dt = \Re e \left(\int t^n e^{i\beta t} dt \right)$.

* Observer la substitution de $\int_t^{t+h} f(\tau) d\tau$ à $h f(t)$. Astuce classique !

La formule de changement de variable

Si $F(t) = \int f(t) dt$ f est continue et si φ est de classe C^1 , alors $F(\varphi(\theta)) = \int f(\varphi(\theta)) \varphi'(\theta) d\theta$.

Exercice g: Calculer $\int \frac{dt}{(1+t^2)^2}$ en faisant le changement de variable $\theta = \text{Arctg } t$.

Les fractions rationnelles:

Les primitives d'une fraction rationnelle s'expriment au moyen des fractions rationnelles et des fonctions Log et Arctg. Le calcul de ces primitives comporte deux étapes.

Calcul des éléments simples: A toute fraction $\frac{P(t)}{Q(t)}$, on associe :

Sa partie entière: c'est le quotient de la division euclidienne de P par Q.

Les éléments simples relatifs aux racines réelles de Q. A une racine α d'ordre k, est associée une expression du type:

$$\pi_{\alpha}(t) = \frac{a_k}{(t-\alpha)^k} + \dots + \frac{a_1}{(t-\alpha)}$$

appelée partie polaire de $\frac{P}{Q}$ en α . On notera que $\frac{P(t)}{Q(t)} - \pi_{\alpha}(t)$ est une fonction qui se prolonge par continuité en α . C'est à dire que $\pi_{\alpha}(t)$ est formé des termes de degré négatif du développement de Laurent en α de $\frac{P(t)}{Q(t)}$.

Les éléments simples correspondant aux racines non réelles de Q. A tout facteur $(t^2+at+b)^k$ (où $a^2-4b < 0$) de la décomposition de Q en facteurs irréductibles de Q, on associe une partie polaire du type:

$$\frac{\alpha_k + t\beta_k}{(t^2+at+b)^k} + \dots + \frac{\alpha_1 + t\beta_1}{(t^2+at+b)}$$

La fraction rationnelle est la somme de ses parties polaires, et de sa partie entière.

Calcul des primitives des éléments simples: Seuls les termes en $\frac{\alpha+t\beta}{(t^2+at+b)^r}$ posent problème. Par changement de variable linéaire ($\theta = ut+v$) on se ramène à $\frac{u'\theta+v'}{(1+\theta^2)^r}$, dont les primitives se calculent, par exemple, en faisant de changement de variable $\varphi = \text{Arctg } \theta$ (cf: Ex g).

Exercice h: Calculer une primitive de $\frac{1-t^7}{(1-t)^2(1+t^2)^2}$.

Les polynômes trigonométriques: Pour calculer une primitive de $P(\cos t, \sin t)$: au moyen des formules:

$$\cos t = \frac{e^{it} + e^{-it}}{2} \quad \text{et} \quad \sin t = \frac{e^{it} - e^{-it}}{2i}$$

on transforme P en une somme de termes de la forme e^{ikt} . Puis en utilisant les formules :

$$e^{ikt} = \cos kt + i \sin kt \quad \text{et} \quad e^{-ikt} = \cos kt - i \sin kt$$

on obtient une somme de termes en $\cos kt$ et $\sin kt$ ($k = 1, \dots$).

Exercice i: Calculer $\int \cos^5 t dt$ et $\int \cos^4 t \sin^2 t dt$.

Les fractions rationnelles en sin et cos: $F(t) = \frac{P(\cos t, \sin t)}{Q(\cos t, \sin t)}$ est une fraction rationnelle en $\text{tg}(t/2)$. Mais dans bien des cas on peut utiliser des changements de variables donnant des calculs plus simples:

Si $F(t) = G(t) \sin t$, où G est une fonction paire, le changement de variable $t = \arccos \theta$ nous ramènera à une fraction rationnelle.

Si $F(t) = G(t) \cos t$, où $G(t) = G(\pi - t)$, le changement de variable $t = \arcsin \theta$ nous ramènera à une fraction rationnelle.

Exercice j: Calculer $\int \frac{\cos t + \cos 2t}{\sin t + \sin 2t} dt$

Les fonctions comportant des racines carrées de polynômes du second degré.

Par un changement de variable linéaire on se ramène au cas des racines suivantes :

$\sqrt{1-t^2}$ On fait le changement de variable $\theta = \arccos t$ (ou $\theta = \arcsin t$, ou $t = \sin \theta$)

$\sqrt{t^2-1}$ On fait le changement de variable $\theta = \operatorname{Argch} t$

$\sqrt{t^2+1}$ On fait le changement de variable $\theta = \operatorname{Argsh} t$ (ou $\theta = \operatorname{Arctg} t$)

Exercice k: Calculer des primitives de $\frac{t}{\sqrt{1+t+t^2}}$, $\frac{t}{\sqrt{1-3t+t^2}}$ et $\frac{t}{\sqrt{-1+3t-t^2}}$

Complément 1: Utilisation de l'intégrale dans le calcul des grandeurs physiques

Comment on passe d'une formule algébrique à une intégrale

Supposons un fil de résistance R , parcouru par un courant d'intensité I constante. Alors l'énergie calorifique dissipée entre les instants a et b , est $W = R I^2 (b-a)$. Si I est fonction continue du temps cette même énergie sera donnée par la formule $W = \int_a^b R I^2(t) dt$. Comment passe-t-on de la première à la

seconde formule ? Il y a deux méthodes; l'une consiste à démontrer que la fonction qui à t associe l'énergie $W(t)$ dissipée entre les instants a et t , est une primitive de $R I^2(t)$; l'autre consiste à écrire les sommes de Riemann.

Première méthode: Pour $a \leq t \leq t+h \leq b$, $W(t+h) - W(t)$ est l'énergie dissipée entre les instants t et $t+h$. Si $J_m(h) = \inf_{\theta \in [t, t+h]} I^2(\theta)$ et $J_M(h) = \sup_{\theta \in [t, t+h]} I^2(\theta)$, on a :

$$h R J_m(h) \leq W(t+h) - W(t) \leq h R J_M(h)$$

En effet "plus I^2 est grand, plus l'énergie dissipée est grande".

Mais puisque I^2 est une fonction continue en t , nous savons que, quel que soit $\varepsilon > 0$, il existe $\eta > 0$ tel que $(h < \eta)$ implique que pour θ appartenant à $[t-h, t+h]$:

$$I^2(t) - \frac{\varepsilon}{R} \leq I^2(\theta) \leq I^2(t) + \frac{\varepsilon}{R}$$

Donc: $I^2(t) - \frac{\varepsilon}{R} \leq \inf_{\theta \in [t, t+h]} I^2(\theta) = J_m(h)$ et $J_M(h) = \sup_{\theta \in [t, t+h]} I^2(\theta) \leq I^2(t) + \frac{\varepsilon}{R}$

Et ainsi $(h < \eta)$ implique:

$$h R I^2(t) - h \varepsilon \leq W(t+h) - W(t) \leq h R I^2(t) + h \varepsilon$$

Soit :

$$-\varepsilon \leq \frac{W(t+h) - W(t) - h R I^2(t)}{h} \leq \varepsilon$$

Ce qui prouve que $R I^2(t)$ est la dérivée à droite de W en t . Il faudrait reprendre le raisonnement entre t et $t-h$, pour montrer que c'est aussi la dérivée à gauche.

Bien sûr dans un cours de physique on n'insiste pas sur ces arguments de continuité, et on déclare seulement: lorsque h est assez petit, on peut faire comme si I était constante et égale à $I(t)$; donc $W(t+h) - W(t)$ est à peu près égal à $R I^2(t) h$.

Seconde méthode : Soit J^2 une fonction en escalier (positive) définie sur $[a, b]$ et supérieure à I^2 . Alors puisque "plus I^2 est grand, plus l'énergie dissipée est grande", l'énergie dissipée par un courant

d'intensité J , est supérieure à l'énergie dissipée par I . Il en résulte que $W \leq R \int_a^b J^2(t) dt$. Et puisque ceci est vrai pour toute fonction en escalier plus grande que I^2 , nous aurons $W \leq R I(I^2, +)$.

Un raisonnement analogue nous montre que $W \geq R I(I^2, -)$. Et par conséquent si I^2 est intégrable au sens de Riemann, nous aurons $W = \int_a^b R I^2(t) dt$.

Bien sûr dans un cours de physique, on se contentera de dire 'approchons I par une fonction en escalier J , pour cette fonction la formule est vraie. Et puisque'il existe des fonctions en escalier aussi proches que l'on veut de I , on peut passer à la limite'.

On notera que les deux méthodes utilisent le même principe physique: 'Plus I^2 est grand, plus l'énergie dissipée est grande'. Mais les arguments simplifiés que l'on emploie habituellement passent ce principe sous silence.

Moyenne d'une fonction sur un intervalle: Nous allons associer à chaque fonction f définie sur un intervalle $[a, b]$ (ou du moins à certaines d'entre elles) un nombre $\mu(f, a, b)$ que nous appellerons la moyenne de f sur l'intervalle $[a, b]$. Nous voulons que cette construction vérifie les conditions suivantes:

$$1) \text{ si } a < b < c, \text{ alors : } \mu(f, a, c) = \frac{1}{c-a} [(b-a) \mu(f, a, b) + (c-b) \mu(f, b, c)]$$

Autrement dit $\mu(f, a, c)$ est la moyenne pondérée (ou barycentre) de $\mu(f, a, b)$ et de $\mu(f, b, c)$, les coefficients de pondération étant les longueurs des intervalles.

$$2) \text{ Si } f \leq g, \text{ alors } \mu(f, a, b) \leq \mu(g, a, b).$$

$$3) \text{ Si } f \text{ est constante et égale à } m \text{ sur } [a, b], \text{ alors } \mu(f, a, b) = m.$$

Supposons définie une telle notion, et essayons de calculer $\beta(f, a, b) = (b-a)\mu(f, a, b)$. D'après le principe de majoration 2) ci-dessus, pour toute fonction en escalier g supérieure à f , nous devons avoir $\beta(g, a, b) \geq \beta(f, a, b)$. Or d'après 1) et 3), $\beta(g, a, b) = \int_a^b g(t) dt$. Donc $\beta(f, a, b)$ est inférieur aux intégrales des fonctions en escalier supérieures à f ; autrement dit $\beta(f, a, b) \leq I(f, +)$. De même $\beta(f, a, b) \geq I(f, -)$. Donc, si f est intégrable au sens de Riemann, $\beta(f, a, b)$ est l'intégrale de f sur $[a, b]$.

C'est pourquoi on définit la moyenne de f (lorsque f est intégrable) par :

$$\mu(f, a, b) = \frac{1}{b-a} \int_a^b f(t) dt$$

On vérifie facilement que les conditions (1), (2) et (3) sont vérifiées. De plus :

4) Si f est continue sur $[a, b]$, et si $m = \inf_{t \in [a, b]} f(t)$ et $M = \sup_{t \in [a, b]} f(t)$, alors (d'après (2)) la moyenne μ de f sur $[a, b]$ est comprise entre m et M . Donc (th des valeurs intermédiaires) il existe (au moins) un nombre t_0 dans $[a, b]$ tel que $f(t_0) = \mu$. Nous obtenons ainsi un résultat connu sous le nom de "théorème de la moyenne": *la moyenne de f sur l'intervalle $[a, b]$ est (si f est continue) la valeur de f en un point de l'intervalle.*

Moyennes quadratiques: Revenons à la situation physique du début du paragraphe. Soit $\mu(I)$ la moyenne de $I(t)$ sur $[a,b]$, alors $\int_a^b R I^2(t) dt$ n'est (en général) pas égal à $R[\mu(I)]^2(b-a)$. On s'en convaincra en calculant les deux expressions lorsque $a=0$, $b=1$ et $I(t)=t$. Autrement dit dans le calcul de l'énergie dissipée, on ne peut pas remplacer I par sa moyenne.

C'est pour remédier de telles situations que l'on introduit la moyenne quadratique de f (en physique on dit aussi "valeur efficace") qui est définie par:

$$\mu_2(f,a,b) = \left[\frac{1}{b-a} \int_a^b f^2(t) dt \right]^{1/2}$$

Elle existe dès que f^2 est intégrable (donc chaque fois que f est continue). On vérifiera facilement que, dans notre exemple, on peut, dans le calcul de l'énergie dissipée, remplacer la fonction I par sa moyenne quadratique.

Complément 2 : Le calcul approché des intégrales

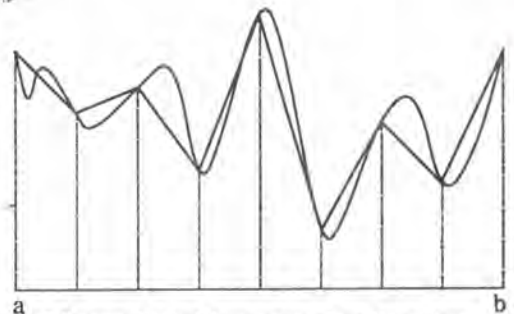
Soit à calculer l'intégrale $I = \int_a^b f(t) dt$ où f est une fonction dont nous ne connaissons pas de primitive. Nous allons en chercher des valeurs approchées. Toutes les méthodes classiques utilisent le même principe: remplacer f par une fonction "voisine" de f et dont on connaît une primitive. En fait on découpera $[a,b]$ en intervalles, et cette fonction voisine sera, sur chacun de ces intervalles, un polynôme de degré plus ou moins grand.

La méthode des trapèzes: Partageons l'intervalle $[a,b]$ en n parties égales, en posant $t_i = a + \frac{i}{n}(b-a)$. Nous avons $t_0=a < t_1 < \dots < t_n=b$. Et sur chaque intervalle $[t_i, t_{i+1}]$ remplaçons f par la fonction affine φ_i qui prend les mêmes valeurs que f en t_i et en t_{i+1} . Nous avons clairement:

$$\int_{t_i}^{t_{i+1}} \varphi_i(t) dt = \frac{1}{2}(t_{i+1}-t_i)[f(t_i)+f(t_{i+1})] = \frac{b-a}{2n} [f(t_i)+f(t_{i+1})]$$

Cette intégrale sera considérée comme une valeur approchée de $\int_{t_i}^{t_{i+1}} f(t) dt$, et nous obtiendrons une valeur approchée de I , en faisant la somme pour $i = 1, \dots, n-1$. D'où la valeur approchée de I :

$$I_n = \sum_{i=0}^{n-1} \frac{b-a}{2n} [f(t_i)+f(t_{i+1})] = \frac{b-a}{2n} \sum_{i=0}^{n-1} [f(t_i)+f(t_{i+1})]$$



L'intervalle $[a,b]$ est partagé en 8 parties égales. Sur chacune d'elles, la fonction est remplacée par une fonction affine.

Cette méthode est appelée "méthode des trapèzes" parceque, graphiquement (et lorsque f est positive), remplacer f par φ_i revient à remplacer l'aire du domaine curviligne limité par f , par l'aire d'un trapèze.

Exercice 1: Au moyen d'une calculatrice programmable, calculer par la méthode des trapèzes, l'intégrale $\int_0^{\pi/2} \sin t dt$. On fera par exemple $n = 100$. Comparer le résultat obtenu et le résultat exact.

Le calcul d'incertitude:

En fait une telle méthode de calcul approché n'a d'intérêt que lorsque l'on ne sait pas calculer de primitive de f . Et elle n'est alors acceptable que si nous sommes capables de donner une majoration de l'erreur faite en remplaçant I par I_n . Pour calculer une telle majoration, nous supposons f suffisamment régulière, par exemple C^∞ .

Nous allons, pour tout i , majorer $\delta_i = \int_{t_i}^{t_{i+1}} f(t) dt - (t_{i+1} - t_i) \frac{f(t_{i+1}) + f(t_i)}{2}$. Pour cela nous intégrons par parties (en choisissant $t + \alpha$ comme primitive^(*) de 1), donc

$$\delta_i = \left\{ [(t+\alpha)f(t)]_{t_i}^{t_{i+1}} - \int_{t_i}^{t_{i+1}} (t+\alpha) f'(t) dt \right\} - (t_{i+1} - t_i) \frac{f(t_{i+1}) + f(t_i)}{2}$$

Dans cette formule α est quelconque; nous choisirons $\alpha = -\frac{1}{2}(t_i + t_{i+1})$, de façon à obtenir:

$$\delta_i = - \int_{t_i}^{t_{i+1}} \left(t - \frac{1}{2}(t_i + t_{i+1}) \right) f'(t) dt.$$

Il nous reste à majorer cette intégrale. Pour cela nous intégrons par parties, en choisissant $\frac{1}{2}(t - \frac{t_i + t_{i+1}}{2})^2 + \beta$ comme primitive de $t - \frac{1}{2}(t_i + t_{i+1})$; et nous obtenons:

$$\delta_i = - \left[\left(\frac{1}{2} \left(t - \frac{t_i + t_{i+1}}{2} \right)^2 + \beta \right) f'(t) \right]_{t_i}^{t_{i+1}} + \int_{t_i}^{t_{i+1}} \left(\frac{1}{2} \left(t - \frac{t_i + t_{i+1}}{2} \right)^2 + \beta \right) f''(t) dt$$

En choisissant $\beta = -\frac{1}{2} \frac{(t_{i+1} - t_i)^2}{2}$, la partie toute intégrée est nulle, et il reste:

$$\delta_i = \int_{t_i}^{t_{i+1}} \left(\frac{1}{2} \left(t - \frac{t_i + t_{i+1}}{2} \right)^2 - \frac{1}{2} \frac{(t_{i+1} - t_i)^2}{2} \right) f''(t) dt = \int_{t_i}^{t_{i+1}} \frac{1}{2} (t - t_i)(t - t_{i+1}) f''(t) dt$$

Soit M_2 un majorant de $|f''|$ sur $[a, b]$. Puisque la fonction $t \rightarrow \frac{1}{2}(t - t_i)(t - t_{i+1})$ est négative ou nulle sur $[t_i, t_{i+1}]$, nous aurons $\frac{1}{2}(t - t_i)(t - t_{i+1}) M_2 \leq \frac{1}{2}(t - t_i)(t - t_{i+1}) f''(t) \leq -\frac{1}{2}(t - t_i)(t - t_{i+1}) M_2$. Donc:

$$\int_{t_i}^{t_{i+1}} \frac{1}{2} (t - t_i)(t - t_{i+1}) M_2 dt \leq \delta_i \leq \int_{t_i}^{t_{i+1}} -\frac{1}{2} (t - t_i)(t - t_{i+1}) M_2 dt$$

Soit (compte tenu du fait que $t_{i+1} - t_i = \frac{b-a}{n}$):

$$-\frac{(b-a)^3 M_2}{12 n^3} \leq \delta_i \leq \frac{(b-a)^3 M_2}{12 n^3}$$

Ou encore $|\delta_i| \leq \frac{(b-a)^3 M_2}{12 n^3}$. Et puisque l'erreur faite en remplaçant I par I_n est la somme des δ_i , nous aurons :

$$|I - I_n| \leq n \frac{(b-a)^3 M_2}{12 n^3} = \frac{(b-a)^3 M_2}{12 n^2}$$

Nous retiendrons que l'erreur faite en remplaçant l'intégrale $\int_a^b f(t) dt$ par la formule des trapèzes est au plus égale à $\frac{(b-a)^3 M_2}{12 n^2}$, pourvu que $(\forall t \in [a, b]) |f''(t)| \leq M_2$.

Un calcul plus précis, que nous ne ferons pas, donnerait :

(*) Notons que le présent calcul repose sur le fait que, dans une intégration par parties, un choix astucieux de la primitive que l'on utilise, permet d'obtenir des simplifications dans les calculs. C'est une méthode que nous pouvons retenir.

$$I - I_n = \frac{(b-a)^2}{12n^2} [f'(b) - f'(a)] + O(1/n^3)$$

Ce qui est le début d'un développement en $1/n$ de $I - I_n$, et nous permet de calculer la "partie principale de l'erreur" si nous connaissons la dérivée de f .

Exercice m: On veut calculer $\int_0^1 e^{-t^2} dt$. Notons qu'il s'agit d'un cas où il n'existe aucune primitive simple.

α) Calculer les dérivées première, deuxième et troisième de la fonction e^{t^2} . Tracer le graphique de cette dernière pour trouver un nombre M_2 qui majore sa valeur absolue sur l'intervalle $[0,1]$.

β) En déduire une majoration de l'erreur que l'on fait en remplaçant l'intégrale cherchée par I_n . Quelle valeur de n faut-il prendre pour obtenir une valeur approchée à 10^{-4} près de l'intégrale. Calculer alors un tel I_n .

La méthode de Simpson : Partageons l'intervalle en $2n$ parties égales, en posant $s_i = a + i \frac{b-a}{2n}$. Puis sur chaque intervalle $[s_{2i}, s_{2i+2}]$, remplaçons f par le polynôme Q_i de degré 2 qui prend les mêmes valeurs que f en s_{2i} , s_{2i+1} et s_{2i+2} . Le polynôme $Q_i(t)$ vaut:

$$\frac{(t-s_{2i+1})(t-s_{2i+2})}{(s_{2i}-s_{2i+1})(s_{2i}-s_{2i+2})} f(s_{2i}) + \frac{(t-s_{2i})(t-s_{2i+2})}{(s_{2i+1}-s_{2i})(s_{2i+1}-s_{2i+2})} f(s_{2i+1}) + \frac{(t-s_{2i})(t-s_{2i+1})}{(s_{2i+2}-s_{2i})(s_{2i+2}-s_{2i+1})} f(s_{2i+2})$$

D'où :

$$\int_{s_{2i}}^{s_{2i+2}} Q_i(t) dt = \frac{1}{3} \frac{b-a}{2n} f(s_{2i}) + \frac{4}{3} \frac{b-a}{2n} f(s_{2i+1}) + \frac{1}{3} \frac{b-a}{2n} f(s_{2i+2})$$

En faisant la somme pour les n intervalles $[s_{2i}, s_{2i+2}]$, nous obtenons une valeur approchée de l'intégrale de f sur $[a,b]$:

$$J_n = \sum_{i=0}^{n-1} \int_{s_{2i}}^{s_{2i+2}} Q_i(t) dt = \sum_{i=0}^{n-1} \left[\frac{1}{3} \frac{b-a}{2n} f(s_{2i}) + \frac{4}{3} \frac{b-a}{2n} f(s_{2i+1}) + \frac{1}{3} \frac{b-a}{2n} f(s_{2i+2}) \right]$$

$$J_n = \frac{b-a}{6n} \left[f(a) + f(b) + 4 \sum_{i=0}^{n-1} f(s_{2i+1}) + 2 \sum_{i=1}^{n-1} f(s_{2i}) \right]$$

Comme dans le cas de la méthode des trapèzes, ce calcul ne présente un intérêt que si nous pouvons majorer l'erreur faite en remplaçant l'intégrale de f par J_n . Nous ne ferons pas le calcul, mais nous retiendrons:

Si f est de classe C^4 , et si $(\forall t \in [a,b]) |f^{(4)}(t)| \leq M_4$, l'erreur est au plus égale à $M_4 \frac{(b-a)^5}{2880 n^4}$.

Exercices sur le chapitre 7

- * **Exercice 1:** Soit $\chi: [0,1] \rightarrow \mathbb{R}$ la fonction caractéristique de $\mathbb{Q} \cap [0,1]$. En remarquant que tout intervalle (non vide) de $\mathbb{R}$ contient des rationnels, et aussi des irrationnels, calculer $I(\chi, +)$ et $I(\chi, -)$. En déduire que χ n'est pas intégrable au sens de Riemann.
- * **Exercice 2:** Soit f une fonction croissante sur $[0,1]$. Calculer le maximum et le minimum de f sur $[\frac{1}{n}, \frac{i+1}{n}]$. En déduire que f est intégrable au sens de Riemann.
- ** **Exercice 3:** On considère la suite $(u_n)_{n \geq 1}$ définie par $u_n = \int_{n\pi}^{n\pi+\pi} \frac{\sin t}{t^2} dt$.
- a) Démontrer les inégalités $\frac{2}{n^2\pi^2} \geq |u_n| \geq \frac{2}{(n+1)^2\pi^2}$. Montrer que $(u_n)_{n \geq 1}$ est équivalente à $((-1)^n \frac{2}{n^2\pi^2})_{n \geq 1}$.
- b) On pose $v_n = \int_{n\pi}^{n\pi+\pi} \frac{\cos t}{t^3} dt$. Au moyen d'une intégration par parties, trouver une relation entre u_n et v_n . En déduire un développement du type $u_n = \frac{A_n}{n^2} + \frac{B_n}{n^3} + o(\frac{1}{n^3})$.
- c) Au moyen d'intégrations par parties trouver un développement à l'ordre 5 de u_n en fonction de $\frac{1}{n}$.
- * **Exercice 4:** Pour tout entier $n \geq 0$, on pose $I_n = \int_0^1 x^n \sin \pi x dx$.
- a) Calculer I_0 et I_1 . Trouver une relation entre I_n et I_{n+2} .
- b) Montrer que la suite $(I_n)_{n \in \mathbb{N}}$ est bornée, puis montrer qu'elle converge vers 0. Démontrer que $(I_n)_{n \in \mathbb{N}}$ est équivalente à une suite de la forme $(\frac{K}{n^2})_{n \in \mathbb{N}}$.
- * **Exercice 5:** a) Soit f de classe C^1 sur $[0, \pi]$, on pose $u_n = \int_0^\pi f(t) \sin nt dt$. Montrer que $(u_n)_{n \in \mathbb{N}} = O(\frac{1}{n})_{n \geq 1}$. On pourra faire une intégration par parties.
- b) On suppose maintenant f de classe C^2 . A quelle condition peut-on affirmer que $(u_n)_{n \in \mathbb{N}} = O(\frac{1}{n^2})_{n \geq 1}$. Montrer que, si cette condition n'est pas remplie, il existe A tel que $(u_n)_{n \in \mathbb{N}}$ soit équivalente à $(\frac{A}{n})_{n \geq 1}$.
- * **Exercice 6:** La fonction $\frac{\sin t}{t}$ "est" continue en 0 (Justifier cette affirmation). Donc $I = \int_0^{\pi/4} \frac{\sin t}{t} dt$ existe.
- a) Montrer que $(\forall x \geq 0) \sin x \leq x$. En déduire que $(\forall x \geq 0) 1 - \cos x \leq \frac{x^2}{2}$ (c'est à dire $\cos x \geq 1 - \frac{x^2}{2}$).
- b) Démontrer les inégalités (valables pour tout x positif ou nul):

$$x - \frac{x^3}{6} \leq \sin x \leq x - \frac{x^3}{6} + \frac{x^5}{120}$$
- c) En déduire une valeur approchée à $3 \cdot 10^{-4}$ près de $I = \int_0^{\pi/4} \frac{\sin t}{t} dt$.
- d) Calculer une valeur approchée à 10^{-6} près de I .
- ** **Exercice 7:** Soit $(u_n)_{n \in \mathbb{N}}$ une suite de points de $[0,1]$, qui converge vers λ ($\in [0,1]$). On note χ la fonction caractéristique de cette suite ($\chi(t) = 1$ si t est l'un des u_n , et $\chi(t) = 0$ sinon). On se propose de montrer que χ est négligeable.

a) Montrer que $I(\chi, -) = 0$.

b) Montrer que, quel que soit ε , il existe une réunion d'intervalles dont la somme des longueurs est au plus égale à ε , et qui contient tous les termes de la suite. (on remarquera que $[\lambda - \frac{\varepsilon}{4}, \lambda + \frac{\varepsilon}{4}]$ contient tous les termes de la suite sauf un nombre fini). En déduire que $(\forall \varepsilon > 0) I(\chi, +) \leq \varepsilon$. Conclure.

