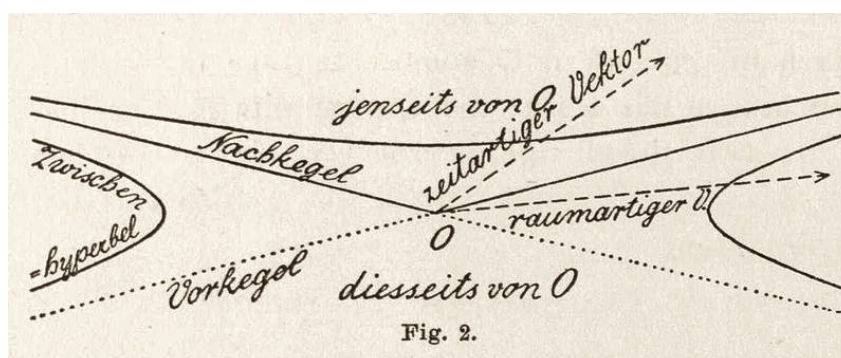
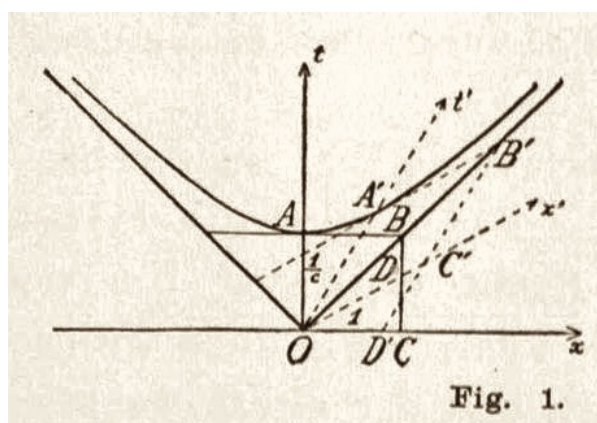


Algèbre Linéaire Avancée 2 pour physiciens

Marc Troyanov



Ce document contient les notes du cours d'algèbre linéaire avancée 2 enseigné à l'EPFL, printemps 2024 (document en cours de révision).

(© marc.troyanov, epfl)

Table des matières

9	Structure des endomorphismes	2
9.1	Triangulation des matrices et des endomorphismes	3
9.2	Polynômes d'endomorphismes et de matrices	6
9.3	Polynômes annulateurs et polynôme minimal d'un endomorphisme	7
9.4	Le théorème de Cayley-Hamilton	9
9.5	Une autre preuve du théorème de Cayley-Hamilton	12
9.6	Vecteurs propres généralisés et théorème de réduction primaire	14
9.7	Complément sur les multiplicités	18
9.8	Lemme des noyaux et preuve du théorème de réduction primaire	19
9.9	Décomposition de Dunford	21
9.10	Sous-espaces cycliques d'un endomorphisme	23
9.11	La forme normale de Jordan d'un endomorphisme	25
9.12	Conséquences du théorème de réduction de Jordan	28
9.13	Réduction pratique d'une matrice à sa forme normale de Jordan	30
9.14	Sur les endomorphismes d'espaces vectoriels réels.	34
10	Espace dual et Formes Bilinéaires	42
10.1	Espace Dual	42
10.1.1	Interpolation de Lagrange	46
10.2	Couplage entre deux espaces vectoriels	47
10.3	Formes bilinéaires sur un espace vectoriel	49
10.4	Formes bilinéaires symétriques et antisymétriques	52
10.5	Formes quadratiques	53
10.5.1	Réduction d'une forme quadratique à une somme de carré (méthode de complétion des carrés de Gauss)	55
11	Produits scalaires et espaces vectoriels euclidiens	57
11.1	Définitions fondamentales.	57
11.2	Orthogonalité dans un espace vectoriel euclidien	60
11.2.1	Projections orthogonales sur un sous-espace vectoriel	61
11.2.2	Symétries orthogonales	64
11.3	Le procédé d'orthonormalisation de Gram-Schmidt	65
11.4	Isométries d'un espace vectoriel euclidien.	66
11.5	Le groupe orthogonal	68
11.6	Espace-temps Galiléen et référentiels inertiels	71

11.7	Le théorème spectral (première version)	76
11.7.1	Application aux séries de Taylor des fonctions de n variables.	78
11.7.2	Application : le tenseur d'inertie et les moments d'inerties principaux . . .	79
12	Espaces vectoriels pseudo-euclidiens	82
12.1	Formes quadratiques sur un espace vectoriel réel, théorème de Sylvester.	82
12.2	Espaces pseudo-euclidiens	85
12.3	Base de Sylvester et espaces pseudo-euclidiens modèles.	87
12.4	Indicatrices et cône isotrope	88
12.5	L'espace-temps de Lorentz-Minkowski $\mathbb{E}^{1,d}$	90
12.6	L'inégalité de Cauchy-Schwarz inversée et quelques conséquences	92
13	Espaces hermitiens, opérateurs normaux et le théorème spectral	95
13.1	Formes sesquilineaires et formes hermitiennes sur un espace vectoriel complexe. .	95
13.2	Espaces vectoriels hermitiens	97
13.3	Opérateurs dans les espaces hermitiens	100
13.3.1	L'adjoint d'un opérateur	101
13.4	Endomorphismes normaux et le théorème spectral	103
13.5	Opérateurs auto-adjoints et unitaires	105
13.6	Espaces de Hilbert et opérateurs auto-adjoints	107
13.7	Applications en mécanique quantique.	108

Chapitre 9

Structure des endomorphismes

Introduction : Position du problème

Etant donné un endomorphisme $f \in \mathcal{L}(V)$ d'un espace vectoriel V de dimension finie sur un corps K , il est naturel de chercher à en analyser la structure. Le mot analyser vient du grec $\alpha\nu\alpha\lambda\upsilon\omega$ (*analuô*), qui signifie *délier*, *décomposer*. Pour analyser la structure d'un endomorphisme on cherche à le réduire en une somme directe d'endomorphismes les plus simples possibles.

De façon plus précise, nous allons chercher à décomposer l'espace V en somme directe de sous-espaces vectoriels qui sont invariants par f :

$$V = W_1 \oplus W_2 \oplus \cdots \oplus W_q, \quad f(W_i) \subset W_i \quad \text{pour tout } i.$$

de façon telle que $f_i = f|_{W_i} \in \mathcal{L}(W_i)$ soit un endomorphisme aussi simple que possible (noter que l'invariance de W_i est nécessaire pour que l'endomorphisme f_i soit bien défini).

Supposons une telle décomposition donnée, on peut alors choisir une base $\mathcal{B}_i$ de chaque sous-espace W_i et la réunion $\mathcal{B} = \mathcal{B}_1 \cup \cdots \cup \mathcal{B}_q$ de ces bases forme une base de V (car V est somme directe des W_i). Dans cette base la matrice de f prend la forme d'une matrice diagonale par blocs

$$M_{\mathcal{B}}(f) = A = \begin{pmatrix} \boxed{A_1} & & & \\ & \boxed{A_2} & & \\ & & \ddots & \\ & & & \boxed{A_q} \end{pmatrix},$$

où chaque A_i est la matrice de f_i dans la base $\mathcal{B}_i$.

Remarque. Nous nous permettons de noter parfois une telle matrice sous l'une des deux formes suivantes :

$$A = \text{Diag}(A_1, A_2, \dots, A_q) \quad \text{ou} \quad A = A_1 \oplus A_2 \oplus \cdots \oplus A_q.$$

La première notation nous rappelle que A est une matrice "diagonale par bloc" et la seconde notation nous rappelle que A est la matrice d'un endomorphisme qui laisse invariante une décomposition de V en somme directe dans une base adaptée.

Par exemple

$$\begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix} \oplus (5) \oplus (6) = \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix} \oplus \begin{pmatrix} 5 & 0 \\ 0 & 6 \end{pmatrix} = \begin{pmatrix} 1 & 2 & 0 & 0 \\ 3 & 4 & 0 & 0 \\ 0 & 0 & 5 & 0 \\ 0 & 0 & 0 & 6 \end{pmatrix}$$

Observons que deux endomorphismes conjugués ont la même structure dans le sens suivant : Supposons que $f, f' \in \mathcal{L}(V)$ sont conjugués par un automorphisme g , i.e. $f' = g \circ f \circ g^{-1}$ et que V admet une décomposition comme somme directe de sous-espaces $W_1, \dots, W_q$ invariants par f , alors les sous-espaces $W'_i = g(W_i)$ sont invariants par f' et l'espace V est aussi somme directe des W'_i . De plus si $\mathcal{B}_i$ est une base de W_i , alors $g(\mathcal{B}_i)$ est une base de W'_i et f et f' ont même matrice dans les bases respectives $\mathcal{B} = \mathcal{B}_1 \cup \dots \cup \mathcal{B}_q$ et $\mathcal{B}' = \mathcal{B}'_1 \cup \dots \cup \mathcal{B}'_q$:

$$M_{\mathcal{B}'}(f') = M_{\mathcal{B}}(f) = A = A_1 \oplus A_2 \oplus \dots \oplus A_q, \quad \text{avec} \quad A_i = M_{\mathcal{B}'_i}(f'_i) = M_{\mathcal{B}_i}(f_i).$$

Ces considérations s'étendent aux matrices carrées. Analyser la structure d'une matrice $B \in M_n(K)$ revient à analyser l'endomorphisme correspondant $L_B : K^n \rightarrow K^n$, défini par $L_B(X) = BX$, et à décomposer K^n en somme directe de sous-espaces vectoriels invariants par L_B , puis choisir une base $\mathcal{B}$ de K^n adaptée à cette décomposition de K^n et finalement à faire le changement de base pour obtenir une matrice diagonale par blocs $A = PBP^{-1}$.

Noter que dans ce cas, la matrice de changement de bases P est la matrice dont la $j^{\text{ème}}$ colonne est donnée par les composantes dans la base canonique du $j^{\text{ème}}$ vecteur de la base $\mathcal{B}$.

De même que deux endomorphismes conjugués ont la même structure, deux matrices semblables ont aussi la même structure, et donc les même écritures comme matrices diagonales par blocs (après changement de base).

Le cas le plus simple est celui d'un endomorphisme (ou d'une matrice) diagonalisable. La structure d'un tel endomorphisme est simplement donnée par la décomposition de l'espace V en somme directe de sous-espaces invariants de dimension 1 (la base obtenue étant une base formée de vecteurs propres).

Dans ce qui suit nous étudions la structure des endomorphismes dont le polynôme caractéristique est scindé (c'est toujours le cas si $K = \mathbb{C}$). Les invariants déjà connus pour analyser un tel endomorphisme (ou une telle matrice) sont : son *polynôme caractéristique*, son *spectre* ainsi que les *multiplicités algébrique* et *géométrique* de chaque valeur propre. Nous verrons dans ce chapitre d'autres invariants tels que le *polynôme minimal* et les multiplicités généralisées.

9.1 Triangulation des matrices et des endomorphismes

Rappelons qu'une matrice carrée $A = (a_{ij})$ est dite *triangulaire supérieure* si $a_{ij} = 0$ pour $i > j$, une matrice triangulaire est donc de la forme

$$A = \begin{pmatrix} a_{11} & * & * & \cdots & * \\ 0 & a_{22} & * & & \vdots \\ \vdots & 0 & a_{33} & & \vdots \\ \vdots & \vdots & & \ddots & * \\ 0 & 0 & 0 & \cdots & a_{nn} \end{pmatrix} \quad (9.1)$$

La matrice $A = (a_{ij})$ est *triangulaire inférieure* si $a_{ij} = 0$ pour $i < j$. La transposée d'une matrice triangulaire inférieure est une matrice triangulaire supérieure et, dans la suite, on dira simplement qu'une matrice est *triangulaire* lorsqu'elle est triangulaire supérieure.

Définitions. La matrice $A \in M_n(K)$ est dite *triangulable*¹ si elle est semblable à une matrice triangulaire, i.e. s'il existe $P \in \text{GL}_n(K)$ tel que $P^{-1}AP$ est triangulaire.

On dit qu'un endomorphisme $f \in \mathcal{L}(V)$ d'un K -espace vectoriel de dimension finie est *triangulable* si sa matrice dans une base adéquate est triangulaire.

La proposition suivante est une reformulation de cette définition :

Proposition 9.1.1. *Soit V un K -espace vectoriel de dimension finie. L'endomorphisme $f \in \mathcal{L}(V)$ est triangulable si et seulement si il existe une base $\mathcal{B} = \{v_1, \dots, v_n\}$ de V et des scalaires $a_{ij} \in K$ tels que*

$$f(v_j) = \sum_{i \leq j} a_{ij} v_i = \sum_{i=1}^j a_{ij} v_i. \quad (9.2)$$

Noter que la condition (9.2) peut aussi s'écrire

$$f(v_j) \in \text{Vec}(\{v_1, v_2, \dots, v_j\}).$$

Théorème 9.1.2. *Soit V un K -espace vectoriel de dimension finie. Un endomorphisme $f \in \mathcal{L}(V)$ est triangulable si et seulement si son polynôme caractéristique $\chi_f(t)$ est scindé.*

Rappelons qu'un polynôme est dit *scindé* s'il est produit de polynômes de degré 1.

Preuve. Si f est triangulable, alors il existe une base dans laquelle la matrice de f prend la forme (9.1). Le polynôme caractéristique de f est donc

$$(t - a_{11})(t - a_{22}) \dots (t - a_{nn}),$$

en particulier ce polynôme est scindé.

On démontre la réciproque par récurrence sur la dimension n du K -espace vectoriel V . Si $n = 1$, il n'y a rien à démontrer car toute matrice de taille 1×1 est triangulaire. Soit $n = \dim(V) > 1$ et supposons le théorème démontré pour tout espace vectoriel de dimension $n - 1$.

Par hypothèse, le polynôme caractéristique $\chi_f(t)$ est scindé. En particulier il existe au moins une racine $\lambda_1 \in K$ de $\chi_f(t)$. Soit $v_1 \in V$ un vecteur propre associé à cette valeur propre, i.e. $v_1 \neq 0$ et $f(v_1) = \lambda_1 v_1$. Choisissons maintenant $v_2, \dots, v_n \in V$ tels que $\mathcal{B} = \{v_1, v_2, \dots, v_n\}$ est une base de V . Notons $W_1 = \text{Vec}(v_1) = Kv_1$ et $W_2 = \text{Vec}(v_2, \dots, v_n)$. Alors $W_2 \subset V$ est un sous-espace vectoriel de dimension $n - 1$ tel que $V = W_1 \oplus W_2$. Notons encore $\pi_1 : V \rightarrow W_1$ et $\pi_2 : V \rightarrow W_2$ les projections canoniques, alors $f = f_1 + f_2$, avec $f_i = \pi_i \circ f$. La matrice de f dans la base $\mathcal{B}$ prend la forme

$$\begin{pmatrix} \lambda_1 & * & \cdots & * \\ 0 & & & \\ \vdots & & S & \\ 0 & & & \end{pmatrix}$$

1. On dit parfois que la matrice est *trigonalisable*.

où S est la matrice de l'endomorphisme $g_2 = f_2|_{W_2} \in \mathcal{L}(W_2)$ dans la base $\{v_2, \dots, v_n\}$. Le polynôme caractéristique de f est égal au polynôme caractéristique de la matrice ci-dessus, donc

$$\chi_f(t) = (t - \lambda_1) \cdot \chi_S(t) = (t - \lambda_1) \cdot \chi_{g_2}(t)$$

Ceci entraîne en particulier que le polynôme $\chi_{g_2}(t)$ est également scindé. Par hypothèse de récurrence, g_2 est triangulable et on peut donc trouver une (nouvelle) base $\{v'_2, \dots, v'_n\}$ de W_2 dans laquelle la matrice S' de g_2 est triangulaire.

Observons que pour $j \geq 2$ on a $f(v'_j) = f_1(v'_j) + f_2(v'_j) = f_1(v'_j) + g_2(v'_j)$. Or $f_1(v'_j)$ est un multiple de v_1 , donc la matrice de f dans la base $\{v_1, v'_2, \dots, v'_n\}$ prend la forme

$$\begin{pmatrix} \lambda_1 & * & \cdots & * \\ 0 & & & \\ \vdots & & S' & \\ 0 & & & \end{pmatrix}$$

où $S' \in M_{n-1}(K)$ est triangulaire. Il s'agit donc d'une matrice triangulaire de $M_n(K)$. □

Corollaire 9.1.3. *Toute matrice $A \in M_n(\mathbb{C})$ est triangulable.*

Preuve. Sur $\mathbb{C}$ tout polynôme est scindé par le théorème fondamental de l'algèbre. □

Corollaire 9.1.4. *Le coefficient d'ordre $n - 1$ du polynôme caractéristique d'une matrice $A \in M_n(\mathbb{C})$ est égal à l'opposé de sa trace.*

Preuve. On sait que deux matrices semblables ont la même trace et le même polynôme caractéristique. Par le corollaire précédent, on peut donc supposer que la matrice A est de la forme (9.1). On a alors

$$\begin{aligned} \chi_A(t) &= \prod_{i=1}^n (t - a_{ii}) = t^n - \left(\sum_{i=1}^n a_{ii} \right) t^{n-1} + \dots + (-1)^n \prod_{i=1}^n a_{ii} \\ &= t^n - \text{Tr}(A) t^{n-1} + \dots + (-1)^n \det(A). \end{aligned}$$
□

Remarque. On peut aussi prouver ce corollaire directement en examinant la définition du polynôme caractéristique.

Corollaire 9.1.5. (A) *La trace d'une matrice $A \in M_n(\mathbb{C})$ est la somme de ses valeurs propres comptées selon leur multiplicité algébrique :*

$$\text{Tr}(A) = \sum_{\lambda \in \sigma(A)} \lambda \cdot \text{multalg}_A(\lambda).$$

(B) *Le déterminant d'une matrice $A \in M_n(\mathbb{C})$ est le produit de ses valeurs propres comptées selon leur multiplicité algébrique :*

$$\det(A) = \prod_{\lambda \in \sigma(A)} \lambda^{\text{multalg}_A(\lambda)}.$$

Preuve. (A) Cette formule est évidente pour une matrice triangulaire T . Le théorème précédent nous dit que toute matrice $A \in M_n(\mathbb{C})$ est triangulable. Il existe donc $P \in GL_n(\mathbb{C})$ telle que $T = P^{-1}AP$ est triangulaire. Cela termine la preuve car A et T ont les mêmes valeurs propres et $\text{Tr}(A) = \text{Tr}(P^{-1}AP)$.

(B) Le raisonnement est le même pour le déterminant. □

Remarque. Le résultat du corollaire précédent reste valable, avec la même preuve, pour toute matrice carrée A à coefficients dans un corps K quelconque, à condition que son polynôme caractéristique soit scindé.

9.2 Polynômes d'endomorphismes et de matrices

La philosophie pour la suite de ce chapitre est la suivante : pour analyser la structure d'un endomorphisme $f \in \mathcal{L}(V)$, on essaye de décomposer V en somme directe de sous-espaces invariants et on analyse la structure de f sur les sous-espaces invariants.

Polynôme d'un endomorphisme

Une opération qui jouera un rôle fondamental est la suivante : soit $f \in \mathcal{L}(V)$ un endomorphisme de V et $p(t) = a_0 + a_1t + \cdots + a_k t^k$ un polynôme à coefficients dans le corps K . Alors on note $p(f) \in \mathcal{L}(V)$ l'endomorphisme obtenu en substituant l'endomorphisme f à l'indéterminée t :

$$p(f) = a_0 \cdot \text{Id}_V + a_1 \cdot f + \cdots + a_k \cdot f^k \in \mathcal{L}(V),$$

où par définition f^m signifie $f \circ f \circ \cdots \circ f$ (m fois) et $f^0 = \text{Id}_V$. De même, si $A \in M_n(K)$ alors on définit $p(A) \in M_n(K)$ par

$$p(A) = a_0 \mathbf{I}_n + a_1 A + \cdots + a_k A^k.$$

On vérifie alors facilement les résultats suivants :

Théorème 9.2.1. (a) Pour un endomorphisme $f \in \mathcal{L}(V)$ donné, l'application $K[t] \rightarrow \mathcal{L}(V)$ donnée par $p(t) \mapsto p(f)$ est un homomorphisme de K -algèbres. En particulier si $p, q \in K[t]$, alors

$$(p \cdot q)(f) = p(f) \circ q(f).$$

(b) Si $W \subset V$ est un sous espace invariant par f , alors ce sous-espace est aussi invariant par $p(f)$.

(c) Si $f \in \mathcal{L}(V)$ et $g \in GL(V)$, alors pour tout $p \in K[t]$, on a

$$p(g^{-1}fg) = g^{-1}p(f)g.$$

(d) Si $v \in V$ est un vecteur propre de f et λ est la valeur propre associée, alors v est aussi un vecteur propre de $p(f)$ et la valeur propre associée est $p(\lambda)$

Remarquons en particulier que la propriété (a) implique que pour tous polynômes $p(t), q(t) \in K[t]$, l'endomorphisme $p(f)$ commute avec $q(f)$:

$$p(f) \circ q(f) = q(f) \circ p(f).$$

La preuve de cette proposition consiste simplement à vérifier les définitions. Démontrons par exemple l'assertion (d). Observons d'abord que si $f(v) = \lambda v$, alors pour tout entier k on a $f^k(v) = \lambda^k v$. Soit maintenant $p(t) = \sum_{k=0}^m a_k t^k$ un polynôme quelconque, alors on a

$$p(f)(v) = \sum_{k=0}^m a_k f^k(v) = \sum_{k=0}^m a_k \lambda^k v = p(\lambda)v.$$

Les propriétés correspondantes sont aussi vraies pour les matrices, en remplaçant la composition par la multiplication matricielle.

Remarque. La réciproque de la propriété (d) est fausse. Voici un contre-exemple : considérons la matrice $A = \begin{pmatrix} 0 & -2 \\ 2 & 0 \end{pmatrix}$ et le polynôme $p(t) = t^4$. Alors $p(A) = A^4 = \begin{pmatrix} 16 & 0 \\ 0 & 16 \end{pmatrix}$, donc $16 = 2^4 \in \sigma(A)$, mais 2 n'est pas valeur propre de A (la matrice A n'a aucune valeur propre réelle et ses valeurs propres complexes sont $\pm 2i$).

9.3 Polynômes annulateurs et polynôme minimal d'un endomorphisme

Définition 9.3.1. On dit qu'un polynôme $p(t)$ *annule* la matrice $A \in M_n(K)$, ou que c'est une *polynôme annulateur* de A si $p(A) = 0$.

De même, si $f \in \mathcal{L}(V)$ est un endomorphisme d'un K -espace vectoriel, on dit que $p \in K[t]$ *annule* f , ou que c'est un *polynôme annulateur* de f si $p(f) = 0 \in \mathcal{L}(V)$, i.e. $p(f)$ est l'endomorphisme nul. On remarque que $p(t)$ est un polynôme annulateur de f si et seulement si $p(f)(v) = 0$ pour tout $v \in V$; de façon équivalente $\text{Ker}(p(f)) = V$.

Exemples 1. L'endomorphisme f est dit *nilpotent* s'il existe m tel que $f^m = 0$. Dans ce cas le polynôme t^m est un polynôme annulateur de f .

2. La matrice $A = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}$ vérifie $A^4 = I_2$. Par conséquent $t^4 - 1$ est un polynôme annulateur de A .

3. L'opérateur de dérivation $D = \frac{d}{dx}$ est un endomorphisme de $\mathbb{R}[x]$ qui n'admet aucun polynôme annulateur non nul.

Proposition 9.3.2. (a) Si $f \in \mathcal{L}(V)$ et $g \in GL(V)$, alors tout polynôme qui annule f annule aussi $g^{-1} \circ f \circ g$.

(b) Deux matrices semblables ont les mêmes polynômes annulateurs.

(c) Si A est la matrice de f dans une base quelconque de V , alors $p(t)$ est un polynôme annulateur de f si et seulement si c'est un polynôme annulateur de A .

Preuve. Exercice. □

Lemme 9.3.3. *Tout endomorphisme $f \in \mathcal{L}(V)$ d'un espace vectoriel de dimension finie admet des polynômes annulateurs non nuls.*

Preuve. Puisque $\dim \mathcal{L}(V) < \infty$, il existe un entier k tel que la famille d'endomorphismes

$$\{\text{Id}_V, f, f^2, \dots, f^k\} \subset \mathcal{L}(V)$$

est liée. Soit k le plus petit entier avec cette propriété, il existe alors des scalaires $\alpha_0, \alpha_1, \dots, \alpha_{k-1}$, uniquement définis et tels que

$$f^k + \alpha_{k-1}f^{k-1} + \dots + \alpha_1 f + \alpha_0 \text{Id}_V = 0 \in \mathcal{L}(V),$$

ce qui signifie que le polynôme

$$\mu_f(t) = t^k + \sum_{i=0}^{k-1} \alpha_i t^i \tag{9.3}$$

annule f . □

Définition. Le polynôme construit en (9.3) s'appelle le *polynôme minimal* de f . On définit le polynôme minimal d'une matrice de la même manière.

Le polynôme minimal d'un endomorphisme possède les propriétés importantes suivantes :

Proposition 9.3.4. *Soit V un K -espace vectoriel de dimension finie.*

- (a) *Le polynôme minimal $\mu_f(t)$ d'un endomorphisme f est l'unique polynôme unitaire de plus petit degré qui annule f .*
- (b) *Ce polynôme divise tout polynôme qui annule f .*
- (c) *Deux endomorphismes conjugués ont le même polynôme minimal, i.e. si $f \in \mathcal{L}(V)$ et $g \in \text{Aut}(V)$ alors $\mu_{g^{-1} \circ f \circ g} = \mu_f(t)$.*

Preuve. La propriété (a) vient de la définition de $\mu_f(t)$. Observer que l'unicité vient de l'indépendance linéaire des endomorphismes $\text{Id}_V, f, f^2, \dots, f^{k-1} \in \mathcal{L}(V)$.

Pour prouver (b), on considère un autre polynôme $p(t)$ annulant f . Si $p(t)$ est non nul, alors $\deg(p) \geq \deg(\mu)$. En appliquant la division polynomiale, il existe deux polynômes $q(t)$ et $r(t)$ tels que

$$p(t) = q(t)\mu_f(t) + r(t) \quad \text{et} \quad \deg r(t) < \deg \mu_f(t).$$

Observons que $r(f) = p(f) - q(f) \circ \mu_f(f) = 0$, donc $r(t)$ est le polynôme nul par minimalité du degré de $\mu_f(t)$. On a donc montré que tout polynôme annulateur de f est un multiple de $\mu_f(t)$. La propriété (c) est conséquence du fait que les endomorphismes f et $g^{-1} \circ f \circ g$ ont les mêmes polynômes annulateurs, ils ont donc le même polynôme minimal. □

Les mêmes définitions s'appliquent aux matrices, et on peut énoncer en particulier le résultat suivant :

Proposition 9.3.5. *Soient $A, B \in M_n(K)$. Si B est semblable à A alors $\mu_B(t) = \mu_A(t)$.*

9.4 Le théorème de Cayley-Hamilton

Le théorème de Cayley-Hamilton dit que tout endomorphisme d'un espace vectoriel de dimension finie est annulé par son polynôme caractéristique.

Théorème 9.4.1 (Cayley-Hamilton). *Pour tout endomorphisme $f \in \mathcal{L}(V)$ d'un espace vectoriel de dimension finie, on a*

$$\chi_f(f) = 0.$$

De même, pour toute matrice $A \in M_n(K)$ on a $\chi_A(A) = 0$. Autrement dit le polynôme caractéristique de A est un polynôme annulateur de A .

Exemple. On rappelle que le polynôme caractéristique de la matrice $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$ est $\chi_A(t) = t^2 - (a + d)t + (ad - bc)$, en appliquant ce polynôme à la matrice elle-même, on calcule que

$$\begin{aligned} \chi_A(A) &= A^2 - (a + d)A + (ad - bc)I_2 \\ &= \begin{pmatrix} a & b \\ c & d \end{pmatrix}^2 - (a + d) \begin{pmatrix} a & b \\ c & d \end{pmatrix} + (ad - bc) \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \\ &= \begin{pmatrix} a^2 + bc & ab + bd \\ ca + dc & cb + d^2 \end{pmatrix} - (a + d) \begin{pmatrix} a & b \\ c & d \end{pmatrix} + (ad - bc) \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \\ &= \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix} \end{aligned}$$

Remarquons qu'on pourrait penser que le théorème de Cayley-Hamilton est trivial et être tenté de le prouver en posant simplement $\chi_A(A) = \det(AI_n - A) = \det(A - A) = \det(0) = 0$. Cet argument n'est pas valide car le théorème de Cayley-Hamilton dit que $\chi_A(A) = 0$ en tant que matrice ou en tant qu'endomorphisme, alors que $\det(AI_n - A) = \det(0) = 0$ est une information de type scalaire.

Pour la preuve du théorème de Cayley-Hamilton, nous aurons besoin du lemme suivant, dont nous laissons la preuve en exercice :

Lemme 9.4.2. *Si f est un endomorphisme d'un espace vectoriel V de dimension finie, et si $W \subset V$ est un sous-espace vectoriel invariant par f , alors le polynôme caractéristique de la restriction $f|_W$ de f à W divise le polynôme caractéristique de f .*

Preuve. Nous laissons la preuve de ce lemme en exercice.

Démonstration du théorème de Cayley-Hamilton. Nous allons prouver que pour tout $v \in V$ on a $\chi_f(f)(v) = 0$. Si $v = 0$ il n'y a rien à montrer. On suppose donc que $v \neq 0$ et on note ce vecteur par $v_1 = v$. On considère ensuite le plus grand entier k tel que les vecteurs

$$v_1, v_2 = f(v_1), v_3 = f(v_2) = f^2(v_1), \dots, v_k = f(v_{k-1}) = f^{k-1}(v_1)$$

sont linéairement indépendants.

Notons $\mathcal{B} = \{v_1, \dots, v_k\}$ et $W = \text{Vec}(\mathcal{B}) \subset V$ le sous-espace vectoriel engendré par ces vecteurs. L'ensemble $\mathcal{B}$ est une base de W puisque ces vecteurs sont supposés linéairement indépendants et qu'ils engendrent W .

Par construction, on sait que $f(v_k)$ est combinaison linéaire des éléments de $\mathcal{B}$, il existe donc $\alpha_1, \dots, \alpha_k \in K$ tels que

$$f(v_k) = \alpha_1 v_1 + \dots + \alpha_k v_k,$$

et comme $f(v_j) = v_{j+1}$ pour $j < k$, on conclut que le sous-espace W est invariant par f .

Notons $g = f|_W$ la restriction de f au sous-espace W . Alors g est un endomorphisme de W et sa matrice dans la base $\mathcal{B}$ est donnée par

$$A = M_{\mathcal{B}}(g) = \begin{pmatrix} 0 & 0 & \cdots & 0 & \alpha_1 \\ 1 & 0 & & 0 & \alpha_2 \\ 0 & 1 & \ddots & \vdots & \vdots \\ \vdots & \vdots & \ddots & 0 & \alpha_{k-1} \\ 0 & 0 & \cdots & 1 & \alpha_k \end{pmatrix}.$$

On a vu aux exercices que le polynôme caractéristique de cette matrice est

$$\chi_A(t) = \det(t\mathbf{I}_n - A) = -\alpha_1 - \alpha_2 t - \dots - \alpha_k t^{k-1} + t^k.$$

Par conséquent $\chi_g(f) = \chi_A(f)$ est l'endomorphisme de V défini par

$$\chi_g(f) = -\alpha_1 \text{Id} - \alpha_2 f - \dots - \alpha_k f^{k-1} + f^k.$$

Si on applique cet endomorphisme à v_1 , on trouve

$$\begin{aligned} \chi_g(f)(v_1) &= -\alpha_1 v_1 - \alpha_2 f(v_1) - \alpha_3 f^2(v_1) - \dots - \alpha_k f^{k-1}(v_1) + f^k(v_1) \\ &= -\alpha_1 v_1 - \alpha_2 v_2 - \dots - \alpha_k v_k + f(v_k) \\ &= 0. \end{aligned}$$

Il reste à prouver que $\chi_f(f)(v_1) = 0$. Or g est la restriction de f au sous-espace invariant $W \subset V$. Le lemme précédent implique alors que $\chi_g(t)$ est un facteur de $\chi_f(t)$, i.e. $\chi_f(t) = q(t) \cdot \chi_g(t)$ pour un certain polynôme $q(t) \in K[t]$. Par conséquent

$$\chi_f(f)(v_1) = q(f) \circ \chi_g(f)(v_1) = q(f)(0) = 0.$$

On a ainsi montré que pour tout vecteur $v_1 \in V$ non nul on a $\chi_f(f)(v_1) = 0$. Cela signifie que χ_f est l'endomorphisme nul. □

Corollaire 9.4.3. *Soit f un endomorphisme d'un espace vectoriel de dimension finie V . Alors le polynôme minimal $\mu_f(t)$ divise le polynôme caractéristique $\chi_f(t)$. De plus ces deux polynômes ont exactement les mêmes racines (qui sont les valeurs propres de f).*

Démonstration. On sait par le corollaire 9.3.4 que le polynôme minimal divise tout polynôme annulateur de f . En particulier $\mu_f(t)$ divise $\chi_f(t)$ puisque $\chi_f(f) = 0$ par le théorème de Cayley-Hamilton.

Cela signifie que $\chi_f(t)$ est un multiple de $\mu_f(t)$, par conséquent toute racine de $\mu_f(t)$ est aussi une racine de $\chi_f(t)$ (car $\mu_f(\lambda) = 0 \Rightarrow \chi_f(\lambda) = 0$).

Pour prouver le sens inverse, on suppose que $\lambda \in K$ est une racine de $\chi_f(t)$, c'est donc une valeur propre de f . On a vu que pour tout polynôme $p(t)$, si λ est une valeur propre de f , alors $p(\lambda)$ est valeur propre de $p(f)$ (théorème 9.2.1). En particulier $\mu_f(\lambda)$ est une valeur propre de $\mu_f(f)$, donc $\mu_f(\lambda) = 0$ car $\mu_f(f) = 0 \in V$.

□

Une conséquence de ce corollaire est que si $\chi_f(t)$ est scindé, alors $\mu_f(t)$ est aussi scindé. Plus précisément, si $\chi_f(t) = \prod_{i=1}^r (t - \lambda_i)^{m_i}$, avec $\lambda_1, \dots, \lambda_r$ distincts, alors le polynôme minimal est du type $\mu_f(t) = \prod_{i=1}^r (t - \lambda_i)^{k_i}$, où les exposants $k_i \in \mathbb{N}$ vérifient $1 \leq k_i \leq m_i$ pour tout i .

Ceci nous conduit à une *méthode effective* pour trouver le polynôme minimal d'un endomorphisme f ou d'une matrice A , dont le polynôme caractéristique est scindé :

- (1) On calcule le polynôme caractéristique $\chi_f(t)$ et on le factorise.
- (2) Si ce polynôme est scindé, il s'écrit $\chi_f(t) = \prod_{i=1}^r (t - \lambda_i)^{m_i}$, où $\sigma(f) = \{\lambda_1, \dots, \lambda_r\}$ est l'ensemble des valeurs propres. Noter que $1 \leq r \leq n = \dim(V)$.
- (3) On considère tous les polynômes de type $p(t) = \prod_{i=1}^r (t - \lambda_i)^{s_i}$ avec $1 \leq s_i \leq m_i$ en commençant par $s_j = 1$, et on vérifie si $p(f) = 0$.
- (4) Le polynôme $p(t)$ de l'étape précédente dont le degré est minimal est le polynôme minimal $\mu_f(t)$.

Exemple 9.4.4. On considère la matrice

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

Le polynôme caractéristique de cette matrice est $\chi(t) = (t - 3)(t - 2)^3$. On constate que ce polynôme est scindé. Le polynôme minimal est donc l'un des polynômes suivants :

$$p_1(t) = (t - 3)(t - 2), \quad p_2(t) = (t - 3)(t - 2)^2, \quad \text{ou} \quad p_3(t) = \chi(t) = (t - 3)(t - 2)^3.$$

Pour décider lequel de ces polynômes est le polynôme minimal, on calcule $p_1(A)$ et $p_2(A)$ (on sait déjà par Cayley-Hamilton que $p_3(A) = 0$). Le calcul nous donne :

$$p_1(A) = \begin{pmatrix} 0 & 0 & 1 & 0 \\ 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 \end{pmatrix} \quad \text{et} \quad p_2(A) = \begin{pmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix}$$

Par conséquent le polynôme minimal est $\mu_A(t) = p_2(t) = (t - 3)(t - 2)^2$.

Il sera commode de considérer, en plus des polynômes caractéristique et minimal, un troisième polynôme :

Définition 9.4.5. Le *polynôme spectral*² d'un endomorphisme f de l'espace vectoriel V de dimension finie est défini par $\nu_f(t) = 1$ si f n'a aucune valeur propre et

$$\nu_f(t) = \prod_{\lambda \in \sigma(f)} (t - \lambda) = \prod_{i=1}^r (t - \lambda_i)$$

si le spectre $\sigma(f) = \{\lambda_1, \dots, \lambda_r\}$ de f est non vide.

Par le corollaire 9.4.3, on sait que le polynôme spectral de f divise le polynôme minimal et le polynôme minimal divise le polynôme caractéristique, ce qu'on peut noter par

$$\nu_f(t) \mid \mu_f(t) \mid \chi_f(t),$$

de plus ces trois polynômes ont les mêmes racines, qui sont les valeurs propres de F . Dans l'exemple précédent, on a

$$\nu_A(t) = (t - 3)(t - 2), \quad \mu_A(t) = (t - 3)(t - 2)^2 \quad \text{et} \quad \chi_A(t) = (t - 3)(t - 2)^3.$$

9.5 Une autre preuve du théorème de Cayley-Hamilton

Dans ce paragraphe, on propose une autre preuve du théorème de Cayley-Hamilton. Cette preuve se place dans le cadre du calcul matriciel et utilise la notion de *polynôme matriciel*.

Définition 9.5.1. Un *polynôme matriciel* de taille n sur un corps K , est une expression formelle

$$P(t) = A_0 + A_1 t + \dots + A_k t^k,$$

où $A_j \in M_n(K)$ pour tout $j \in \{0, \dots, k\}$. Le symbole t s'appelle l'*indéterminée* du polynôme et on note $M_n(K)[t]$ l'ensemble de ces polynômes matriciels.

Étant donné $P(t) \in M_n(K)[t]$, on peut ou bien substituer un scalaire $x \in K$ à l'indéterminée t , ou bien une matrice $X \in M_n(K)$. Dans les deux cas on obtient une matrice $P(x)$ ou $P(X)$. Par exemple si

$$P(t) = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} + \begin{pmatrix} 0 & -1 \\ 0 & 0 \end{pmatrix} t + \begin{pmatrix} 0 & 0 \\ 0 & 5 \end{pmatrix} t^2,$$

alors

$$P(2) = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} + 2 \begin{pmatrix} 0 & -1 \\ 0 & 0 \end{pmatrix} + 4 \begin{pmatrix} 0 & 0 \\ 0 & 5 \end{pmatrix} = \begin{pmatrix} 1 & -2 \\ 0 & 20 \end{pmatrix},$$

et

$$\begin{aligned} P \begin{pmatrix} 2 & 0 \\ 1 & 0 \end{pmatrix} &= \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} + \begin{pmatrix} 0 & -1 \\ 0 & 0 \end{pmatrix} \begin{pmatrix} 2 & 0 \\ 1 & 0 \end{pmatrix} + \begin{pmatrix} 0 & 0 \\ 0 & 5 \end{pmatrix} \begin{pmatrix} 1 & -2 \\ 0 & 20 \end{pmatrix}^2 \\ &= \begin{pmatrix} 0 & 0 \\ 10 & 0 \end{pmatrix} \end{aligned}$$

2. Ce polynôme ne semble pas avoir de nom particulier dans la littérature, on peut aussi l'appeler le *polynôme des valeurs propres*.

Si $P(t) = A_0 + A_1t + \cdots A_k t^k$ et $Q(t) = B_0 + B_1t + \cdots B_m t^m$, sont deux polynômes matriciels, leur produit est défini par la formule usuelle :

$$(P \cdot Q)(t) = \sum_{r=0}^{k+m} \left(\sum_{i=0}^r A_i B_{r-i} \right) t^r,$$

où on considère que $A_i = 0$ si $i > k$ et $B_j = 0$ si $j > m$. En travaillant avec des polynômes matriciels, il y a lieu de prendre certaines précautions. En particulier, si $P(t), Q(t) \in M_n(K)[t]$ et $X \in M_n(K)$, alors, généralement on a

$$(P \cdot Q)(X) \neq P(X) \cdot Q(X).$$

Par exemple si

$$P(t) = A_0 + A_1t \text{ et } Q(t) = B_0 + B_1t,$$

alors

$$(P \cdot Q)(t) = A_0B_0 + (A_0B_1 + A_1B_0)t + A_1B_1t^2.$$

On voit donc que pour tout $X \in M_n(K)$ on a

$$(P \cdot Q)(X) = A_0B_0 + (A_0B_1 + A_1B_0)X + A_1B_1X^2,$$

et

$$P(X)Q(X) = A_0B_0 + A_0B_1X + A_1XB_0 + A_1XB_1X.$$

Si X ne commute pas avec B_1 ou B_0 , alors, généralement on aura $(P \cdot Q)(X) \neq P(X) \cdot Q(X)$.

Lemme 9.5.2. Si $P(t) = A_0 + A_1t + \cdots A_k t^k$ et $Q(t) = B_0 + B_1t + \cdots B_m t^m$, sont deux polynômes matriciels, et $X \in M_n(K)$ est une matrice qui commute avec chaque B_j , alors, la relation

$$(P \cdot Q)(X) = P(X) \cdot Q(X)$$

est vérifiée.

