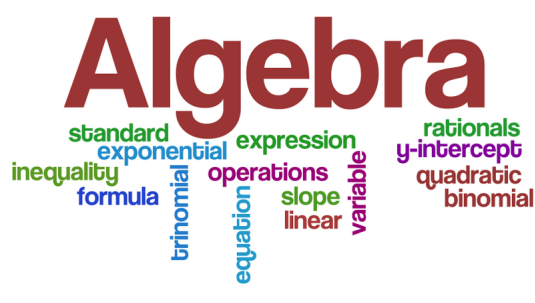




# ALGÈBRE BILINÉAIRE

Licence L2  
Emmanuel Hebey  
Année 2024-2025



# TABLE DES MATIÈRES

|                                                    |           |
|----------------------------------------------------|-----------|
| 1. INTRODUCTION                                    | p. 04     |
| <br>CHAPITRE 1. FORMES BILINÉAIRES ET QUADRATIQUES | <br>p. 05 |
| 2. PREMIÈRES DÉFINITIONS                           | p. 05     |
| 3. SYMÉTRIE ET ANTISYMÉTRIE                        | p. 09     |
| 4. CHANGEMENT DE BASES                             | p. 10     |
| 5. DUALITÉ                                         | p. 12     |
| 6. RANG D'UNE FORME BILINÉAIRE                     | p. 13     |
| 7. FORMES BILINÉAIRES ET FORMES QUADRATIQUES       | p. 14     |
| 8. ORTHOGONALITÉ                                   | p. 17     |
| 9. LE THÉORÈME DE SYLVESTER                        | p. 19     |
| <br>CHAPITRE 2. ESPACES PRÉHILBERTIENS             | <br>p. 25 |
| 10. NORMES ET PRODUITS SCALAIRES                   | p. 25     |
| 11. ORTHOGONALITÉ                                  | p. 29     |
| 12. RETOUR SUR LA SIGNATURE                        | p. 30     |
| 13. LE CRITÈRE DE SYLVESTER                        | p. 31     |
| 14. BASES ORTHONORMÉES                             | p. 33     |
| 15. ORTHOGONALITÉ ET SOUS ESPACES VECTORIELS       | p. 38     |
| 16. PROJECTEURS ORTHOGONAUX                        | p. 39     |
| 17. SYMÉTRIES ORTHOGONALES                         | p. 41     |
| <br>CHAPITRE 3. LE THÉORÈME SPECTRAL               | <br>p. 43 |
| 18. ENDOMORPHISME ADJOINT                          | p. 43     |
| 19. LE THÉORÈME SPECTRAL                           | p. 45     |
| 20. PREUVE LORSQUE $n = 3$                         | p. 47     |
| 21. LE POINT DE VUE DES MATRICES                   | p. 48     |
| 22. DÉCOMPOSITIONS D'IWASAWA ET DE CHOLESKY        | p. 52     |
| 23. DÉCOMPOSITION SVD                              | p. 55     |

|                                                                 |           |
|-----------------------------------------------------------------|-----------|
| 24. PREUVE DU THÉORÈME SPECTRAL - CAS GÉNÉRAL                   | p. 56     |
| 25. INÉGALITÉ DE KANTOROVITCH                                   | p. 61     |
| 26. CARACTÉRISATION DES VALEURS PROPRES                         | p. 63     |
| 27. CARACTÉRISATION DE COURANT-FISCHER                          | p. 64     |
| 28. PRINCIPE DU MAXIMUM DE KY FAN                               | p. 65     |
| 29. TRACE D'UN PRODUIT DE MATRICES                              | p. 67     |
| 30. MAJORATION DE SCHUR                                         | p. 67     |
| 31. INÉGALITÉS DE KY FAN ET VON NEUMANN                         | p. 68     |
| 32. LE LEMME D'ENTRELAÇEMENT DE CAUCHY                          | p. 70     |
| 33. ENTRELAÇEMENTS D'HERMITE ET DE CAUCHY                       | p. 71     |
| <br>CHAPITRE 4. ESPACES HERMITIENS                              | <br>p. 73 |
| 34. PREMIÈRES DÉFINITIONS                                       | p. 73     |
| 35. BASES ORTHONORMÉES                                          | p. 74     |
| 36. RÉDUCTION DES ENDOMORPHISMES NORMAUX                        | p. 75     |
| <br>CHAPITRE 5. LA DIMENSION INFINIE                            | <br>p. 78 |
| 37. COMPACTITÉ ET DIMENSION                                     | p. 80     |
| 38. LE SECOND THÉORÈME DE RIESZ                                 | p. 81     |
| 39. ESPACES DE BANACH ET DE HILBERT                             | p. 82     |
| 40. APPLICATIONS LINÉAIRES CONTINUES                            | p. 84     |
| 41. PROJECTIONS ORTHOGONALES                                    | p. 87     |
| 42. LE PREMIER THÉORÈME DE RIESZ                                | p. 88     |
| 41. APPLICATIONS MULTILINÉAIRES CONTINUES                       | p. 89     |
| 42. L'ISOMÉTRIE NATURELLE $L_c(E, L_c(F, G)) \sim L_c(E, F; G)$ | p. 92     |
| 43. PRINCIPE DE CONTRACTION DE BANACH-PICARD                    | p. 93     |

# ALGÈBRE BILINÉAIRE

EMMANUEL HEBEY

## 1. INTRODUCTION

Ce cours fait suite au cours d'algèbre linéaire 3 du premier semestre. Il s'agira ici de développer les bases de la théorie bilinéaire. Deux théorèmes marquants sont au programme de ce cours: le théorème de Sylvester qui classe les formes quadratiques, et le théorème spectral qui revient sur la problématique de la diagonalisation des endomorphismes mais avec une approche nouvelle. Au passage nous aborderons la théorie préhilbertienne, première étape avant le développement de la théorie hilbertienne et de l'analyse sur les espaces vectoriels normés.

Cinq chapitres forment ce cours. La théorie des formes bilinéaires et quadratiques, qui reprend l'analogie réelle  $(x, y) \rightarrow xy$  et  $x \rightarrow x^2$ , est développée au Chapitre 1. La théorie préhilbertienne des espaces vectoriels munis d'un produit scalaire est développée au Chapitre 2. Le théorème spectral et la théorie des endomorphismes symétriques sont traités au Chapitre 3. Les espaces hermitiens et la réduction des endomorphismes normaux sont traités au Chapitre 4. La théorie infinie est traitée au Chapitre 5.

# CHAPITRE 1

## FORMES BILINÉAIRES ET QUADRATIQUES

On commence avec la théorie générale des formes bilinéaires et quadratiques avec, comme objectif dans ce chapitre, de présenter le théorème de Sylvester.

### 2. PREMIÈRES DÉFINITIONS

On commence avec la définition d'une forme bilinéaire.

**Définition 2.1.** Soit  $E$  un  $\mathbb{R}$ -espace vectoriel. Par définition, une forme bilinéaire sur  $E$  est une application  $B : E \times E \rightarrow \mathbb{R}$  qui vérifie que pour tout  $x_0 \in E$  et tout  $y_0 \in E$ , les applications

$$x \rightarrow B(x, y_0) \quad \text{et} \quad y \rightarrow B(x_0, y)$$

sont linéaires de  $E$  dans  $\mathbb{R}$ .

En d'autres termes, une fonction  $B : E \times E \rightarrow \mathbb{R}$  est une forme bilinéaire si la fonction est linéaire par rapport à chacune de ses variables. La plus simple des formes bilinéaires est la fonction  $B : \mathbb{R} \times \mathbb{R} \rightarrow \mathbb{R}$  donnée par  $B(x, y) = xy$ . Attention, une forme bilinéaire (non identiquement nulle) n'est jamais linéaire par rapport au couple  $(x, y)$ . Par exemple, pour  $B : \mathbb{R} \times \mathbb{R} \rightarrow \mathbb{R}$  donnée par  $B(x, y) = xy$ , on a que

$$B(x + x', y + y') \neq B(x, y) + B(x', y')$$

dès que  $x'y + xy' \neq 0$  car

$$B(x + x', y + y') - B(x, y) - B(x', y') = x'y + xy'.$$

En dimension finie, ce qui sera notre cas, les formes bilinéaires sont par contre, comme les applications linéaires, représentables par des matrices. Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie  $n$ , et soit  $\mathcal{B} = (e_1, \dots, e_n)$  une base de  $E$ . Soit de plus  $B : E \times E \rightarrow \mathbb{R}$  une forme bilinéaire sur  $E$ , et soit  $B_{ij}$  les réels

$$B_{ij} = B(e_i, e_j)$$

où  $i, j = 1, \dots, n$ . Si  $x$  et  $y$  sont deux vecteurs de  $E$  de coordonnées respectives  $(x_1, \dots, x_n)$  et  $(y_1, \dots, y_n)$  dans  $\mathcal{B}$ , donc  $x = \sum_{i=1}^n x_i e_i$  et  $y = \sum_{j=1}^n y_j e_j$ , alors

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

Pour démontrer cette relation on écrit que

$$\begin{aligned} B(x, y) &= B\left(\sum_{i=1}^n x_i e_i, \sum_{j=1}^n y_j e_j\right) \\ &= \sum_{i=1}^n x_i B\left(e_i, \sum_{j=1}^n y_j e_j\right) \\ &= \sum_{i=1}^n \sum_{j=1}^n x_i y_j B(e_i, e_j) \end{aligned}$$

puisque pour tout  $x_0 \in E$  et tout  $y_0 \in E$ , les applications

$$x \rightarrow B(x, y_0) \quad \text{et} \quad y \rightarrow B(x_0, y)$$

sont linéaires sur  $E$ . D'un point de vue matriciel cela donne lieu à la définition suivante.

**Définition 2.2.** Soient  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie  $n$ , et  $\mathcal{B} = (e_1, \dots, e_n)$  une base de  $E$ . Soit de plus  $B : E \times E \rightarrow \mathbb{R}$  une forme bilinéaire sur  $E$ . Par définition, la matrice de  $B$  dans la base  $\mathcal{B}$ , notée  $M_{\mathcal{B}}(B)$ , est la matrice  $M_{\mathcal{B}}(B) = (B_{ij})_{i,j}$  carrée d'ordre  $n$  composée des

$$B_{ij} = B(e_i, e_j) .$$

Elle caractérise  $B$  en ce sens que pour tous  $x$  et  $y$  dans  $E$ ,

$$B(x, y) = {}^tV_{\mathcal{B}}(x)M_{\mathcal{B}}(B)V_{\mathcal{B}}(y) ,$$

où  $V_{\mathcal{B}}(x)$  et  $V_{\mathcal{B}}(y)$  représentent les matrices colonnes composées respectivement des coordonnées de  $x$  et  $y$  dans  $\mathcal{B}$ , et où  ${}^tV_{\mathcal{B}}(x)$  est la matrice ligne transposée de  $V_{\mathcal{B}}(x)$ .

D'un point de vue matriciel on a écrit que

$$B(x, y) = (x_1 \quad \dots \quad x_n) \begin{pmatrix} B_{11} & \dots & B_{1n} \\ \vdots & & \vdots \\ B_{n1} & \dots & B_{nn} \end{pmatrix} \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix}$$

et, comme on s'en convaincra facilement, on retrouve la relation déjà vue

$$B(x, y) = \sum_{i,j=1}^n B_{ij}x_iy_j .$$

**Exercice:** Soit  $\mathbb{R}_2[X]$  l'espace des polynômes réels d'ordre 2 et soit  $\mathcal{B} = (P_1, P_2, P_3)$  la base canonique de  $\mathbb{R}_2[X]$  constituée des polynômes  $P_1(X) = 1$ ,  $P_2(X) = X$  et  $P_3(X) = X^2$ . On considère  $B : \mathbb{R}_2[X] \times \mathbb{R}_2[X] \rightarrow \mathbb{R}$  donnée par

$$B(P, Q) = \int_0^1 P(t)Q(t)dt .$$

Montrer que  $B$  est une forme bilinéaire sur  $\mathbb{R}_2[X]$ . Donner sa matrice dans la base canonique  $\mathcal{B} = (P_1, P_2, P_3)$  de  $\mathbb{R}_2[X]$ .

**Solution:** Il est clair que  $\forall P, Q, R \in \mathbb{R}_2[X]$  et  $\forall \lambda \in \mathbb{R}$ ,

$$\begin{aligned} \int_0^1 (P(t) + \lambda Q(t))R(t)dt &= \int_0^1 P(t)R(t)dt + \lambda \int_0^1 Q(t)R(t)dt , \text{ et} \\ \int_0^1 P(t)(Q(t) + \lambda R(t))dt &= \int_0^1 P(t)Q(t)dt + \lambda \int_0^1 P(t)R(t)dt . \end{aligned}$$

Par suite,

$$\begin{aligned} B(P + \lambda Q, R) &= B(P, R) + \lambda B(Q, R) \text{ et} \\ B(P, Q + \lambda R) &= B(P, Q) + \lambda B(P, R) \end{aligned}$$

et donc  $B$  est une forme bilinéaire sur  $\mathbb{R}_2[X]$ . La forme est aussi symétrique, i.e.  $B(P, Q) = B(Q, P)$  pour tous  $P$  et  $Q$ , et

$$\begin{aligned} B(P_1, P_1) &= \int_0^1 dt = 1 , \quad B(P_1, P_2) = \int_0^1 t dt = \frac{1}{2} \\ B(P_1, P_3) &= \int_0^1 t^2 dt = \frac{1}{3} , \quad B(P_2, P_1) = B(P_1, P_2) = \frac{1}{2} \end{aligned}$$

De même,

$$\begin{aligned} B(P_2, P_2) &= \int_0^1 t^2 dt = \frac{1}{3}, \quad B(P_2, P_3) = \int_0^1 t^3 dt = \frac{1}{4} \\ B(P_3, P_1) &= B(P_1, P_3) = \frac{1}{3}, \quad B(P_3, P_2) = B(P_2, P_3) = \frac{1}{4} \\ B(P_3, P_3) &= \int_0^1 t^4 dt = \frac{1}{5}. \end{aligned}$$

Par suite,

$$M_{\mathcal{B}}(B) = \begin{pmatrix} 1 & \frac{1}{2} & \frac{1}{3} \\ \frac{1}{2} & \frac{1}{3} & \frac{1}{4} \\ \frac{1}{3} & \frac{1}{4} & \frac{1}{5} \end{pmatrix}.$$

□

**Exercice:** Soit  $\mathcal{M}_2(\mathbb{R})$  l'espace des matrices carrées réelles d'ordre 2. On note  $\mathcal{B} = (A_1, A_2, A_3, A_4)$  la base canonique de  $\mathcal{M}_2(\mathbb{R})$  donnée par

$$A_1 = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}, \quad A_2 = \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}, \quad A_3 = \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix}, \quad A_4 = \begin{pmatrix} 0 & 0 \\ 0 & 1 \end{pmatrix}.$$

Montrer que l'application  $\Phi : \mathcal{M}_2(\mathbb{R}) \times \mathcal{M}_2(\mathbb{R}) \rightarrow \mathbb{R}$  donnée par

$$\Phi(A, B) = \text{trace}({}^tAB)$$

est une forme bilinéaire sur  $\mathcal{M}_2(\mathbb{R})$ . Donner sa matrice dans la base canonique  $(A_1, A_2, A_3, A_4)$  de  $\mathcal{M}_2(\mathbb{R})$ .

**Solution:** Il est clair que  $\forall A, B, C \in \mathcal{M}_2(\mathbb{R})$  et  $\forall \lambda \in \mathbb{R}$ ,

$$\begin{aligned} \text{trace}({}^t(A + \lambda B), C) &= \text{trace}({}^tAC) + \lambda \text{trace}({}^tBC), \quad \text{et} \\ \text{trace}({}^tA(B + \lambda C)) &= \text{trace}({}^tAB) + \lambda \text{trace}({}^tAC). \end{aligned}$$

Donc  $\Phi$  est bien bilinéaire. On calcule

$$\Phi(A_i, A_j) = \delta_{ij}$$

où les  $\delta_{ij}$  sont les symboles de Kroenecker (1 si  $i = j$ , 0 sinon). Par exemple,

$$\begin{aligned} \Phi(A_2, A_3) &= \text{trace}({}^tA_2A_3) \\ &= \text{trace}^t \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix} \\ &= \text{trace} \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} 0 & 0 \\ 1 & 0 \end{pmatrix} \\ &= \text{trace} \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix} \\ &= 0. \end{aligned}$$

Ou encore,

$$\begin{aligned}
 \Phi(A_1, A_2) &= \text{trace}({}^t A_1 A_2) \\
 &= \text{trace}({}^t \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}) \\
 &= \text{trace} \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix} \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} \\
 &= \text{trace} \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix} \\
 &= 0 .
 \end{aligned}$$

Par suite

$$M_{\mathcal{B}}(\Phi) = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} .$$

□

Si une forme bilinéaire sur  $E$  s'écrit  $B(x, y) = \sum_{i,j=1}^n a_{ij} x_i y_j$  pour tous  $x, y \in E$ , où les  $x_i$  sont les coordonnées de  $x$  dans une base  $\mathcal{B}$ , et les  $y_i$  sont les coordonnées de  $y$  dans  $\mathcal{B}$ , alors  $a_{ij} = B(e_i, e_j)$  pour tous  $i, j = 1, \dots, n$ . La raison en est le résultat suivant.