**** Exercice 8:** On définit une suite par $I_n = \int_1^a (\text{Log } t)^n dt$

a) On suppose $1 < a < e$. Déterminer un nombre α tel que $(\forall n) |I_n| \leq (a-1)\alpha^n$. En déduire la limite de $(I_n)_{n \in \mathbb{N}}$.

b) On suppose $a = e$. Montrer que la suite $(I_n)_{n \in \mathbb{N}}$ est décroissante, et qu'elle a une limite λ positive ou nulle. En utilisant la question a, montrer que, quel que soit $b < e$, on a $\lambda \leq e - b$. Que vaut λ ?

c) On suppose $a > e$. Montrer que $(\forall n) I_n \geq \frac{a-e}{2} [\text{Log}(\frac{a+e}{2})]^n$. En déduire que la suite $(I_n)_{n \in \mathbb{N}}$ tend vers $+\infty$.

*** Exercice 9:** Etudier la fonction $f: x \rightarrow \int_x^{2x} e^{-t^2} dt$

a) Calculer sa dérivée.

b) Montrer que $|f(x)| \leq |x| e^{-x^2}$; en déduire la limite de f à l'infini.

c) Donner un développement à l'ordre 4 de f en 0.

*** Exercice 10:** Etudier la fonction $f: x \rightarrow \int_x^{2x} \frac{\text{Log } t}{1+t^4} dt$ ($x > 0$).

a) Calculer sa dérivée.

b) Quelles sont ses limites en 0 et à l'infini ?

c) Donner un équivalent de f au voisinage de 0. En déduire que f se prolonge par continuité en 0, et préciser la direction de la tangente au point limite.

*** Exercice 11:** Etudier la fonction $f: t \rightarrow \int_t^{1-t} e^{u^2} du$. Déterminer sa dérivée, son sens de variation, ses limites en $\pm\infty$.

**** Exercice 12:** On pose $I_n = \int_0^{\pi/2} \cos^n t dt$.

a) Etablir une relation de récurrence liant I_n et I_{n-2} . On pourra faire une intégration par parties dans I_n .

b) Calculer I_0 et I_1 . Puis donner des formules permettant de calculer I_{2n} et I_{2n+1} .

c) En majorant $\int_0^\varepsilon \cos^n t dt$ et $\int_\varepsilon^{\pi/2} \cos^n t dt$ montrer que $(\forall \varepsilon > 0) 0 \leq I_n \leq \varepsilon + \frac{\pi}{2} \cos^n \varepsilon$. En déduire que la suite

$(u_n)_{n \in \mathbb{N}}$ converge vers 0.

***** Exercice 13:** Soit une courbe donnée en coordonnées polaires par $\rho = f(\theta)$ ($\theta \in [0, \pi]$ et f à valeurs dans $[0, \infty[$). On se propose de calculer l'aire $\mathcal{A}(\omega)$ comprise entre la courbe et les rayons vecteurs d'angles polaires 0 et ω .

1) On pose $m = \inf_{t \in [\theta, \theta+h]} f(t)$ et $M = \sup_{t \in [\theta, \theta+h]} f(t)$. Majorer et minorer $\mathcal{A}(\theta+h) - \mathcal{A}(\theta)$. En s'inspirant du §3, en déduire que la fonction $\mathcal{A}$ est dérivable à droite, et calculer cette dérivée.

2) Montrer que $\mathcal{A}$ est dérivable à gauche. Puis montrer que $\mathcal{A}$ est dérivable, et donner une formule permettant de calculer $\mathcal{A}'(\theta)$.

3) On considère la courbe $\rho = \sqrt{\cos 2\theta}$. La tracer, puis calculer l'aire intérieure à l'une de ses boucles.

*** Exercice 14: Lorsqu'une masse $m(\text{kg})$ a une vitesse $v(\text{mètres par seconde})$, son énergie cinétique est $\frac{1}{2}mv^2(\text{joules})$.

Une barre homogène de longueur $L(\text{mètres})$ a une extrémité fixe; elle tourne autour de cette extrémité A, avec une vitesse angulaire $\omega(\text{radians par seconde})$. Sa masse spécifique est $\mu(\text{grammes par mètre})$.

a) On considère la partie de la barre formée des points situés à une distance de A comprise entre x et $x+h$. Majorer et minorer l'énergie cinétique de cette partie de la barre.

b) Soit $E(x)$ l'énergie cinétique de la partie de la barre située à une distance de A inférieure à x . Montrer que E est une fonction dérivable à droite. Montrer que E est une fonction dérivable. Quelle est sa dérivée? Donner une formule permettant de calculer $E(L)$ par primitivation.

*** Exercice 15: Soit f une fonction positive continue de $[a,b]$ dans $\mathbb{R}$, on appelle γ la courbe d'équation $\rho = f(z)$ dans le demi-plan des (ρ, z) ($\rho > 0$). On note A le domaine plan défini par $(a \leq z \leq b)$ et $(0 \leq \rho \leq f(z))$. On note $\mathcal{V}$ le volume de révolution engendré par $\mathcal{V}$ en tournant autour de Oz.

a) Soit ϕ une fonction en escalier définie sur $[a,b]$ et qui minore f . Montrer que le volume V de $\mathcal{V}$ est supérieur à π fois l'intégrale de ϕ^2 . Montrer que si ψ est une fonction en escalier définie sur $[a,b]$ et qui majore f , le volume V de $\mathcal{V}$ est inférieur à π fois l'intégrale de ψ^2 .

b) En déduire que $V = \int_a^b \pi f^2(t) dt$.

* Exercice 16: Soit $I = \int_0^{\pi/2} \sqrt{t} \sin t dt$

a) Au moyen de la méthode de Simpson, calculer une valeur approchée à 10^{-5} près de I .

b) Même question mais en utilisant la méthode des trapèzes.

Calculs de primitives

$$A = \int \cos t e^{2t} dt$$

$$B = \int t \sin t \sin \pi t dt$$

$$C = \int \frac{t^6}{(1-t)^3(1+t^2)} dt$$

$$D = \int \frac{1+t^2}{1+t^4} dt$$

$$E = \int \frac{1+t}{\sqrt{1+t^2}} dt$$

$$F = \int \cos^5 t \cos 2t dt$$

$$G = \int t^2 \sqrt{1-2t^2} dt$$

$$H = \int \frac{t^2 dt}{\sqrt{(t-2)(t-3)}}$$

$$I = \int \frac{t^2 dt}{\sqrt{(t-2)(3-t)}}$$

$$J = \int \frac{\cos t \sin t}{2 + \operatorname{tg}^2 t} dt$$

$$K = \int \frac{t^3 dt}{\sqrt{1+t+t^2}} dt$$

$$L = \int \frac{\cos^3 t dt}{\cos^2 t + 3 \sin 2t}$$

$$M = \int \frac{\sqrt{1+\cos^2 t} dt}{(1+\sin^2 t) \sin t}$$

$$N = \int \frac{\cos^2 t dt}{\cos t + \sin t}$$

$$O = \int t^n (\operatorname{Log} t)^m dt$$

$$P = \int \cos 3t \cos \pi t e^{3t} dt$$

$$Q = \int \frac{t dt}{(1+t)^4(1-t)^2}$$

$$R = \int t^7 e^{-t^2} dt$$

$$S = \int \frac{t^3 dt}{(t+1)\sqrt{1-t^2}}$$

$$T = \int \frac{t^6 dt}{(1+t^2)^2(1-t^2)^2}$$

$$U = \int \frac{t^7 dt}{1-t^5}$$

$$V = \int_{-1}^1 t^3 \sqrt{1-t^2} dt$$

$$W = \int_0^1 \frac{1+3t}{(1+t)^6(1+t^2)} dt$$

$$X = \int_0^{\pi/6} \frac{\cos 3t}{\cos^2 2t} dt$$

$$Y = \int \frac{dt}{1-t^6}$$

$$Z = \int \frac{dt}{(1-t^4)(1+t^2)}$$

$$AA = \int \cos^3 t \sin^2 3t dt$$

$$AB = \int t \sqrt{1-t} dt$$

$$AC = \int \frac{\cos t}{2 + \cos t} dt$$

$$AD = \int \frac{1+t}{1-t^3} dt$$

$$AE = \int \frac{\operatorname{ch} t}{1 + \operatorname{ch} t} dt$$

$$AF = \int \operatorname{ch}^3 t \operatorname{sh} 3t dt$$

$$AG = \int \operatorname{ch} t \cos^3 t dt$$

La convergence "à priori" et les séries numériques

§ 1 : Les propriétés fondamentales des nombres réels.

Nous admettrons les deux propriétés suivantes de l'ensemble des nombres réels:

Toute suite de nombres réels croissante et majorée est convergente (ou - ce qui est équivalent - toute suite décroissante minorée est convergente).

Tout ensemble de nombres réels qui est majoré a une borne supérieure() .*

Vous avez appris ces propriétés au cours des années précédentes; on ne vous les a jamais démontrées. Bien sûr ! On ne vous a jamais défini les nombres réels. Essayons d'approfondir un peu la question, d'abord par une brève étude historique.

L'origine du nombre est contemporaine des premières civilisations. On inventa d'abord les nombres entiers (positifs) pour compter les collections finies d'objets. De même au cours préparatoire, vous avez appris à compter votre collection de billes. Pour mesurer (des longueurs, des aires,...) avec une unité donnée, les nombres entiers ne suffisent plus; il faut utiliser les nombres fractionnaires, ou les nombres décimaux. C'est au cours élémentaire que vous avez étudié cette première extension de la notion de nombre. Toutes les grandes civilisations de l'antiquité possédaient des systèmes d'unités et de sous-unités utilisées essentiellement dans le négoce.

Notre civilisation scientifique est directement issue de la Grèce ancienne. Il est donc intéressant de savoir que Pythagore (-585,-500) utilisait les nombres entiers et les nombres fractionnaires (avec des restrictions philosophiques en vertu desquelles ces derniers n'avaient pas un véritable statut de nombre). C'est en découvrant les propriétés du triangle rectangle, que les grecs constatèrent que la diagonale du carré de côté 1, ne pouvait être mesurée au moyen des nombres fractionnaires (puisque'il n'existe aucun nombre fractionnaire dont le carré soit 2). Il fallait revoir toute la conception que l'on se faisait des notions de nombre et de proportion. Il fallut plusieurs siècles aux mathématiciens grecs pour sortir de ce dilemme. Ce fut Eudoxe (-408,-355) qui étendit la notion de proportion, pour en faire la première ébauche de ce que nous appelons les nombres réels.

Pour Eudoxe le rapport d'une longueur L à l'unité λ (nous dirions la mesure de L lorsque l'unité est λ) est définie par d'une part l'ensemble des nombres fractionnaires r tels que $r\lambda < L$, d'autre part l'ensemble de ceux qui sont tels que $r\lambda > L$. C'est la propriété fondamentale qu'un nombre réel λ est connu dès que l'on connaît les rationnels supérieurs à λ , et les rationnels inférieurs à λ . L'outil ainsi construit évolua peu jusqu'au début du 19-ème siècle.

A partir de la fin du 17-ème siècle, on utilisa systématiquement les suites. En fait la notion de convergence n'était pas vraiment définie, mais ceci ne semble pas avoir gêné les mathématiciens de l'époque. Il suffisait d'un peu de flair pour ne pas dire de bêtises. La situation devint intenable au début du 19-ème siècle. Essentiellement parce que les suites qui au départ n'étaient qu'un outil pour calculer des nombres comme $(\pi, \sqrt{2}, \dots)$ dont l'existence était indiscutable, servaient à construire des solutions de problèmes de plus en plus compliqués, dont l'existence n'était pas évidente. Il fallut inventer des règles permettant de dire si une suite a ou

(*) Rappelons que λ est borne supérieure de l'ensemble X si, d'une part $(\forall x \in X)$ on a $x \leq \lambda$, d'autre part $(\forall \epsilon > 0)$ $(\exists x \in X)$ tel que $\lambda - \epsilon \leq x \leq \lambda$

non une limite, lorsque l'on ne connaît pas cette limite (cf: § 2 ci-dessous). En fait on ne fit que formaliser un certain nombre de procédures que l'on utilisait depuis longtemps (ce qui explique qu'au siècle précédent on n'avait pas raconté trop d'aneries), et dont les principales sont:

- * La propriété des suites croissantes majorées.
- * La propriété de la borne supérieure.
- * Le critère de Cauchy (§ 2 ci-dessous).

Toutes ces propriétés sont équivalentes.

On avait ainsi l'outil permettant de développer l'étude des fonctions^(**). Mais on n'avait toujours pas défini les nombres réels. Ce n'est que dans les années 1880-90 que Cantor (1845,1918) et Dedekind (1831,1916) donnèrent une véritable construction de $\mathbb{R}$, à partir des rationnels.

L'objet du présent chapitre est de démontrer à partir des propriétés rappelées ci-dessus, le critère de Cauchy, puis de construire l'outil qui permet l'utilisation commode de ce critère, c'est la théorie des séries. Les étudiants curieux trouveront en appendice une nouvelle note historique, et une esquisse d'une construction logique de $\mathbb{R}$ à partir de $\mathbb{Q}$.

Les critères de convergence "à priori", pour quoi faire ?

Étudions d'abord un exemple. Soit f une fonction continue sur $[a,b]$ telle que $f(a) < 0 < f(b)$. Nous savons qu'il existe λ ($a < \lambda < b$) tel que $f(\lambda) = 0$. Définissons une suite $(u_n)_{n \in \mathbb{N}}$ par $u_0 = a$ et la procédure récurrente:

$$\begin{cases} \text{Si } f(u_n + \frac{b-a}{2^{n+1}}) \leq 0, \text{ alors } u_{n+1} = u_n + \frac{b-a}{2^{n+1}} \\ \text{Si } f(u_n + \frac{b-a}{2^{n+1}}) > 0, \text{ alors } u_{n+1} = u_n \end{cases}$$

Exercice a : Démontrer par récurrence que, pour tout n , $f(u_n) \leq 0$; et que, pour tout n , $f(u_n + \frac{b-a}{2^n}) > 0$.

Supposons que f soit strictement croissante, la relation $f(u_n) \leq 0 < f(u_n + \frac{b-a}{2^n})$, équivaut à la relation $u_n \leq \lambda < u_n + \frac{b-a}{2^n}$. Ce qui entraîne $|u_n - \lambda| \leq \frac{b-a}{2^n}$; et ainsi la suite $(u_n)_{n \in \mathbb{N}}$ converge vers λ . Nous avons ainsi décrit l'une des procédures les plus simples, qui permette de calculer les racines des équations numériques, au moyen d'une calculatrice programmable.

Oublions maintenant le théorème de la valeur intermédiaire.

Nous pouvons affirmer "à priori" que $(u_n)_{n \in \mathbb{N}}$ est convergente puisque c'est une suite croissante (c'est clair) majorée (par b). Notons μ sa limite. Puisque f est continue en μ , la suite $(f(u_n))_{n \in \mathbb{N}}$ converge vers $f(\mu)$; ce qui prouve que $f(\mu) \leq 0$. Posons alors $v_n = u_n + \frac{b-a}{2^n}$; la suite $(v_n)_{n \in \mathbb{N}}$ a la même limite μ que $(u_n)_{n \in \mathbb{N}}$; donc $(f(v_n))_{n \in \mathbb{N}}$ converge vers $f(\mu)$. Puisque $(\forall n) f(v_n) > 0$, ceci prouve que $f(\mu) \geq 0$. Nous avons ainsi démontré que $f(\mu) = 0$. Autrement dit nous avons démontré le théorème de la valeur intermédiaire.

Nous voyons que le même processus de construction d'une suite, peut être interprété de deux façons. D'une part comme une méthode de calcul numérique permettant de calculer effectivement des valeurs approchées de λ . D'autre part pour montrer que l'équation $f(x) = 0$ a une racine. C'est parceque nous avons pu affirmer à priori la convergence de $(u_n)_{n \in \mathbb{N}}$, que nous avons pu donner cette seconde interprétation. Ainsi les critères de convergence à priori sont des outils pour démontrer des théorèmes d'existence. Dans toute la suite il convient de garder à l'esprit ces deux points de vue.

^(**) A l'exception près des nombres négatifs, qui posaient encore des problèmes de nature métaphysique, et ne reçurent le statut de nombre que vers 1890.

§ 2 : Les suites de Cauchy dans $\mathbb{R}$.

Le critère des suites croissantes majorées est la seule règle que vous connaissiez et qui vous permet de montrer qu'une suite est convergente sans connaître d'avance sa limite. Il est peu utilisable car l'hypothèse que la suite est monotone est fort restrictive. Nous allons donner un critère aux hypothèses moins contraignantes.

Nous dirons que la suite $(u_n)_{n \in \mathbb{N}}$ est une suite de Cauchy, si :

$$(\forall \varepsilon > 0) (\exists N) \text{ tel que } (\forall n) (\forall p) (n \geq N \text{ et } p \geq N) \text{ implique } |u_n - u_p| \leq \varepsilon$$

Nous allons démontrer qu'une suite de nombres réels est convergente si et seulement si elle est de Cauchy.

Montrons que toute suite convergente est de Cauchy. Si $(u_n)_{n \in \mathbb{N}}$ converge vers λ , $\varepsilon > 0$ étant donné, nous pouvons trouver N tel que $n \geq N$ implique $|u_n - \lambda| \leq \varepsilon/2$. Alors si n et p sont tous les deux au moins égaux à N , nous aurons $|u_n - \lambda| \leq \varepsilon/2$ et $|u_p - \lambda| \leq \varepsilon/2$; donc :

$$|u_n - u_p| \leq |u_n - \lambda| + |u_p - \lambda| \leq \varepsilon.$$

Réciproquement si $(u_n)_{n \in \mathbb{N}}$ est une suite de Cauchy :

* D'abord $(u_n)_{n \in \mathbb{N}}$ est une suite bornée. En effet (en faisant $\varepsilon = 1$) il existe un nombre N tel que $n \geq N$ implique $|u_n - u_N| \leq 1$; c'est à dire $-1 + u_N \leq u_n \leq 1 + u_N$. Donc quel que soit n :

$$|u_n| \leq \sup(|u_0|, |u_1|, \dots, |u_{N-1}|, |-1 + u_N|, |1 + u_N|) = M$$

* Donc pour tout n , l'ensemble $E_n = \{u_n, u_{n+1}, \dots, u_p, \dots\}$ formé des termes de la suite de rang supérieur ou égal à n , est majoré et minoré. Ainsi E_n a une borne inférieure m_n et une borne supérieure M_n . Puisque $E_{n+1} \subset E_n$, nous avons $m_{n+1} \geq m_n$. Donc la suite $(m_n)_{n \in \mathbb{N}}$ est une suite croissante majorée par M ; elle converge vers un certain λ .

* Puisque la suite $(u_n)_{n \in \mathbb{N}}$ est de Cauchy, quel que soit $\varepsilon > 0$, il existe N tel que $n \geq N$ implique que deux éléments quelconques de E_n sont distants de moins de ε . Donc (lorsque $n \geq N$) $0 \leq M_n - m_n \leq \varepsilon$. Il en résulte que la suite $(M_n - m_n)_{n \in \mathbb{N}}$ tend vers 0, et que $(M_n)_{n \in \mathbb{N}}$ converge vers λ . Mais pour tout n nous avons $m_n \leq u_n \leq M_n$, et ainsi la suite $(u_n)_{n \in \mathbb{N}}$ est encadrée par deux suites qui convergent vers λ , donc elle a pour limite λ .

§ 3 : Le langage des séries

Le critère de Cauchy montre que dans les problèmes liés à la convergence "à priori" d'une suite $(U_n)_{n \in \mathbb{N}}$, les différences $U_n - U_p$ jouent un rôle fondamental. C'est pourquoi il est commode de changer de langage pour faire apparaître dans les notations les différences $U_{n+1} - U_n$.

Une série numérique est une application de $\mathbb{N}$ dans $\mathbb{R}$; donc formellement il n'y a aucune différence entre une suite et une série, ce sont les questions que l'on pose à leur sujet qui différencient les deux notions.

A toute série $(u_n)_{n \in \mathbb{N}}$ on associe la suite $(U_n)_{n \in \mathbb{N}}$ définie par $U_n = \sum_{i=0}^n u_i$. Les U_n sont appelés les sommes partielles de la série. Nous dirons que la série est sommable^(*) si la suite $(U_n)_{n \in \mathbb{N}}$ a une limite. Celle-ci sera alors appelée la somme de la série, et notée $\sum_{i=0}^{\infty} u_i$.

(*) On dit aussi, assez improprement, que la série est convergente.

La différence $\sum_{i=0}^{\infty} u_i - \sum_{i=0}^n u_i$, sera aussi notée $\sum_{i=n+1}^{\infty} u_i$, elle est appelée le reste d'ordre n de la série.

Assez souvent les u_n ne seront bien définis que pour $n \geq 1$, ou $n \geq 2$... Nous poserons alors $U_n = \sum_{i=1}^n u_i$, ou $U_n = \sum_{i=2}^n u_i$, ... Remarquons que si nous oublions de définir les premiers des u_n , les

U_n sont modifiés: on retranche le même nombre à chacun d'eux. Ceci n'influe pas sur l'existence de la limite, mais modifie cette limite. Ainsi chaque fois que nous chercherons seulement à savoir si une série est sommable, nous pourrons oublier ses premiers termes. Ce qui ne sera plus possible si nous voulons calculer sa somme.

Notation: Lorsque nous sommes en présence d'une famille de nombres $(u_n)_{n \in \mathbb{N}}$ il est désormais important que nous précisions si nous la considérons comme une suite ou comme une série. Dans la suite de ce cours nous écrirons alors "la suite $(u_n)_{n \in \mathbb{N}}$ " ou "la série $((u_n))_{n \in \mathbb{N}}$ " (avec un double parenthésage).

§ 4 : Séries à termes positifs

Comment démontrer qu'une série positive est sommable ? La suite $(U_n)_{n \in \mathbb{N}}$ est croissante si et seulement si les u_n sont positifs. Le critère des suites convergentes majorées se traduit donc par: Si les termes de la série $((u_n))_{n \in \mathbb{N}}$ sont positifs et s'il existe un nombre M qui majore ses sommes partielles, alors elle est sommable. Il suffit d'ailleurs que les u_n soient positifs à partir d'un certain rang.

Pour construire un majorant des sommes partielles de $((u_n))_{n \in \mathbb{N}}$ nous construirons souvent une seconde série à termes positifs $((v_n))_{n \in \mathbb{N}}$ qui est sommable et qui majore $((u_n))_{n \in \mathbb{N}}$, c'est à dire telle que $(\forall n)$ (ou, plus généralement, quelque soit n assez grand) $v_n \geq u_n$ (≥ 0). La somme de $((v_n))_{n \in \mathbb{N}}$ est alors un majorant de toutes les sommes partielles de $((u_n))_{n \in \mathbb{N}}$ puisque:

$$\sum_{i=0}^n u_i \leq \sum_{i=0}^n v_i \leq \lim_{n \rightarrow \infty} \sum_{i=0}^n v_i = \sum_{i=0}^{\infty} v_i.$$

Nous aurons alors $\sum_{i=0}^{\infty} u_i \leq \sum_{i=0}^{\infty} v_i$.

Exemple: Montrons que la série $((\frac{1}{n^2}))_{n \geq 1}$ est sommable. Nous utiliserons la majoration $\frac{1}{n^2} \leq \frac{1}{n(n-1)}$. La

série $((\frac{1}{n(n-1)}))_{n \geq 2}$ est sommable, de somme 1, car $\sum_{i=2}^n \frac{1}{i(i-1)} = 1 - \frac{1}{n}$ (pour s'en convaincre, faire un raisonnement par récurrence). Donc la série $((\frac{1}{n^2}))_{n \geq 1}$ est sommable. Essayons de calculer une valeur approchée de sa somme. Pour tout $n \geq 2$ et tout $p > n$, nous pouvons écrire:

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

D'où, en passant à la limite sur p :

$$\sum_{i=1}^{\infty} \frac{1}{i^2} - \sum_{i=1}^n \frac{1}{i^2} \leq \frac{1}{n}$$

Ainsi en sommant les n premiers termes de notre série nous obtiendrons une valeur approchée (par défaut évidemment) à $\frac{1}{n}$ près de sa somme.

Exercice b: Soit α ($0 < \alpha < 1$), on pose $u_n = \sin \alpha^n$. En utilisant la majoration $|\sin t| \leq |t|$, montrer qu'elle est sommable. On suppose $\alpha = 1/2$, au moyen d'une calculatrice, calculer une valeur approchée à 10^{-3} près de sa limite.

Remarque: Calculer une somme partielle d'une série c'est calculer la somme d'un nombre en général grand de termes, et ces termes vont le plus souvent en décroissant. Dans ce calcul les erreurs d'arrondi vont s'ajouter, pour donner sur la somme cherchée une incertitude inacceptable. Une façon d'y remédier consiste à commencer le calcul par la fin, c'est à dire de calculer $u_n + u_{n-1} + u_{n-2} + \dots + u_1 + u_0$ plutôt que $u_0 + u_1 + \dots + u_{n-1} + u_n$. Ceci s'explique ainsi: la précision de calcul de la machine dépend des nombres qu'elle manipule; elle est d'autant meilleure que ces nombres sont petits, puisque le calcul se fait sur un nombre fixe (8, 10, 16, ...) de chiffres significatifs. Nous avons donc intérêt à ce que dans notre suite d'additions, la somme déjà calculée soit, le plus longtemps possible, très petite; ce qui a lieu si nous commençons le calcul par la fin (puisque les premiers termes sont les plus grands).