Preuve. En utilisant que X commute avec chaque B_j , on a

$$\begin{aligned} P(X)Q(X) &= \left(\sum_{i=0}^k A_i X^i \right) \cdot \left(\sum_{j=0}^m B_j X^j \right) = \sum_{i=0}^k \sum_{j=0}^m A_i X^i B_j X^j \\ &= \sum_{i=0}^k \sum_{j=0}^m A_i B_j X^{i+j} \\ &= \sum_{r=0}^{k+m} \left(\sum_{i=0}^r A_i B_{r-i} \right) X^r \\ &= (P \cdot Q)(X). \end{aligned}$$

□

Preuve alternative du théorème de Cayley-Hamilton.

La preuve de la formule de Laplace du §7.5 (chapitre 7 du polycopié 1) s'applique non seulement à une matrice à coefficients dans un corps K mais aussi, et sans changement, à une matrice à coefficients dans l'anneau des polynômes $K[t]$ (ou dans un autre anneau commutatif quelconque). Donc pour tout polynôme matriciel $Q(t) \in M_n(K[t])$ on a

$$\det(Q(t)) \cdot I_n = \text{Cof}(Q(t))^\top \cdot Q(t).$$

On applique ce qui précède au polynôme matriciel $Q(t) = (tI_n - A)$, et on note

$$P(t) = \text{Cof}(tI_n - A)^\top = C_0 + C_1t + \cdots + C_{n-1}t^{n-1},$$

où les C_j sont des matrices de taille $n \times n$. On a donc l'identité

$$\chi_A(t) \cdot I_n = P(t) \cdot (tI_n - A),$$

et on sait par le lemme précédent que dans une telle égalité on peut substituer à t toute matrice $X \in M_n(K)$ qui commute avec A . On a donc pour une telle matrice

$$\chi_A(X) = P(X) \cdot (X - A).$$

Or il est trivial que A commute avec A et on a donc prouvé que

$$\chi_A(A) = P(A) \cdot (A - A) = 0 \in M_n(K).$$

□

9.6 Vecteurs propres généralisés et théorème de réduction primaire

Les notions suivantes joueront un rôle central dans la suite de ce chapitre :

Définition 9.6.1. Soit $f \in \mathcal{L}(V)$ un endomorphisme d'un K -espace vectoriel et $\lambda \in K$.

- (i) On dit qu'un vecteur $v \in V$ est un *vecteur propre généralisé* de f associé à λ si $v \neq 0$ et s'il existe un entier $m \in \mathbb{N}$ tel que $v \in \text{Ker}((\lambda \text{Id}_V - f)^m)$, i.e.

$$(\lambda \text{Id}_V - f)^m v = 0.$$

- (ii) Le plus petit entier m tel que $^3 (f - \lambda)^m v = 0$ s'appelle l'*ordre* du vecteur propre généralisé.

- (iii) Pour tout $k \in \mathbb{N}$, l'entier

$$\delta_{f,\lambda}(k) = \dim \left(\text{Ker}(f - \lambda)^k \right)$$

s'appelle la *multiplicité généralisée d'ordre m de λ pour f* (lorsque $k = 0$, on convient que $\delta_{f,\lambda}(k) = 0$).

Si $A \in M_n(K)$, on notera de même $\delta_{A,\lambda}(k) = \dim \left(\text{Ker}(A - \lambda \text{Id}_n)^k \right)$. Lorsque l'endomorphisme f (ou la matrice A) été fixé, on notera simplement $\delta_\lambda(k)$.

3. Dans la suite, on s'autorisera pour simplifier à noter l'endomorphisme $(f - \lambda \text{Id}_V)$ par $(f - \lambda)$.

Exemples.

- (i) Tout vecteur propre est un vecteur propre généralisé d'ordre 1.
- (ii) Si f est nilpotent, i.e. s'il existe m tel que f^m est nul, alors tout élément non nul de V est un vecteur propre généralisé (associé à la valeur propre $\lambda = 0$).
- (iii) Le vecteur $\begin{pmatrix} 0 \\ 1 \end{pmatrix}$ n'est pas un vecteur propre de la matrice $\begin{pmatrix} \alpha & 1 \\ 0 & \alpha \end{pmatrix}$, mais c'est un vecteur propre généralisé associé à la valeur propre α .
- (iv) La fonction $\varphi \in C^\infty(\mathbb{R})$ définie par $\varphi(x) = x^{m-1} \cdot e^{\lambda x}$ est un vecteur propre généralisé de l'opérateur $\frac{d}{dx}$ car

$$\left(\frac{d}{dx} - \lambda\right)^m \left(x^{m-1} \cdot e^{\lambda x}\right) = 0.$$

Nous laissons la vérification de ces exemples en exercice.

Lemme 9.6.2. *S'il existe un vecteur propre généralisé pour $f \in \mathcal{L}(V)$ associé à $\lambda \in K$, alors λ est une valeur propre de f .*

Ce lemme nous dit que s'il existe une notion de vecteur propre généralisé, il n'y a pas de notion de *valeur propre généralisée*, qui serait différente des valeurs propres usuelles.

Preuve. Si v est un vecteur propre associé à λ il n'y a rien à montrer. Sinon, il existe $m \geq 2$ tel que $(f - \lambda)^m v = 0$. Supposons m minimal avec cette propriété et posons $w = (f - \lambda)^{m-1} v$. Alors $w \neq 0$ par hypothèse et $(f - \lambda)w = (f - \lambda)^m v = 0$. Donc w est vecteur propre et la valeur propre associée est λ .

□

Remarque 9.6.3. On montre facilement que deux endomorphismes conjugués (ou deux matrices semblables) ont les mêmes multiplicités généralisées pour chaque valeur propre. De plus, si A est la matrice de f dans une base quelconque, alors $\delta_{f,\lambda}(m) = \delta_{A,\lambda}(m)$ pour toute valeur propre λ et tout entier m .

Les multiplicités généralisées d'une matrice $A \in M_n(K)$ peuvent se calculer au moyen de la formule du rang :

$$\delta_{A,\lambda}(m) = n - \text{rang}(A - \lambda)^m.$$

Rappelons que le rang d'une matrice est le nombre maximal de colonne (ou de lignes) qui sont linéairement indépendantes. Il peut se calculer avec la méthode de Gauss-Jordan.

Le théorème suivant est d'une importance majeure, il nous dit en particulier que si f est un endomorphisme d'un espace vectoriel de dimension finie qui admet un polynôme annulateur scindé, alors tout vecteur de V se décompose de façon unique comme somme de vecteurs propres généralisés.

Théorème 9.6.4 (Théorème de réduction primaire). *Soit f un endomorphisme du K -espace vectoriel de dimension finie V . Considérons le polynôme*

$$p(t) = \prod_{i=1}^r (t - \lambda_i)^{s_i},$$

où $\{\lambda_1, \dots, \lambda_r\} = \sigma(f)$ est l'ensemble des valeurs propres et $s_i \in \mathbb{N}$.

Alors on a les propriétés suivante :

- (i) *Le sous-espace $U_i = \text{Ker}(\lambda_i - f)^{s_i}$ est invariant par f pour tout $i \in \{1, \dots, r\}$.*
- (ii) *La restriction de $(\lambda_i - f)$ à U_i est un endomorphisme nilpotent.*
- (iii) *Le noyau de $p(f)$ est somme directe des U_i :*

$$\text{Ker}(p(f)) = U_1 \oplus \dots \oplus U_r, \quad (9.4)$$

et sa dimension est

$$\dim \text{Ker}(p(f)) = \sum_{i=1}^r \delta_{f, \lambda_i}(s_i). \quad (9.5)$$

Une première conséquence intéressante de ce théorème est le résultat suivant :

Corollaire 9.6.5. *Soit $f \in \mathcal{L}(V)$ un endomorphisme d'un K -espace vectoriel de dimension n . Notons $\sigma(f) = \{\lambda_1, \dots, \lambda_r\}$ l'ensemble de ses valeurs propres et $\nu_f(t) = \prod_{\lambda \in \sigma(f)} (t - \lambda) \in K[t]$ son polynôme spectral. Alors le noyau de $\nu_f(f)$ est le sous-espace vectoriel de V engendré par les vecteurs propres de f . Plus précisément, on a*

$$\text{Ker}(\nu_f(f)) = E_{\lambda_1}(f) \oplus \dots \oplus E_{\lambda_r}(f). \quad (9.6)$$

où $E_{\lambda_i}(f)$ est l'espace propre associé à la valeur propre λ_i .

Preuve. C'est une application directe du théorème 9.6.4. □

Corollaire 9.6.6. *L'endomorphisme f est diagonalisable si et seulement si le polynôme spectral $\nu_f(t)$ annule f .*

Preuve. C'est une conséquence immédiate de la proposition précédente, puisqu'un endomorphisme f d'un espace vectoriel V est diagonalisable si et seulement si V est somme directe des espaces propres de f . □

Remarque. Le corollaire précédent implique que pour un endomorphisme f d'un espace vectoriel V de dimension finie, les conditions suivantes sont équivalentes :

- (i) f est diagonalisable.
- (ii) $\nu_f(t)$ est un polynôme annulateur de f .
- (iii) Le polynôme minimal coïncide avec le polynôme spectral : $\mu_f(t) = \nu_f(t)$.
- (iv) Le polynôme minimal $\mu_f(t)$ est scindé et toutes ses racines sont simples.
- (v) Il existe un polynôme scindé à racines simples qui annule f .

En particulier, si $\mu_f(t)$ admet une racine multiple, alors f n'est pas diagonalisable.

Nous laissons la preuve de cette remarque en exercice.

Exemples.

1. Le polynôme caractéristique de la matrice

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

est $\chi_A(t) = t^2 + 1$. Ce polynôme n'a pas de racines réelles, donc A n'a aucun vecteur propre dans $\mathbb{R}^2$ et n'est donc pas diagonalisable dans $M_2(\mathbb{R})$. Par contre si on regarde A comme matrice complexe, alors $\chi_A(t) = (t+i)(t-i)$ et A est donc diagonalisable dans $M_2(\mathbb{C})$.

2. Le polynôme caractéristique de la matrice

$$B = \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$$

est $\chi_B(t) = (t-1)^2$. Donc le polynôme minimal est ou bien $(t-1)$ ou bien $(t-1)^2$. Or on vérifie immédiatement que $(t-1)$ n'annule pas la matrice B , par conséquent $\mu_B(t) = (t-1)^2$ possède une racine double et B n'est pas diagonalisable dans $M_2(\mathbb{C})$.

3. Le polynôme minimal de la matrice

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

est $(t-3)(t-2)^3$. (cf. exemple 9.4.4). Cette matrice n'est donc pas diagonalisable.

Pour énoncer la seconde conséquence importante du théorème de décomposition primaire, on introduit la notion de sous-espace caractéristique par rapport à un endomorphisme.

Définition 9.6.7. Soit f un endomorphisme d'un espace vectoriel de dimension finie V et $\lambda \in \sigma(f)$ une valeur propre de f de multiplicité algébrique $m_\lambda = \text{multalg}_\lambda(f)$ (rappelons qu'il s'agit du plus grand entier m tel que $(t-\lambda)^m$ divise le polynôme caractéristique $\chi_f(t)$). Le sous-espace vectoriel

$$N_\lambda(f) = \text{Ker}(f - \lambda)^{m_\lambda} \subset V$$

s'appelle le *sous-espace caractéristique*, ou *sous-espace propre généralisé* associé à la valeur propre λ .

Nous avons alors le corollaire suivant du théorème de réduction primaire

Corollaire 9.6.8. Soit f un endomorphisme d'un espace vectoriel de dimension finie V dont le polynôme caractéristique est scindé, alors

(i) On a la décomposition suivante de V en somme directe :

$$V = N_{\lambda_1}(f) \oplus \cdots \oplus N_{\lambda_r}(f), \quad (9.7)$$

cela signifie que tout vecteur $v \in V$ peut s'écrire de façon unique comme somme de vecteurs propres généralisés.

- (ii) Si $\lambda \in \sigma(f)$ une valeur propre de f , alors l'ensemble des vecteurs propres généralisés associés à λ est l'ensemble des vecteurs non nuls de $N_\lambda(f)$.

Preuve. Dans le cas où le polynôme caractéristique de f est scindé, le théorème de Cayley-Hamilton nous dit alors que $\text{Ker}(\chi_f(f)) = V$. On peut donc appliquer le théorème de réduction primaire 9.6.4 au polynôme $p(t) = \chi_f(t)$ et on conclut que V est somme directe des sous-espaces propres généralisés, ce qui prouve la première affirmation. La seconde affirmation est une conséquence immédiate de la première. $\square$

La seconde affirmation de ce corollaire peut se reformuler en disant que si $m > m_\lambda = \text{multalg}_f(\lambda)$, alors $\text{Ker}(f - \lambda \text{Id}_V)^m = \text{Ker}(f - \lambda)^{m_\lambda}$, de façon équivalente, l'ordre de tout vecteur propre généralisé est au plus égal à la multiplicité algébrique de la valeur propre associée.

9.7 Complément sur les multiplicités

Nous démontrons ici deux résultats complémentaires : le premier concerne les multiplicités algébriques d'un endomorphisme :

Proposition 9.7.1. *Soit f un endomorphisme d'un espace vectoriel de dimension finie. Supposons que le polynôme caractéristique de f est scindé, alors*

- (a) *La somme des multiplicités algébriques de toutes les valeurs propres est égale à la dimension de l'espace vectoriel V :*

$$\sum_{\lambda \in \sigma(f)} \text{multalg}_\lambda(f) = \dim(V)$$

- (b) *La multiplicité algébrique de toute valeur propre $\lambda \in \sigma(f)$ est égale à la dimension de l'espace caractéristique associé :*

$$\text{multalg}_\lambda(f) = \dim(N_\lambda(f)).$$

- (c) *Pour tout $\lambda \in \sigma(f)$ on a*

$$1 \leq \text{multgeom}_\lambda(f) \leq \text{multalg}_\lambda(f) \leq n.$$

Rappelons que le sous-espace caractéristique de f associé à la valeur propre λ est le noyau $N_\lambda(f) = \text{Ker}(f - \lambda)^{m_\lambda}$, où $m_\lambda = \text{multalg}_\lambda(f)$.

Preuve. a.) Le polynôme caractéristique de f est scindé, il s'écrit donc $\chi_f(t) = \prod_{i=1}^r (t - \lambda_i)^{m_i}$, où on a noté $m_i = \text{multalg}_{\lambda_i}(f)$ et on a

$$\sum_{i=1}^r m_i = \deg(\chi_f(t)) = \dim(V).$$

b.) Notons $N_i = N_{\lambda_i}(f)$ le $i^{\text{ème}}$ sous-espace caractéristique de f et $m_i = \text{multalg}_{\lambda_i}(f)$. Le sous-espace N_i est invariant par f , donc la restriction de f à N_i définit un endomorphisme $f_i \in \mathcal{L}(N_i)$. Le théorème de réduction primaire 9.6.4 avec le théorème de Cayley-Hamilton nous dit que les sous-espaces N_i sont invariants et $V = N_1 \oplus \cdots \oplus N_r$. Le polynôme caractéristique admet donc la factorisation suivante :

$$\chi_f(t) = \chi_{f_1}(t) \cdots \chi_{f_r}(t).$$

Le lemme 9.6.2 entraîne que $f_i \in \mathcal{L}(N_i)$ n'a qu'une valeur propre qui est λ_i , car tout vecteur non nul de N_i est un vecteur propre généralisé associé à λ_i donc le polynôme $\chi_{f_i}(t)$ est du type $(t - \lambda_i)^{k_i}$ pour tout $i = 1, \dots, r$. On a donc

$$\prod_{i=1}^r (t - \lambda_i)^{k_i} = \chi_f(t) = \prod_{i=1}^r (t - \lambda_i)^{m_i},$$

et par unicité de la décomposition d'un polynôme en facteurs irréductibles, on en déduit que

$$m_i = k_i = \deg(\chi_{f_i}(t)) = \dim(N_i).$$

pour tout i .

c.) Rappelons que par définition $\text{multgeom}_\lambda(f) = \dim(E_\lambda(f))$. Pour toute valeur propre λ de f , on a

$$\{0\} \neq E_\lambda(f) = \text{Ker}(f - \lambda) \subset \text{Ker}(f - \lambda)^{m_\lambda} \subset V,$$

d'où les inégalités voulues. □

Le second résultat concerne les multiplicités généralisées :

Proposition 9.7.2. *Si le polynôme caractéristique de $f \in \mathcal{L}(V)$ est scindé et $\lambda \in \sigma(f)$, alors*

(i) *Les multiplicités généralisées associées à chaque valeur propre forment une suite monotone :*

$$\delta_{f,\lambda}(k) \leq \delta_{f,\lambda}(k+1) \text{ pour tout } k \in \mathbb{N}.$$

(ii) $\delta_{f,\lambda}(1) = \text{multgeom}_\lambda(f)$.

(iii) $\delta_{f,\lambda}(k) = \text{multalg}_\lambda(f)$, pour tout $k \geq \text{multalg}_\lambda(f)$.

Preuve. Les deux premières propriétés sont immédiates à partir de la définition $\delta_{f,\lambda}(k) = \dim \text{Ker}(f - \lambda)^{m_k}$. La troisième propriété se déduit du point (ii) du Corollaire 9.6.8. □

Cette proposition implique en particulier qu'il existe $m \leq \text{multalg}_\lambda(f)$ tel que

$$\text{multgeom}_\lambda(f) = \delta_\lambda(1) \leq \delta_\lambda(2) \leq \dots \leq \delta_\lambda(m) = \delta_\lambda(m+1) = \text{multalg}_\lambda(f).$$

Le plus petit m ayant cette propriété est l'ordre maximal d'un vecteur propre associé à λ .

9.8 Lemme des noyaux et preuve du théorème de réduction primaire

Dans cette section, nous démontrons le théorème de décomposition primaire. La preuve repose sur le résultat suivant, qui en est une généralisation, et qui s'appelle le *lemme des noyaux*.

Théorème 9.8.1 (Lemme des noyaux). *Soit V un K -espace vectoriel, $f \in \mathcal{L}(V)$ un endomorphisme de V et $p(t) \in K[t]$ un polynôme. Supposons que $p(t)$ se factorise sous la forme d'un produit*

$$p(t) = q_1(t)q_2(t) \cdots q_r(t)$$

où les polynômes $q_i(t)$ sont deux-à-deux premiers entre eux. Notons $W_i = \text{Ker}(q_i(f))$, pour $i = 1, \dots, r$. Alors les W_i sont invariants par f et le noyau de $p(f)$ se décompose comme somme directe

$$\text{Ker}(p(f)) = W_1 \oplus \cdots \oplus W_r.$$

Rappelons que les polynômes $q_i(t), q_j(t), \dots, q_r(t)$ sont *premiers entre eux* s'ils n'admettent pas de facteur commun non constant.

Démonstration. La preuve se fait par récurrence sur r . Pour $r = 1$ il n'y a rien à démontrer, démontrons le théorème pour $r = 2$, i.e. $p(t) = q_1(t)q_2(t)$. Comme q_1 et q_2 sont supposés premiers entre eux, l'identité de Bézout (cf. Appendice A du polycopié du premier semestre) nous dit qu'il existe $s_1(t), s_2(t) \in K[t]$ tels que

$$q_1(t)s_1(t) + q_2(t)s_2(t) = 1.$$

On a aussi en substituant f à l'indéterminée la relation

$$q_1(f)s_1(f) + q_2(f)s_2(f) = \text{Id}_V. \quad (9.8)$$

- On montre d'abord que la somme est directe. Soit donc $v \in \text{Ker}(q_1(f)) \cap \text{Ker}(q_2(f))$. On veut montrer que $v = 0$. En utilisant la relation (9.8), on peut écrire

$$v = (\text{Id}_V)(v) = (q_1(f)s_1(f) + q_2(f)s_2(f))(v).$$

Or ce vecteur est nul car $q_1(f)s_1(f)(v) = s_1(f)q_1(f)(v) = s_1(f)(0) = 0$ puisque on a supposé $v \in \text{Ker}(q_1(f))$. De même $q_2(f)s_2(f)(v) = 0$. Remarquons qu'on a utilisé (et qu'on réutilisera dans la suite) que deux polynômes en f commutent.

- On prouve maintenant que $\text{Ker}(p(f)) = \text{Ker}(q_1(f)) + \text{Ker}(q_2(f))$. Il s'agit de montrer deux inclusions.
 - Supposons d'abord $v \in \text{Ker}(q_1(f)) + \text{Ker}(q_2(f))$. Cela signifie que v s'écrit $v = v_1 + v_2$ où $v_i \in \text{Ker}(q_i(f))$ (pour $i = 1, 2$). Montrons que dans ce cas on a $v \in \text{Ker}(p(f))$:

$$\begin{aligned} p(f)(v) &= q_1(f)q_2(f)(v) \\ &= q_1(f)q_2(f)(v_1 + v_2) \\ &= q_2(f)q_1(f)(v_1) + q_1(f)q_2(f)(v_2) \\ &= 0. \end{aligned}$$

- Supposons maintenant que $v \in \text{Ker}(p(f))$, i.e. $p(f)(v) = 0$. A l'aide de l'expression (9.8), on peut écrire v sous la forme

$$v = (q_1(f)s_1(f) + q_2(f)s_2(f))(v).$$

Notons alors $v_1 = q_2(f)s_2(f)(v)$ et $v_2 = q_1(f)s_1(f)(v)$ et montrons que cette écriture permet de voir v comme élément de $\text{Ker}(q_1(f)) + \text{Ker}(q_2(f))$. En effet

$$q_1(f)(v_1) = q_1(f)q_2(f)s_2(f)(v) = s_2(f)p(f)(v) = s_2(f)(0) = 0.$$

On montre de même que $q_2(f)(v_2) = 0$.

Pour conclure la preuve par récurrence on se ramène au cas $r = 2$ en écrivant

$$p(t) = q_1(t)\varphi(t)$$

avec $\varphi(t) = q_2(t) \cdots q_r(t)$. □

La démonstration du théorème de décomposition primaire est maintenant très courte. Rappelons d'abord l'énoncé :

Théorème. Soit f un endomorphisme du K -espace vectoriel de dimension finie V et $p(t)$ le polynôme

$$p(t) = \prod_{i=1}^r (t - \lambda_i)^{s_i},$$

où $\{\lambda_1, \dots, \lambda_r\} = \sigma(f)$ est l'ensemble des valeurs propres et $s_i \in \mathbb{N}$. Alors on a les propriétés suivantes :

- (i) Le sous-espace $U_i = \text{Ker}(\lambda_i - f)^{s_i}$ est invariant par f .
- (ii) La restriction de $(f - \lambda_i)$ à U_i est un endomorphisme nilpotent.
- (iii) Le noyau de $p(f)$ est somme directe des U_i :

$$\text{Ker}(p(f)) = U_1 \oplus \dots \oplus U_r,$$

et sa dimension est

$$\dim(\text{Ker}(p(f))) = \sum_{i=1}^r \delta_{f, \lambda_i}(s_i).$$

Démonstration. Observons d'abord que $x \in U_i$ si et seulement si $(\lambda_i - f)^{s_i}(x) = 0$, donc

$$(\lambda_i - f)^{s_i}(f(x)) = f((\lambda_i - f)^{s_i}(x)) = 0,$$

ce qui signifie que $f(x) \in U_i$ et prouve l'affirmation (i).

La preuve de (ii) est évidente puisque la restriction de $(\lambda_i - f)^{s_i}$ à $U_i = \text{Ker}(\lambda_i - f)^{s_i}$ est nulle.

Pour prouver (iii), on remarque que les polynômes $(t - \lambda_i)^{s_i}$ sont premiers entre eux car on suppose que $\lambda_i \neq \lambda_j$ pour $i \neq j$. Un argument par récurrence basé sur le lemme de noyaux entraîne alors immédiatement que

$$\text{Ker}(p(f)) = \text{Ker}(\lambda - f)^{s_1} \oplus \dots \oplus \text{Ker}(\lambda - f)^{s_r},$$

La dernière équation se déduit maintenant du fait que, par définition, $\dim(U_i) = \delta_{f, \lambda_i}(s_i)$. □

9.9 Décomposition de Dunford

Dans ce paragraphe, nous allons prouver que tout endomorphisme d'un espace vectoriel complexe de dimension finie est somme d'un endomorphisme diagonalisable et d'un endomorphisme nilpotent. Commençons par un résultat particulièrement simple :

Lemme 9.9.1. Soit $f \in \mathcal{L}(V)$ un endomorphisme d'un espace vectoriel V de dimension finie dont le polynôme caractéristique est scindé. Supposons que f n'admet qu'une valeur propre λ , alors $(f - \lambda \text{Id}_V)$ est nilpotent.

Preuve. Les hypothèses entraînent que $\chi_f(t) = (t - \lambda)^n$ où $n = \dim(V)$. Par le théorème de Cayley-Hamilton, on a alors $\chi_f(f) = (f - \lambda \text{Id}_V)^n = 0$, ce qui signifie précisément que $(f - \lambda \text{Id}_V)$ est nilpotent. □

Théorème 9.9.2. *Toute matrice $A \in M_n(K)$ dont le polynôme caractéristique est scindé peut s'écrire sous la forme $A = D + N$, où D est une matrice diagonalisable et N est une matrice nilpotente qui commute avec D , i.e. $DN = ND$.*

On peut montrer que cette décomposition est unique, on l'appelle la *décomposition de Dunford* de A . Cette décomposition s'appelle aussi la décomposition de *Jordan-Chevalley*.

Preuve. Soit $A \in M_n(K)$ et supposons que le polynôme caractéristique $\chi_A(t)$ est scindé, avec racines $\lambda_1, \dots, \lambda_r$. Notons $N_i = N_{\lambda_i}(A) = \text{Ker}(A - \lambda_i I_n)^{m_i}$ le sous-espace caractéristique associé à la valeur propre λ_i (où m_i est la multiplicité algébrique de λ_i). Le théorème de réduction primaire implique que N_i est invariant par A et

$$K^n = N_1 \oplus \dots \oplus N_r.$$

Choisissons une base de K^n dont les m_1 premiers vecteurs forment une base de N_1 , puis les m_2 vecteurs suivants forment une base de N_2 etc. Alors la matrice de l'opérateur A prend la forme par blocs suivante dans cette base :

$$B = \begin{pmatrix} B_1 & 0 & \dots & 0 \\ 0 & B_2 & \dots & 0 \\ \vdots & \ddots & \ddots & \vdots \\ 0 & \dots & 0 & B_r \end{pmatrix},$$

où le bloc B_i est une matrice de taille $m_i \times m_i$. Plus précisément on a $B = P^{-1}AP$ où P est la matrice de passage de la base canonique dans la nouvelle base.

En utilisant le théorème de triangulation, on peut au moyen d'un changement de base supplémentaire se ramener au cas où chaque bloc B_i est une matrice triangulaire supérieure dont les coefficients diagonaux sont tous égaux à la valeur propre λ_i .

Chaque sous-matrice B_i peut alors s'écrire $B_i = \lambda_i I_{m_i} + T_i$ où T_i est une matrice strictement triangulaire (i.e. avec 0 sur la diagonale). La matrice T_i est nilpotente et commute avec $\lambda_i I_{m_i}$, ce qui complète la preuve du théorème. □

Remarques. 1. La décomposition de Dunford permet de calculer les puissances de toute matrice carrée à coefficient complexe, car le polynôme caractéristique d'une telle matrice est scindé et on a

$$A^m = (D + N)^m = \sum_{j=0}^m \binom{m}{j} D^j N^{m-j}.$$

Les puissances de la matrice diagonale D et de la matrice nilpotente N sont évidemment faciles à calculer. Notons que le développement binomial ci-dessus pour $(D + N)^m$ est valide car les deux matrices commutent.

2. On peut démontrer que la décomposition de Dunford d'une matrice est unique. Dans le paragraphe suivant on est effective (calculable). De plus D et N s'obtiennent comme polynômes de A .

9.10 Sous-espaces cycliques d'un endomorphisme

La structure d'un endomorphisme d'un espace vectoriel de dimension finie dont le polynôme caractéristique est scindé est décrite d'une manière très complète par sa *forme normale de Jordan*⁴. Nous abordons ce thème par la proposition suivante, qui jouera un rôle fondamental dans la suite.

Proposition 9.10.1. *Soit $f \in \mathcal{L}(V)$ un endomorphisme d'un espace vectoriel sur le corps K , et $u \in V$ un vecteur propre généralisé d'ordre m de f pour la valeur propre λ . Alors les vecteurs $u_1, \dots, u_m \in V$ définis par $u_j = (f - \lambda)^{m-j}(u)$, i.e.*

$$u_1 = (f - \lambda)^{m-1}(u), \quad u_2 = (f - \lambda)^{m-2}(u), \quad \dots, \quad u_{m-1} = (f - \lambda)(u), \quad u_m = u, \quad (9.9)$$

sont linéairement indépendants. De plus, le sous-espace vectoriel $U \subset V$ engendré par $\{u_1, \dots, u_m\}$ est invariant par f .

Définition 9.10.2. Un sous-espace vectoriel U de l'espace vectoriel V est dit *cyclique* pour l'endomorphisme $f \in \mathcal{L}(V)$ s'il contient un vecteur propre généralisé u d'ordre $m = \dim(U)$. Dans ce cas, les vecteurs $\{u_1, \dots, u_m\} \subset U$ définis par (9.9) forment la *base cyclique* de U associée au vecteur $u = u_m$. On dit aussi que $u = u_m$ est une *racine* (ou un *générateur*) du cycle.

Remarque 9.10.3. Les vecteurs définis par (9.9) vérifient $(f - \lambda)(u_1) = (f - \lambda)^m(u) = 0$ et $(f - \lambda)(u_j) = u_{j-1}$ pour $j \geq 2$. En appliquant $(f - \lambda)^k$ à ces vecteurs, on trouve inductivement que

$$(f - \lambda)^k(u_j) = \begin{cases} u_{j-k}, & \text{si } j > k, \\ 0, & \text{si } j \leq k. \end{cases} \quad (9.10)$$

Preuve de la proposition. Rappelons que $u \in V$ est un vecteur propre généralisé d'ordre $m \geq 2$ de f pour la valeur propre λ si $(f - \lambda)^m(u) = 0$ et $(f - \lambda)^{m-1}(u) \neq 0$.

Si $m = 1$, nous avons $u_1 = u_m = u$ qui est non nul car c'est un vecteur propre ; il n'y a donc rien à démontrer dans ce cas et on suppose pour la suite de la preuve que $m \geq 2$.

Observons que par définition $(f - \lambda)(u_1) = 0$ et $(f - \lambda)(u_j) = u_{j-1}$ pour $j \geq 2$, par conséquent

$$f(u_1) = \lambda u_1 \quad \text{et} \quad f(u_j) = u_{j-1} + \lambda u_j, \quad \text{pour } j = 2, \dots, m, \quad (9.11)$$

ce qui entraîne en particulier que le sous-espace $U \subset V$ est invariant par f .

Pour montrer que les vecteurs $\{u_1, \dots, u_m\}$ sont linéairement indépendants, on observe d'abord que ces vecteurs sont non nuls en raison de la condition $(f - \lambda)^{j-1}(u_j) = u_1 \neq 0$.

Supposons maintenant que $\sum_{j=1}^m \alpha_j u_j = 0$ avec $\{\alpha_1, \dots, \alpha_m\} \subset K$, et appliquons $(f - \lambda)^{m-1}$ à cette relation. On trouve à partir de (9.10) que

$$0 = (f - \lambda)^{m-1} \left(\sum_{j=1}^m \alpha_j u_j \right) = \sum_{j=1}^m \alpha_j (f - \lambda)^{m-1}(u_j) = \alpha_m u_1,$$

4. On dit aussi *forme canonique* de Jordan, ou *forme réduite* de Jordan ; ces expressions sont synonymes.

par conséquent $\alpha_m = 0$. En appliquant $(f - \lambda)^{m-2}$, on trouve maintenant

$$0 = (f - \lambda)^{m-2} \left(\sum_{j=1}^m \alpha_j u_j \right) = \sum_{j=1}^m \alpha_j (f - \lambda)^{m-2}(u_j) = \alpha_{m-1} u_1 + \alpha_m u_2 = \alpha_{m-1} u_1,$$

ce qui implique que par conséquent $\alpha_{m-1} = 0$. En répétant l'argument, on trouve que $\alpha_j = 0$ pour tout j . □

Rappelons la relation (9.11) qui dit que $f(u_1) = \lambda u_1$ et $f(u_j) = u_{j-1} + \lambda u_j$ pour $k = 2, \dots, m$. La matrice de la restriction de f au sous-espace cyclique $U \subset V$ dans la base cyclique prend la forme

$$J_m(\lambda) = \begin{pmatrix} \lambda & 1 & 0 & \cdots & 0 \\ 0 & \lambda & 1 & \ddots & \vdots \\ \vdots & 0 & \lambda & \ddots & 0 \\ \vdots & \vdots & & \ddots & 1 \\ 0 & 0 & 0 & \cdots & \lambda \end{pmatrix} \quad (9.12)$$

Une telle matrice s'appelle un *bloc de Jordan de taille m* . Par exemple

$$J_1(\lambda) = (\lambda), \quad J_2(\lambda) = \begin{pmatrix} \lambda & 1 \\ 0 & \lambda \end{pmatrix} \quad \text{et} \quad J_3(\lambda) = \begin{pmatrix} \lambda & 1 & 0 \\ 0 & \lambda & 1 \\ 0 & 0 & \lambda \end{pmatrix}.$$

Remarquons aussi qu'un bloc de Jordan $J_m(0)$ de valeur propre $\lambda = 0$ est une matrice nilpotente et que, d'une manière générale, tout bloc de Jordan est somme d'une matrice scalaire et d'une matrice nilpotente puisque

$$J_m(\lambda) = \lambda I_m + J_m(0).$$

Lorsque $U = V$ dans la définition précédente, on dit que V est un espace vectoriel cyclique pour l'endomorphisme f . C'est donc le cas si et seulement s'il existe un vecteur propre généralisé dont l'ordre est égale à la dimension de V . Le lemme suivant nous donne une information sur ces endomorphismes qui sera très utile dans la suite.

Lemme 9.10.4. *Soit $f \in \mathcal{L}(U)$ un endomorphisme λ -cyclique d'ordre $m = \dim(U)$, alors pour tout $k \in \{1, \dots, m\}$ on a*

- (a) $\{u_1, \dots, u_k\}$ est une base de $\text{Ker}(f - \lambda)^k$.
- (b) $\{u_1, \dots, u_{m-k}\}$ est une base de $\text{Im}(f - \lambda)^k$.
- (c) Les multiplicités généralisées de λ valent

$$\delta_{f,\lambda}(k) = \min\{k, m\}.$$

Preuve. Les affirmations (a) et (b) découlent immédiatement des équations (9.10). L'affirmation (c) se déduit de (a) et de la définition $\delta_{f,\lambda}(k) = \dim(\text{Ker}(f - \lambda)^k)$. □

Remarque. On peut aussi prouver ce lemme par le calcul matriciel. La matrice de $(f - \lambda)$ dans la base cyclique est un bloc de Jordan $J_m(0)$ pour $\lambda = 0$; il suffit donc de calculer les puissances de $J_m(0)$ pour

voir quel est le rang. Il est facile de voir que si $1 \leq k \leq m$, alors $J_m(0)^k$ est la matrice qui a le coefficient 1 en position $(i, i+k)$ et 0 partout ailleurs. Les k premières colonnes de cette matrice sont nulles et les $m-k$ dernières colonnes sont linéairement indépendantes. Donc le rang de $J_m(0)^k$ est égale à $m-k$.

Voyons quelques conséquences immédiates du lemme précédent : Soit $f \in \mathcal{L}(V)$ un endomorphisme λ -cyclique d'ordre m , alors

- (i) $(f - \lambda)$ est nilpotent d'ordre m .
- (ii) Tout les vecteurs de V sont des vecteurs propres généralisés de f .
- (iii) λ est l'unique valeur propre de f et sa multiplicité géométrique est 1.
- (iv) On a $\chi_f(t) = \mu_f(t) = (t - \lambda)^m$, en particulier f n'est pas diagonalisable si $m \geq 2$.

9.11 La forme normale de Jordan d'un endomorphisme

Théorème 9.11.1 (Théorème de réduction de Jordan.). *Soient V un espace vectoriel de dimension finie sur un corps K et $f \in \mathcal{L}(V)$ un endomorphisme de V . Si le polynôme caractéristique $\chi_f(t)$ est scindé, alors V peut se décomposer en somme directe de sous-espaces vectoriels cycliques qui sont invariants par f . De plus, le nombre de sous-espaces cycliques associés à une valeur propre λ est égale à la multiplicité géométrique de cette valeur propre.*

Ce théorème dit qu'il existe des sous-espaces vectoriels $V_1, \dots, V_q \subset V$ tels que

- (i) $V = V_1 \oplus \dots \oplus V_q$.
- (ii) Les sous-espaces V_j sont invariants par f , i.e. $f(V_j) \subset V_j$ pour $j = 1, \dots, q$.
- (iii) La restriction $f_j = f|_{V_j}$ est un endomorphisme cyclique de V_j pour une valeur propre λ_j .
- (iv) Pour chaque valeur propre λ_k il y a m_k sous-espaces cycliques, où m_k est la multiplicité géométrique de λ_k .

En particulier V_j est un sous-espace vectoriel d'un sous-espace caractéristique $N_{\lambda_k}(f) \subset V$ (car chaque élément de V_j est un vecteur propre généralisé pour une valeur propre λ_k).

Preuve. La preuve est assez longue et se décompose en plusieurs étapes.

Première étape : réduction au cas d'un endomorphisme n'ayant qu'une valeur propre.

Soit $\sigma(f) = \{\lambda_1, \dots, \lambda_r\}$ le spectre de f . Pour tout i , on note $N_{\lambda_i}(f) \subset V$ le sous-espace caractéristique associé à λ_i . Le théorème de décomposition primaire, avec le théorème de Cayley-Hamilton, nous dit que l'espace V est somme directe des $N_{\lambda_i}(f)$ et que ces espaces sont invariants par f . Pour démontrer le théorème de Jordan, nous pouvons donc supposer que f n'a qu'une seule valeur propre λ (i.e. $V = N_\lambda(f)$).

Deuxième étape : réduction au cas d'un endomorphisme nilpotent.

On suppose donc que $f \in \mathcal{L}(V)$ est un endomorphisme qui n'a qu'une valeur propre λ et dont le polynôme caractéristique est scindé. On sait par le lemme 9.9.1 que ces hypothèses impliquent que $g = (f - \lambda \text{Id}_V)$ est nilpotent. Il est clair que $W \subset V$ est un sous-espace vectoriel invariant par f si et seulement si W est invariant par g . De plus, W est λ -cyclique pour f si et seulement si ce sous-espace est cyclique pour g (associé à l'unique valeur propre de g , qui est 0). Pour démontrer le théorème de Jordan, nous pouvons donc supposer sans perte de généralité que f est nilpotent.

Troisième étape : le cas nilpotent.

Il suffit donc de démontrer que si $f \in \mathcal{L}(V)$ est un endomorphisme nilpotent d'un espace vectoriel V de dimension finie, alors V est somme directe de q sous-espaces invariants cycliques, où $q = \dim(\text{Ker}(f))$.

(On rappelle qu'un endomorphisme nilpotent n'a qu'une valeur propre, qui est 0, la multiplicité géométrique de cette valeur propre est la dimension de $\text{Ker}(f)$).

La preuve se fait par récurrence sur l'ordre de nilpotence m de f . Si f est nilpotent d'ordre 1, alors f est l'endomorphisme nul. On peut choisir une base quelconque $\{v_1, \dots, v_n\}$ de V et noter $V_j = Kv_j = \text{Vec}(v_i)$ le sous-espace vectoriel de dimension 1 engendré par v_j . Alors chaque V_j est trivialement invariant par f et cyclique d'ordre 1, et on a $V = V_1 \oplus \dots \oplus V_n$. De plus $n = \dim(V) = \dim(\text{Ker}(f))$, la preuve est donc complète pour le cas $m = 1$.

On suppose maintenant que l'affirmation est démontrée pour les endomorphismes nilpotents d'ordre $m - 1$ avec $m \geq 2$ et on considère le cas d'un endomorphisme $f \in \mathcal{L}(V)$ nilpotent d'ordre $m \geq 2$. Notons $W = \text{Im}(f)$ et choisissons un sous-espace de $\text{Ker}(f)$ complémentaire à $W \cap \text{Ker}(f)$ qu'on note U . On a donc

$$\text{Ker}(f) = (\text{Ker}(f) \cap W) \oplus U.$$

Les sous-espaces U et W de V sont invariants par f (car le noyau et l'image d'un endomorphisme sont toujours des sous-espaces invariants). La restriction de f à U est l'endomorphisme nul et la restriction de f à W est un endomorphisme nilpotent d'ordre $(m - 1)$.

Par hypothèse de récurrence, il existe une décomposition de W en somme directe de sous-espaces invariants cycliques :

$$W = \text{Im}(f) = W_1 \oplus \dots \oplus W_q, \quad f(W_j) \subset W_j \text{ et } W_j \text{ est cyclique pour } f,$$

de plus $q = \dim(\text{Ker}(f|_W)) = \dim(\text{Ker}(f) \cap W)$.

Chaque sous-espace W_j admet donc une base cyclique $\mathcal{C}_j = \{w_{j,1}, \dots, w_{j,m_j}\}$, où $m_j = \dim W_j$. Rappelons que cela signifie que $f(w_{j,1}) = 0$ et $f(w_{j,i}) = w_{j,i-1}$ pour $i > 1$.

Puisque $W_j \subset W = \text{Im}(f)$, il existe un vecteur $v_j \in V$ tel que $f(v_j) = w_{j,m_j}$; on note alors

$$\mathcal{B}_j = \mathcal{C}_j \cup \{v_j\} = \{w_{j,1}, \dots, w_{j,m_j}, v_j\} \quad \text{et} \quad V_j = W_j + Kv_j = \text{Vec}(\mathcal{B}_j).$$

Nous affirmons que les sous-espaces V_j et les familles $\mathcal{B}_j$ possèdent les propriétés suivantes :

- (a) La réunion $\mathcal{B} = \mathcal{B}_1 \cup \dots \cup \mathcal{B}_q$ est une famille libre de V .
- (b) $\mathcal{B}_j = \{w_{j,1}, \dots, w_{j,m_j}, v_j\}$ est une base de V_j .
- (c) V_j est invariant par f .
- (d) V_j est cyclique d'ordre $m_j + 1$ pour f .