**Lemme 2.1.** Soient  $E$  un  $\mathbb{R}$ -espace vectoriel,  $\mathcal{B} = (e_1, \dots, e_n)$  une base de  $E$  et  $(a_{ij})$ ,  $(b_{ij})$  deux matrices  $n \times n$ . On suppose que pour tous  $x, y \in E$ ,

$$\sum_{i,j=1}^n a_{ij} x_i y_j = \sum_{i,j=1}^n b_{ij} x_i y_j ,$$

où les  $x_i$  sont les coordonnées de  $x$  dans  $\mathcal{B}$  et les  $y_i$  sont les coordonnées de  $y$  dans  $\mathcal{B}$ . Alors  $a_{ij} = b_{ij}$  pour tous  $i, j = 1, \dots, n$ .

*Démonstration.* Soient  $i_0$  et  $j_0$  fixés. On prend  $x = (0, \dots, 0, 1, 0, \dots, 0)$  (à la  $i_0$ ème place) et  $y = (0, \dots, 0, 1, 0, \dots, 0)$  (à la  $j_0$ ème place). Alors

$$\sum_{i,j=1}^n a_{ij} x_i y_j = a_{i_0 j_0} \text{ et } \sum_{i,j=1}^n b_{ij} x_i y_j = b_{i_0 j_0} .$$

Donc  $a_{i_0 j_0} = b_{i_0 j_0}$ , et comme  $i_0, j_0$  sont quelconques, on récupère que  $a_{ij} = b_{ij}$  pour tous  $i, j = 1, \dots, n$ . □

Le résultat suivant a lieu.

**Théorème 2.1.** Soient  $E$  un  $\mathbb{R}$ -espace vectoriel,  $\mathcal{B} = (e_1, \dots, e_n)$  une base de  $E$  et  $(a_{ij})$ , une matrice  $n \times n$ . On définit

$$B(x, y) = \sum_{i,j=1}^n a_{ij} x_i y_j$$

pour tous  $x, y \in E$ . Alors  $B$  est bilinéaire sur  $E \times E$  et la matrice de représentation de  $B$  dans  $\mathcal{B}$  est la matrice des  $a_{ij}$ .

*Démonstration.* Soit  $x \in E$  quelconque fixé. Pour tous  $\lambda, \mu \in \mathbb{R}$ , et pour tous  $y, y' \in E$ ,

$$\begin{aligned} B(x, \lambda y + \mu y') &= \sum_{i,j=1}^n a_{ij} x_i (\lambda y_j + \mu y'_j) \\ &= \lambda \sum_{i,j=1}^n a_{ij} x_i y_j + \mu \sum_{i,j=1}^n a_{ij} x_i y'_j \\ &= \lambda B(x, y) + \mu B(x, y') . \end{aligned}$$

On obtient donc la linéarité par rapport à la seconde variable. On démontre de même la linéarité par rapport à la première variable.  $\square$

### 3. SYMÉTRIE ET ANTISYMMÉTRIE

On aborde la symétrie et l'antisymétrie des formes bilinéaires et on montre que toute forme bilinéaire peut s'écrire comme la somme d'une forme bilinéaire symétrique et d'une forme bilinéaire antisymétrique.

**Définition 3.1.** Soit  $E$  un  $\mathbb{R}$ -espace vectoriel et  $B : E \times E \rightarrow \mathbb{R}$  une forme bilinéaire sur  $E$ . La forme bilinéaire  $B$  est dite *symétrique* si  $B(y, x) = B(x, y)$  pour tous  $x, y \in E$ . Elle est dite *antisymétrique* si  $B(y, x) = -B(x, y)$  pour tous  $x, y \in E$ .

En dimension finie, si  $\mathcal{B} = (e_1, \dots, e_n)$  est une base de  $E$ , on a vu que les  $B_{ij} = B(e_i, e_j)$  caractérisent  $B$ . Ils caractérisent aussi sa symétrie ou son antisymétrie.

**Lemme 3.1.** Soit  $E$  un  $\mathbb{R}$ -espace vectoriel,  $\mathcal{B} = (e_1, \dots, e_n)$  une base de  $E$  et  $B : E \times E \rightarrow \mathbb{R}$  une forme bilinéaire sur  $E$ . Alors  $B$  est *symétrique* si et seulement si  $B_{ji} = B_{ij}$  pour tous  $i, j$ , et  $B$  est *antisymétrique* si et seulement si  $B_{ji} = -B_{ij}$  pour tous  $i, j$ .

*Démonstration.* On a que

$$B(x, y) = \sum_{i,j=1}^n B_{ij} x_i y_j$$

pour tous  $x, y \in E$ , où les  $x_i$  et  $y_i$  sont les coordonnées de  $x$  et  $y$  dans  $\mathcal{B}$ . Il est alors clair que si  $B_{ji} = B_{ij}$  pour tous  $i, j$ , alors  $B(y, x) = B(x, y)$  pour tous  $x, y$  et que si  $B_{ji} = -B_{ij}$  pour tous  $i, j$ , alors  $B(y, x) = -B(x, y)$  pour tous  $x, y$ . Réciproquement si  $B(y, x) = B(x, y)$  pour tous  $x, y$ , et si  $i, j$  sont dans  $\{1, \dots, n\}$ , alors comme dans la preuve du Lemme 2.1, en prenant  $x$  tel que  $x_k = 0$  pour tout  $k \neq i$  et  $x_i = 1$ , et en prenant  $y$  tel que  $y_k = 0$  pour tout  $k \neq j$  et  $y_j = 1$ , on voit que  $B_{ji} = B_{ij}$ . De même, si on avait  $B(y, x) = -B(x, y)$  alors on aurait  $B_{ji} = -B_{ij}$ . Le lemme est démontré.  $\square$

Attention, une forme bilinéaire n'est pas soit symétrique soit antisymétrique. Elle peut très bien n'être ni l'un ni l'autre. Le théorème suivant a lieu.

**Théorème 3.1.** Soit  $E$  un  $\mathbb{R}$ -espace vectoriel. Toute forme bilinéaire sur  $E$  s'écrit comme somme d'une forme bilinéaire symétrique et d'une forme bilinéaire antisymétrique.

*Démonstration.* Soit  $B : E \times E \rightarrow \mathbb{R}$  une forme bilinéaire sur  $E$ . On pose

$$B_s(x, y) = \frac{1}{2} (B(x, y) + B(y, x)) \text{ et } B_a(x, y) = \frac{1}{2} (B(x, y) - B(y, x))$$

pour tous  $x, y \in E$ . On vérifie facilement que  $B_s$  et  $B_a$  sont deux formes bilinéaires sur  $E$ . On vérifie tout aussi facilement que  $B_s$  est symétrique et que  $B_a$  est antisymétrique. Toujours sans aucune difficulté, on vérifie pour finir que  $B(x, y) = B_s(x, y) + B_a(x, y)$  pour tous  $x, y$ . Le théorème est démontré.  $\square$

#### 4. CHANGEMENT DE BASE

Le théorème qui suit répond à la question du changement de matrice de représentation par changement de base.

**Théorème 4.1.** Soient  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie  $n$ ,  $\mathcal{B} = (e_1, \dots, e_n)$  et  $\tilde{\mathcal{B}} = (\tilde{e}_1, \dots, \tilde{e}_n)$  deux bases de  $E$ , et  $B : E \times E \rightarrow \mathbb{R}$  une forme bilinéaire sur  $E$ . Soit  $P = M_{\mathcal{B} \rightarrow \tilde{\mathcal{B}}}$  la matrice de passage de  $\mathcal{B}$  à  $\tilde{\mathcal{B}}$ . Alors

$$M_{\tilde{\mathcal{B}}}(B) = {}^t P M_{\mathcal{B}}(B) P ,$$

où  $M_{\mathcal{B}}(B)$  est la matrice de  $B$  dans la base  $\mathcal{B}$ , et  $M_{\tilde{\mathcal{B}}}(B)$  est la matrice de  $B$  dans la base  $\tilde{\mathcal{B}}$ .

*Démonstration.* Etant donné  $x \in E$ , soient  $V_{\mathcal{B}}(x)$  et  $V_{\tilde{\mathcal{B}}}(x)$  les matrices colonnes composées respectivement des coordonnées de  $x$  dans  $\mathcal{B}$  et des coordonnées de  $x$  dans  $\tilde{\mathcal{B}}$ . Pour  $x$  et  $y$  deux vecteurs quelconques de  $E$ , on a

$$B(x, y) = {}^t V_{\mathcal{B}}(x) M_{\mathcal{B}}(B) V_{\mathcal{B}}(y) ,$$

et

$$B(x, y) = {}^t V_{\tilde{\mathcal{B}}}(x) M_{\tilde{\mathcal{B}}}(B) V_{\tilde{\mathcal{B}}}(y) .$$

De plus,  $V_{\mathcal{B}}(x) = P V_{\tilde{\mathcal{B}}}(x)$ , et  $V_{\mathcal{B}}(y) = P V_{\tilde{\mathcal{B}}}(y)$ . On peut ainsi écrire que

$$\begin{aligned} B(x, y) &= {}^t V_{\mathcal{B}}(x) M_{\mathcal{B}}(B) V_{\mathcal{B}}(y) \\ &= {}^t V_{\tilde{\mathcal{B}}}(x) ({}^t P M_{\mathcal{B}}(B) P) V_{\tilde{\mathcal{B}}}(y) \\ &= {}^t V_{\tilde{\mathcal{B}}}(x) M_{\tilde{\mathcal{B}}}(B) V_{\tilde{\mathcal{B}}}(y) . \end{aligned}$$

Puisque  $x$  et  $y$  sont quelconques, cela impose

$$M_{\tilde{\mathcal{B}}}(B) = {}^t P M_{\mathcal{B}}(B) P ,$$

comme on le voit en prenant successivement des  $x$  et  $y$  de coordonnées dans  $\tilde{\mathcal{B}}$  des  $n$ -uplets du type  $(0, \dots, 0, 1, 0, \dots, 0)$ . Le théorème est démontré.  $\square$

**Exercice:** Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension 3, et  $\mathcal{B} = (e_1, e_2, e_3)$  une base de  $E$ . On note  $B$  la forme bilinéaire sur  $E$  définie par

$$\begin{aligned} B(x, y) &= x_1 y_1 - x_2 y_2 + x_3 y_3 - \frac{1}{2} (x_1 y_2 + x_2 y_1) \\ &\quad - \frac{7}{2} (x_1 y_3 + x_3 y_1) + \frac{1}{2} (x_2 y_3 + x_3 y_2) \end{aligned}$$

pour tous  $x = \sum_{i=1}^3 x_i e_i$  et  $y = \sum_{i=1}^3 y_i e_i$ .

(1) Ecrire la matrice  $M_{\mathcal{B}}(B)$  de  $B$  dans  $\mathcal{B}$ .

(2) Montrer que les vecteurs  $\tilde{e}_1 = e_1$ ,  $\tilde{e}_2 = e_1 + 2e_2$ , et  $\tilde{e}_3 = 3e_1 - e_2 + e_3$  forment une base  $\tilde{\mathcal{B}}$  de  $E$ . Ecrire la matrice de passage de la base  $\mathcal{B}$  à la base  $\tilde{\mathcal{B}}$ .

(3) Calculer la matrice  $M_{\tilde{\mathcal{B}}}(B)$  de  $B$  dans  $\tilde{\mathcal{B}}$ .

**Solution:** (1) La matrice  $M_{\mathcal{B}}(B)$  de  $B$  dans  $\mathcal{B}$  est composée des  $B(e_i, e_j)$ . On trouve

$$M_{\mathcal{B}}(B) = \begin{pmatrix} 1 & -\frac{1}{2} & -\frac{7}{2} \\ -\frac{1}{2} & -1 & \frac{1}{2} \\ -\frac{7}{2} & \frac{1}{2} & 1 \end{pmatrix}$$

(2) On a trois vecteurs en dimension 3. Il suffit donc de montrer que la famille  $(\tilde{e}_1, \tilde{e}_2, \tilde{e}_3)$  est libre. On a

$$\begin{aligned} \lambda_1 \tilde{e}_1 + \lambda_2 \tilde{e}_2 + \lambda_3 \tilde{e}_3 &= 0 \\ \Leftrightarrow \lambda_1 e_1 + \lambda_2(e_1 + 2e_2) + \lambda_3(3e_1 - e_2 + e_3) &= 0 \\ \Leftrightarrow (\lambda_1 + \lambda_2 + 3\lambda_3)e_1 + (2\lambda_2 - \lambda_3)e_2 + \lambda_3 e_3 &= 0 \end{aligned}$$

et puisque  $(e_1, e_2, e_3)$  est une base, on doit avoir que

$$\begin{cases} \lambda_1 + \lambda_2 + 3\lambda_3 = 0 \\ 2\lambda_2 - \lambda_3 = 0 \\ \lambda_3 = 0 \end{cases}$$

En remontant de système de la dernière équation à la première, on trouve que  $\lambda_1 = \lambda_2 = \lambda_3 = 0$ . La famille  $(\tilde{e}_1, \tilde{e}_2, \tilde{e}_3)$  est donc libre, et puisque  $E$  est de dimension 3, il s'agit d'une base de  $E$ . Les colonnes de la matrice de passage de la base  $\mathcal{B}$  à la base  $\tilde{\mathcal{B}}$  sont constituées des coordonnées dans  $\mathcal{B}$  des vecteurs de  $\tilde{\mathcal{B}}$ . Si  $P = M_{\mathcal{B} \rightarrow \tilde{\mathcal{B}}}$  est cette matrice de passage, on a donc

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

(3) Par formule de changement de base  $M_{\tilde{\mathcal{B}}}(B) = {}^t P M_{\mathcal{B}}(B) P$ . On calcule

$$\begin{aligned} M_{\mathcal{B}}(B)P &= \begin{pmatrix} 1 & -\frac{1}{2} & -\frac{7}{2} \\ -\frac{1}{2} & -1 & \frac{1}{2} \\ -\frac{7}{2} & \frac{1}{2} & 1 \end{pmatrix} \begin{pmatrix} 1 & 1 & 3 \\ 0 & 2 & -1 \\ 0 & 0 & 1 \end{pmatrix} \\ &= \begin{pmatrix} 1 & 0 & 0 \\ -\frac{1}{2} & -\frac{5}{2} & 0 \\ -\frac{7}{2} & -\frac{5}{2} & -10 \end{pmatrix} \end{aligned}$$

On a

$${}^t P = \begin{pmatrix} 1 & 0 & 0 \\ 1 & 2 & 0 \\ 3 & -1 & 1 \end{pmatrix}$$

Par suite

$$\begin{aligned} M_{\tilde{\mathcal{B}}}(B) &= {}^t P M_{\mathcal{B}}(B) P \\ &= \begin{pmatrix} 1 & 0 & 0 \\ 1 & 2 & 0 \\ 3 & -1 & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 & 0 \\ -\frac{1}{2} & -\frac{5}{2} & 0 \\ -\frac{7}{2} & -\frac{5}{2} & -10 \end{pmatrix} \\ &= \begin{pmatrix} 1 & 0 & 0 \\ 0 & -5 & 0 \\ 0 & 0 & -10 \end{pmatrix}. \end{aligned}$$

□

## 5. DUALITÉ

Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie. Le dual de  $E$ , noté  $E^*$ , est par définition l'espace  $L(E, \mathbb{R})$  des formes linéaires sur  $E$ .