Comment démontrer qu'une série positive n'est pas sommable ? Si les u_n sont positifs, pour montrer que la série $((u_n))_{n \in \mathbb{N}}$ n'est pas sommable, il suffit de montrer que la suite $(U_n)_{n \in \mathbb{N}}$ tend vers l'infini. Pour cela nous pouvons minorer les u_n par les termes d'une série $((v_n))_{n \in \mathbb{N}}$ positive qui est elle

même non sommable. Nous aurons alors $U_n = \sum_{i=0}^n u_i \geq \sum_{i=0}^n v_i$, et la suite $(U_n)_{n \in \mathbb{N}}$ sera ainsi minorée par une suite qui tend vers l'infini.

Exercice c: Calculer les sommes partielles de la série $((v_n))_{n \geq 1} = ((\log(n+1) - \log n))_{n \geq 1}$, en déduire qu'elle n'est pas sommable. En déduire alors que la série $((\frac{1}{n}))_{n \geq 1}$ n'est pas sommable.

§ 5 : Séries sommables et séries absolument sommables.

Nous allons étudier maintenant les séries dont les termes ne sont pas tous de même signe; et pour cela nous allons essayer de nous ramener aux séries à termes positifs.

Séries absolument sommables: Nous dirons qu'une série $((u_n))_{n \in \mathbb{N}}$ est absolument sommable si la série des valeurs absolues $((|u_n|))_{n \in \mathbb{N}}$ est une série (à termes positifs) sommable.

Toute série absolument sommable est sommable.

C'est une conséquence du critère de Cauchy. En effet si $U_n = \sum_{i=0}^n u_i$ et $V_n = \sum_{i=0}^n |u_i|$ nous aurons (pour $p > n$):

$$|U_p - U_n| = \left| \sum_{i=n+1}^p u_i \right| \leq \sum_{i=n+1}^p |u_i| \leq V_p - V_n$$

Puisque la suite $(V_n)_{n \in \mathbb{N}}$ est convergente, elle a la propriété de Cauchy. Donc

$$(\forall \varepsilon > 0) (\exists N > 0) (\forall n) (\forall p) (n \geq N \text{ et } p \geq N) \text{ implique } |V_p - V_n| \leq \varepsilon$$

L'inégalité ci-dessus nous permet alors d'affirmer que

$$(\forall \varepsilon > 0) (\exists N > 0) (\forall n) (\forall p) (n \geq N \text{ et } p \geq N) \text{ implique } |U_p - U_n| \leq \varepsilon$$

Donc la suite $(U_n)_{n \in \mathbb{N}}$ est une suite de Cauchy, donc elle converge.

Pour montrer qu'une série est sommable, il suffit donc de majorer les valeurs absolues de ses termes par une série (à termes positifs) sommable. Pour cela il faut que nous ayons un stock de séries dont nous connaissons la nature; c'est ce que nous fabriquerons aux § 7 et 8.

Si l'on peut minorer les valeurs absolues des $|u_n|$ par les termes d'une série (à termes positifs) non sommable, alors la série $((u_n))_{n \in \mathbb{N}}$ n'est pas absolument sommable, ce qui ne l'empêche pas d'être sommable comme on le verra ci-dessous.

Remarque: De l'inégalité $|U_p - U_n| \leq |V_p - V_n|$, résulte (en passant à la limite sur p):

$$\left| \sum_{i=0}^{\infty} u_i - U_n \right| \leq \sum_{i=0}^{\infty} |u_i| - V_n$$

Ainsi la valeur absolue du reste de la série $((u_n))_{n \in \mathbb{N}}$, est inférieure ou égale au reste de la série $((|u_n|))_{n \in \mathbb{N}}$.

Les règles de comparaison des suites dans l'étude de la sommabilité absolue

Pour étudier la sommabilité absolue d'une série $((u_n))_{n \in \mathbb{N}}$ nous devons essayer de majorer ou de minorer la valeur absolue de ses termes par ceux d'une série positive $((v_n))_{n \in \mathbb{N}}$. Pour cela nous pouvons utiliser les critères de comparaison du chapitre 6.

La relation O: Si $(|u_n|)_{n \in \mathbb{N}} = O((v_n)_{n \in \mathbb{N}})$, et si $((v_n))_{n \in \mathbb{N}}$ est une série positive sommable, alors la série $((u_n))_{n \in \mathbb{N}}$ est absolument sommable. De plus si $(r_n)_{n \in \mathbb{N}}$ et $(s_n)_{n \in \mathbb{N}}$ sont les suites des restes de $((u_n))_{n \in \mathbb{N}}$ et $((v_n))_{n \in \mathbb{N}}$, alors $(r_n)_{n \in \mathbb{N}} = O((s_n)_{n \in \mathbb{N}})$

En effet si $(|u_n|)_{n \in \mathbb{N}} = O((v_n)_{n \in \mathbb{N}})$, il existe N et A tels que pour $n \geq N$, on ait $|u_n| \leq A v_n$. Il en résulte que les valeurs absolues des u_n sont majorées (à partir du rang N) par ceux de la série positive sommable $((A v_n))_{n \in \mathbb{N}}$; donc $((u_n))_{n \in \mathbb{N}}$ est absolument sommable. De plus:

Pour $p \geq n \geq N$

$$\left| \sum_{i=n}^p u_i \right| \leq \sum_{i=n}^p A v_i = A \sum_{i=n}^p v_i$$

Donc, pour $n \geq N$:

$$\left| \sum_{i=n}^{\infty} u_i \right| \leq \sum_{i=n}^{\infty} A v_i = A \sum_{i=n}^{\infty} v_i$$

Ce qui donne la majoration annoncée du reste.

La relation $\mathcal{O}$: Si $(u_n)_{n \in \mathbb{N}} = \mathcal{O}((v_n)_{n \in \mathbb{N}})$, et si $((v_n))_{n \in \mathbb{N}}$ est une série positive sommable, alors la série $((u_n))_{n \in \mathbb{N}}$ est absolument sommable. De plus si $(r_n)_{n \in \mathbb{N}}$ et $(s_n)_{n \in \mathbb{N}}$ sont les suites des restes de $((u_n))_{n \in \mathbb{N}}$ et $((v_n))_{n \in \mathbb{N}}$, alors $(r_n)_{n \in \mathbb{N}} = \mathcal{O}((s_n)_{n \in \mathbb{N}})$.

En effet si $(u_n)_{n \in \mathbb{N}} = \mathcal{O}((v_n)_{n \in \mathbb{N}})$, quelque soit ε , il existe N tel que pour $n \geq N$, on ait $|u_n| \leq \varepsilon v_n$. Il en résulte que les valeurs absolues des u_n sont majorées (à partir du rang N) par ceux de la série positive sommable $((\varepsilon v_n))_{n \in \mathbb{N}}$; donc $((u_n))_{n \in \mathbb{N}}$ est absolument sommable. De plus:

Pour $p \geq n \geq N$

$$\left| \sum_{i=n}^p u_i \right| \leq \sum_{i=n}^p \varepsilon v_i = \varepsilon \sum_{i=n}^p v_i$$

Donc, pour $n \geq N$:

$$\left| \sum_{i=n}^{\infty} u_i \right| \leq \sum_{i=n}^{\infty} \varepsilon v_i = \varepsilon \sum_{i=n}^{\infty} v_i$$

Ce qui donne la majoration annoncée du reste.

Séries équivalentes: Si $(|u_n|)_{n \in \mathbb{N}} \sim (|v_n|)_{n \in \mathbb{N}}$, alors la série $((u_n))_{n \in \mathbb{N}}$ est absolument sommable, si et seulement si $((v_n))_{n \in \mathbb{N}}$ est absolument sommable.

En effet si $(|u_n|)_{n \in \mathbb{N}} \sim (|v_n|)_{n \in \mathbb{N}}$, quelque soit ε , il existe N tel que pour $n \geq N$, on ait les inégalités $(1-\varepsilon)|v_n| \leq |u_n| \leq (1+\varepsilon)|v_n|$. Il en résulte que les séries $((|v_n|))_{n \in \mathbb{N}}$ et $((|u_n|))_{n \in \mathbb{N}}$ sont toutes deux sommables, ou toutes deux non sommables.

On prendra garde que les règles ci-dessus sont valables pour la sommabilité absolue, et n'ont pas d'équivalent dans le cas de la sommabilité non absolue; comme en témoignent les exemples suivants.

Exercice d : a) Soit $u_n = \frac{(-1)^n}{\log n}$ et $v_n = \frac{(-1)^n}{\log n} + \frac{1}{n}$. Montrer que les suites $(u_n)_{n \geq 2}$ et $(v_n)_{n \geq 2}$ sont équivalentes, mais que la série $((u_n))_{n \geq 2}$ est sommable (on pourra s'inspirer du §6 ci-dessous), tandis que $((v_n))_{n \geq 2}$ ne l'est pas (on pourra utiliser le fait que $((v_n - u_n))_{n \geq 2}$ n'est pas sommable - voir Ex c).

b) On pose $w_n = u_n - v_n$. Montrer que $(w_n)_{n \geq 2} = o((u_n)_{n \geq 2})$, mais que la série $((u_n))_{n \geq 2}$ est sommable, tandis que $((w_n))_{n \geq 2}$ ne l'est pas.

Exercice e : Montrer que si $((u_n))_{n \in \mathbb{N}}$ et $((v_n))_{n \in \mathbb{N}}$ sont deux séries sommables, alors $((u_n + v_n))_{n \in \mathbb{N}}$ est une série sommable. Montrer que si $((u_n))_{n \in \mathbb{N}}$ et $((v_n))_{n \in \mathbb{N}}$ sont absolument sommables, alors $((u_n + v_n))_{n \in \mathbb{N}}$ est absolument sommable.

§ 6 : Séries qui ne sont pas absolument sommables.

Il existe des séries qui sont sommables mais pas absolument sommables. L'exemple le plus simple est la série $((\frac{(-1)^n}{n}))_{n \geq 1}$. Nous savons qu'elle n'est pas absolument sommable (exercice c). Pour montrer qu'elle est sommable il suffit de lui appliquer le critère suivant:

Critère des séries alternées: Si la suite $(u_n)_{n \in \mathbb{N}}$ est à termes positifs, décroissante (éventuellement à partir d'un certain rang), et converge vers zéro, alors la série $(((-1)^n u_n))_{n \in \mathbb{N}}$ est sommable.

En effet en notant U_n ses sommes partielles nous avons $U_{2n+2} = U_{2n} + (u_{2n+2} - u_{2n+1})$, donc la suite $(U_{2n})_{n \in \mathbb{N}}$ est décroissante. De même $U_{2n+1} = U_{2n-1} + (u_{2n} - u_{2n+1})$, donc la suite $(U_{2n+1})_{n \in \mathbb{N}}$ est croissante. Par ailleurs pour tout n : $U_{2n} = U_{2n-1} + u_{2n} \geq U_{2n-1} \geq U_{2n-3} \geq \dots \geq U_1$; ainsi la suite décroissante $(U_{2n})_{n \in \mathbb{N}}$ est minorée par U_1 , donc elle est convergente. Notons λ sa limite. Puisque la suite $(U_{2n} - U_{2n-1})_{n \in \mathbb{N}} = (u_{2n})_{n \in \mathbb{N}}$ converge vers zéro, la suite $(U_{2n-1})_{n \in \mathbb{N}}$ a aussi pour limite λ . Il en résulte que la suite $(U_n)_{n \in \mathbb{N}}$ converge vers λ .

Le critère d'Abel: Si la suite $(a_n)_{n \in \mathbb{N}}$ est positive décroissante (à partir d'un certain rang) et converge vers 0, alors:

1) $(\forall t)$ la série $((a_n \sin nt))_{n \in \mathbb{N}}$ est sommable.

2) $(\forall t \neq 2k\pi)$ la série $((a_n \cos nt))_{n \in \mathbb{N}}$ est sommable.

On remarquera que le critère des séries alternées est un cas particulier de l'assertion 2 (celui où $t = \pi$).

Démonstration : Posons $S_n(t) = \sum_{k=0}^n \sin kt$ et $T_n(t) = \sum_{k=0}^n \cos kt$. Si $t \neq 0 \pmod{2\pi}$ nous avons alors:

$$T_n(t) + iS_n(t) = \sum_{k=0}^n e^{ikt} = \frac{1 - e^{i(n+1)t}}{1 - e^{it}}$$

Donc $(T_n(t))^2 + (S_n(t))^2 = \left| \frac{1 - e^{i(n+1)t}}{1 - e^{it}} \right|^2 \leq \left| \frac{2}{1 - e^{it}} \right|^2 = \frac{4}{2 - 2\cos t} = \frac{1}{\sin^2(\frac{t}{2})}$. Par conséquent $|T_n(t)|$ et

$|S_n(t)|$ sont $(\forall n)$ au plus égaux à $K = \frac{1}{\sin \frac{t}{2}}$

Nous pouvons alors écrire:

$$\begin{aligned} \sum_{k=0}^n a_k \sin kt &= S_0(t)(a_0 - a_1) + S_1(t)(a_1 - a_2) + S_2(t)(a_2 - a_3) + \dots + S_{n-1}(t)(a_{n-1} - a_n) + S_n(t) a_n \\ \sum_{k=0}^n a_k \cos kt &= T_0(t)(a_0 - a_1) + T_1(t)(a_1 - a_2) + T_2(t)(a_2 - a_3) + \dots + T_{n-1}(t)(a_{n-1} - a_n) + T_n(t) a_n \end{aligned}$$

Posons $s_n(t) = S_n(t)(a_n - a_{n+1})$ et $t_n(t) = T_n(t)(a_n - a_{n+1})$. Les séries $((s_n(t)))_{n \in \mathbb{N}}$ et $((t_n(t)))_{n \in \mathbb{N}}$ sont absolument sommables, puisque $(\forall n) |s_n(t)| \leq K(a_n - a_{n+1})$ et $|t_n(t)| \leq K(a_n - a_{n+1})$, et que la série $((a_n - a_{n+1}))_{n \in \mathbb{N}}$ est une série positive sommable.

Ainsi la suite $(\sum_{k=0}^n a_k \sin kt)_{n \in \mathbb{N}}$ des sommes partielles de la série $((a_n \sin nt))_{n \in \mathbb{N}}$ est la

somme d'une suite qui converge vers 0 (la suite $(a_n S_n(t))_{n \in \mathbb{N}}$), et d'une suite convergente (la suite des sommes partielles de la série $((s_n(t)))_{n \in \mathbb{N}}$); ce qui montre que la série $((a_n \sin nt))_{n \in \mathbb{N}}$ est sommable.

On notera que le raisonnement ne s'applique pas lorsque $t = 0 \pmod{2\pi}$, car il est alors impossible de majorer S_n et T_n . Dans ce cas la série $((a_n \sin nt))_{n \in \mathbb{N}}$ est sommable puisqu'elle est nulle; tandis que la série $((a_n \cos nt))_{n \in \mathbb{N}}$ n'a aucune raison d'être sommable.

Comment démontrer qu'une série à termes de signe quelconque n'est pas sommable?

Pour terminer nous donnerons un critère permettant de démontrer qu'une série n'est pas sommable. Si $((u_n))_{n \in \mathbb{N}}$ est sommable, la suite $(U_n)_{n \in \mathbb{N}}$ a une limite λ ; et puisque la suite $(U_{n-1})_{n \geq 1}$ converge aussi vers λ ; et la suite $(U_n - U_{n-1})_{n \geq 1} = (u_n)_{n \geq 1}$ converge vers 0.

Ainsi lorsque la suite $(u_n)_{n \in \mathbb{N}}$ ne converge pas vers zéro, la série $((u_n))_{n \in \mathbb{N}}$ n'est certainement pas sommable.

§ 7 : Les séries $((\frac{1}{n^\alpha}))_{n \geq 1}$

Pour utiliser les règles de comparaison ci-dessus, il nous faut connaître le comportement de quelques séries de référence. Parmi celles ci nous retiendrons les séries $((\frac{1}{n^\alpha}))_{n \in \mathbb{N}}$.

Si $\alpha > 1$ la série $((\frac{1}{n^\alpha}))_{n \in \mathbb{N}}$ est sommable et son reste à l'ordre n est $O(\frac{1}{n^{\alpha-1}})$

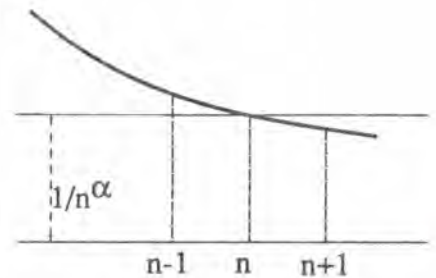
Si $\alpha \leq 1$ la série $((\frac{1}{n^\alpha}))_{n \in \mathbb{N}}$ n'est pas sommable.

En effet pour tout $n \geq 2$, nous avons les inégalités:

$$v_n = \int_{n-1}^n \frac{1}{t^\alpha} dt \geq \frac{1}{n^\alpha} \geq \int_n^{n+1} \frac{1}{t^\alpha} dt = w_n \quad (\alpha > 0)$$

$$\text{Or } \sum_{i=2}^n v_i = \int_1^n \frac{1}{t^\alpha} dt = \begin{cases} \log n & \text{si } \alpha = 1 \\ \frac{1}{1-\alpha} \frac{1}{n^{\alpha-1}} - \frac{1}{1-\alpha} & \text{si } \alpha \neq 1 \end{cases}$$

$$\text{Et } \sum_{i=2}^n w_i = \int_2^{n+1} \frac{1}{t^\alpha} dt = \begin{cases} \log(n+1) - \log 2 & \text{si } \alpha = 1 \\ \frac{1}{1-\alpha} \frac{1}{(n+1)^{\alpha-1}} - \frac{1}{1-\alpha} \frac{1}{2^{\alpha-1}} & \text{si } \alpha \neq 1 \end{cases}$$



Donc pour $\alpha > 1$, la série $((v_n))_{n \in \mathbb{N}}$ est sommable, et $((\frac{1}{n^\alpha}))_{n \in \mathbb{N}}$ l'est également. Tandis que si $\alpha \leq 1$, la série $((w_n))_{n \in \mathbb{N}}$ n'est pas sommable, et $((\frac{1}{n^\alpha}))_{n \in \mathbb{N}}$ ne l'est pas non plus.

Quelques conséquences:

Si $(|u_n|)_{n \in \mathbb{N}} = O((\frac{1}{n^\alpha}))_{n \in \mathbb{N}}$ avec $\alpha > 1$, alors la série $((u_n))_{n \in \mathbb{N}}$ est absolument sommable. Donc

à fortiori, si $(|u_n|)_{n \in \mathbb{N}} = o((\frac{1}{n^\alpha}))_{n \in \mathbb{N}}$ avec $\alpha > 1$, alors la série $((u_n))_{n \in \mathbb{N}}$ est absolument sommable.

Si $(|u_n|)_{n \in \mathbb{N}} \sim ((\frac{k}{n^\alpha}))_{n \in \mathbb{N}}$ avec $\alpha \leq 1$, alors la série $((u_n))_{n \in \mathbb{N}}$ n'est pas absolument sommable.

Exercice f : Démontrer que, pour tout $n \geq 3$, $\frac{1}{n \log n} \geq \int_n^{n+1} \frac{dt}{t \log t}$. En déduire que la série $((\frac{1}{n \log n}))_{n \in \mathbb{N}}$ n'est pas sommable.

§ 8 : Les séries géométriques

Ce sont les séries de la forme $((\alpha^n))_{n \in \mathbb{N}}$. Nous connaissons leurs sommes partielles :

$$\sum_{i=0}^n \alpha^i = \frac{1 - \alpha^{n+1}}{1 - \alpha}$$

Elles sont sommables si $|\alpha| < 1$, et ne le sont pas si $|\alpha| \geq 1$. Deux règles de convergence classiques résultent de la comparaison aux séries géométriques.

Règle de d'Alembert.

Si la suite $(\frac{u_n}{u_{n-1}})_{n \in \mathbb{N}}$ est minorée à partir d'un certain rang par un nombre $k \geq 1$

(à fortiori si elle converge vers une limite $\lambda > 1$), alors la série $((u_n))_{n \in \mathbb{N}}$ n'est pas sommable

Si la suite $(\frac{u_n}{u_{n-1}})_{n \in \mathbb{N}}$ est majorée à partir d'un certain rang par un nombre $k < 1$ (à fortiori si

elle converge vers une limite $\lambda < 1$), alors la série $((u_n))_{n \in \mathbb{N}}$ est absolument sommable.

En effet si pour $n \geq N$, $|\frac{u_n}{u_{n-1}}| \leq k < 1$, alors $|u_n| \leq u_N k^{n-N} = (u_N k^{-N}) k^n$. Ainsi $|u_n|$ est, à partir du rang N , majorée par une série géométrique sommable, donc elle est sommable.

Inversement si $|\frac{u_n}{u_{n-1}}| \geq k \geq 1$, alors $|u_n| \geq (u_N k^{-N}) k^n$, donc la suite $(|u_n|)_{n \in \mathbb{N}}$ ne converge pas vers 0.

Exercice g : L'énoncé ci-dessus affirme que si $|\frac{u_n}{u_{n-1}}|$ converge vers $\lambda < 1$, alors il existe un nombre $k < 1$ tel que, à partir d'un certain rang, $|\frac{u_n}{u_{n-1}}| < k$. Expliquer cette affirmation.

Règle de Cauchy.

Si la suite $(\sqrt[n]{|u_n|})_{n \in \mathbb{N}}$ reste, à partir d'un certain rang, inférieure à $k < 1$ (à fortiori si elle converge vers $\lambda < 1$), alors la série $((u_n))_{n \in \mathbb{N}}$ est absolument sommable.

Si la suite $(\sqrt[n]{|u_n|})_{n \in \mathbb{N}}$ reste, à partir d'un certain rang, supérieure à $k \geq 1$ (à fortiori si elle converge vers $\lambda > 1$), alors la série $((u_n))_{n \in \mathbb{N}}$ n'est pas sommable.

En effet si $\sqrt[n]{|u_n|} \leq k$, c'est que $|u_n| \leq k^n$, et puisque $k < 1$, la série géométrique $((k^n))_{n \in \mathbb{N}}$ est sommable, donc $((|u_n|))_{n \in \mathbb{N}}$ est sommable. Inversement si $\sqrt[n]{|u_n|} \geq k$, c'est que $|u_n| \geq k^n$, et puisque $k \geq 1$, la suite $(|u_n|)_{n \in \mathbb{N}}$ ne converge pas vers 0.

Complément : Esquisse d'une construction de $\mathbb{R}$.

Jusqu'au début du 19^{ème} siècle, la notion de nombre avait pour origine la mesure des grandeurs physiques, et en particulier la mesure des longueurs. L'ensemble des nombres était défini par un système d'axiomes, qui n'était que la formalisation des propriétés physiques évidentes d'une droite orientée. Les mathématiciens trouvaient ainsi leur unité autour de la géométrie, telle qu'Euclide (-323, -285) l'avait formalisée dans un texte ('les éléments') qui apparaissait à tous comme la base intangible et définitive.

Nous avons vu au chapitre 1 comment, pour démontrer des théorèmes d'existence, il fallut ajouter un nouvel axiome. Quelle que soit la formulation que l'on en donne (borne supérieure, suites majorées, suites de

Cauchy,...), ce dernier ne correspond à aucune propriété évidente de la droite physique. À la même époque la confiance que l'on avait dans l'œuvre d'Euclide fut ébranlée; depuis l'antiquité l'axiomatique des 'Eléments' posait problème: le postulat des parallèles semblait 'physiquement peu évident'; on cherchait donc à le démontrer, mais (vers 1845) on s'aperçut que c'était impossible; ceci fut considéré comme une faille dans le texte qui servait de fondement à toutes les mathématiques. Ce sont quelques unes des raisons qui amenèrent la question qui nous intéresse: est-il possible de construire logiquement l'ensemble des nombres réels sans aucune référence à une axiomatique d'origine géométrique? L'étape essentielle de cette construction (la construction de $\mathbb{R}$ à partir de $\mathbb{Q}$) fut donnée (vers 1880) par Cantor et Dedekind.

Qu'est ce qu'un réel ?: Supposons que nous connaissions l'ensemble $\mathbb{Q}$ des nombres rationnels; comment construire l'ensemble $\mathbb{R}$ (qui doit être plus gros que $\mathbb{Q}$, et doit vérifier la propriété de la borne supérieure, ou une des propriétés équivalentes)?

Notons $\mathcal{C}$ l'ensemble des suites de Cauchy de nombres rationnels. Chaque élément de $\mathcal{C}$ doit avoir une limite dans $\mathbb{R}$. Nous allons donc décider que, pour tout élément $(u_n)_{n \in \mathbb{N}}$ de $\mathcal{C}$, il existe un nombre réel $\lambda(u_n)$ (qui apparaîtra plus tard comme sa limite).

Mais deux suites distinctes peuvent avoir la même limite. C'est pourquoi nous décidons que si la suite $(u_n \cdot v_n)_{n \in \mathbb{N}}$ converge vers 0, réels $\lambda(u_n)_{n \in \mathbb{N}}$ et $\lambda(v_n)_{n \in \mathbb{N}}$ doivent être identiques.

Tout nombre réel est obtenu ainsi, donc λ apparaît comme une application surjective de $\mathcal{C}$ dans $\mathbb{R}$. Elle n'est pas injective, et $\lambda(u_n) = \lambda(v_n)$ signifie que $(u_n \cdot v_n)_{n \in \mathbb{N}}$ a pour limite 0.

Les nombres réels apparaissent ainsi comme des entités abstraites, introduites pour les besoins de notre construction. Mais on peut pousser l'analyse de la situation un peu plus loin. Le nombre réel $\lambda(u_n)$ est associé à tout un ensemble de suites: tous les éléments $(v_n)_{n \in \mathbb{N}}$ de $\mathcal{C}$ tels que $(v_n \cdot u_n)_{n \in \mathbb{N}}$ converge vers 0. On peut considérer que $\lambda(u_n)$ est cet ensemble de suites. Les nombres réels sont alors des sous-ensembles de $\mathcal{C}$. Tous ces sous-ensembles sont deux à deux disjoints. Leur réunion est $\mathcal{C}$ tout entier. On dit qu'ils forment une partition de $\mathcal{C}$.

Nous avons ainsi défini $\mathbb{R}$ comme ensemble, il reste à définir les lois de composition qui en feront un corps, et à démontrer ses propriétés.