Pour prouver (a), il est commode de renoter par w_{j,m_j+1} le vecteur v_j . Supposons que

$$\sum_{j=1}^q \sum_{i=1}^{m_j+1} \alpha_{j,i} \cdot w_{j,i} = 0,$$

en appliquant f à cette relation de dépendance linéaire, on obtient que

$$\sum_{j=1}^q \sum_{i=2}^{m_j+1} \alpha_{j,i} \cdot w_{j,i-1} = \sum_{j=1}^q \sum_{i=1}^{m_j+1} \alpha_{j,i} \cdot f(w_{j,i}) = 0,$$

(on utilise que $f(w_{j,1}) = 0$). En décalant l'indice i d'une unité on peut écrire

$$\sum_{j=1}^q \sum_{i=1}^{m_j} \alpha_{j,i+1} \cdot w_{j,i} = 0.$$

Or cette identité implique que chaque $\alpha_{j,i+1} = 0$ pour tout $i \geq 1$ car la réunion des $\mathcal{C}_j$ est une base de W . Mais alors on a aussi

$$\sum_{j=1}^q \alpha_{j,1} \cdot w_{j,1} = 0,$$

et donc chaque $\alpha_{j,1} = 0$ car $\{w_{1,1}, \dots, w_{q,1}\}$ est une partie libre (c'est un sous-ensemble de $\mathcal{C}_j$).

L'affirmation (b) est maintenant immédiate puisque $\mathcal{B}_j$ est une famille libre qui engendre V_j .

Les affirmations (c) et (d) découlent du fait que $f(\mathcal{B}_j) \subset \mathcal{B}_j \cup \{0\}$ et que $\mathcal{B}_j$ est une base cyclique de V_j par construction.

Nous pouvons maintenant conclure la preuve du théorème. L'affirmation (a) et la définition de U entraînent que $U \oplus V_1 \oplus \dots \oplus V_q \subset V$ est une somme directe. Nous affirons que

$$V = U \oplus V_1 \oplus \dots \oplus V_q. \quad (9.13)$$

En effet, par définition de U on a

$$\dim(\text{Ker}(f)) = \dim(U) + \dim(\text{Ker}(f) \cap W) = \dim(U) + q,$$

et d'autre part $\dim(V_j) = \dim(W_j) + 1$ pour tout $j = 1, \dots, q$. Par conséquent

$$\begin{aligned} \dim(U \oplus V_1 \oplus \dots \oplus V_q) &= (\dim(\text{Ker}(f)) - q) + \sum_{j=1}^q (\dim(W_j) + 1) \\ &= \dim(\text{Ker}(f)) + \sum_{j=1}^q \dim(W_j) \\ &= \dim(\text{Ker}(f)) + \dim(W) \\ &= \dim(V), \end{aligned}$$

puisque $W = \text{Im}(f)$. L'égalité de ces dimensions impliquent la somme directe (9.13). Finalement, la restriction de f à U est l'endomorphisme nul. On peut donc décomposer U en sous-espaces de dimension 1 (en choisissant une base quelconque), disons $U = U_1 \oplus \dots \oplus U_p$ avec $p = \dim(U)$. On a donc la décomposition

$$V = U_1 \oplus \dots \oplus U_p \oplus V_1 \oplus \dots \oplus V_q,$$

en $(p + q)$ -sous-espaces cycliques invariants, et

$$p + q = \dim(U) + \dim(\text{Ker}(f) \cap W) = \dim(\text{Ker}(f)).$$

La preuve du théorème est complète. □

9.12 Conséquences du théorème de réduction de Jordan

On peut également énoncer le théorème 9.11.1 sous la forme suivante :

Théorème 9.12.1. *Soit $f \in \mathcal{L}(V)$ un endomorphisme d'un espace vectoriel V de dimension finie. Si le polynôme caractéristique de f est scindé, alors il existe une base $\mathcal{B}$ de V dans laquelle la matrice de f est diagonale par blocs et chaque bloc est une matrice de Jordan.*

$$\begin{pmatrix} \boxed{J_{m_1}(\lambda_{i_1})} & & & \\ & \boxed{J_{m_2}(\lambda_{i_2})} & & \\ & & \ddots & \\ & & & \boxed{J_{m_q}(\lambda_{i_q})} \end{pmatrix}. \quad (9.14)$$

Le nombre de blocs de Jordan associé à chaque valeur propre est égale à la multiplicité géométrique de cette valeur propre.

Une telle base $\mathcal{B}$ de V s'appelle une *base de Jordan* pour l'endomorphisme f , et la matrice (9.14) s'appelle la *forme normale de Jordan*, ou *forme canonique de Jordan* pour f .

Le résultat suivant calcule le nombre de blocs de Jordan de chaque taille :

Proposition 9.12.2. *Soit $f \in \mathcal{L}(V)$ comme dans le théorème précédent et λ une valeur propre de f . On note $\alpha_\lambda(m)$ le nombre de blocs de Jordan de taille m dans la matrice (9.14). Alors pour tout $m \geq 1$ on a*

$$\boxed{\alpha_\lambda(m) = 2\delta_\lambda(m) - \delta_\lambda(m+1) - \delta_\lambda(m-1)}, \quad (9.15)$$

où les $\delta_\lambda(k) = \dim(\text{Ker}(f - \lambda)^k)$ sont les multiplicités généralisées de f .

Preuve. On rappelle que si $J = J_m(\lambda)$ est un bloc de Jordan de taille m , alors $\dim(\text{Ker}(J - \lambda)^k) = \min\{m, k\}$. On a donc pour tout k

$$\delta_\lambda(k) = \sum_{q=1}^n \alpha_\lambda(q) \min\{q, k\}, \quad \text{où } n = \dim(V),$$

et par conséquent :

$$2\delta_\lambda(m) - \delta_\lambda(m+1) - \delta_\lambda(m-1) = \sum_{q=1}^n \alpha_\lambda(q) (2\min\{q, m\} - \min\{q, m-1\} - \min\{q, m+1\}).$$

L'égalité 9.15 découle maintenant de l'identité suivante, valide pour des entiers naturels q, m quelconque et dont nous laissons la vérification en exercice :

$$2\min\{q, m\} - \min\{q, m-1\} - \min\{q, m+1\} = \begin{cases} 1, & \text{si } q = m, \\ 0, & \text{si } q \neq m. \end{cases}$$

□

Corollaire 9.12.3. *Toute matrice $A \in M_n(\mathbb{C})$ est semblable à une matrice de type (9.14). Cette matrice est unique à l'ordre des blocs près.*

Preuve. L'existence d'une forme canonique (9.14) pour toute matrice $A \in M_n(\mathbb{C})$ se déduit immédiatement du théorème de réduction de Jordan. L'unicité à l'ordre des blocs près provient de la proposition précédente qui calcule le nombre de bloc de chaque ordre associé à chaque valeur propre en fonction des multiplicités généralisées $\delta_{A,\lambda}(k)$ dimensions des noyaux $\text{Ker}(A - \lambda)^k$ pour tout k , en observant que ces dimensions sont les mêmes pour deux matrices semblables. $\square$

Définition. Si $A \in M_n(\mathbb{C})$ et si $A' = P^{-1}AP$ est de type (9.14), alors on dit que A' est la *forme canonique de Jordan* de la matrice A . On notera parfois $A' = J[A]$ la forme de Jordan de la matrice A , bien qu'elle ne soit unique qu'à permutation des blocs de Jordan près.

La proposition suivante nous donne des informations détaillées sur les blocs de Jordan :

Proposition 9.12.4. *Supposons que*

$$\chi_f(t) = \prod_{i=1}^r (t - \lambda_i)^{m_i} \quad \text{et} \quad \mu_f(t) = \prod_{i=1}^r (t - \lambda_i)^{s_i},$$

alors on a les propriétés suivantes de la forme normale de Jordan :

- (i) La taille de chaque bloc $J_p(\lambda_i)$ est au plus égale à s_i
- (ii) Pour tout i , il existe au moins un bloc de Jordan $J_{s_i}(\lambda_i)$ de taille s_i .
- (iii) Le nombre total de blocs de Jordan pour λ_i est égal à la multiplicité géométrique de λ_i .
- (iv) La somme des tailles des blocs de Jordan pour λ_i est égale à la multiplicité algébrique m_i de λ_i .
- (v) La dimension de V est la somme des tailles de tous les blocs de Jordan.
- (vi) Le nombre de blocs de Jordan de chaque taille pour la valeur propre λ_i est déterminé par les multiplicités généralisées $\delta_{\lambda_i}(k)$ ($1 \leq k \leq s_i$), selon l'équation (9.15).

Chaque propriété est une conséquence assez-simple des résultats précédents. Nous laissons la vérification en exercice.

Exemple. La proposition précédente implique immédiatement que si $A \in M_3(K)$ est une matrice telle que $\mu_A(t) = (t - \lambda)^2$, alors sa forme normale de Jordan est

$$J[A] = J_2(\lambda) \oplus J_1(\lambda) = \begin{pmatrix} \lambda & 1 & 0 \\ 0 & \lambda & 0 \\ 0 & 0 & \lambda \end{pmatrix}$$

Une conséquence des résultats précédents est le

Théorème 9.12.5. *Soit V un espace vectoriel de dimension finie sur un corps K quelconque et $f, g \in \mathcal{L}(V)$ deux endomorphismes de V dont les polynômes caractéristiques sont scindés. Alors les conditions suivantes sont équivalentes :*

- (i) f et g sont conjugués.
- (ii) On a $\sigma(f) = \sigma(g)$ et égalité de toutes les multiplicités généralisées : $\delta_{f,\lambda}(k) = \delta_{g,\lambda}(k)$ pour toute valeur propre λ et tout entier k .

(iii) f et g ont la même forme de Jordan (à l'ordre des blocs près).

De même, deux matrices $A, B \in M_n(K)$ dont les polynômes caractéristiques sont scindés sont semblables si et seulement si elles ont les mêmes multiplicités généralisées et donc la même forme de Jordan (à l'ordre des blocs près).

La condition (ii) de ce théorème s'exprime parfois en disant que le spectre $\sigma(f)$, avec la famille de toutes les multiplicités généralisées $\delta_{f,\lambda}(k)$ forment un *système complet d'invariants* pour la classe de conjugaison d'un endomorphisme f dont le polynôme caractéristique est scindé. Une application intéressante de ce théorème est donnée dans l'exercice suivant :

Exercice. Prouver que toute matrice $A \in M_n(\mathbb{C})$ est semblable à sa transposée $A^\top$.

9.13 Réduction pratique d'une matrice à sa forme normale de Jordan

Examinons comment réduire concrètement une matrice A à sa forme normale de Jordan. On parle parfois de “jordanisation” de la matrice, comprise comme une généralisation de la diagonalisation.

On se donne donc une matrice $A \in M_n(K)$ où K est un corps quelconque, et on suppose que le polynôme caractéristique $\chi_A(t)$ est scindé⁵ :

$$\chi_A(t) = \prod_{i=1}^r (t - \lambda_i)^{m_i}, \quad \sigma(A) = \{\lambda_1, \dots, \lambda_r\} \subset K, \quad \{m_1, \dots, m_r\} \subset \mathbb{N}.$$

L'entier m_i est la multiplicité algébrique de la valeur propre λ_i .

La *forme normale de Jordan* de A est alors une matrice $J[A] \in M_n(K)$ qui vérifie les trois conditions suivantes :

- (i) $J[A]$ est semblable à A , i.e. il existe $P \in GL_n(K)$ tel que $J[A] = P^{-1}AP$.
- (ii) $J[A]$ est une matrice diagonale par bloc.
- (iii) Chaque bloc est un bloc de Jordan $J_m(\lambda)$ associé à une valeur propre $\lambda \in \sigma(A)$.

L'existence et l'unicité de la matrice $J[A]$ (à l'ordre des blocs de Jordan près) ont été démontrées au paragraphe précédent. Les colonnes de la matrice P forment une *base de Jordan* pour A .

Remarque. Un bloc de Jordan de taille 1 est simplement un scalaire car $J_1(\lambda)$ est la 1×1 matrice (λ) . Lorsqu'une matrice est diagonalisable, sa jordanisation n'est rien d'autre que sa diagonalisation (et chaque bloc de Jordan est de taille 1).

D'un point de vue pratique, la réduction d'une matrice A à sa forme normale de Jordan se décompose en deux problèmes :

Problème 1. Déterminer la forme de Jordan $J[A]$ de A .

Problème 2. Trouver la matrice de changement de base P telle que $J[A] = P^{-1}AP$.

La solution du problème 1 révèle la structure de l'endomorphisme associé à A . Pour jordaniser une matrice, il est préférable de résoudre le problème 1 *avant* de résoudre le problème 2. Pour trouver la matrice P il suffit de construire une base de Jordan de A .

5. Rappelons que c'est toujours le cas si $K = \mathbb{C}$.

Problème 1. Déterminer la forme normale de Jordan d'une matrice.

Rappelons la proposition 9.12.2 nous permet de déduire la forme normale de Jordan $J[A]$ d'une matrice A à partir des multiplicités généralisées. Toutefois l'équation (9.15) est assez lourde à utiliser et la proposition 9.12.4 nous donne des informations qui sont parfois suffisantes lorsqu'on connaît le polynôme caractéristique et le polynôme minimal d'une matrice A (que l'on suppose scindés).

Exemple 1. (a) On veut déterminer toutes les formes normales de Jordan possibles d'une matrice A dont le polynôme caractéristique est $\chi_A(t) = (t - 2)^4$ et le polynôme minimal est $\mu_A(t) = (t - 2)^2$.

Pour répondre à cette question, on observe que A est une matrice de taille 4×4 car $\deg(\chi_A(t)) = 4$. Il n'y a qu'une valeur propre, qui est $\lambda = 2$. On sait aussi $(A - \lambda I_4)$ est nilpotent d'ordre 2 car $\mu_A(t) = (t - 2)^2$. Donc la forme normale de Jordan de A peut ou bien contenir deux blocs de Jordan $J_2(2)$ ou un bloc $J_2(2)$ et deux blocs $J_1(2)$. Les formes normales de Jordan possibles pour A sont donc

$$J_2(2) \oplus J_2(2) = \begin{pmatrix} 2 & 1 & 0 & 0 \\ 0 & 2 & 0 & 0 \\ 0 & 0 & 2 & 1 \\ 0 & 0 & 0 & 2 \end{pmatrix} \quad (\text{si } \text{multgeom}_2(A) = 2)$$

et

$$J_2(2) \oplus J_1(2) \oplus J_1(2) = \begin{pmatrix} 2 & 1 & 0 & 0 \\ 0 & 2 & 0 & 0 \\ 0 & 0 & 2 & 0 \\ 0 & 0 & 0 & 2 \end{pmatrix} \quad (\text{si } \text{multgeom}_2(A) = 3)$$

(b) Supposons que les polynômes caractéristique et minimal de la matrice A sont respectivement $\chi_A(t) = (t - 3)^3(t + 1)^4$ et $\mu_A(t) = (t - 3)^2(t + 1)^3$. Alors la seule forme canonique de Jordan possible pour A s'écrit, à permutation des blocs près, sous la forme suivante :

$$J[A] = J_1(3) \oplus J_2(3) \oplus J_1(-1) \oplus J_3(-1) = \begin{pmatrix} 3 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 3 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 3 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & -1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & -1 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & -1 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & -1 & 1 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & -1 \end{pmatrix}$$

En effet par la condition (ii) énoncée plus haut il existe au moins un bloc de Jordan $J_2(3)$ et un bloc $J_3(-1)$ car $\mu_A = (t - 3)^2(t + 1)^3$. Mais la condition (iii) dit que la somme des tailles des blocs de Jordan pour la valeur propre 3 est égale à 3 (i.e. la multiplicité algébrique) et la somme des tailles des blocs de Jordan pour la valeur propre -1 est égale à 4. Donc il n'y qu'une façon de compléter, et c'est d'ajouter un bloc $J_1(3)$ et un bloc $J_1(-1)$.

(c) Considérons une variante où le polynôme caractéristique et minimal sont $\chi_A(t) = (t - 3)^3(t + 1)^4$ et $\mu_A(t) = (t - 3)^2(t + 1)^2$. Alors il y a deux formes de Jordan possibles :

$$A' = J_1(3) \oplus J_2(3) \oplus J_2(-1) \oplus J_2(-1) \quad \text{et} \quad A'' = J_1(3) \oplus J_2(3) \oplus J_2(-1) \oplus J_1(-1) \oplus J_1(-1)$$

La multiplicité géométrique de la valeurs propre -1 est égale à 2 dans le premier cas et à 3 dans le second cas.

Remarque. Il n'est pas toujours possible de déterminer la forme de Jordan uniquement à partir des polynômes caractéristique et minimal et des multiplicités géométriques. Dans un tel cas il faut calculer quelques multiplicités généralisées (au pire les calculer toutes).

Problème 2. Trouver une base de Jordan.

Pour trouver une base de Jordan d'une matrice $A \in M_n(K)$, il faut d'abord déterminer sa forme normale de Jordan (sinon on ne sait pas ce qu'on cherche). Une base de Jordan est une base formée de vecteurs propres généralisé qui forment une famille de cycles. Il y a autant de cycles que de blocs de Jordan et la longueur de chaque cycle correspond à la taille du bloc de Jordan correspondant. Un cycle associé à la valeur propre λ est une suite de vecteurs $\{u_1, \dots, u_m\}$ telle que

$$(A - \lambda)(u_m) = u_{m-1}, (A - \lambda)(u_{m-1}) = u_{m-2}, \dots, (A - \lambda)(u_2) = u_1, (A - \lambda)(u_1) = 0.$$

En particulier u_1 est vecteur propre. Dans une base de Jordan, chaque cycle doit être maximal, i.e. $u_m \notin \text{Im}(A - \lambda)^{m-1}$. La base de Jordan est construite lorsqu'on ne peut plus construire de nouveau cycles linéairement indépendant des précédents. La matrice de changement de base P est alors la matrice dont les colonnes sont les coordonnées des vecteurs de notre base. Il est important de vérifier que $J[A] = P^{-1}AP$ (ou si on préfère $PJ[A] = AP$, ce qui permet d'économiser le calcul de P^{-1}).

Exemple 2. On demande de jordaniser la matrice

$$A := \begin{pmatrix} 7 & 0 & 2 \\ 3 & 7 & 2 \\ 0 & 0 & 3 \end{pmatrix}$$

Le polynôme caractéristique est $\chi_A(t) = (t - 7)^2(t - 3)$, en particulier il est scindé et les valeurs propres sont 3 et 7. On a

$$(A - 3I_3) = \begin{pmatrix} 4 & 0 & 2 \\ 3 & 4 & 2 \\ 0 & 0 & 0 \end{pmatrix}, \quad (A - 7I_3) = \begin{pmatrix} 0 & 0 & 2 \\ 3 & 0 & 2 \\ 0 & 0 & -4 \end{pmatrix}, \quad (A - 7I_2)^2 = \begin{pmatrix} 0 & 0 & -8 \\ 0 & 0 & -2 \\ 0 & 0 & 16 \end{pmatrix}.$$

On voit facilement que $\text{multgeom}_A(7) = 1 < \text{multalg}_A(7) = 2$. Ceci implique que A n'est pas diagonalisable et que le polynôme minimal est donc $\mu_A(t) = (t - 7)^2(t - 3)$ et la forme normale de Jordan de A est

$$J[A] = J_1(3) \oplus J_2(7) = \begin{pmatrix} 3 & 0 & 0 \\ 0 & 7 & 1 \\ 0 & 0 & 7 \end{pmatrix}$$

Pour construire une base de Jordan on doit donc trouver un vecteur propre pour la valeur propre 3 et un cycle de longueur 2 pour la valeur propre 7.

Un vecteur propre pour $\lambda = 3$ est $\{X = (4, 1, -8)\}$. Pour construire un cycle associé à la valeur propre $\lambda = 7$ on cherche d'abord un vecteur $Y_2 \in \text{Ker}(A - 7)^2 \setminus \text{Ker}(A - 7)$. On peut choisir

$Y_2 = (1, 0, 0)$. Le deuxième vecteur du cycle est alors $Y_1 = (A - 7)Y_2 = (0, 3, 0)$. La base de Jordan cherchée est

$$\{X, Y_1, Y_2\} = \{(4, 1, -8), (0, 3, 0), (1, 0, 0)\}.$$

Pour finaliser la jordanisation, on pose

$$P = \begin{pmatrix} 4 & 0 & 1 \\ 1 & 3 & 0 \\ -8 & 0 & 0 \end{pmatrix}, \text{ alors } P^{-1} = \begin{pmatrix} 0 & 0 & -1/8 \\ 0 & 1/3 & 1/24 \\ 1 & 0 & 1/2 \end{pmatrix} \text{ et } P^{-1}AP = \begin{pmatrix} 3 & 0 & 0 \\ 0 & 7 & 1 \\ 0 & 0 & 7 \end{pmatrix}.$$

Exemple 3. On demande de jordaniser la matrice $B = \begin{pmatrix} 0 & 1 & 0 & 2 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 2 & 0 \end{pmatrix}$.

Le polynôme caractéristique est $\chi_B(t) = t^4$ et cette matrice est donc nilpotente. On calcule que

$$B^2 = \begin{pmatrix} 0 & 0 & 4 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix} \text{ et } B^3 = \begin{pmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix},$$

ainsi B est nilpotente d'ordre 3. Le rang de B est 2 et le rang de B^2 est 1. On en déduit que la forme normale de Jordan est

$$J[B] = J_3(0) \oplus J_1(0) = \begin{pmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix}$$

Cherchons une base de Jordan, elle doit contenir un cycle de longueur 3 et un cycle de longueur 1. Pour construire le cycle de longueur 3, on cherche un vecteur $X_3 \in \text{Ker}(B^3) \setminus \text{Ker}(B^2)$, par exemple $X_3 = (0, 0, 1)$. On complète le cycle en posant $X_2 = BX_3 = (0, 0, 0, 2)$ et $X_1 = BX_2 = B^2X_3 = (4, 0, 0)$.

Le cycle de longueur 1 est maintenant donné par un vecteur du noyau de B et qui est linéairement indépendant de X_1 . On peut prendre $Y_1 = (0, 2, 0, -1)$. On a donc construit une base de Jordan pour B :

$$\{X_1, X_2, X_3, Y_1\} = \{(4, 0, 0, 0), (0, 0, 0, 2), (0, 0, 1, 0), (0, 2, 0, -1)\}.$$

On vérifie qu'il s'agit d'une base de Jordan en posant $Q = \begin{pmatrix} 4 & 0 & 0 & 0 \\ 0 & 0 & 0 & 2 \\ 0 & 0 & 1 & 0 \\ 0 & 2 & 0 & -1 \end{pmatrix}$ et en vérifiant que

$$Q^{-1}BQ = J[B] \text{ (ou si on préfère } BQ = J[B]Q).$$

Exemple 4. Soit à jordaniser la matrice

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

Le polynôme caractéristique est $\chi_C(t) = (t-2)^3(t-3)$, le spectre est $\{2, 3\}$ et on a

$$(C - 3I_4) = \begin{pmatrix} -1 & 1 & 0 & -1 \\ 1 & -1 & -\frac{1}{2} & 1 \\ 0 & 2 & -1 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix} \quad \text{et} \quad (C - 2I_4) = \begin{pmatrix} 0 & 1 & 0 & -1 \\ 1 & 0 & -\frac{1}{2} & 1 \\ 0 & 2 & 0 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

La valeur propre $\lambda = 3$ est de multiplicité algébrique 1 ; l'espace propre $E_3(C) = \text{Ker}(C - 3I_4)$ associé à la valeur propre 3 est de dimension 1 et il est engendré par le vecteur $Y = (1, 0, 0, -1)$. La valeur propre $\lambda = 2$ est de multiplicité algébrique 3 et de multiplicité géométrique égale à $1 = \dim \text{Ker}(C - 2I_4)$. la forme normale de C possède un unique bloc de Jordan pour chaque valeur propre et on peut déjà conclure que

$$J[C] = J_3(2) \oplus J_1(3) = \begin{pmatrix} 2 & 1 & 0 & 0 \\ 0 & 2 & 1 & 0 \\ 0 & 0 & 2 & 0 \\ 0 & 0 & 0 & 3 \end{pmatrix}.$$

Nous devons construire un cycle de longueur 3 pour la valeur propre $\lambda = 2$. On a on a

$$(C - 2I_4)^2 = \begin{pmatrix} 1 & 0 & -\frac{1}{2} & 0 \\ 0 & 0 & 0 & 0 \\ 2 & 0 & -1 & 2 \\ 0 & 0 & 0 & 1 \end{pmatrix} \quad \text{et} \quad (C - 2I_4)^3 = \begin{pmatrix} 0 & 0 & 0 & -1 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}.$$

On doit choisir un vecteur $X_3 \in \text{Ker}((C - 2I_4)^3) \setminus \text{Ker}((C - 2I_4)^2)$, prenons le vecteur $X_3 = (0, 0, -2, 0)$ (on pourrait prendre $(0, 0, 1, 0)$, mais notre choix va simplifier un peu les calculs). On définit ensuite les vecteurs $X_2 = (C - 2I_4)X_3 = (0, 1, 0, 0)$ et $X_1 = (C - 2I_4)X_2 = (1, 0, 2, 0)$. Nous avons construit notre base de Jordan $\{X_1, X_2, X_3, Y\}$. La matrice de changement de base est donnée par

$$P = \begin{pmatrix} 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & 0 \\ 2 & 0 & -2 & 0 \\ 0 & 0 & 0 & -1 \end{pmatrix}, \quad P^{-1} = \begin{pmatrix} 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & 0 \\ 1 & 0 & -\frac{1}{2} & 1 \\ 0 & 0 & 0 & -1 \end{pmatrix}$$

On vérifie que $CP = PJ[C]$.

9.14 Sur les endomorphismes d'espaces vectoriels réels.

Dans ce paragraphe, nous étudions la structure des endomorphismes d'un espace vectoriel V de dimension finie sur le corps $\mathbb{R}$ des réels. Rappelons que, grâce au théorème fondamental de l'algèbre, le polynôme caractéristique de tout endomorphisme d'un $\mathbb{C}$ -espace vectoriel est scindé et donc il admet une base de Jordan.

Dans le cas réel, le polynôme caractéristique n'est pas toujours scindé mais nous rappelons le résultat suivant :

Lemme 9.14.1 (Décomposition d'un polynôme à coefficients réels en facteurs irréductibles).
Tout polynôme $p(t) \in \mathbb{R}[t]$ peut s'écrire

$$p(t) = \prod_{i=1}^s (t - \lambda_i)^{m_i} \cdot \prod_{j=1}^r q_j(t)^{n_j},$$

où $\lambda_j \in \mathbb{R}$ pour tout $j = 1, \dots, r$ et $q_i(t) \in \mathbb{R}[t]$ est un polynôme irréductible de degré 2 pour tout $i = 1, \dots, r$.

On dit que les termes $(t - \lambda_i)$ sont les *facteurs linéaires* (ou facteurs du premier degré) et les $q_j(t)$ sont les *facteurs quadratiques* de $p(t)$. Ils sont uniquement déterminés par le polynôme $p(t)$ à permutation des facteurs près. Remarquons que ce lemme implique en particulier que tout polynôme de $\mathbb{R}[t]$ de degré impair admet au moins une racine réelle⁶.

Rappelons rapidement la preuve de ce lemme (qui a été vue aux exercices) : Il est clair que $\mathbb{R}[t] \subset \mathbb{C}[t]$, i.e. tout polynôme à coefficients réels est aussi un polynôme à coefficients complexes. En particulier $p(z)$ est un nombre complexe bien défini pour tout $z \in \mathbb{C}$. Mais pour un polynôme à coefficients réels on a $p(\bar{z}) = \overline{p(z)}$. En particulier

$$p(\lambda) = 0 \quad \Rightarrow \quad p(\bar{\lambda}) = 0.$$

En utilisant le théorème fondamental de l'algèbre, on peut maintenant factoriser $p(t)$ comme polynôme à coefficients complexes :

$$p(t) = \prod_{i=1}^d (t - \lambda_i)^{m_i},$$

et la remarque précédente nous dit que les racines complexes apparaissent par paires de racines conjuguées. Nous pouvons donc renuméroter les racines de la façon suivante :

$$\lambda_1, \dots, \lambda_s, \lambda_{s+1}, \bar{\lambda}_{s+1}, \dots, \lambda_{s+r}, \bar{\lambda}_{s+r},$$

où $\lambda_i \in \mathbb{R}$ pour $i \leq s$ et $\lambda_i \in \mathbb{C} \setminus \mathbb{R}$ pour $i > s$. On peut alors noter $\lambda_{s+j} = \alpha_j + \sqrt{-1}\beta_j$ (avec $\alpha_j, \beta_j \in \mathbb{R}$ et $\beta_j \neq 0$). On a finalement

$$p(t) = \prod_{i=1}^s (t - \lambda_i)^{m_i} \cdot \prod_{j=1}^r [(t - \lambda_j)(t - \bar{\lambda}_j)]^{n_j} = \prod_{i=1}^s (t - \lambda_i)^{m_i} \cdot \prod_{j=1}^r q_j(t)^{n_j},$$

où les facteurs $q_j(t)$ sont des polynômes quadratiques à coefficients réels. Plus précisément

$$\begin{aligned} q_i(t) &= (t - \lambda_j)(t - \bar{\lambda}_j) = t^2 - (\lambda_j + \bar{\lambda}_j)t + \lambda_j \cdot \bar{\lambda}_j \\ &= t^2 - 2\alpha_j t + (\alpha_j^2 + \beta_j^2) \\ &= (t - \alpha_j)^2 + \beta_j^2. \end{aligned}$$

□

Considérons maintenant un endomorphisme f d'un espace vectoriel réel V de dimension finie. On note

$$\sigma_{\mathbb{R}}(f) = \text{L'ensemble des racines réelles du polynôme caractéristique } \chi_f(t),$$

et

$$\sigma_{\mathbb{C}}(f) = \text{L'ensemble des racines complexes de } \chi_f(t),$$

6. Cela se démontre aussi facilement à partir du théorème de la valeur intermédiaire.

On sait que $\sigma_{\mathbb{R}}(f)$ est l'ensemble des valeurs propres de f (i.e. $\lambda \in \sigma_{\mathbb{R}}(f)$ si et seulement si $\lambda \in \mathbb{R}$ et il existe un vecteur non nul v de V tel que $f(v) = \lambda v$), on sait aussi que

$$\lambda \in \sigma_{\mathbb{C}}(f) \quad \Leftrightarrow \quad \bar{\lambda} \in \sigma_{\mathbb{C}}(f).$$

Les éléments de $\sigma_{\mathbb{C}}(f) \setminus \mathbb{R}$ seront appelés les *valeurs propres complexes* de f .

Question : Quelle est la signification réelle des valeurs propres complexes ?

La notion suivante est la clé pour répondre à cette question :

Définition. Soit V un espace vectoriel réel (on ne suppose pas $\dim(V) < \infty$). On appelle *complexifié* de V , et on note $V_{\mathbb{C}}$ l'espace vectoriel ainsi défini :

(i) Comme groupe abélien $V_{\mathbb{C}} = V \times V$ avec la loi de groupe du produit direct :

$$(u_1, v_1) + (u_2, v_2) = (u_1 + u_2, v_1 + v_2).$$

(ii) La multiplication d'un vecteur $(u, v) \in V_{\mathbb{C}}$ par un nombre complexe $\lambda = \alpha + i\beta \in \mathbb{C}$ est définie par

$$(\alpha + i\beta) \cdot (u, v) = (\alpha u - \beta v, \beta u + \alpha v).$$

Proposition 9.14.2. *L'espace $V_{\mathbb{C}}$ est un espace vectoriel sur le corps $\mathbb{C}$ pour les opérations ainsi définies.*

Nous laissons la preuve en exercice (rappelons qu'il s'agit de vérifier 8 axiomes qui ne sont que des règles de calculs, les 4 premiers axiomes rappellent simplement le fait connu que $V \times V$ est un groupe abélien pour la somme définie selon les composantes).

La structure de cet espace vectoriel est plus intuitive si l'on note un élément $w = (u, v) \in V_{\mathbb{C}}$ sous la forme $w = u + iv = u + \sqrt{-1}v$. Alors la multiplication par un scalaire complexe est

$$\lambda v = (\alpha + i\beta)(u + iv) = (\alpha u - \beta v) + i(\beta u + \alpha v).$$

Avec cette notation la preuve de la proposition est très facile.

Exemples 1.) Le complexifié de l'espace numérique $V = \mathbb{R}^n$ est $V_{\mathbb{C}} = \mathbb{C}^n$.

2.) Le complexifié de l'espace vectoriel $\mathbb{R}[t]$ des polynômes à coefficients réels est l'espace vectoriel $\mathbb{C}[t]$ des polynômes à coefficients complexes.

3.) Le complexifié de l'espace $C^k([a, b], \mathbb{R})$ des fonctions de classe C^k à valeurs réelles sur un intervalle $[a, b]$ est l'espace $C^k([a, b], \mathbb{C})$ des fonctions C^k à valeurs complexes sur $[a, b]$.

Sur l'espace vectoriel complexe $V_{\mathbb{C}}$ on peut définir l'opération de *conjugaison complexe* par $\overline{(u, v)} = (u, -v)$, ou si on préfère :

$$w = u + iv \in V_{\mathbb{C}} \quad \Rightarrow \quad \bar{w} = u - iv.$$

Cette opération est $\mathbb{C}$ -*antilinéaire*, c'est-à-dire qu'on a les propriétés :

$$\overline{w_1 + w_2} = \bar{w}_1 + \bar{w}_2 \quad \text{et} \quad \overline{\lambda w} = \bar{\lambda} \bar{w}.$$

Un vecteur w de $V_{\mathbb{C}}$ est réel (i.e. c'est un élément de V) si et seulement si $w = \bar{w}$. On peut alors définir les *parties réelles* et *imaginaires* d'un vecteur $w \in V_{\mathbb{C}}$ du complexifié par

$$\operatorname{Ré}(w) = \frac{w + \bar{w}}{2} \quad \operatorname{Im}(w) = \frac{w - \bar{w}}{2i}.$$

Observer que si $w = u + iv$, avec $u, v \in V$, alors $\operatorname{Ré}(w) = u$ et $\operatorname{Im}(w) = v$.

A tout endomorphisme $f \in \mathcal{L}(V)$ d'un espace vectoriel réel V on associe un endomorphisme noté $f_{\mathbb{C}} \in \mathcal{L}(V_{\mathbb{C}})$ de l'espace vectoriel complexifié $V_{\mathbb{C}}$ de V . Cet endomorphisme est simplement défini par

$$f_{\mathbb{C}}(w) = f_{\mathbb{C}}(u + iv) = f(u) + if(v).$$

On vérifie alors facilement la proposition suivante :

Proposition 9.14.3. *Si $\mathcal{B} = \{v_1, \dots, v_n\}$ est une base de l'espace vectoriel réel V , alors $\mathcal{B}$ est aussi une base du complexifié $V_{\mathbb{C}} = V + iV$. De plus si $A = M_{\mathcal{B}}(f)$ est la matrice de f dans la base $\mathcal{B}$, alors A est aussi la matrice de $f_{\mathbb{C}}$ dans cette base :*

$$M_{\mathcal{B}}(f_{\mathbb{C}}) = M_{\mathcal{B}}(f).$$

En particulier $V_{\mathbb{C}}$ a la même dimension (comme espace vectoriel complexe) que V (comme espace vectoriel réel) :

$$\dim_{\mathbb{C}}(V_{\mathbb{C}}) = \dim_{\mathbb{R}}(V).$$

Remarquons cependant que $V_{\mathbb{C}}$ est aussi un espace vectoriel sur le corps $\mathbb{R}$ et qu'on a

$$\dim_{\mathbb{R}}(V_{\mathbb{C}}) = 2 \dim_{\mathbb{R}}(V).$$

Revenons à la question de la signification réelle d'une valeur propre complexe. Soit f un endomorphisme d'un espace vectoriel réel V de dimension finie et soit $\lambda = \alpha + i\beta$ une valeur propre complexe (i.e. $\chi_f(\lambda) = 0$). Par la proposition précédente, on déduit que λ est une valeur propre du complexifié $f_{\mathbb{C}} \in \mathcal{L}(V_{\mathbb{C}})$, il existe donc un vecteur $w = u + iv \in V_{\mathbb{C}}$ tel que $f_{\mathbb{C}}(w) = \lambda w$. En prenant les parties réelles et imaginaires, on a donc

$$f(u) + if(v) = f(u + iv) = (\alpha + i\beta)(u + iv) = (\alpha u - \beta v) + i(\beta u + \alpha v).$$

On a alors le résultat suivant :

Théorème 9.14.4. *Si V est un espace vectoriel réel de dimension finie et si $\lambda = \alpha + i\beta$ est une valeur propre complexe ($\beta \neq 0$), alors il existe deux vecteurs $u, v \in V$ linéairement indépendants tels que*

$$\begin{cases} f(u) &= \alpha u - \beta v, \\ f(v) &= \beta u + \alpha v. \end{cases}$$

En particulier le sous-espace U de V engendré par u et v est de dimension 2 et il est invariant par f .

Preuve. Il ne reste qu'à démontrer l'indépendance linéaire de u et v . Observons d'abord que $u \neq 0$. En effet, si on avait $u = 0$, alors $\beta v = \alpha u - f(u) = 0$ par la première équation. Or nous supposons que $\beta \neq 0$, donc $v = 0$ ce qui contredit l'hypothèse que $w = u + iv \neq 0$. Pour montrer l'indépendance linéaire de u et v , on suppose maintenant par l'absurde qu'il existe $\gamma \in \mathbb{R}$ tel que $v = \gamma u$, alors

$$(\beta + \alpha\gamma)u = \beta u + \alpha v = f(v) = f(\gamma u) = \gamma f(u) = \gamma(\alpha u - \beta v) = \gamma(\alpha - \beta\gamma)u.$$

Cela implique que $\beta u = -\gamma^2 \beta u$, et puisque $u \neq 0$ et $\beta \neq 0$, on en déduit que $\gamma^2 = -1$. Mais ceci est impossible puisque $\gamma \in \mathbb{R}$. □

Remarque. Le théorème reste vrai pour un espace vectoriel de dimension infinie à condition que $\lambda \in \mathbb{C}$ soit une valeur propre complexe du complexifié de l'endomorphisme considéré.

Exemple. Le complexifié de $V = C^\infty(\mathbb{R}, \mathbb{R})$ est l'espace vectoriel $V_{\mathbb{C}} = C^\infty(\mathbb{R}, \mathbb{C})$ des fonctions infiniment différentiables sur $\mathbb{R}$ à valeurs dans $\mathbb{C}$, et l'opérateur de dérivation $D = \frac{d}{dx}$ s'étend naturellement à $C^\infty(\mathbb{R}, \mathbb{C})$. Tout nombre complexe $\lambda \in \mathbb{C}$ est valeur propre car

$$D(e^{\lambda x}) = \lambda e^{\lambda x}.$$

Si $\lambda = \alpha + i\beta$, avec $\alpha, \beta \in \mathbb{R}$ et $\beta \neq 0$, alors

$$e^{\lambda x} = e^{(\alpha+i\beta)x} = e^{\alpha x} e^{i\beta x} = e^{\alpha x} (\cos(\beta x) + i \sin(\beta x)).$$

Posons

$$\begin{aligned}\varphi(x) &= \operatorname{Ré}(e^{\lambda x}) = e^{\alpha x} \cos(\beta x) \\ \psi(x) &= \operatorname{Im}(e^{\lambda x}) = e^{\alpha x} \sin(\beta x).\end{aligned}$$

Alors on a bien

$$\begin{aligned}D(\varphi(x)) &= \varphi'(x) = D(e^{\alpha x} \cos(\beta x)) = \alpha e^{\alpha x} \cos(\beta x) - \beta e^{\alpha x} \sin(\beta x) \\ &= \alpha \varphi(x) - \beta \psi(x),\end{aligned}$$

et

$$\begin{aligned}D(\psi(x)) &= \psi'(x) = D(e^{\alpha x} \sin(\beta x)) = \beta e^{\alpha x} \cos(\beta x) + \alpha e^{\alpha x} \sin(\beta x) \\ &= \alpha \psi(x) + \beta \varphi(x).\end{aligned}$$

Remarque. Si dans le théorème précédent on suppose $\dim(V) = 2$, alors $\{u, v\}$ est une base de V et on vérifie très simplement que la matrice de f dans cette base est la matrice

$$K(\alpha, \beta) = \begin{pmatrix} \alpha & \beta \\ -\beta & \alpha \end{pmatrix}.$$

Nous avons donc le résultat suivant

Théorème 9.14.5. *Tout endomorphisme f d'un espace vectoriel réel V de dimension 2 admet une base $\mathcal{B}$ dans laquelle la matrice de f prend l'une des trois formes suivantes :*

$$\text{Diag}(\alpha, \beta) = \begin{pmatrix} \alpha & 0 \\ 0 & \beta \end{pmatrix}, \quad J_2(\alpha) = \begin{pmatrix} \alpha & 1 \\ 0 & \alpha \end{pmatrix} \quad \text{ou} \quad K(\alpha, \beta) = \begin{pmatrix} \alpha & \beta \\ -\beta & \alpha \end{pmatrix}. \quad (9.16)$$

avec $\alpha, \beta \in \mathbb{R}$ (et $\beta \neq 0$ dans le troisième cas).