**Théorème 5.1.** *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie  $n$ , et soit  $(e_1, \dots, e_n)$  une base de  $E$ . Pour  $i = 1, \dots, n$ , on définit les éléments  $e_i^* \in E^*$  par:*

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

*pour tout  $j = 1, \dots, n$ , où les  $\delta_{ij}$  sont les symboles de Kronecker. La famille  $(e_1^*, \dots, e_n^*)$  est alors une base de  $E^*$ , dite base duale de la base  $(e_1, \dots, e_n)$ . En particulier,  $E$  et  $E^*$  ont même dimension.*

*Proof.* Il suffit de montrer que la famille  $(e_1^*, \dots, e_n^*)$  est à la fois libre et génératrice dans  $E^*$ . On montre tout d'abord que la famille est libre. Supposons que pour des réels  $\lambda_1, \dots, \lambda_n$ ,

$$\lambda_1 e_1^* + \dots + \lambda_n e_n^* = 0.$$

Alors pour tout  $i = 1, \dots, n$ ,

$$\lambda_1 e_1^*(e_i) + \dots + \lambda_n e_n^*(e_i) = 0,$$

ce qui donne que  $\lambda_i = 0$  pour tout  $i = 1, \dots, n$  puisque  $e_j^*(e_i) = \delta_{ij}$ . La famille est donc bien libre. On vérifie par ailleurs que la famille est génératrice en remarquant que pour tout  $f \in E^*$ ,

$$f = \sum_{i=1}^n f(e_i) e_i^*.$$

Le membre de gauche et le membre de droite de cette identité sont en effet deux applications linéaires qui coïncident sur la base  $(e_1, \dots, e_n)$ , et si deux applications linéaires coïncident sur une même base, alors elles sont identiquement égales. Le théorème est démontré.  $\square$

Le théorème qui suit aborde la question du bi-dual  $E^{**} = (E^*)^*$  de  $E$ .

**Théorème 5.2.** *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie. Alors  $E^{**}$  s'assimile naturellement à  $E$  au sens où il existe un isomorphisme naturel de  $E$  dans  $E^{**}$ . Cet isomorphisme  $\Phi$  est défini par: pour tout  $x \in E$ ,  $\Phi(x)$  est l'élément de  $E^{**} = L(E^*, \mathbb{R})$  défini par la relation  $\Phi(x).(f) = f(x)$  pour tout  $f \in E^*$ .*

*Démonstration.* Il suit du théorème précédent que  $E^{**}$  a même dimension que  $E^*$ , qui a à son tour même dimension que  $E$ . Donc  $E$  et  $E^{**}$  ont mêmes dimensions. Pour montrer que  $\Phi : E \rightarrow E^{**}$  est un isomorphisme de  $E$  sur  $E^{**}$  il nous suffit donc de montrer que  $\Phi$  est injective (on passe sur le fait que  $\Phi$  est bien linéaire, ce qui est immédiat). Notons  $\mathcal{B} = (e_1, \dots, e_n)$  une base de  $E$ , et soit  $(e_1^*, \dots, e_n^*)$  sa base duale. Si  $x = \sum_i x_i e_i$  est tel que  $\Phi(x) = 0$ , alors  $\Phi(x).(e_i^*) = 0$  pour tout  $i = 1, \dots, n$ . Donc  $e_i^*(x) = 0$  pour tout  $i = 1, \dots, n$ . Or  $e_i^*(x) = x_i$ , de sorte que si  $\Phi(x) = 0$ , alors  $x = 0$ . D'où le théorème.  $\square$

On peut se poser la question de l'antédualité: étant donnée une base  $(\eta_1, \dots, \eta_n)$  de  $E^*$ , provient-elle d'une base  $(e_1, \dots, e_n)$  de  $E$  avec  $e_i^* = \eta_i$  pour tout  $i$ ? Le résultat suivant répond par l'affirmative à cette question.

**Théorème 5.3.** *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie  $n$ . Soit  $(\eta_1, \dots, \eta_n)$  une base de  $E^*$ . Il existe une base  $(e_1, \dots, e_n)$  de  $E$  qui est telle que  $e_i^* = \eta_i$  pour tout  $i$ .*

*Démonstration.* Soit  $\Phi : E \rightarrow E^{**}$  l'isomorphisme canonique du théorème précédent. Soit  $(\eta_1^*, \dots, \eta_n^*)$  la base duale de  $E^{**}$  obtenue à partir de  $(\eta_1, \dots, \eta_n)$ . On pose  $e_i = \Phi^{-1}(\eta_i^*)$ ,  $i = 1, \dots, n$ . Comme  $\Phi$  est un isomorphisme,  $(e_1, \dots, e_n)$  est une base de  $E$ . On a  $\Phi(e_i) = \eta_i^*$  pour tout  $i$  et donc, pour tous  $i, j = 1, \dots, n$ ,  $\Phi(e_i) \cdot (\eta_j) = \eta_i^*(\eta_j) = \delta_{ij}$  où les  $\delta_{ij}$  sont les symboles de Kronecker (vérifiant donc que  $\delta_{ij} = 1$  si  $i = j$  et 0 sinon). Or  $\Phi(e_i) \cdot (\eta_j) = \eta_j(e_i)$  et donc,  $\eta_j = e_j^*$  pour tout  $j$ . C'est exactement ce que l'on voulait démontrer.  $\square$

## 6. RANG D'UNE FORME BILINÉAIRE

Soient  $E$  un  $\mathbb{R}$ -espace vectoriel, et  $B : E \times E \rightarrow \mathbb{R}$  une forme bilinéaire sur  $E$ . On note  $B_1$  et  $B_2$  les applications de  $E$  dans  $E^*$  définies par: pour tout  $x \in E$ ,  $B_1(x)$  est défini par

$$B_1(x) \cdot (y) = B(x, y), \quad \forall y \in E$$

et, pour tout  $y \in E$ ,  $B_2(y)$  est défini par

$$B_2(y) \cdot (x) = B(x, y), \quad \forall x \in E.$$

On vérifie facilement que  $B_1$  et  $B_2$  sont linéaires, ce que l'on peut encore écrire sous la forme  $B_1 \in L(E, E^*)$  et  $B_2 \in L(E, E^*)$ .

Soit  $E$  un espace vectoriel de dimension finie,  $\mathcal{B} = (e_1, \dots, e_n)$  une base de  $E$  et  $\mathcal{B}^* = (e_1^*, \dots, e_n^*)$  la base duale de  $\mathcal{B}$ . Soit  $\eta \in E^*$ , et soient  $\eta_1, \dots, \eta_n$  les coordonnées de  $\eta$  dans  $\mathcal{B}^*$ . On a alors l'égalité  $\eta = \sum_{j=1}^n \eta_j e_j^*$ , et on obtient que  $\eta(e_i) = \eta_i$  pour tout  $i$  (puisque  $e_j^*(e_i) = \delta_{ij}$ ). Les coordonnées d'une forme  $\eta$  dans la base duale  $\mathcal{B}^*$  sont les  $\eta_i = \eta(e_i)$ . On démontre la proposition suivante.

**Proposition 6.1.** *Soient  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie,  $\mathcal{B}$  une base de  $E$ , et  $B : E \times E \rightarrow \mathbb{R}$  une forme bilinéaire sur  $E$ . Les trois quantités  $r_1, r_2, r_3$  définies par:*

- (i)  $r_1$  vaut le rang de  $B_1$ , application linéaire de  $E$  dans  $E^*$ ,
- (ii)  $r_2$  vaut le rang de  $B_2$ , application linéaire de  $E$  dans  $E^*$ , et
- (iii)  $r_3$  vaut le rang de la matrice de  $B$  dans  $\mathcal{B}$ ,

sont en fait égales.

*Démonstration.* Notons  $\mathcal{B}^*$  la base duale de  $\mathcal{B}$ , et  $M_{\mathcal{B}}(B)$  la matrice de  $B$  dans  $\mathcal{B}$ . On vérifie facilement que

$$M_{\mathcal{B}\mathcal{B}^*}(B_1) = {}^t M_{\mathcal{B}}(B),$$

où  $M_{\mathcal{B}\mathcal{B}^*}(B_1)$  est la matrice de représentation de  $B_1$  dans les bases  $\mathcal{B}$  et  $\mathcal{B}^*$ . En effet, les coordonnées de  $B_1(e_i)$  dans  $\mathcal{B}^*$  sont (cf. ci-dessus) les  $B_1(e_i) \cdot (e_j)$ , à savoir exactement les  $B_{ij}$ . De la même façon, on vérifie facilement que

$$M_{\mathcal{B}\mathcal{B}^*}(B_2) = M_{\mathcal{B}}(B),$$

où  $M_{\mathcal{B}\mathcal{B}^*}(B_2)$  est la matrice de représentation de  $B_2$  dans les bases  $\mathcal{B}$  et  $\mathcal{B}^*$ . En remarquant que les rangs d'une matrice et de sa transposée sont égaux (c'est évident avec la définition du rang via le déterminant des sous matrices), on obtient la proposition avec ce qui a été dit au chapitre précédent sur les rangs d'une application linéaire et de ses matrices de représentations.  $\square$

Grâce à cette proposition on peut définir le rang d'une forme bilinéaire.

**Définition 6.1.** Soient  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie,  $\mathcal{B}$  une base de  $E$ , et  $B : E \times E \rightarrow \mathbb{R}$  une forme bilinéaire sur  $E$ . Le rang de  $B$ , noté  $Rg(B)$ , est l'une des trois quantités équivalentes de la proposition précédente. Il ne dépend pas du choix de la base  $\mathcal{B}$ .

C'est souvent (iii) qui est utilisé dans la pratique. Attention pour montrer qu'une matrice  $n \times n$  est de rang  $p < n$  il ne suffit pas de trouver une sous matrice carrée d'ordre  $p$  inversible, il faut aussi montrer que toutes les sous matrices carrées d'ordre  $q > p$  ne le sont pas.

## 7. FORMES BILINÉAIRES ET FORMES QUADRATIQUES

On commence par la définition des formes quadratiques.

**Définition 7.1.** Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie. Une forme quadratique sur  $E$  est une application  $Q : E \rightarrow \mathbb{R}$  qui s'exprime, dans chaque base de  $E$ , sous la forme d'un polynôme homogène de degré 2 des variables coordonnées.

En d'autres termes,  $Q : E \rightarrow \mathbb{R}$  est une forme quadratique si, pour toute base  $\mathcal{B}$  de  $E$ , il existe des réels  $a_{ij}$ ,  $i \leq j$ , tels que

$$Q(x) = \sum_{\substack{i,j=1 \\ i \leq j}}^n a_{ij} x_i x_j, \quad (\star)$$

où  $n$  est la dimension de  $E$ , et les  $x_i$  sont les coordonnées de  $x$  dans  $\mathcal{B}$ . On vérifie facilement que si  $(\star)$  est vérifiée dans une base  $\mathcal{B}$ , elle l'est aussi (avec des  $a_{ij}$  différents !) dans toute autre base  $\tilde{\mathcal{B}}$  de  $E$ . En effet, l'expression de coordonnées  $\tilde{x}_i$  dans une base, en fonction des coordonnées  $x_i$  dans une autre base, est linéaire.

**Exercice:** Soit  $E$  de dimension 2,  $\mathcal{B} = (e_1, e_2)$  une base de  $E$  et  $\tilde{\mathcal{B}} = (\tilde{e}_1, \tilde{e}_2)$  une autre base de  $E$ . On note  $(x_1, x_2)$  les coordonnées d'un vecteur  $X$  dans  $\mathcal{B}$  et  $(\tilde{x}_1, \tilde{x}_2)$  les coordonnées de  $X$  dans  $\tilde{\mathcal{B}}$ . On suppose que  $\tilde{e}_1 = e_1 + 2e_2$  et  $\tilde{e}_2 = 3e_1 + e_2$ .

(1) Exprimer  $x_1$  et  $x_2$  en fonction de  $\tilde{x}_1$  et  $\tilde{x}_2$ .

(2) Soit  $Q$  la forme quadratique donnée dans  $\mathcal{B}$  par

$$Q(X) = 2x_1^2 + 6x_1x_2 - x_2^2.$$

Donner l'expression de  $Q$  dans  $\tilde{\mathcal{B}}$ .

**Solution:**  $Q$  est bien une forme quadratique car elle s'exprime sous la forme d'un polynôme homogène de degré 2 des variables coordonnées  $x_1, x_2$ . Dans  $\mathcal{B}$ ,  $a_{11} = 2$ ,  $a_{12} = 6$  et  $a_{22} = -1$ .

(1) On a

$$\begin{aligned} X &= \tilde{x}_1 \tilde{e}_1 + \tilde{x}_2 \tilde{e}_2 \\ &= \tilde{x}_1 (e_1 + 2e_2) + \tilde{x}_2 (3e_1 + e_2) \\ &= (\tilde{x}_1 + 3\tilde{x}_2)e_1 + (2\tilde{x}_1 + \tilde{x}_2)e_2 \\ &= x_1 e_1 + x_2 e_2. \end{aligned}$$

Donc

$$\begin{cases} x_1 = \tilde{x}_1 + 3\tilde{x}_2 \\ x_2 = 2\tilde{x}_1 + \tilde{x}_2. \end{cases}$$

(2) On a donc

$$\begin{aligned}
Q(X) &= 2x_1^2 + 6x_1x_2 - x_2^2 \\
&= 2(\tilde{x}_1 + 3\tilde{x}_2)^2 + 6(\tilde{x}_1 + 3\tilde{x}_2)(2\tilde{x}_1 + \tilde{x}_2) - (2\tilde{x}_1 + \tilde{x}_2)^2 \\
&= 2(\tilde{x}_1^2 + 9\tilde{x}_2^2 + 6\tilde{x}_1\tilde{x}_2) + 6(2\tilde{x}_1^2 + 3\tilde{x}_2^2 + 7\tilde{x}_1\tilde{x}_2) \\
&\quad - (4\tilde{x}_1^2 + \tilde{x}_2^2 + 4\tilde{x}_1\tilde{x}_2) \\
&= (2 + 12 - 4)\tilde{x}_1^2 + (12 + 42 - 4)\tilde{x}_1\tilde{x}_2 + (18 + 18 - 1)\tilde{x}_2^2 \\
&= 10\tilde{x}_1^2 + 50\tilde{x}_1\tilde{x}_2 + 35\tilde{x}_2^2 .
\end{aligned}$$

La forme quadratique  $Q$  qui était un polynôme homogène de degré 2 des variables coordonnées dans  $\mathcal{B}$  est encore un polynôme homogène de degré 2 des variables coordonnées dans  $\tilde{\mathcal{B}}$ . Par contre, bien sûr, les coefficients changent. On a maintenant, dans  $\tilde{\mathcal{B}}$ ,  $\tilde{a}_{11} = 10$ ,  $\tilde{a}_{12} = 50$  et  $\tilde{a}_{22} = 35$ .  $\square$

Le théorème qui suit établit le lien entre formes quadratiques et formes bilinéaires symétriques.

**Théorème 7.1.** *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie. Une fonction  $Q : E \rightarrow \mathbb{R}$  est une forme quadratique sur  $E$  si et seulement si :*

- (i)  $\forall \lambda \in \mathbb{R}, \forall x \in E, Q(\lambda x) = \lambda^2 Q(x)$ , et
- (ii) l'application  $B : E \times E \rightarrow \mathbb{R}$  définie par
$$B(x, y) = \frac{1}{2}(Q(x+y) - Q(x) - Q(y))$$
est bilinéaire sur  $E$ .

Dans ce cas,  $B$  est symétrique, on a  $B(x, x) = Q(x)$  pour tout  $x$ , et on dit que  $B$  est la forme bilinéaire associée à  $Q$  (encore appelée forme polaire de  $Q$ ).

Un modèle de forme quadratique est  $Q : \mathbb{R} \rightarrow \mathbb{R}$  donnée par  $Q(x) = x^2$ . La forme bilinéaire associée est  $B : \mathbb{R} \times \mathbb{R} \rightarrow \mathbb{R}$  donnée par  $B(x, y) = xy$ . L'identité de (ii) est l'identité remarquable

$$xy = \frac{1}{2}((x+y)^2 - x^2 - y^2) .$$

Une identité comme (ii) nous dit que l'on connaît tous les doubles produits dès lors que l'on connaît tous les carrés.