Comment on fait de $\mathbb{Q}$ un sous-ensemble de $\mathbb{R}$: À tout rationnel α , on peut associer la suite constante $(\forall n) u_n = \alpha$. C'est évidemment une suite de Cauchy. Nous lui avons associé un réel, notons le $\iota(\alpha)$. Nous obtenons une application ι de $\mathbb{Q}$ dans $\mathbb{R}$. Elle est injective (car $\iota(\alpha) = \iota(\beta)$ signifie que la suite constante $\beta - \alpha$ converge vers 0). Si nous décidons d'identifier α à son image $\iota(\alpha)$, nous faisons de $\mathbb{Q}$ un sous-ensemble de $\mathbb{R}$.

L'addition et la multiplication: Soit α et β deux réels. Choisissons $(a_n)_{n \in \mathbb{N}}$ et $(b_n)_{n \in \mathbb{N}}$ tels que $\lambda(a_n) = \alpha$ et $\lambda(b_n) = \beta$. Par définition $\alpha + \beta$ sera le nombre réel $\lambda(a_n + b_n)$. Mais attention !!! Cette construction comporte des choix (celui des deux suites $(a_n)_{n \in \mathbb{N}}$ et $(b_n)_{n \in \mathbb{N}}$); pour qu'elle soit acceptable il faut que le réel $\lambda(a_n + b_n)$ soit indépendant de ces choix (car $\alpha + \beta$ ne doit dépendre que de α et β). Il faut donc vérifier que si nous choisissons $(u_n)_{n \in \mathbb{N}}$ et $(v_n)_{n \in \mathbb{N}}$ telles que $\lambda(u_n) = \alpha$ (i.e.: $(a_n - u_n)_{n \in \mathbb{N}}$ converge vers 0) et que $\lambda(v_n) = \beta$ (i.e.: $(b_n - v_n)_{n \in \mathbb{N}}$ converge vers 0), alors $\lambda(u_n + v_n) = \lambda(a_n + b_n)$. Ce qui est clair puisque

$$((u_n + v_n) - (a_n + b_n))_{n \in \mathbb{N}} = (u_n - a_n)_{n \in \mathbb{N}} + (v_n - b_n)_{n \in \mathbb{N}}.$$

La définition du produit est en tout point analogue: si $\lambda(a_n) = \alpha$ et $\lambda(b_n) = \beta$, alors (par définition) $\alpha\beta$ est $\lambda(a_n b_n)$. Et pour que ceci soit logiquement acceptable, il faut démontrer que si $\lambda(a_n) = \lambda(u_n)$ et $\lambda(b_n) = \lambda(v_n)$, alors $\lambda(a_n b_n) = \lambda(u_n v_n)$. Ceci résulte de l'égalité

$$a_n b_n - u_n v_n = a_n (b_n - v_n) + v_n (a_n - u_n)$$

Les suites $(a_n)_{n \in \mathbb{N}}$ et $(v_n)_{n \in \mathbb{N}}$ sont bornées (car elles sont de Cauchy). Les suites $(b_n - v_n)_{n \in \mathbb{N}}$ et $(a_n - u_n)_{n \in \mathbb{N}}$ convergent vers 0; donc la suite $(a_n b_n - u_n v_n)_{n \in \mathbb{N}}$ converge vers 0.

Nous avons ainsi fait de $\mathbb{R}$ un corps. Pour s'en convaincre il faut vérifier les propriétés énoncées au chap 1 §7. C'est facile mais assez fastidieux. Démontrons par exemple la distributivité de la multiplication par rapport à l'addition. Il faut montrer que $(\alpha + \beta)\gamma = \alpha\gamma + \beta\gamma$. Nous choisissons dans $\mathcal{S}$ des suites telles que $\lambda(a_n) = \alpha$, $\lambda(b_n) = \beta$ et $\lambda(c_n) = \gamma$. Alors $\alpha + \beta = \lambda(a_n + b_n)$ et $(\alpha + \beta)\gamma = \lambda((a_n + b_n)c_n)$. Tandis que $\alpha\gamma + \beta\gamma = \lambda(a_n c_n + b_n c_n)$. Et il est clair que la suite $((a_n + b_n)c_n - (a_n c_n + b_n c_n))_{n \in \mathbb{N}}$ converge vers 0.

Toutes les autres démonstrations sont analogues. Et pour être tout à fait correct, il faut vérifier que l'injection ι de $\mathbb{Q}$ dans $\mathbb{R}$ commute à l'addition et à la multiplication (c'est à dire que la somme (ou le produit) de deux rationnels est la même quand on les considère comme des éléments de $\mathbb{Q}$ ou comme des éléments de $\mathbb{R}$).
La comparaison des réels et la valeur absolue: Nous dirons que $\alpha \leq \beta$ s'il existe des suites telles que $\lambda(a_n) = \alpha$, $\lambda(b_n) = \beta$ et $a_n \leq b_n$ à partir d'un certain rang.

Nous pouvons aussi définir la valeur absolue d'un réel: si $\alpha = \lambda(a_n)$, alors $|\alpha| = \lambda(|a_n|)$ (attention il y a encore quelque chose à vérifier pour que ce soit logiquement correct). Et ainsi nous pourrions parler de suites de réels convergentes, et de suites de Cauchy de réels, puisque tous les termes de la définition ont été définis.

Il reste à démontrer que toute suite de Cauchy (de réels) est convergente, et que la limite d'un élément $(u_n)_{n \in \mathbb{N}}$ de $\mathcal{S}$ est $\lambda(u_n)$. Ceci se fera en trois étapes.

Démontrons d'abord qu'entre deux réels distincts il existe (au moins) un rationnel. Soit $\alpha < \beta$. Puisque $\alpha \leq \beta$ nous pouvons écrire $\alpha = \lambda(u_n)$ et $\beta = \lambda(v_n)$, où la suite $(v_n - u_n)_{n \in \mathbb{N}}$ est positive (à partir d'un certain rang). Puisque $\alpha \neq \beta$, la suite $(v_n - u_n)_{n \in \mathbb{N}}$ ne converge pas vers 0. Il existe donc un rationnel $\varepsilon (> 0)$ et un entier N tels que $n > N$ implique $v_n - u_n \geq \varepsilon$. Puisque les suites $(v_n)_{n \in \mathbb{N}}$ et $(u_n)_{n \in \mathbb{N}}$ sont de Cauchy, il existe $P \geq N$ tel que, quel que soit $n \geq P$, $|u_n - u_P| \leq \frac{\varepsilon}{3}$ et $|v_n - v_P| \leq \frac{\varepsilon}{3}$. Alors le rationnel $u_P + \frac{\varepsilon}{2}$ est compris entre α et β .

Limites des suites de Cauchy de rationnels: Soit $(u_n)_{n \in \mathbb{N}}$ un élément de $\mathcal{S}$; considérons le comme une suite de Cauchy dans $\mathbb{R}$. Montrons que l'on obtient ainsi une suite qui converge (dans $\mathbb{R}$) vers $\lambda(u_n)$. Soit ε un réel strictement positif, nous savons qu'il existe un rationnel ε' tel que $0 < \varepsilon' < \varepsilon$. Puisque la suite $(u_n)_{n \in \mathbb{N}}$ est de Cauchy dans $\mathbb{Q}$, il existe N tel que $(\forall n \geq N \text{ et } (\forall p \geq N) |u_n - u_p| \leq \varepsilon')$. Notons $(v_n^p)_{n \in \mathbb{N}}$ la suite constante égale à u_p . L'inégalité $|u_n - u_p| \leq \varepsilon'$ peut s'écrire $|u_n - v_n^p| \leq \varepsilon'$; ce qui implique $|\lambda(u_n) - \lambda(v_n^p)| \leq \varepsilon'$. En notant u_p le réel $\lambda(v_n^p)$ identifié à u_p , nous obtenons $|\lambda(u_n) - u_p| \leq \varepsilon'$; et puisque ceci est valable pour tout $p \geq N$, nous avons démontré que la suite $(u_n)_{n \in \mathbb{N}}$ (considérée comme une suite de réels) converge vers le réel $\lambda(u_n)$.

Limites des suites de Cauchy de réels: Soit maintenant $(v_n)_{n \in \mathbb{N}}$ une suite de Cauchy de réels, nous supposons qu'elle n'est pas stationnaire (si elle était stationnaire elle aurait clairement une limite). Ainsi pour tout n il existe un $p \geq n$, tel que $u_p \neq u_n$; choisissons un tel p et choisissons un rationnel v_n compris entre u_n et u_p (on sait qu'il en existe). Nous obtenons ainsi une suite $(v_n)_{n \in \mathbb{N}}$ de rationnels. On vérifie facilement qu'elle est de Cauchy (car $(u_n)_{n \in \mathbb{N}}$ est de Cauchy). Elle converge donc vers $\lambda(v_n)$. Et on montre que $(u_n)_{n \in \mathbb{N}}$ converge aussi vers $\lambda(v_n)$.

Ceci termine notre construction. Nous avons, par des constructions logiques, défini $\mathbb{R}$ à partir de $\mathbb{Q}$. Ainsi les nombres réels existent en dehors de tout support géométrique (donc physique).

Exercices sur le chapitre 8

* **Exercice 1:** Montrer la sommabilité et calculer la somme de

$$a = \left(\left(\frac{1}{n(n+1)(n+2)} \right) \right)_{n \geq 1}$$

$$b = \left(\left(\frac{4n-3}{n(n^2-4)} \right) \right)_{n \geq 3}$$

$$c = \left(\left(\text{Log} \frac{(n+1)(n+3)}{n(n+2)} \right) \right)_{n \geq 1}$$

$$d = \left(\left(\text{Arctg} \frac{n+1}{n} - \text{Arctg} \frac{n+2}{n+1} \right) \right)_{n \geq 1}$$

* **Exercice 2:** Etudier la sommabilité et la sommabilité absolue des séries suivantes

$$u = \left(\left(1 - \cos \frac{1}{n} \right) \right)_{n \geq 0}$$

$$v = \left(\left(\frac{2}{\sin \frac{2}{n}} - \frac{1}{\sin \frac{1}{n}} \right) \right)_{n \geq 0}$$

$$w = \left(\left(\frac{1+2+\dots+n}{1+2^2+\dots+n^2} \right) \right)_{n \geq 1}$$

$$x = \left(\left(\text{th} \left(n + \frac{1}{2} \right) - \text{th} n \right) \right)_{n \in \mathbb{N}}$$

$$y = \left(\left((-1)^n \left(\sin \frac{2}{n} - \sin \frac{1}{n} \right) \right) \right)_{n \geq 0}$$

$$z = \left(\left(\frac{1}{\text{Log}(\text{ch}(an))} \right) \right)_{n \in \mathbb{N}}$$

* **Exercice 3:**

a) On considère la série $u = \left(\left(\frac{e^{-n}}{n} \right) \right)_{n \geq 1}$. Montrer qu'elle est sommable.

b) Comparer u à la série $v = (e^{-n})_{n \in \mathbb{N}}$. Sommer la série v .

c) Quelle somme partielle va-t-on calculer pour obtenir une valeur approchée de sa somme à 10^{-5} près ?

* **Exercice 4:** On considère les séries $u = \left(\left(\frac{1}{n^2+n+1} \right) \right)_{n \in \mathbb{N}}$ et $v = \left(\left(\frac{1}{n^2+n} \right) \right)_{n \geq 1}$.

a) Montrer qu'elles sont sommables, et calculer la somme de v .

b) Montrer que $(\forall n \geq 2) u_n \leq v_n \leq u_{n-1}$. Calculer une valeur approchée à 10^{-4} près de la somme de u ?

* **Exercice 5:** a) Montrer que les séries suivantes sont absolument sommables

$$u = \left(\left(\text{Arctg} \left(n + \frac{1}{2} \right) - \text{Arctg} n \right) \right)_{n \in \mathbb{N}}$$

$$v = \left(\left(\text{Arctg}(n+1) - \text{Arctg} \left(n - \frac{1}{2} \right) \right) \right)_{n \in \mathbb{N}}$$

$$w = \left(\left(\sin \left(n + \frac{1}{n^2} \right) - \sin n \right) \right)_{n \geq 1}$$

b) On veut calculer leurs sommes avec une erreur inférieure à 10^{-3} . Quelle somme partielle va-t-on calculer ?

* **Exercice 6:** Soit $((u_n))_{n \in \mathbb{N}}$ une série sommable. On pose $(\forall n) v_n = \sup(u_n, 0)$. Montrer que $((v_n))_{n \in \mathbb{N}}$ est sommable si et seulement si $((u_n))_{n \in \mathbb{N}}$ est absolument sommable.

** **Exercice 7:** Les séries suivantes sont-elles sommables, sont-elles absolument sommables ?

$$a = \left(\left(\frac{n!}{e^n} \right) \right)_{n \in \mathbb{N}}$$

$$b = \left(\left(\frac{n!}{(2n)!} \right) \right)_{n \in \mathbb{N}}$$

$$c = ((\text{sh } n^\alpha))_{n \geq 1} \quad (\alpha \in \mathbb{R})$$

$$d = \left(\left(\left(\frac{1+n}{1+n^2} \right)^{\text{Log } n} \right) \right)_{n \geq 1}$$

$$e = ((\text{Log } n)^{-n})_{n \geq 2}$$

$$f = (((-1)^n \sin n))_{n \in \mathbb{N}}$$

$$g = \left(\left(\left(1 - \frac{1}{n} \right)^{n^2} \right) \right)_{n \geq 1}$$

$$h = \left(\left(\frac{n^2+1}{n^3+2^n} \right) \right)_{n \in \mathbb{N}}$$

- * **Exercice 8:** Les séries suivantes sont-elles sommables, sont-elles absolument sommables:
- $$u = ((\operatorname{Log}(\cos \frac{1}{n}))_{n \geq 1} \quad v = ((e^{-n^2}))_{n \in \mathbb{N}}$$
- $$w = ((\sin(n\pi + e^{-n})))_{n \in \mathbb{N}} \quad x = ((\cos(n + \frac{1}{n^2})))_{n \geq 1}$$
- ** **Exercice 9:** Pour quelle(s) valeur(s) de la constante a , la série définie par $u_n = \sqrt[3]{n^3 + an} - \sqrt{n^2 + 1}$ est-elle sommable ?
- * **Exercice 10:** Montrer que la série définie par $u_n = \frac{n+3}{n(n+1)(n+2)}$ est sommable, et calculer sa somme.
- * **Exercice 11:** Les séries définies par
- $$u_n = \operatorname{Log}(\operatorname{ch}(\sin \frac{1}{n})) \quad v_n = \operatorname{Log}(\cos \frac{1}{n} + \sin \frac{1}{n}) \quad w_n = \operatorname{sh}(\sin \frac{1}{n}) - \sin(\operatorname{sh} \frac{1}{n})$$
- sont-elles sommables ? sont-elles absolument sommables ?
- (Extrait d'un partiel de 1981)
- * **Exercice 12:** Discuter suivant les valeurs de α , la sommabilité et la sommabilité absolue des séries définies par:
- $$u_n = \operatorname{ch}(1 + \frac{1}{n^\alpha}) - \operatorname{ch}(1 - \frac{1}{n^\alpha}) \quad v_n = \frac{\operatorname{ch}(2n+1) - \operatorname{ch} 2n}{e^{n\alpha}}$$
- $$w_n = n^\alpha (\operatorname{Arctg}(2n+1) - \operatorname{Arctg} 2n)$$
- ** **Exercice 13:** La série définie par $u_n = \int_0^n e^{-t^2} \sin nt \, dt$ est-elle sommable ? Est-elle absolument sommable (on pourra modifier l'expression par une ou plusieurs intégrations par parties).
- ** **Exercice 14:** La série définie par $u_n = \int_0^n e^{-t^3} \cos nt \, dt$ est-elle sommable ? Est-elle absolument sommable (on pourra modifier l'expression par une ou plusieurs intégrations par parties).
- (Extrait d'un partiel de 1981)
- * **Exercice 15:** On définit une série $((u_n))_{n \in \mathbb{N}}$ par $u_0 = 1$, $u_1 = 2$, et (pour $n \geq 2$) $u_n = u_{n-1} - \alpha u_{n-2}$. Comment faut-il choisir α pour que cette série soit sommable ?
- ** **Exercice 16:** Soit $F(x) = \frac{P(x)}{Q(x)}$ une fraction rationnelle. On suppose $d^\circ P < d^\circ Q$. Montrer que la série $u = ((F(n) \cos n))_{n \in \mathbb{N}}$ est sommable. A quelle condition est-elle absolument sommable ?
- ** **Exercice 17:** a) Soit $u = ((u_n))_{n \in \mathbb{N}}$ et $v = ((v_n))_{n \in \mathbb{N}}$ deux séries de carré sommable (i.e.: Les séries $((u_n^2))_{n \in \mathbb{N}}$ et $((v_n^2))_{n \in \mathbb{N}}$ sont sommables). Montrer que la série $((u_n v_n))_{n \in \mathbb{N}}$ est sommable (on utilisera l'inégalité $(\forall x) (\forall y) |xy| \leq x^2 + y^2$).
- b) Soit E l'ensemble de toutes les séries numériques de carré sommable. Montrer que E est un espace vectoriel.
- c) Montrer qu'en posant $\langle u, v \rangle = \sum_{n \geq 0} u_n v_n$, on définit un produit scalaire sur E .
- * **Exercice 18:** Au moyen d'une calculatrice, déterminer à 10^{-6} près la somme de la série $((e^{-n^2}))_{n \in \mathbb{N}}$ (on pourra utiliser le fait que, pour $a \geq 1$: $e^{-a^2} \leq e^{-a}$).

** Exercice 19: On veut calculer la somme de la série $((\frac{1}{n^2}))_{n \geq 0}$.

a) Calculer (exactement) la somme de la série $((\frac{1}{n(n+1)}))_{n \geq 0}$.

b) On pose $v_n = \frac{1}{n(n+1)} - \frac{1}{n^2}$. Au moyen d'une calculatrice déterminer à 10^{-3} près la somme de la série $((v_n))_{n \geq 0}$. En déduire une valeur approchée de la somme de $((\frac{1}{n^2}))_{n \geq 0}$.

c) On pose $w_n = \frac{1}{n(n+1)(n+2)}$. Calculer $v_n - w_n$. Calculer (exactement) la somme de la série $((w_n))_{n \geq 0}$. Puis calculer à 10^{-5} près la somme de la série $((v_n - w_n))_{n \geq 0}$. En déduire une valeur approchée de la somme de $((\frac{1}{n^2}))_{n \geq 0}$.

*** Exercice 20: a) Soit $((u_n))_{n \in \mathbb{N}}$ une série sommable. On pose $v_{2n} = u_{2n+1}$ et $v_{2n+1} = u_{2n}$. Comparer les sommes partielles de rang $2n$ et $2n+1$ des séries $((u_n))_{n \in \mathbb{N}}$ et $((v_n))_{n \in \mathbb{N}}$. En déduire que $((v_n))_{n \in \mathbb{N}}$ est sommable de même somme que $((u_n))_{n \in \mathbb{N}}$.

b) Soit $u = ((\frac{(-1)^n}{1+n}))_{n \in \mathbb{N}}$. On définit $\varphi: \mathbb{N} \rightarrow \mathbb{N}$ par $\varphi(3k) = 4k$, $\varphi(3k+1) = 4k+2$ et $\varphi(3k+2) = 2k+1$. Montrer que φ est bijective. On pose $v_n = u_{\varphi(n)}$. Soit U_n et V_n les sommes partielles à l'ordre n de $((u_n))_{n \in \mathbb{N}}$ et $((v_n))_{n \in \mathbb{N}}$. Montrer que $V_{3k+2} = U_{2k+1} + \sum_{i=k+1}^{2k+1} u_{2i}$.

c) Montrer que $\sum_{i=k+1}^{2k+1} u_{2i} \geq \sum_{i=k+1}^{2k+1} \int_1^{i+1} \frac{dt}{2t} \geq \frac{1}{2} \int_{k+1}^{2k+2} \frac{dt}{t}$; puis que $\sum_{i=k+1}^{2k+1} u_{2i} \leq \sum_{i=k+1}^{2k+1} \int_{i-1}^i \frac{dt}{2t} \leq \frac{1}{2} \int_k^{2k+1} \frac{dt}{t}$.

En déduire que v est sommable, mais n'a pas même somme que u .

Intégrales généralisées

Au chapitre 7 nous avons défini les fonctions intégrables au sens de Riemann. Pour cela nous avons décidé à priori (chap 7 § 2) de ne considérer que des fonctions bornées, et (si l'intervalle I est non borné) à support borné. Ces restrictions étaient indispensables à la construction que nous voulions faire; le but du présent chapitre est d'essayer de nous en affranchir.

On remarquera les analogies entre les constructions ci-dessous et celles qui nous mènent à la théorie des séries (chap. 8). Les problématiques sont semblables:

* La théorie des séries consiste à prévoir le comportement d'une suite à partir de renseignements sur les différences de deux termes successifs.

* Ici nous allons apprendre à prévoir le comportement à l'infini d'une fonction $x \rightarrow \int_a^x f(t) dt$,

à partir de renseignements sur sa dérivée f .

Les méthodes employées pour traiter ces problématiques sont tout aussi semblables: majorations, passage aux valeurs absolues, ... Le lecteur aura intérêt à mettre ces analogies en évidence.

Nous considérons un intervalle I . Si $I=[a,b]$, toute fonction sur I est évidemment à support borné, et dès qu'elle est continue, elle est bornée; il n'y a donc pas grand intérêt à étendre la théorie de Riemann. Dans la suite I sera donc soit non borné, soit non fermé (soit les deux); ce sera par exemple $[a, +\infty[$, ou $[a,b]$, ou $]a,b[$, ... Nous rencontrerons différents types de limites, des limites en $+\infty$, en $-\infty$, en a à gauche, en a à droite. Pour ne pas allonger inutilement le texte, on a développé en priorité le cas des limites en $+\infty$; les autres cas sont traités par analogie, les démonstrations étant données en exercice.

§ 1 : Le critère de Cauchy pour les fonctions au voisinage de $+\infty$.

Au chapitre 8, nous avons vu des critères permettant d'affirmer qu'une suite est convergente sans connaître sa limite. De la même façon il existe deux critères permettant d'affirmer qu'une fonction $f: [a, +\infty[\rightarrow \mathbb{R}$ a une limite, sans connaître cette limite.

L'un est valable pour les fonctions monotones: Toute fonction $f: [a, +\infty[\rightarrow \mathbb{R}$ croissante majorée a une limite en $+\infty$. Toute fonction $f: [a, +\infty[\rightarrow \mathbb{R}$ décroissante minorée a une limite en $+\infty$.

L'autre est valable en toute généralité: Pour qu'une fonction $f: [a, +\infty[\rightarrow \mathbb{R}$ ait une limite en $+\infty$, il faut et il suffit qu'elle vérifie la condition (de Cauchy):

$$(\forall \varepsilon > 0) (\exists A \geq a) (\forall x) (\forall y) (x \geq A \text{ et } y \geq A) \text{ implique } |f(x) - f(y)| \leq \varepsilon.$$

Démonstration du premier critère: Considérons la suite $(f(n))_{n \in \mathbb{N}}$. Elle est croissante (parce que f est croissante) et majorée (parce que f est majorée), donc elle a une limite λ . Montrons que λ est la limite de f en $+\infty$. Nous savons que:

$$(\forall \varepsilon > 0) (\exists N) (\forall n \in \mathbb{N}) (n \geq N) \text{ implique } |f(n) - \lambda| \leq \varepsilon \text{ (c'est à dire } \lambda - \varepsilon \leq f(n) \leq \lambda)$$

Choisissons $x \geq N$, alors $E(x)$ (ie: partie entière de x) et $E(x)+1$ sont des entiers supérieurs ou égaux à N , donc $\lambda - \varepsilon \leq f(E(x)) \leq \lambda$ et $\lambda - \varepsilon \leq f(E(x)+1) \leq \lambda$.

Et, puisque $f(E(x)) \leq f(x) \leq f(E(x)+1)$, on a aussi $\lambda - \varepsilon \leq f(x) \leq \lambda$. Nous avons donc montré que, pour $x \geq N$, $|f(x) - \lambda| \leq \varepsilon$, ce qui signifie que λ est limite de f en $+\infty$.

Montrons que si f a une limite λ , elle vérifie la condition de Cauchy: Nous savons que:

$$(\forall \varepsilon) (\exists A) (\forall x) (x \geq A) \text{ implique } |f(x) - \lambda| \leq \varepsilon/2$$

Donc si $x \geq A$ et $y \geq A$, on a $|f(x) - \lambda| \leq \varepsilon/2$ et $|f(y) - \lambda| \leq \varepsilon/2$, ce qui entraîne $|f(x) - f(y)| \leq \varepsilon$ (puisque $|f(x) - f(y)| \leq |f(x) - \lambda| + |f(y) - \lambda|$).

Et ainsi $(\forall x) (\forall y) (x \geq A \text{ et } y \geq A) \text{ implique } |f(x) - f(y)| \leq \varepsilon$; c'est la condition de Cauchy.

Montrons que si f vérifie la condition de Cauchy, elle a une limite: Considérons la suite $(f(n))_{n \in \mathbb{N}}$; puisque f vérifie la condition de Cauchy:

$$(\forall \varepsilon > 0) (\exists A) \text{ tel que } (\forall n) (\forall p) \ n \geq A \text{ et } p \geq A \text{ implique } |f(n) - f(p)| \leq \varepsilon$$

Autrement dit la suite $(f(n))_{n \in \mathbb{N}}$ est une suite de Cauchy, et elle a une limite λ .

Alors pour tout $x \geq A$, nous aurons (quelque soit l'entier $p \geq A$) $|f(x) - f(p)| \leq \varepsilon$. D'où

$$|f(x) - \lim_{p \rightarrow \infty} f(p)| = |f(x) - \lambda| \leq \varepsilon$$

Et en ne lisant que ce qui est souligné on s'aperçoit que f a pour limite λ en $+\infty$.

§ 2 : Les autres formes du critère de Cauchy

Limites en $-\infty$: Soit maintenant f une fonction sur $] -\infty, a]$:

Si f est croissante et minorée, alors elle a une limite en $-\infty$; si f est décroissante et majorée elle a une limite en $-\infty$.

Une condition nécessaire et suffisante pour que f ait une limite en $-\infty$, est que:

$$(\forall \varepsilon > 0) (\exists A \leq a) (\forall x) (\forall y) (x \leq A \text{ et } y \leq A) \text{ implique } |f(x) - f(y)| \leq \varepsilon.$$

On déduit ces énoncés des précédents par considération de la fonction $x \rightarrow f(-x)$.

Limites à gauche de a : Soit f une fonction sur $[b, a[$ (ou sur $]b, a[$, ou sur $] -\infty, a[$):

Si f est croissante et majorée, elle a une limite en a . Si f est décroissante et minorée, elle a une limite en a .