Remarques. (i) On ne suppose pas que $\beta \neq \alpha$. Lorsque $\alpha = \beta$ le premier et le deuxième cas se distinguent par le polynôme minimal ($\mu_f(t) = (t - \alpha)$ dans le premier cas et $\mu_f(t) = (t - \alpha)^2$ dans le deuxième cas).

(ii) Dans le troisième cas le polynôme caractéristique est $\chi_f(t) = (t - \alpha)^2 + \beta^2$ et il n'a pas de racine réelle car on suppose $\beta \neq 0$.

Preuve. Si le polynôme caractéristique $\chi_f(t)$ est scindé, alors il existe une base de Jordan et nous sommes dans le premier ou le second cas. Si $\chi_f(t)$ n'est pas scindé comme polynôme à coefficients réels, alors il existe une paire de valeurs propres complexes conjuguées et le résultat est démontré dans le théorème précédent. □

Corollaire 9.14.6. *Toute matrice $A \in M_2(\mathbb{R})$ est semblable à l'une des matrices de type (9.16), i.e. il existe $P \in GL_2(\mathbb{R})$ telle que $P^{-1}AP$ est l'une des matrices de (9.16).*

En dimension 3 nous avons un résultat semblable :

Théorème 9.14.7. *Tout endomorphisme f d'un espace vectoriel réel V de dimension 3 admet une base $\mathcal{B}$ dans laquelle la matrice de f prend l'une des quatre formes suivante :*

$$\text{Diag}(\alpha, \beta, \gamma) = \begin{pmatrix} \alpha & 0 & 0 \\ 0 & \beta & 0 \\ 0 & 0 & \gamma \end{pmatrix}, \quad J_2(\alpha) \oplus J_1(\gamma) = \begin{pmatrix} \alpha & 1 & 0 \\ 0 & \alpha & 0 \\ 0 & 0 & \gamma \end{pmatrix}, \quad J_3(\alpha) = \begin{pmatrix} \alpha & 1 & 0 \\ 0 & \alpha & 1 \\ 0 & 0 & \alpha \end{pmatrix},$$

ou

$$K(\alpha, \beta) \oplus J_1(\gamma) = \begin{pmatrix} \alpha & \beta & 0 \\ -\beta & \alpha & 0 \\ 0 & 0 & \gamma \end{pmatrix},$$

avec $\alpha, \beta, \gamma \in \mathbb{R}$ (et $\beta \neq 0$ dans le dernier cas).

Ce théorème peut naturellement se reformuler en termes de matrices $A \in M_3(\mathbb{R})$.

Preuve. Si le polynôme caractéristique $\chi_f(t)$ est scindé (comme polynôme à coefficients réel), alors il existe une base de Jordan et la matrice de f dans cette base possède un, deux ou trois blocs de Jordan, ce qui nous donne une des trois premières matrices. Si $\chi_f(t)$ n'est pas scindé sur les réels, alors il admet une racine réelle que nous notons γ et deux racines complexes conjuguées $\alpha \pm i\beta$.

Il existe donc une paire de vecteurs $u, v \in V$ vérifiant les équations du Théorème 9.14.4. Il existe aussi un vecteur propre $z \in V$ pour la valeur propre γ . On vérifie que les vecteurs $u, v, z \in V$ sont nécessairement linéairement indépendants. Ils forment donc une base de V et dans cette base la matrice de f est $K(\alpha, \beta) \oplus J_1(\gamma)$. □

Ce théorème admet la généralisation suivante :

Théorème 9.14.8. Soit f un endomorphisme d'un espace vectoriel réel V de dimension finie. Notons γ_i ses valeurs propres réelles ($i = 1, \dots, r$) et $\alpha_j \pm \sqrt{-1}\beta_j$ ses paires de valeurs propres complexes conjuguées, $j = 1, \dots, s$ (on suppose $\beta_j \neq 0$). On suppose que les valeurs propres complexes sont deux-à-deux distinctes. Alors il existe une base de V pour laquelle la matrice de f prend la forme

$$\tilde{J} \oplus K(\alpha_1, \beta_1) \oplus \dots \oplus K(\alpha_s, \beta_s),$$

où $\tilde{J}$ est une matrice générale de Jordan de valeurs propres $\gamma_1, \dots, \gamma_j$.

Généralisation. Mentionnons pour terminer que si les valeurs propres complexes ont des multiplicités ≥ 1 , alors on peut encore trouver une base dans laquelle la matrice de f prend une forme standard, que nous expliquons maintenant.

Si V est un espace vectoriel de dimension 4 et $f \in \mathcal{L}(V)$ est un endomorphisme dont le polynôme caractéristique s'écrit

$$\chi_f(t) = (t - \lambda)^2(t - \bar{\lambda})^2 = ((t - \alpha)^2 + \beta^2)^2,$$

alors il existe une base de V dans laquelle la matrice de f prend l'une des deux formes suivantes :

$$K(\alpha, \beta) \oplus K(\alpha, \beta) = \begin{pmatrix} \alpha & \beta & 0 & 0 \\ -\beta & \alpha & 0 & 0 \\ 0 & 0 & \alpha & \beta \\ 0 & 0 & -\beta & \alpha \end{pmatrix}$$

ou

$$K_2(\alpha, \beta) = \begin{pmatrix} \alpha & \beta & 1 & 0 \\ -\beta & \alpha & 0 & 1 \\ 0 & 0 & \alpha & \beta \\ 0 & 0 & -\beta & \alpha \end{pmatrix}.$$

La première matrice est la matrice diagonale par blocs $K(\alpha, \beta) \oplus K(\alpha, \beta)$ et la seconde matrice possède les mêmes blocs en diagonale et une 2×2 matrice identité I_2 en sur-diagonale. Les deux matrices ont même polynôme caractéristique $\chi_f(t) = ((t - \alpha)^2 + \beta^2)^2$. Le polynôme minimal de la première matrice est $\mu(t) = (t - \alpha)^2 + \beta^2$ et le polynôme minimal de la seconde matrice est $\mu(t) = ((t - \alpha)^2 + \beta^2)^2$.

Cette structure se généralise en toute dimension : tout endomorphisme d'un espace vectoriel réel admet une base dans laquelle sa matrice est un produit direct de blocs de Jordan et de matrices du type

$$K_{2m}(\alpha, \beta) = \begin{pmatrix} K(\alpha, \beta) & I_2 & & \\ 0_2 & K(\alpha, \beta) & & \\ & & \ddots & \\ & & & K(\alpha, \beta) & I_2 \\ & & & 0_2 & K(\alpha, \beta) \end{pmatrix}$$

La matrice $K_{2m}(\alpha, \beta)$ est la matrice réelle de taille $4m \times 4m$ formée de m blocs qui sont des matrices $K(\alpha, \beta)$ et $(m - 1)$ blocs I_2 en surdiagonale. Son polynôme minimal est égale à son polynôme caractéristique :

$$\chi(t) = \mu(t) = ((t - \alpha)^2 + \beta^2)^m.$$

Annexe : Qui est Jordan ?

Le nom "Jordan" en algèbre peut faire référence à l'une des trois personnes suivantes :

- Wilhelm Jordan (1842– 1899), géographe allemand. La méthode d'échelonnage systématique pour les systèmes linéaire lui est attribuée, il voyait cet algorithme comme un raffinement de la méthode d'élimination qu'il attribuait à Gauss. Toutefois il semble que la méthode de Gauss-Jordan a été d'abord découverte par le mathématicien Belge Clasen en 1888.
- Camille Jordan (1838–1922), mathématicien français, professeur à l'École Polytechnique de Paris. Les formes canoniques de Jordan apparaissent dans son livre *Traité des substitutions et des équations algébriques* (1870).
- Pascual Jordan (1902–1980), physicien et mathématicien allemand. On lui doit d'importantes contributions en mécanique quantique et en théorie quantique des champs. Son nom est aussi attaché aux algèbres de Jordan, qui sont une classe d'algèbres non associatives qui sont utilisées dans la formalisation des observables en mécanique quantique.

Chapitre 10

Espace dual et Formes Bilinéaires

10.1 Espace Dual

Définitions. Le *dual* d'un K -espace vectoriel V est l'espace vectoriel des applications linéaires définies sur V à valeurs dans le corps K . On le note

$$V^* = \mathcal{L}(V, K).$$

Un élément de V^* s'appelle un *covecteur* de V ou une *forme linéaire* sur V .

Exemples. 1.) Toute forme linéaire sur K^n est une application $\varphi : K^n \rightarrow K$ qui peut s'écrire

$$\varphi(x) = \varphi(x_1, \dots, x_n) = a_1x_1 + \dots + a_nx_n, \quad (a_i \in K).$$

2.) La *trace* définit une forme linéaire sur l'espace vectoriel $M_n(K)$ des matrices carrées de taille $n \times n$ sur K .

3.) Si $V = C^0([a, b])$ est l'espace des fonctions continues sur l'intervalle $[a, b]$ et $x_0 \in [a, b]$, alors l'évaluation en x_0 définit une forme linéaire sur V^* . On la note

$$\delta_{x_0} : C^0([a, b]) \rightarrow \mathbb{R}, \quad \delta_{x_0}(g) = g(x_0).$$

On dit parfois que le covecteur δ_{x_0} est la *masse de Dirac* concentrée au point x_0 .

4.) Une autre forme linéaire sur $C^0([a, b])$ est l'intégration :

$$I_{[a, b]} : C^0([a, b]) \rightarrow \mathbb{R}, \quad I_{[a, b]}(g) = \int_a^b g(x)dx.$$

Proposition 10.1.1. *Tout espace vectoriel de dimension finie est isomorphe à son dual.*

Preuve. Soit V un espace vectoriel de dimension finie sur le corps K et V^* son dual. Alors Les espaces V et V^* ont même dimension car

$$\dim(V^*) = \dim(\mathcal{L}(V, K)) = \dim(V) \cdot \dim(K) = \dim(V).$$

Ces deux espaces vectoriels sont donc isomorphes.

□

Remarque. Observons que cette démonstration ne repose pas sur un argument direct mais elle utilise la théorie de la dimension. Pour un espace vectoriel de dimension infinie l'argument ne marche pas et on peut démontrer qu'*un espace vectoriel de dimension infinie n'est jamais isomorphe à son dual*.

La proposition suivante complète la précédente.

Proposition 10.1.2. Soit $\mathcal{B} = \{v_1, \dots, v_n\}$ une base de l'espace vectoriel V . Notons $\varphi_i \in V^*$ le covecteur défini par

$$\varphi_i(v_j) = \delta_{ij} = \begin{cases} 1, & \text{si } i = j, \\ 0, & \text{si } i \neq j. \end{cases} \quad (10.1)$$

Alors $\mathcal{B}^* = \{\varphi_1, \dots, \varphi_n\}$ est une base de V^* .

Preuve. On remarque que

$$\text{Card}(\mathcal{B}^*) = \text{Card}(\mathcal{B}) = n = \dim(V^*),$$

il suffit de prouver que $\mathcal{B}^*$ est une famille libre. Supposons que $\sum_{i=1}^n \lambda_i \varphi_i = 0$, alors on a pour tout j

$$0 = \left(\sum_{i=1}^n \lambda_i \varphi_i \right) (v_j) = \sum_{i=1}^n \lambda_i \varphi_i(v_j) = \sum_{i=1}^n \lambda_i \delta_{ij} = \lambda_j.$$

Donc $\lambda_j = 0$ pour tout j . On a montré que $\mathcal{B}^*$ est une famille libre de V^* , et c'est donc une base puisque son cardinal est égal à la dimension de V . □

Définition. La base $\mathcal{B}^*$ s'appelle la *base duale* de $\mathcal{B}$. On note parfois v_i^* le covecteur φ_i et on dit que v_i^* est le covecteur dual (ou la forme linéaire duale) au vecteur de base v_i . L'équation (10.1) s'appelle la *relation de dualité* entre les deux bases.

Lorsque $\{e_1, \dots, e_n\}$ est la base canonique de K^n , la base duale est notée $\{\varepsilon_1, \dots, \varepsilon_n\} \subset (K^n)^*$ et la relation de dualité s'écrit

$$\varepsilon_i(e_j) = \delta_{ij}.$$

Problème. Pour illustrer ces notions considérons le problème suivant : Soit $\{v_1, \dots, v_n\}$ une base quelconque de $\mathbb{R}^n$ et notons $\theta_1, \dots, \theta_n \subset (\mathbb{R}^n)^*$ la base duale. On demande de déterminer la matrice de transition de la base duale canonique $\{\varepsilon_i\}$ vers la base $\{\theta_i\}$ à partir de la matrice de transition de la base canonique $\{e_i\}$ vers la base $\{v_i\}$.

Solution. Notons P la matrice de transition de la base $\{e_i\}$ vers la base $\{v_i\}$ et P' la matrice de transition de la base $\{\varepsilon_i\}$ vers la base $\{\theta_i\}$. Rappelons que par définition

$$v_j = \sum_{k=1}^n p_{kj} e_k, \quad \text{et} \quad \theta_i = \sum_{l=1}^n p'_{li} \varepsilon_l.$$

Par définition de ε_l on a donc

$$\theta_i(e_k) = \sum_{l=1}^n p'_{li} \varepsilon_l(e_k) = \sum_{l=1}^n p'_{li} \delta_{lk} = p'_{ki},$$

et avec la relation de dualité entre les bases $\{\theta_i\}$ et $\{v_j\}$, nous obtenons.

$$\delta_{ij} = \theta_i(v_j) = \theta_i\left(\sum_{k=1}^n p_{kj}e_k\right) = \sum_{k=1}^n p_{kj}\theta_i(e_k) = \sum_{k=1}^n p_{kj}p'_{ki}.$$

Cette relation s'écrit matriciellement $P^\top P' = I_n$, la matrice cherchée P' est donc la matrice inverse de la transposée de P .

$$P' = (P^\top)^{-1} = (P^{-1})^\top.$$

Cette matrice s'appelle la *matrice contragrédiente* de P . (matrice)

Exemple. Pour trouver la base duale $\mathcal{B}^* \subset (K^2)^*$ de la base $\mathcal{B} = \{v, w\} = \{(2, 1), (-1, 1)\} \subset K^2$ on cherche la matrice de changement de base P et sa contragrédiente P' :

$$P = \begin{pmatrix} 2 & -1 \\ 1 & 1 \end{pmatrix} \quad \text{et} \quad P' = (P^{-1})^\top = \frac{1}{3} \begin{pmatrix} 1 & -1 \\ 1 & 2 \end{pmatrix}.$$

Notons cette base duale $\mathcal{B}^* = \{\varphi, \psi\}$, alors ces covecteurs sont donnés par

$$\varphi = \frac{1}{3}(\varepsilon_1 + \varepsilon_2) \quad \text{et} \quad \psi(x, y) = \frac{1}{3}(-\varepsilon_1 + 2\varepsilon_2).$$

Ces covecteurs sont donc les fonctions $K^2 \rightarrow K$ telles que

$$\varphi(x, y) = \frac{1}{3}(x + y) \quad \text{et} \quad \psi(x, y) = -\frac{1}{3}x + \frac{2}{3}y.$$

On vérifie facilement que $\varphi(v) = \psi(w) = 1$ et $\varphi(w) = \psi(v) = 0$.

Proposition 10.1.3. Soit $\mathcal{B} = \{v_1, \dots, v_n\}$ une base de l'espace vectoriel V et $\mathcal{B}^* = \{\varphi_1, \dots, \varphi_n\}$ la base duale de V^* , alors on a les propriétés suivantes :

(a) Tout vecteur $x \in V$ s'écrit

$$x = \sum_{i=1}^n \varphi_i(x)v_i.$$

(b) Tout covecteur $\psi \in V^*$ s'écrit

$$\psi = \sum_{i=1}^n \psi(v_i)\varphi_i.$$

Preuve. (a) Développons le vecteur x dans la base $\mathcal{B}$, on a $x = \sum_{j=1}^n x_j v_j$, donc

$$\varphi_i(x) = \varphi_i\left(\sum_{j=1}^n x_j v_j\right) = \sum_{j=1}^n x_j \varphi_i(v_j) = \sum_{j=1}^n x_j \delta_{ij} = x_i,$$

par conséquent $x = \sum_{i=1}^n x_i v_i = \sum_{i=1}^n \varphi_i(x)v_i$.

(b) Développons le covecteur ψ dans la base duale $\mathcal{B}^*$, on a $\psi = \sum_{j=1}^n \xi_j \varphi_j$, donc

$$\psi(v_i) = \left(\sum_{j=1}^n \xi_j \varphi_j\right)(v_i) = \sum_{j=1}^n \xi_j \varphi_j(v_i) = \sum_{j=1}^n \xi_j \delta_{ij} = \xi_i,$$

et donc $\psi = \sum_{i=1}^n \xi_i \varphi_i = \sum_{i=1}^n \psi(v_i) \varphi_i$.

□

Définition. Soient V, W deux espaces vectoriels sur le corps K et $f : V \rightarrow W$ une application linéaire. L'application duale $f^* : W^* \rightarrow V^*$ est l'application linéaire définie par

$$f^*(\psi) = \psi \circ f.$$

Ainsi si $v \in V$ et $\psi \in W^*$, alors $f^*(\psi)(v) = \psi(f(v)) \in K$. Il est facile de vérifier que f^* est linéaire :

$$f^*(\alpha_1 \psi_1 + \alpha_2 \psi_2) = \alpha_1 f^*(\psi_1) + \alpha_2 f^*(\psi_2).$$

Théorème 10.1.4. Soit $f \in \mathcal{L}(V, W)$ une application linéaire entre les espaces vectoriels V et W . Donnons-nous des bases $\mathcal{B} = \{v_1, \dots, v_n\}$ et $\mathcal{B}' = \{w_1, \dots, w_m\}$ de V et W respectivement. Si A est la matrice de f dans ces bases, alors la matrice de l'application duale $f^* \in \mathcal{L}(W^*, V^*)$ dans les bases duales $(\mathcal{B}')^*$ et $\mathcal{B}^*$ est la matrice transposée $A^\top$ de A :

Remarque. Observons que A est une matrice de taille $m \times n$ ($m = \dim W$ et $n = \dim V$) et $A^\top$ est une matrice de taille $n \times m$, ce qui est compatible avec le fait que f est une application de W^* dans V^* et $n = \dim V^*$ $m = \dim W^*$.

Preuve. Notons $\mathcal{B}^* = \{\varphi_1, \dots, \varphi_n\} \subset V^*$ la base duale de $\mathcal{B}$ et $\mathcal{B}'^* = \{\psi_1, \dots, \psi_m\} \subset W^*$ la base duale de $\mathcal{B}'$. Rappelons que la matrice $A = (a_{ij})$ de f dans les bases $\mathcal{B}, \mathcal{B}'$ est définie par la relation

$$f(v_j) = \sum_{i=1}^m a_{ij} w_i.$$

Calculons $f^*(\psi_i)(v_j)$ en utilisant la définition de f^* et la propriété $\psi_i(w_k) = \delta_{ik}$:

$$f^*(\psi_i)(v_j) = \psi_i(f(v_j)) = \psi_i\left(\sum_{k=1}^m a_{kj} w_k\right) = \sum_{k=1}^m a_{kj} \psi_i(w_k) = \sum_{k=1}^m a_{kj} \delta_{ik} = a_{ij}.$$

La proposition précédente implique alors

$$f^*(\psi_i) = \sum_{j=1}^n f^*(\psi_i)(v_j) \varphi_j = \sum_{j=1}^n a_{ij} \varphi_j = \sum_{j=1}^n (a_{ji})^\top \varphi_j.$$

Ce qui montre que la matrice de f^* dans les bases duales est la matrice $A^\top$.

□

Corollaire 10.1.5. Si V et W sont des espaces vectoriels de dimension finie, alors pour tout $f \in \mathcal{L}(V, W)$ on a $\text{rang}(f^*) = \text{rang}(f)$.

Preuve. Choisissons des bases $\mathcal{B}$ et $\mathcal{B}'$ de V et W , et notons $A = M_{\mathcal{B}'\mathcal{B}}(f)$ la matrice de f dans ces bases. Alors $A^\top$ est la matrice de f^* dans les bases duales et on a donc

$$\text{rang}(f^*) = \text{rang}(A^\top) = \text{rang}(A) = \text{rang}(f).$$

□

10.1.1 Interpolation de Lagrange

Le *problème de l'interpolation* est le suivant : Soit $\varphi : \mathbb{R} \rightarrow \mathbb{R}$ une fonction quelconque. La fonction n'est pas connue mais on dispose d'un nombre fini de mesures (ou d'observations) qui nous donnent un nombre fini de valeurs, disons

$$\varphi(a_1) = b_1, \varphi(a_2) = b_2, \dots, \varphi(a_n) = b_n, \quad (10.2)$$

et on voudrait, à partir de cette information, reconstruire la fonction φ . Ce problème n'a pas de solution unique en général, mais nous allons montrer qu'il est uniquement résoluble dans l'espace vectoriel des polynômes de degrés $\leq n - 1$.

Pour résoudre ce problème, il est utile de considérer l'espace dual de l'espace vectoriel des polynômes. Rappelons que le *covecteur d'évaluation en a* est la forme linéaire $\delta_a : \mathbb{R}[x] \rightarrow \mathbb{R}$ définie par

$$\delta_a(p) = p(a),$$

pour tout polynôme $p(x)$. Pour résoudre le problème de l'interpolation dans l'espace $\mathcal{P}_{n-1} \subset \mathbb{R}[x]$ des polynômes de degré $\leq n - 1$, on cherche d'abord n polynômes $\varphi_1, \dots, \varphi_n \in \mathcal{P}_{n-1}$ tels que

$$\delta_{a_j}(\varphi_i) = \varphi_i(a_j) = \delta_{ij}.$$

Supposons que ces polynômes ont été construits, alors la solution du problème (10.2) est clairement donnée par

$$\varphi(x) = \sum_{i=1}^n b_i \varphi_i(x),$$

en effet on a

$$\varphi(a_j) = \sum_{i=1}^n b_i \varphi_i(a_j) = \sum_{i=1}^n b_i \delta_{ij} = b_j.$$

Or la construction des polynômes $\varphi_i \in \mathcal{P}_{n-1}$ est élémentaire, il suffit de poser

$$\varphi_i(x) = \prod_{j \neq i} \frac{x - a_j}{a_i - a_j}.$$

Par conséquent la solution du problème d'interpolation (10.2) est donnée explicitement par la formule

$$\varphi(x) = \sum_{i=1}^n b_i \prod_{j \neq i} \frac{x - a_j}{a_i - a_j}.$$

Cette formule s'appelle la *formule d'interpolation de Lagrange*. Par exemple si $n = 3$ cette formule d'écrit

$$\varphi(x) = \frac{b_1(x - a_2)(x - a_3)}{(a_1 - a_2)(a_1 - a_3)} + \frac{b_2(x - a_1)(x - a_3)}{(a_2 - a_1)(a_2 - a_3)} + \frac{b_3(x - a_1)(x - a_2)}{(a_3 - a_1)(a_3 - a_2)}.$$

Remarque. La construction de Lagrange met en évidence le fait que les formes linéaires $\{\delta_{a_1}, \dots, \delta_{a_n}\}$ forment une base de $\mathcal{P}_{n-1}^*$, et que $\{\varphi_1, \dots, \varphi_n\} \subset \mathcal{P}_{n-1}$ est la base duale.

10.2 Couplage entre deux espaces vectoriels

Définition. Un *couplage*¹ entre deux espaces vectoriels V et W sur un corps K est une application :

$$\beta : V \times W \rightarrow K,$$

qui est bilinéaire, c'est-à-dire linéaire en chaque variable :

$$\beta(x, \mu_1 y_1 + \mu_2 y_2) = \mu_1 \beta(x, y_1) + \mu_2 \beta(x, y_2), \quad \beta(\lambda_1 x_1 + \lambda_2 x_2, y) = \lambda_1 \beta(x_1, y) + \lambda_2 \beta(x_2, y).$$

La condition de bilinéarité peut également s'écrire :

$$\beta \left(\sum_{i=1}^m \lambda_i x_i, \sum_{j=1}^n \mu_j y_j \right) = \sum_{i=1}^m \sum_{j=1}^n \lambda_i \mu_j \beta(x_i, y_j).$$

Exemples 1. L'intégration définit un couplage entre $V = \mathbb{R}[x]$ et $W = C^0([a, b])$:

$$I : V \times W \rightarrow \mathbb{R}, \quad I(p, f) = \int_a^b p(x) f(x) dx.$$

2. On note ℓ_1 l'espace vectoriel des *suites réelles absolument sommables*, qui est défini par

$$\ell_1 = \{ \xi = (x_k)_{k=1}^\infty \mid x_k \in \mathbb{R}, \sum_{k=1}^\infty |x_k| < \infty \}.$$

On note aussi ℓ_∞ l'espace vectoriel des suites réelles bornées

$$\ell_\infty = \{ \xi = (x_k)_{k=1}^\infty \mid x_k \in \mathbb{R}, \sup_k |x_k| < \infty \}.$$

Alors le *couplage de sommation* est défini par

$$\sigma : \ell_1 \times \ell_\infty \rightarrow \mathbb{R}, \quad \sigma(\xi, \eta) = \sum_{k=1}^\infty x_k y_k.$$

3. L'exemple qui suit est une variante de l'exemple précédent. On fixe $p \in (0, \infty)$ et on note ℓ_p l'espace vectoriel des *suites p -sommables* :

$$\ell_p = \{ \xi = (x_k)_{k=1}^\infty \mid x_k \in \mathbb{R}, \sum_{k=1}^\infty |x_k|^p < \infty \}.$$

Alors le couplage de sommation $\sigma : \ell_p \times \ell_q \rightarrow \mathbb{R}$ est bien défini à condition que $1/p + 1/q = 1$ (c'est une conséquence de l'*inégalité de Hölder*, démontrée au cours d'analyse).

4. Un couplage entre les espaces de matrices $M_{m,n}(K)$ et $M_{n,m}(K)$ est défini par la trace du produit matriciel :

$$\begin{array}{ccc} M_{m,n}(K) \times M_{n,m}(K) & \rightarrow & K \\ (A, B) & \mapsto & \text{Trace}(A \cdot B) \end{array}$$

1. 'couplage' se dit 'pairing' en anglais.

5. Tout espace vectoriel V admet un couplage avec son dual :

$$\begin{aligned} V^* \times V &\rightarrow K \\ (\phi, v) &\mapsto \phi(v). \end{aligned}$$

On l'appelle le *couplage canonique de V avec son dual*. Ce couplage est universel (il est défini pour tout espace vectoriel) et ne dépend pas du choix d'une base, ni d'aucun autre choix.

Couplage et dualité

A tout couplage $\beta : V \times W \rightarrow K$ entre deux K -espaces vectoriels on peut associer une application linéaire entre chacun des espaces vectoriels et le dual de l'autre. Ces applications linéaires

$$\beta_g : V \rightarrow W^* \quad \text{et} \quad \beta_d : W \rightarrow V^*$$

sont définies de la façon suivante :

$$\beta_g(v) \in W^* \text{ est le covecteur tel que } \beta_g(v)(w) = \beta(v, w) \text{ pour tout } w \in W.$$

De même

$$\beta_d(w) \in V^* \text{ est le covecteur tel que } \beta_d(w)(v) = \beta(v, w) \text{ pour tout } v \in V.$$

Les lettres 'g' et 'd' signifient que β_g agit sur la variable de gauche et β_d agit sur la variable de droite.

Définition 10.2.1. Le couplage $\beta : V \times W \rightarrow K$ est *non dégénéré* si

$$\forall v \in V, \text{ on a } [(\beta(v, y) = 0 \ \forall y \in W) \Leftrightarrow v = 0]$$

et

$$\forall w \in W, \text{ on a } [(\beta(x, w) = 0 \ \forall x \in V) \Leftrightarrow w = 0].$$

Lemme 10.2.2. $\beta : V \times W \rightarrow K$ est non dégénéré si et seulement si β_g et β_d sont injectives.

Preuve. La première condition de la définition précédente dit exactement que $\text{Ker}(\beta_g) = \{0\}$ et la deuxième condition dit que $\text{Ker}(\beta_d) = \{0\}$. □

Corollaire 10.2.3. Soit $\beta : V \times W \rightarrow K$ un couplage entre deux espaces vectoriels de dimension finies. Alors β est non dégénéré si et seulement si β_g et β_d sont des isomorphismes.

Preuve. Observons d'abord que si β_g est injective, alors $\dim(V) \leq \dim(W^*) = \dim(W)$ et si β_d est injective, alors $\dim(W) \leq \dim(V^*) = \dim(V)$, par conséquent, si β est non dégénéré, alors $\dim(V) = \dim(W)$. On conclut la preuve en rappelant qu'une application linéaire injective entre deux espaces vectoriels de même dimension finie est un isomorphisme. □

La notation “ $\langle \text{bra} | \text{ket} \rangle$ ” de Dirac

Si $\beta : V \times W \rightarrow K$ est un couplage entre deux K -espaces vectoriels, il est commode de noter

$$\langle v | w \rangle_\beta = \beta(v, w),$$

ou simplement $\langle v | w \rangle$ lorsque le couplage β a été fixé. Les homomorphismes définis précédemment peuvent alors s’écrire de façon plus concise

$$\beta_g(v) = \langle v | \cdot \rangle \in W^* \quad \text{et} \quad \beta_d(w) = \langle \cdot | w \rangle \in V^*.$$

Le vecteur $v \in V$ est alors vu comme un covecteur de W (i.e. un élément de W^*) et $w \in W$ est vu comme un covecteur de V (i.e. un élément de V^*). En mécanique quantique on utilise souvent la variante suivante de cette notation :

$$\beta_g(v) = \langle v | \quad \text{et} \quad \beta_d(w) = | w \rangle.$$

10.3 Formes bilinéaires sur un espace vectoriel

Définition. Une *forme bilinéaire* sur un K -espace vectoriel V est une application

$$g : V \times V \rightarrow K$$

qui est bilinéaire. Une forme bilinéaire est donc un couplage de V avec lui même. La bilinéarité signifie que g est linéaire en chacune de ses deux variables, ce qu’on peut aussi écrire sous la forme

$$g \left(\sum_{i=1}^m \lambda_i x_i, \sum_{j=1}^n \mu_j y_j \right) = \sum_{i,j=1}^m \lambda_i \mu_j g(x_i, y_j).$$

Exemples 1. Le produit scalaire standard sur $\mathbb{R}^n$ défini par

$$x \cdot y = x_1 y_1 + \cdots + x_n y_n$$

est une forme bilinéaire sur $\mathbb{R}^n$.

2. On définit une forme bilinéaire sur l’espace des $m \times n$ matrices sur le corps K par la formule

$$\begin{array}{ccc} M_{m,n}(K) \times M_{m,n}(K) & \rightarrow & K \\ (A, B) & \mapsto & \text{Trace}(A^\top \cdot B) \end{array}$$

3. Si $C \in M_n(K)$ est une matrice carrée quelconque, on peut lui associer une forme bilinéaire sur K^n définie par

$$g(x, y) = \sum_{i,j=1}^n c_{ij} x_i y_j.$$

Définition. Soit g une forme bilinéaire définie sur un espace vectoriel V de dimension $n < \infty$, et soit $\mathcal{B} = \{v_1, \dots, v_n\}$ une base de V . La matrice

$$G = (g_{ij}) \in M_n(K) \quad \text{définie par} \quad g_{ij} = g(v_i, v_j)$$

s'appelle la *matrice de Gram*² de g relativement à la base $\mathcal{B}$.

Exemple. La matrice de Gram de la forme bilinéaire g définie sur $\mathbb{R}^2$ par

$$g(x, y) = ax_1y_1 + bx_1y_2 + cx_2y_1 + dx_2y_2 \quad \text{est} \quad G = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$$

On fera attention à ne pas confondre une matrice de Gram avec la matrice d'un endomorphisme. Dans les deux cas il s'agit d'une matrice carrée, mais leur signification est très différente. L'interprétation de la matrice de Gram vient de la proposition suivante :

Proposition 10.3.1. Si $x = \sum_{i=1}^n x_i v_i$ et $y = \sum_{j=1}^n y_j v_j$, alors

$$g(x, y) = \sum_{i,j=1}^n g_{ij} x_i y_j. \quad (10.3)$$

La preuve est une application immédiate de la bilinéarité de g . □

Remarque. Si $\mathcal{B} = \{v_1, \dots, v_n\}$ est une base de V et si $X \in K^n$ et $Y \in K^n$ sont les vecteurs colonnes associés respectivement aux vecteurs $x = \sum_{i=1}^n x_i v_i \in V$ et $y = \sum_{j=1}^n y_j v_j \in V$:

$$X = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}, \quad Y = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix},$$

alors

$$g(x, y) = X^\top G Y = (x_1 \cdots x_n) \cdot G \cdot \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix}, \quad (10.4)$$

où G est la matrice de Gram de la forme bilinéaire g dans la base $\mathcal{B}$.

Corollaire 10.3.2. Si $\mathcal{B} = \{v_1, \dots, v_n\}$ et $\mathcal{B}' = \{v'_1, \dots, v'_n\}$ sont deux bases de V , et si P est la matrice de changement de base (i.e. $P = M_{\mathcal{B}\mathcal{B}'}(\text{Id}_V)$), alors les matrices de Gram de la forme bilinéaire g dans ces deux bases sont reliées par

$$G' = P^\top G P.$$

En particulier les matrices G et G' ont le même rang.

Définition. On appellera *rang* de la forme bilinéaire g le rang de sa matrice de Gram dans une base quelconque.

Preuve. Rappelons que si $x = \sum_{i=1}^n x_i v_i = \sum_{i=1}^n x'_i v'_i$, alors les vecteurs colonnes X et X' correspondants sont reliés par

$$X = P X'.$$

2. Jørgen Pedersen Gram, mathématicien danois 1850–1916.

Par conséquent nous avons d'une part

$$g(x, y) = X^\top GY = (PX')^\top G(PY') = (X'^\top P^\top)G(PY') = X'^\top (P^\top GP)Y';$$

et d'autre part $g(x, y) = X'^\top G'Y'$ pour tous $x, y \in K^n$. Ceci implique $G' = P^\top GP$. □

Définition. Deux matrices carrées G_1, G_2 sont dites *congruentes* s'il existe une matrice inversible P telle que

$$G_2 = P^\top G_1 P.$$

Exercice. Montrer que la relation de congruence est une relation d'équivalence.

Attention de ne pas confondre la relation de congruence $G \sim P^\top GP$ avec la relation de similitude $G \sim P^{-1}GP$. Deux matrices sont congruentes si et seulement si elles représentent la même forme bilinéaire dans des bases différentes alors que deux matrices sont semblables si et seulement si elles représentent le même endomorphisme dans des bases différentes.

Définition. Le *produit tensoriel* de deux co-vecteurs $\phi, \psi \in V^*$ est la forme bilinéaire $\phi \otimes \psi : V \times V \rightarrow K$ définie par la formule

$$(\phi \otimes \psi)(x, y) = \phi(x)\psi(y).$$

Proposition 10.3.3. Si $\mathcal{B} = \{v_1, \dots, v_n\}$ est une base de V et $\mathcal{B}^* = \{\varphi_1, \dots, \varphi_n\} \subset V^*$ est la base duale, alors toute forme bilinéaire $g : V \times V \rightarrow K$ s'écrit

$$g = \sum_{i,j=1}^n g_{ij} \varphi_i \otimes \varphi_j,$$

où $G = (g_{ij})$ est la matrice de Gram de g dans la base $\mathcal{B}$.

Preuve. Notons $h = \sum_{i,j=1}^n g_{ij} \varphi_i \otimes \varphi_j$. Il faut montrer que $h = g$. Or par définition du produit tensoriel on a

$$h(v_\mu, v_\nu) = \sum_{i,j=1}^n g_{ij} \varphi_i(v_\mu) \varphi_j(v_\nu) = g_{\mu\nu} = g(v_\mu, v_\nu),$$

car $\varphi_i(v_\mu) = \delta_{\mu,i}$ et $\varphi_j(v_\nu) = \delta_{\nu,j}$. Ceci montre que g et h coïncident sur la base $\mathcal{B}$, donc $g = h$ par bilinéarité. □

Corollaire 10.3.4. L'ensemble des formes bilinéaires sur un espace vectoriel V est un espace vectoriel de dimension n^2 (si $n = \dim V$) et

$$\{\varphi_i \otimes \varphi_j \mid 1 \leq i, j \leq n\}$$

est une base de cet espace vectoriel.

La proposition précédente nous dit que la matrice de Gram (g_{ij}) représente les composantes de la forme bilinéaire g dans la base $\{\varphi_i \otimes \varphi_j\}$.

10.4 Formes bilinéaires symétriques et antisymétriques

Soit V un espace vectoriel sur un corps K . On suppose, dans ce paragraphe et le suivant, que K n'est pas de caractéristique 2 (c'est-à-dire $1 + 1 \neq 0$ dans K).

Définition. Une forme bilinéaire $\alpha : V \times V \rightarrow K$ est dite *symétrique* si $\alpha(y, x) = \alpha(x, y)$ pour tous $x, y \in V$. Elle est *antisymétrique*³ si $\alpha(y, x) = -\alpha(x, y)$ pour tous $x, y \in V$.

On observe que toute forme bilinéaire α sur V s'écrit de façon unique comme somme d'une forme bilinéaire symétrique et d'une forme bilinéaire antisymétrique. On peut en effet écrire $\alpha(v, w) = \phi(v, w) + \psi(v, w)$ avec $\phi(v, w) := \frac{1}{2}(\alpha(v, w) + \alpha(w, v))$ et $\psi(v, w) := \frac{1}{2}(\alpha(v, w) - \alpha(w, v))$. Il est clair que ϕ est bilinéaire et symétrique et ψ est bilinéaire et antisymétrique.

Lorsque V est de dimension finie, on peut relier ces notions à la matrice de Gram. En effet la matrice de Gram G d'une forme bilinéaire dans une base quelconque est une matrice symétrique (i.e. $G^\top = G$) si et seulement si la forme bilinéaire est symétrique et elle est antisymétrique (i.e. $G^\top = -G$) si et seulement si la forme bilinéaire est antisymétrique.

Théorème 10.4.1. Soit β une forme bilinéaire symétrique sur un K -espace vectoriel V de dimension finie. Alors il existe une base $\{v_1, \dots, v_n\}$ de V telle que $\beta(v_i, v_j) = 0$ si $i \neq j$.

Définition. Une telle base est dite *orthogonale pour le forme bilinéaire β* (ou β -orthogonale).

Preuve. On raisonne par récurrence sur $n = \dim(V)$. Si $n = 1$, il n'y a rien à démontrer. Supposons donc que le théorème est démontré pour tout espace vectoriel de dimension $(n - 1)$, et soit $\beta : V \times V \rightarrow K$ une forme bilinéaire symétrique sur un espace vectoriel V de dimension n .

Nous affirmons d'abord que si $\beta(v, v) = 0$ pour tout $v \in V$, alors on a aussi $\beta(u, v) = 0$ pour tous $u, v \in V$. Cela découle par exemple du raisonnement suivant :

$$0 = \beta(u + v, u + v) = \beta(u, u) + \beta(u, v) + \beta(v, u) + \beta(v, v) = \beta(u, v) + \beta(v, u) = 2\beta(u, v),$$

donc $\beta(u, v) = 0$. Dans ce cas toute base est orthogonale. Supposons donc qu'il existe $v_1 \in V$ tel que $\beta(v_1, v_1) \neq 0$ et définissons

$$W := \{w \in V \mid \beta(v_1, w) = 0\}.$$

Alors W est un sous-espace vectoriel de V de dimension $n - 1$ (c'est le noyau du covecteur non nul $x \mapsto \beta(v_1, x)$). Par hypothèse de récurrence, il existe une base β -orthogonale de W , que nous notons $\{v_2, \dots, v_n\} \subset W$. Il est clair que $\beta(v_1, v_j) = 0$ pour tout $j = 2, \dots, n$ par définition de W et $\beta(v_i, v_j) = 0$ pour tous $i, j = 2, \dots, n$ par choix des vecteurs w_j . On a donc obtenu une base β -orthogonale $\{v_1, \dots, v_n\}$ de V . □

Remarque. Soit $\{v_1, \dots, v_n\}$ une base β -orthogonale de V et $x, y \in V$, alors

$$\beta(x, y) = \sum_{i=1}^n \alpha_i x_i y_i, \quad \text{avec } \alpha_i = \beta(v_i, v_i),$$

3. En anglais on dit *skew symmetric*.

où les x_i, y_i sont les composantes de x et y dans cette base. La matrice de Gram de β dans cette base est donc la matrice diagonale

$$B = (\beta(v_i, v_j)) = \begin{pmatrix} \alpha_1 & & 0 \\ & \ddots & \\ 0 & & \alpha_n \end{pmatrix}$$

Corollaire 10.4.2. *Soit β une forme bilinéaire symétrique sur un K -espace vectoriel V de dimension finie. Alors il existe des formes linéaires $\phi_1, \dots, \phi_r \in V^*$ linéairement indépendantes et des scalaires non nuls $\alpha_1, \dots, \alpha_r \in K$ tels que*

$$\beta = \sum_{i=1}^r \alpha_i \phi_i \otimes \phi_i, \quad \text{i.e.} \quad \beta(x, y) = \sum_{i=1}^r \alpha_i \phi_i(x) \phi_i(y) \quad \text{pour tous } x, y \in V.$$

De plus r est le rang de la matrice de Gram de β . Ce rang est donc indépendant de la base orthogonale choisie.

Preuve. Soit $\{v_1, \dots, v_n\}$ une base β -orthogonale de V et notons $\alpha_i = \beta(v_i, v_i)$. On a alors $\beta(v_i, v_j) = \alpha_i \delta_{ij}$. Soit maintenant $\{\phi_1, \dots, \phi_n\} \subset V^*$ la base duale de $\{v_1, \dots, v_n\}$, alors on a

$$\beta = \sum_{i=1}^n \alpha_i \phi_i \otimes \phi_i.$$

Il suffit en effet de vérifier cette égalité sur des vecteurs de bases :

$$\left(\sum_{i=1}^n \alpha_i \phi_i \otimes \phi_i \right) (v_j, v_k) = \sum_{i=1}^n \alpha_i \phi_i(v_j) \phi_i(v_k) = \sum_{i=1}^n \alpha_i \delta_{ij} \delta_{ik} = \alpha_j \delta_{jk} = \beta(v_j, v_k).$$

Quitte à réordonner les vecteurs de base, on peut supposer que $\alpha_i \neq 0$ si et seulement si $1 \leq i \leq r$, on a donc finalement

$$\beta = \sum_{i=1}^r \alpha_i \phi_i \otimes \phi_i.$$

□

10.5 Formes quadratiques

Dans ce paragraphe, on suppose que V est un espace vectoriel de dimension finie sur un corps K de caractéristique $\neq 2$.