*Démonstration du théorème.* Si  $Q$  est une forme quadratique, (i) est bien évidemment vérifiée. Par ailleurs, si  $\mathcal{B}$  est une base de  $E$ , et si  $Q$  s'écrit dans  $\mathcal{B}$  sous la forme  $(\star)$ , alors le  $B$  donné par (ii) vaut

$$\begin{aligned}
B(x, y) &= \frac{1}{2} \sum_{\substack{i,j=1 \\ i \leq j}}^n a_{ij}(x_i + y_i)(x_j + y_j) - \frac{1}{2} \sum_{\substack{i,j=1 \\ i \leq j}}^n a_{ij}x_i x_j - \frac{1}{2} \sum_{\substack{i,j=1 \\ i \leq j}}^n a_{ij}y_i y_j \\
&= \frac{1}{2} \sum_{\substack{i,j=1 \\ i \leq j}}^n a_{ij}x_i y_j + \frac{1}{2} \sum_{\substack{i,j=1 \\ i \leq j}}^n a_{ij}x_j y_i
\end{aligned}$$

où les  $x_i$  et  $y_i$  sont les coordonnées de  $x$  et  $y$  dans  $\mathcal{B}$ . Il suit que (ii) est elle aussi vérifiée, à savoir que  $B$  est bien bilinéaire, si  $Q$  est quadratique. Réciproquement, supposons que (i) et (ii) soient vérifiées. En écrivant que

$$B(x, x) = \frac{1}{2}(Q(2x) - 2Q(x)) ,$$

et en utilisant (i), on voit que  $Q(x) = B(x, x)$ . Etant donnée une base  $\mathcal{B}$  de  $E$ , composée de vecteurs  $e_i$ , et puisque  $B$  est bilinéaire, il existe des  $B_{ij} = B(e_i, e_j)$  tels que

$$B(x, y) = \sum_{i,j} B_{ij} x_i y_j ,$$

où les  $x_i$  et  $y_i$  sont les coordonnées de  $x$  et  $y$  dans  $\mathcal{B}$ . Il s'ensuit que

$$\begin{aligned} Q(x) &= \sum_{i,j=1}^n B_{ij} x_i x_j , \\ &= \sum_{\substack{i,j=1 \\ i \leq j}}^n a_{ij} x_i x_j \end{aligned}$$

pour des  $a_{ij}$  convenables obtenus en regroupant les  $B_{ij}$  avec les  $B_{ji}$  pour  $i < j$ . Cela montre que  $Q$  est bien une forme quadratique. Le théorème est démontré.  $\square$

Avec ce théorème on obtient une autre définition possible des formes quadratiques.

**Théorème 7.2.** *Une application  $Q : E \rightarrow \mathbb{R}$  est une forme quadratique sur  $E$  s'il existe une forme bilinéaire symétrique  $B$  sur  $E$  pour laquelle  $Q(x) = B(x, x)$  pour tout  $x \in E$ . Dans ce cas  $B$  est la forme polaire associée à  $Q$ .*

*Démonstration du théorème.* Soit donc  $B$  bilinéaire symétrique et  $Q$  définie par  $Q(x) = B(x, x)$  pour tout  $x$ . Comme  $B$  est bilinéaire,

$$Q(\lambda x) = B(\lambda x, \lambda x) = \lambda B(x, \lambda x) = \lambda^2 B(x, x) = \lambda^2 Q(x)$$

pour tout  $\lambda \in \mathbb{R}$  et tout  $x \in E$ . De plus, pour tous  $x, y \in E$ ,

$$\begin{aligned} Q(x+y) - Q(x) - Q(y) &= B(x+y, x+y) - B(x, x) - B(y, y) \\ &= B(x, x) + B(x, y) + B(y, x) + B(y, y) - B(x, x) - B(y, y) \\ &= B(x, y) + B(y, x) \\ &= 2B(x, y) \end{aligned}$$

et donc, en raison du Théorème 7.1,  $Q$  est une forme quadratique et  $B$  est sa forme polaire associée.  $\square$

Indépendamment, de la preuve du Théorème 7.1 ci-dessus on tire le lemme très pratique suivant.

**Lemme 7.1.** *Soit  $E$  de dimension  $n$  et  $\mathcal{B}$  une base de  $E$ . Soit  $Q$  une forme quadratique sur  $E$  dont l'expression dans  $\mathcal{B}$  est*

$$Q(x) = \sum_{\substack{i,j=1 \\ i \leq j}}^n a_{ij} x_i x_j .$$

*Les coefficients  $B_{ij}$  de la forme bilinéaire symétrique associée à  $Q$  sont alors donnés par  $B_{ij} = \frac{a_{ij}}{2}$  si  $i < j$ ,  $B_{ii} = a_{ii}$  pour tout  $i$  et  $B_{ij} = \frac{a_{ji}}{2}$  si  $i > j$ .*

**Exercice:** Soit  $E$  de dimension 2,  $\mathcal{B} = (e_1, e_2)$  une base de  $E$ . On note  $(x_1, x_2)$  les coordonnées d'un vecteur  $X$  dans  $\mathcal{B}$ . Soit  $Q$  la forme quadratique donnée dans  $\mathcal{B}$  par

$$Q(x) = 2x_1^2 + 6x_1x_2 - x_2^2 .$$

Donner l'expression dans  $\mathcal{B}$  de la forme bilinéaire symétrique  $B$  associée à  $Q$ .

**Solution:** Ici  $a_{11} = 2$ ,  $a_{12} = 6$  et  $a_{22} = -1$ . On a donc  $B_{11} = 2$ ,  $B_{12} = 3$ ,  $B_{21} = 3$  et  $B_{22} = -1$ .

Pour rappel, on vérifie facilement qu'une forme bilinéaire  $B$  sur  $E$  est symétrique si et seulement si, pour toute base de  $E$ , sa matrice dans cette base est symétrique.

**Définition 7.2.** Soit  $E$  un  $\mathbb{R}$ -espace vectoriel et  $Q$  une forme quadratique sur  $E$ . Soit  $B$  la forme bilinéaire symétrique associée à  $Q$ . Les vecteurs  $x \in E$  qui vérifient que  $Q(x) = 0$  (donc  $B(x, x) = 0$ ) sont dits isotropes (sous entendu pour  $Q$  ou pour  $B$ ). L'ensemble  $\mathcal{C}$  des vecteurs isotropes est appelé le cône isotrope. Le noyau de  $Q$  est l'ensemble des  $x \in E$  pour lesquels  $B(x, y) = 0$  pour tout  $y \in E$ . On le note  $\text{Ker}(Q)$ .

On parle de cône car si  $x \in \mathcal{C}$  alors, pour tout  $\lambda \in \mathbb{R}$ ,  $\lambda x \in \mathcal{C}$ . Soient par exemple  $E$  de dimension 2,  $\mathcal{B}$  une base de  $E$  et  $Q$  la forme quadratique donnée dans  $\mathcal{B}$  par

$$Q(x) = x_1^2 + 2x_1x_2 - 3x_2^2.$$

On a

$$Q(x) = (x_1 + x_2)^2 - 4x_2^2$$

et on voit donc que le cône isotrope  $\mathcal{C}$  est constitué des  $(x_1, x_2)$  qui sont tels que  $(x_1 + x_2)^2 = 4x_2^2$ . Soit donc  $x_1 + x_2 = 2x_2$  ou  $x_1 + x_2 = -2x_2$  et on trouve donc que

$$\mathcal{C} = \{(x_1, x_2) / x_1 = x_2 \text{ ou } x_1 = -3x_2\}.$$

La forme bilinéaire symétrique associée est donnée par

$$\begin{aligned} B(x, y) &= x_1y_1 + x_1y_2 + x_2y_1 - 3x_2y_2 \\ &= (x_1 + x_2)y_1 + (x_1 - 3x_2)y_2. \end{aligned}$$

On a donc  $B(x, y) = 0$  pour tout  $y$  si et seulement si  $x_1 + x_2 = 0$  et  $x_1 - 3x_2 = 0$ . On trouve alors  $x_1 = x_2 = 0$ . Et donc  $\text{Ker}(Q) = \{0\}$ .

Pour finir, le rang d'une forme quadratique se définit comme étant le rang de la forme bilinéaire symétrique qui lui est associée.

**Définition 7.3.** Si  $Q$  est une forme quadratique sur un  $\mathbb{R}$ -espace vectoriel de dimension finie, on appelle rang de  $Q$  le rang de la forme bilinéaire symétrique qui lui est associée.

## 8. ORTHOGONALITÉ

On aborde maintenant l'orthogonalité relative à une forme bilinéaire donnée. Elle diffère de l'orthogonalité classique, notamment en raison de la possible présence de vecteurs isotropes.

**Définition 8.1.** Soit  $E$  un  $\mathbb{R}$ -espace vectoriel et  $B$  une forme bilinéaire symétrique sur  $E$ . On dit que deux vecteurs  $x$  et  $y$  de  $E$  sont  $B$ -orthogonaux si  $B(x, y) = 0$ . Si  $A$  est un sous ensemble de  $E$ , le  $B$ -orthogonal de  $A$ , noté  $A^{\perp_B}$ , est défini par

$$A^{\perp_B} = \{y \in E / \forall x \in A, B(x, y) = 0\}.$$

En d'autres termes,  $A^{\perp_B}$  est l'ensemble des vecteurs de  $E$  qui sont  $B$ -orthogonaux à tous les vecteurs de  $A$ .

On vérifie facilement que pour tout sous ensemble  $A$  de  $E$ ,  $A^{\perp_B}$  est un sous espace vectoriel de  $E$ .

**Lemme 8.1.** *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel et  $B$  une forme bilinéaire symétrique sur  $E$ . Soit  $F \subset E$  un sous espace vectoriel de  $E$ . Si tous les vecteurs de  $F$  sont isotropes, alors  $B = 0$  sur  $F \times F$ .*

*Démonstration.* Soient  $x, y \in F$ . Si  $F$  est un espace vectoriel,  $x + y \in F$ . On a alors  $Q(x) = 0$ ,  $Q(y) = 0$  et  $Q(x + y) = 0$  où  $Q$  est la forme quadratique associée à  $B$ . Avec le Théorème 7.1 on obtient alors que  $B(x, y) = 0$ . Comme  $x$  et  $y$  sont quelconques,  $B = 0$  sur  $F \times F$ .  $\square$

On pourra vérifier que les relations suivantes sont vraies:

- (i)  $\forall A \subset B \subset E, B^{\perp_B} \subset A^{\perp_B}$ ,
- (ii)  $\forall A, B \subset E, (A \cup B)^{\perp_B} = A^{\perp_B} \cap B^{\perp_B}$ ,
- (iii)  $\forall A, B \subset E, A^{\perp_B} + B^{\perp_B} \subset (A \cap B)^{\perp_B}$ ,
- (iv)  $\forall A \subset E, A^{\perp_B} = \text{Vect}(A)^{\perp_B}$ , et
- (v)  $\forall A \subset E, A \subset (A^{\perp_B})^{\perp_B}$ .

Les notions de  $B$ -orthogonalité et de bases donnent lieu à la définition suivante.

**Définition 8.2.** *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie,  $(e_1, \dots, e_n)$  une base de  $E$ , et  $B$  une forme bilinéaire symétrique sur  $E$ . On dit que la base  $(e_1, \dots, e_n)$  est  $B$ -orthogonale si pour tous  $i \neq j \in \{1, \dots, n\}$ ,  $B(e_i, e_j) = 0$ . La base est dite  $B$ -orthonormale si pour tous  $i, j \in \{1, \dots, n\}$ ,  $B(e_i, e_j) = \delta_{ij}$ , où les  $\delta_{ij}$  sont les symboles de Kronecker.*

Une base  $B$ -orthonormale est une base  $B$ -orthogonale pour laquelle on a en plus que  $B(e_i, e_i) = 1$  pour tout  $i$ . Attention, en raison de la possibilité d'avoir des vecteurs isotropes parmi les  $e_i$  on ne peut pas passer d'une famille orthogonale à une famille orthonormale en divisant par les  $B(e_i, e_i)$ . Imaginons le cas extrême  $B = 0$ . Alors toute base est  $B$ -orthogonale, mais aucune n'est  $B$ -orthonormale.

**Théorème 8.1.** *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie. Pour toute forme bilinéaire symétrique  $B$  sur  $E$ , il existe une base de  $E$  qui est  $B$ -orthogonale.*

*Démonstration.* Soit  $n$  la dimension de  $E$ . On démontre le théorème par récurrence sur  $n$ . Si  $n = 1$ , le théorème est trivialement vrai. On suppose maintenant que pour tout  $\mathbb{R}$ -espace vectoriel de dimension  $n$ , et toute forme bilinéaire symétrique  $B$  sur cet espace, il existe une base de l'espace qui soit  $B$ -orthogonale. On considère alors  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension  $n + 1$ , et  $B$  une forme bilinéaire symétrique sur  $E$ . On veut montrer que  $E$  possède une base  $B$ -orthogonale. Sans perdre en généralité, on pourra supposer que  $B$  n'est pas la forme bilinéaire nulle (auquel cas le résultat est vrai, toute base de  $E$  étant  $B$ -orthogonale). Si  $B$  n'est pas identiquement nulle, il existe un vecteur  $e_{n+1}$  de  $E$  pour lequel  $Q(e_{n+1}) \neq 0$ , où  $Q$  est la forme quadratique associée à  $B$ . En effet, si  $Q \equiv 0$  alors  $B \equiv 0$ . On note  $F$  le sous espace vectoriel (la droite vectorielle) de  $E$  de base  $e_{n+1}$ , et on note  $F^{\perp_B}$  le  $B$ -orthogonal de  $F$ . Alors

$$F^{\perp_B} = \{x \in E / B(e_{n+1}, x) = 0\},$$

de sorte que  $F^{\perp_B} = \text{Ker} B(e_{n+1}, \cdot)$ , où  $B(e_{n+1}, \cdot)$  est la forme linéaire sur  $E$  qui à  $x$  associe  $B(e_{n+1}, x)$ . Cette forme linéaire est surjective, puisque non nulle, et donc de rang 1. Il suit du théorème du rang que son noyau est de dimension  $n$ , et donc que  $F^{\perp_B}$  est de dimension  $n$ . On vérifie facilement que  $F \cap F^{\perp_B} = \{0\}$ . En effet, si  $x \in F \cap F^{\perp_B}$ , alors  $x \in F$  entraîne que  $x = \lambda e_{n+1}$ , tandis que  $x \in F^{\perp_B}$  entraînent que  $x$  doit être  $B$ -orthogonal à lui-même, et donc que  $B(x, x) = 0$ . Or  $B(\lambda e_{n+1}, \lambda e_{n+1}) = \lambda^2 B(e_{n+1}, e_{n+1})$  de sorte que nécessairement  $\lambda = 0$ , et donc  $x = 0$  (puisque  $Q(e_{n+1}) = B(e_{n+1}, e_{n+1}) \neq 0$ ). Comme  $\dim F = 1$ ,  $\dim F^{\perp_B} = n$ ,  $\dim E = n + 1$ , et  $F \cap F^{\perp_B} = \{0\}$ , on a que

$$E = F \oplus F^{\perp_B} .$$

Par hypothèse de récurrence, il existe une base  $(e_1, \dots, e_n)$  de  $F^{\perp_B}$  qui est  $B$ -orthogonale. Puisque  $E = F \oplus F^{\perp_B}$ ,  $(e_1, \dots, e_{n+1})$  est une base de  $E$ . Par ailleurs, cette base est  $B$ -orthogonale puisque  $e_{n+1} \in F$  et les  $e_i$ ,  $i = 1, \dots, n$ , sont dans  $F^{\perp_B}$ . On a donc montré que si la propriété d'existence d'une base  $B$ -orthogonale est vraie en dimension  $n$ , alors elle l'est aussi en dimension  $n + 1$ . Cela démontre le théorème par récurrence.  $\square$

## 9. LE THÉORÈME DE SYLVESTER

Le théorème de Sylvester est un théorème de classification des formes quadratiques. Une forme quadratique est ainsi (quand on la lit convenablement) toujours une somme/différence de carrés.

**Théorème 9.1** (Loi d'inertie de Sylvester). *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie  $n$ , et  $Q$  une forme quadratique de rang  $r$  sur  $E$ . Il existe alors un entier  $p \leq r$ , et une base  $\mathcal{B}$  de  $E$ , de sorte que, pour tout  $x \in E$  de coordonnées  $(x_1, \dots, x_n)$  dans  $\mathcal{B}$ ,*

$$Q(x) = x_1^2 + \dots + x_p^2 - x_{p+1}^2 - \dots - x_r^2 .$$

*De plus,  $p$  ne dépend que de  $Q$  et pas de  $\mathcal{B}$  au sens où s'il existe une autre base  $\tilde{\mathcal{B}}$  de  $E$ , et un entier  $\tilde{p} \leq r$ , tels que pour tout  $x \in E$  de coordonnées  $(\tilde{x}_1, \dots, \tilde{x}_n)$  dans  $\tilde{\mathcal{B}}$ ,*

$$Q(x) = \tilde{x}_1^2 + \dots + \tilde{x}_{\tilde{p}}^2 - \tilde{x}_{\tilde{p}+1}^2 - \dots - \tilde{x}_r^2 ,$$

*alors, nécessairement,  $\tilde{p} = p$ .*

En d'autres termes, pour toute forme quadratique sur un  $\mathbb{R}$ -espace vectoriel  $E$  de dimension finie, il existe une base  $\mathcal{B}$  de  $E$  dans laquelle  $Q$  s'écrit sous la forme

$$Q(x) = \sum_{\text{sur certains } i} x_i^2 - \sum_{\text{sur d'autres } i} x_i^2$$

et le nombre de carrés en positif et le nombre de carrés en négatif ne dépendent que de  $Q$  et pas de  $\mathcal{B}$ . En ce qui concerne la terminologie, l'entier  $p$  du théorème permet de définir ce que l'on appelle la signature de  $Q$ .