Une condition nécessaire et suffisante pour que f ait une limite en a , est que:

$$(\forall \varepsilon > 0) (\exists \eta > 0) (\forall x) (\forall y) (a - x \leq \eta \text{ et } a - y \leq \eta) \text{ implique } |f(x) - f(y)| \leq \varepsilon.$$

Exercice a: Ecrire la démonstration de ces énoncés en s'inspirant des démonstrations du §1.

Limites à droite de a : Soit f une fonction sur $]a, b]$ (ou sur $]a, b[$, ou sur $]a, +\infty[$):

Si f est décroissante et majorée, elle a une limite en a . Si f est croissante et minorée, elle a une limite en a .

Une condition nécessaire et suffisante pour que f ait une limite en a , est que:

$$(\forall \varepsilon > 0) (\exists \eta > 0) (\forall x) (\forall y) (a < x \leq a + \eta \text{ et } a < y \leq a + \eta) \text{ implique } |f(x) - f(y)| \leq \varepsilon.$$

§ 3 : Intégrales convergentes en $+\infty$.

Soit f une fonction de $[a, +\infty[$ dans $\mathbb{R}$. Supposons que $(\forall x \geq a)$ f soit intégrable sur $[a, x]$ (dans les exemples que nous rencontrerons f sera presque toujours continue, donc intégrable sur tout $[a, x]$),

nous pouvons définir une fonction F sur $[a, +\infty[$ par $F(x) = \int_a^x f(t) dt$.

Si F a une limite en $+\infty$, nous dirons que l'intégrale de f converge en $+\infty$. Cette limite sera alors appelée l'intégrale de f sur $[a, +\infty[$, et notée $\int_a^{+\infty} f(t) dt$.

Notons que s'il existe $b > a$ tel que $\int_b^{+\infty} f(t) dt$ existe, alors quel que soit $c \geq a$, $\int_c^{+\infty} f(t) dt$ existe.

Exercice b: Montrer que l'intégrale de $\frac{1}{1+t^2}$ converge en $+\infty$. Calculer $\int_0^{+\infty} \frac{dt}{1+t^2}$.

Exercice c: Soit $f: [0, +\infty[\rightarrow \mathbb{R}$ continue, on suppose que $\int_1^{+\infty} f(t) dt$ existe. Montrer que, quelque soit $a \geq 0$, $\int_a^{+\infty} f(t) dt$ existe.

Cas des fonctions positives: Si f est une fonction positive ou nulle, F est croissante donc pour que l'intégrale de f converge en $+\infty$, il suffit qu'il existe M tel que $(\forall x \geq a) F(x) = \int_a^x f(t) dt \leq M$. En parti-

culier si $(\forall t) 0 \leq f(t) \leq g(t)$, et si l'intégrale de g est convergente, alors l'intégrale de f est convergente (C'est alors $\int_a^{+\infty} g(t) dt$ qui joue le rôle de M); de plus $0 \leq \int_a^{+\infty} f(t) dt \leq \int_a^{+\infty} g(t) dt$.

Exercice d: Montrer que l'intégrale de $f(t) = \frac{1}{1+t^2+\text{Log}t}$ converge en $+\infty$, en comparant à $g(t) = \frac{1}{1+t^2}$. Déterminer un nombre A tel que $\int_1^{+\infty} f(t) dt - 10^{-3} \leq \int_1^A f(t) dt \leq \int_1^{+\infty} f(t) dt$.

Exemples à connaître: Les intégrales des fonctions $t \rightarrow a^t$ ($a < 1$) et $t \rightarrow \frac{1}{t^\alpha}$ ($\alpha > 1$) convergent en $+\infty$ (Pour s'en convaincre calculer des primitives et montrer qu'elles ont une limite en $+\infty$). Les intégrales des fonctions $t \rightarrow a^t$ ($a \geq 1$) et $t \rightarrow \frac{1}{t^\alpha}$ ($\alpha \leq 1$) ne convergent pas en $+\infty$.

Liens entre convergence des intégrales et sommabilité des séries

Soit $f: [0, +\infty[\rightarrow \mathbb{R}$, posons $F(x) = \int_0^x f(t) dt$ et $u_n = \int_n^{n+1} f(t) dt = F(n+1) - F(n)$. Si l'intégrale de f converge en $+\infty$ (c'est à dire si F a une limite λ en $+\infty$), alors la suite $(F(n))_{n \in \mathbb{N}}$ a une limite. Autrement dit la série $((u_n))_{n \in \mathbb{N}}$ est sommable (de somme λ). Nous avons déjà utilisé ce type de raisonnement au chapitre 8 §7.

Il est possible que la série $((u_n))_{n \in \mathbb{N}}$ soit sommable, et que l'intégrale de f ne soit pas convergente en $+\infty$ (ainsi pour $f(t) = \cos 2\pi t$, $F(x) = \frac{1}{2\pi} \sin 2\pi x$ n'a pas de limite, tandis que $F(n) = 0$ pour tout entier n). Mais si f a pour limite 0 en $+\infty$, et si la série $((u_n))_{n \in \mathbb{N}}$ est sommable, alors l'intégrale de f est convergente en $+\infty$.

En effet si f a pour limite 0 en $+\infty$:

$$(\forall \varepsilon > 0) (\exists A) (\forall t) (t \geq A) \text{ implique } |f(t)| \leq \varepsilon/2$$

Il en résulte que $x \geq A+1$ implique $\left| \int_{E(x)}^x f(t) dt \right| \leq \varepsilon/2$ (où $E(x)$ désigne la partie entière de x).

Et puisque la série $((u_n))_{n \in \mathbb{N}}$ est sommable (de somme S):

$$(\forall \epsilon > 0) (\exists B) (\forall n) (n \geq B) \text{ implique } \left| \sum_{i=0}^n u_i - S \right| \leq \epsilon/2$$

Donc si $x \geq \sup(A+1, B+1)$, on a:

$$\left| \int_0^x f(t) dt - S \right| = \left| \sum_{i=0}^{E(x)-1} u_i - S + \int_{E(x)}^x f(t) dt \right| \leq \left| \sum_{i=0}^{E(x)-1} u_i - S \right| + \left| \int_{E(x)}^x f(t) dt \right| \leq \epsilon$$

Ce qui prouve que S est la limite de F en $+\infty$.

Exercice e: On suppose maintenant que f est une fonction à valeurs positives, montrer que la sommabilité de la série $((u_n))_{n \in \mathbb{N}}$ implique la convergence de l'intégrale de f (même si f n'a pas pour limite 0 en $+\infty$).

On remarquera que le fait que l'intégrale de f converge en $+\infty$ n'implique pas que f a pour limite 0 en $+\infty$ (pour un exemple voir l'exercice 11). Inversement il est possible que f tende vers 0 en $+\infty$, et que l'intégrale de f ne converge pas en $+\infty$ (penser à $x \rightarrow \frac{1}{x}$).

§ 4 : Intégrales absolument convergentes en $+\infty$.

Soit f une fonction continue de $[a, +\infty[$ dans $\mathbb{R}$, nous dirons que l'intégrale de f est absolument convergente en $+\infty$, si l'intégrale de $|f|$ converge en $+\infty$. Ceci implique que l'intégrale de f est convergente en $+\infty$.

En effet $(\forall x < y) \quad \left| \int_x^y f(t) dt \right| \leq \int_x^y |f(t)| dt.$

Donc $(\forall x < y) \quad \left| \int_a^y f(t) dt - \int_a^x f(t) dt \right| \leq \int_a^y |f(t)| dt - \int_a^x |f(t)| dt$

Puisque l'intégrale de $|f|$ est convergente, la fonction $x \rightarrow \int_a^x |f(t)| dt$ vérifie la condition de Cauchy;

l'inégalité ci-dessus montre que la fonction $x \rightarrow \int_a^x f(t) dt$ vérifie aussi la condition de Cauchy, elle a donc une limite en $+\infty$.

Pour montrer que l'intégrale d'une fonction f est absolument convergente, nous comparerons donc $|f|$ aux fonctions classiques dont nous connaissons l'intégrale (cf: §3).

Comparaison aux fonctions puissances: S'il existe $\alpha > 1$ tel que, au voisinage de $+\infty$, $|f|$ soit $O(1/t^\alpha)$, alors l'intégrale de f est absolument convergente en $+\infty$. A fortiori s'il existe $\alpha > 1$ tel que, au voisinage de $+\infty$, $|f|$ soit $o(1/t^\alpha)$, alors l'intégrale de f est absolument convergente en $+\infty$. De même s'il existe $\alpha > 1$ et $B > 0$ tels que, au voisinage de $+\infty$, f soit équivalente à B/t^α , alors l'intégrale de f est absolument convergente en $+\infty$. Car dans les trois cas il existe A et K tels que $t > A$ implique $|f(t)| \leq K/t^\alpha$, et l'intégrale de $1/t^\alpha$ converge puisque $\alpha > 1$.

S'il existe $\alpha \leq 1$ et $K > 0$ tels que, au voisinage de $+\infty$, $f(t)$ soit minorée par K/t^α , alors l'intégrale de f ne converge pas en $+\infty$. S'il existe $\alpha \leq 1$ et $B > 0$ tels que, au voisinage de $+\infty$, $f(t)$ soit équivalente à B/t^α , alors l'intégrale de f ne converge pas en $+\infty$. En effet il existe alors A et $B(>0)$ tels que, $t > A$ implique $f(t) \geq B/2t^\alpha$, et l'intégrale de $B/2t^\alpha$ diverge.

Comparaison aux exponentielles: S'il existe $\alpha < 1$ tel que, au voisinage de $+\infty$, $|f|$ soit $O(\alpha^t)$, alors l'intégrale de f est absolument convergente en $+\infty$. A fortiori s'il existe $\alpha < 1$ tel que, au voisinage de $+\infty$, $|f|$ soit $\mathcal{O}(\alpha^t)$, alors l'intégrale de f est absolument convergente en $+\infty$. De même s'il existe $\alpha < 1$ et B tels que, au voisinage de $+\infty$, f soit équivalente à $B\alpha^t$, alors l'intégrale de f est absolument convergente en $+\infty$. Car dans les trois cas il existe A et K tels que $t > A$ implique $|f(t)| \leq K\alpha^t$.

Exercice f: Reprendre le contenu de ces deux paragraphes en remplaçant $+\infty$ par $-\infty$. Qu'est ce qui reste vrai ? Qu'est ce qui change ?

§ 5 : Intégrales (absolument) convergentes en a

Cette fois nous considérons des fonctions sur un intervalle de la forme $]a, b]$ (ou $]a, +\infty[$) qui sont intégrables sur tout intervalle de la forme $[x, b]$ ($x > a$). Nous dirons que l'intégrale de f converge en a , si la fonction F définie par $F(x) = \int_x^b f(t) dt$, a une limite en a . Cette limite sera appelée l'intégrale de f sur $]a, b]$, et notée $\int_a^b f(t) dt$.

Si l'intégrale de $|f|$ est convergente en a , on dit que l'intégrale de f est absolument convergente en a . Ceci implique que l'intégrale de f est convergente en a .

Si $|f|$ est majorée par une fonction positive g dont l'intégrale converge en a , alors l'intégrale de f est absolument convergente.

Exercice g : Montrer que $(\forall c) \int_a^c f(t) dt$ existe, si et seulement si $\int_a^b f(t) dt$ existe. Montrer que si f est bornée au voisinage de a , alors l'intégrale de f est convergente en a .

Exercice h: En s'inspirant de ce qui a été fait (pour la convergence en $+\infty$) au §4, démontrer que si l'intégrale de $|f|$ converge en a , alors l'intégrale de f converge en a .

Exercice i: En s'inspirant de ce qui a été fait (pour la convergence en $+\infty$) au § 3 démontrer que si $|f| \leq g$, et si l'intégrale de g converge en a , alors l'intégrale de $|f|$ converge en a .

Ainsi pour démontrer que l'intégrale de f converge il nous suffit de comparer $|f|$ à des fonctions simples dont nous connaissons le comportement. En général il s'agit des fonctions puissances. L'intégrale de $t \rightarrow \frac{1}{(t-a)^\alpha}$ converge en a si et seulement si $\alpha < 1$ (pour s'en convaincre, calculer une primitive). Donc:

S'il existe $\alpha < 1$ tel que, au voisinage de a , $|f|$ soit $O(1/(t-a)^\alpha)$, alors l'intégrale de f est absolument convergente en a . A fortiori s'il existe $\alpha < 1$ tel que, au voisinage de a , $|f|$ soit $\mathcal{O}(1/(t-a)^\alpha)$, alors l'intégrale de f est absolument convergente en a . De même s'il existe $\alpha < 1$ et B tels que, au voisinage de a , f soit équivalente à $B/(t-a)^\alpha$, alors l'intégrale de f est absolument convergente en a . Car il existe alors ε et K tels que $a+\varepsilon \leq t < a$ implique $|f(t)| \leq K/(t-a)^\alpha$, et l'intégrale de $1/(t-a)^\alpha$ converge puisque $\alpha < 1$.

S'il existe $\alpha \geq 1$ et $K > 0$ tels que, au voisinage de a , $f(t)$ soit minorée par $K/(t-a)^\alpha$, alors l'intégrale de f ne converge pas en a . S'il existe $\alpha \geq 1$ et B tels que, au voisinage de a , $f(t)$ soit équivalente à $B/(t-a)^\alpha$, alors l'intégrale de f ne converge pas en a . En effet il existe alors ε et $K > 0$ tels que, $a+\varepsilon \leq t < a$ implique $f(t) \geq K/2(t-a)^\alpha$.

Tout ceci peut évidemment être repris, sans modification notable, pour les fonctions définies à gauche de a .

§ 6 : Théories de l'intégrale sur les intervalles non compacts^(*)

Appelons fonctions faiblement intégrables sur l'intervalle $[a, +\infty[$ les fonctions continues sauf en un nombre fini de points, dont l'intégrale converge en $+\infty$. Elles forment un espace vectoriel de fonctions $\mathcal{F}_f([a, +\infty[)$.

* L'application qui à toute f associe $\int_a^{+\infty} f(t) dt$, est une forme linéaire sur $\mathcal{F}_f([a, +\infty[)$.

* Elle est croissante, c'est à dire $f \leq g$ implique $\int_a^{+\infty} f(t) dt \leq \int_a^{+\infty} g(t) dt$.

Et bien sûr les fonctions caractéristiques des intervalles bornés sont dans $\mathcal{F}_f([a, b])$, et ont pour intégrale la longueur desdits intervalles.

Nous retrouvons ainsi toutes les propriétés d'une théorie de l'intégrale (au sens du chap. 7).

Mais nous pouvons aussi considérer les fonctions dont l'intégrale converge absolument (nous les dirons intégrables). Nous obtenons une nouvelle théorie de l'intégrale définie sur un espace de fonctions $\mathcal{F}([a, +\infty[)$ (qui est strictement plus petit que $\mathcal{F}_f([a, +\infty[)$).

Nous pourrions de même définir sur tout intervalle I non compact deux théories de l'intégrale; l'une formée des fonctions dont l'intégrale converge à l'extrémité non fermée ou non bornée de I (ou aux deux extrémités de I); l'autre formée des fonctions dont l'intégrale converge absolument.

§ 7 : Convergence à l'infini de l'intégrale des fonctions oscillantes

Nous nous intéressons ici aux fonctions données par une formule du type $f(t) = g(t) \sin t$ (ou bien $f(t) = g(t) \cos t$). Si g garde un signe constant sur $[A, +\infty[$, alors f a sur cet intervalle le signe de $\sin t$ (ou de $\cos t$); d'où le nom de fonctions oscillantes.

Si g est monotone au voisinage de $+\infty$, et tend vers 0 en $+\infty$, alors l'intégrale de f converge en $+\infty$. Si g est monotone au voisinage de $-\infty$, et tend vers 0 en $-\infty$, alors l'intégrale de f converge en $-\infty$. Cette règle est connue sous le nom de critère d'Abel.

Pour nous en convaincre considérons d'abord la série définie par $u_n = \int_{n\pi}^{(n+1)\pi} g(t) \sin t dt$. Puisque

g est monotone au voisinage de $+\infty$, et converge vers 0 en $+\infty$, elle garde un signe constant au voisinage de $+\infty$; il en résulte que le signe de u_n est celui de $(-1)^n$ si $g > 0$ (et $(-1)^{n+1}$ si $g < 0$). Nous allons donc montrer la sommabilité de $((u_n))_{n \in \mathbb{N}}$ grâce à la règle des séries alternées (chap. 8 §6). Pour cela nous devons démontrer que la suite $(|u_n|)_{n \in \mathbb{N}}$ est décroissante. Ce qui résulte du calcul suivant:

$$\begin{aligned} |u_n| - |u_{n+1}| &= \left| \int_{n\pi}^{(n+1)\pi} g(t) \sin t dt \right| - \left| \int_{(n+1)\pi}^{(n+2)\pi} g(t) \sin t dt \right| \\ |u_n| - |u_{n+1}| &= \int_{n\pi}^{(n+1)\pi} |g(t)| |\sin t| dt - \int_{(n+1)\pi}^{(n+2)\pi} |g(t)| |\sin t| dt \quad (\text{car } g \text{ a un signe constant}) \end{aligned}$$

^(*) Nous dirons qu'un intervalle est compact s'il est fermé et borné. Un intervalle non compact est donc soit non borné (par exemple $[1, +\infty[$), soit non fermé (par exemple $]0, 1[$), soit les deux (cf: chap 10 §7).

$$|u_n| - |u_{n+1}| = \int_0^\pi [|g(t+n\pi)| - |g(t+(n+1)\pi)|] \sin t \, dt$$

Et ceci est positif puisque $[|g(t+n\pi)| - |g(t+(n+1)\pi)|] \geq 0$ (du fait que $|g|$ est décroissante).

La convergence de l'intégrale de f résulte alors du § 3 (lien entre convergence des intégrales et sommabilité des séries).

Exercice j: Démontrer de même que si f est décroissante et a pour limite 0 en $+\infty$, alors les intégrales de $t \rightarrow f(t) \cos \alpha t$ et de $t \rightarrow f(t) \sin \alpha t$ convergent en $+\infty$.

Exercices sur le chapitre 9

** Exercice 1: Les intégrales des fonctions suivantes sont elles convergentes en $+\infty$? Sont elles absolument convergentes ?
 $f(t) = t e^{-t}$ $g(t) = \frac{\cos t}{1+t^2}$ $h(t) = \frac{\sin t}{\operatorname{Log} t}$

* Exercice 2: Les intégrales des fonctions suivantes sont elles convergentes en 0? Sont elles absolument convergentes ?
 $a(t) = \operatorname{Log} t$ $b(t) = t \operatorname{Log} t$ $c(t) = \frac{\operatorname{Log} t}{t}$
 $d(t) = \frac{\operatorname{Log} t}{\sin t}$ $e(t) = e^{-1/t}$ $f(t) = \sin(1/t) \operatorname{Log} t$

** Exercice 3: On considère les deux fonctions $f(t) = \frac{\sin t}{t} + \frac{\sin^2 t}{t \operatorname{Log} t}$ et $g(t) = \frac{\sin t}{t}$. Montrer qu'elles sont équivalentes au voisinage de $+\infty$; mais que la seconde a une intégrale convergente en $+\infty$, et la première une intégrale divergente en $+\infty$. On pourra remarquer que $\frac{2 \sin^2 t}{t \operatorname{Log} t} = \frac{1}{t \operatorname{Log} t} - \frac{\cos 2t}{t \operatorname{Log} t}$.

* Exercice 4: Les intégrales suivantes sont elles convergentes? absolument convergentes ?

$$I = \int_0^{+\infty} \frac{\operatorname{Log} t}{1+t^2} dt$$

$$J = \int_0^{+\infty} \frac{\sin t}{t\sqrt{t}} dt$$

$$K = \int_0^1 \frac{\sqrt{\sin \pi t}}{t^2-t} dt$$

* Exercice 5: Les intégrales suivantes sont elles convergentes? absolument convergentes ?

$$A = \int_{-\infty}^0 \frac{e^t \operatorname{Log} |t|}{t} dt$$

$$B = \int_{-\infty}^{+\infty} \frac{dt}{\operatorname{ch} t}$$

$$C = \int_{-\infty}^0 \frac{dt}{\sqrt[3]{\operatorname{sh} t}}$$

** Exercice 6: Les intégrales suivantes sont elles convergentes? absolument convergentes ?

a) $I = \int_0^{+\infty} \sin t^2 dt$ (on pourra faire le changement de variable $t^2 = u$)

b) $J = \int_0^{\infty} \sin(t^2+t+1) dt$

c) $K = \int_0^{\infty} \sin(at^2+bt+c) dt$

*** Exercice 7:

Soit P et Q des polynômes sans racine commune. Montrer que l'on peut donner un sens au symbole $\int_{-\infty}^{+\infty} \sqrt{\frac{P(t)}{Q(t)}} dt$, si

et seulement si d'une part $3+\deg P \leq \deg Q$, d'autre part Q n'a que des racines simples.

* Exercice 8: Les intégrales suivantes sont elles convergentes? absolument convergentes ?

$$M = \int_0^1 \frac{\operatorname{Log} x}{(1-x)^\alpha} dx \quad (\text{discuter suivant les valeurs de } \alpha).$$

$$N = \int_0^1 \frac{dx}{x(x-1)(x-2)}$$

$$P = \int_0^{\infty} \frac{\cos x}{1+x} dx$$

* Exercice 9: Les intégrales des fonctions suivantes sont elles convergentes en $+\infty$?

$$j(t) = \frac{1}{t \operatorname{Log} t}$$

$$k(t) = \frac{1}{t (\operatorname{Log} t)^\alpha} \quad (\alpha \in \mathbb{R})$$

$$m(t) = \frac{\operatorname{Log} t}{t^2}$$

Exercice 10: Les intégrales suivantes sont-elles convergentes ? absolument convergentes ?

$$A = \int_0^1 \operatorname{Log} \frac{1}{1-t} dt$$

$$B = \int_0^1 \frac{\operatorname{Log} (\sin \pi t)}{\sqrt{t(1-t)}} dt$$

$$C = \int_0^1 \left(\sin \frac{1}{t} \right) dt$$

***** Exercice 11:** Soit $(\alpha_n)_{n \in \mathbb{N}}$ une suite d'entiers positifs, on définit $f: [0, +\infty[\rightarrow \mathbb{R}$ par:

$$\begin{cases} f(t) = (t-2n)^{\alpha_n} & \text{si } 2n \leq t \leq 2n+1 \\ f(t) = (2n+2-t)^{\alpha_n} & \text{si } 2n+1 \leq t \leq 2n+2 \end{cases}$$

a) Calculer $\int_{2n}^{2n+2} f(t) dt$,

b) Montrer que l'intégrale de f est convergente en $+\infty$, si et seulement si la série $\left(\frac{1}{1+\alpha_n}\right)_{n \in \mathbb{N}}$ est sommable. Ceci

nous donne de nombreux exemples de fonctions qui ne tendent pas vers 0 en $+\infty$, et dont l'intégrale converge en $+\infty$.

c) Construire une fonction non bornée au voisinage de l'infini, et dont l'intégrale converge à l'infini.

**** Exercice 12:** a) Montrer que l'intégrale $\int_0^{+\infty} \frac{\sin t}{t} dt$ existe. Est-elle absolument convergente en $+\infty$?

b) Montrer que l'intégrale $\int_0^{+\infty} \frac{1-\cos t}{t^2} dt$ existe. Est-elle absolument convergente en $+\infty$?

c) Montrer, par une intégration par parties, que $\int_0^{+\infty} \frac{\sin t}{t} dt = \int_0^{+\infty} \frac{1-\cos t}{t^2} dt$.

d) Par une nouvelle intégration par parties, trouver une fonction $f(t)$ telle que $\int_0^{+\infty} \frac{\sin t}{t} dt = \int_0^{+\infty} \frac{f(t)}{t^3} dt$

**** Exercice 13:** Soit $f(t) = \frac{\sin t}{(\operatorname{Log} t)^2}$.

a) Montrer que l'intégrale de f converge en $+\infty$.

b) Minorer $\int_{n\pi}^{(n+1)\pi} f^2(t) dt$ par une quantité de la forme $\frac{K}{(\operatorname{Log} (n+1))^4}$. En déduire que l'intégrale de f^2 n'est pas convergente en $+\infty$.

**** Exercice 14:** a) Calculer $I_k = \int_1^\infty \frac{dt}{t^{2k}}$ et $J = \int_1^\infty \frac{dt}{1+t^2}$

b) Montrer que $(\forall k) \frac{(-1)^{k+1}}{t^{2k}(1+t^2)} = \frac{1}{t^2} - \frac{1}{t^4} + \frac{1}{t^6} - \dots + \frac{(-1)^{k+1}}{t^{2k}} - \frac{1}{1+t^2}$.

c) Montrer que $|I_1 - I_2 + I_3 - I_4 + \dots + (-1)^{k+1} I_k - J| \leq I_{k+1}$. Quelle est la somme de la série $\left(\frac{(-1)^n}{2n+1}\right)_{n \in \mathbb{N}}$?

*** Exercice 15: On considère $\int_0^{+\infty} (\sin^2 t)^{kt} dt$ (k entier au moins égal à 1).

a) Démontrer que cette intégrale existe si et seulement si la série $((u_n)_{n \in \mathbb{N}}$ définie par $u_n = \int_0^\pi \sin^{2kn} t dt$, est sommable.

b) On pose $I_p = \int_0^{\pi/2} \sin^{2p} t dt$. Au moyen d'une intégration par parties, calculer I_p/I_{p+1} .

c) Montrer que $(\forall p > 0) \frac{2p-1}{2p} \geq \frac{p-1}{p}$; en déduire que la suite $(I_p/p)_{p \geq 1}$ est croissante, et que l'intégrale donnée n'existe pas pour $k = 1$.

d) Montrer que $\frac{I_{2p-2}}{I_{2p}} \leq \frac{(p-1)^2}{p^2}$ pour p grand. En déduire que pour $k = 2$, la suite $(\frac{u_n}{n^2})_{n \geq 1}$ décroît, et que l'intégrale existe.

e) Montrer que l'intégrale existe si $k > 2$.