Définition. Une *forme quadratique* sur V est une application $Q : V \rightarrow K$ pour laquelle il existe une forme bilinéaire symétrique $\beta : V \times V \rightarrow K$ telle que

$$Q(v) = \beta(v, v).$$

Lemme 10.5.1. *La forme bilinéaire symétrique β est déterminé par Q de façon unique. Plus précisément, on a la formule de polarisation :*

$$\beta(v, w) = \frac{1}{4}(Q(v + w) - Q(v - w)). \quad (10.5)$$

Remarque. Les formules suivantes permettent également de retrouver la forme bilinéaire symétrique β à partir de la forme quadratique Q :

$$\beta(v, w) = \frac{1}{2}(Q(v + w) - Q(v) - Q(w)) \quad (10.6)$$

$$\beta(v, w) = \frac{1}{2}(Q(v) + Q(w) - Q(v - w)) \quad (10.7)$$

Les formules (10.5), (10.6) et (10.7) s'appellent les *formules de polarisation*. De la forme quadratique Q .

Lorsque $V = K^n$, une forme quadratique Q sur $V = K^n$ s'écrit

$$Q(x) = \beta(x, x) = \sum_{i,j=1}^n b_{ij}x_i x_j, \quad (10.8)$$

où $b_{ij} = \beta(e_i, e_j)$. Ainsi une forme quadratique sur K^n n'est rien d'autre qu'un *polynôme homogène de degré 2 en n variables*.

Le corollaire 10.4.2 peut se reformuler pour les formes quadratiques de la façon suivante :

Théorème 10.5.2. *Soit Q une forme quadratique sur un K -espace vectoriel V de dimension finie. Alors il existe des formes linéaires $\phi_1, \dots, \phi_r \in V^*$ linéairement indépendantes et des scalaires $\alpha_1, \dots, \alpha_r \in K$ non nuls tels que*

$$Q = \sum_{i=1}^r \alpha_i \phi_i^2, \quad \text{i.e.} \quad Q(x) = \sum_{i=1}^r \alpha_i \phi_i(x)^2, \quad \text{pour tout } x \in V.$$

De plus l'entier r ne dépend que de la forme quadratique Q .

Définitions. 1.) L'entier r s'appelle le *rang* de la forme quadratique Q , observons que nécessairement $r \leq \dim(V^*) = \dim(V)$.

2.) On dit que la forme quadratique Q est *non dégénérée*, si elle est de rang maximal, i.e. si $r = \dim(V)$.

3.) La matrice de Gram de la forme quadratique Q par rapport à une base donnée est par définition la matrice de Gram de la forme bilinéaire symétrique associée.

Remarque. On dit qu'une base $\{v_1, \dots, v_n\}$ de V *orthogonalise* la forme quadratique Q si dans cette base on a

$$Q(x) = \sum_{i=1}^n \alpha_i x_i^2.$$

Dans ce cas on peut constater directement que la matrice de Gram de la forme bilinéaire symétrique associée β est une matrice diagonale. On a en effet

$$\beta(v_i, v_j) = \frac{1}{4}(Q(v_i + v_j) - Q(v_i - v_j)) = \frac{1}{4}((\alpha_i + \alpha_j) - (\alpha_i - \alpha_j)) = 0, \quad \text{si } i \neq j,$$

et

$$\beta(v_i, v_i) = \frac{1}{4} (Q(v_i + v_i) - Q(v_i - v_i)) = \frac{1}{4} Q(2v_i) = \alpha_i.$$

Cela signifie que la matrice de Gram de β dans la base $\{v_i\}$ est la matrice diagonale dont les coefficients sont $\beta_{ij} = \alpha_i \delta_{ij}$.

Remarque. Une forme quadratique sur K^n est un polynôme homogène de n -variables à coefficients dans le corps K . Orthogonaliser cette forme quadratique revient à faire un changement de variables qui l'exprime comme somme pondérée de carrés.

10.5.1 Réduction d'une forme quadratique à une somme de carré (méthode de complétion des carrés de Gauss)

La méthode élémentaire suivante, qu'on attribue à Gauss, permet de construire un changement linéaire de variables qui orthogonalise une forme quadratique donnée. Soit $Q(x) = \sum_{i,j=1}^n b_{ij} x_i x_j$ une forme quadratique non nulle à n variables. Plusieurs cas peuvent se présenter :

- (i) Si Q contient le terme x_1^2 , alors cette forme quadratique peut s'écrire sous la forme

$$Q(x_1, \dots, x_n) = a \left(x_1^2 + 2 \sum_{i=2}^n b_i x_1 x_i \right) + \widehat{Q}_1(x_2, \dots, x_n),$$

où $a \neq 0$ et $\widehat{Q}_1$ est une forme quadratique en $(n-1)$ variables (qui ne contient pas la variable x_1). L'idée est alors d'ajouter le terme $a \left(\sum_{i=2}^n b_i x_i \right)^2$ pour compléter le carré de la partie de Q qui contient x_1 , puis de soustraire ce terme (pour conserver l'égalité). On écrit donc

$$\begin{aligned} Q(x_1, \dots, x_n) &= a \left(x_1 + \sum_{i=2}^n b_i x_i \right)^2 - a \left(\sum_{i=2}^n b_i x_i \right)^2 + \widehat{Q}_1(x_2, \dots, x_n) \\ &= a \left(x_1 + \sum_{i=2}^n b_i x_i \right)^2 + \widehat{Q}_2(x_2, \dots, x_n), \end{aligned}$$

où $\widehat{Q}_2$ est la forme quadratique en $(n-1)$ variables définie par

$$\widehat{Q}_2(x_2, \dots, x_n) = \widehat{Q}_1(x_2, \dots, x_n) - a \left(\sum_{i=2}^n b_i x_i \right)^2$$

- (ii) Si Q ne contient pas le terme x_1^2 , mais qu'il contient un terme x_j^2 avec $j \geq 2$, alors on procède comme dans le cas (i) mais avec le terme x_j .
- (iii) Si Q ne contient aucun terme carré, alors c'est une somme de terme mixtes $x_i x_j$. Dans ce cas peut utiliser l'identité

$$x_i x_j = \frac{1}{4} \left((x_i + x_j)^2 - (x_i - x_j)^2 \right),$$

qui nous ramène au cas précédent.

On itère le procédé jusqu'à ce que la forme $Q(x)$ apparaisse comme somme de carrés. Voyons quelques exemples concrets.

Exemples.

1. La forme quadratique sur $\mathbb{R}^2$ définie par

$$Q_1(x_1, x_2) = 2x_1^2 - 5x_2^2 + 4x_1x_2$$

peut se réduire ainsi

$$\begin{aligned} Q_1(x_1, x_2) &= 2x_1^2 - 5x_2^2 + 4x_1x_2 \\ &= 2(x_1^2 + 2x_1x_2 + x_2^2) - 2x_2^2 - 5x_2^2 \\ &= 2(x_1 + x_2)^2 - 7x_2^2. \end{aligned}$$

2. La forme quadratique

$$Q_3(x, y, z) = x^2 + 2xy + 10y^2 - 6yz + 6z^2$$

peut s'écrire de la façon suivante comme somme de carrés

$$Q_3 = (x + y)^2 + (3y - z)^2 + 5z^2$$

Les étapes pour cet exemple sont :

$$\begin{aligned} Q_3(x, y, z) &= x^2 + 2xy + 10y^2 - 6yz + 6z^2 \\ &= (x + y)^2 + 9y^2 - 6yz + 6z^2 \\ &= (x + y)^2 + (3y - z)^2 + 5z^2 \end{aligned}$$

3. La forme quadratique sur $\mathbb{R}^3$ définie par

$$Q_2(x, y, z) = 6x^2 + 12xy - 12xz + 7y^2 - 8yz + 10z^2$$

peut se réduire ainsi

$$\begin{aligned} Q_2(x, y, z) &= 6(x + y - z)^2 - 6(y - z)^2 + 7y^2 - 8yz + 10z^2 \\ &= 6(x + y - z)^2 - 6(y^2 - 2yz + z^2) + 7y^2 - 8yz + 10z^2 \\ &= 6(x + y - z)^2 + y^2 + 4yz + 4z^2 \\ &= 6(x + y - z)^2 + (y + 2z)^2. \end{aligned}$$

Notons que cette méthode ne donne pas une façon unique d'écrire une forme quadratique comme somme de carrés, car elle dépend de l'ordre donné aux variables.

Chapitre 11

Produits scalaires et espaces vectoriels euclidiens

11.1 Définitions fondamentales.

On considère un espace vectoriel réel V sur le corps des réels.

Définitions. Un *produit scalaire (généralisé)* sur V est une application

$$g : V \times V \rightarrow \mathbb{R}$$

qui est *bilinéaire, symétrique et définie-positive* :

- (i.) g est bilinéaire, c'est-à-dire linéaire en chaque variable.
- (ii.) g est symétrique, c'est-à-dire $g(x, y) = g(y, x)$ pour tous $x, y \in V$.
- (iii.) g est *positive*, c'est-à-dire $g(x, x) \geq 0$ pour tout $x \in V$.
- (iv.) g est *définie*, c'est-à-dire $g(x, x) = 0 \Leftrightarrow x = 0$.

Un *espace euclidien* est un espace vectoriel réel de dimension finie muni d'un produit scalaire.

Remarque. Lorsqu'on s'est donné un produit scalaire g sur V , on note souvent $\langle x, y \rangle_g = g(x, y)$ (ou simplement $\langle x, y \rangle$ s'il n'y a pas de risque d'ambiguïté).

Exemples 1. Le *produit scalaire standard* sur $\mathbb{R}^n$ est défini par

$$\langle x, y \rangle = x_1 y_1 + \cdots + x_n y_n = \sum_{i=1}^n x_i y_i.$$

2. Sur l'espace $M_n(\mathbb{R})$ on a un produit scalaire défini par

$$\langle A, B \rangle = \text{Trace}(A^\top \cdot B).$$

3. On note ℓ_2 l'espace vectoriel des suites réelles de carré sommable

$$\ell_2 = \left\{ \xi = (x_k)_{k=1}^\infty \mid x_k \in \mathbb{R}, \sum_{k=1}^\infty |x_k|^2 < \infty \right\}.$$

Un produit scalaire naturel sur cet espace est défini par

$$(\xi \mid \eta)_{\ell_2} = \sum_{i=1}^{\infty} x_i y_i.$$

Observons qu'il s'agit d'une généralisation en dimension infinie du produit scalaire standard de $\mathbb{R}^n$.

4. Un produit scalaire naturel est défini sur l'espace vectoriel $C^0([a, b])$ des fonctions continues sur l'intervalle $[a, b]$ par l'intégration :

$$(f \mid h)_{L^2} = \int_a^b f(x)h(x)dx,$$

on l'appelle¹ le "produit scalaire L^2 ".

Définition 11.1.1. Si g est un produit scalaire sur V , on définit la *norme* d'un vecteur $x \in V$ associée à g par :

$$\|x\|_g = \sqrt{g(x, x)} = \sqrt{\langle x, x \rangle}.$$

La norme est bien définie car $\langle x, x \rangle \geq 0$ pour tout x . Le résultat suivant est une propriété fondamentale des produits scalaires.

Proposition 11.1.2 (Inégalité de Cauchy-Schwarz.). *Soit V est un espace vectoriel réel muni d'un produit scalaire $\langle \cdot, \cdot \rangle$. Pour tous $x, y \in V$, on a*

$$|\langle x, y \rangle| \leq \|x\| \|y\|.$$

De plus il y a égalité si et seulement si x et y sont colinéaires.

Preuve. On note $p(t) = \|xt + y\|^2$ et on observe qu'il s'agit un polynôme à coefficients réel de degré 2. Plus précisément on a

$$\begin{aligned} p(t) &= \|tx + y\|^2 \\ &= \langle tx + y, tx + y \rangle = t^2 \langle x, x \rangle + t \langle x, y \rangle + t \langle y, x \rangle + \langle y, y \rangle \\ &= \|x\|^2 t^2 + 2 \langle x, y \rangle t + \|y\|^2. \end{aligned}$$

Il est d'autre part que $p(t) \geq 0$ pour tout $t \in \mathbb{R}$, le polynôme $p(t)$ a donc au plus une racine réelle, ce qui implique que son discriminant Δ est ≤ 0 . On a donc

$$\Delta = (\langle x, y \rangle)^2 - \|x\|^2 \|y\|^2 \leq 0,$$

c'est-à-dire $|\langle x, y \rangle| \leq \|x\| \|y\|$. De plus on a égalité si et seulement si $\Delta = 0$. Dans ce cas il existe $t \in \mathbb{R}$ tel que $p(t) = 0$, ce qui signifie que $y = -tx$.

□

1. Cette terminologie est justifiée par le fait que ce produit scalaire peut être défini sur l'espace des fonctions de carré intégrable au sens de Lebesgue. La lettre L fait référence à Lebesgue et l'exposant 2 à la condition d'intégrabilité du carré de f . Il s'agit donc de l'espace vectoriel des fonctions $f : [a, b] \rightarrow \mathbb{R}$ qui vérifient $\int_a^b f(x)^2 dx < \infty$, l'intégrale étant prise au sens de Lebesgue (qui est plus générale que la notion d'intégrabilité au sens de Riemann).

Proposition 11.1.3. *La norme vérifie les propriétés suivantes pour tous $x, y \in V$ et $\lambda \in \mathbb{R}$:*

- (a) $\|x\| \geq 0$ et $\|x\| = 0$ si et seulement si $x = 0$.
- (b) $\|\lambda x\| = |\lambda| \|x\|$.
- (c) $\|x + y\| \leq \|x\| + \|y\|$.

Preuve. Les deux premières propriétés suivent facilement des définitions. La troisième propriété est une conséquence de l'inégalité de Cauchy-Schwarz. On a en effet

$$\|x + y\|^2 = \|x\|^2 + 2\langle x, y \rangle + \|y\|^2 \leq \|x\|^2 + 2\|x\|\|y\| + \|y\|^2 = (\|x\| + \|y\|)^2.$$

Comme les normes de x , y et $x + y$ sont positives ou nulles, on peut prendre la racine carrée dans l'inégalité ci-dessus, ce qui nous donne $\|x + y\| \leq \|x\| + \|y\|$. □

Remarquons qu'on a également $\|x - y\| \leq \|x\| + \|y\|$, car $\|-y\| = \|y\|$ par la première propriété, et donc

$$\|x - y\| = \|x + (-y)\| \leq \|x\| + \|-y\| = \|x\| + \|y\|.$$

Définition. Si g est un produit scalaire sur l'espace vectoriel réel V , on définit :

- (1.) La *distance* entre deux éléments x et y de V est

$$d(x, y) = \|y - x\|.$$

- (2.) L'*angle* $\alpha \in [0, \pi]$ entre deux vecteurs non nuls $x, y \in V$ est défini par

$$\cos(\alpha) = \frac{\langle x, y \rangle}{\|x\|\|y\|}.$$

Cette notion est bien définie car d'une part $\|x\|\|y\| \neq 0$ lorsque x et y sont non nuls et d'autre part on a

$$-1 \leq \frac{\langle x, y \rangle}{\|x\|\|y\|} \leq +1$$

par l'inégalité de Cauchy-Schwarz. Notons que le produit scalaire est parfois défini géométriquement à partir de la notion d'angle via la formule

$$\langle x, y \rangle = \|x\|\|y\| \cos(\alpha),$$

mais du point de vue de l'algèbre linéaire, c'est le produit scalaire qui est la notion de base et l'angle est une notion dérivée, et non l'inverse.

- (3.) L'*aire* du parallélogramme $\mathcal{P}(x, y)$ construit sur les vecteurs x et y est définie par

$$\text{Aire}(x, y) = \sqrt{\|x\|^2\|y\|^2 - \langle x, y \rangle^2}.$$

A nouveau, l'inégalité de Cauchy-Schwarz justifie aussi que $\text{Aire}(x, y)$ est bien définie. On vérifie facilement que

$$\text{Aire}(x, y) = \|x\|\|y\| \sin(\alpha).$$

Proposition 11.1.4. Soit E un espace vectoriel euclidien. Alors pour tous $x, y, z \in E$ on a

(i.) $d(x, z) \leq d(x, y) + d(y, z)$ (inégalité du triangle).

(ii.) Si x et y sont non nuls, alors l'angle θ entre x et y est égal à $\pi/2$ si et seulement si

$$\|x + y\|^2 = \|x\|^2 + \|y\|^2 \quad (\text{théorème de Pythagore}).$$

Preuve. (i) En utilisant la proposition 11.1.3 (c) on voit que

$$d(x, z) = \|z - x\| = \|(z - y) + (y - x)\| \leq \|(z - y)\| + \|(y - x)\| = d(x, y) + d(y, z).$$

(ii) Le théorème de Pythagore est une conséquence de la bilinéarité du produit scalaire et de la définition de l'angle. En effet on a d'une part

$$\theta = \pi/2 \quad \Leftrightarrow \quad \cos(\theta) = 0 \quad \Leftrightarrow \quad \langle x, y \rangle = 0.$$

D'autre part

$$2\langle x, y \rangle = \|x + y\|^2 - (\|x\|^2 + \|y\|^2)$$

□

Proposition 11.1.5. Soient a, b deux vecteurs non nuls d'un espace vectoriel euclidien E . Alors il existe deux vecteurs c et d tels que a et d forment un angle de $\pi/2$, c est colinéaire à a , et $b = c + d$.

Preuve. On cherche c et d sous la forme $c = \lambda a$ et $d = b - c$. On demande que a et d forment un angle droit, on a donc

$$0 = \langle d, a \rangle = \langle b - c, a \rangle = \langle b - \lambda a, a \rangle = \langle b, a \rangle - \lambda \langle a, a \rangle.$$

Par conséquent :

$$\lambda = \frac{\langle b, a \rangle}{\langle a, a \rangle} = \frac{\langle b, a \rangle}{\|a\|^2},$$

et donc

$$c = \frac{\langle b, a \rangle}{\|a\|^2} a \quad \text{et} \quad d = b - \frac{\langle b, a \rangle}{\|a\|^2} a.$$

□

Observons que dans cette décomposition $b = c + d$, le vecteur c représente la composante de b en direction de a et d représente la composante de b normale à a .

11.2 Orthogonalité dans un espace vectoriel euclidien

Définitions. 1. Deux vecteurs $x, y \in E$ sont dit *orthogonaux* si $\langle x, y \rangle = 0$. On note cette relation $x \perp y$.

2. Deux sous-espaces vectoriels $W_1, W_2 \subset E$ sont dit *orthogonaux* si $x \perp y$ pour tout $x \in W_1$ et tout $y \in W_2$. On note cette relation $W_1 \perp W_2$.

3. Une base $\{v_1, \dots, v_n\}$ de E est dite *orthogonale* si $v_i \perp v_j$ pour tous $1 \leq i, j \leq n$ tels que $i \neq j$.

4. Une base $\{v_1, \dots, v_n\}$ de E est dite *orthonormée* si elle est orthogonale et si $\|v_i\| = 1$ pour tout i .

Lemme 11.2.1. Soit (E, g) un espace vectoriel euclidien et $\{v_1, \dots, v_n\}$ une base de E . Les conditions suivantes sont équivalentes :

- (a) $\{v_1, \dots, v_n\}$ est une base orthonormée de E .
- (b) $g(v_i, v_j) = \langle v_i, v_j \rangle = \delta_{ij}$.
- (c) La matrice de Gram G de g dans cette base est la matrice identité.

Ce lemme ne fait que traduire les définitions. □

Théorème 11.2.2. Tout espace vectoriel euclidien E admet des bases orthonormées.

Preuve. La preuve se fait par récurrence sur $n = \dim(E)$. Si $n = 1$ il suffit de choisir un vecteur $w \in E$ non nul. Alors $e = \frac{w}{\|w\|}$ est une base de E .

Admettons le théorème démontré pour $n - 1$ et supposons que $\dim(E) = n$. Choisissons de nouveau un vecteur $w \in E$ non nul et définissons un covecteur $\theta \in E^*$ par

$$\theta(x) = \langle w, x \rangle.$$

Alors $\theta : E \rightarrow \mathbb{R}$ est une application linéaire surjective (car $\theta(w) = \|w\|^2 \neq 0$) donc, par le théorème du rang, on a $\dim \text{Ker}(\theta) = n - 1$. Notons ce sous-espace

$$E_1 = w^\perp = \{x \in E \mid \langle w, x \rangle = 0\} = \text{Ker}(\theta) \subset E.$$

Par hypothèse de récurrence, il existe une base orthonormée de E_1 . Notons $\{e_1, \dots, e_{n-1}\} \subset E_1$ cette base et $e_n = \frac{w}{\|w\|}$. Alors on a

$$\langle e_i, e_j \rangle = \delta_{ij}, \quad \text{pour } 1 \leq i, j \leq n.$$

par conséquent $\{e_1, \dots, e_n\}$ est une base orthonormée de E . Le théorème est démontré. □

Remarque. Si $\{e_1, \dots, e_n\}$ est une base orthonormée de l'espace vectoriel euclidien E , et si $x, y \in E$ sont des vecteurs de composantes x_i, y_j dans cette base, alors on a

$$\langle x, y \rangle = \left\langle \sum_{i=1}^n x_i e_i, \sum_{j=1}^n y_j e_j \right\rangle = \sum_{i,j=1}^n x_i y_j \langle e_i, e_j \rangle = \sum_{i,j=1}^n x_i y_j \delta_{ij} = \sum_{i=1}^n x_i y_i.$$

Ainsi dans une base orthonormée, le produit scalaire se calcule de la même manière que le produit scalaire standard de $\mathbb{R}^n$.

11.2.1 Projections orthogonales sur un sous-espace vectoriel

Soit V un espace vectoriel euclidien. dont on note $\langle \cdot, \cdot \rangle$ le produit scalaire. Tout sous-espace vectoriel $W \subset V$ est alors lui-même un espace euclidien pour le même produit scalaire restreint à W . En particulier W possède des bases orthonormées. Le théorème suivant nous permet de construire la *projection orthogonale* de V sur W .

Théorème 11.2.3. Soit $\{w_1, \dots, w_m\}$ une base orthonormée du sous-espace vectoriel W de l'espace vectoriel euclidien V . On note $P_W : V \rightarrow V$ l'application définie par

$$P_W(x) = \sum_{i=1}^m \langle w_i, x \rangle w_i. \quad (11.1)$$

Cette application possède les propriétés suivantes :

- (i) P_W est linéaire.
- (ii) $P_W(x) = x$ si et seulement si $x \in W$.
- (iii) $W = \text{Im}(P_W)$.
- (iv) $P_W \circ P_W = P_W$.
- (v) Le noyau de P_W est l'ensemble des vecteurs de V qui sont orthogonaux à tous les vecteurs de W . On note

$$W^\perp = \text{Ker}(P_W) = \{v \in V \mid \langle v, w \rangle = 0 \ \forall w \in W\}.$$

- (vi) W et $W^\perp$ sont supplémentaires dans V , i.e. $V = W \oplus W^\perp$.

Définitions. On dit qu'un sous-espace vectoriel W d'un espace vectoriel euclidien V est la somme directe orthogonale de W_1 et W_2 si $W_1, W_2 \subset V$ sont des sous-espaces vectoriels tels que $W = W_1 \oplus W_2$ et $W_1 \perp W_2$. Dans ce cas on note

$$W = W_1 \boxplus W_2.$$

La proposition précédente nous dit en particulier que pour tout sous-espace vectoriel $W \subset V$ on a

$$V = W \boxplus W^\perp.$$

Preuve du théorème. (i) La linéarité de P_W découle de la linéarité de l'application $x \mapsto \langle w_i, x \rangle$.
(ii) Supposons $x \in W$, alors on peut s'écrire $x = \sum_{j=1}^m x_j w_j$ et donc

$$\langle w_i, x \rangle = \langle w_i, \sum_{j=1}^m x_j w_j \rangle = x_i.$$

Par conséquent

$$P_W(x) = \sum_{i=1}^m x_i w_i = x.$$

Inversément, si $y \notin W$, alors on a $y \neq P_W(y)$ (car clairement $P_W(y) \in W$).

(iii) Il est clair par construction que $\text{Im}(P_W) \subset W$. D'autre part $W \subset \text{Im}(P_W)$ par la condition précédente, car tout $x \in W$ vérifie $x = P_W(x) \in \text{Im}(P_W)$.

(iv) Cette condition découle immédiatement de (ii) car pour tout x on a $P_W(x) \in W$, donc

$$P_W^2(x) = P_W(P_W(x)) = P_W(x).$$

(v) Montrons d'abord que $W^\perp \subset \text{Ker}(P_W)$. Soit donc $x \in W^\perp$, alors $\langle w, x \rangle = 0$ pour tout $w \in W$, en particulier $\langle w_i, x \rangle = 0$ pour tout $i = 1, \dots, m$, et donc $P_W(x) = \sum_{i=1}^m \langle w_i, x \rangle w_i = 0$.

Montrons maintenant l'inclusion réciproque. Supposons $x \in \text{Ker}(P_W)$, alors

$$P_W(x) = \sum_{i=1}^m \langle w_i, x \rangle w_i = 0,$$

et donc $\langle w_i, x \rangle = 0$ pour tout $i = 1, \dots, m$ car les vecteurs w_i sont linéairement indépendants. Pour montrer que $x \in W^\perp$ on doit prouver que $\langle w, x \rangle = 0$ pour tout $w \in W$. Mais si $w \in W$, alors on peut écrire $w = \sum_{i=1}^m \lambda_i w_i$, et on a donc

$$\langle w, x \rangle = \left\langle \sum_{i=1}^m \lambda_i w_i, x \right\rangle = \sum_{i=1}^m \lambda_i \langle w_i, x \rangle = 0.$$

(vi) On remarque tout d'abord que $W \cap W^\perp = \{0\}$. En effet, si $x \in W \cap W^\perp$, alors $x \perp x$, c'est-à-dire $\|x\|^2 = \langle x, x \rangle = 0$ ce qui implique que $x = 0$ car tout produit scalaire est défini positif. Il nous reste à prouver que $V = W + W^\perp$. Or pour tout vecteur $x \in V$ on a $(x - P_W(x)) \in \text{Ker}(P_W)$ car $P_W(x - P_W(x)) = P_W(x) - P_W^2(x) = P_W(x) - P_W(x) = 0$. On a donc

$$x = P_W(x) + \underbrace{(x - P_W(x))}_{\in \text{Ker}(P_W) = W^\perp} \in W + W^\perp.$$

□

Remarques.

- (i) L'application P_W ne dépend que du sous-espace $W \subset V$ et pas du choix de la base orthonormée $\{w_1, \dots, w_m\} \subset W$. Cette application s'appelle la *projection orthogonale de V sur W* . On dit aussi que P_W est un *projecteur orthogonal*.
- (ii) Si on note $P_{W^\perp}$ la projection sur $W^\perp$, alors on a

$$P_W + P_{W^\perp} = \text{Id}_V, \quad P_W \circ P_{W^\perp} = P_{W^\perp} \circ P_W = 0.$$

- (iii) Si $\{w_1, \dots, w_m\} \subset W$ est une base de W et $\{w_{m+1}, \dots, w_n\} \subset W^\perp$ est une base de $W^\perp$, alors $\{w_1, \dots, w_n\}$ est une base de V et la matrice de P_W dans cette base est

$$M(P_W) = I_m \oplus 0_{n-m} = \begin{pmatrix} 1 & & & & \\ & \ddots & & & \\ & & 1 & & \\ & & & 0 & \\ & & & & \ddots \\ & & & & & 0 \end{pmatrix}$$

Un autre raisonnement pour obtenir cette matrice est le suivant : La relation $P_W^2 = P_W$ nous dit que le polynôme minimal de P_W est $\mu(t) = t^2 - t = t(t - 1)$, il est scindé à racine simple donc P_W est diagonalisable et les multiplicités géométriques des valeurs propres sont m pour $\lambda = 1$ et $(n - m)$ pour $\lambda = 0$.

Proposition 11.2.4. Soit V un espace vectoriel euclidien et W un sous-espace vectoriel. Alors pour tout $x \in V$, le point $P_W(x)$ est le point de W le plus proche de x . Plus précisément, si on note $x' = P_W(x)$, alors $x' \in W$ et

$$\|x - x'\| \leq \|x - y\| \quad \text{pour tout } y \in W,$$

avec égalité si et seulement si $y = x'$.

Preuve. Observons que $P_W(x - x') = P_W(x) - P_W(P_W(x)) = 0$, donc $(x - x') \in \text{Ker}(P_W) = W^\perp$. On a donc pour tout $y \in W$, en utilisant le théorème de Pythagore

$$\|x - y\|^2 = \|(x - x') + (x' - y)\|^2 = \|x - x'\|^2 + \|x' - y\|^2.$$

On a donc $\|x - y\|^2 \geq \|x - x'\|^2$, avec égalité si et seulement si $\|x' - y\| = 0$, c'est-à-dire si $y = x'$. $\square$

Corollaire 11.2.5. La distance d'un point x d'un espace vectoriel euclidien E à un sous-espace vectoriel $W \subset E$ est donnée par

$$\text{dist}(x, W) = \|x - P_W(x)\| = \left\| x - \sum_{i=1}^m \langle w_i, x \rangle w_i \right\|,$$

où $\{w_1, \dots, w_m\}$ une base orthonormée de W .

11.2.2 Symétries orthogonales

Soit V un espace vectoriel euclidien et $W \subset V$ est un sous-espace vectoriel,

Définition. On appelle *symétrie orthogonale* à travers W l'endomorphisme $S_W : V \rightarrow V$ défini par

$$S_W = 2P_W - \text{Id}_V. \quad (11.2)$$

Si $\{w_1, \dots, w_m\}$ est une base orthonormée de W , alors on peut écrire explicitement

$$S_W(x) = -x + 2 \sum_{i=1}^m \langle w_i, x \rangle w_i. \quad (11.3)$$

Le théorème 11.2.3 implique le corollaire suivant dont la preuve est très simple :

Corollaire 11.2.6. La symétrie orthogonale S_W possède les propriétés suivantes :

- (i) S_W est linéaire.
- (ii) $S_W(x) = x$ pour tout $x \in W$ et $S_W(y) = -y$ pour tout $y \in W^\perp$.
- (iii) $S_W^2 = \text{Id}_V$, en particulier S_W est inversible et égale à son propre inverse.

$\square$

Remarquons que la décomposition de V en somme orthogonale $V = W \oplus W^\perp$ signifie que tout vecteur $v \in V$ s'écrit d'une manière unique $v = x + y$ avec $x \in W$ et $y \in W^\perp$. On a alors

$$S_W(v) = S_W(x + y) = x - y,$$

En effet, si $x \in W$ et $y \in W^\perp$, alors

$$(2P_W - \text{Id}_V)(x + y) = 2P_W(x + y) - \text{Id}_V(x + y) = 2x - (x + y) = x - y.$$

Si $\{w_1, \dots, w_n\}$ est une base de V telle que $\{w_1, \dots, w_m\} \subset W$ et $\{w_{m+1}, \dots, w_n\} \subset W^\perp$ alors la matrice de S_W dans cette base est

$$M(S_W) = I_m \oplus (-I_{n-m}) = \begin{pmatrix} 1 & & & & & \\ & \ddots & & & & \\ & & 1 & & & \\ & & & -1 & & \\ & & & & \ddots & \\ & & & & & -1 \end{pmatrix}.$$

11.3 Le procédé d'orthonormalisation de Gram-Schmidt

La proposition suivante nous donne une méthode explicite pour construire une base orthonormée d'un espace euclidien.

Proposition 11.3.1. *Soit $\{v_1, \dots, v_m\}$ des vecteurs linéairement indépendants d'un espace vectoriel euclidien V . Alors il existe des vecteurs $\{u_1, \dots, u_m\}$ tels que*

- (i) $\{u_1, \dots, u_m\}$ est un système de vecteurs orthonormé, i.e. $\langle u_i, u_j \rangle = \delta_{ij}$ pour tous $i, j \in \{1, \dots, m\}$.
- (ii) Pour tout $k = 1, \dots, m$ on a

$$u_k \in \text{Vec}(\{v_1, \dots, v_k\}), \quad \text{c'est à dire } u_k \text{ est combinaison linéaire de } v_1, \dots, v_k.$$

- (iii) $\langle u_i, v_i \rangle > 0$ pour tout $i = 1, \dots, m$.

De plus cette famille $\{u_1, \dots, u_m\}$ est unique et la construction est algorithmique.

Preuve. Le premier vecteur u_1 doit être un multiple positif de v_1 et on doit avoir $\|u_1\| = 1$. On a donc

$$u_1 = \frac{v_1}{\|v_1\|}.$$

Supposons qu'on a construit les vecteurs $u_1, \dots, u_{k-1}$, et notons

$$W_{k-1} = \text{Vec}(\{v_1, \dots, v_{k-1}\}) = \text{Vec}(\{u_1, \dots, u_{k-1}\}).$$

On note alors

$$\widehat{v}_k = v_k - P_{W_{k-1}}(v_k) = v_k - \sum_{i=1}^{k-1} \langle u_i, v_k \rangle u_i,$$

On vérifie facilement les propriétés suivantes :

- (i) $\widehat{v}_k \perp W_{k-1}$.
- (ii) $\{u_1, \dots, u_{k-1}, \widehat{v}_k\}$ est une famille libre de V , en particulier $\|\widehat{v}_k\|$ est non nul.

Le vecteur u_k cherché est alors défini par

$$u_k = \frac{\widehat{v}_k}{\|\widehat{v}_k\|},$$

le procédé s'arrête après m étapes. □

Notons que $\{u_1, \dots, u_m\}$ est une base orthonormée du sous-espace $W = \text{Vec}(\{v_1, \dots, v_m\})$ engendré par les vecteurs donnés.

Définition 11.3.2. On dit que cette base orthonormée a été obtenue à partir de $\{v_1, \dots, v_m\}$ par le *procédé d'orthonormalisation de Gram-Schmidt*.

Le procédé d'orthonormalisation peut se résumer dans les formules suivantes :

$$\widehat{v}_1 = v_1, \quad \widehat{v}_2 = v_2 - \frac{\langle \widehat{v}_1, v_2 \rangle}{\|\widehat{v}_1\|^2} \widehat{v}_1, \quad \dots, \quad \widehat{v}_k = v_k - \sum_{i=1}^{k-1} \frac{\langle \widehat{v}_i, v_k \rangle}{\|\widehat{v}_i\|^2} \widehat{v}_i,$$

puis pour tout k on pose $u_k = \frac{\widehat{v}_k}{\|\widehat{v}_k\|}$.

11.4 Isométries d'un espace vectoriel euclidien.

Définition. Soit (E, g) un espace vectoriel euclidien. Une *isométrie* de E est une application bijective $f : E \rightarrow E$ qui respecte les distances, c'est-à-dire

$$\|f(y) - f(x)\| = \|y - x\|, \quad \forall x, y \in E,$$

où $\|\cdot\| = \|\cdot\|_g$ est la norme associée au produit scalaire g .

Nous laissons au lecteur le soin de vérifier à partir de cette définition que les isométries de $\mathbb{E}^n$ forment un groupe.

Théorème 11.4.1. *L'application $f : E \rightarrow E$ est une isométrie si et seulement s'il existe un vecteur $b \in \mathbb{E}^n$ et une application linéaire $f_0 : E \rightarrow E$ tels que $f(x) = f_0(x) + b$ pour tout $x \in E$ et*

$$\|f_0(x)\| = \|x\| \quad \forall x \in E.$$

On dit que f_0 est la *partie linéaire* de l'isométrie f et b est le *vecteur de translation* de f . Remarquons que ce vecteur est donné par $b = f(0)$.

Preuve. Nous démontrons d'abord le théorème dans le cas particulier où f est une isométrie fixant l'origine, i.e. $f(0) = 0$. Nous devons prouver que dans ce cas, f est linéaire.

On remarque d'abord que pour tout $x \in E$, on a

$$\|f(x)\| = \|f(x) - f(0)\| = d(f(x), f(0)) = d(x, 0) = \|x\|.$$

On montre maintenant que f respecte les produit scalaires, i.e. $\langle f(x), f(y) \rangle = \langle x, y \rangle$ pour tous $x, y \in E$. Cela découle du calcul suivant :

$$\begin{aligned} 2\langle f(x), f(y) \rangle &= \|f(x)\|^2 + \|f(y)\|^2 - \|f(y) - f(x)\|^2 \\ &= d(f(x), 0)^2 + d(f(y), 0)^2 - d(f(x), f(y))^2 \\ &= d(x, 0)^2 + d(y, 0)^2 - d(x, y)^2 \\ &= \|x\|^2 + \|y\|^2 - \|x - y\|^2 \\ &= 2\langle x, y \rangle. \end{aligned}$$

Nous pouvons maintenant montrer la linéarité de f . Soient $x \in E$ un vecteur quelconque et $\alpha \in \mathbb{R}$, alors

$$\begin{aligned}\|f(\alpha x) - \alpha f(x)\|^2 &= \|f(\alpha x)\|^2 - 2\langle f(\alpha x), \alpha f(x) \rangle + \alpha^2 \|f(x)\|^2 \\ &= \|f(\alpha x)\|^2 - 2\alpha \langle f(\alpha x), f(x) \rangle + \alpha^2 \|f(x)\|^2 \\ &= \|\alpha x\|^2 - 2\alpha \langle \alpha x, x \rangle + \alpha^2 \|x\|^2 \\ &= 0,\end{aligned}$$

ce qui prouve que $f(\alpha x) = \alpha f(x)$.

D'autre part, si $x, y \in E$ sont deux vecteurs, alors

$$\begin{aligned}\|f(x) + f(y) - f(x+y)\|^2 &= \langle f(x) + f(y) - f(x+y), f(x) + f(y) - f(x+y) \rangle \\ &= \|f(x)\|^2 + \|f(y)\|^2 + \|f(x+y)\|^2 + 2\langle f(x), f(y) \rangle - 2\langle f(x), f(x+y) \rangle - 2\langle f(y), f(x+y) \rangle \\ &= \|x\|^2 + \|y\|^2 + \|x+y\|^2 + 2\langle x, y \rangle - 2\langle x, x+y \rangle - 2\langle y, x+y \rangle \\ &= \langle x+y - (x+y), x+y - (x+y) \rangle = 0,\end{aligned}$$

ce qui prouve que $f(x+y) = f(x) + f(y)$. On a donc démontré qu'une isométrie de E qui fixe l'origine est une application linéaire.

Pour le cas d'une isométrie générale, on définit une application $f_0 : E \rightarrow E$ par $f_0(x) = f(x) - f(0)$. Alors il est clair que $f_0(0) = 0$ et f_0 est une isométrie car

$$\begin{aligned}d(f_0(x), f_0(y)) &= \|f_0(x) - f_0(y)\| \\ &= \|(f(x) - f(0)) - (f(y) - f(0))\| \\ &= \|f(x) - f(y)\| \\ &= d(x, y).\end{aligned}$$

On a donc montré que l'application f s'écrit $f(x) = f_0(x) + b$, où $b = f(0) \in E$ est constant et f_0 est une isométrie linéaire. □

Corollaire 11.4.2. *Si g est un produit scalaire sur $\mathbb{R}^n$ et $f : \mathbb{R}^n \rightarrow \mathbb{R}^n$ est une isométrie pour la distance associée à ce produit scalaire, alors on a*

$$f(x) = Ax + b,$$

où $b = f(0)$ et $A \in GL_n(\mathbb{R})$ est une matrice vérifiant $A^\top GA = G$ (où G est la matrice de Gram de g dans la base canonique de $\mathbb{R}^n$).

Preuve. On a vu dans la preuve du théorème précédent que l'application $x \mapsto Ax$ préserve le produit scalaire, i.e. on a

$$g(Ax, Ay) = g(x, y)$$

pour tous $x, y \in \mathbb{R}^n$. D'autre part $Ae_r = \sum_{i=1}^n a_{ir}e_i$ et $Ae_s = \sum_{j=1}^n a_{js}e_j$. Rappelons que par définition de la matrice de Gram, on a $g_{ij} = g(e_i, e_j)$, par conséquent

$$g_{rs} = g(e_r, e_s) = g(Ae_r, Ae_s) = g\left(\sum_{i=1}^n a_{ir}e_i, \sum_{j=1}^n a_{js}e_j\right) = \sum_{i,j=1}^n a_{ir}g_{ij}a_{js} = \left(A^\top GA\right)_{rs},$$

ce qui prouve que $G = A^\top GA$.

Voici une autre preuve : on peut représenter les vecteurs de $\mathbb{R}^n$ par des matrice-colonnes et écrire $g(x, y) = X^\top GY$. L'égalité $g(x, y) = g(AX, AY)$ s'écrit alors $X^\top GY = (AX)^\top G(AY)$ et donc

$$X^\top GY = (AX)^\top G(AY) = (X^\top A^\top) G(AY) = X^\top (A^\top GA)Y,$$

pour tous X, Y . Ceci implique que $G = A^\top GA$. □

Ce résultat justifie la définition importante suivante :

Définition 11.4.3. Une matrice $A \in M_n(\mathbb{R})$ est *G-orthogonale* si $A^\top GA = G$. On note

$$O(G) = \{A \in M_n(\mathbb{R}) \mid A^\top GA = G\}.$$

Remarques.

1. Il est facile de vérifier que $\det(A) = \pm 1$ pour tout $A \in O(G)$. De plus $O(G)$ est un sous-groupe de $GL_n(\mathbb{R})$.