**Définition 9.1.** *Une forme quadratique  $Q$  est dite de signature  $(p, q)$  si  $p + q = \text{Rg}(Q)$ , et s'il existe une base  $\mathcal{B}$  de  $E$  dans laquelle*

$$Q(x) = x_1^2 + \dots + x_p^2 - x_{p+1}^2 - \dots - x_r^2 .$$

*La signature est indépendante de la base.*

On démontre maintenant le théorème de Sylvester. Il s'agit essentiellement de la voir comme une conséquence de l'existence des bases orthogonales.

*Démonstration.* (1) On montre pour commencer qu'il existe un entier  $p \leq r$  et une base de  $E$  de sorte que, dans cette base,  $Q$  s'écrive comme dans le théorème de Sylvester. Pour cela, on note  $B$  la forme bilinéaire associée à  $Q$ , et on considère  $\mathcal{B} = (e_1, \dots, e_n)$  une base  $B$ -orthogonale de  $E$ . La matrice de  $B$  dans  $\mathcal{B}$  est alors diagonale, puisque  $B(e_i, e_j) = 0$  si  $i \neq j$ . Notons  $\lambda_i = Q(e_i)$ ,  $i = 1, \dots, n$ , les termes sur la diagonale de la matrice  $M_{\mathcal{B}}(B)$  de  $B$  dans  $\mathcal{B}$ . Il est clair que

$$\text{Rg}(M_{\mathcal{B}}(B)) = \text{Nombre de } \lambda_i \text{ non nuls}.$$

Si on note  $r = \text{Rg}(B)$ , et quitte à permuter les vecteurs  $e_i$ , on peut supposer que  $\lambda_i \neq 0$  si  $i \leq r$ , et  $\lambda_i = 0$  si  $i = r+1, \dots, n$ . Si  $x$  a pour coordonnées  $(x_1, \dots, x_n)$  dans  $\mathcal{B}$ , on a que

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

Soit  $p$  le nombre de  $\lambda_i$  strictement positifs. Toujours quitte à permuter les  $e_i$ , on peut supposer que

$$\begin{aligned} \lambda_i &> 0 \text{ si } 1 \leq i \leq p, \\ \lambda_i &< 0 \text{ si } p+1 \leq i \leq r. \end{aligned}$$

On pose

$$\begin{aligned} \mu_i &= \sqrt{\lambda_i} \text{ si } 1 \leq i \leq p, \\ \mu_i &= \sqrt{-\lambda_i} \text{ si } p+1 \leq i \leq r. \end{aligned}$$

et on pose ensuite

$$\begin{aligned} \tilde{e}_i &= \frac{1}{\mu_i} e_i \text{ si } 1 \leq i \leq r, \\ \tilde{e}_i &= e_i \text{ si } r+1 \leq i \leq n. \end{aligned}$$

Alors  $\tilde{\mathcal{B}} = (\tilde{e}_1, \dots, \tilde{e}_n)$  est une base  $B$ -orthogonale de  $E$ , ce qui se vérifie facilement en se souvenant que  $\mathcal{B}$  l'est. Par suite, si  $x$  a pour coordonnées  $(\tilde{x}_1, \dots, \tilde{x}_n)$  dans  $\tilde{\mathcal{B}}$ , alors,

$$Q(x) = \sum_{i=1}^p \tilde{x}_i^2 - \sum_{i=p+1}^r \tilde{x}_i^2$$

puisque  $Q(\tilde{e}_i) = 1$  si  $1 \leq i \leq p$ , et  $Q(\tilde{e}_i) = -1$  si  $p+1 \leq i \leq r$ . Cela démontre la première partie du théorème de Sylvester.

(2) On démontre maintenant la seconde partie du théorème de Sylvester. On suppose qu'il existe deux entiers  $p, \tilde{p}$ , et deux bases  $\mathcal{B} = (e_1, \dots, e_n)$  et  $\tilde{\mathcal{B}} = (\tilde{e}_1, \dots, \tilde{e}_n)$  de  $E$  tels que, dans  $\mathcal{B}$  et dans  $\tilde{\mathcal{B}}$ ,

$$Q(x) = x_1^2 + \dots + x_p^2 - x_{p+1}^2 - \dots - x_r^2$$

et

$$Q(x) = \tilde{x}_1^2 + \dots + \tilde{x}_{\tilde{p}}^2 - \tilde{x}_{\tilde{p}+1}^2 - \dots - \tilde{x}_{\tilde{r}}^2.$$

On note

$E_1$  le sous espace vectoriel de  $E$  de base  $(e_1, \dots, e_p)$ ,

$E_2$  le sous espace vectoriel de  $E$  de base  $(e_{p+1}, \dots, e_n)$ ,

$\tilde{E}_1$  le sous espace vectoriel de  $E$  de base  $(\tilde{e}_1, \dots, \tilde{e}_{\tilde{p}})$ ,

$\tilde{E}_2$  le sous espace vectoriel de  $E$  de base  $(\tilde{e}_{\tilde{p}+1}, \dots, \tilde{e}_n)$ .

On a clairement que

$$Q(x) > 0 \text{ si } x \in E_1 \setminus \{0\},$$

$$Q(x) \leq 0 \text{ si } x \in E_2,$$

$$Q(x) > 0 \text{ si } x \in \tilde{E}_1 \setminus \{0\},$$

$$Q(x) \leq 0 \text{ si } x \in \tilde{E}_2.$$

On en déduit que

$$E_1 \cap \tilde{E}_2 = \{0\} \quad \text{et} \quad \tilde{E}_1 \cap E_2 = \{0\},$$

de sorte que, par propriété de la dimension d'une somme directe,

$$\dim E_1 + \dim \tilde{E}_2 \leq n \quad \text{et} \quad \dim \tilde{E}_1 + \dim E_2 \leq n.$$

La première de ces deux inégalités entraîne que  $p \leq \tilde{p}$ . La seconde entraîne que  $\tilde{p} \leq p$ . On en déduit que  $\tilde{p} = p$ . La seconde partie du théorème de Sylvester est démontrée.  $\square$

Un corollaire au théorème de Sylvester est le suivant.

**Corollaire 9.1.** *Si  $Q_1$  et  $Q_2$  sont deux formes quadratiques sur un  $\mathbb{R}$ -espace vectoriel de dimension finie, et si  $Q_1$  et  $Q_2$  ont même signature, alors il existe  $\Phi \in \text{End}(E)$  un isomorphisme de  $E$  tel que  $Q_2(x) = Q_1(\Phi(x))$  pour tout  $x \in E$ . Réciproquement, s'il existe un isomorphisme  $\Phi$  de  $E$  tel que  $Q_2(x) = Q_1(\Phi(x))$  pour tout  $x \in E$ , alors  $Q_1$  et  $Q_2$  ont même signature.*

*Démonstration.* Si  $Q_1$  et  $Q_2$  ont même signature  $(p, q)$ , il existe, en vertu du théorème de Sylvester, deux bases  $\mathcal{B} = (e_1, \dots, e_n)$  et  $\tilde{\mathcal{B}} = (\tilde{e}_1, \dots, \tilde{e}_n)$  de  $E$  telles que, dans ces bases,

$$Q_1(x) = x_1^2 + \dots + x_p^2 - x_{p+1}^2 - \dots - x_r^2,$$

$$Q_2(x) = \tilde{x}_1^2 + \dots + \tilde{x}_p^2 - \tilde{x}_{p+1}^2 - \dots - \tilde{x}_r^2.$$

Soit  $\Phi$  l'isomorphisme de  $E$  qui envoie  $\tilde{e}_i$  sur  $e_i$ . Alors pour tout  $x = \sum_i \tilde{x}_i \tilde{e}_i$  dans  $E$ ,

$$Q_1(\Phi(x)) = Q_2(x).$$

Cela démontre la première partie du corollaire. Réciproquement, on considère  $Q_1$  et  $Q_2$  deux formes quadratiques sur  $E$ , et  $\Phi$  un isomorphisme de  $E$  tels que  $Q_2(x) = Q_1(\Phi(x))$  pour tout  $x \in E$ . Soit  $(p, q)$  la signature de  $Q_2$ . Soit de plus  $\tilde{\mathcal{B}} = (\tilde{e}_1, \dots, \tilde{e}_n)$  une base de  $E$  telle que, dans cette base,

$$Q_2(x) = \tilde{x}_1^2 + \dots + \tilde{x}_p^2 - \tilde{x}_{p+1}^2 - \dots - \tilde{x}_r^2.$$

On note  $\mathcal{B}$  la base de  $E$  formée des  $e_i = \Phi(\tilde{e}_i)$ . C'est bien une base puisque (cf. les cours précédents) les isomorphismes envoient les bases sur des bases. On a alors

$$\begin{aligned} Q_1\left(\sum_i x_i e_i\right) &= Q_1\left(\Phi\left(\sum_i x_i \tilde{e}_i\right)\right) \\ &= Q_2\left(\sum_i x_i \tilde{e}_i\right) \\ &= x_1^2 + \dots + x_p^2 - x_{p+1}^2 - \dots - x_r^2. \end{aligned}$$

On en déduit que, dans  $\mathcal{B}$ ,  $Q_1$  s'écrit sous la forme

$$Q_1(x) = x_1^2 + \cdots + x_p^2 - x_{p+1}^2 - \cdots - x_r^2,$$

et donc que  $Q_1$  est aussi de signature  $(p, q)$ . Cela démontre la seconde partie du corollaire.  $\square$

Deux formes quadratiques  $Q_1$  et  $Q_2$  sur un  $\mathbb{R}$ -espace vectoriel  $E$  sont dites *équivalentes* s'il existe un isomorphisme  $\Phi \in \text{End}(E)$  de  $E$  qui est tel que  $Q_2(x) = Q_1(\Phi(x))$  pour tout  $x \in E$ . Une formulation équivalente du corollaire est que: *deux formes quadratiques sur un  $\mathbb{R}$ -espace vectoriel de dimension finie sont équivalentes si et seulement si elles ont même signature.*

**Remarque 9.1.** Une remarque bien pratique est la suivante. Supposons que dans une base  $\mathcal{B}$  une forme quadratique  $Q$  s'écrit

$$Q(x) = a_1x_1^2 + \cdots + a_px_p^2 - a_{p+1}x_{p+1}^2 - \cdots - a_rx_r^2$$

avec  $a_1 > 0, \dots, a_r > 0$ . La signature de  $Q$  est alors  $(p, q)$  où  $q$  est tel que  $p + q = r$ , autrement dit la même que s'il n'y avait pas les  $a_i$ . Pour le voir il suffit de considérer la base  $\tilde{\mathcal{B}} = (\frac{1}{\sqrt{a_1}}e_1, \dots, \frac{1}{\sqrt{a_r}}e_r, e_{r+1}, \dots, e_n)$ . On a alors  $\tilde{x}_1 = \sqrt{a_1}x_1, \dots, \tilde{x}_r = \sqrt{a_1}x_r, \tilde{x}_{r+1} = x_{r+1}, \dots, \tilde{x}_n = x_n$ . Dans cette base on a alors que

$$Q(x) = \tilde{x}_1^2 + \cdots + \tilde{x}_p^2 - \tilde{x}_{p+1}^2 - \cdots - \tilde{x}_r^2$$

et la signature de  $Q$  est donc bien comme annoncée. Plus généralement, si dans une base  $\mathcal{B}$  une forme quadratique  $Q$  s'écrit

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

alors la signature  $(p, q)$  de  $Q$  est donnée par  $p =$  nombre de  $a_i$  strictement positifs et  $q =$  nombre de  $a_i$  strictement négatifs. Le rang est  $r = p + q$ . Même analyse que ci-dessus mais il faut en plus permuer les coordonnées pour mettre les strictement positifs en premiers, suivis des strictements négatifs, suivis des nuls.

**Exercice :** Soit  $Q : \mathbb{R}^3 \rightarrow \mathbb{R}$  l'application définie par

$$Q(x_1, x_2, x_3) = x_1^2 + x_2^2 - 5x_3^2 + 2x_2x_3 - 2x_1x_3 - 6x_1x_2$$

- (1) Pourquoi  $Q$  est-elle une forme quadratique ?
- (2) Quelle est la forme bilinéaire symétrique  $B$  associée à  $Q$  ? Quelle est la matrice  $M$  de  $B$  dans la base canonique de  $\mathbb{R}^3$  ?
- (3) Décomposer  $Q$  en sommes et différences de carrés de façon à écrire  $Q$  sous la forme  $Q = \tilde{Q} \circ \Phi$  où  $\tilde{Q}$  est une forme quadratique facile à manipuler et  $\Phi$  est un isomorphisme de  $\mathbb{R}^3$ .
- (4) Quelle est la signature de  $Q$  ?
- (5) Caractériser les vecteurs isotropes de  $Q$ , à savoir les vecteurs  $x$  pour lesquels  $Q(x) = 0$ .

**Solution:** (1)  $Q$  est une forme quadratique car  $Q$  est un polynôme homogène de degré 2 des variables coordonnées (dans la base canonique de  $\mathbb{R}^3$ ).

(2) On écrit

$$\begin{aligned} Q(X) &= x_1^2 + x_2^2 - 5x_3^2 + 2x_2x_3 - 2x_1x_3 - 6x_1x_2 \\ &= x_1^2 + x_2^2 - 5x_3^2 + (x_2x_3 + x_3x_2) \\ &\quad - (x_1x_3 + x_3x_1) - 3(x_1x_2 + x_2x_1) . \end{aligned}$$

On en déduit que

$$\begin{aligned} B(X, Y) &= x_1y_1 + x_2y_2 - 5x_3y_3 + (x_2y_3 + x_3y_2) \\ &\quad - (x_1y_3 + x_3y_1) - 3(x_1y_2 + x_2y_1) \end{aligned}$$

La matrice de  $B$  dans la base canonique de  $\mathbb{R}^3$  est la matrice des  $a_{ij}$ . Elle est donnée par

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

(3) On rappelle que

$$Q(X) = x_1^2 + x_2^2 - 5x_3^2 + 2x_2x_3 - 2x_1x_3 - 6x_1x_2$$

On essaye d'éliminer  $x_1$  de l'équation. On remarque que

$$(x_1 - x_3 - 3x_2)^2 = \mathbf{x}_1^2 + x_3^2 + 9x_2^2 - 2\mathbf{x}_1\mathbf{x}_3 - 6\mathbf{x}_1\mathbf{x}_2 + 6x_2x_3$$

On peut donc écrire que

$$\begin{aligned} Q(X) &= (x_1 - x_3 - 3x_2)^2 - (x_3^2 + 9x_2^2 + 6x_2x_3) + x_2^2 - 5x_3^2 + 2x_2x_3 \\ &= (x_1 - 3x_2 - x_3)^2 - 8x_2^2 - 6x_3^2 - 4x_2x_3 \end{aligned}$$

On élimine maintenant  $x_2$  des trois termes  $8x_2^2 + 6x_3^2 + 4x_2x_3$ . On écrit pour cela que

$$8x_2^2 + 6x_3^2 + 4x_2x_3 = 8 \left( x_2 + \frac{1}{4}x_3 \right)^2 + \frac{11}{2}x_3^2$$

Par suite

$$Q(X) = (x_1 - 3x_2 - x_3)^2 - 8 \left( x_2 + \frac{1}{4}x_3 \right)^2 - \frac{11}{2}x_3^2$$

Soit  $\tilde{Q}$  la forme quadratique

$$\tilde{Q}(X) = x_1^2 - 8x_2^2 - \frac{11}{2}x_3^2$$

et soit  $\Phi : \mathbb{R}^3 \rightarrow \mathbb{R}^3$  l'endomorphisme de  $\mathbb{R}^3$  donné par

$$\Phi(x_1, x_2, x_3) = \left( x_1 - 3x_2 - x_3, x_2 + \frac{1}{4}x_3, x_3 \right) .$$

Alors

$$Q(X) = \tilde{Q}(\Phi(X))$$

pour tout  $X \in \mathbb{R}^3$ . Il reste à montrer que  $\Phi$  est un isomorphisme. Si on note  $\mathcal{B}$  la base canonique de  $\mathbb{R}^3$ , alors

$$M_{\mathcal{B}\mathcal{B}}(\Phi) = \begin{pmatrix} 1 & -3 & -1 \\ 0 & 1 & \frac{1}{4} \\ 0 & 0 & 1 \end{pmatrix}$$

Cette matrice est triangulaire supérieure. Son déterminant vaut (cf. cours précédents)  $\det(M) = 1$ . Donc  $M$  est inversible ( $1 \neq 0$ ) et donc  $\Phi$  est un isomorphisme de  $\mathbb{R}^3$ .