* Exercice 16: Etudier l'existence des intégrales suivantes

$$I = \int_{-\infty}^{+\infty} e^{-t^2} \cos^2 t dt$$

$$J = \int_{-\infty}^0 e^t \operatorname{Log}(1+t^2) dt$$

* Exercice 17: Les intégrales suivantes sont-elles convergentes ? absolument convergentes ?

$$I = \int_{-\infty}^0 \frac{1}{\operatorname{ch} t} (1+2\cos t) dt$$

$$J = \int_{-\infty}^0 \frac{\sin t}{\operatorname{sh} t} dt$$

* Exercice 18: Les intégrales suivantes sont-elles convergentes ? sont-elles absolument convergentes ?

$$I = \int_0^{+\infty} \frac{e^t}{\operatorname{ch} t} dt$$

$$J = \int_0^{+\infty} \operatorname{Log} t \sin t dt$$

Espaces normés

(Suites convergentes et applications continues)

Toutes les notions d'analyse que nous connaissons, sont fondées sur la possibilité de donner un sens (lorsque x et y sont des nombres) à l'expression ' x est proche de y ', ou plus précisément 'la distance de x à y est plus petite que...'. Pour étendre ces notions aux espaces $\mathbb{R}^n$ et aux espaces de fonctions, et en particulier pour définir la convergence des suites de vecteurs et des suites de fonctions, nous utiliserons la notion de norme. Au chapitre 4 on a défini une norme sur tout espace muni d'un produit scalaire. Le concept général de norme est un peu plus général.

La norme d'un vecteur du plan est sa longueur, c'est à dire que la norme de $\vec{OA}$ est la longueur OA , ou encore la distance de A à l'origine. De façon analogue dans un espace vectoriel, la norme d'un vecteur doit être considérée comme 'sa distance au vecteur nul'. On gardera cette interprétation géométrique à l'esprit, en particulier lors de l'étude des espaces de fonctions.

§ 1 : Normes sur un espace vectoriel

On appelle norme sur un espace vectoriel E (réel ou complexe), toute application N de E dans $]0, \infty[$ telle que:

$$1) N(x) = 0 \Leftrightarrow x = 0$$

$$2) (\forall x \in E), (\forall \lambda \in \mathbb{R} \text{ ou } \mathbb{C}) : N(\lambda x) = |\lambda| N(x)$$

$$3) (\forall x \text{ et } y \text{ dans } E) : N(x+y) \leq N(x) + N(y)$$

$$3bis) (\forall x \text{ et } y \text{ dans } E) : N(x-y) \leq N(x) + N(y)$$

A titre d'exemple, on pense bien sûr à la norme (ou longueur) des vecteurs de la géométrie. D'autres exemples nous sont fournis par les normes associées aux produits scalaires (ou hermitiens) (cf: chap 4 § 2). Mais il y en a bien d'autres; voici quelques unes de celles que nous aurons l'occasion de rencontrer dans la suite de ce cours.

Normes sur les espaces de dimension finie: Soit E un espace vectoriel muni d'une base $\{e_1, \dots, e_n\}$. Pour tout $x = \sum x_i e_i$ nous pouvons poser:

$$N_1(x) = \sum_i |x_i|$$

$$N_2(x) = \sqrt{\sum_i |x_i|^2}$$

$$N_\infty(x) = \sup (|x_1|, \dots, |x_n|)$$

On reconnaît en N_2 la norme euclidienne associée au (seul) produit scalaire (ou hermitien) pour lequel $\{e_1, \dots, e_n\}$ est une base orthonormée.

Exercice a : Démontrer que N_1 et N_∞ sont des normes sur E .

Normes équivalentes.

Il existe sur tout espace vectoriel, une infinité de normes distinctes.

On dit que deux normes N et N' sur le même espace E sont équivalentes si il existe des constantes K et L strictement positives, telles que

$$(\forall x \in E) : K N(x) \leq N'(x) \leq L N(x)$$

Exercice b : Montrer que, dans les exemples ci-dessus, on a :

$$(\forall x) \quad N_{\infty}(x) \leq N_1(x) \leq n N_{\infty}(x) \quad \text{et} \quad N_{\infty}(x) \leq N_2(x) \leq \sqrt{n} N_{\infty}(x)$$

Trouver K et L tels que $(\forall x) \quad K N_1(x) \leq N_2(x) \leq L N_1(x)$

Ainsi les trois normes que nous venons de définir sur E sont équivalentes.

On démontre que deux normes quelconques sur un espace de dimension finie, sont équivalentes.

Nous admettrons ce résultat dans toute la suite.

Normes sur les espaces de fonctions.

Considérons l'espace $\mathcal{C}([a,b], \mathbb{R})$ (ou $\mathcal{C}([a,b], \mathbb{C})$) des fonctions continues sur $[a,b]$; nous pouvons poser:

$$N_{\infty}(f) = \sup_{t \in [a,b]} |f(t)|$$

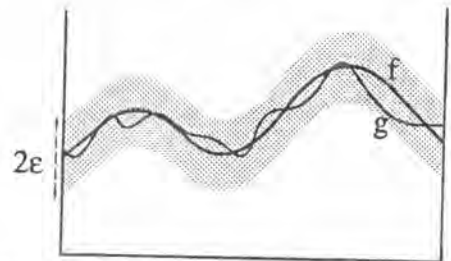
Notons que $N_{\infty}(f)$ existe puisque toute fonction continue sur un intervalle fermé borné est majorée, et atteint ses bornes. Ceci définit une norme qui est appelée la norme uniforme.

Exercice b : Soit f et g deux fonctions continues sur $[a,b]$. Démontrer que

$$(\forall t \in [a,b]) : |(f+g)(t)| \leq N_{\infty}(f) + N_{\infty}(g)$$

Puis montrer que N_{∞} est une norme.

Remarque: L'inégalité $N_{\infty}(f) \leq k$ équivaut à $(\forall x \in [a,b]) \quad |f(x)| \leq k$.



Dire que la norme uniforme de $f-g$ est au plus égale à ϵ , c'est dire que le graphique de g est situé dans une bande de largeur 2ϵ entourant celui de f

Sur les mêmes espaces nous pouvons poser
$$N_2(f) = \left[\int_a^b |f(t)|^2 dt \right]^{1/2}.$$

C'est la norme associée au produit scalaire $\langle f, g \rangle = \int_a^b f(t) \overline{g(t)} dt$ (dans le cas complexe au produit

hermitien : $(f|g) = \int_a^b f(t) \overline{g(t)} dt$). Cette norme est appelée la norme de la convergence en moyenne quadratique

Nous pouvons aussi poser:
$$N_1(f) = \int_a^b |f(t)| dt$$

Exercice c : Montrer que N_1 est une norme.

Cette norme est appelée la norme de la convergence en moyenne

Contrairement aux espaces de dimension finie, sur ces espaces de fonctions nous pouvons définir des normes qui ne sont pas équivalentes (cf: Ex. d et e).

Exercice d : Démontrer les inégalités

$$\forall f \in \mathcal{C}([a,b], \mathbb{R}) \quad N_1(f) \leq (b-a) N_\infty(f)$$

$$\forall f \in \mathcal{C}([a,b], \mathbb{R}) \quad N_2(f) \leq \sqrt{b-a} N_\infty(f)$$

$$\forall f \in \mathcal{C}([a,b], \mathbb{R}) \quad N_1(f) \leq \sqrt{b-a} N_2(f)$$

Pour démontrer la 3^{ème}, appliquer l'inégalité de Schwarz au

$$\text{produit scalaire } \langle f, 1 \rangle = \int_a^b |f(t)| \cdot 1 \, dt$$

Exercice e : Notons f_n la fonction de $\mathcal{C}([0,1], \mathbb{R})$

$$f_n(t) = nt \text{ si } 0 \leq t \leq \frac{1}{n}$$

$$f_n(t) = 2 - nt \text{ si } \frac{1}{n} \leq t \leq \frac{2}{n}$$

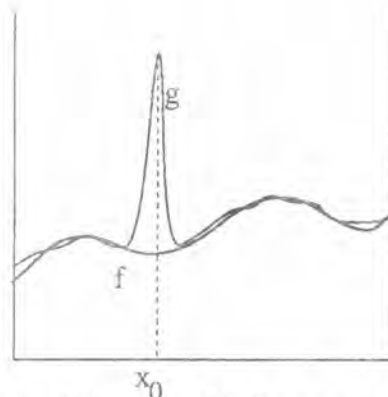
$$\text{et } f_n(t) = 0 \text{ si } t > \frac{2}{n}$$

Calculer $N_1(f_n)$, $N_\infty(f_n)$ et $N_2(f_n)$.

En déduire qu'il n'existe aucun nombre K tel que

$$\forall f \in \mathcal{C}([0,1], \mathbb{R}) \quad N_\infty(f) \leq K N_1(f)$$

(resp: $N_\infty(f) \leq K N_2(f)$; resp: $N_2(f) \leq K N_1(f)$).



Au sens de la norme N_2 (ou au sens de la norme N_1), aussi près qu'on veut de f on trouve des fonctions qui, au point x_0 , diffèrent beaucoup de f . Mais la mesure de l'ensemble des points où elles diffèrent beaucoup de f est petite.

§ 2 : Suites convergentes dans un espace normé

Soit E un espace muni d'une norme N soit $(\varphi_n)_{n \in \mathbb{N}}$ une suite d'éléments de E ; on dit que $(\varphi_n)_{n \in \mathbb{N}}$ converge vers λ ($\in E$) pour la norme N si la suite numérique $(N(\varphi_n - \lambda))_{n \in \mathbb{N}}$ converge vers 0. Autrement dit $(\varphi_n)_{n \in \mathbb{N}}$ converge vers λ pour la norme N si:

$$(\forall \varepsilon > 0) (\exists n_0) \text{ tel que } (\forall n) (n > n_0) \text{ implique } N(\varphi_n - \lambda) \leq \varepsilon$$

Nous avons tout simplement recopié la définition des suites numériques convergentes (chap 6 § 1) en remplaçant $|\varphi_n - \lambda|$ par $N(\varphi_n - \lambda)$. Cette analogie est l'assurance que nous allons pouvoir démontrer que les suites vectorielles convergentes ont des propriétés fort analogues à celles que nous connaissons.

Exercice f : Reprenons les fonctions f_n de l'exercice e.

1) La suite $(f_n)_{n \in \mathbb{N}}$ converge-t-elle vers 0 pour N_1 (resp : pour N_2 ; pour N_∞) ?

2) La suite $(\sqrt{n} f_n)_{n \in \mathbb{N}}$ converge-t-elle vers 0 pour N_1 (resp : pour N_2 ; pour N_∞) ?

S'il existe un nombre K , tel que $(\forall f \in E) N(f) \leq K N'(f)$, alors toute suite $(\varphi_n)_{n \in \mathbb{N}}$ qui converge vers λ pour N' , converge aussi vers λ pour N .

En effet : Il existe $K(>0)$ tel que $(\forall n) N(\varphi_n - \lambda) \leq K N'(\varphi_n - \lambda)$. Donc la convergence vers 0 de $(N'(\varphi_n - \lambda))_{n \in \mathbb{N}}$ entraîne la convergence vers 0 de $(N(\varphi_n - \lambda))_{n \in \mathbb{N}}$.

Donc si les normes N et N' sont équivalentes, les suites qui convergent pour N , sont les mêmes que celles qui convergent pour N' ; et elles ont même limite pour N et N' .

En particulier dans un espace de dimension finie, lorsque l'on parle de suite convergente, il est inutile de préciser pour quelle norme. Au contraire sur un espace de dimension infinie, lorsqu'on parle de suite convergente il faut préciser quelle norme on utilise (cf: Exercice f). Nous retiendrons que toute suite de $\mathcal{C}([a,b], \mathbb{R})$ qui converge uniformément (i.e. pour N_∞) converge pour N_2 ; et que toute suite qui converge pour N_2 converge pour N_1 . Au contraire il existe des suites qui convergent pour N_1 ou pour N_2 et qui ne convergent pas uniformément (exercice f).

Propriétés algébriques. L'espace E étant muni d'une norme N .

1) Si les suites $(u_n)_{n \in \mathbb{N}}$ et $(v_n)_{n \in \mathbb{N}}$ d'éléments de E convergent vers u et v pour N , alors (quelque soit le scalaire λ) la suite $(u_n + \lambda v_n)_{n \in \mathbb{N}}$ converge vers $u + \lambda v$ pour N .

En effet : $N((u_n + \lambda v_n) - (u + \lambda v)) \leq N(u_n - u) + |\lambda| N(v_n - v)$. Donc si les suites $(N(u_n - u))_{n \in \mathbb{N}}$ et $(N(v_n - v))_{n \in \mathbb{N}}$ convergent vers 0, alors $N((u_n + \lambda v_n) - (u + \lambda v))_{n \in \mathbb{N}}$ converge aussi vers zéro.

2) Si la suite $(u_n)_{n \in \mathbb{N}}$ d'éléments de E converge vers u pour N , et si la suite numérique $(\lambda_n)_{n \in \mathbb{N}}$ converge vers λ , alors la suite $(\lambda_n u_n)_{n \in \mathbb{N}}$ converge vers λu pour N .

En effet : $\lambda_n u_n - \lambda u = \lambda_n (u_n - u) + (\lambda_n - \lambda)u$. Donc :

$$N(\lambda_n u_n - \lambda u) \leq (|\lambda| + |\lambda_n - \lambda|)N(u_n - u) + |\lambda_n - \lambda| N(u)$$

Il en résulte que si les suites $(|\lambda_n - \lambda|)_{n \in \mathbb{N}}$ et $(N(u_n - u))_{n \in \mathbb{N}}$ convergent vers zéro, la suite $(N(\lambda_n u_n - \lambda u))_{n \in \mathbb{N}}$ converge aussi vers zéro.

3) Supposons que E soit muni d'un produit scalaire, et que N soit la norme associée à ce produit scalaire. Si dans E les suites $(u_n)_{n \in \mathbb{N}}$ et $(v_n)_{n \in \mathbb{N}}$ convergent vers u et v pour N , alors $(\langle u_n, v_n \rangle)_{n \in \mathbb{N}}$ converge vers $\langle u, v \rangle$.

En effet : $|\langle u_n, v_n \rangle - \langle u, v \rangle| = |\langle u_n, v_n - v \rangle + \langle u_n - u, v \rangle| \leq |\langle u_n, v_n - v \rangle| + |\langle u_n - u, v \rangle|$
 $\leq N(u_n) N(v_n - v) + N(u_n - u) N(v)$ (d'après Schwarz)

$$|\langle u_n, v_n \rangle - \langle u, v \rangle| \leq [N(u_n - u) + N(u)] N(v_n - v) + N(u_n - u) N(v)$$

Donc si $(N(u_n - u))_{n \in \mathbb{N}}$ et $(N(v_n - v))_{n \in \mathbb{N}}$ convergent vers 0, $(|\langle u_n, v_n \rangle - \langle u, v \rangle|)_{n \in \mathbb{N}}$ converge aussi vers 0.

4) Supposons que E soit muni d'un produit (on notera ef le produit de e et f) et qu'il existe K tel que $(\forall e)$ et $(\forall f) : N(ef) \leq K N(e) N(f)$; alors si $(e_n)_{n \in \mathbb{N}}$ et $(f_n)_{n \in \mathbb{N}}$ sont deux suites qui convergent vers e et f pour la norme N , la suite $(e_n f_n)_{n \in \mathbb{N}}$ converge vers ef pour N .

Exercice g : Donner une démonstration de cette affirmation en s'inspirant de la précédente.

Conséquences :

Dans l'espace de la géométrie, si $(\vec{u}_n)_{n \in \mathbb{N}}$ et $(\vec{v}_n)_{n \in \mathbb{N}}$ convergent vers $\vec{u}$ et vers $\vec{v}$, alors $(\vec{u}_n \wedge \vec{v}_n)_{n \in \mathbb{N}}$ converge vers $\vec{u} \wedge \vec{v}$.

Dans $\mathbb{C}$ si deux suites convergent vers z et Z , alors la suite produit converge vers zZ .

Dans $\mathcal{C}([a, b], \mathbb{R})$ si $(f_n)_{n \in \mathbb{N}}$ converge vers f , et $(g_n)_{n \in \mathbb{N}}$ vers g pour la norme N_∞ , alors $(f_n g_n)_{n \in \mathbb{N}}$ converge vers fg pour la norme N_∞ (cf: Ex 25).

Dans l'espace des matrices $n \times n$, si $(A_n)_{n \in \mathbb{N}}$ converge vers A , et $(B_n)_{n \in \mathbb{N}}$ vers B , alors $(A_n B_n)_{n \in \mathbb{N}}$ converge vers AB (pour n'importe quelle norme).

Exercice h : Dans $\mathcal{C}([0, 1], \mathbb{R})$, on considère la suite $(g_n)_{n \in \mathbb{N}} = (f_n \sqrt{n})_{n \in \mathbb{N}}$ (où les f_n sont définies à l'exercice e). Montrer que cette suite converge vers 0 pour la norme N_1 . Montrer que la suite $((g_n)^2)_{n \in \mathbb{N}}$ ne converge pas vers 0 pour N_1 .

Suites dans un espace de dimension finie.

Dans un espace de dimension finie la notion de suite convergente ne dépend pas de la norme avec laquelle on travaille, puisque sur un tel espace toutes les normes sont équivalentes.

Soit E de dimension finie, et soit $\{e_1, \dots, e_k\}$ une base de E , à toute suite $(u_n)_{n \in \mathbb{N}}$ d'éléments de E , nous pouvons associer ses suites coordonnées $((u_n)_i)_{n \in \mathbb{N}}$ (définies par $(\forall n) u_n = \sum_i (u_n)_i e_i$).

Alors la suite $(u_n)_{n \in \mathbb{N}}$ converge vers $u = \sum_i u_i e_i$, si et seulement si $(\forall i)$ la suite $((u_n)_i)_{n \in \mathbb{N}}$ converge vers u_i . En effet: Si $(u_n)_{n \in \mathbb{N}}$ converge vers u , $(N_\infty(u_n - u))_{n \in \mathbb{N}}$ converge vers 0. Nous pouvons choisir

de travailler ici avec N_∞ . Mais $(\forall i)$ nous avons $N_\infty(u_n - u) \geq |(u_n)_i - u_i|$, donc les suites $((u_n)_i - u_i)_{n \in \mathbb{N}}$ convergent vers 0.

Réciproquement si les suites $((u_n)_i - u_i)_{n \in \mathbb{N}}$ convergent vers 0, la somme $(\sum_i |(u_n)_i - u_i|)_{n \in \mathbb{N}}$ converge vers zéro. Mais cette somme n'est autre que la suite $(N_1(u_n - u))_{n \in \mathbb{N}}$; donc $(u_n)_{n \in \mathbb{N}}$ converge vers u .

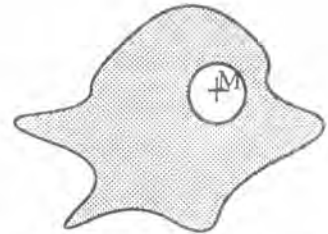
On notera l'utilisation de normes différentes dans les deux parties de la démonstration. C'est l'utilisation systématique du fait que, en dimension finie, toutes les normes sont équivalentes. Bien choisir la norme avec laquelle on travaille permet quelque fois de simplifier une démonstration de façon spectaculaire.

§ 3 : Éléments de topologie

Vocabulaire de la topologie

On dit qu'un sous-ensemble U d'un espace E muni d'une norme N , est ouvert pour N , si pour tout point u de U , il existe un nombre $\varepsilon > 0$, tel que $N(x-u) \leq \varepsilon$ implique que x est dans U .

On dit qu'un sous-ensemble X d'un espace E muni d'une norme N , est fermé pour N , si son complémentaire est ouvert pour N .



U est un ouvert, donc tout point M de U est centre d'un disque contenu dans U .

Exercice i : On considère l'espace vectoriel de dimension 2 formé d'un plan P muni d'une origine O . Les ensembles suivants sont-ils ouverts (pour la norme habituelle de P)? Sont-ils fermés? Sont-ils ni ouverts ni fermés?

- a) L'intérieur du cercle de centre O et de rayon 1.
- b) Une droite D de P .
- c) Le segment $[AB]$
- d) Le même segment privé de ses extrémités A et B .

Cas de deux normes équivalentes:

Supposons que l'espace E soit muni de deux normes N et N' et qu'il existe un nombre $K > 0$ tel que $(\forall x \in E)$ on ait $N(x) \leq K N'(x)$. Alors tout sous-ensemble U de E qui est ouvert pour N , est ouvert pour N' .

En effet : Quel que soit u dans U :

* $(\exists \varepsilon > 0)$ tel que $(\forall x) N(x-u) \leq \varepsilon$ implique $x \in U$.

* $N'(x-u) \leq \frac{\varepsilon}{K}$ entraîne $N(x-u) \leq \varepsilon$

* Donc $N'(x-u) \leq \varepsilon' = \frac{\varepsilon}{K}$ implique $x \in U$.

Il en résulte que si deux normes sur E sont équivalentes elles définissent les mêmes ouverts, et les mêmes fermés. Et ainsi sur un espace de dimension finie, toutes les normes donnent les mêmes ouverts et les mêmes fermés.

Exercice j : Démontrer que l'intersection de deux sous-ensembles U et V de E qui sont ouverts pour N , est un sous-ensemble ouvert pour N .

Que peut-on dire de $U \cup V$?

Exercice k : Soit X et Y deux sous-ensembles de E qui sont fermés pour N . Que peut-on dire de $X \cup Y$? De $X \cap Y$?

Soit $(u_n)_{n \in \mathbb{N}}$ une suite de E qui converge vers λ pour N , et soit U un sous-ensemble de E qui est ouvert pour N , et contient λ . Alors il existe N_0 tel que $(n \geq N_0)$ implique que u_n est dans U .

En effet:

- * $(\exists \varepsilon > 0)$ tel que $(\forall u) N(u - \lambda) \leq \varepsilon$ implique $(u \in U)$ (car U est ouvert).
- * $(\exists N_0)$ tel que $(\forall n) n \geq N_0$ implique $(N(u_n - \lambda) \leq \varepsilon)$ (car $\lim_{n \in \mathbb{N}} (u_n) = \lambda$)
- * Donc $(n \geq N_0)$ implique $(u_n \in U)$.

Soit un sous-ensemble F de E qui est fermé pour N . Et soit une suite $(u_n)_{n \in \mathbb{N}}$ d'éléments de F . Si $(u_n)_{n \in \mathbb{N}}$ converge, sa limite est dans F .

En effet soit U le complémentaire de F ; c'est un ouvert. Nous raisonnerons par l'absurde. Si la limite λ était dans U , il existerait (cf: ci-dessus) N_0 tel que tous les termes de la suite d'indice plus grand que N_0 soient dans U ; ce qui est contradictoire avec l'hypothèse qu'ils sont tous dans F .

Le théorème de Bolzano-Weierstrass

Soit $I = [a, b]$ un intervalle fermé et borné de $\mathbb{R}$. De toute suite de points de I , on peut extraire une suite convergente. Ce qui signifie que pour toute suite $(u_n)_{n \in \mathbb{N}}$ dans I il existe une application strictement croissante $\varphi: \mathbb{N} \rightarrow \mathbb{N}$ telle que la suite $(u_{\varphi(n)})_{n \in \mathbb{N}}$ converge.

Démonstration: Nous allons d'abord construire une suite d'intervalles I_n contenant chacun une infinité de termes de la suite, et tels que $(\forall n) I_n \supset I_{n+1}$. Cette construction se fait par récurrence:

On pose $I_0 = [a, b]$

Si $I_n = [a_n, b_n]$ est construit, on le partage en $[a_n, \frac{a_n+b_n}{2}]$ et $[\frac{a_n+b_n}{2}, b_n]$.

L'un au moins de ces deux intervalles contient une infinité de termes de la suite, on le choisit pour I_{n+1} . Ainsi $I_{n+1} = [a_{n+1}, b_{n+1}]$ est de longueur moitié de celle de I_n . Et de plus $a_{n+1} \geq a_n$ et $b_{n+1} \leq b_n$.

La suite a_n est donc croissante; comme elle est majorée par b , elle converge vers un nombre λ . Pour tout n , et tout $p > n$:

$$a_n \leq a_p \leq b_p \leq b_n$$

Donc :

$$a_n \leq \lambda = \lim_{p \in \mathbb{N}} (a_p) \leq b_n$$

Donc :

$$0 \leq b_n - \lambda \leq b_n - a_n = \frac{b-a}{2^n}$$

Donc b_n converge aussi vers λ .

Construisons maintenant la fonction φ . On pose $\varphi(0) = 0$; et $\varphi(n)$ est le plus petit entier tel que, d'une part $u_{\varphi(n)}$ soit dans I_n , d'autre part $\varphi(n) > \varphi(n-1)$ (c'est cette seconde condition qui assure que φ est strictement croissante). Alors pour tout n , nous avons $a_n \leq u_{\varphi(n)} \leq b_n$. Donc :

$$|u_{\varphi(n)} - \lambda| \leq b_n - a_n = \frac{b-a}{2^n}$$

Donc $(u_{\varphi(n)})_{n \in \mathbb{N}}$ converge vers λ .

Plus généralement, si X est un sous-ensemble fermé et borné^(*) d'un espace E de dimension finie, alors de toute suite de points de X , on peut extraire une suite convergente. Un sous-ensemble fermé et borné d'un espace E de dimension finie, est appelé un compact de E . Nous ferons la démonstration lorsque E est de dimension 2.

Choisissons une base $\{e_1, e_2\}$ de E . La suite vectorielle $(u_n)_{n \in \mathbb{N}}$ donnée, a deux suites coordonnées $(u^1_n)_{n \in \mathbb{N}}$ et $(u^2_n)_{n \in \mathbb{N}}$. Puisque X est borné, il est contenu dans un carré $[-A, +A] \times [-A, +A]$. La première suite coordonnée est contenue dans $[-A, +A]$; il existe donc $\varphi: \mathbb{N} \rightarrow \mathbb{N}$ strictement croissante telle que $(u^1_{\varphi(n)})_{n \in \mathbb{N}}$ converge.

(*) Un sous-ensemble B d'un espace normé, est dit borné, s'il existe un nombre M , tel que $(\forall x \in B) N(x) \leq M$.

La suite $(u^2_{\varphi(n)})_{n \in \mathbb{N}}$ n'a aucune raison de converger, mais elle est contenue dans $[-A, A]$, il existe donc $\psi: \mathbb{N} \rightarrow \mathbb{N}$ croissante telle que $(u^2_{\varphi(\psi(n))})_{n \in \mathbb{N}}$ converge.

Mais alors $(u^1_{\varphi(\psi(n))})_{n \in \mathbb{N}}$ converge aussi (c'est une sous-suite d'une suite convergente). Donc les deux suites coordonnées de la suite vectorielle $(u_{\varphi(\psi(n))})_{n \in \mathbb{N}}$ sont convergentes; ceci suffit à assurer que cette suite vectorielle converge.