2. Lorsque g est le produit scalaire standard de $\mathbb{R}^n$, alors on a $G = I_n$ et on note le groupe orthogonal correspondant simplement

$$O(n) = O(I_n) = \{A \in M_n(\mathbb{R}) \mid A^\top A = I_n\}.$$

Observer que $A \in O(n)$ si et seulement si A est inversible et $A^\top = A^{-1}$.

11.5 Le groupe orthogonal

Dans cette section, nous étudions les matrices orthogonales en détail.

Proposition 11.5.1. Pour toute matrice $A \in M_n(\mathbb{R})$ les propriétés suivantes sont équivalentes :

- (i) $A \in O(n)$, c'est-à-dire $A^\top A = I_n$.
- (ii) A est inversible et $A^{-1} = A^\top$.
- (iii) $\|Ax\| = \|x\|$ pour tout $x \in \mathbb{R}^n$.
- (iv) $\langle Ax, Ay \rangle = \langle x, y \rangle$ pour tous $x, y \in \mathbb{R}^n$.
- (v) Les colonnes de A forment une base orthonormée de $\mathbb{R}^n$.
- (vi) Les lignes de A forment une base orthonormée de $\mathbb{R}^n$.
- (vii) Pour tout vecteur $b \in \mathbb{R}^n$, l'application affine $f : \mathbb{R}^n \rightarrow \mathbb{R}^n$ définie par $f(x) = Ax + b$ est une isométrie.

De plus $O(n)$ est un sous-groupe de $GL_n(\mathbb{R})$ et pour tout $A \in O(n)$ on a $\det(A) = \pm 1$.

Dans cette proposition, le produit scalaire est le produits scalaire standard de $\mathbb{R}^n$ et la norme et la distance sont associées à ce produit scalaire. Nous laissons la preuve de cette proposition en exercice.

Remarquons que l'application déterminant définit un homomorphisme de groupes

$$\det : O(n) \rightarrow \{\pm 1\}.$$

Le noyau de cet homomorphisme est le *groupe spécial orthogonal* :

$$SO(n) = O(n) \cap SL_n(\mathbb{R}) = \{A \in M_n(\mathbb{R}) \mid A^\top A = I_n \text{ et } \det(A) = +1\}.$$

La proposition suivante décrit les 2×2 matrices orthogonales.

Proposition 11.5.2. *Pour toute matrice $A \in O(2)$, il existe un angle θ tel que*

$$A = R_\theta = \begin{pmatrix} \cos(\theta) & -\sin(\theta) \\ \sin(\theta) & \cos(\theta) \end{pmatrix}, \text{ si } \det(A) = +1,$$

et

$$A = S_{\theta/2} = \begin{pmatrix} \cos(\theta) & \sin(\theta) \\ \sin(\theta) & -\cos(\theta) \end{pmatrix}, \text{ si } \det(A) = -1.$$

La matrice R_θ représente une rotation d'angle θ et $S_{\theta/2}$ représente la réflexion à travers la droite vectorielle formant un angle $\theta/2$ avec le premier vecteur e_1 de la base canonique.

Preuve. Les colonnes d'une matrices orthogonale $A \in O(2)$ doivent former une base orthonormée de $\mathbb{R}^2$. Il existe donc $\theta \in (-\pi, \pi]$ tel que la première colonne s'écrive $\begin{pmatrix} \cos(\theta) \\ \sin(\theta) \end{pmatrix}$. La deuxième colonnes de A doit-être un vecteur de norme 1 orthogonal à la première colonne, c'est-à dire $\pm \begin{pmatrix} \sin(\theta) \\ \cos(\theta) \end{pmatrix}$. Ceci démontre que ou bien $A = R_\theta$ ou bien $A = S_{\theta/2}$.

Finalement R_θ est une matrice de rotation car l'angle entre tout vecteur non nul $\mathbf{x}$ et $R_\theta(\mathbf{x})$ est égal à θ et $S_{\theta/2}$ est une symétrie car cette matrice possède deux vecteurs propres orthogonaux de valeurs propre $+1$ et -1 respectivement. Ces vecteurs propres sont

$$\begin{pmatrix} \cos(\theta/2) \\ \sin(\theta/2) \end{pmatrix} \text{ et } \begin{pmatrix} \sin(\theta/2) \\ -\cos(\theta/2) \end{pmatrix}.$$

Nous laissons la vérification de ces deux dernières affirmations en exercice.

□

Décrivons maintenant les 3×3 matrices orthogonales.

Proposition 11.5.3. *Toute matrice $A \in O(3)$ est semblable à une matrice du type*

$$\begin{pmatrix} \pm 1 & 0 & 0 \\ 0 & \cos(\theta) & -\sin(\theta) \\ 0 & \sin(\theta) & \cos(\theta) \end{pmatrix}.$$

Nous allons démontrer le résultat plus général suivant qui donne la structure des isométries linéaires d'un espace euclidien.

Théorème 11.5.4. *Soit $f : V \rightarrow V$ une isométrie linéaire d'un espace vectoriel euclidien de dimension n . Alors il existe une base orthonormée de V dans laquelle la matrice de f prend la*

forme

$$M(f) = I_r \oplus (-I_s) \oplus R_{\theta_1} \oplus \cdots \oplus R_{\theta_m}$$

$$= \begin{pmatrix} I_r & & & & \\ & -I_s & & & \\ & & \begin{pmatrix} \cos(\theta_1) & -\sin(\theta_1) \\ \sin(\theta_1) & \cos(\theta_1) \end{pmatrix} & & \\ & & & \ddots & \\ & & & & \begin{pmatrix} \cos(\theta_m) & -\sin(\theta_m) \\ \sin(\theta_m) & \cos(\theta_m) \end{pmatrix} \end{pmatrix}$$

Pour la preuve, nous aurons besoin du lemme suivant :

Lemme 11.5.5. *Soit $f : V \rightarrow V$ une isométrie linéaire. Supposons que $W \subset V$ est invariant par f . Alors $W^\perp$ est aussi invariant par f .*

Preuve du lemme. Observons d'abord que si $W \subset V$ est invariant par f , i.e. $f(W) \subset W$, alors $f(W) = W$, i.e. la restriction de f à W est un isomorphisme de W vers W . Cela découle du théorème du rang et du fait que $\text{Ker}(f) = 0$ pour toute isométrie linéaire d'un espace euclidien. Pour montrer que $W^\perp$ est invariant par f , on se donne $y \in W^\perp$ et $x \in W$ quelconque. On a alors $f^{-1}(x) \in W$, d'où l'on déduit que

$$\langle f(y), x \rangle = \langle f(y), f(f^{-1}(x)) \rangle = \langle y, f^{-1}(x) \rangle = 0.$$

Cela prouve que $f(y) \in W^\perp$.

□

Démonstration du théorème. Nous démontrons le théorème par récurrence sur la dimension de V . Si V est de dimension 1, alors le théorème affirme simplement que les seules isométries linéaires de V sont $f(v) = v$ et $f(v) = -v$, ce qui est évident. Le cas de la dimension 2 a été étudié dans la proposition 11.5.2. Supposons maintenant le théorème démontré pour toute isométrie d'un espace euclidien de dimension inférieure à $n = \dim(V)$. On sait par le théorème 9.14.4 que tout endomorphisme $f : V \rightarrow V$ d'un espace vectoriel réel de dimension finie admet un sous-espace invariant $W \subset V$ de dimension 1 ou 2. Par le lemme précédent, on sait que $W^\perp$ est alors aussi invariant par f . On distingue alors trois cas :

Cas 1. $\dim(W) = 1$. Par hypothèse de récurrence on peut trouver une base orthonormée $\{e_2, \dots, e_n\}$ de $W^\perp$ telle que la matrice de la restriction de f à $W^\perp$ dans cette base prenne la forme

$$I_{r'} \oplus (-I_{s'}) \oplus R_{\theta_1} \oplus \cdots \oplus R_{\theta_m},$$

avec $r' + s' + 2m = n - 1$. Soit $e_1 \in W$ un vecteur de norme 1. Alors on a $f(e_1) = \pm e_1$ et la matrice de f dans la base $\{e_1, e_2, \dots, e_n\}$ prend la forme

$$I_r \oplus (-I_s) \oplus R_{\theta_1} \oplus \cdots \oplus R_{\theta_m},$$

où $(r, s) = (r' + 1, s')$ si $f(e_1) = e_1$ et $(r, s) = (r', s' + 1)$ si $f(e_1) = -e_1$.

Cas 2. $\dim(W) = 2$ et la restriction de f à W est une symétrie. Alors il existe une droite $W_1 \subset W$ invariante par f et nous sommes ramenés au cas 1.

Cas 3. $\dim(W) = 2$ et la restriction de f à W est une rotation d'angle θ . Par hypothèse de récurrence on peut trouver une base orthonormée $\{e_3, \dots, e_n\}$ de $W^\perp$ telle que la matrice de la restriction de f à $W^\perp$ dans cette base prenne la forme

$$I_{r'} \oplus (-I_{s'}) \oplus R_{\theta_1} \oplus \dots \oplus R_{\theta_m},$$

avec $r' + s' + 2m = n - 2$. Soit $\{e_1, e_2\} \in W$ une base orthonormée de W , alors la matrice de la restriction de f au plan invariant W est la matrice de rotation R_θ et la matrice de f dans la base $\{e_1, e_2, \dots, e_n\}$ prend la forme

$$R_\theta \oplus I_r \oplus (-I_s) \oplus R_{\theta_1} \oplus \dots \oplus R_{\theta_m}.$$

Le théorème est démontré. □

On peut reformuler le théorème de la façon suivante : *Pour tout $A \in O(n)$, il existe une matrice $Q \in O(n)$ telle que*

$$A' := Q^\top A Q = Q^{-1} A Q = I_r \oplus (-I_s) \oplus R_{\theta_1} \oplus \dots \oplus R_{\theta_m}.$$

11.6 Espace-temps Galiléen et référentiels inertiels

La mécanique classique, telle que développée depuis Galilée et Newton, et jusqu'à la fin du 19ème siècle, étudie des phénomènes telles que le mouvement dans un espace et pendant un intervalle de temps qui sont considérés comme des absolus. Le philosophe Emmanuel Kant considère que le temps et l'espace sont des formes pures de l'intuition, des *catégories synthétiques a priori* de la connaissance. Pour Kant, dont l'un des projets est d'établir un cadre philosophique permettant d'intégrer la mécanique Newtonienne, l'espace et le temps sont donnés à notre intuition de façon indépendante de toute expérimentation (c'est ici le sens du mot *a priori*).

En mécanique classique, les événements sont donc étudiés dans un espace-temps de dimension 4 correspondant à une dimension temporelle et 3 dimensions spatiales. La mesure du temps est considérée comme absolue, pouvant se dérouler selon un axe réel $-\infty < t < \infty$, et l'espace est un espace euclidien de dimension 3, que nous identifions à $\mathbb{R}^3$. Il est commode d'appeler *événement* un élément ² (x, y, z, t) de $\mathbb{R}^4$. On dit alors que $\mathbb{R}^4 = \mathbb{R}^3 \times \mathbb{R}$ est un *espace-temps galiléen* lorsqu'on le muni des deux structures suivantes :

- (i) La *mesure du temps*, qui est la fonction $\mathbb{R}^4 \rightarrow \mathbb{R}$ donnée par $(x, y, z, t) \rightarrow t$. Elle est concrètement réalisée par l'horloge de référence de l'expérimentateur.
- (ii) La *norme euclidienne*, qui est la fonction $\mathbb{R}^4 \rightarrow \mathbb{R}$ donnée par $(x, y, z, t) \rightarrow \sqrt{x^2 + y^2 + z^2}$. Cette norme est associée au produit scalaire euclidien de $\mathbb{R}^3$ et permet de reconstruire mathématiquement toute la géométrie euclidienne de l'espace (les distances, les angles, les aires etc.). Cette norme représente donc toutes les informations géométriques que fournissent les instruments de mesure disponibles à l'expérimentateur.

2. On dit alors que (x, y, z, t) est le *quadrivecteur* représentant les coordonnées spatio-temporelles de l'événement considéré. Notons qu'on peut parfois ignorer une ou deux coordonnées spatiales, par exemple si on étudie un mouvement dans un plan ou un mouvement rectiligne. L'événement est alors représenté par un élément $(x, y, t) \in \mathbb{R}^3$ ou $(x, t) \in \mathbb{R}^2$.

Nous avons alors les définitions naturelles suivantes :

Définitions 1. On appelle *durée* ou *intervalle temporel* entre deux événements $\xi_1 = (x_1, y_1, z_1, t_1)$ et $\xi_2 = (x_2, y_2, z_2, t_2)$ la quantité

$$\tau(\xi_1, \xi_2) = |t_1 - t_2|.$$

La durée est indépendante de la position spatiale des événements ξ_1 et ξ_2 .

2. La *distance* entre ces événements est la quantité

$$d(\xi_1, \xi_2) = \sqrt{(x_2 - x_1)^2 + (y_2 - y_1)^2 + (z_2 - z_1)^2}.$$

La distance est indépendante du temps.

Définition. Une transformation $f : \mathbb{R}^4 \rightarrow \mathbb{R}^4$ est une *transformation galiléenne* si elle est vérifie les conditions suivantes :

- (i) La transformation f est bijective.
- (ii) La transformation préserve les durées : pour tous $\xi_1, \xi_2 \in \mathbb{R}^4$ on a

$$\tau(f(\xi_1), f(\xi_2)) = \tau(\xi_1, \xi_2).$$

Cette condition correspond à l'hypothèse d'un temps absolu.

- (iii) La transformation est *isométrique*, c'est-à-dire qu'elle préserve les distances : pour tous $\xi_1, \xi_2 \in \mathbb{R}^4$ on a

$$d(f(\xi_1), f(\xi_2)) = d(\xi_1, \xi_2).$$

Cette condition reflète l'hypothèse d'un espace absolu et uniforme, le même pour tous les observateurs.

- (iv) La transformation est *inertielle* : elle transforme un mouvement rectiligne uniforme en un mouvement rectiligne uniforme.

Exemple. L'exemple le plus simple de transformation galiléenne est donné par la formule :

$$\begin{pmatrix} x \\ y \\ z \\ t \end{pmatrix} \mapsto \begin{pmatrix} x + tv_1 \\ y + tv_2 \\ z + tv_3 \\ t + t_0 \end{pmatrix}$$

Théorème 11.6.1. L'application $f : \mathbb{R}^4 \rightarrow \mathbb{R}^4$ est une transformation galiléenne si et seulement s'il existe $t_0 \in \mathbb{R}$, deux vecteurs $b, v \in \mathbb{R}^3$ et une matrice orthogonale $A \in O(3)$ telles que f se décompose sous la forme suivante :

$$\begin{pmatrix} x \\ y \\ z \end{pmatrix} \mapsto A \begin{pmatrix} x \\ y \\ z \end{pmatrix} + \begin{pmatrix} b_1 \\ b_2 \\ b_3 \end{pmatrix} + t \begin{pmatrix} v_1 \\ v_2 \\ v_3 \end{pmatrix} \quad \text{et} \quad t \mapsto t_0 \pm t.$$

Démonstration. La condition (ii) nous permet de décrire une transformation galiléenne en séparant la coordonnée temporelle et les coordonnées spatiales. La coordonnées spatiale se transforme selon la règle $t \mapsto t_0 \pm t$, et pour les coordonnées spatiales, nous avons à chaque instant t une bijection de $f_t : \mathbb{R}^3 \rightarrow \mathbb{R}^3$ qui doit être une isométrie pour la norme euclidienne standard.

Si on écrit les coordonnées spatiotemporelle d'un événement sous la forme $(x, t) = (x_1, x_2, x_3, t)$, alors f_t n'agit que sur les coordonnées spatiales $x = (x_1, x_2, x_3)$ par la formule

$$x \rightarrow f_t x = A(t)x + a(t),$$

où $A(t) \in O(3)$ est une matrice orthogonale qui dépend du temps et $a(t)$ est le vecteur de translation, qui dépend lui aussi du temps.

Admettons pour simplifier que la transformation temporelle est l'identité $t \mapsto t$ et supposons que $t \mapsto x(t)$ représente la trajectoire d'une particule. Notons $y(t) = f_t(x(t))$, nous avons alors

$$y(t) = A(t)x(t) + a(t), \quad \dot{y}(t) = A(t)\dot{x}(t) + \dot{A}(t)x(t) + \dot{a}(t),$$

ou le point représente la dérivée par rapport au temps ; et donc aussi

$$\ddot{y}(t) = A(t)\ddot{x}(t) + 2\dot{A}(t)\dot{x}(t) + \ddot{A}(t)x(t) + \ddot{a}(t). \quad (11.4)$$

L'hypothèse (iv) que la transformation est inertielle dit que pour toute trajectoire telle que $\ddot{x} = 0$ on doit aussi avoir $\ddot{y} = 0$. Cette condition nous dit que si $x(t) = c + tw$, où c et w sont des vecteurs constants quelconques de $\mathbb{R}^3$, alors on a $\dot{x}(t) = w$ et $\ddot{x}(t) = 0$. La condition (11.4) entraîne alors que

$$0 = \ddot{y}(t) = 2\dot{A}(t)w + \ddot{A}(t)(c + tw) + \ddot{a}(t).$$

Cette condition est valable pour tous $c, w \in \mathbb{R}^3$ constants. En posant $c = w = 0$ on obtient que $\ddot{a}(t) = 0$; il existe donc deux vecteurs constants $b, v \in \mathbb{R}^3$ tels que $a(t) = b + tv$. En posant $w = 0$, on voit que $\ddot{A}(t)c = 0$ pour tout vecteur c et donc $\ddot{A}(t) = 0$. Finalement en posant $c = 0$ on voit que $\dot{A}(t)w = 0$ pour tout vecteur w et donc $\dot{A}(t) = 0$.

On a ainsi prouvé que $A(t)$ est un élément constant de $O(3)$ et $a(t) = b + tv$, c'est-à-dire que

$$f_t(x) = Ax + b + tv,$$

avec $A \in O(3)$ et $b, v \in \mathbb{R}^3$.

□

Terminons ce paragraphe par quelques remarques :

Remarques. (i) Puisque la matrice A est constante et $\ddot{a} = 0$, l'équation (11.4) nous dit que l'accélération de $y(t) = f_t(x(t))$ vérifie

$$\ddot{y}(t) = A\ddot{x}(t).$$

Et comme $A \in O(3)$, nous avons en particulier $\|\ddot{y}(t)\| = \|\ddot{x}(t)\|$. Ceci implique que les transformations galiléennes respectent l'équation Newtonienne du mouvement (force = masse $\times$ accélération), ce que l'on formule sous la forme du principe de relativité galiléenne : *Le temps et l'espace sont absolus et les lois de la mécanique sont les mêmes dans tous les référentiels inertiels.*

(ii) Le lecteur attentif aura remarqué que les transformations galiléennes autorisent l'inversion du temps, i.e. la transformation $t \mapsto -t$ de la coordonnée temporelle. Cette convention est justifiée par le fait que l'inversion du temps est compatible avec la loi d'évolution Newtonienne d'une particule se mouvant dans un champ de force. Mais nous savons bien empiriquement que le temps s'écoule du passé vers l'avenir, et que la plupart des évolutions ne sont pas réversibles.

Ajoutons que l'inversion du temps contredit le second principe de la thermodynamique. Il est donc raisonnable de ne considérer que les transformations galiléennes qui respectent l'orientation temporelle (on dit d'une telle transformation qu'elle respecte la chronologie).

(iii) Une remarque similaire s'applique à l'orientation de l'espace. Si on considère que l'orientation de l'espace est une donnée essentielle de la physique, alors nous devons restreindre les transformations galiléennes au cas où $A \in SO(3)$. Notons que l'orientation de l'espace joue un rôle significatif en électromagnétisme (loi de Biot-Savard, force de Lorentz...)

(iv) Mentionnons pour finir que l'inverse d'une transformation galiléenne et la composition de deux transformations galiléennes sont encore des transformations galiléennes. Ces transformations forment donc un groupe, qu'on appelle le *groupe de Galilée*. C'est un sous-groupe du groupe des transformations affines de $\mathbb{R}^4$.

“Le mouvement est comme rien” (un texte de Galilée)

La physique est une science expérimentale dont les lois sont formulées mathématiquement. Toutefois les expérimentations de laboratoire et les développements mathématiques sont complétés par un autre type de raisonnement, qu'on appelle des « expériences de pensées » (Gedankenexperimente). Ces raisonnements mettent en scène des schémas pseudo-expérimentaux et permettent de dégager les grands principes de la physique. Le grand génie des expériences de pensées est Albert Einstein, mais le procédé remonte à Galilée. Dans son Dialogue sur les deux grands systèmes du monde (1632), Galilée propose une expérience de pensée célèbre. Dans ce texte majeur de l'histoire des sciences, il nous donne la première formulation historique du principe de relativité sous la forme suivante.

« Enfermez-vous avec un ami dans la plus vaste cabine d'un grand navire, Et apportez des mouches, des papillons et d'autres petits animaux semblables. Amenez aussi un grand bocal d'eau contenant des poissons, suspendez au plafond un petit seau dont l'eau tombe goutte à goutte par un orifice étroit et tombe dans un vase posé sur le sol.

Puis, alors que le navire est à l'arrêt, observez attentivement, comment ces petits animaux volent avec des vitesses égales quel que soit l'endroit de la cabine vers lequel ils se dirigent. Les poissons nagent indifféremment dans toutes les directions. Les gouttelettes d'eau tombent régulièrement dans le vase situé sur le sol. Si vous lancez un objet à un ami, vous n'avez pas besoin de le lancer plus fort dans une direction que dans une autre, si les distances sont égales, et si vous sautez à pieds joints, vous franchissez des distances égales dans toutes les directions. Il ne fait aucun doute que si le navire est à l'arrêt les choses doivent se passer ainsi.

Une fois que vous aurez observé attentivement tout cela, faites avancer le bateau à l'allure qui vous plaira, pour autant que la vitesse soit uniforme et ne fluctue pas de-ci de-là. Vous ne discernerez alors aucun changement dans tous les effets précédents, et aucune observation ne vous renseignera si le navire est en marche ou s'il est arrêté.

Si vous sautez, vous franchirez sur le plancher les mêmes distances qu'auparavant et, si le navire se déplace, vous n'en ferez pas pour autant des sauts plus grands vers la poupe que vers la proue, bien que, pendant que vous êtes en l'air, le plancher qui est en dessous ait glissé dans la direction opposée à celle de votre saut. Les gouttes d'eau tomberont comme précédemment dans le vase inférieur. Les poissons dans leur eau, et sans plus de fatigue, nageront d'un côté comme de l'autre.

Enfin les papillons et les mouches continueront leur vol indifférent dans n'importe quel sens, sans être influencé par la marche et la direction du navire, on ne les verra pas s'accumuler du côté de la cloison qui fait face à la poupe ; ce qui ne manquerait pas d'arriver s'ils devaient s'épuiser à suivre le navire dans sa course rapide.

La cause de la permanence de tous ces effets, c'est que le mouvement uniforme est commun au navire et à ce qu'il contient, y compris l'air. Le mouvement est mouvement, et agit comme mouvement, en tant et seulement qu'il est en rapport avec les choses qui en sont privées ; mais en ce qui concerne celles qui y participent toutes également, il est sans effet ; il est comme s'il n'était pas. Le mouvement est comme rien ! »

Ajoutons quelques commentaires à ce beau texte. Le Dialogue sur les deux grands systèmes du monde est rédigé sous forme de dialogue entre trois protagonistes. Dans ce texte, écrit en italien et non en latin, Galilée cherche à convaincre le lecteur de la supériorité du système héliocentrique (plaçant le soleil au centre de l'Univers) sur le géocentrisme qui place la terre au centre. L'héliocentrisme a été avancé par Copernic comme un modèle permettant de rendre compte du mouvement des planètes en décrivant leur orbites autour du soleil et non de la terre. Le modèle copernicien a été amélioré par Kepler qui, suite aux observations de Tycho-Brahé, énonça les trois lois qui portent son nom (les orbites des planètes sont des ellipses etc.). L'héliocentrisme fut condamné par l'Eglise en 1616 comme contraire aux enseignements de l'Ecriture. L'un des arguments contre l'héliocentrisme était que si la terre tourne autour du soleil, nous devrions nous en rendre compte (par exemple la chute des corps ne serait pas verticale etc.) Le texte ci-dessus réfute cet argument, mais une polémique s'en est suivie et Galilée fut condamné par l'inquisition à renier ses thèses et son ouvrage fut interdit de publication. C'était en 1633, un an après la publication du Dialogue (l'interdiction fut levée en 1741 par le pape Benoît XIV).

Le principe d'inertie (ou de relativité) de Galilée, est à rapprocher de la première Loi de Newton (énoncée dans les *Principia* en 1687). *Tout corps persévère dans l'état de repos ou de mouvement uniforme en ligne droite dans lequel il se trouve, à moins que quelque force n'agisse sur lui, et ne le contraigne à changer d'état.*

En langage contemporain, le principe de relativité Galiléenne s'énonce ainsi : *Il existe une famille de référentiels qui sont mutuellement au repos ou en mouvement rectiligne uniforme. Les lois de la mécanique (classique) ont la même forme dans tout référentiel inertiel.* Le navire de Galilée est un exemple de référentiel inertiel.

11.7 Le théorème spectral (première version)

Le but de ce paragraphe est de démontrer le théorème suivant :

Théorème 11.7.1 (Théorème Spectral). *Soit $H \in M_n(\mathbb{R})$ une $n \times n$ matrice symétrique (i.e. $H^\top = H$) à coefficients réel. Alors on a*

- (1.) *Les valeurs propres de H sont réelles.*
- (2.) *Les espaces propres de H associés à des valeurs propres distinctes sont deux-à-deux orthogonaux.*
- (3.) *Il existe une base orthonormée de $\mathbb{R}^n$ formée de vecteurs propres.*

Remarques. (i.) Dans les points (2) et (3) l'orthogonalité fait référence au produit scalaire standard de $\mathbb{R}^n$, que nous noterons ici $\langle \cdot, \cdot \rangle$.

(ii.) Le théorème nous dit en particulier que pour une matrice réelle symétrique, la multiplicité géométrique de chaque valeur propre est égale à sa multiplicité algébrique.

(iii.) Ce théorème est l'une des formes du théorème spectral, nous en verrons d'autres.

Le lemme suivant sera utilisé dans la preuve du théorème.

Lemme 11.7.2. *Pour toute matrice $A \in M_n(\mathbb{R})$ et tous $X, Y \in \mathbb{R}^n$ on a*

$$\langle X, AY \rangle = \langle A^\top X, Y \rangle = X^\top AY = \sum_{i,j=1}^n x_i a_{ij} y_j,$$

en particulier $A^\top = A$ si et seulement si $\langle AX, Y \rangle = \langle X, AY \rangle$ pour tous $X, Y \in \mathbb{R}^n$. (on regarde les éléments $X, Y \in \mathbb{R}^n$ comme des vecteurs-colonnes).

Preuve : Exercice.

Démonstration du théorème spectral.

(1.) On démontre d'abord que les valeurs propres de H sont réelles. Considérons donc une valeur propre $\lambda \in \mathbb{C}$ et montrons que $\lambda \in \mathbb{R}$.

En effet, si $\lambda \in \sigma_{\mathbb{C}}(H)$, alors il existe $Z \in \mathbb{C}^n$ tel que $HZ = \lambda Z$. On a alors aussi $H\bar{Z} = \bar{\lambda}\bar{Z}$, car $\bar{H} = H$ puisque H est une matrice à coefficients réels. On a donc

$$H\bar{Z} = \bar{H}\bar{Z} = \overline{HZ} = \overline{\lambda Z} = \bar{\lambda}\bar{Z}.$$

Par conséquent on a

$$Z^\top (H\bar{Z}) = Z^\top (\bar{\lambda}\bar{Z}) = \bar{\lambda} \sum_{i=1}^n z_i \bar{z}_i = \bar{\lambda} \sum_{i=1}^n |z_i|^2,$$

Mais on suppose que $H^\top = H$, donc

$$(Z^\top H)\bar{Z} = (H^\top Z)^\top \bar{Z} = (HZ)^\top \bar{Z} = (\lambda Z)^\top \bar{Z} = \lambda Z^\top \bar{Z} = \lambda \sum_{i=1}^n z_i \bar{z}_i = \lambda \sum_{i=1}^n |z_i|^2.$$

Ainsi

$$\lambda \sum_{i=1}^n |z_i|^2 = Z^\top H \bar{Z} = \bar{\lambda} \sum_{i=1}^n |z_i|^2.$$

Mais comme on a supposé que $Z \neq 0$, cette égalité implique $\lambda = \bar{\lambda}$. Cela prouve que toute valeur propre de H est réelle.

(2.) Montrons maintenant que les espaces propres associés à des valeurs propres distinctes sont orthogonaux, i.e.

$$\lambda, \mu \in \sigma(H), \lambda \neq \mu \Rightarrow E_\lambda \perp E_\mu.$$

En effet, si $x \in E_\lambda$ et $y \in E_\mu$, alors

$$\langle x, Hy \rangle = \langle x, \mu y \rangle = \mu \langle x, y \rangle.$$

Mais on a aussi

$$\langle Hx, y \rangle = \langle \lambda x, y \rangle = \lambda \langle x, y \rangle.$$

Et donc, en utilisant le lemme précédent :

$$\lambda \langle x, y \rangle = \langle Hx, y \rangle = \langle x, Hy \rangle = \mu \langle x, y \rangle,$$

ce qui implique $\langle x, y \rangle = 0$ si $\lambda \neq \mu$. On a prouvé que des espaces propres associés à des valeurs propres distinctes sont orthogonaux.

(3.) Nous pouvons maintenant prouver qu'il existe une base orthonormée formée de vecteurs propres. Notons $\sigma(H) = \{\lambda_1, \dots, \lambda_r\}$ l'ensemble des valeurs propres de H . On peut choisir une base orthonormée $\{u_1, \dots, u_{m_1}\}$ de l'espace propre $E_{\lambda_1} \subset \mathbb{R}^n$ (où m_1 est la multiplicité géométrique de λ_1).

Par le point (2), on sait que pour tout $j \geq 2$, l'espace propre E_{λ_j} est orthogonal E_{λ_1} . C'est-à-dire

$$j \geq 2 \Rightarrow E_{\lambda_j} \subset E_{\lambda_1}^\perp = \{y \in \mathbb{R}^n \mid \langle x, y \rangle = 0, \forall x \in E_{\lambda_1}\}.$$

Affirmation. $E_{\lambda_1}^\perp$ est invariant par H .

Cette affirmation signifie que $y \in E_{\lambda_1}^\perp \Rightarrow Hy \in E_{\lambda_1}^\perp$, et se vérifie facilement grâce au lemme précédent. En effet, supposons que $y \in E_{\lambda_1}^\perp$, alors on a pour tout $x \in E_{\lambda_1}$

$$\langle x, Hy \rangle = \langle Hx, y \rangle = \lambda_1 \langle x, y \rangle = 0,$$

ce qui entraîne que $Hy \in E_{\lambda_1}^\perp$. L'affirmation est démontrée.

Nous pouvons maintenant conclure la démonstration. Par hypothèse de récurrence, il existe alors une base orthonormée $\{u_{m_1+1}, \dots, u_n\}$ de $E_{\lambda_1}^\perp$ formée de vecteurs propres pour H . La réunion $\{u_1, \dots, u_n\}$ des deux bases est la base propre orthonormée de $\mathbb{R}^n$ cherchée.

□

Diagonalisation Orthogonale

Définition 11.7.3. Deux matrices $A, A' \in M_n(\mathbb{R})$ sont dites *orthogonalement congruentes* s'il existe $P \in O(n)$ telle que

$$A' = P^\top AP.$$

Remarquer que dans ce cas les matrices sont aussi semblables car $P^\top = P^{-1}$. Le théorème spectral peut se reformuler ainsi :

Théorème 11.7.4. *Pour une matrice réelle $A \in M_n(\mathbb{R})$, les deux conditions suivantes sont équivalentes :*

- (a) *A est orthogonalement diagonalisable, c'est-à-dire orthogonalement congruente à une matrice diagonale.*
- (b) *A est symétrique, i.e. $A^\top = A$.*

Concrètement, cela signifie que si $A \in M_n(\mathbb{R})$ est symétrique, alors il existe une matrice orthogonale $P \in O(n)$ telle que

$$D = P^{-1}AP = P^\top AP$$

est diagonale.

Pour diagonaliser orthogonalement une matrice symétrique $A \in M_n(\mathbb{R})$, on procède selon les étapes suivantes :

1. On calcule le polynôme caractéristique de A et on cherche les valeurs propres $\{\lambda_1, \dots, \lambda_r\}$. Celles-ci sont réelles car la matrice est symétrique.
2. $\mathbb{R}^n$ est somme directe des espaces propres E_{λ_i} car A est diagonalisable par le théorème spectral. De plus les espaces propres sont deux-à-deux orthogonaux car A est symétrique.
3. Pour chaque valeur propre λ_i on cherche une base orthonormée de l'espace propre E_{λ_i} .
4. La réunion des bases obtenues en (3) est une base propre orthonormée de l'espace vectoriel.

Des exemples seront vus aux exercices.

11.7.1 Application aux séries de Taylor des fonctions de n variables.

Soit $\Omega \subset \mathbb{R}^n$ un domaine ouvert de $\mathbb{R}^n$ et $f : \Omega \rightarrow \mathbb{R}$ une fonction continue ayant des dérivées continues jusqu'à l'ordre 3. Le *développement de Taylor d'ordre 1* au voisinage d'un point $p \in \Omega$ s'écrit

$$f(p+v) = f(p) + df_p(v) + O(\|v\|^2).$$

où $df_p \in (\mathbb{R}^n)^*$ est la *différentielle de f en p* . C'est le covecteur défini par

$$df_p(v) = \sum_{i=1}^n \frac{\partial f}{\partial x_i}(p) v_i.$$

Ce covecteur représente l'approximation linéaire de f pour un petit "accroissement" $p+v$ de p . Le *développement de Taylor d'ordre 2* en p s'écrit

$$f(p+v) = f(p) + df_p(v) + \frac{1}{2} h_p(v, v) + o(\|v\|^2),$$

où h_p est le *hessien* de f en p . C'est la forme bilinéaire symétrique définie par

$$h_p(v, w) = \sum_{i,j=1}^n \frac{\partial^2 f}{\partial x_i \partial x_j}(p) v_i v_j.$$

La matrice de Gram de la forme bilinéaire h_p est la *matrice hessienne* de f en p , c'est-à-dire la matrice des dérivées secondes :

$$H_p = \left(\frac{\partial^2 f(p)}{\partial x_i \partial x_j} \right).$$

On peut réécrire le développement de Taylor sous la forme suivante :

$$f(x) = f(p) + \sum_{i=1}^n \frac{\partial f(p)}{\partial x_i} (x_i - p_i) + \frac{1}{2} \sum_{i,j=1}^n \frac{\partial^2 f(p)}{\partial x_i \partial x_j} (x_i - p_i)(x_j - p_j) + o(\|x - p\|^3).$$

C'est une formule importante à comprendre et connaître.

Le théorème spectral, nous dit qu'on peut réduire le hessien à une forme diagonale en faisant un changement de coordonnées orthonormé. De façon plus précise, il existe un système de coordonnées $y_i = \sum_j p_{ij} x_j$, où $P = (p_{ij})$ est une matrice orthogonale, telles que dans ces coordonnées le développement de Taylor à l'ordre 2 s'écrit

$$\tilde{f}(y) = \tilde{f}(q) + \sum_{i=1}^n \frac{\partial \tilde{f}}{\partial y_i}(q)(y_i - q_i) + \frac{1}{2} \sum_{i=1}^n \frac{\partial^2 \tilde{f}}{\partial y_i^2}(q)(y_i - q_i)^2 + o(\|y - q\|^3).$$

Cette écriture permet de déterminer les points où la fonction f atteint un maximum local ou un minimum local (il faut que la différentielle en ce point possède soit nulle, puis on examine les signes des valeurs propres de H).

11.7.2 Application : le tenseur d'inertie et les moments d'inerties principaux

En mécanique, l'étude de la rotation d'un solide indéformable autour d'un axe met en évidence l'importance de la notion de *moment cinétique*. Considérons un solide indéformable qui est représenté par un domaine borné $D \subset \mathbb{R}^3$. La masse totale de ce solide est donnée par l'intégrale³

$$M = \int_D \rho(x) dx,$$

où la fonction $\rho : D \rightarrow \mathbb{R}_+$ représente la distribution (ou densité) de masse (si la masse se mesure en kg, alors l'unité de la fonction ρ est kg/m³). Le *centre de gravité*, ou *barycentre* du solide D est le point de $\mathbb{R}^3$ défini par

$$C = \frac{1}{M} \int_D x \rho(x) dx.$$

3. Il s'agit ici d'une notation allégée pour l'intégrale triple

$$\iiint_D \rho(x_1 x_2 x_3) dx_1 dx_2 dx_3.$$

Ses coordonnées sont $C = (c_1, c_2, c_3)$, avec $c_i = \frac{1}{M} \int_D x_i \rho(x) dx$. Dans la suite on supposera que $C = 0$ (l'origine des coordonnées). Cette hypothèse revient à translater si nécessaire le solide D pour ramener son centre de gravité à l'origine des coordonnées.

Définition. On appelle *moment d'inertie*⁴ du solide indéformable D en direction du vecteur unité u la quantité

$$\mathcal{I}_D(u) = \int_D \delta_u(x)^2 \rho(x) dx,$$

où $\delta_u(x)$ est la distance entre le point x et l'axe $\mathbb{R} \cdot u$.

Lemme 11.7.5. *Le moment d'inertie de D en direction de u peut aussi s'écrire*

$$\mathcal{I}_D(u) = \mathcal{J}_D(u, u),$$

où $\mathcal{J}_D$ est la forme bilinéaire définie sur $\mathbb{R}^3$ par

$$\mathcal{J}_D(u, v) = \int_D (\|x\|^2 \langle u, v \rangle - \langle x, u \rangle \langle x, v \rangle) \rho(x) dx. \quad (11.5)$$

Preuve. On sait que la composante normale de x selon le vecteur u est $x - \langle x, u \rangle u$ (car on suppose $\|u\| = 1$), on a donc par le théorème de Pythagore

$$\|x\|^2 = \delta_u(x)^2 + \langle x, u \rangle^2.$$

Le moment d'inertie peut donc s'écrire

$$\mathcal{I}_D(u) = \int_D \delta_u(x)^2 \rho(x) dx = \int_D (\|x\|^2 - \langle x, u \rangle^2) \rho(x) dx = \mathcal{J}_D(u, u).$$

□

Définition. La forme bilinéaire $\mathcal{J}_D : \mathbb{R}^3 \times \mathbb{R}^3 \rightarrow \mathbb{R}$ définie par (11.5) s'appelle le *tenseur d'inertie* du solide D . Il ne dépend que de la forme de D et de la distribution de masse ρ .

Le tenseur d'inertie de D est clairement une forme bilinéaire symétrique sur $\mathbb{R}^3$. Le théorème spectral nous dit par conséquent qu'il existe une base orthonormée $\{u_1, u_2, u_3\}$ de $\mathbb{R}^3$ qui orthogonalise $\mathcal{J}_D$, c'est-à-dire :

$$\mathcal{J}_D(u_1, u_1) = \mathcal{J}_1, \quad \mathcal{J}_D(u_2, u_2) = \mathcal{J}_2, \quad \mathcal{J}_D(u_3, u_3) = \mathcal{J}_3, \quad \mathcal{J}_D(u_i, u_j) = 0 \quad \text{si } i \neq j.$$

On appelle $\mathcal{J}_1, \mathcal{J}_2, \mathcal{J}_3$ les *moments principaux d'inertie* et les directions de u_1, u_2, u_3 les *axes principaux d'inertie* du solide indéformable D . Dans la base $\{u_1, u_2, u_3\}$, la matrice de Gram de $\mathcal{J}_D$ est la matrice diagonale

$$\begin{pmatrix} \mathcal{J}_1 & 0 & 0 \\ 0 & \mathcal{J}_2 & 0 \\ 0 & 0 & \mathcal{J}_3 \end{pmatrix}.$$

Supposons maintenant que le solide D est en rotation autour d'un axe passant par le centre de gravité, et notons ω le vecteur de rotation instantanée. Alors le *moment cinétique* est le vecteur défini par

$$L_D(\omega) = \int_D (\|x\|^2 \omega - \langle x, \omega \rangle x) \rho(x) dx.$$

4. Comparer avec la formule (1.57), page 54 du livre de mécanique de J.P. Ansermet. Noter que dans ce livre la formule est écrite pour un système fini de masses ponctuelles et non une densité continue de masse

Observons que pour tout vecteur $v \in \mathbb{R}^3$ on a

$$\mathcal{J}_D(\omega, v) = \langle L_D(\omega), v \rangle$$

Notons $\omega = \omega_1 u_1 + \omega_2 u_2 + \omega_3 u_3$, alors on a

$$L_D(\omega) = \sum_{i=1}^3 \langle L_D(\omega), u_i \rangle u_i = \sum_{i=1}^3 \mathcal{J}_D(\omega, u_i) u_i = \sum_{j=1}^3 \sum_{i=1}^3 \omega_j \mathcal{J}_D(u_j, u_i) u_i$$

Le moment cinétique se calcule donc dans la base $\{u_1, u_2, u_3\}$ à partir des moments principaux d'inertie par la formule

$$L_D(\omega) = \sum_{j=1}^3 \mathcal{J}_i \cdot \omega_j.$$

En l'absence de force extérieure, le moment cinétique est conservé. Si des forces extérieures sont appliquées alors la variation du moment cinétique est égale à la somme des moments de forces. Cette relation s'écrit

$$\frac{dL_D}{dt} = \int_D x \times F(x) dx.$$

Cette équation détermine la dynamique du solide en rotation autour de son centre de gravité, qui est supposé fixe. Si le centre de gravité est lui aussi en mouvement, alors on doit ajouter l'équation de Newton :

$$M \frac{d^2 x_0}{dt^2} = \int_D F(x) dx = \text{force extérieure totale agissant sur le solide},$$

où $x_0(t)$ représente la position du centre de gravité au temps t . Ces deux dernières équations donnent une description complète de l'évolution d'un solide indéformable en mouvement.