(4) Comme  $\Phi$  est un isomorphisme,  $Q$  et  $\tilde{Q}$  sont équivalentes. Par suite elles ont même signature. La signature de  $\tilde{Q}$  (cf. la remarque ci-dessus) est  $(1, 2)$ . La signature de  $Q$  est donc aussi  $(1, 2)$ .

(5) Les vecteurs isotropes de  $\tilde{Q}$  sont facilement caractérisables. Il s'agit du cône  $\mathcal{C}$  de  $\mathbb{R}^3$  donné par

$$\mathcal{C} = \left\{ X \text{ de coordonnées } x_1, x_2, x_3 \text{ dans } \hat{\mathcal{B}} \text{ tels que } 8x_2^2 + \frac{11}{2}x_3^2 = x_1^2 \right\}$$

Si on note  $\mathcal{I}$  l'ensemble des vecteurs isotropes de  $Q$ , puisque  $Q = \tilde{Q} \circ \Phi$ , on voit que

$$Q(X) = 0 \Leftrightarrow \Phi(X) \in \mathcal{C}$$

et on récupère donc, puisque  $\Phi$  est un isomorphisme, que

$$\mathcal{I} = \Phi^{-1}(\mathcal{C})$$

Cette relation caractérise les vecteurs isotropes de  $Q$ .

**Exercice :** Soit  $Q : \mathbb{R}^4 \rightarrow \mathbb{R}$  l'application définie par

$$Q(x_1, x_2, x_3, x_4) = x_1^2 + 2x_2^2 + x_3^2 + 2x_4^2 + 2x_1x_2 - 2x_1x_3 + 4x_1x_4 + 2x_2x_4 - 8x_3x_4.$$

Déterminer la signature de  $Q$ .

**Solution:** On utilisait les développements de  $(a + b + c)^2$  et  $(a + b)^2$  en dimension 3. On va ici utiliser les développements de  $(a + b + c + d)^2$ ,  $(a + b + c)^2$  et  $(a + b)^2$ . Pour  $X = (x_1, x_2, x_3, x_4)$  on écrit

$$\begin{aligned} Q(X) &= \mathbf{x}_1^2 + 2x_2^2 + x_3^2 + 2x_4^2 + 2\mathbf{x}_1\mathbf{x}_2 - 2\mathbf{x}_1\mathbf{x}_3 + 4\mathbf{x}_1\mathbf{x}_4 + 2x_2x_4 - 8x_3x_4 \\ &= (\mathbf{x}_1 + \mathbf{x}_2 - \mathbf{x}_3 + 2\mathbf{x}_4)^2 - \mathbf{x}_2^2 - \mathbf{x}_3^2 - 4\mathbf{x}_4^2 + 2\mathbf{x}_2\mathbf{x}_3 - 4\mathbf{x}_2\mathbf{x}_4 + 4\mathbf{x}_3\mathbf{x}_4 \\ &\quad + 2x_2^2 + x_3^2 + 2x_4^2 + 2x_2x_4 - 8x_3x_4 \\ &= (x_1 + x_2 - x_3 + 2x_4)^2 + \mathbf{x}_2^2 - 2x_4^2 + 2\mathbf{x}_2\mathbf{x}_3 - 2\mathbf{x}_2\mathbf{x}_4 - 4x_3x_4 \\ &= (x_1 + x_2 - x_3 + 2x_4)^2 + (\mathbf{x}_2 + \mathbf{x}_3 - \mathbf{x}_4)^2 - \mathbf{x}_3^2 - \mathbf{x}_4^2 + 2\mathbf{x}_3\mathbf{x}_4 \\ &\quad - 2x_4^2 - 4x_3x_4 \\ &= (x_1 + x_2 - x_3 + 2x_4)^2 + (x_2 + x_3 - x_4)^2 - \mathbf{x}_3^2 - 3x_4^2 - 2\mathbf{x}_3\mathbf{x}_4 \\ &= (x_1 + x_2 - x_3 + 2x_4)^2 + (x_2 + x_3 - x_4)^2 - (x_3 + x_4)^2 - 2x_4^2. \end{aligned}$$

Soit  $\tilde{Q}$  la forme quadratique

$$\tilde{Q}(X) = x_1^2 + x_2^2 - x_3^2 - 2x_4^2$$

et soit  $\Phi : \mathbb{R}^4 \rightarrow \mathbb{R}^4$  l'endomorphisme de  $\mathbb{R}^4$  donné par

$$\Phi(x_1, x_2, x_3, x_4) = (x_1 + x_2 - x_3 + 2x_4, x_2 + x_3 - x_4, x_3 + x_4, x_4).$$

Alors  $Q(X) = \tilde{Q}(\Phi(X))$  pour tout  $X \in \mathbb{R}^4$ . La matrice de  $\Phi$  dans la base canonique de  $\mathbb{R}^4$  est triangulaire supérieure à diagonale composée de 1. Son déterminant vaut donc  $1 \neq 0$  et  $\Phi$  est ainsi un isomorphisme. Donc  $Q$  et  $\tilde{Q}$  sont équivalentes. Elles ont donc même signature. La signature de  $\tilde{Q}$  vaut  $(2, 2)$ . La signature de  $Q$  vaut donc elle aussi  $(2, 2)$ .

## CHAPITRE 2

### ESPACES PRÉHILBERTIENS

Un espace préhilbertien est un espace vectoriel muni d'un produit scalaire. Les premières définitions sont données dans la section qui suit.

#### 10. NORMES ET PRODUITS SCALAIRES.

Soit  $E$  un  $\mathbb{R}$ -espace vectoriel. On distingue deux structures sur  $E$ : celle donnée par une norme et celle donnée par un produit scalaire:

(1)  $N : E \rightarrow \mathbb{R}^+$  est une norme sur  $E$  si:

$$(N1) \quad \forall x \in E, N(x) = 0 \Leftrightarrow x = 0,$$

$$(N2) \quad \forall \lambda \in \mathbb{R}, \forall x \in E, N(\lambda x) = |\lambda|N(x),$$

$$(N3) \quad \forall x, y \in E, N(x + y) \leq N(x) + N(y).$$

(2)  $B : E \times E \rightarrow \mathbb{R}$  est un produit scalaire sur  $E$  si:

$$(S1) \quad B \text{ est bilinéaire symétrique,}$$

$$(S2) \quad \forall x \in E, B(x, x) \geq 0,$$

$$(S3) \quad \forall x \in E, B(x, x) = 0 \Leftrightarrow x = 0.$$

On se réfère à la propriété (N3) sous les termes d'inégalité triangulaire. Une forme bilinéaire symétrique  $B$  qui vérifie (S2) et (S3) est dite définie positive.

**Définition 10.1.** *Un espace préhilbertien est un espace vectoriel  $E$  muni d'un produit scalaire.*

Dans la pratique, ce que l'on fera dans la suite, on note plutôt  $\|\cdot\|$  pour une norme, et  $\langle \cdot, \cdot \rangle$  pour un produit scalaire. Donc  $N(x) = \|x\|$  et  $B(x, y) = \langle x, y \rangle$ . A titre d'exemples,

$$\|(x_1, \dots, x_n)\|_1 = \sqrt{\sum_{i=1}^n x_i^2} \quad \text{et} \quad \|(x_1, \dots, x_n)\|_2 = \max_{i=1, \dots, n} |x_i|$$

sont des normes sur  $\mathbb{R}^n$ , tandis que

$$\langle (x_1, \dots, x_n), (y_1, \dots, y_n) \rangle = \sum_{i=1}^n x_i y_i$$

est un produit scalaire sur  $\mathbb{R}^n$ . On remarque que

$$\|(x_1, \dots, x_n)\|_1 = \sqrt{\langle (x_1, \dots, x_n), (x_1, \dots, x_n) \rangle}$$

de sorte que  $\|\cdot\|_1^2$  est la forme quadratique associée à la forme bilinéaire  $\langle \cdot, \cdot \rangle$ .

**Lemme 10.1.** *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie  $n$  et soit  $B$  une forme bilinéaire symétrique sur  $E$ . Si  $Q$  est la forme quadratique associée, définie par  $Q(x) = B(x, x)$  pour tout  $x$ , alors  $B$  est un produit scalaire si et seulement si  $Q$  est de signature  $(n, 0)$ .*

*Démonstration.* Clairement  $Q$  est positive ou nulle sans vecteurs isotropes autres que 0 si et seulement si  $Q$  est de signature  $(n, 0)$ , car toute autre décomposition de Cayley-Hamilton va produire soit des vecteurs isotropes non nuls (si le rang est strictement inférieur à  $n$ ), soit des valeurs strictement négatives de  $Q$  si dans la signature  $(p, q)$  on a  $q \geq 1$ . Donc  $B$  est un produit scalaire si et seulement si  $p + q = n$  et  $q = 0$ , soit si et seulement si  $Q$  est de signature  $(n, 0)$ .  $\square$

Les structures de produits scalaires sont plus “fortes” que les structures de normes (produit scalaire  $\Rightarrow$  norme), comme on le verra juste après l’exercice qui suit.

**Exercice:** On note  $B$  la forme bilinéaire sur  $\mathbb{R}^2$  définie par

$$B(x, y) = 2x_1y_1 + 3x_2y_2 - 2(x_1y_2 + x_2y_1)$$

pour tous  $x = (x_1, x_2)$  et  $y = (y_1, y_2)$ . Montrer que  $B$  est un produit scalaire sur  $E$ .

**Solution:** On vérifie facilement que  $B$  est bilinéaire symétrique. Reste à montrer que  $B$  est définie positive. On a

$$\begin{aligned} B(x, x) &= 2x_1^2 + 3x_2^2 - 4x_1x_2 \\ &= 2(x_1 - x_2)^2 + x_2^2. \end{aligned}$$

Il s’ensuit que  $B(x, x) \geq 0$  pour tout  $x$  et que  $B(x, x) = 0$  si et seulement si  $x_1 = x_2$  et  $x_2 = 0$ , et donc si et seulement si  $x_1 = x_2 = 0$ . Donc  $B$  est bien définie positive. Par suite  $B$  est un produit scalaire.  $\square$

A tout produit scalaire est associé une norme qui est la racine carrée de la forme quadratique associé à la forme bilinéaire symétrique produit scalaire. Le fait qu’un produit scalaire soit défini positif garantit qu’il n’y a pas de vecteurs isotropes autres que le vecteur trivial nul.

**Théorème 10.1.** *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel. A tout produit scalaire  $\langle \cdot, \cdot \rangle$  est associé une norme  $\| \cdot \|$  en posant  $\|x\| = \sqrt{\langle x, x \rangle}$  pour tout  $x$ .*

La preuve de ce théorème passe par la preuve de l’inégalité dite de Cauchy-Schwarz.

**Théorème 10.2** (Inégalité de Cauchy-Schwarz). *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel. Pour tous  $x, y \in E$ ,*

$$|\langle x, y \rangle| \leq \|x\| \times \|y\| ,$$

où  $\| \cdot \|$  est définie par  $\|x\| = \sqrt{\langle x, x \rangle}$ .

A titre de remarque, vous avez déjà vue cette inégalité en théorie de l’intégration. L’espace  $E = C^0([0, 1], \mathbb{R})$  est un  $\mathbb{R}$ -espace vectoriel (même si de dimension infinie...). La forme bilinéaire symétrique

$$\langle f, g \rangle = \int_0^1 f(t)g(t)dt$$

est un produit scalaire sur  $E$ . Sa norme associée est donnée par

$$\|f\| = \sqrt{\int_0^1 f^2(t)dt} ,$$

et on a bien que  $\forall f, g \in E$ ,

$$\left| \int_0^1 f(t)g(t)dt \right| \leq \sqrt{\int_0^1 f^2(t)dt} \sqrt{\int_0^1 g^2(t)dt} ,$$

ce qui illustre l'inégalité de Cauchy-Schwarz dans ce cas précis.

*Démonstration de l'inégalité de Cauchy-Schwarz.* On écrit que pour tout  $\lambda \in \mathbb{R}$ , et tous  $x, y \in E$ ,

$$\langle x + \lambda y, x + \lambda y \rangle \geq 0 .$$

Or

$$\begin{aligned} \langle x + \lambda y, x + \lambda y \rangle &= \langle x + \lambda y, x \rangle + \lambda \langle x + \lambda y, y \rangle \\ &= \langle x, x \rangle + 2\lambda \langle x, y \rangle + \lambda^2 \langle y, y \rangle . \end{aligned}$$

Par suite, dire que pour tout  $\lambda \in \mathbb{R}$ ,  $\langle x + \lambda y, x + \lambda y \rangle \geq 0$ , entraîne que la fonction polynôme du second degré en  $\lambda$  ci-dessus ne peut avoir 2 racines distinctes, et donc entraîne pour le discriminant  $\Delta$  du polynôme du second degré en question que

$$\Delta = 4\langle x, y \rangle^2 - 4\langle x, x \rangle \langle y, y \rangle \leq 0 .$$

Il s'ensuit que

$$\langle x, y \rangle^2 \leq \langle x, x \rangle \langle y, y \rangle ,$$

et on retrouve l'inégalité de Cauchy-Schwarz. D'où le théorème.  $\square$

On démontre maintenant le Théorème 10.1.

*Démonstration du Théorème 10.1.* Etant donné  $x \in E$ , on pose

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

Il nous faut montrer que  $\|\cdot\|$  vérifie les points (N1), (N2), et (N3) de la définition des normes. Le point (N1) est automatiquement vérifié dans la mesure où

$$\langle x, x \rangle = 0 \Leftrightarrow x = 0 .$$

Le point (N2) est aussi vérifié puisque

$$\begin{aligned} \langle \lambda x, \lambda x \rangle &= \lambda \langle \lambda x, x \rangle \\ &= \lambda^2 \langle x, x \rangle . \end{aligned}$$

Reste donc à montrer l'inégalité triangulaire, à savoir que pour tous  $x, y \in E$ ,

$$\|x + y\| \leq \|x\| + \|y\| .$$

On écrit ici que

$$\begin{aligned} \|x + y\|^2 &= \langle x + y, x + y \rangle \\ &= \langle x, x \rangle + 2\langle x, y \rangle + \langle y, y \rangle \\ &= \|x\|^2 + \|y\|^2 + 2\langle x, y \rangle \\ &\leq \|x\|^2 + \|y\|^2 + 2\|x\|\|y\| \text{ (d'après Cauchy-Schwarz)} \\ &= (\|x\| + \|y\|)^2 . \end{aligned}$$

Donc  $\forall x, y \in E$ ,  $\|x + y\| \leq \|x\| + \|y\|$ . D'où le résultat. Le théorème est démontré.  $\square$

Plusieurs identités remarquables sont associées à cette notion de produit scalaire.

**Théorème 10.3** (Identités remarquables). *Soit  $\langle \cdot, \cdot \rangle$  un produit scalaire sur un espace vectoriel réel  $E$ , et soit  $\| \cdot \|$  la norme qui lui est associée. Alors,*

- (1)  $\forall x, y \in E, \|x + y\|^2 = \|x\|^2 + \|y\|^2 + 2\langle x, y \rangle,$
- (2)  $\forall x, y \in E, \|x + y\|^2 + \|x - y\|^2 = 2(\|x\|^2 + \|y\|^2),$
- (3)  $\forall x, y \in E, \|x + y\|^2 - \|x - y\|^2 = 4\langle x, y \rangle .$

*En particulier, et d'après (1) ou (3), le produit scalaire est entièrement déterminé par la norme.*

*Démonstration.* On a

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

d'où la première relation. De cette relation on déduit maintenant que

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

Par suite, en couplant les deux relations que l'on vient de démontrer,

$$\|x + y\|^2 + \|x - y\|^2 = 2(\|x\|^2 + \|y\|^2) .$$

De même,

$$\|x + y\|^2 - \|x - y\|^2 = 4\langle x, y \rangle .$$

D'où les relations (2) et (3), et le résultat.  $\square$

De (3), on tire que si  $\| \cdot \|$  est une norme sur  $E$ , alors  $\| \cdot \|$  provient d'un produit scalaire si et seulement si l'application

$$(x, y) \longrightarrow \frac{1}{4} (\|x + y\|^2 - \|x - y\|^2)$$

est un produit scalaire sur  $E$ . On constate assez facilement ici que la seule chose qu'il y ait à montrer est que cette application est bilinéaire. Si

$$B : E \times E \rightarrow \mathbb{R}$$

$$(x, y) \longrightarrow \frac{1}{4} (\|x + y\|^2 - \|x - y\|^2)$$

est cette application, il est en effet clair que  $B(x, x) \geq 0$  pour tout  $x$ , et que  $B(x, x) = 0$  si et seulement si  $x = 0$ , puisque pour tout  $x$ ,  $B(x, x) = \|x\|^2$ .