Exercice l : Dans la démonstration ci-dessus préciser l'argument qui permet d'affirmer que l'un au moins des deux intervalles $[a_n, \frac{a_n+b_n}{2}]$, $[\frac{a_n+b_n}{2}, b_n]$ contient une infinité de termes de la suite.

Exercice m : Préciser l'argument qui permet d'affirmer que toute suite extraite d'une suite convergente, est elle-même convergente.

Exercice n : Soit $(u_n)_{n \in \mathbb{N}}$ une suite qui converge vers λ , montrer qu'on peut trouver une application croissante $\varphi: \mathbb{N} \rightarrow \mathbb{N}$ telle que $(u_{\varphi(n)} - \lambda)_{n \in \mathbb{N}} = o(\frac{1}{10^n})$. Construire explicitement une telle φ lorsque $u_n = \frac{1}{2^n}$.

§ 4 : La notion de continuité.

Soit X une partie d'un espace normé E (on notera N_E la norme de E), soit f une application de X dans un espace normé F (on notera N_F la norme de F), on dira que f est continue (pour les normes N_E et N_F) au point $x \in X$, si pour toute suite $(x_n)_{n \in \mathbb{N}}$ de points de X , qui converge vers x (pour N_E), la suite image $(f(x_n))_{n \in \mathbb{N}}$ converge vers $f(x)$ (pour N_F).

Notons que f est continue au point x si et seulement si :

$$(\forall \varepsilon > 0) (\exists \eta > 0) \text{ tel que } (\forall y) (N_E(y-x) \leq \eta) \text{ implique } (N_F(f(y) - f(x)) \leq \varepsilon)$$

Démontrons d'abord que cette seconde condition entraîne que l'image d'une suite $(x_n)_{n \in \mathbb{N}}$ qui converge vers x , est une suite qui converge vers $f(x)$.

* Soit $\varepsilon > 0$, nous lui associons $\eta > 0$ tel que $(N_E(y-x) \leq \eta)$ implique $(N_F(f(y) - f(x)) \leq \varepsilon)$

* Puisque $(x_n)_{n \in \mathbb{N}}$ converge vers x , il existe n_0 tel que $n \geq n_0$ implique $N_E(x_n - x) \leq \eta$.

* Donc $n \geq n_0$ entraîne : $N_F(f(x_n) - f(x)) \leq \varepsilon$.

Pour démontrer la réciproque, nous démontrerons que si la seconde condition n'est pas vérifiée (i.e. si $(\exists \varepsilon > 0)$ tel que $(\forall \eta > 0) (\exists y_\eta)$ tel que $(N_E(y_\eta - x) \leq \eta)$ et $(N_F(f(y_\eta) - f(x)) \geq \varepsilon)$, alors il existe une suite qui converge vers x et dont l'image ne converge pas vers $f(x)$. Pour cela choisissons $\eta = 1/n$; nous obtenons $y_n (= y_{1/n})$ tel que $N_E(y_n - x) \leq 1/n$ et tel que $N_F(f(y_n) - f(x)) \geq \varepsilon$; la suite $(y_n)_{n \in \mathbb{N}}$ est la suite cherchée.

On dit que $f: X \rightarrow F$ est continue (pour N_E et N_F) si elle est continue en tout point x de X (pour N_E et N_F). Ceci revient à dire que l'image par f de toute suite convergente (pour N_E) est une suite convergente (pour N_F).

Exercice p : Soit $f: X \subset E \rightarrow F$ telle que l'image par f de toute suite convergente dans X (c'est à dire de toute suite de points de X qui converge vers un point de X), soit une suite convergente dans F . On veut montrer que f est continue sur X . Pour cela on va raisonner par l'absurde. Supposons qu'il existe un point u dans X où f n'est pas continue :

a) Si f n'est pas continue en u , il existe une suite $(u_n)_{n \in \mathbb{N}}$ qui converge vers u et dont l'image par f ne converge pas vers $f(u)$. Montrer que la suite $(v_n)_{n \in \mathbb{N}}$ définie par $v_{2n} = u_n$ et $v_{2n+1} = u$, converge vers u .

b) Montrer que la suite $(f(v_n))_{n \in \mathbb{N}}$ n'a pas de limite.

c) Conclure.

Exercice q : 1) Soit N et N' des normes équivalentes sur E , et soit $f: E \rightarrow F$ continue pour N et pour une certaine norme N_F . Montrer que f est continue pour N' et pour N_F .

2) Soit v et v' deux normes équivalentes sur F , et soit $f: E \rightarrow F$ continue pour v et pour une certaine norme N_E sur E ; montrer que f est continue pour v' et pour N_E .

Exercice r : On munit $E = \mathcal{C}^\infty([0,1], \mathbb{R})$ de la norme N_∞ .

1) Soit $\phi: E \rightarrow E$ qui à toute u associe sa dérivée. Montrer que ϕ n'est continue en aucun point de E , pour la norme N_∞ (on montrera par exemple que $(\forall f \in E)$ la suite des fonctions $(t \rightarrow f(t) + \frac{1}{n} \sin nt)$ converge vers f pour N_∞ , et que la suite de leurs dérivées ne converge pas vers f').

2) Soit $\psi: E \rightarrow E$ qui à toute u associe la fonction $x \rightarrow \int_0^x u(t) dt$. Montrer que ψ est continue, pour la norme N_∞ .

On montrera d'abord que $(\forall u) (\forall x) \left| \int_0^x u(t) dt \right| \leq N_\infty(u)$. Puis que $(\forall u) N_\infty(\psi(u)) \leq N_\infty(u)$.

Les objets que l'on considère en analyse sont presque toujours définis comme des limites de suites, et ne sont pas connus exactement. C'est le cas des plus familiers d'entre eux comme $\sqrt{2}$, π , la fonction exponentielle, ... Une application continue est une fonction qui transforme tout objet de ce type, en un objet du même type. C'est donc une fonction dont on connaît raisonnablement la valeur sur ces objets familiers de l'analyse. Au contraire si une fonction n'est pas continue, il est souvent très difficile de calculer sa valeur sur ces objets familiers. Par exemple calculer la valeur en π de la fonction caractéristique de $\mathbb{Q}$ (qui vaut 1 sur les rationnels, et 0 sur les irrationnels) est un problème fort compliqué (connu jadis sous le nom de 'problème de la quadrature du cercle') qui ne fut résolu qu'au début du 19-ème siècle.

§ 5 : Les propriétés générales des fonctions continues

1) Soit $f: U \rightarrow F$ et $g: V \rightarrow G$ (où $U \subset E$ et $f(U) \subset V \subset F$) des applications continues (pour des normes N_E, N_F, N_G), alors $g \circ f: U \rightarrow G$ est continue (pour N_E et N_G).

En effet si $(u_n)_{n \in \mathbb{N}}$ est une suite convergente dans U , son image $(f(u_n))_{n \in \mathbb{N}}$ est convergente (puisque f est continue); et puisque g est continue, la suite $(g(f(u_n)))_{n \in \mathbb{N}}$ est aussi convergente; mais ceci prouve que $g \circ f$ est continue.

2) Si $f: U \rightarrow F$ et $g: U \rightarrow F$ sont continues (pour des normes N_E et N_F), alors $f+g: U \rightarrow F$ est continue (pour N_E et N_F).

Exercice s : Ecrire la démonstration. Pour cela montrer que l'image par $f+g$ de toute suite $(u_n)_{n \in \mathbb{N}}$ qui converge dans U , est une suite convergente dans F .

3) Si $f: U \rightarrow F$ et $\lambda: U \rightarrow \mathbb{R}$ sont continues (pour des normes de N_E, N_F de E et F , et pour la valeur absolue dans $\mathbb{R}$), alors $\lambda f: U \rightarrow F$ est continue (pour N_E et N_F).

Exercice t : Ecrire la démonstration.

4) Supposons que F soit muni d'un produit scalaire $\langle \cdot, \cdot \rangle$, et que $f: U \rightarrow F$ et $g: U \rightarrow F$ soient continues (pour une norme N_E sur E et pour la norme associée à ce produit scalaire), alors la fonction $\langle f, g \rangle: U \rightarrow \mathbb{R}$ (qui à u associe $\langle f(u), g(u) \rangle$) est continue (pour la norme N_E).

5) Supposons que F soit muni d'un produit et qu'il existe $K(>0)$ tel que $(\forall u \in F) (\forall v \in F) N_F(uv) \leq K N_F(u) N_F(v)$. Alors si $f: U \rightarrow F$ et $g: U \rightarrow F$ sont continues (pour une norme N_E sur E et pour N_F), la fonction $f \cdot g: U \rightarrow F$ (définie par $f \cdot g(u) = f(u) g(u)$) est continue. En particulier:

(Si $F = \mathbb{R}$ ou $\mathbb{C}$) Le produit de deux fonctions continues à valeurs réelles ou complexes est continue.

(Si F est l'espace des vecteurs de la géométrie) Le produit vectoriel de deux applications continues, est une application continue.

(Si F est l'espace des matrices $n \times n$) Le produit de deux applications continues à valeurs matricielles, est une application continue.

§ 6 : Fonctions continues sur les espaces de dimension finie.

Dans un espace de dimension finie toutes les normes sont équivalentes, donc (cf: ex q) si $f: X(\subset E) \rightarrow F$

Si E est de dimension finie, la continuité de f est indépendante de la norme choisie sur E .

Si F est de dimension finie, la continuité de f est indépendante de la norme choisie sur F .

Si E et F sont de dimension finie, la continuité de f est indépendante des normes choisies sur E et sur F .

Supposons F de dimension finie, et donnons nous une base de F . Toute $f: X \rightarrow F$, est définie par ses fonctions coordonnées f_j . Supposons d'abord que f soit continue: pour toute suite convergente $(x_n)_{n \in \mathbb{N}}$ dans X , la suite $(f(x_n))_{n \in \mathbb{N}}$ est convergente dans F , donc ses suites coordonnées $(f_i(x_n))_{n \in \mathbb{N}}$ sont convergentes; ce qui démontre que les f_i sont continues. Réciproquement, supposons que les f_i soient continues: pour toute suite convergente $(x_n)_{n \in \mathbb{N}}$ dans X , les suites $(f_i(x_n))_{n \in \mathbb{N}}$ sont convergentes dans $\mathbb{R}$, donc la suite $(f(x_n))_{n \in \mathbb{N}}$ est convergente dans F ; ce qui démontre que f est continue.

Nous retiendrons: Si F est de dimension finie une application $f: X \rightarrow F$ est continue si et seulement si ses fonctions coordonnées sont continues.

Si E est de dimension finie, et si $\{e_1, \dots, e_p\}$ est une base de E , l'application $f: X \rightarrow F$ apparaît comme une "fonction de p variables réelles"; c'est à dire que si $x = \sum x_j e_j$, nous pouvons noter $f(x_1, \dots, x_p)$ l'élément $f(x)$ de F . Nous pouvons alors considérer les fonctions partielles de f au point x , c'est à dire les fonctions (d'une seule variable numérique):

$$t \rightarrow f(t) = f(x_1, \dots, x_{j-1}, t, x_{j+1}, \dots, x_p)$$

Si f est continue elles sont continues (car composées de deux fonctions continues: f d'une part, et $t \rightarrow (x_1, \dots, x_{j-1}, t, x_{j+1}, \dots, x_p)$ d'autre part).

Mais les fonctions partielles de f en x peuvent être continues sans que f le soit (Ex. u).

Exercice u: On considère l'application $f: \mathbb{R}^2 \rightarrow \mathbb{R}$ définie par $f(x, y) = \frac{xy}{x^2 + y^2}$ si $(x, y) \neq (0, 0)$, et $f(0, 0) = 0$. Quelles sont ses fonctions partielles à l'origine? Montrer qu'elles sont continues. Puis montrer qu'il existe une suite $(u_n)_{n \in \mathbb{N}}$ de points de $\mathbb{R}^2$ qui converge vers $(0, 0)$, et dont l'image par f converge vers $1/2$.

§ 7 : Fonctions continues sur un compact

Nous savons que toute fonction continue sur un intervalle fermé et borné de $\mathbb{R}$ est majorée et minorée, et atteint ses bornes. Nous allons généraliser ces propriétés aux fonctions sur un espace vectoriel de dimension finie.

Soit K un sous-ensemble d'un espace E de dimension finie; on suppose que K est fermé et borné. Alors toute fonction continue de K dans $\mathbb{R}$:

a) est majorée et minorée (i.e.: il existe m et M tels que $(\forall x \in K) \quad m \leq f(x) \leq M$.

b) atteint ses bornes (i.e.: il existe $x_0 \in K$ et $x_1 \in K$ tels que $(\forall x \in K) \quad f(x_0) \leq f(x) \leq f(x_1)$.

Démontrons qu'il existe M tel que $(\forall x \in K) \quad f(x) \leq M$. Nous raisonnerons par l'absurde. S'il n'en était pas ainsi, pour tout $n \in \mathbb{N}$ il existerait x_n dans K , tel que $f(x_n) \geq n$. Puisque K est compact, nous pourrions extraire de la suite $(x_n)_{n \in \mathbb{N}}$ une sous-suite convergente $(x_{\phi(n)})_{n \in \mathbb{N}}$ (notons y sa limite).

Puisque f est continue en y , la suite $(f(x_{\varphi(n)}))_{n \in \mathbb{N}}$ convergerait vers $f(y)$, ce qui est absurde, puisque $f(x_{\varphi(n)}) \geq \varphi(n)$, et que la suite $(\varphi(n))_{n \in \mathbb{N}}$ converge vers l'infini.

Exercice v : Démontrer de même qu'il existe m tel que $(\forall x \in K) f(x) \geq m$.

Démontrons qu'il existe x_0 tel que $(\forall x \in K) f(x) \geq f(x_0)$. L'ensemble des valeurs de la fonction f est un ensemble de nombres réels minoré par m ; il a donc une borne inférieure b . Rappelons que b est un nombre tel que:

d'une part $(\forall x \in K) b \leq f(x)$

d'autre part $(\forall \varepsilon > 0) (\exists x \in K) \text{ tel que } b \leq f(x) \leq b + \varepsilon$.

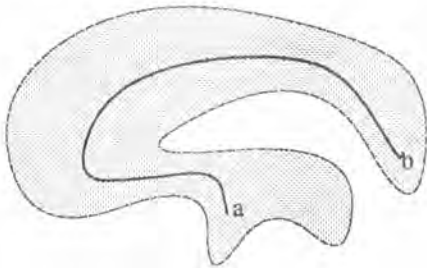
La seconde propriété nous permet de construire une suite $(y_n)_{n \in \mathbb{N}}$ telle que $(\forall n) b \leq f(y_n) \leq b + \frac{1}{n}$.

Puisque K est compact nous pouvons extraire de $(y_n)_{n \in \mathbb{N}}$ une suite convergente $(y_{\varphi(n)})_{n \in \mathbb{N}}$ (notons y sa limite). Puisque f est continue en y , la suite $(f(y_{\varphi(n)}))_{n \in \mathbb{N}}$ converge vers $f(y)$. Mais par ailleurs $(\forall n) b \leq f(y_{\varphi(n)}) \leq b + \frac{1}{\varphi(n)}$, donc la suite $(f(y_{\varphi(n)}))_{n \in \mathbb{N}}$ converge aussi vers b . Ce qui prouve que $f(y) = b$, et que y est le point x_0 cherché.

Exercice w : Démontrer de même qu'il existe x_1 tel que $(\forall x \in K) f(x) \leq f(x_1)$.

Le théorème de la valeur intermédiaire.

Il reste une propriété des fonctions définies sur un segment, que nous n'avons pas généralisée. Il s'agit du théorème de la valeur intermédiaire. Nous pouvons le faire de la façon suivante.



Nous dirons qu'un sous ensemble X d'un espace vectoriel E est connexe^() si pour tout couple (a, b) de points de X , il existe une fonction continue γ de $[0, 1]$ dans X telle que $\gamma(0) = a$ et $\gamma(1) = b$ (une telle fonction est aussi appelée un arc joignant a et b).*

Alors si f est une application continue d'un ensemble connexe X dans $\mathbb{R}$, et si A et B sont des valeurs de f , quel que soit C compris entre A et B , il existe (au moins) un point x dans X , tel que $f(x) = C$.

En effet si a et b sont des points tels que $f(a) = A$ et $f(b) = B$, considérons un arc γ joignant a et b . La fonction $f \circ \gamma$ est continue (comme composée d'applications continues) sur $[0, 1]$; donc nous pouvons lui appliquer le classique théorème de la valeur intermédiaire. Ainsi il existe $c \in [0, 1]$ tel que $f(\gamma(c)) = C$. Il suffit de poser $x = \gamma(c)$.

Notons qu'il en résulte que l'image d'un ensemble connexe par une fonction (à valeurs dans $\mathbb{R}$) continue est un intervalle (fermé, ouvert, ou semi-fermé; borné ou non borné).

§ 8 : La continuité uniforme

Soit $f : U \rightarrow F$, dire que f est continue sur U ($\subseteq E$), c'est dire que:

$$(\forall u \in U)(\forall \varepsilon > 0)(\exists \eta > 0)(\forall y \in U) N_E(y - u) \leq \eta \text{ implique } N_F(f(y) - f(u)) \leq \varepsilon$$

Nous dirons que f est uniformément continue sur U si:

$$(\forall \varepsilon > 0)(\exists \eta > 0)(\forall u \in U)(\forall y \in U) N_E(y - u) \leq \eta \text{ implique } N_F(f(y) - f(u)) \leq \varepsilon$$

(*) Ou plutôt "connexe par arcs".

Nous avons seulement changé la place de $(\forall u \in U)$. Dans la première formulation il était en tête, donc η dépendait de u ; dans la nouvelle formulation η ne dépend plus de u : il est possible de trouver un η qui convient quel que soit le choix de u (cf. chap 6b).

La continuité uniforme est une condition technique qui intervient dans de nombreuses démonstrations (par exemple pour montrer qu'une fonction continue sur $[a,b]$ est intégrable au sens de Riemann). Nous retiendrons:

Toute fonction dérivable sur un intervalle I de $\mathbb{R}$ dont la dérivée est bornée (i.e. il existe M tel que, quel que soit x , $|f'(x)| \leq M$), est uniformément continue.

En effet, d'après le théorème des accroissements finis, pour que $|f(y)-f(u)| \leq \varepsilon$ il suffit que $|y-u| \leq \frac{\varepsilon}{M}$. Donc on peut prendre (quel que soit u !) $\eta = \frac{\varepsilon}{M}$.

Toute fonction continue sur un compact est uniformément continue.

Exercice x: Soit K un sous-ensemble fermé borné d'un espace E de dimension finie, et soit f une application continue de K dans F . Pour démontrer qu'elle est uniformément continue nous allons raisonner par l'absurde. Supposons donc que f ne soit pas uniformément continue.

1) Ecrire formellement que f n'est pas uniformément continue. En déduire qu'il existe $\varepsilon > 0$, et deux suites $(y_n)_{n \in \mathbb{N}}$ et $(u_n)_{n \in \mathbb{N}}$ telles que, quel que soit n , $|y_n - u_n| \leq \frac{1}{n}$ et $|f(y_n) - f(u_n)| > \varepsilon$.

2) Montrer qu'il existe une application strictement croissante de $\mathbb{N}$ dans $\mathbb{N}$ telle que les suites $(y_{\phi(n)})_{n \in \mathbb{N}}$ et $(u_{\phi(n)})_{n \in \mathbb{N}}$ convergent.

3) Utiliser le fait que ces deux suites ont alors même limite λ pour montrer que f ne pourrait alors être continue en λ .

Exercices sur le chapitre 10

- * **Exercice 1:** Soit I un intervalle non fermé (ou non borné). Soit $\mathcal{C}_b(I, \mathbb{R})$ l'espace vectoriel des fonctions bornées de I dans $\mathbb{R}$; on pose :

$$N_\infty(f) = \sup_{t \in I} |f(t)|$$

- 1) Justifier l'existence de $N_\infty(f)$.
- 2) Montrer que l'on a ainsi défini une norme sur $\mathcal{C}_b(I, \mathbb{R})$.

- ** **Exercice 2:** Soit E un espace de dimension finie, et deux bases $\{e_1, \dots, e_n\}$ et $\{f_1, \dots, f_n\}$ de E . On note $A = (a_{ij})$ la matrice des coordonnées des f_j dans la base $\{e_i\}$. On note K un majorant des nombres $|a_{ij}|$.

- a) On note N_{e_∞} et N_{f_∞} les normes sur E obtenues en prenant le maximum des valeurs absolues des coordonnées dans chacune de ces deux bases. Montrer que $(\forall u) \ N_{e_\infty}(u) \leq nK \ N_{f_\infty}(u)$. En déduire que ces deux normes sont équivalentes.
- b) Montrer que les normes N_{e_2} et N_{f_2} que l'on peut définir à partir de ces deux bases, sont équivalentes.

- * **Exercice 3:** Soit E l'espace des fonctions continues sur $[0, \pi]$. Pour toute f dans E , on pose:

$$v_1(f) = \int_0^\pi |f(t)| \sin t \, dt \quad \text{et} \quad v_2(f) = \sqrt{\int_0^\pi (f(t))^2 \sin t \, dt}$$

- a) Définir sur E un produit scalaire $\langle \cdot, \cdot \rangle$, tel que v_2 soit la norme associée à ce produit scalaire.
- b) Démontrer que v_1 est une norme sur E .
- c) En appliquant l'inégalité de Schwarz au produit scalaire de $|f|$ et de 1, déterminer un nombre k , tel que $(\forall f \in E) \ v_1(f) \leq k \ v_2(f)$.
- d) En utilisant les fonctions $t \rightarrow \cos^n t$, montrer qu'il n'existe aucun nombre k' tel que $(\forall f \in E) \ v_2(f) \leq k' v_1(f)$.

- *** **Exercice 4:** Soit P le plan de la géométrie muni d'une origine O (et de son produit scalaire naturel). Soit H l'hexagone régulier de côté 1, centré en O . Pour tout M dans P , on pose:

$$v(M) = \left\{ \sup_{t \in H} t \right\}^{-1}$$

(pour $M = O$ cette formule n'a pas de sens, on pose $v(O) = 0$).

- a) Montrer que v est une norme.
- b) Trouver des constantes K et L telles que $(\forall M) \ K \|M\| \leq v(M) \leq L \|M\|$.

- ** **Exercice 5:** Pour toute f dans l'espace E des fonctions continues sur $[0, 2\pi]$, on pose:

$$v(f) = \left[\int_0^{2\pi} (k - \cos t) f^2(t) \, dt \right]^{1/2}$$

- a) Montrer que si $k \geq 1$, on a ainsi défini une norme sur E .
- b) Montrer que si $k > 1$, cette norme est équivalente à la norme de la moyenne quadratique.
- c) Soit f_n la fonction qui vaut $1 - nx$ si $x \leq 1/n$ et qui est nulle pour $x \geq 1/n$. Pour $k = 1$, calculer $v(f_n)$ et $N_2(f_n)$. En déduire que, pour $k = 1$, la norme v n'est pas équivalente à la norme de la moyenne quadratique.