Chapitre 12

Espaces vectoriels pseudo-euclidiens

12.1 Formes quadratiques sur un espace vectoriel réel, théorème de Sylvester.

Soit Q une forme quadratique sur un espace vectoriel réel de dimension finie V . Le théorème 10.4.2 nous dit qu'il existe une base $\{v_1, \dots, v_n\}$ de V qui est orthogonale pour Q . Dans cette base on a

$$x = \sum_{i=1}^n x_i v_i \quad \Rightarrow \quad Q(x) = \sum_{i=1}^n \alpha_i x_i^2.$$

Il est habituel de noter p le nombre de coefficients α_i qui sont positifs et q le nombre de coefficients α_i qui sont négatifs.

Définition 12.1.1. (1) Le couple (p, q) s'appelle la *signature*¹ de Q .

(2) La somme $r = p + q$ s'appelle le *rang* de Q .

(3) La forme quadratique Q est *non dégénérée* si $r = \dim(V)$.

(4) Q est *positive* si $q = 0$ et *négative* si $p = 0$.

(5) La forme quadratique Q est *définie positive* si elle est positive et non dégénérée.

(6) Q est *définie négative* si elle est négative et non dégénérée.

Remarques 1. Une notation utile est la suivante : On dit que $Q > 0$ sur le sous-espace vectoriel $W \subset V$ si la restriction de Q à W est définie positive, c'est-à-dire $Q(x) > 0$ pour tout $x \in W \setminus \{0\}$. On dira dans ce cas que le sous-espace vectoriel W est défini positif pour Q . De même on dit $Q < 0$ si $(-Q) > 0$ sur W .

2. On dit aussi que Q est *semi-définie positive* (resp. semi-définie négative) si $Q(x) \geq 0$ pour tout $x \in V$ (respectivement $Q(x) \leq 0$ pour tout $x \in V$). Il est clair qu'une forme quadratique de signature (p, q) est semi-définie positive si et seulement si $q = 0$ et semi-définie négative si et seulement si $p = 0$.

1. ne pas confondre avec la signature d'une permutation.

Le résultat suivant justifie ces définitions.

Théorème 12.1.2 (Théorème d'inertie de Sylvester). *La signature (p, q) ne dépend que de la forme quadratique Q et non de la base choisie. Plus précisément, nous avons la caractérisation suivante de la signature :*

- (i) p est la dimension maximale d'un sous-espace vectoriel de V sur lequel Q est définie positive.
- (ii) q est la dimension maximale d'un sous-espace vectoriel de V sur lequel Q est définie négative.

Preuve. Nous allons démontrer que le coefficient p est la dimension maximale d'un sous-espace vectoriel sur lequel Q est définie positive. Observons tout d'abord que, quitte à permuter les vecteurs de la base orthogonale, on peut supposer que

$$\begin{cases} \alpha_i > 0 & \text{pour } 1 \leq i \leq p, \\ \alpha_i < 0 & \text{pour } p < i \leq r = p + q, \\ \alpha_i = 0 & \text{pour } i > r. \end{cases}$$

Dans cette base on peut écrire

$$Q(x) = \sum_{i=1}^r \alpha_i x_i^2 = \sum_{i=1}^p \alpha_i x_i^2 - \sum_{j=p+1}^r |\alpha_j| x_j^2.$$

Soit maintenant $W \subset V$ un sous-espace vectoriel de dimension maximale tel que la forme quadratique Q restreinte à W est définie positive. Il est clair que $\dim(W) \geq p$ car la restriction de Q au sous-espace $\text{Vec}\{v_1, \dots, v_p\}$ est définie positive.

Nous allons prouver par que $\dim(W) = p$. Supposons par l'absurde que $\dim(W) > p$ et notons $U = \text{Vec}\{v_{p+1}, \dots, v_n\}$. Alors $\dim(U) = n - p$ et donc $\dim(W \cap U) > 0$ car

$$\dim(W \cap U) = \dim(W) + \dim(U) - \dim(W + U) > p + (n - p) - n = 0$$

(on utilise que $\dim(W + U) \leq n$ puisque $W + U \subset V$). Il existe donc un vecteur non nul $x \in W \cap U$, mais ceci est impossible car

$$x \in W \setminus \{0\} \Rightarrow Q(x) > 0 \quad \text{et} \quad x \in U \Rightarrow Q(x) \leq 0.$$

(rappelons que par hypothèse Q est définie positive sur W). Cette contradiction montre que

$$p = \max\{\dim(W) \mid W \subset V \text{ est un sous-espace vectoriel tel que } Q > 0 \text{ sur } W\}.$$

Un argument similaire montre que q est la dimension maximale d'un sous-espace vectoriel sur lequel $Q < 0$. On a ainsi obtenu une caractérisation de la signature (p, q) d'une forme quadratique qui ne dépend pas du choix d'une base orthogonale, ce qui prouve le théorème de Sylvester. $\square$

Le concept de signature peut aussi se définir pour les formes bilinéaires symétriques :

Définition 12.1.3. La *signature* (p, q) d'une forme bilinéaire symétrique $g : V \times V \rightarrow \mathbb{R}$ sur un espace vectoriel réel de dimension finie V est la signature de la forme quadratique associée $Q(x) = g(x, x)$. On définit de même les notions de *rang*, de *forme bilinéaires symétrique (non) dégénérée* et *définie positive (négative)*.

Remarque. Observons que la forme quadratique Q est définie positive (i.e. de signature $(n, 0)$) si et seulement si la forme bilinéaire associée est un produit scalaire. La norme associée à ce produit scalaire est alors $\|x\| = \sqrt{Q(x)}$.

Du point de vue matriciel, le théorème de Sylvester s'énonce ainsi :

Corollaire 12.1.4. Toute matrice symétrique $A \in M_n(\mathbb{R})$ est congruente à une matrice diagonale de type

$$H_{p,q} = \begin{pmatrix} I_p & 0 & 0 \\ 0 & -I_q & 0 \\ 0 & 0 & 0 \end{pmatrix}.$$

Les nombres (p, q) ne dépendent que de A et non de la base orthogonale choisie.

Définition. Le couple (p, q) s'appelle alors la *signature* (p, q) de la matrice symétrique $A \in M_n(\mathbb{R})$. On dit que la matrice A est *définie positive* si $(p, q) = (n, 0)$.

Preuve. Notons β la forme bilinéaire sur $\mathbb{R}^n$ dont la matrice de Gram est A . On peut choisir une base $\{v_1, \dots, v_n\}$ de $\mathbb{R}^n$ telle que $\beta(v_i, v_j) = 0$ si $i \neq j$ et

$$\beta(v_i, v_i) > 0, \quad \beta(v_j, v_j) < 0, \quad \beta(v_k, v_k) = 0,$$

pour $1 \leq i \leq p < j \leq (p+q) < k \leq n$. On définit alors une nouvelle base $\{w_1, \dots, w_n\}$ par

$$w_i = \frac{v_i}{\sqrt{\beta(v_i, v_i)}}, \quad w_j = \frac{v_j}{\sqrt{-\beta(v_j, v_j)}}, \quad w_k = v_k,$$

(avec $1 \leq i \leq p < j \leq (p+q) < k \leq n$). Il est facile de vérifier que la matrice de β dans cette base est la matrice $H_{p,q}$.

□

Définition 12.1.5. Soit g une forme bilinéaire symétrique sur un espace vectoriel réel V de dimension n . On dit qu'une base $\{e_1, \dots, e_n\}$ de V est une *base de Sylvester*, ou une *base orthonormale généralisée* si

$$g(e_i, e_j) = \begin{cases} +1, & \text{si } 1 \leq i = j \leq p \\ -1, & \text{si } p+1 \leq i = j \leq p+q \\ 0, & \text{sinon,} \end{cases}$$

où (p, q) est la signature de g .

Le corollaire 12.1.4 nous garantit l'existence d'une base de Sylvester (non unique) pour toute forme bilinéaire symétrique sur un espace vectoriel réel V de dimension finie.

Proposition 12.1.6. (i) Deux matrices symétriques réelles $A, B \in M_n(\mathbb{R})$ sont congruentes si et seulement si elles ont la même signature.

(ii) Si $A \in M_n(\mathbb{R})$ est une matrice symétrique de signature (p, q) , alors p est le nombre de valeurs propres strictement positives de A comptées avec multiplicités et q est le nombre de valeurs propres strictement négatives comptées avec multiplicités.

Preuve. (i) L'affirmation découle du fait que toute matrice symétrique $A \in M_n(\mathbb{R})$ est congruente à sa forme de Sylvester et que la congruence est une relation d'équivalence.

(ii) Par le théorème spectral on sait que toute matrice symétrique $A \in M_n(\mathbb{R})$ est orthogonalement diagonalisable, donc à la fois semblable et congruente à une matrice diagonale D qui contient donc les valeurs propres de A , répétées autant de fois que leur multiplicité. □

Rappelons que pour une matrice réelle symétrique, la multiplicité géométrique de chaque valeur propre est égale à sa multiplicité algébrique (c'est une application du théorème spectral, qui nous dit en particulier qu'une telle matrice est diagonalisable).

Voici un exemple pour le point (ii) de cette proposition. Considérons la matrice symétrique

$$A = \begin{pmatrix} 3 & -2 & 4 \\ -2 & 6 & 2 \\ 4 & 2 & 3 \end{pmatrix}.$$

Son polynôme caractéristique est $\chi_A(t) = (t-7)^2(t+2)$, en particulier $\det(A) \neq 0$ et donc A est de rang 3. Les valeurs propres sont $+7$ avec multiplicité 2 et -2 avec multiplicité 1. La signature de A est donc $(p, q) = (2, 1)$. En particulier A n'est pas définie positive.

12.2 Espaces pseudo-euclidiens

Définitions 1. On appelle *espace vectoriel pseudo-euclidien* un espace vectoriel V sur le corps $\mathbb{R}$ de dimension finie muni d'une forme quadratique Q non dégénérée. L'espace (V, Q) est dit *euclidien* si Q est définie positive (i.e. $Q(x) > 0$ pour tout $x \in V$ non nul).

2. Une application affine $f : V_1 \rightarrow V_2$ entre deux espaces pseudo-euclidiens (V_1, Q_1) et (V_2, Q_2) est une *isométrie* si pour tous $x, y \in V_1$ on a

$$Q_2(f(y) - f(x)) = Q_1(y - x).$$

Lorsque f est linéaire, cette condition peut s'écrire $Q_1 = Q_2 \circ f$, c'est-à-dire $Q_2(f(x)) = Q_1(x)$ pour tout $x \in V_1$.

Remarque. Lorsque (V, Q) est euclidien, on a les notions de norme d'un vecteur $\|x\| = \sqrt{Q(x)}$ et de distance $d(x, y) = \|y - x\|$ entre deux points. Une isométrie entre deux espaces euclidiens est une bijection qui respecte les distances. Dans le cas général $Q(x)$ n'est pas forcément positif, toutefois même lorsqu'il n'y a pas de norme ou de distance associée à une forme quadratique, celle-ci peut représenter des quantités géométriques intéressantes. Lorsque deux espaces pseudo-euclidiens sont isométriques, on considère que leurs géométries sont équivalentes (par exemple tous les espaces euclidiens de même dimension sont isométriques, leur géométrie sont donc équivalentes à celle de l'espace $\mathbb{R}^n$ muni de son produit scalaire standard).

Les développements du chapitre 10 nous ont appris qu'il y a équivalence entre les trois points de vues suivants :

- (1) La théorie des espaces pseudo-euclidiens.
- (2) La théorie des formes bilinéaires symétriques sur un espace vectoriel réel de dimension finie.
- (3) L'étude des matrices carrées symétriques à coefficients réels de déterminant non nul.

Rappelons ces équivalences :

- Si Q est une forme quadratique sur V , alors il existe une unique forme bilinéaire symétrique, que l'on notera par $g : V \times V \rightarrow \mathbb{R}$, telle que $Q(x) = g(x, x)$ pour tout $x \in V$. Cette forme bilinéaire est donnée par les formules de polarisation, par exemple

$$g(x, y) = \frac{1}{4} (Q(x + y) - Q(x - y)).$$

- Si $\{v_1, \dots, v_n\}$ est une base de V , alors la *matrice de Gram* de Q (ou de g) dans cette base est la matrice $G \in M_n(\mathbb{R})$ définie par

$$G = (g_{ij}), \quad \text{avec } g_{ij} = g(v_i, v_j).$$

Cette matrice est clairement symétrique, i.e. $G^\top = G$ car $g_{ji} = g(v_j, v_i) = g(v_i, v_j) = g_{ij}$. La matrice de Gram est donc déterminée par g .

- Inversement on peut calculer $g(x, y)$ à partir de la formule

$$g(x, y) = \sum_{i,j=1}^n g_{ij} x_i y_j = X^\top G Y$$

où $(x_1, \dots, x_n)$ sont les composantes de x dans la base $\{v_1, \dots, v_n\}$ (c'est-à-dire $x = \sum_{i=1}^n x_i v_i$) et $X \in \mathbb{R}^n$ est le vecteur-colonne de $\mathbb{R}^n$ associé, (de même pour $(y_1, \dots, y_n)$ et Y).

- La condition de non dégénérescence de Q (ou de g) signifie que pour tout $x \in V \setminus \{0\}$ on peut trouver $y \in V$ tel que $g(x, y) \neq 0$.
- Il est facile de vérifier que g est non dégénéré si et seulement $\det(G) \neq 0$, i.e. la matrice de Gram G est inversible.

Nous avons alors le résultat suivant :

Proposition 12.2.1. *Soient (V_1, Q_1) et (V_2, Q_2) deux espaces pseudo-euclidiens de dimension n et $f : V_1 \rightarrow V_2$ un isomorphisme linéaire. Alors les conditions suivantes sont équivalentes :*

- (1) f est une isométrie, i.e. $Q_2 \circ f = Q_1$.
- (2) $g_2(f(x), f(y)) = g_1(x, y)$ pour tous $x, y \in V_1$, où g_i est la forme bilinéaire associée à Q_i (pour $i = 1, 2$).
- (3) Les matrices de Gram de g_1 et g_2 dans des bases $\mathcal{B}_1 \subset V_1$ et $\mathcal{B}_2 \subset V_2$ sont reliées par

$$G_1 = A^\top G_2 A,$$

où $A = M_{\mathcal{B}_2, \mathcal{B}_1}(f)$ est la matrice de f dans ces bases. En particulier les matrices G_1 et G_2 sont congruentes et la congruence est réalisée par la matrice de l'endomorphisme f .

La preuve est un simple exercice.

Remarque. Lorsque $V_1 = V_2 = \mathbb{R}^n$ et $\mathcal{B}_1 = \mathcal{B}_2 =$ la base canonique, la formule de congruence est évidente car

$$g_2(f(X), f(Y)) = (AX)^\top G_2 AX = X^\top (A^\top G_2 A) Y = X^\top G_1 Y.$$

Si (V, Q) est un espace pseudo-euclidien, l'ensemble des isométries linéaires de V dans lui même forme un groupe, que l'on appelle le *groupe orthogonal de Q* . On le note

$$O(Q) = \{f \in GL(V) \mid Q \circ f = Q\}.$$

Dans le cas où $V = \mathbb{R}^n$ on peut identifier tout endomorphisme $f \in \mathcal{L}(\mathbb{R}^n)$ avec sa matrice A dans la base canonique. On peut donc écrire

$$O(Q) = \{A \in GL_n(\mathbb{R}) \mid A^\top G A = G\}.$$

Le *groupe spécial orthogonal de Q* est le sous-groupe

$$SO(Q) = O(Q) \cap SL_n(\mathbb{R}) = \{A \in GL_n(\mathbb{R}) \mid A^\top G A = G \text{ et } \det(A) = 1\}.$$

12.3 Base de Sylvester et espaces pseudo-euclidiens modèles.

Rappelons qu'une base $\{v_1, \dots, v_n\}$ d'un espace vectoriel pseudo-euclidien (V, Q) est une *base de Sylvester* (ou une *base orthonormée généralisée*) si $Q(v_i) = \pm 1$ pour tout i et $v_i \perp_Q v_j$ si $i \neq j$. Le théorème de Sylvester nous dit que tout espace vectoriel pseudo-euclidien admet des bases de Sylvester. De plus le nombre p d'éléments de la base tels que $Q(v_i) = +1$ et le nombre q d'éléments tels que $Q(v_j) = -1$ ne dépendent pas de la base choisie. Le couple (p, q) est la *signature* de la forme quadratique Q et nous avons $p + q = n$ car Q est supposée non dégénérée. Il suit du théorème de Sylvester que tout espace pseudo-euclidien est isométrique à l'espace vectoriel $\mathbb{R}^n$ muni de la forme quadratique standard de signature $(p, q) = (p, n - p)$:

$$Q(x) = \sum_{i=1}^p x_i^2 - \sum_{j=p+1}^n x_j^2.$$

On note $\mathbb{E}^{p,q}$ cet espace et on considère que c'est l'*espace pseudo-euclidien modèle* (ou standard) de signature (p, q) . La forme bilinéaire symétrique associée est

$$g(x, y) = \sum_{i=1}^p x_i y_i - \sum_{j=p+1}^n x_j y_j,$$

et la matrice de Gram dans la base canonique est la matrice

$$H_{p,q} = I_p \oplus (-I_q) = \begin{pmatrix} 1 & & & & & 0 \\ & \ddots & & & & \\ & & 1 & & & \\ & & & -1 & & \\ & & & & \ddots & \\ 0 & & & & & -1 \end{pmatrix} \quad (12.1)$$

c'est-à-dire

$$H_{p,q} = (\eta_{ij}), \quad \text{avec} \quad \eta_{ij} = \begin{cases} 1 & \text{si } i = j \leq p, \\ -1 & \text{si } i = j > p, \\ 0 & \text{si } i \neq j. \end{cases} \quad (12.2)$$

(*symbole de Kronecker généralisé* de signature (p, q)). Remarquons que $\mathbb{E}^{n,0}$ est l'espace euclidien $\mathbb{R}^n$ muni de son produit scalaire standard, on le note simplement $\mathbb{E}^n$.

Le groupe des isométries linéaires de $\mathbb{E}^{p,q}$ se note $O(p, q)$:

$$O(p, q) = \{A \in GL_n(\mathbb{R}) \mid A^\top H_{p,q} A = H_{p,q}\}.$$

Le sous-groupe des isométries de déterminant 1 est

$$SO(p, q) = O(p, q) \cap SL_n(\mathbb{R}) = \{A \in O(p, q) \mid \det A = +1\}.$$

12.4 Indicatrices et cône isotrope

Définition. Soit (V, Q) un espace pseudo-euclidien.

(1) On appelle *cône isotrope* de (V, Q) l'ensemble des vecteurs isotropes. On le note

$$S_0(V, Q) = \{x \in V \mid Q(x) = 0\}.$$

(2) On appelle *indatrice positive* de (V, Q) l'ensemble

$$S_+(V, Q) = \{x \in V \mid Q(x) = 1\}.$$

(3) L' *indatrice négative* de (V, Q) est l'ensemble

$$S_-(V, Q) = \{x \in V \mid Q(x) = -1\}.$$

Remarques.

- (i) Le cône isotrope et les indicatrices ne sont pas des sous-espaces vectoriels de V .
- (ii) Si $x \in S_0(V, Q)$, alors $\lambda x \in S_0(V, Q)$ pour tous $\lambda \in \mathbb{R}$.
- (iii) Si $x \in S_+(V, Q)$, alors $\lambda x \in S_+(V, Q)$ si et seulement si $\lambda = \pm 1$. La même propriété est vraie pour $S_-(V, Q)$.
- (iv) Les ensembles $S_0(V, Q)$, $S_+(V, Q)$ et $S_-(V, Q)$ déterminent complètement la forme quadratique Q , i.e. si Q_1 et Q_2 sont deux formes quadratiques telles que

$$S_0(V, Q_1) = S_0(V, Q_2), \quad S_+(V, Q_1) = S_+(V, Q_2), \quad S_-(V, Q_1) = S_-(V, Q_2),$$

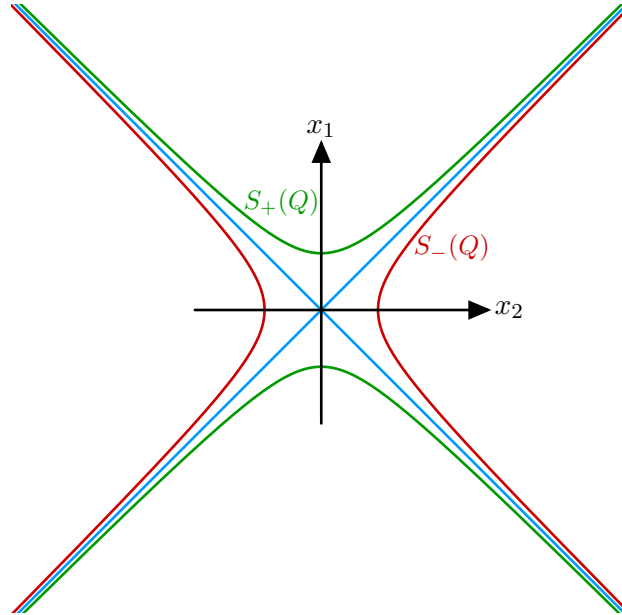
alors $Q_1 = Q_2$.

- (v) Les ensembles $S_0(V, Q)$, $S_+(V, Q)$ et $S_-(V, Q)$ sont invariants par l'action du groupe $O(Q)$, c'est-à-dire que si $f \in O(Q)$, alors $x \in S_+(V, Q)$ si et seulement si $f(x) \in S_+(V, Q)$.

Exemples.

- (1) Pour le plan euclidien, on a $S_-(\mathbb{E}^2) = \emptyset$, $S_0(\mathbb{E}^2) = \{0\}$ et $S_+(\mathbb{E}^2)$ est le cercle unité.

- (2) Plus généralement l'indicatrice S_+ d'un espace euclidien $\mathbb{E}^n$ est la *sphère unité* ; c'est l'ensemble des points de $\mathbb{E}^n$ dont la distance à l'origine 0 est égale à 1. L'indicatrice négative est l'ensemble vide et $S_0(\mathbb{E}^n) = \{0\}$.
- (3) Pour $\mathbb{E}^{1,1}$, $S_0(\mathbb{E}^{1,1})$ est la réunion des deux droites $\{x_2 = \pm 1\}$ et $S_{\pm}(\mathbb{E}^{1,1})$ sont deux hyperboles dont les asymptotes sont les droites du cône isotrope.
- (4) Pour $\mathbb{E}^{1,2}$, $S_0(\mathbb{E}^{1,2})$ est le cône circulaire droit $\{x_1^2 = x_2^2 + x_3^2\}$, l'indicatrice négative $S_-(\mathbb{E}^{1,2})$ est une hyperboloïde de révolution à deux nappes et l'indicatrice positive $S_+(\mathbb{E}^{1,2})$ est une hyperboloïde de révolution à une nappe.



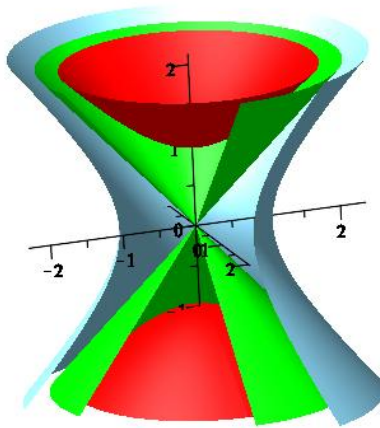


FIGURE 12.1 – Cône isotrope et indicatrices de $\mathbb{E}^{1,2}$.

12.5 L'espace-temps de Lorentz-Minkowski $\mathbb{E}^{1,d}$

Définition. On appelle *espace-temps de Lorentz-Minkowski* (ou simplement espace de Minkowski) de dimension $d + 1$ tout espace pseudo-euclidien de signature $(1, d)$. Dans une base adaptée, on peut donc écrire la forme quadratique

$$Q(x) = c^2 t^2 - x_1^2 - \dots - x_d^2.$$

Un élément $x = (t, x_1, \dots, x_d) \in \mathbb{E}^{1,d}$ s'appelle un *événement*.

L'interprétation physique est la suivante : Les coordonnées $x_1, \dots, x_d$ représentent des coordonnées de l'espace (on supposera en général que $d \leq 3$) et t représente une coordonnée temporelle. Le paramètre c est la vitesse maximale de propagation d'un signal dans l'espace temps. L'expérience nous apprend que c est la vitesse de la lumière.

Il est commode de noter $x_0 = ct$, alors la forme quadratique fondamentale s'écrit

$$Q(x) = x_0^2 - x_1^2 - \dots - x_d^2,$$

cela revient essentiellement à choisir des unités telles que $c = 1$.

Le groupe des isométries linéaires de l'espace de Minkowski est le groupe $O(1, d)$. On l'appelle le *groupe de Lorentz*. Le *principe de relativité d'Einstein* dit que les lois de la physique ne doivent pas dépendre du référentiel choisi. Ce principe se traduit mathématiquement de la manière suivante :

Les notions attachées à l'espace-temps qui ont une signification physique doivent être invariantes sous l'action du groupe de Lorentz.

Les définitions suivantes modélisent les notions liées à la causalité dans l'espace-temps de Lorentz-Minkowski.

Définitions.

- (1) On dit que deux événements $x, y \in \mathbb{E}^{1,d}$ sont *en relation de causalité* si un signal émis depuis l'un des événements peut atteindre l'autre événement. La condition s'écrit mathématiquement par

$$Q(y - x) \geq 0.$$

- (2) On considère qu'un événement ne peut pas être cause d'un événement passé mais seulement d'un événement futur. On dira en conséquence qu'un événement y est dans le *futur causal* de x si

$$Q(y - x) \geq 0 \quad \text{et} \quad y_0 \geq x_0.$$

L'ensemble des événements dans le futur causal de x s'appelle le *cône futur* de x et se note

$$\mathcal{C}_x = \{y \in \mathbb{E}^{1,d} \mid Q(y - x) \geq 0 \text{ et } y_0 \geq x_0\}.$$

- (3) Le *cône isotrope* (ou cône de lumière) issu de x est l'ensemble des y tels que $Q(y - x) = 0$.
(4) Si y est dans le futur de x , on appelle *temps propre* ou *intervalle spatio-temporel* la quantité

$$\tau(x, y) = \sqrt{Q(y - x)} = \sqrt{(y_0 - x_0)^2 - (y_1 - x_1)^2 - \dots - (y_d - x_d)^2}.$$

On utilise aussi la terminologie suivante :

- (i) Un vecteur x est de *type temps* si $Q(x) > 0$.
- (ii) Un vecteur x est de *type lumière*, ou *isotrope* si $Q(x) = 0$.
- (iii) Un vecteur x est de *type espace* si $Q(x) < 0$.

Ces notions ont un sens physique car elles sont invariantes sous l'action du groupe de Lorentz. Par exemple si $f \in O(1, d)$, alors x est de type temps si et seulement si $f(x)$ est de type temps.

La notion de ligne d'univers

Considérons une particule, ou un objet quelconque, qui se déplace au cours du temps. Sa trajectoire dans l'espace est représentée dans un certain référentiel par la fonction $t \mapsto x(t) = (x_1(t), \dots, x_d(t)) \in \mathbb{R}^d$. On appelle *ligne d'univers* de cette particule la fonction

$$t \mapsto x(t) = (t, x_1(t), \dots, x_d(t)) \in \mathbb{E}^{1,d}$$

Proposition 12.5.1. *La ligne d'univers de toute trajectoire physiquement réalisable vérifie la condition suivante :*

$$\text{pour tout } t_2 > t_1, \text{ on a } x(t_2) \in \mathcal{C}_{x(t_1)}.$$

On dit qu'une ligne d'univers est *inertielle*, si elle représente une droite de $\mathbb{E}^{1,d}$ cette droite doit être contenue dans le cône de lumière de chacun de ses points.

12.6 L'inégalité de Cauchy-Schwarz inversée et quelques conséquences

Théorème 12.6.1 (Inégalité de Cauchy-Schwarz inversée). *Si $x, y \in \mathbb{E}^{1,d}$ sont isotropes ou de type temps, alors*

$$|g(x, y)| \geq \sqrt{Q(x)}\sqrt{Q(y)}, \quad (12.3)$$

où g est la forme bilinéaire associée à Q . On a égalité si et seulement si x et y sont colinéaires.

Preuve. Observons tout d'abord que si x ou y est isotrope, alors $Q(x)Q(y) = 0$ et l'inégalité est triviale. Il est par ailleurs facile de vérifier que si x et y sont colinéaires alors $Q(x)Q(y) = g(x, y)^2$. On suppose donc que x et y sont de type temps et linéairement indépendants et nous donnons deux démonstrations de l'inégalité stricte (12.3).

Première preuve : On suppose donc x et y linéairement indépendants et on note $W = \text{Vec}x, y \subset \mathbb{E}^{1,d}$ le sous-espace vectoriel engendré par x et y . La restriction de g à W peut à priori être une forme bilinéaire symétrique de signature $(p, q) = (0, 2)$ ou $(p, q) = (1, 1)$. Or la signature $(0, 2)$ est exclue car nous avons supposé que x et y sont de type temps, i.e. $Q(x) > 0$ et $Q(y) > 0$. Donc g est de signature $(1, 1)$ sur W . Par conséquent le déterminant de la matrice de Gram associée à la base $\{x, y\}$ de W est négatif, on a donc prouvé que

$$Q(x)Q(y) - g(x, y)^2 = \det \begin{pmatrix} g(x, x) & g(x, y) \\ g(x, y) & g(y, y) \end{pmatrix} < 0,$$

ce qui est équivalent à l'inégalité (12.3).

Pour la deuxième preuve, on considère la droite affine qui passe par x et de vecteur directeur y , qui est l'ensemble

$$L = \{x + sy \mid s \in \mathbb{R}\} \subset \mathbb{E}^{1,d}.$$

Si x et y ne sont pas colinéaires, alors L n'est pas contenue dans l'intérieur du cône isotrope de $\mathbb{E}^{1,d}$ et cette droite contient donc des vecteurs de type temps, de type espaces et deux vecteurs isotropes. Par conséquent la fonction

$$f(s) = Q(x + sy) = g(x + sy, x + sy) = Q(x) + 2sg(x, y) + s^2Q(y)$$

est un polynôme du second degré qui s'annule pour exactement deux valeurs de s , ce qui implique que le discriminant de $f(s)$ est strictement positif. On a donc

$$\Delta = g(x, y)^2 - Q(x)Q(y) > 0.$$

□

Lemme 12.6.2. *Si $x, y \in \mathcal{C}_0$ sont deux vecteurs du cône futur de 0, alors $g(x, y) \geq 0$.*

Preuve. L'hypothèse $x, y \in \mathcal{C}_0$ signifie que

$$\sqrt{\sum_{i=1}^d x_i^2} \leq x_0, \quad \text{et} \quad \sqrt{\sum_{i=1}^d y_i^2} \leq y_0.$$

On a donc par l'inégalité de Cauchy-Schwarz classique dans $\mathbb{R}^d$

$$\sum_{i=1}^d x_i y_i \leq \sqrt{\sum_{i=1}^d x_i^2} \sqrt{\sum_{i=1}^d y_i^2} \leq x_0 y_0,$$

ce qui implique

$$g(x, y) = x_0 y_0 - \sum_{i=1}^d x_i y_i \geq 0.$$

Corollaire 12.6.3. *Soient $x, y \in \mathcal{C}_0$ deux vecteurs du cône futur de 0, alors $Q(x + y) \geq 0$ et*

$$\sqrt{Q(x + y)} \geq \sqrt{Q(x)} + \sqrt{Q(y)}$$

Preuve. Le lemme précédent et l'inégalité de Cauchy-Schwarz inversée impliquent que

$$g(x, y) \geq \sqrt{Q(x)} \sqrt{Q(y)},$$

par conséquent on a

$$\begin{aligned} Q(x + y) &= g(x + y, x + y) \\ &= (g(x, x) + 2g(x, y) + g(y, y)) \\ &\geq (Q(x) + 2\sqrt{Q(x)}\sqrt{Q(y)} + Q(y)) \\ &= (\sqrt{Q(x)} + \sqrt{Q(y)})^2. \end{aligned}$$

□

Le paradoxe des jumeaux

Si A et B sont deux événements de l'espace-temps, on note $A \preceq B$ si B est dans le futur causal de A , i.e. si le vecteur $(B - A) \in \mathcal{C}_0$.

Lemme 12.6.4. *La relation $\preceq$ est une relation d'ordre partiel sur $\mathbb{E}^{1,d}$.*

Preuve. Il faut montrer que si B est dans le futur causal de A et C est dans le futur causal de B alors C est dans le futur causal de A , ce qui est tout-à-fait intuitif. Mathématiquement cela signifie

$$A \preceq B \text{ et } B \preceq C \Rightarrow A \preceq C.$$

Pour le démontrer, on note $x = (B - A)$, $y = (C - B)$ et $z = (C - A)$. Alors $z = x + y$ et par hypothèse on a $x, y \in \mathcal{C}_0$. On doit prouver que $z \in \mathcal{C}_0$.

$$x_0 \leq y_0 \text{ et } y_0 \leq z_0 \Rightarrow x_0 \leq z_0,$$

par conséquent $z_0 > 0$. D'autre part le corollaire précédent implique que $Q(z) \geq 0$, ce qui complète la preuve.

□

Nous pouvons maintenant énoncer le

Paradoxe des jumeaux

Si A , B et C sont trois événements tels que B est dans le futur causal de A et C est dans le futur causal de B , alors C est dans le futur causal de A et les temps propres associés sont reliés par l'*inégalité causale*, aussi appelée l'*inégalité du triangle inversée* :

$$\tau(A, C) \geq \tau(A, B) + \tau(B, C). \quad (12.4)$$

On a égalité si et seulement si l'événement B est situé sur la ligne d'univers inertielle de A vers C .

Rappelons que si A, B, C sont trois points de l'espace Euclidien $\mathbb{E}^n$, alors leurs distances respectives vérifient l'inégalité du triangle :

$$d(A, C) \leq d(A, B) + d(B, C).$$

d'où le nom de inégalité du triangle inversée pour l'inégalité (12.4)

Preuve. On pose de nouveau $x = (B - A)$, $y = (C - B)$ et $z = (C - A)$, alors on a $z = x + y$. Le corollaire 12.6.3 entraîne alors que

$$\tau(A, C) = \sqrt{Q(z)} = \sqrt{Q(x + y)} \geq \sqrt{Q(x)} + \sqrt{Q(y)} = \tau(A, B) + \tau(B, C).$$

□

Ce résultat a été appelé le “paradoxe des jumeaux” par Paul Langevin², qui l’a formulé de la manière suivante : Si deux frères jumeaux se rencontrent en un lieu précis au même moment, et si l’un des jumeaux s’en va faire un voyage cosmique puis rejoint son frère, alors à son retour l’un des jumeaux (celui qui a voyagé) est plus jeune que l’autre. Ce paradoxe a rendu les commentateurs perplexes, mais il ne signifie pas qu’il y a une contradiction dans la théorie car la situation des deux frères n’est pas symétrique. L’un a une ligne d’univers inertielle et l’autre non.

2. Paul Langevin (1872-1946) est un physicien français qui fut parmi les premiers à admettre et propager la relativité restreinte en France (avec l’exception notable de Henri Poincaré). Il a formulé le paradoxe des jumeaux à ses collègues physiciens et philosophes en 1911. Ce paradoxe a engendré une certaine perplexité et provoqué d’intéressantes discussions sur la relativité einsteinienne et son interprétation.

Chapitre 13

Espaces hermitiens, opérateurs normaux et le théorème spectral

Rappelons qu'un espace vectoriel euclidien est un espace vectoriel de dimension finie sur le corps $\mathbb{R}$. Dans ce chapitre, nous étudions une notion analogue pour les espaces vectoriels de dimension finie sur le corps $\mathbb{C}$. Nous étudions en particulier une classe d'opérateurs (i.e. d'endomorphismes) sur de tels espaces, appelés opérateurs *normaux*, pour lesquels nous pouvons généraliser le théorème spectral vu au chapitre 10. Parmi les opérateurs normaux, on peut citer les opérateurs *unitaires* et les *opérateurs auto-adjoints*. Ces opérateurs sont importants notamment en mécanique quantique pour décrire les observables d'un système quantique.

13.1 Formes sesquilineaires et formes hermitiennes sur un espace vectoriel complexe.

Définition 13.1.1. Soient V et W deux espaces vectoriels sur le corps $\mathbb{C}$ des nombres complexes. On dit qu'une application $f : V \rightarrow W$ est *antilinéaire* si elle vérifie

$$f(v + w) = f(v) + f(w) \quad \text{et} \quad f(\lambda v) = \bar{\lambda} f(v).$$

pour tous $v, w \in V$ et tous $\lambda \in \mathbb{C}$.

Remarquons que les applications anti-linéaires sont $\mathbb{R}$ -linéaires, c'est-à-dire qu'elles sont linéaires lorsqu'on regarde V et W comme des espaces vectoriels sur $\mathbb{R}$.

Exemples 1.) L'application $s : \mathbb{C}^n \rightarrow \mathbb{C}^n$ qui conjugue toutes les coordonnées, i.e. $s(z_1, \dots, z_n) = (\bar{z}_1, \dots, \bar{z}_n)$ est anti-linéaire.

2.) Si $\theta : V \rightarrow \mathbb{C}$ est une forme linéaire, alors l'application $v \mapsto \overline{\theta(v)}$ est anti-linéaire.

3.) On appelle *adjointe* d'une matrice $A \in M_n(\mathbb{C})$ la transposée de la matrice conjuguée, et on note $A^* = \bar{A}^\top$. Il est clair que $A \mapsto A^*$ est une application anti-linéaire de l'espace vectoriel $M_n(\mathbb{C})$ dans lui-même.

Définition 13.1.2. (1) Une application $h : V \times V \rightarrow \mathbb{C}$ est dite *sesquilinéaire* si elle est antilinéaire en la première variable et linéaire en la deuxième variable¹ :

$$h(\lambda v_1 + v_2, w) = \bar{\lambda}h(v_1, w) + h(v_2, w) \quad \text{et} \quad h(v, \mu w_1 + w_2) = \mu h(v, w_1) + h(v, w_2).$$

(2) L'application $h : V \times V \rightarrow \mathbb{C}$ est une *forme hermitienne*² si elle est sesquilinéaire et elle vérifie de plus

$$h(w, v) = \overline{h(v, w)}$$

pour tous $v, w \in V$.

(3) La *forme quadratique* associée à une forme hermitienne h sur V est la fonction $Q : V \rightarrow \mathbb{R}$ définie par

$$Q(w) = h(w, w).$$

Observons que $Q(w)$ est en effet un nombre réel pour tout $w \in V$, car $Q(w) = h(w, w) = \overline{h(w, w)} = \overline{Q(w)}$.

Exemple. Si $\{\varphi_1, \dots, \varphi_n\}$ est une base de l'espace dual V^* , et $a_1, \dots, a_n \in \mathbb{R}$, alors la fonction $h : V \times V \rightarrow \mathbb{C}$ définie par

$$h(x, y) = \sum_{j=1}^n a_j \overline{\varphi_j(x)} \varphi_j(y)$$

est une forme hermitienne sur V . On peut prouver que toutes les formes hermitiennes sur un espace vectoriel de dimension finie sont de ce type.

Lemme 13.1.3. La forme quadratique associée à la forme hermitienne h vérifie les propriétés suivantes :

(i) $Q(\lambda w) = |\lambda|^2 Q(w)$ pour tout $\lambda \in \mathbb{C}$ et tout $w \in V$ (en particulier $Q(\sqrt{-1}w) = Q(w)$).

(ii) On peut retrouver h à partir de Q par la formule de polarisation :

$$h(x, y) = \frac{1}{4} (Q(x + y) - Q(x - y)) - \frac{i}{4} (Q(x + iy) - Q(x - iy)), \quad (13.1)$$

où $i = \sqrt{-1}$.

Preuve. La preuve de la première affirmation est élémentaire :

$$Q(\lambda w) = h(\lambda w, \lambda w) = \bar{\lambda} \lambda h(w, w) = |\lambda|^2 Q(w).$$

Pour prouver la formule de polarisation, on considère séparément les parties réelles et imaginaires. Dans les calculs qui suivent, on utilise les identités

$$h(y, x) = \overline{h(x, y)}, \quad h(x, iy) = ih(x, y), \quad h(iy, x) = -ih(y, x).$$

1. La convention opposée est également utilisée, i.e. certains auteurs demandent que h soit linéaire en la première variable et antilinéaire en la deuxième. Ici nous suivons la convention la plus usuelles parmi les physiciens.

2. En référence au mathématicien français Charles Hermite (1822-1901).

On a d'une part

$$\begin{aligned}
Q(x+y) - Q(x-y) &= h(x+y, x+y) - h(x-y, x-y) \\
&= (h(x, x) + h(x, y) + h(y, x) + h(y, y)) - (h(x, x) - h(x, y) - h(y, x) + h(y, y)) \\
&= 2(h(x, y) + h(y, x)) \\
&= 2(h(x, y) + \overline{h(x, y)}) \\
&= 4\operatorname{Ré}(h(x, y)),
\end{aligned}$$

et d'autre part

$$\begin{aligned}
Q(x-iy) - Q(x+iy) &= h(x-iy, x-iy) - h(x+iy, x+iy) \\
&= (h(x, x) - ih(x, y) + ih(y, x) + h(y, y)) - (h(x, x) + ih(x, y) - ih(y, x) + h(y, y)) \\
&= -2i(h(x, y) - h(y, x)) \\
&= -2i(h(x, y) - \overline{h(x, y)}) \\
&= 4\operatorname{Im}(h(x, y)).
\end{aligned}$$

Donc

$$\begin{aligned}
h(x, y) &= \operatorname{Ré}(h(x, y)) + i \operatorname{Im}(h(x, y)) \\
&= \frac{1}{4}(Q(x+y) - Q(x-y)) + \frac{i}{4}(Q(x-iy) - Q(x+iy)).
\end{aligned}$$

□

13.2 Espaces vectoriels hermitiens

Définitions. Soit V un espace vectoriel complexe. On appelle *produit scalaire hermitien* sur V la donnée d'une forme hermitienne $h : V \times V \rightarrow \mathbb{C}$ qui est définie positive. Cette condition signifie que $Q(v) = h(v, v) > 0$ pour tout vecteur non nul de V .