**Exercice:** Soit  $\mathcal{M}_n(\mathbb{R})$  l'espace vectoriel des matrices réelles carrées  $n \times n$ . Montrer que

$$\langle M, N \rangle = \text{tr}({}^t M N)$$

est un produit scalaire sur  $\mathcal{M}_n(\mathbb{R})$  où  $\text{tr}$  représente la trace.

**Solution:** Notons  $M = (a_{ij})$ . Alors  ${}^t M = (a_{ji})$  et

$$M = \begin{pmatrix} a_{11} & \dots & a_{1n} \\ \vdots & & \vdots \\ a_{n1} & \dots & a_{nn} \end{pmatrix}, \quad {}^t M = \begin{pmatrix} a_{11} & \dots & a_{n1} \\ \vdots & & \vdots \\ a_{1n} & \dots & a_{nn} \end{pmatrix} .$$

Si  ${}^t M M = (b_{ij})$  alors, pour tous  $i, j$ ,

$$b_{ij} = \sum_{\alpha=1}^n a_{\alpha i} a_{\alpha j} .$$

Il est clair que  $\langle \cdot, \cdot \rangle$  est une forme bilinéaire. Elle est symétrique car  $\text{tr} M = \text{tr}^t M$  pour toute matrice carrée  $M$ , tandis que

$${}^t({}^t NM) = {}^t MN$$

pour toutes matrices carrées  $M, N$  de  $\mathcal{M}_n(\mathbb{R})$ . La forme est enfin définie positive car (cf. ci-dessus)

$$\begin{aligned} \langle M, M \rangle &= \sum_{i=1}^n \sum_{\alpha=1}^n a_{\alpha i}^2 \\ &= \sum_{i,j=1}^n a_{ij}^2 \end{aligned}$$

de sorte que  $\langle M, M \rangle \geq 0$  pour toute matrice carrée  $M$  et de sorte que  $\langle M, M \rangle = 0$  si et seulement si  $M = 0$ . Donc  $\langle \cdot, \cdot \rangle$  est bien un produit scalaire sur  $\mathcal{M}_n(\mathbb{R})$ .  $\square$

## 11. ORTHOGONALITÉ

On aborde maintenant les notions d'orthogonalité et d'orthonormalité déjà rencontrées pour les formes bilinéaires quelconques, mais cette fois-ci donc dans le cadre des produits scalaires pour lesquels il n'y a plus de vecteur isotrope non trivial. Soit  $E$  est un  $\mathbb{R}$ -espace vectoriel muni d'un produit scalaire  $\langle \cdot, \cdot \rangle$ . On dit que deux vecteurs  $x$  et  $y$  de  $E$  sont orthogonaux (sous entendu pour  $\langle \cdot, \cdot \rangle$ ) si  $\langle x, y \rangle = 0$ . Si maintenant  $X$  est un sous ensemble de  $E$ , on définit l'orthogonal de  $X$  (toujours sous entendu pour  $\langle \cdot, \cdot \rangle$  et que l'on note  $X^\perp$ ) par

$$X^\perp = \left\{ y \in E \mid \forall x \in X, \langle x, y \rangle = 0 \right\}.$$

En d'autres termes,  $X^\perp$  est l'ensemble des vecteurs de  $E$  qui sont orthogonaux à tous les vecteurs de  $X$ . On vérifie facilement que pour tout sous ensemble  $X$  de  $E$ ,  $X^\perp$  est un sous espace vectoriel de  $E$ . De même, on vérifie facilement que:

$$(P1) \quad \forall X \subset Y, Y^\perp \subset X^\perp,$$

$$(P2) \quad \forall X, Y, (X \cup Y)^\perp = X^\perp \cap Y^\perp,$$

$$(P3) \quad \forall X, X \subset (X^\perp)^\perp.$$

Le théorème suivant a lieu.

**Théorème 11.1** (Théorème de Pythagore). *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel muni d'un produit scalaire  $\langle \cdot, \cdot \rangle$ , et soit  $\| \cdot \|$  la norme associée au produit scalaire  $\langle \cdot, \cdot \rangle$ . Deux vecteurs  $x$  et  $y$  de  $E$  sont alors orthogonaux si et seulement si  $\|x + y\|^2 = \|x\|^2 + \|y\|^2$ .*

*Démonstration.* Pour tous  $x, y \in E$ ,

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

où  $\| \cdot \|$  est la norme associée au produit scalaire  $\langle \cdot, \cdot \rangle$ . Il s'ensuit clairement que  $\langle x, y \rangle = 0$  si et seulement si  $\|x + y\|^2 = \|x\|^2 + \|y\|^2$ . D'où le résultat.  $\square$

## 12. RETOUR SUR LA SIGNATURE

On revient dans cette section sur les considérations du premier chapitre et plus précisément sur la signature des formes quadratiques. Le premier résultat que l'on démontre est le suivant.

**Lemme 12.1.** *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie,  $F$  un sous espace vectoriel de  $E$  et  $B$  une forme bilinéaire symétrique définie sur  $E \times E$ . On suppose que la restriction  $B|_F$  de  $B$  à  $F$  est définie positive. On note  $F^{\perp_B}$  l'orthogonal de  $F$  pour  $B$ . Alors  $E = F \oplus F^{\perp_B}$ .*

*Démonstration.* On commence par montrer que  $F$  et  $F^{\perp_B}$  sont d'intersection réduite au vecteur nul. Soit  $x \in F \cap F^{\perp_B}$ . Alors  $B(x, x) = 0$ . Or  $x \in F$  et  $B|_F$  est supposée définie positive. Donc  $x = 0$ . D'où la relation  $F \cap F^{\perp_B} = \{0\}$ . Soit maintenant  $x \in E$ . On note  $(e_1, \dots, e_k)$  une base orthogonale de  $F$  pour  $B|_F$  (qui existe d'après le Théorème 8.1). On pose

$$x_i = \frac{B(x, e_i)}{B(e_i, e_i)}$$

pour tout  $i = 1, \dots, k$  (ce qui fait sens puisque  $B|_F$  est définie positive de sorte que  $B(e_i, e_i) > 0$  pour tout  $i = 1, \dots, k$ ). Soit

$$\bar{x}_1 = \sum_{i=1}^k x_i e_i \quad \text{et} \quad \bar{x}_2 = x - \sum_{i=1}^k x_i e_i .$$

On a  $x = \bar{x}_1 + \bar{x}_2$  et  $\bar{x}_1 \in F$ . De plus, puisque  $(e_1, \dots, e_k)$  est orthogonale pour la restriction  $B|_F$  de  $B$  à  $F$ , donc en fait pour  $B$ ,

$$B(\bar{x}_2, e_j) = B(x, e_j) - x_j B(e_j, e_j) = 0$$

pour tout  $j = 1, \dots, k$ . Donc  $\bar{x}_2 \in F^{\perp_B}$ . D'où le lemme.  $\square$

On démontre ici une caractérisation dimensionnelle de la signature. Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie et  $B$  une forme bilinéaire symétrique définie sur  $E \times E$ . On note dans ce qui suit

$$E_B^+ = \{F \text{ sev de } E / B|_F \text{ est définie positive}\}$$

et

$$E_B^- = \{F \text{ sev de } E / B|_F \text{ est définie négative}\} ,$$

où  $B|_F$  est la restriction de  $B$  à  $F$  et où l'on dit d'une forme bilinéaire symétrique  $\tilde{B}$  sur un espace vectoriel  $F$  qu'elle est définie négative si  $\hat{B} = -\tilde{B}$  est définie positive. Le Théorème 12.1 est parfois utilisé comme définition de la signature d'une forme quadratique.

**Théorème 12.1.** *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie et soit  $Q$  une forme quadratique définie sur  $E$ . Soit  $B$  la forme polaire de  $Q$ . On définit*

$$\begin{cases} p_B^+ = \sup_{F \in E_B^+} (\dim F) \\ p_B^- = \sup_{F \in E_B^-} (\dim F) . \end{cases}$$

*Le couple  $(p_B^+, p_B^-)$  n'est alors rien d'autre que la signature de  $Q$ .*

*Démonstration.* Notons  $(p, q)$  la signature de  $Q$ . Du théorème de Sylvester on tire facilement que  $p \leq p_B^+$  et que  $q \leq p_B^-$ . Supposons que  $p < p_B^+$ . Alors  $p_B^+ \geq p+1$  et il existe ainsi un  $F \in E_B^+$  qui est de dimension  $p+1$ . Du Lemme 12.1 on tire que

$$E = F \oplus F^{\perp_B}.$$

Soit  $(e_1, \dots, e_{p+1})$  une base orthogonale de  $F$ . Comme  $F \in E_B^+$ , les  $B(e_i, e_i)$  sont tous strictement positifs. Quitte à renormaliser (donc à remplacer les  $e_i$  par  $e_i/\sqrt{B(e_i, e_i)}$ ) on peut supposer que  $B(e_i, e_i) = 1$  pour tout  $i = 1, \dots, p+1$ . Soit maintenant  $(\tilde{e}_{p+2}, \dots, \tilde{e}_n)$  une base de  $F^{\perp_B}$ . En renormalisant comme dans la preuve de Sylvester, et quitte à réordonner les  $\tilde{e}_i$ , il existe  $\tilde{p} \geq 0$  et  $\tilde{q} \geq 0$  pour lesquels, dans  $F^{\perp_B}$ ,

$$Q(x) = \sum_{i=p+2}^{p+1+\tilde{p}} x_i^2 - \sum_{i=p+2+\tilde{p}}^{p+1+\tilde{p}+\tilde{q}} x_i^2,$$

où les  $x_i$  sont les coordonnées de  $x$  dans la base  $(\tilde{e}_{p+2}, \dots, \tilde{e}_n)$ . Pour  $x \in E$  on a alors

$$Q(x) = \sum_{i=1}^{p+1+\tilde{p}} x_i^2 - \sum_{i=p+2+\tilde{p}}^{p+1+\tilde{p}+\tilde{q}} x_i^2,$$

où les  $x_i$  sont les coordonnées de  $x$  dans la base de  $E$  constituée des  $e_1, \dots, e_{p+1}$  suivis des  $\tilde{e}_{p+2}, \dots, \tilde{e}_n$ . Le théorème de Sylvester 9.1 impose alors que  $p+1+\tilde{p} = p$ , ce qui est impossible. Donc  $p = p_B^+$ . Soit maintenant  $\tilde{B} = -B$ . La signature de  $\tilde{B}$  vaut alors  $(q, p)$ . Donc  $q = p_B^+$  en raison de ce qui vient d'être démontré. Or  $E_B^+ = E_B^-$ . Donc on a aussi  $q = p_B^-$ . Le théorème est démontré.  $\square$

### 13. LE CRITÈRE DE SYLVESTER

On démontre ici un critère pour décider si oui ou non une forme bilinéaire symétrique est un produit scalaire. On commence avec la définition des mineurs principaux d'une matrice.

**Définition 13.1.** Soit  $M \in \mathcal{M}_n(\mathbb{R})$  une matrice réelle  $n \times n$ . On note  $M = (a_{ij})_{1 \leq i, j \leq n}$  et pour  $k = 1, \dots, n$ , on note  $M_k = (a_{ij})_{1 \leq i, j \leq k}$  les sous matrices principales de  $M$ . On appelle alors mineurs principaux de  $M$  les déterminants  $\Delta_k = \det(M_k)$  des sous matrices principales  $M_k$  pour  $k = 1, \dots, n$ .

Le critère de Sylvester qui permet de déterminer si oui ou non une forme bilinéaire symétrique est un produit scalaire est donné par le théorème suivant.

**Théorème 13.1** (Le critère de Sylvester). Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie  $n$  et  $B$  une forme bilinéaire symétrique définie sur  $E \times E$ . Soit  $\mathcal{B}$  une base de  $E$  et  $M = M_{\mathcal{B}}(B)$  la matrice de représentation de  $B$  dans  $\mathcal{B}$ . Alors  $B$  est un produit scalaire sur  $E$  si et seulement si les  $n$  mineurs principaux  $\Delta_k$  de  $M$ ,  $k = 1, \dots, n$ , sont tous strictement positifs.

*Démonstration.* On sait déjà que  $B$  est symétrique. Supposons que  $B$  est un produit scalaire, donc que  $B$  est définie positive. Notons  $\mathcal{B} = (e_1, \dots, e_n)$ . Soit  $F_k = \text{Vect}(e_1, \dots, e_k)$ . La restriction de  $B$  à  $F_k$  est un produit scalaire sur  $F_k$ . Il existe donc  $\tilde{\mathcal{B}}_k = (\tilde{e}_1^k, \dots, \tilde{e}_k^k)$  une base  $B$ -orthogonale de  $F_k$ . Si  $\lambda_i = B(\tilde{e}_i^k, \tilde{e}_i^k)$ , alors  $\lambda_i > 0$  pour tout  $i = 1, \dots, k$ . La matrice de représentation dans la base  $(e_1, \dots, e_k)$  de la restriction de  $B$  à  $F_k$  n'est rien d'autre que la sous matrice principale  $M_k$ .

Par ailleurs, si  $P_k$  est la matrice de passage de  $(e_1, \dots, e_k)$  à  $(\tilde{e}_1^k, \dots, \tilde{e}_k^k)$ , alors, par changement de bases,

$$M_{\tilde{B}_k}(B) = {}^t P_k M_k P_k .$$

Comme  $M_{\tilde{B}_k}(B)$  est la matrice diagonale des  $\lambda_i$ , en passant au déterminant on voit que

$$\det(P_k)^2 \det(M_k) = \prod_{i=1}^k \lambda_i$$

et donc  $\det(M_k) > 0$ . Comme  $k$  est ici quelconque, les mineurs principaux de  $M$  sont tous strictement positifs.

Pour démontrer la réciproque on procède par récurrence sur la dimension  $n$  de  $E$ . Notre hypothèse de récurrence  $\mathcal{H}_n$  est donc que pour tout  $\mathbb{R}$ -espace vectoriel  $E$  de dimension  $n$  et toute forme bilinéaire symétrique  $B$  définie sur  $E \times E$ , si les  $n$  mineurs principaux  $\Delta_k$  de la matrice de représentation  $M = M_{\mathcal{B}}(B)$  de  $B$  dans une base  $\mathcal{B}$  de  $E$  sont tous strictement positifs, alors  $B$  est un produit scalaire sur  $E$ . Le cas  $n = 1$  est immédiat. On aborde maintenant la question de l'hérédité. On suppose que notre hypothèse de récurrence est vraie à l'ordre  $n$ . Soit alors  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension  $n + 1$ ,  $B$  une forme bilinéaire symétrique définie sur  $E \times E$  et  $\mathcal{B}$  une base de  $E$  pour laquelle les mineurs de  $M = M_{\mathcal{B}}(B)$  sont tous strictement positifs. On écrit que  $\mathcal{B} = (e_1, \dots, e_{n+1})$  et on note  $E_n = \text{Vect}(e_1, \dots, e_n)$ . Si  $\hat{M}$  est la matrice de représentation dans  $(e_1, \dots, e_n)$  de la restriction de  $B$  à  $E_n$  alors  $\hat{M} = M_n$ . Les mineurs  $\Delta_k$  de  $M$  et  $M_n$  sont les mêmes pour  $k = 1, \dots, n$ . Par hypothèse de récurrence, la restriction de  $B$  à  $E_n$  est donc définie positive. En vertu du Lemme 12.1 on a donc que

$$E = E_n \oplus E_n^{\perp B} .$$

En particulier,  $\dim E_n^{\perp B} = 1$ . Soit  $(\tilde{e}_1, \dots, \tilde{e}_n)$  une base  $B$ -orthogonale de  $E_n$ . Quitte à normaliser on peut supposer que  $B(\tilde{e}_i, \tilde{e}_i) = 1$  pour tout  $i = 1, \dots, n$ . Soit  $\tilde{e}_{n+1} \neq 0$  un vecteur directeur de  $E_n^{\perp B}$ . Alors  $\tilde{\mathcal{B}} = (\tilde{e}_1, \dots, \tilde{e}_n, \tilde{e}_{n+1})$  est une base de  $E$ . On a alors que

$$M_{\tilde{\mathcal{B}}}(B) = \begin{pmatrix} I_n & 0_c \\ 0_\ell & \Lambda \end{pmatrix} , \quad (13.1)$$

où  $I_n$  est la matrice identité  $n \times n$ ,  $\Lambda = B(\tilde{e}_{n+1}, \tilde{e}_{n+1})$ ,  $0_c$  est une colonne de  $n$  zéros et  $0_\ell$  est une ligne de  $n$  zéros. Si  $P$  est la matrice de passage de  $\mathcal{B}$  à  $\tilde{\mathcal{B}}$  et si  $\tilde{M} = M_{\tilde{\mathcal{B}}}(B)$ , alors, par changement de bases,

$$\tilde{M} = {}^t P M P$$

et donc

$$\Lambda = \det(P)^2 \Delta_{n+1} .$$

On en déduit que  $\Lambda > 0$  puisque  $\det M = \Delta_{n+1} > 0$ . Clairement il suit de (13.1) que  $B$  est définie positive puisque, dans  $\tilde{\mathcal{B}}$ ,

$$B(x, x) = \sum_{i=1}^n \tilde{x}_i^2 + \Lambda \tilde{x}_{n+1}^2 ,$$

où les  $\tilde{x}_i$  sont les coordonnées de  $x$  dans  $\tilde{\mathcal{B}}$ . Donc  $B$  est un produit scalaire sur  $E$ . Par récurrence on a montré notre réciproque. Le théorème est démontré.  $\square$

## 14. BASES ORTHONORMÉES

On traite de la notion de base orthonormale et du procédé d'orthonormalisation de Gram-Schmidt.