- * **Exercice 6:** Pour toute f dans $\mathcal{C}([0,1], \mathbb{R})$ on pose $v(f) = N_\infty(f \cos)$. Montrer que v est une norme. Trouver des nombres K et K' tels que $(\forall f) \quad KN_\infty(f) \leq v(f) \leq K'N_\infty(f)$.
- * **Exercice 7:** On considère dans $\mathcal{C}([0,1], \mathbb{R})$ la suite $(f_n)_{n \in \mathbb{N}}$ définie par $f_n(t) = \sin(t + \frac{1}{n})$.
- Montrer que $(\forall t \in [0,1]) \quad |f_n(t) - \sin t| \leq \frac{1}{n}$. En déduire que $N_\infty(f_n - \sin) \leq \frac{1}{n}$, et que la suite $(f_n)_{n \in \mathbb{N}}$ converge vers $\sin$ pour la norme uniforme.
 - La suite $(f_n)_{n \in \mathbb{N}}$ converge-t-elle vers $\sin$ pour la norme N_2 ? Pour la norme N_1 ?
- ** **Exercice 8:** On considère dans $\mathcal{C}([0,1], \mathbb{R})$ la suite $(f_n)_{n \in \mathbb{N}}$ définie par $f_n(t) = \sin \pi t^n$.
- Calculer $N_\infty(f_n)$. En déduire que la suite $(f_n)_{n \in \mathbb{N}}$ ne converge pas vers 0 pour la norme uniforme.
 - En utilisant l'inégalité $|\sin u| \leq |u|$, montrer que $N_2(f_n) \leq \frac{\pi}{\sqrt{2n+1}}$. En déduire que $(f_n)_{n \in \mathbb{N}}$ converge vers 0 pour N_2 et pour N_1 .
 - Montrer que la suite $(f_n)_{n \in \mathbb{N}}$ n'a aucune limite pour la norme N_∞ .
- * **Exercice 9:**
- Dans l'espace $\mathcal{C}([0,1], \mathbb{R})$ on considère la suite $(f_n)_{n \in \mathbb{N}}$ définie par $f_n(t) = (1-t^2)^n$. Montrer que la suite $(f_n)_{n \in \mathbb{N}}$ ne converge pas vers 0 pour la norme uniforme.
 - On choisit arbitrairement α dans $]0,1[$, et dans l'espace $\mathcal{C}([\alpha,1], \mathbb{R})$ on considère la suite $(g_n)_{n \in \mathbb{N}}$ définie par $g_n(t) = (1-t^2)^n$. Montrer que la suite $(g_n)_{n \in \mathbb{N}}$ converge vers 0 pour la norme uniforme.
- * **Exercice 10:** On définit une suite $(f_n)_{n \in \mathbb{N}}$ dans $\mathcal{C}([0,1], \mathbb{R})$ en posant $f_n(t) = \frac{t^n}{n}$.
- Montrer que la suite $(f_n)_{n \in \mathbb{N}}$ converge vers la constante 1 pour la norme N_2 .
 - Montrer que la suite $(f_n)_{n \in \mathbb{N}}$ converge pour la norme N_1 . Quelle est sa limite?
 - Montrer que la suite $(f_n)_{n \in \mathbb{N}}$ n'a pas de limite pour la suite N_∞ .
- ** **Exercice 11:** On définit une suite $(u_n)_{n \geq 2}$ de nombres compris entre 0 et 1, de la façon suivante: si p est le plus grand facteur premier de l'entier n , $u_n = \frac{p}{n}$. En utilisant le fait que l'ensemble des nombres premiers est infini:
- Trouver une sous-suite de la suite $(u_n)_{n \in \mathbb{N}}$ qui converge vers 0.
 - Trouver une sous-suite de la suite $(u_n)_{n \in \mathbb{N}}$ qui converge vers 1.
 - Trouver une sous-suite de la suite $(u_n)_{n \in \mathbb{N}}$ qui converge vers $1/2$.
- * **Exercice 12:** Montrer que tout sous-ensemble fini de $\mathbb{R}^2$ est fermé.
- * **Exercice 13:** Trouver un sous-ensemble de $\mathbb{R}^2$ qui n'est ni ouvert ni fermé.
- * **Exercice 14:**
- Le sous-ensemble X de $\mathbb{R}^2$ formé des points $M(x,y)$ tels que $(x+y > 1 \text{ et } x-2y \leq 0 \text{ et } 3x+y > 2)$ est-il ouvert? est-il fermé? n'est-il ni l'un ni l'autre?
 - Mêmes questions pour $Y(x-y \leq 0 \text{ et } x+4y > 1 \text{ et } 2x-y < 0)$
 - Mêmes questions pour $Z(x+2y > 0 \text{ ou } x-y < 0)$

- ** Exercice 15:** a) Dans $\mathbb{R}$ on considère l'ensemble X formé de 0 et des points de la forme $\frac{1}{n}$ (où $n \in \mathbb{N}$). Soit u un point de $\mathbb{R}$ n'appartenant pas à X , déterminer un nombre k (>0) tel que l'intervalle $[u-k, u+k]$ ne rencontre pas X (on pourra distinguer plusieurs cas suivant que $u < 0$, $u > 1$ ou que $0 < u < 1$). Qu'en résulte-t-il pour X ?
 b) Dans $\mathbb{R}^2$ on considère le sous-ensemble Y formé des points d'une suite convergente $(u_n)_{n \in \mathbb{N}}$ et de sa limite λ . Montrer qu'il est fermé.
- * Exercice 16:** Soit E un espace muni d'une norme N ; soit $f: E \rightarrow \mathbb{R}$ une application continue pour N .
 a) Soit $U = f^{-1}(]0, \infty[)$ l'ensemble des points de E dont l'image par f est strictement positive. Montrer que U est ouvert dans E pour la norme N .
 b) Soit $F = f^{-1}([0, \infty[)$ l'ensemble des points de E dont l'image par f est positive ou nulle. Montrer que F est fermé dans E pour la norme N .
- * Exercice 17:** Soit E l'ensemble des séries numériques absolument sommables.
 a) Montrer que E est un espace vectoriel.
 b) Pour toute série $((u_n))_{n \in \mathbb{N}}$ dans E , on pose $N((u_n))_{n \in \mathbb{N}} = \sum_{i=0}^{\infty} |u_n|$. Montrer que ceci définit une norme sur E .
 c) A toute série on associe sa somme. Montrer que l'on obtient ainsi une application continue de E dans $\mathbb{R}$.
- * Exercice 18:** On considère dans $\mathcal{C}([a, b], \mathbb{R})$ une suite $(f_n)_{n \in \mathbb{N}}$ qui converge vers une limite g pour la norme N_{∞} .
 a) En utilisant le théorème des accroissements finis, montrer que $(\forall t) |\sin f_n(t) - \sin g(t)| \leq N_{\infty}(f_n - f)$. En déduire que la suite $(\sin \circ f_n)_{n \in \mathbb{N}}$ converge vers $\sin \circ g$ pour N_{∞} .
 b) On sait que l'image de la fonction g est un segment, on le note $[\alpha, \beta]$. Montrer que

$$(\forall \epsilon > 0) (\exists N) \text{ tel que } (\forall n) (\forall t \in [a, b]) \quad n \geq N \text{ implique } \alpha - \epsilon \leq f_n(t) \leq \beta + \epsilon$$
 En déduire qu'il existe A et B tels que $(\forall n) (\forall t \in [a, b]) \quad A \leq f_n(t) \leq B$.
 c) Montrer que la suite $(\exp \circ f_n)_{n \in \mathbb{N}}$ converge vers $\exp \circ g$ (où $\exp$ désigne la fonction exponentielle).
- ** Exercice 19:** Soit E l'espace des fonctions continument dérivables de $[0, 1]$ dans $\mathbb{R}$. On le munit de la norme uniforme. Montrer que l'application de E dans $\mathbb{R}$ qui à toute fonction associe sa dérivée en $1/2$, n'est pas continue.
- ** Exercice 20:** Soit E l'espace des fonctions continument dérivables de $[0, 2\pi]$ dans $\mathbb{R}$. Soit Φ l'application de E dans $\mathbb{R}$ qui à toute fonction f associe la longueur de son graphe.
 a) On munit E de la norme uniforme, montrer que Φ n'est pas continue en 0 (utiliser les fonctions $f_n(t) = \frac{1}{n} \sin nt$).
 b) On pose $v(f) = N_{\infty}(f) + N_{\infty}(f')$. Montrer que l'on a ainsi défini une norme sur E . Montrer que Φ est continue pour cette norme.
- *** Exercice 21:** Soit E l'espace des fonctions continues de $[0, 1]$ dans $\mathbb{R}$ muni de la norme uniforme,
 A) Les applications suivantes de E dans E sont-elles continues:
 a) $f \rightarrow f^2$
 b) $f \rightarrow e^f$
 c) $(t \rightarrow f(t)) \rightarrow (t \rightarrow f(t^2))$
 B) Reprendre la question précédente en munissant E de la norme de la moyenne quadratique.

** Exercice 22: Dans l'espace $\mathcal{C}([0, \pi], \mathbb{R})$ on considère les fonctions f_n ($n \in \mathbb{N}$) définies par $f_n(t) = \sin^n t$.

a) On pose $u_n = N_1(f_n)$. Montrer que la suite $(u_n)_{n \in \mathbb{N}}$ est décroissante et a une limite. En écrivant, pour un α quelconque dans $]0, \pi/2[$

$$\int_0^\pi f_n(t) dt = \int_0^\alpha f_n(t) dt + \int_\alpha^{\pi-\alpha} f_n(t) dt + \int_{\pi-\alpha}^\pi f_n(t) dt$$

montrer que $|u_n| \leq 2\alpha \sin^n \alpha + \pi - 2\alpha$. En déduire que la limite de $(u_n)_{n \in \mathbb{N}}$ est nulle.

b) Quelle est la limite de la suite numérique $(f_n(\pi/2))_{n \in \mathbb{N}}$?

Qu'en résulte-t-il pour l'application $ev_{\pi/2}: \mathcal{C}([0, \pi], \mathbb{R}) \rightarrow \mathbb{R}$ qui à toute fonction associe sa valeur en $\pi/2$?

c) L'application $ev_{\pi/2}: \mathcal{C}([0, \pi], \mathbb{R}) \rightarrow \mathbb{R}$ est-elle continue pour la norme N_2 ?

* Exercice 23: On identifie tout nombre réel α au couple $(\alpha, 0)$; ceci permet de considérer $\mathbb{R}$ comme un sous-ensemble de $\mathbb{R}^2$. Soit F un sous-ensemble de $\mathbb{R}$.

a) On suppose que F est un sous-ensemble fermé de $\mathbb{R}$. Montrer que F est un sous-ensemble fermé de $\mathbb{R}^2$.

b) On suppose que F est un sous-ensemble ouvert non vide de $\mathbb{R}$. Montrer que F n'est pas un sous-ensemble ouvert de $\mathbb{R}^2$.

* Exercice 24: Sur l'espace E des matrices $n \times n$, on pose $N(a_j^i) = \sup_{i,j} (|a_j^i|)$. Montrer que l'on a ainsi défini une norme.

Déterminer un nombre K tel que $(\forall A \text{ et } B) \quad N(AB) \leq K N(A) N(B)$.

* Exercice 25: Soit f et g des fonctions définies sur $[0, 1]$. Montrer que $(\forall t) \quad |f(t)g(t)| \leq N_\infty(f)N_\infty(g)$.

En déduire que $N_\infty(fg) \leq N_\infty(f)N_\infty(g)$.

Outils fondamentaux du raisonnement

Troisième partie: Ensembles ordonnés

§ 1 : Exemples de relations d'ordre.

On dit qu'un ensemble X est ordonné si on a y a défini une relation d'ordre; c'est à dire une relation, souvent notée $\leq$ et lue "inférieur ou égal à", qui vérifie les trois propriétés suivantes:

$$\alpha) (\forall x) \quad x \leq x$$

$$\beta) (\forall x) (\forall y) (\forall z) \text{ si } x \leq y \text{ et } y \leq z \text{ alors } x \leq z$$

$$\gamma) (\forall x) (\forall y) \text{ si } x \leq y \text{ et } y \leq x \text{ alors } x=y$$

Premier type d'exemples: $\mathbb{R}$ et tous ses sous-ensembles, en particulier $\mathbb{N}, \mathbb{Z}, \mathbb{Q} \dots$ Notons que deux nombres quelconques sont toujours comparables, ce qu'on exprime en disant que l'ordre de $\mathbb{R}$ (ou $\mathbb{N}, \mathbb{Z}, \mathbb{Q} \dots$) est un ordre total.

Deuxième type d'exemples: Soit E un ensemble, alors la relation "est inclus dans" est une relation d'ordre sur l'ensemble des parties de E (que l'on notera $\mathcal{P}(E)$). Ici il ne s'agit pas d'un ordre total.

Troisième type d'exemples: Si X est ordonné, alors, quelque soit Y , l'ensemble $\mathcal{F}(Y, X)$ des applications de Y dans X est ordonné; on dit que $f \leq g$ si $(\forall y \in Y) f(y) \leq g(y)$. C'est ainsi que les ensembles de fonctions numériques sont habituellement ordonnés. Ici encore il ne s'agit pas d'un ordre total.

Remarque: Lorsque X est ordonné au lieu d'écrire $y \leq x$, on se permet souvent d'écrire $x \geq y$. Et la relation $\geq$ est alors une nouvelle relation d'ordre sur X (c'est à dire qu'elle vérifie les propriétés α, β et γ ; c'est la relation opposée à $\leq$).

Remarque: Avec notre définition la relation "strictement plus petit que" (sur $\mathbb{R}$ par exemple) n'est pas une relation d'ordre (car elle ne vérifie pas α). En fait à toute relation d'ordre est associée une relation "strictement plus petit que"; l'expérience prouve qu'on l'utilise moins souvent que la relation "inférieure ou égal à"; c'est pourquoi on fait jouer le rôle principal à cette dernière.

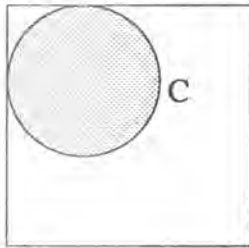
§ 2 : A la recherche d'un plus grand élément.

On dit que x_m est le plus grand élément de l'ensemble ordonné X si $x_m \in X$ et si $(\forall x \in X) x \leq x_m$; on dit aussi que x_m est l'élément maximum de X . Noter qu'on emploie aussi le mot maximum dans un sens très différent lorsqu'on parle du "maximum d'une fonction".

Le plus souvent les ensembles ordonnés n'ont pas de plus grand élément, c'est pourquoi on introduit les notions suivantes:

* Un élément y de X est dit maximal si aucun élément de X n'est strictement plus grand que y ; autrement dit si $y \in X$, et si

$$(\forall x \in X) \quad y \leq x \text{ implique } x=y.$$



Tout élément maximum est un élément maximal; la réciproque est vraie si X est totalement ordonné, elle ne l'est pas en général. Par exemple sur la figure ci contre, le cercle C est maximal (pour l'inclusion) parmi les cercles contenus dans le carré; mais il existe des cercles contenus dans le carré et qui ne sont pas contenus dans C .

* Si X est un sous-ensemble de l'ensemble ordonné E , on dit que l'élément μ de E est un majorant de X s'il est supérieur ou égal à tous les éléments de X ; autrement dit si $(\forall x \in X)$ alors $x \leq \mu$.

* Si X est un sous-ensemble de l'ensemble totalement ordonné E , on dit que l'élément m de E est la borne supérieure de X , si m est le plus petit majorant de X ; autrement dit si:

D'une part $(\forall x \in X)$ on a $x \leq m$.

D'autre part $(\forall \mu \in E)$ si μ est un majorant de X alors $m \leq \mu$.

Dans les exercices ces deux propriétés seront souvent utilisées de la façon suivante:

D'une part $(\forall x \in X)$ on a $x \leq m$.

D'autre part(*) si $(x' < m)$ alors $(\exists x \in X)$ tel que $(x' < x)$.

A titre d'exercice, formuler les définitions de "minimum", "minimal", "minorant" et "borne inférieure".

Exercices:

- * Exercice 1: Sur l'ensemble des suites numériques $(u_n)_{n \in \mathbb{N}}$ nous avons défini les relations "majorer", "majorer à partir d'un certain rang", $\circ$ et $\circ$. Parmi ces relations il n'y a qu'une relation d'ordre, laquelle ?
- * Exercice 2: La relation "avoir une plus grande aire que" est elle une relation d'ordre sur l'ensemble des disques du plan ? Est elle une relation d'ordre sur l'ensemble des disques centrés à l'origine ?
- * Exercice 3: Sur l'ensemble des suites numériques $(u_n)_{n \in \mathbb{N}}$ on définit une relation par:

$$(u_n)_{n \in \mathbb{N}} << (v_n)_{n \in \mathbb{N}} \Leftrightarrow \begin{cases} \text{ou bien } (\forall n) u_n = v_n \\ \text{ou bien le premier terme non nul de la suite } (v_n - u_n)_{n \in \mathbb{N}} \text{ est positif} \end{cases}$$
 Est ce une relation d'ordre ?
- * Exercice 4: Sur l'ensemble des fonctions polynômes la relation "majorer au voisinage de $+\infty$ ", est elle une relation d'ordre ?
- * Exercice 5: Montrer que tout ensemble fini totalement ordonné, a un maximum. On fera une récurrence sur le nombre des éléments. Ecrire un sous-ensemble fini de $\mathcal{P}(\mathbb{R})$ qui n'a pas de plus grand élément (pour la relation d'inclusion).
- * Exercice 6: Combien existe-t-il de relations d'ordre total sur un ensemble à n éléments ?
- * Exercice 7: Sur l'ensemble des rectangles du plan dont les cotés sont parallèles aux axes, on considère la relation

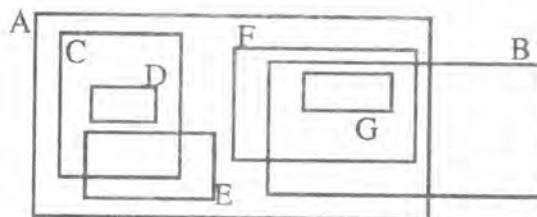
$$R \leq R' \text{ si et seulement si } R \text{ est contenu dans } R'.$$
 a) Montrer que c'est une relation d'ordre.

(*) Cette assertion peut être formulée en n'utilisant que les signes $\leq$ et $=$. Comment ?

b) On a dessiné ci-contre un ensemble de 5 rectangles. Quels sont les éléments maximaux de cet ensemble ? Quels sont ses éléments minimaux

c) Recommencer l'exercice en utilisant la relation d'ordre

$$R \leq R' \text{ si et seulement si } R \text{ contient } R'$$



Exercice 8: Ecrire toutes les relations d'ordre que l'on peut mettre sur l'ensemble à trois éléments $\{a, b, c\}$. Préciser celles qui sont des relations d'ordre total.

Préciser pour chacune d'elles les éléments maximaux de l'ensemble.

Exercice 9: Sur l'ensemble Δ des disques du plan, on considère la relation

$$D \leq D' \text{ si et seulement si } D \text{ est contenu dans } D'.$$

- Montrer que c'est une relation d'ordre, mais pas une relation d'ordre total.
- Dessiner un ensemble de deux disques D_1 et D_2 dont l'ensemble des majorants n'a pas de plus petit élément.
- Existe-t-il des sous-ensembles de Δ qui n'ont pas de minorant.
- Ces résultats sont-ils conservés si l'on remplace "disques" par "carré de cotés parallèles aux axes".

Exercice 10: On considère l'ensemble des fonctions sur $[0, 1]$ muni de sa relation d'ordre habituelle.

a) Soit X le sous-ensemble formé des fonctions $x \rightarrow x^2$, $x \rightarrow \text{Log}(1+x)$, $x \rightarrow e^x$, $x \rightarrow \sin x$. L'ensemble des majorants de X a un plus petit élément. Dessiner son graphe.

b) On considère le sous-ensemble Y formé des fonctions $x \rightarrow 1 - x^n$ ($n \geq 1$). L'ensemble des majorants de Y a-t-il un plus petit élément ?

- Que deviennent ces résultats si l'on remplace "fonctions" par "fonctions continues" ?
- Que deviennent ces résultats si l'on remplace "fonctions" par "fonctions de classe C^1 " ?

Exercice 11: Dans le complément au chapitre 8 nous avons défini la relation d'ordre sur $\mathbb{R}$. Démontrer que c'est bien une relation d'ordre.

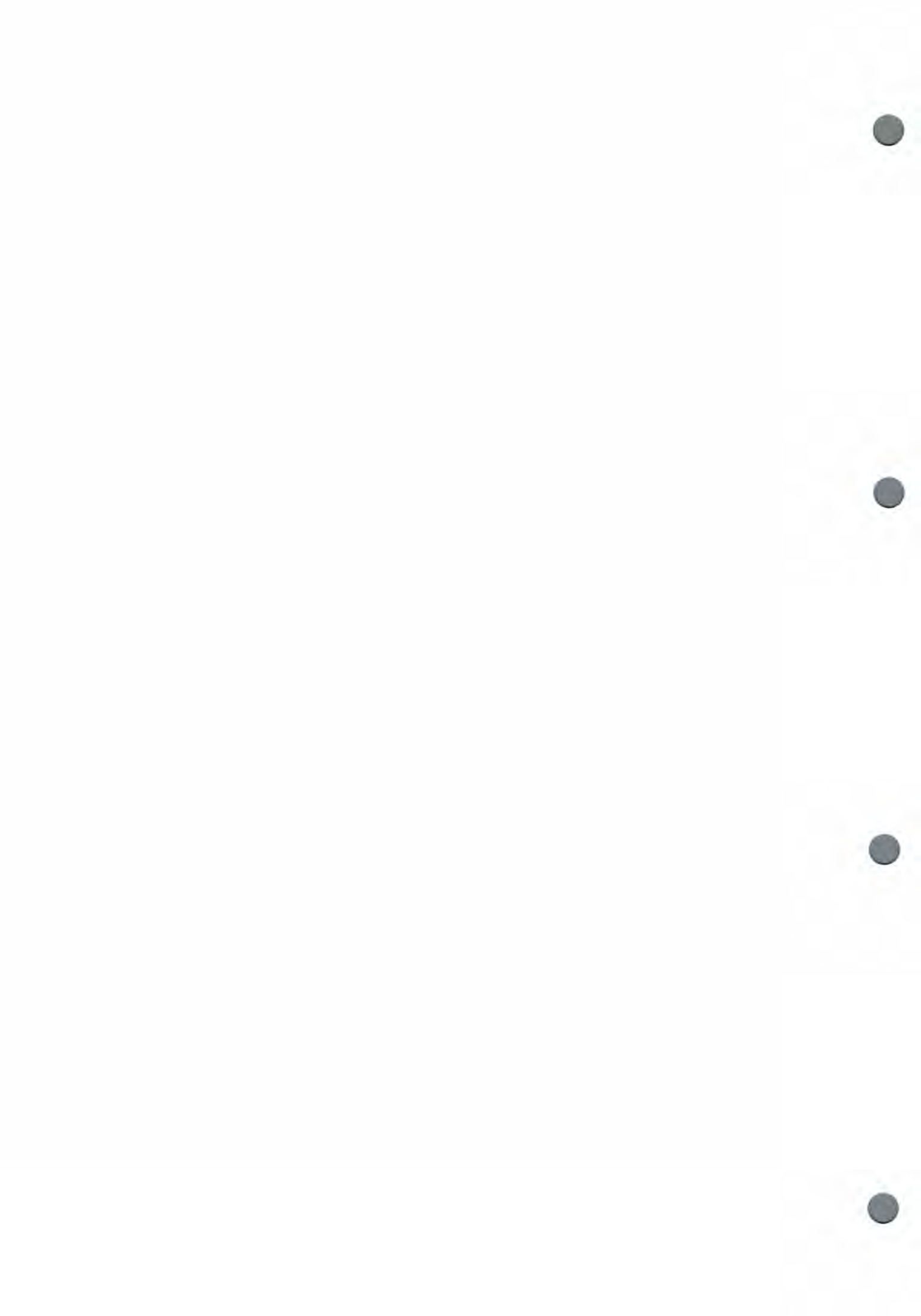


Table des matières

Premier fascicule.

Chap 1. Algèbre linéaire.

§1: Espaces vectoriels - sous-espaces vectoriels.	p. 5
§2: Familles libres - Familles génératrices - Bases.	p. 5
§3: Applications linéaires - Equations linéaires.	p. 7
§4: Le calcul matriciel.	p. 7
§5: Le dual.	p. 9
§6: Sous-espaces supplémentaires et sommes directes.	p. 10
§7: Algèbre linéaire sur un corps K .	p. 11
28 exercices sur le chapitre 1	p. 13

Chap 1b. Outils fondamentaux du raisonnement: Les opérations ensemblistes.

§1: Rappels des notations classiques de la théorie ensembliste.	p. 17
§2: Injections - Bijections - Surjections.	p. 17
14 exercices sur le chapitre 1b.	p. 18

Chap 2. Méthodes de calcul dans les espaces de dimension finie.

§1: La méthode du pivot.	p. 21
§2: Utilisation de la méthode du pivot en algèbre linéaire.	p. 22
§3: Le calcul des déterminants.	p. 23
§4: Utilisation des déterminants en algèbre linéaire.	p. 25
§5: Quand doit-on utiliser l'une au l'autre méthode ?	p. 27
25 exercices sur le chapitre 2.	p. 29

Chap 3. Utilisation de l'algèbre linéaire en analyse.

§1: Quelques espaces vectoriels de fonctions.	p. 33
§2: Quelques applications linéaires en analyse.	p. 34
§3: Equations différentielles linéaires du premier ordre.	p. 35
§4: Equations différentielles linéaires du second ordre à coefficients constants.	p. 36
Complément: Eléments d'une théorie de l'interpolation.	p. 36
21 exercices sur le chapitre 3.	p. 38

Chap 4. Produits scalaires.

§1: La notion de produit scalaire.	p. 41
§2: Norme euclidienne et angles.	p. 42
§3: Familles orthogonales et orthonormées.	p. 43
§4: Sous-espaces orthogonaux.	p. 44
§5: Endomorphismes orthogonaux et matrices orthogonales.	p. 45
§6: Extension au cas des espaces vectoriels sur $\mathbb{C}$.	p. 47
Complément 1: La méthode des moindres carrés.	p. 47
Complément 2: Orientation et produit vectoriel.	p. 49
22 exercices sur le chapitre 4	p. 51

Chap 5. Vecteurs propres et valeurs propres.

§1: Définitions générales.	p. 55
§2: Recherche des vecteurs propres et des valeurs propres en dimension finie.	p. 55
§3: Le polynôme caractéristique.	p. 56
§4: Diagonalisation.	p. 57
§5: Diagonalisation des matrices orthogonales et des matrices symétriques.	p. 58
27 exercices sur le chapitre 5.	p. 62

Chap 6. Etude des comportements asymptotiques.

§1: Les suites de nombres réels.	p. 65
§2: Etude des fonctions au voisinage de $+\infty$.	p. 66
§3: Etude des fonctions au voisinage de a .	p. 67
§4: Développements de Taylor.	p. 68
§5: Développements limités et développements de Laurent.	p. 68
21 exercices sur le chapitre 6	p. 70

Chap 6b. Outils fondamentaux du raisonnement:**Quantificateurs et expressions formalisées.**

§1: Pourquoi formalise-t-on l'écriture des mathématiques ?	p. 73
§2: Les règles d'écriture des expressions formalisées.	p. 73
11 exercices sur le chapitre 6b.	p. 75

Chap 7. L'intégrale définie.

§1: Les propriétés fondamentales d'une théorie de l'intégrale.	p. 78
§2: L'intégrale de Riemann.	p. 80
§3: L'intégrabilité des fonctions continues.	p. 81
§4: Quelques fonctions intégrables non continues.	p. 82
§5: Le retour aux primitives.	p. 83
§6: Les principales méthodes du calcul des primitives.	p. 84
Complément 1: Utilisation de l'intégrale dans le calcul des grandeurs physiques.	p. 86
Complément 2: Méthodes de calcul approché des intégrales définies.	p. 88
16 exercices sur le chapitre 7, et des calculs de primitives.	p. 91

Chap 8. La convergence à priori et les séries numériques.

§1: Les propriétés fondamentales des nombres réels.	p. 95
§2: Les suites de Cauchy dans $\mathbb{R}$.	p. 96
§3: Le langage des séries.	p. 97
§4: Séries à termes positifs.	p. 98
§5: Séries sommables et séries absolument sommables.	p. 99
§6: Séries qui ne sont pas absolument sommables.	p. 101
§7: Les séries $((\frac{1}{n^\alpha}))_{n \geq 1}$	p. 102
§8: Les séries géométriques.	p. 103
Complément: La construction de $\mathbb{R}$	p. 103
20 exercices sur le chapitre 8.	p. 106

Chap 9. Intégrales généralisées.

§1: Le critère de Cauchy pour les fonctions au voisinage de $+\infty$.	p. 109
§2: Les autres formes du critère de Cauchy.	p. 110
§3: Intégrales convergentes en $+\infty$.	p. 110
§4: Intégrales absolument convergentes en $+\infty$.	p. 112
§5: Intégrales absolument convergentes en a .	p. 113
§6: Théorie de l'intégrale sur les intervalles non compacts.	p. 114
§7: Convergence à l'infini de l'intégrale des fonctions oscillantes.	p. 114
18 exercices sur le chapitre 9.	p. 116

Chap 10. Espaces normés (Suites convergentes et applications continues).

§1: Normes sur un espace vectoriel.	p. 119
§2: Suites convergentes dans un espace normé.	p. 121
§3: Eléments de topologie.	p. 123
§4: La notion de continuité.	p. 125
§5: Les propriétés générales des fonctions continues.	p. 126
§6: Fonctions continues sur les espaces de dimension finie.	p. 126
§7: Fonctions continues sur un compact.	p. 127
§8: La continuité uniforme.	p. 128
25 exercices sur le chapitre 10.	p. 130

Chap 10b. Outils fondamentaux du raisonnement : Ensembles ordonnés.

§1: Relations d'ordre	p. 135
§2: A la recherche d'un plus grand élément.	p. 135
11 exercices sur le chapitre 10b.	p. 136