Un *espace vectoriel hermitien* est un espace vectoriel complexe de dimension finie muni d'un produit scalaire hermitien.

On peut résumer la définition de produit scalaire hermitien dans les quatre propriétés suivantes :

- (i) $h : V \times V \rightarrow \mathbb{C}$ est $\mathbb{R}$ -bilinéaire.
- (ii) $h(x, iy) = ih(x, y) = -h(ix, y)$ (où $i = \sqrt{-1}$).
- (iii) $h(y, x) = \overline{h(x, y)}$.
- (iv) $h(x, x) > 0 \quad \forall x \in V \setminus \{0\}$,

pour tous $x, y \in V$.

Exemples 1. Le produit scalaire hermitien standard sur $\mathbb{C}^n$ est défini par

$$\langle z, w \rangle = \bar{z}_1 w_1 + \dots + \bar{z}_n w_n,$$

la forme quadratique associée est

$$Q(w) = \langle w, w \rangle = \sum_{i=1}^n \bar{w}_i w_i = \sum_{i=1}^n |w_i|^2.$$

2. Le produit scalaire L^2 sur l'espace vectoriel $\mathcal{F}$ des fonctions continues $f : [a, b] \rightarrow \mathbb{C}$ est défini par

$$(f | g) = \int_a^b \overline{f(x)} g(x) dx,$$

la forme quadratique associée est

$$Q(f) = \int_a^b |f(x)|^2 dx.$$

3. L'espace des suites complexes de carré intégrable est l'espace vectoriel

$$\ell_2(\mathbb{C}) = \{\zeta = (z_i)_{i \in \mathbb{N}} \mid z_i \in \mathbb{C}, \sum_{i=1}^{\infty} |z_i|^2 < \infty\}.$$

On munit cet espace vectoriel du produit scalaire hermitien : $\langle \zeta, \xi \rangle = \sum_{i=1}^{\infty} \bar{z}_i x_i$. La forme quadratique associée est $Q(\zeta) = \sum_{i=1}^{\infty} |z_i|^2$.

4. Un produit scalaire hermitien est défini sur $M_n(\mathbb{C})$ par

$$\langle A, B \rangle = \text{Trace}(A^* B),$$

où $A^* = \bar{A}^\top$ est la *matrice adjointe* de A .

Définition. Lorsque h est un produit scalaire hermitien sur V , on définit la *norme* d'un vecteur $w \in V$ par

$$\|w\| = \sqrt{Q(w)} = \sqrt{h(w, w)}.$$

Nous démontrons maintenant que l'inégalité de Cauchy-Schwarz est encore valable pour un produit scalaire hermitien.

Proposition 13.2.1 (Inégalité de Cauchy-Schwarz hermitienne). *Si V est un espace vectoriel complexe muni d'un produit scalaire hermitien $\langle \cdot, \cdot \rangle$, alors on a pour tous $x, y \in V$*

$$|\langle x, y \rangle| \leq \|x\| \|y\|.$$

De plus on a égalité si et seulement si x et y sont colinéaires.

Preuve. la preuve donnée dans le cas réel (voir Proposition 11.1.2) ne marche pas et doit être légèrement modifiée. On note

$$a = \langle x, y \rangle = \overline{\langle y, x \rangle},$$

et on veut montrer que $|a| \leq \|x\| \|y\|$. Si $a = 0$ il n'y a rien à montrer, sinon on pose $p(t) = \|axt + y\|^2$ où t est un paramètre réel. En utilisant les propriétés du produit scalaire hermitien, on calcule

$$p(t) = \|tax + y\|^2 = \langle tax + y, tax + y \rangle = t^2 \bar{a} a \langle x, x \rangle + t \bar{a} \langle x, y \rangle + t a \langle y, x \rangle + \langle y, y \rangle.$$

En particulier $p(t)$ est un polynôme à coefficients réel de degré 2, qui s'écrit

$$p(t) = \|x\|^2 |a|^2 t^2 + 2|a|^2 t + \|y\|^2,$$

dont le discriminant $\Delta = 4|a|^2 (|a|^2 - \|x\|^2 \|y\|^2)$ doit être négatif, c'est-à-dire $|\langle x, y \rangle| = |a| \leq \|x\| \|y\|$. De plus on a égalité si et seulement s'il existe $t \in \mathbb{R}$ tel que $y = -tax$. $\square$

On a alors les propriétés suivantes, analogues au cas du produit scalaire sur les espaces vectoriels réels.

Proposition 13.2.2. *La norme associée à un produit scalaire hermitien sur un espace vectoriel complexe V vérifie*

- (i) $\|w\| \geq 0$ pour tout $w \in V$ et $\|w\| = 0$ si et seulement si $w = 0$.
- (ii) $\|\lambda w\| = |\lambda| \|w\|$ pour tous $w \in V$ et $\lambda \in \mathbb{C}$.
- (iii) $\|v + w\| \leq \|v\| + \|w\|$.

On a aussi une version du théorème de Pythagore : si $v \perp w$, i.e. si $\langle v, w \rangle = 0$, alors

$$\|v + w\|^2 = \|v\|^2 + \|w\|^2.$$

La notion d'espace vectoriel hermitien représente donc le pendant complexe de celle d'espace vectoriel euclidien. En particulier nous avons les propositions suivantes :

Proposition 13.2.3. *Sur tout espace vectoriel hermitien (V, h) , on peut construire des bases orthonormales dans V , i.e. des bases $\{u_1, \dots, u_n\}$ telles que*

$$h(u_i, u_j) = \delta_{ij}.$$

Une telle base s'appelle aussi une base unitaire de (V, h) .

Ce résultat se démontre comme dans le cas d'un espace vectoriel euclidien, voir le théorème 11.2.2. En répétant la preuve du point (vi) du théorème 11.2.3), on obtient aussi la

Proposition 13.2.4. *Soit W un sous-espace vectoriel d'un espace vectoriel hermitien (V, h) . Alors son complément orthogonal, défini par*

$$W^\perp = \{x \in V \mid h(x, w) = 0 \ \forall w \in W\},$$

est un sous-espace vectoriel complexe de V . De plus, on a $V = W \oplus W^\perp$.

Le procédé de Gram-Schmidt peut également être étendu aux espaces hermitiens. Rappelons l'algorithme :

Étant donné une base $\{v_1, v_2, \dots, v_m\}$ de l'espace hermitien V , on construit une base unitaire $\{u_1, u_2, \dots, u_n\} \subset V$ en suivant les étapes suivantes :

1. On pose $u_1 = \frac{v_1}{\|v_1\|}$.

2. Pour $k \geq 2$, on suppose que les vecteurs $\{u_1, u_2, \dots, u_{k-1}\}$ ont été construits et on considère le vecteur $\hat{u}_k$ obtenu en soustrayant de v_k les projections des vecteurs précédents :

$$\hat{u}_k = v_k - \sum_{j=1}^{k-1} \langle v_k, u_j \rangle u_j.$$

3. On normalise le vecteur $\hat{u}_k$ pour obtenir u_k :

$$u_k = \frac{\hat{u}_k}{\|\hat{u}_k\|}.$$

13.3 Opérateurs dans les espaces hermitiens

Un endomorphisme $\mathbb{C}$ -linéaire $T : V \rightarrow V$ d'un espace vectoriel hermitien V s'appelle aussi un *opérateur* de V . Il est utile d'observer que la matrice $A = (a_{ij})$ d'un opérateur T de V dans une base unitaire $\{e_1, \dots, e_n\}$ est donnée par

$$a_{ij} = \langle e_i, T e_j \rangle. \quad (13.2)$$

La vérification se fait par un simple calcul : on a par définition $T e_j = \sum_{k=1}^n a_{kj} e_k$, par conséquent

$$\langle e_i, T e_j \rangle = \langle e_i, \sum_{k=1}^n a_{kj} e_k \rangle = \sum_{k=1}^n a_{kj} \underbrace{\langle e_i, e_k \rangle}_{=\delta_{ik}} = a_{ij}.$$

Il faut être attentif à la place de l'opérateur T dans l'équation (13.2), on a en effet $\langle T e_i, e_j \rangle = \overline{a_{ji}}$.

A tout opérateur T de V , on peut associer une application $\phi_T : V \rightarrow \mathbb{C}$ définie par

$$\phi_T(x) = \langle x, T x \rangle.$$

L'application ϕ_T vérifie $\phi_T(\lambda x) = |\lambda|^2 \phi_T(x)$ pour tout $x \in V$, ça n'est donc pas une forme quadratique au sens classique³ à valeur dans le corps $\mathbb{C}$.

Lemme 13.3.1. *L'application ϕ_T détermine l'opérateur T , i.e. si les opérateurs $T_1, T_2 \in \mathcal{L}(V)$ vérifient $\langle x, T_1 x \rangle = \langle x, T_2 x \rangle$ pour tout $x \in V$, alors $T_1 = T_2$.*

Preuve. Notons $T = (T_1 - T_2)$. Nous devons montrer que si $\langle w, T w \rangle = 0$ pour tout $w \in V$, alors $T = 0$. Fixons $\alpha \in \mathbb{C}$ et calculons

$$\langle (\alpha x + y), T(\alpha x + y) \rangle = |\alpha|^2 \langle x, T x \rangle + \bar{\alpha} \langle x, T y \rangle + \alpha \langle y, T x \rangle + \langle y, T y \rangle.$$

Par hypothèse, on a

$$\langle x, T x \rangle = \langle y, T y \rangle = \langle (\alpha x + y), T(\alpha x + y) \rangle = 0,$$

donc nous avons

$$\bar{\alpha} \langle x, T y \rangle + \alpha \langle y, T x \rangle = 0,$$

3. On dit que ϕ est une *forme quadratique hermitienne*.

pour tout $\alpha \in \mathbb{C}$, ce qui n'est possible que si $\langle x, Ty \rangle = \langle y, Tx \rangle = 0$ (prendre par exemple $\alpha = 1$, puis $\alpha = i$ pour le voir). □

Remarque. L'analogie du lemme correspondant est faux dans le cas des opérateurs $\mathbb{R}$ -linéaires dans un espace vectoriel euclidien. Dans le cas euclidien, la condition $\langle w, Tw \rangle = 0$ pour tout $w \in V$ entraîne que l'opérateur est antisymétrique. La preuve ci-dessus ne fonctionne pas car le corps de base est $\mathbb{R}$ et donc $\bar{\alpha} = \alpha$.

13.3.1 L'adjoint d'un opérateur

Définition. Soit V un espace vectoriel complexe muni d'un produit scalaire hermitien que l'on note $\langle \cdot | \cdot \rangle$ et soit $T : V \rightarrow V$ un opérateur (c'est-à-dire un endomorphisme $\mathbb{C}$ -linéaire de V). On dit que l'opérateur $T^* : V \rightarrow V$ est l'*adjoint* de T si

$$\langle Tx | y \rangle = \langle x | T^*y \rangle$$

pour tous $x, y \in V$. Par le lemme 13.3.1, on sait que l'adjoint d'un opérateur, s'il existe, est unique.

Proposition 13.3.2. *L'adjoint possède les propriétés suivantes :*

- (a) *L'adjoint de l'adjoint de T est l'opérateur T lui-même : $T^{**} = T$.*
- (b) $(TS)^* = S^*T^*$
- (c) $(S + T)^* = S^* + T^*$
- (d) $(\lambda T)^* = \bar{\lambda}T^*$ pour tout $\lambda \in \mathbb{C}$.
- (e) *Si T est inversible, alors l'inverse de l'adjoint de T est égale à l'adjoint de l'inverse de T .*

Preuve. La preuve de ces propriétés est un simple jeu formel :

- (a) On a pour tous $x, y \in V$:

$$\langle x | Ty \rangle = \overline{\langle Ty | x \rangle} = \overline{\langle y | T^*x \rangle} = \langle T^*x | y \rangle = \langle x | T^{**}y \rangle.$$

- (b) On a pour tous $x, y \in V$:

$$\langle x | (TS)^*y \rangle = \langle TSx | y \rangle = \langle Sx | T^*y \rangle = \langle x | S^*T^*y \rangle$$

- (c) La propriété (c) vient de l'additivité du produit scalaire hermitien en chaque variable :

$$\langle (S + T)x | y \rangle = \langle Sx | y \rangle + \langle Tx | y \rangle = \langle x | S^*y \rangle + \langle x | T^*y \rangle = \langle x | (S^* + T^*)y \rangle$$

- (d) L'argument est semblable :

$$\langle (\lambda T)x | y \rangle = \bar{\lambda} \langle Tx | y \rangle = \bar{\lambda} \langle x | T^*y \rangle = \langle x | \bar{\lambda}T^*y \rangle$$

- (e) Notons $S = T^{-1}$ et I l'identité de V , alors $S^* = (T^*)^{-1}$ car

$$I = I^* = (ST)^* = T^*S^*.$$

□

Proposition 13.3.3. *Si $\dim(V) < \infty$, alors tout opérateur $T : V \rightarrow V$ admet un adjoint, qui est unique.*

Remarque. En dimension infinie, il existe des opérateurs qui n'ont pas d'adjoint.

Preuve. Nous avons déjà mentionné que l'unicité suit du lemme 13.3.1. Pour montrer l'existence, on se donne une base orthonormée $\{e_1, \dots, e_n\}$ de V . Alors on a $Te_j = \sum_{i=1}^n a_{ij}e_i$ avec $a_{ij} = \langle e_i, Te_j \rangle$. Par définition de l'adjoint, en supposant son existence, on doit avoir

$$\langle e_j | T^*e_k \rangle = \langle Te_j | e_k \rangle = \left\langle \sum_{i=1}^n a_{ij}e_i | e_k \right\rangle = \sum_{i=1}^n \bar{a}_{ij} \langle e_i | e_k \rangle = \bar{a}_{kj}.$$

Par conséquent, l'opérateur $T^* : V \rightarrow V$ défini par

$$T^*e_k = \sum_{j=1}^n \bar{a}_{kj}e_j, \quad (13.3)$$

est l'adjoint de T .

□

Corollaire 13.3.4. *Soit T un opérateur d'un espace hermitien V .*

- (a) *Si $A \in M_n(\mathbb{C})$ est la matrice de l'opérateur T dans une base orthonormée $\{e_1, \dots, e_n\} \subset V$, alors la matrice transposée-conjuguée $\bar{A}^\top$ est la matrice de l'opérateur T^* dans la même base.*
- (b) *Si λ est valeur propre de T , alors $\bar{\lambda}$ est valeur propre de T^* .*

Preuve. La première affirmation découle immédiatement de (13.3). Pour prouver (b), on observe que le point (a) nous apprend que le polynôme caractéristique de l'adjoint T^* est

$$\chi_{T^*}(t) = \chi_{\bar{A}^\top}(t) = \chi_A(t),$$

c'est donc le polynôme dont les coefficients sont les conjugués complexes du polynôme $\chi_A(t) = \chi_T(t)$. Par conséquent

$$\chi_T(\lambda) = 0 \quad \Leftrightarrow \quad \chi_{T^*}(\bar{\lambda}) = 0.$$

□

Cette proposition justifie la notation suivante pour la matrice conjuguée d'une matrice $A \in M_n(\mathbb{C})$:

$$A^* = \bar{A}^\top,$$

et on dit que A^* est la *matrice adjointe* de A .

13.4 Endomorphismes normaux et le théorème spectral

Définition 13.4.1. Un opérateur T d'un espace hermitien V est dit *normal* s'il commute avec son adjoint : $TT^* = T^*T$.

Proposition 13.4.2. (a) Pour un opérateur T de V , les conditions suivantes sont équivalentes :

- (i) T est normal.
- (ii) $\langle Tx, Ty \rangle = \langle T^*x, T^*y \rangle$ pour tous $x, y \in V$.
- (iii) $\|Tx\| = \|T^*x\|$ pour tout $x \in V$.
- (b) Si T est normal et v est un vecteur propre de T pour la valeur propre $\lambda \in \mathbb{C}$, alors v est aussi un vecteur propre de T^* pour la valeur propre $\bar{\lambda}$.
- (c) Si $v, w \in V$ sont des vecteurs propres d'un opérateur normal T associés à des valeurs propres distinctes $\lambda, \mu \in \mathbb{C}$, alors $v \perp w$.

Preuve. (a) On montre d'abord (i) $\Rightarrow$ (ii). Supposons que T est normal, alors

$$\langle Tx, Ty \rangle = \langle x, T^*Ty \rangle = \langle x, TT^*y \rangle = \langle T^*x, T^*y \rangle.$$

Pour montrer (ii) $\Rightarrow$ (i), on remarque que si $\langle Tx, Ty \rangle = \langle T^*x, T^*y \rangle$ pour tous $x, y \in V$. Alors le calcul ci-dessus montre que $\langle x, T^*Ty \rangle = \langle x, TT^*y \rangle$ pour tous $x, y \in V$. Mais ceci n'est possible que si $TT^* = T^*T$.

L'implication (ii) $\Rightarrow$ (iii) est évidente et l'implication inverse découle de la formule de polarisation (13.1).

(b) Il est facile de vérifier que $(T - \lambda I_V)$ est normal si et seulement si T est normal. En utilisant le point (a) on a donc pour tout vecteur non nul $v \in V$,

$$\begin{aligned} Tv = \lambda v &\Leftrightarrow \|(T - \lambda I_V)v\| = 0 \\ &\Leftrightarrow \|(T - \lambda I_V)^*v\| = 0 \\ &\Leftrightarrow \|(T^* - \bar{\lambda} I_V)v\| = 0 \\ &\Leftrightarrow T^*v = \bar{\lambda}v. \end{aligned}$$

Ce qui prouve que v est un vecteur propre de T pour la valeur propre λ si et seulement si v est aussi un vecteur propre de T^* pour la valeur propre $\bar{\lambda}$.

(c) Supposons que $Tv = \lambda v$ et $Tw = \mu w$ avec $\mu \neq \lambda$. Alors d'après le point précédent on sait que $T^*v = \bar{\lambda}v$. On a donc

$$\mu \langle v | w \rangle = \langle v | Tw \rangle = \langle T^*v | w \rangle = \langle \bar{\lambda}v | w \rangle = \bar{\lambda} \langle v | w \rangle.$$

Ainsi $(\lambda - \bar{\lambda}) \langle v | w \rangle = 0$, et puisque $\mu \neq \lambda$ on conclut que $\langle v | w \rangle = 0$.

□

Théorème 13.4.3 (Théorème spectral). Un opérateur $T : V \rightarrow V$ d'un espace hermitien V est normal si et seulement s'il est orthogonalement diagonalisable, c'est à dire qu'il existe une base unitaire $\{e_1, \dots, e_n\} \subset V$ et $\lambda_1, \dots, \lambda_n \in \mathbb{C}$ tels que $Te_i = \lambda_i e_i$ pour $i = 1, \dots, n$.

Preuve. Supposons que T est orthogonalement diagonalisable, alors on sait par la formule (13.3) que son adjoint est défini par $T^*e_i = \bar{\lambda}_i e_i$, pour $i = 1, \dots, n$. On alors

$$T(T^*e_i) = T(\bar{\lambda}_i e_i) = \bar{\lambda}_i T(e_i) = \bar{\lambda}_i \lambda_i e_i = |\lambda_i|^2 e_i,$$

et de même

$$T^*(Te_i) = T^*(\lambda_i e_i) = \lambda_i T^*(e_i) = \lambda_i \bar{\lambda}_i e_i = |\lambda_i|^2 e_i.$$

On a donc $T(T^*e_i) = T^*(Te_i)$ pour tous les vecteurs de la base $\{e_i\}$ de V , ce qui entraîne que $TT^* = T^*T$, i.e. T est un opérateur normal.

On démontre la réciproque par récurrence sur $n = \dim(V)$, en supposant que $n \geq 2$ car il n'y a rien à prouver si $n = 1$. Tout endomorphisme d'un espace vectoriel de dimension finie possède au moins un vecteur propre, que l'on peut supposer de norme 1. Notons e_1 ce vecteur propre et λ_1 la valeur propre correspondante. On a donc $Te_1 = \lambda_1 e_1$ et on a vu plus haut que $T^*e_1 = \bar{\lambda}_1 e_1$.

Notons $W = e_1^\perp = \{x \in V \mid \langle e_1, x \rangle = 0\}$ l'orthogonal de e_1 . On sait que W est un sous-espace vectoriel de dimension $n - 1$ de V , montrons que W est invariant par T : en effet supposons que $x \in W$, alors

$$\langle e_1, Tx \rangle = \langle T^*e_1, x \rangle = \langle \bar{\lambda}_1 e_1, x \rangle = \lambda_1 \langle e_1, x \rangle = 0,$$

ce qui signifie que $Tx \in W$. On a ainsi montré que $T(W) \subset W$ et on note $T_W = T|_W : W \rightarrow W$ l'opérateur obtenu par restriction de T .

La première affirmation de la Proposition 13.4.2 implique T_W est un opérateur normal de W et par hypothèse de récurrence, il existe donc une base orthonormée $\{e_2, \dots, e_n\}$ de W et $\lambda_2, \dots, \lambda_n \in \mathbb{C}$ tels que $T_W e_j = \lambda_j e_j$ pour $j = 2, \dots, n$.

Il est clair que la famille de vecteurs $\{e_1, e_2, \dots, e_n\}$ est une base unitaire de V et que c'est une base propre pour T .

□

Nous pouvons reformuler le théorème spectral sous la forme importante suivante :

Théorème 13.4.4 (Théorème spectral, variante). *Soit V un espace vectoriel hermitien et T un opérateur linéaire de V . Alors T est normal si et seulement s'il peut être écrit sous la forme suivante :*

$$T = \sum_{j=1}^r \lambda_j P_j, \quad (13.4)$$

où $\sigma(T) = \{\lambda_1, \dots, \lambda_r\} \subset \mathbb{C}$ est le spectre de T et $P : V \rightarrow E_j$ est le projecteur orthogonal sur le sous-espace propre $E_j = \text{Ker}(T - \lambda_j I_V)$.

Notons que V se décompose en somme directe orthogonale :

$$V = E_1 \oplus \dots \oplus E_r, \quad (E_i \perp E_j, \text{ si } i \neq j),$$

et que les projecteurs P_j vérifient les propriétés suivantes :

- (1) $P_j^2 = P_j$.
- (2) $P_i \circ P_j = 0$ si $i \neq j$.
- (3) $\text{Im}(P_j) = E_j$ et $\text{ker}(P_j) \perp E_j$.

$$(4) \quad I_V = \sum_{j=1}^r P_j.$$

On formule ces propriétés en disant que $\{P_k\}$ est un *système complet de projecteurs orthogonaux*.

Lorsque l'opérateur T possède n valeurs propres distinctes, alors la décomposition spectrale prend la forme simplifiée suivante :

$$T(x) = \sum_{j=1}^n \lambda_j \langle e_j, x \rangle e_j. \quad (13.5)$$

13.5 Opérateurs auto-adjoints et unitaires

Définition 13.5.1. Un opérateur $T : V \rightarrow V$ d'un espace hermitien V est dit :

- (1) Autoadjoint (ou hermitien), si $T = T^*$.
- (2) Anti-autoadjoint, si $T = -T^*$.
- (3) Unitaire, si $TT^* = I_V$.

Ces trois types d'opérateurs sont clairement normaux.

Proposition 13.5.2. *Un opérateur T de V est autoadjoint si et seulement $\langle x, Tx \rangle$ est réel pour tout $x \in V$.*

Preuve. Supposons que $T : V \rightarrow V$ est auto-adjoint, alors

$$\langle x, Tx \rangle = \langle T^*x, x \rangle = \langle Tx, x \rangle = \overline{\langle x, Tx \rangle},$$

ce qui implique que $\langle x, Tx \rangle$ est réel. Supposons inversement que $\langle x, Tx \rangle$ est réel pour tout vecteur $x \in V$, alors

$$\langle x, Tx \rangle = \overline{\langle x, Tx \rangle} = \langle Tx, x \rangle = \langle x, T^*x \rangle,$$

et le lemme 13.3.1 entraîne alors que $T = T^*$. □

Le théorème spectral permet de prouver facilement les caractérisation suivantes :

Corollaire 13.5.3. (a) *Un opérateur de V est autoadjoint si et seulement s'il est normal et toutes ses valeurs propres sont réelles.*

(b) *Un opérateur de V est anti-autoadjoints si et seulement s'il est normal et toutes ses valeurs propres sont imaginaires.*

(c) *Un opérateur de V est unitaire si et seulement s'il est normal et toutes ses valeurs propres sont des nombres complexes de modules 1.*

Remarque. On peut utiliser la Proposition 13.4.2 pour montrer (a) qu'un opérateur auto-adjoint a des valeurs propres réelles. On a en effet pour tout vecteur propre v (de valeur propre λ

$$\lambda v = Tv = T^*v = \bar{\lambda}v.$$

donc $\lambda = \bar{\lambda}$. Même argument pour (b) et (c).

Proposition 13.5.4. (a) *L'ensemble des opérateurs autoadjoints de V forme un sous-espace vectoriel réel de $\mathcal{L}(V)$.*

(b) *L'ensemble des opérateurs unitaires de V forme un sous-groupe de $\mathrm{GL}(V)$, que l'on note $\mathrm{U}(V)$ et qui s'appelle le groupe unitaire de V .*

Le groupe unitaire de $\mathbb{C}^n$, pour le produit scalaire hermitien standard est noté $\mathrm{U}(n)$.

Noter par contraste que l'ensemble des opérateurs normaux ne forme pas un sous-espace vectoriel.

Si S et T sont normaux alors $S + T$ n'est en général pas un opérateur normal.

13.6 Espaces de Hilbert et opérateurs auto-adjoints

Cette section est facultative

Définition 13.6.1. On appelle *espace de Hilbert* un espace vectoriel $\mathcal{H}$ sur le corps $\mathbb{R}$ ou $\mathbb{C}$ muni d'un produit scalaire $\langle \cdot, \cdot \rangle$ (qui est un produit scalaire hermitien dans le cas complexe) et qui est complet pour la norme $\| \cdot \|$ associée au produit scalaire. Cette condition signifie que toute suite de Cauchy de $\mathcal{H}$ doit converger.

Exemples 1. Tout espace vectoriel de dimension finie sur $\mathbb{R}$ ou $\mathbb{C}$ muni d'un produit scalaire (produit scalaire hermitien dans le cas complexe) est un espace de Hilbert car le théorème de Bolzano-Weierstrass entraîne que toute suite de Cauchy de $\mathbb{R}^n$ ou $\mathbb{C}^n$ converge. La notion d'espace de Hilbert généralise donc à la fois les espaces euclidiens et les espaces hermitiens et autorise la dimension infinie.

2. L'espace $\ell_2(\mathbb{N})$ est un espace de Hilbert (c'est le prototype d'espace de Hilbert en dimension infinie).

3. L'espace $C^0([a, b], \mathbb{C})$ des fonctions continues sur un intervalle $[a, b]$ n'est pas un espace de Hilbert pour le produit L^2 (car la limite d'une suite de fonctions continues n'est pas toujours continue si la convergence n'est pas uniforme, ce qui implique que l'espace $C^0([a, b], \mathbb{C})$ n'est pas complet pour la norme associée au produit scalaire L^2).

Définition 1.) Soit $\mathcal{H}$ un espace de Hilbert. Un *opérateur* de $\mathcal{H}$ est une application $\mathbb{R}$ -linéaire ou $\mathbb{C}$ -linéaire de $\mathcal{H}$ dans lui-même qui est continue pour la norme associée.

2.) L'opérateur $T : \mathcal{H} \rightarrow \mathcal{H}$ est dit *auto-adjoint* si la condition suivante est vérifiée :

$$\langle Tx, y \rangle = \langle x, Ty \rangle,$$

pour tous $x, y \in \mathcal{H}$.

Exemple. On peut associer à toute fonction continue $\varphi : [a, b] \times [a, b] \rightarrow \mathbb{C}$ un opérateur T sur $C^0([a, b], \mathbb{C})$ par la formule

$$(Tf)(x) = \int_a^b f(y)\varphi(x, y)dy.$$

Cet opérateur est auto-adjoint si et seulement si $\varphi(y, x) = \overline{\varphi(x, y)}$ pour tous $x, y \in [a, b]$.

Définition. L'opérateur T sur l'espace de Hilbert $\mathcal{H}$ admet une *décomposition spectrale finie* si on peut écrire

$$T = \lambda_1 P_1 + \cdots + \lambda_r P_r$$

où $P_1, \dots, P_r$ sont des projecteurs de $\mathcal{H}$ qui sont deux-à-deux orthogonaux.

Le théorème spectral nous dit que tout espace de Hilbert de dimension finie admet une décomposition spectrale. Un opérateur autoadjoint d'un espace de Hilbert de dimension infinie admet aussi une décomposition spectrale, qui peut être infinie. Cela donne une écriture du type

$$T = \sum_{i=1}^{\infty} \lambda_i P_i,$$

lorsqu'il y a un nombre dénombrable de valeurs propres. Dans le cas le plus général, le spectre n'est pas forcément un sous-ensemble discret de $\mathbb{R}$ et la décomposition spectrale s'écrit sous la forme d'une intégrale que l'on étend au spectre continu de T :

$$T = \int_{\sigma(T)} \lambda P_\lambda d\lambda$$

Un chapitre important de l'*analyse fonctionnelle* est de donner un sens précis à ce genre de formules et de les démontrer rigoureusement.

13.7 Applications en mécanique quantique.

Cette section est facultative

En mécanique Newtonienne, l'état d'un système évolue selon une loi déterministe et les quantités observables sont des fonctions des variables d'état que l'on peut (en principe) connaître exactement à chaque instant.

L'exemple le plus simple est la mécanique classique du point matériel. L'état du système au temps t est déterminé par la position $x(t) \in \mathbb{R}^3$ et le moment $p = m\dot{x} \in \mathbb{R}^3$ de la particule. La loi d'évolution est prescrite par l'équation de Newton : $\frac{dp}{dt} = F$.

La formalisation mathématique de la mécanique quantique repose sur un certain nombre de postulats que nous décrivons brièvement ci-dessous sans chercher à être ni rigoureux ni exhaustif :

Premier postulat : L'état d'un système quantique (par exemple une particule) au temps t est représenté par un élément non nul $\psi = \psi(t)$ d'un espace de Hilbert $\mathcal{H}$, appelé *vecteur d'état* du système (typiquement une *fonction d'onde*).

Deux vecteurs $\psi_1, \psi_2 \in \mathcal{H}$ représentent le même état si l'un est multiple de l'autre. On supposera donc souvent que le vecteur d'état est normalisé $\|\psi\| = 1$.

Deuxième postulat : On associe à chaque observable du système un opérateurs autoadjoints T de l'espace de Hilbert $\mathcal{H}$.

Par le théorème spectral, Il existe alors une décomposition de $\mathcal{H}$ en somme directe orthogonale de sous-espaces propres pour T . Les vecteurs propres de T s'appellent des *états propres* pour l'opérateur T , et tout état ψ est une "superposition" d'états propres (le mot "superposition" veut dire ici "combinaison linéaire", la décomposition de ψ en combinaison linéaire de vecteurs propres de T s'appelle en mécanique quantique une *superposition d'états propres*).

Troisième postulat : Lors d'une expérience, on ne peut observer que les états propres. En conséquence, ce qui est mesuré est une valeur propre de l'opérateur T . Puisque cet opérateur est autoadjoint, ses valeurs propres sont des nombres réels.

Si on a la décomposition spectrale $T = \sum_i \lambda_j P_j$, alors toute expérimentation visant à mesurer l'observable T a pour effet de projeter un état ψ (c'est-à-dire un vecteur $\psi \in \mathcal{H}$) sur un sous-espace propre de T :

$\psi \mapsto P_j(\psi)$, de plus, la mesure obtenue est la valeur propre λ_i , avec une probabilité

$$p_j = \left| \frac{\langle \psi | P_j(\psi) \rangle}{\|\psi\| \|P_j(\psi)\|} \right|^2.$$

Remarques.

- (a) Si $\psi, \phi \in \mathcal{H}$ sont deux vecteurs non nuls de $\mathcal{H}$ qui représentent deux états possibles d'un système quantique, alors on dit que le nombre complexe

$$\frac{\langle \psi | \phi \rangle}{\|\psi\| \|\phi\|} \in \mathbb{C}$$

représente l'*amplitude de probabilité* que le système préparé dans l'état ψ soit observé dans l'état ϕ . La probabilité elle-même est le carré du module de cette amplitude.

L'importance de cette notion vient du fait qu'en mécanique quantique, les calculs de transitions et de comportement des systèmes se font en manipulant algébriquement les amplitudes de probabilité et non directement les probabilités.

- (b) Si l'opérateur auto-adjoint T possède $n = \dim(\mathcal{H})$ valeurs propres distinctes, alors la décomposition spectrale prend la forme simplifiée dans une base unitaire (voir (13.5)) :

$$T(\psi) = \sum_{j=1}^n \lambda_j \langle e_j, \psi \rangle e_j.$$

Les amplitudes de probabilités sont alors données par $\langle e_j, \psi \rangle$ et les probabilités de transition par $p_j = |\langle e_j, \psi \rangle|^2$ (on suppose $\|\psi\| = 1$), cette quantité représente la probabilité que le système dans l'état ψ soit observé dans l'état propre e_j .

- (c) Si l'espace des états est un espace de Hilbert de dimension infinie, alors le spectre de l'opérateur T (l'ensemble des valeurs propres) peut-être continu ou discret. Lorsque le spectre est discret, cela implique que la variable observée ne peut prendre que des valeurs discrètes (principe de quantification).

Quatrième postulat : Le dernier postulat nous dit que loi d'évolution du système est prescrite par l'*équation de Schrödinger* :

$$i\hbar \frac{d}{dt} \psi = H(\psi)$$

où $i = \sqrt{-1}$, $\hbar = \frac{ih}{2\pi}$ est la constante de Plank réduite et H est un opérateur autoadjoint de l'espace de Hilbert $\mathcal{H}$ qu'on appelle le *Hamiltonien du système*.

Remarques.

- (1) Ces postulats n'ont clairement rien d'intuitifs, ils ont été développés dans les années 1920-1940 par les Pères fondateurs de la mécanique quantique (Bohr, Heisenberg, Dirac...). L'histoire de la mécanique quantique est l'un des chapitres les plus complexes et passionnants de l'histoire des science.
- (2) Suivant les auteurs, l'ordre des postulats, leur nombre et leurs formulations exactes peuvent présenter des variations assez importantes.

- (3) En principe l'espace de Hilbert est de dimension infinie, mais des modèles simplifiés peuvent être développés sur la base d'un espace de Hilbert de dimension finie (en négligeant donc une partie de l'information). Cette approche est semblable à celle qui consiste à réduire le nombre de dimensions (par exemple de 3 à 2 ou à 1) dans un problème de mécanique classique. Ces modèles de dimension finie ont l'avantage de ramener la mécanique quantique à de l'algèbre linéaire. Ils sont très efficaces en chimie pour modéliser les orbitales des électrons dans un modèle quantique d'atome.
- (4) Si l'espace de Hilbert est de dimension finie, l'équation de Schrödinger peut en principe se résoudre directement en calculant l'exponentielle de la matrice du Hamiltonien H :

$$i\hbar \frac{d}{dt}\psi = H(\psi) \quad \Rightarrow \quad \psi(t) = e^{-\frac{i}{\hbar}tH} \cdot \psi_0$$

- (5) Le hamiltonien est lui-même un opérateur autoadjoint. Selon le deuxième postulat il lui correspond donc une quantité observable. La théorie montre que cette observable est l'*énergie totale* du système (qui est donc constante au cours du temps, conformément au principe de conservation de l'énergie).

Sur la fonction d'onde

Dans le modèle de Schrödinger, l'état de la particule-onde est représenté à chaque instant par une *fonction d'onde* ψ_t , qui est une fonction

$$\psi_t : \mathbb{R}^3 \rightarrow \mathbb{C} \quad \text{telle que} \quad \int_{\mathbb{R}^3} |\psi_t|^2 = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} |\psi_t(x, y, z)|^2 dx dy dz < \infty.$$

On considère que deux fonctions d'onde ψ et ϕ représentent le même état si $\phi = \alpha\psi$ où α est un nombre complexe non nul. Cela nous permet de normaliser la fonction d'onde, i.e. de supposer que

$$\int_{\mathbb{R}^3} |\psi_t|^2 = 1.$$

On a alors l'interprétation suivante : *la probabilité de présence de la particule dans une région $\Omega \subset \mathbb{R}^3$ est donnée par l'intégrale*

$$\text{Prob}(\psi_t \mid \Omega) = \int_{\Omega} |\psi_t|^2.$$

Ainsi la quantité $|\psi_t(x)|^2$ représente une *densité de probabilité de présence de la particule au point x* . Par comparaison l'argument complexe de ψ représente une simple phase et n'a pas d'interprétation physique (car les fonctions d'onde ψ et $e^{i\theta}\psi$ représentent le même état de la particule).

L'ensemble des fonctions intégrables⁴ $\phi : \mathbb{R}^3 \rightarrow \mathbb{C}$ vérifiant

$$\int_{\mathbb{R}^3} |\phi|^2 < \infty$$

4. intégrable au sens de Lebesgue... mais ne nous inquiétons pas de cela.

forme un espace de Hilbert, que l'on note $L^2(\mathbb{R}^3, \mathbb{C})$. Le produit scalaire hermitien sur cet espace est défini par

$$\langle \phi, \psi \rangle = \int_{\mathbb{R}^3} \overline{\phi(x)} \psi(x) dx.$$

A priori une fonction $\phi \in L^2(\mathbb{R}^3, \mathbb{C})$ n'est pas forcément continue, elle n'a a fortiori pas de dérivée partielle. Mais on démontre en analyse fonctionnelle que l'espace $L^2(\mathbb{R}^3, \mathbb{C})$ contient un sous-espace vectoriel $\mathcal{H} \subset L^2(\mathbb{R}^3, \mathbb{C})$ qui est dense (c'est-à-dire que tout élément de $L^2(\mathbb{R}^3, \mathbb{C})$ peut s'approximer par une suite d'éléments de $\mathcal{H}$, de même que tout nombre réel peut s'approximer par une suite de nombres rationnels) et qui a la propriété suivante :

$$\psi \in \mathcal{H} \quad \Rightarrow \quad \Delta \psi \in \mathcal{H},$$

où Δ est l'opérateur de Laplace défini par

$$\Delta \psi = \frac{\partial^2 \psi}{\partial x^2} + \frac{\partial^2 \psi}{\partial y^2} + \frac{\partial^2 \psi}{\partial z^2}.$$

Le hamiltonien est alors donné par

$$H(\psi) = \left(V - \frac{\hbar^2}{2m} \Delta \right) \cdot \psi,$$

où m est la masse de la particule et $V : \mathbb{R}^3 \rightarrow \mathbb{R}$ est une fonction qui représente l'énergie potentielle. Ce hamiltonien représente la version quantique de la quantité "énergie totale = énergie cinétique + énergie potentielle".

Exercice. Prouver que H est autoadjoint, i.e.

$$\int_{\mathbb{R}^3} \overline{\phi} H(\psi) dx = \int_{\mathbb{R}^3} \overline{H(\phi)} \psi dx$$

(on peut supposer pour cet exercice que ϕ et ψ sont des fonctions de classe C^2 qui tendent vers 0, ainsi que leur dérivées, lorsque x tend vers l'infini).

En conclusion, l'état de la particule est représenté par sa fonction d'onde $\psi_t \in \mathcal{H}$ qui est normalisée, i.e. $\int_{\mathbb{R}^3} |\phi|^2 = 1$, et qui vérifie l'équation de Schrödinger :

$$i\hbar \frac{\partial \psi}{\partial t} = \left(V - \frac{\hbar^2}{2m} \Delta \right) \psi.$$

Index

- Adjoint d'un opérateur, 101
- Amplitude de probabilité, 108
- Base unitaire, 99
- Cône de lumière, 91
- Causalité, 91
- Congruence, 51
- Contragrédiente, 44
- Couplage, 47
- Couplage non dégénéré, 48
- Dirac (notation de), 49
- Dual d'un espace vectoriel, 42
- Duale (application), 45
- Duale (base), 43
- Espace pseudo-euclidien, 85
- Espace de Lorentz-Minkowski, 90
- Espace de Minkowski, 90
- Espace hermitien, 97
- Espace-temps, 90
- Forme bilinéaire, 49
- Forme canonique de Jordan, 29
- Forme hermitienne, 96
- Forme quadratique, 53
- Formes bilinéaires symétriques/asymétriques.,
52
- Gram (matrice de), 49
- Groupe unitaire, 105
- Inégalité de Cauchy-Schwarz, 58, 98, 99
- Inégalité du triangle, 60
- Inégalité du triangle inversée (inégalité causale),
94
- Isométrie, 66, 85
- Matrice adjointe, 102
- Multiplicité généralisée, 14
- Polarisation (formules de), 54
- Polynôme annulateur, 7
- polynôme spectral, 12
- Produit scalaire, 57
- Produit scalaire hermitien, 97
- Produit tensoriel, 51
- Sous espace cyclique, 23
- Sous-espace caractéristique, 17
- Sous-espace propre généralisé, 17
- Symétrie orthogonale, 64
- Système complet de projecteurs orthogonaux,
105
- Temps propre, 91
- Théorème de Pythagore, 60
- Théorème spectral, 104
- Vecteur propre généralisé, 14