**Définition 14.1.** Soit  $E$  un  $\mathbb{R}$ -espace vectoriel muni d'un produit scalaire  $\langle \cdot, \cdot \rangle$ , et soit  $\|\cdot\|$  la norme associée au produit scalaire  $\langle \cdot, \cdot \rangle$ . On dit qu'une famille  $(e_1, \dots, e_n)$  de vecteurs de  $E$  est:

- (i) orthogonale si pour tous  $i \neq j \in \{1, \dots, n\}$ ,  $\langle e_i, e_j \rangle = 0$ ,
- (ii) orthonormale si pour tous  $i, j \in \{1, \dots, n\}$ ,  $\langle e_i, e_j \rangle = \delta_{ij}$ ,

où les symboles de Kroenecker  $\delta_{ij}$  sont définis par  $\delta_{ij} = 1$  si  $i = j$ , et  $\delta_{ij} = 0$  si  $i \neq j$ . Si  $E$  est de dimension finie, et si  $(e_1, \dots, e_n)$  est une base de  $E$ , on parle alors de base orthogonale et de base orthonormale.

Une famille orthonormale est une famille orthogonale dont les vecteurs sont de normes 1 pour la norme associée au produit scalaire. Etant donnée  $(e_1, \dots, e_n)$  une famille orthogonale de vecteurs non nuls, on obtient une famille orthonormale  $(\tilde{e}_1, \dots, \tilde{e}_n)$  en posant pour tout  $i$ ,

$$\tilde{e}_i = \frac{1}{\|e_i\|} e_i ,$$

où  $\|\cdot\|$  désigne la norme associée au produit scalaire. Lorsqu'il y a un seul produit scalaire, la terminologie est bien choisie. Sinon, on parlera de famille orthogonale pour le produit scalaire  $\langle \cdot, \cdot \rangle$ , et de famille orthonormale pour le produit scalaire  $\langle \cdot, \cdot \rangle$ .

**Proposition 14.1.** Soit  $E$  un  $\mathbb{R}$ -espace vectoriel muni d'un produit scalaire. Toute famille orthogonale de vecteurs non nuls de  $E$  est libre. En particulier, toute famille orthonormale de vecteurs de  $E$  est libre, et si  $E$  est de dimension finie  $n$ , et si  $\mathcal{B}$  est une famille orthonormale de  $E$  composée de  $n$  vecteurs, alors  $\mathcal{B}$  est une base de  $E$ .

*Démonstration.* Soit  $\langle \cdot, \cdot \rangle$  le produit scalaire placé sur  $E$ . Il suffit de démontrer que si  $(e_1, \dots, e_n)$  est une famille orthogonale de  $E$  composée de vecteurs non nuls (i.e.  $e_i \neq 0$  pour tout  $i$ ) alors  $(e_1, \dots, e_n)$  est libre. Une famille orthonormale étant forcément composée de vecteurs non nuls (puisque de normes 1), on en déduira que toute famille orthonormale de vecteurs de  $E$  est libre. Sachant par ailleurs qu'une famille libre de  $n$  vecteurs dans un espace vectoriel de dimension  $n$  est une base, on en déduira de plus que si  $E$  est de dimension finie  $n$ , et si  $\mathcal{B}$  est une famille orthonormale de  $E$  composée de  $n$  vecteurs, alors  $\mathcal{B}$  est une base de  $E$ . Soit donc maintenant  $(e_1, \dots, e_n)$  une famille orthogonale de  $E$  composée de vecteurs non nuls. Supposons que pour des  $\lambda_1, \dots, \lambda_n$  réels,

$$\lambda_1 e_1 + \dots + \lambda_n e_n = 0 .$$

Alors, pour tout  $i_0 = 1, \dots, n$ ,

$$\left\langle \sum_{i=1}^n \lambda_i e_i, e_{i_0} \right\rangle = 0 .$$

Or, la famille étant orthogonale,

$$\begin{aligned} \left\langle \sum_{i=1}^n \lambda_i e_i, e_{i_0} \right\rangle &= \sum_{i=1}^n \lambda_i \langle e_i, e_{i_0} \rangle \\ &= \lambda_{i_0} \|e_{i_0}\|^2, \end{aligned}$$

où  $\|\cdot\|$  est la norme associée au produit scalaire  $\langle \cdot, \cdot \rangle$ . Comme  $e_{i_0} \neq 0$ ,  $\|e_{i_0}\| \neq 0$ , et donc, forcément  $\lambda_{i_0} = 0$ . Puisque  $i_0$  est quelconque dans  $\{1, \dots, n\}$ , on en déduit que  $\lambda_i = 0$  pour tout  $i = 1, \dots, n$ . D'où le fait que la famille  $(e_1, \dots, e_n)$  est libre. La proposition est démontrée.  $\square$

On sait déjà (voir le chapitre précédent en algèbre bilinéaire) que si  $E$  est un  $\mathbb{R}$ -espace vectoriel de dimension finie munie d'un produit scalaire  $\langle \cdot, \cdot \rangle$ , alors  $E$  possède une base orthogonale  $(e_1, \dots, e_n)$  pour  $\langle \cdot, \cdot \rangle$ . Les vecteurs d'une base étant forcément non nuls, on en déduit l'existence d'une base orthonormale  $(\tilde{e}_1, \dots, \tilde{e}_n)$  en posant

$$\tilde{e}_i = \frac{1}{\|e_i\|} e_i,$$

où  $\|\cdot\|$  est la norme associée au produit scalaire  $\langle \cdot, \cdot \rangle$ . Tout espace vectoriel de dimension finie possède donc une base orthonormale. On propose ici une "autre" preuve de ce résultat basée sur ce que l'on appelle le procédé d'orthonormalisation de Gram-Schmidt. Partant d'une base quelconque, on la "transforme" en une base orthonormale...

**Théorème 14.1.** *Tout espace vectoriel de dimension finie muni d'un produit scalaire possède une base orthonormale.*

*Démonstration.* On démontre le théorème avec le procédé d'orthonormalisation de Gram-Schmidt. Soient  $E$  l'espace vectoriel en question et  $\langle \cdot, \cdot \rangle$  le produit scalaire placé sur  $E$ . Soit  $(e_1, \dots, e_n)$  une base quelconque de  $E$ . A partir de  $(e_1, \dots, e_n)$  on construit une base orthonormale  $(\tilde{e}_1, \dots, \tilde{e}_n)$  en appliquant ce que l'on appelle le procédé d'orthonormalisation de Gram-Schmidt. Pour commencer, on pose

$$\tilde{e}_1 = \frac{1}{\|e_1\|} e_1,$$

où  $\|\cdot\|$  désigne la norme associée au produit scalaire  $\langle \cdot, \cdot \rangle$ . On a  $\text{Vect}(\tilde{e}_1) = \text{Vect}(e_1)$ . On construit maintenant le second vecteur. On pose

$$x = e_2 - \langle e_2, \tilde{e}_1 \rangle \tilde{e}_1.$$

On vérifie que  $\langle x, \tilde{e}_1 \rangle = 0$ . On a  $x \neq 0$  car  $e_2$  n'est pas colinéaire à  $e_1$ . On pose alors

$$\tilde{e}_2 = \frac{1}{\|x\|} x.$$

On a ainsi  $\|\tilde{e}_1\| = \|\tilde{e}_2\| = 1$ , et  $\langle \tilde{e}_1, \tilde{e}_2 \rangle = 0$ , de sorte que  $(\tilde{e}_1, \tilde{e}_2)$  est une famille orthonormale. On a aussi  $\text{Vect}(\tilde{e}_1, \tilde{e}_2) = \text{Vect}(e_1, e_2)$ . Si  $n = 2$ , le procédé s'arrête. Sinon,  $n \geq 3$ , et on cherche à déterminer  $\lambda$  et  $\mu$  réels pour que le vecteur

$$x = e_3 + \lambda \tilde{e}_1 + \mu \tilde{e}_2$$

soit orthogonal à  $\tilde{e}_1$  et  $\tilde{e}_2$ . On trouve ici que

$$\langle x, \tilde{e}_1 \rangle = \langle e_3, \tilde{e}_1 \rangle + \lambda \text{ et que } \langle x, \tilde{e}_2 \rangle = \langle e_3, \tilde{e}_2 \rangle + \mu.$$

Le vecteur

$$x = e_3 - \langle e_3, \tilde{e}_1 \rangle \tilde{e}_1 - \langle e_3, \tilde{e}_2 \rangle \tilde{e}_2$$

est alors orthogonal aux vecteurs  $\tilde{e}_1$  et  $\tilde{e}_2$ . Il est non nul car  $e_3 \notin \text{Vect}(e_1, e_2)$ . On pose

$$\tilde{e}_3 = \frac{1}{\|x\|} x ,$$

de sorte que  $(\tilde{e}_1, \tilde{e}_2, \tilde{e}_3)$  est une famille orthonormale. On a encore  $\text{Vect}(\tilde{e}_1, \tilde{e}_2, \tilde{e}_3) = \text{Vect}(e_1, e_2, e_3)$ . Si  $n = 3$ , le procédé s'arrête, on a là une base orthonormale  $(\tilde{e}_1, \tilde{e}_2, \tilde{e}_3)$  de  $E$ . Sinon,  $n \geq 4$ , et on continue jusqu'à obtenir une famille orthonormale composée d'autant de vecteurs que la dimension  $n$  de  $E$ . De façon un peu plus précise, et si les  $(\tilde{e}_1, \dots, \tilde{e}_k)$  sont construits comme ci-dessus, avec  $k < n$ ,  $n$  la dimension de  $E$ , et donc aussi  $\text{Vect}(\tilde{e}_1, \dots, \tilde{e}_k) = \text{Vect}(e_1, \dots, e_k)$ , on cherche des réels  $\lambda_1, \dots, \lambda_k$  pour que le vecteur

$$x = e_{k+1} + \sum_{i=1}^k \lambda_i \tilde{e}_i$$

soit orthogonal aux vecteurs  $\tilde{e}_i$ ,  $i = 1, \dots, k$ . On trouve ici que

$$\lambda_i = -\langle e_{k+1}, \tilde{e}_i \rangle .$$

On sait que  $x \neq 0$  car  $e_{k+1} \notin \text{Vect}(e_1, \dots, e_k)$ . On pose

$$\tilde{e}_{k+1} = \frac{1}{\|x\|} x ,$$

de sorte que  $(\tilde{e}_1, \dots, \tilde{e}_k, \tilde{e}_{k+1})$  est une famille orthonormale. On a là encore que  $\text{Vect}(\tilde{e}_1, \dots, \tilde{e}_{k+1}) = \text{Vect}(e_1, \dots, e_{k+1})$ . L'inclusion  $\subset$  s'obtient par construction et il reste à remarquer que les deux familles étant libres les deux espaces sont de même dimension  $k + 1$  pour obtenir l'égalité. En répétant ce procédé on construit une famille orthonormale composée de  $n$  vecteurs,  $n$  la dimension de  $E$ . Cette famille est forcément une base orthonormale de  $E$ . D'où le résultat.  $\square$

On remarquera que le procédé d'orthonormalisation de Gram-Schmidt assure de la relation  $\text{Vect}(\tilde{e}_1, \dots, \tilde{e}_k) = \text{Vect}(e_1, \dots, e_k)$  pour tout  $k$ . Il existe par ailleurs une formule simple donnant les coordonnées d'un vecteur dans une base orthonormale. C'est l'objet du résultat important suivant.

**Lemme 14.1.** *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie  $n$  muni d'un produit scalaire  $\langle \cdot, \cdot \rangle$ . Soit de plus  $\mathcal{B} = (e_1, \dots, e_n)$  une base orthonormale de  $E$ . Si  $x \in E$  a pour coordonnées  $(x_1, \dots, x_n)$  dans  $\mathcal{B}$ , alors  $x_i = \langle x, e_i \rangle$  pour tout  $i = 1, \dots, n$ .*

*Proof.* Par définition des coordonnées,

$$x = x_1 e_1 + \dots + x_n e_n .$$

Par suite, pour tout  $i$ ,

$$\langle x, e_i \rangle = \sum_{j=1}^n x_j \langle e_j, e_i \rangle ,$$

et puisque  $\langle e_j, e_i \rangle = \delta_{ij}$ , on obtient que  $\langle x, e_i \rangle = x_i$ . Le résultat est démontré.  $\square$

Si  $\mathcal{B} = (e_1, \dots, e_n)$  est une base orthonormale pour un produit scalaire  $\langle \cdot, \cdot \rangle$ , le produit scalaire (et la norme associée) dans cette base prend l'expression du produit scalaire euclidien (et de la norme euclidienne). Avec les notations usuelles,

$$\begin{aligned} \langle x, y \rangle &= \left\langle \sum_{i=1}^n x_i e_i, \sum_{j=1}^n y_j e_j \right\rangle \\ &= \sum_{i=1}^n \sum_{j=1}^n x_i y_j \langle e_i, e_j \rangle \\ &= \sum_{i=1}^n \sum_{j=1}^n x_i y_j \delta_{ij} \\ &= \sum_{i=1}^n x_i y_i . \end{aligned}$$

En particulier, on récupère aussi que  $\|x\| = \sqrt{\sum_{i=1}^n x_i^2}$ , où  $\|\cdot\|$  est la norme associée au produit scalaire  $\langle \cdot, \cdot \rangle$ .

**Exercice:** Soit  $\mathbb{R}_2[X]$  l'espace des polynômes réels de degré au plus 2. On place sur  $\mathbb{R}_2[X]$  la forme bilinéaire symétrique

$$\langle P, Q \rangle = \int_0^1 P(x)Q(x)dx .$$

- (1) Montrer que  $\langle \cdot, \cdot \rangle$  est un produit scalaire sur  $\mathbb{R}_2[X]$ .  
 (2) Trouver une base orthonormale de  $\mathbb{R}_2[X]$ .

**Solution:** (1) Il est clair que  $\langle \cdot, \cdot \rangle$  est bilinéaire symétrique et que  $\langle P, P \rangle \geq 0$  pour tout  $P \in \mathbb{R}_2[X]$ . Si

$$\langle P, P \rangle = \int_0^1 P(x)^2 dx = 0$$

alors  $P(x) = 0$  pour tout  $x \in [0, 1]$  puisqu'un polynôme est une fonction continue. En particulier les coefficients du polynôme sont nuls puisqu'un polynôme de  $\mathbb{R}_2[X]$  dont les coefficients ne sont pas nuls a au plus deux racines. Il en aurait ici une infinité, d'où la nullité des coefficients. En bref,  $\langle \cdot, \cdot \rangle$  est un produit scalaire.

(2)  $(1, X, X^2)$  est une base de  $\mathbb{R}_2[X]$ . On va appliquer Gram-Schmidt à cette base. On a  $\|1\| = 1$ . On pose  $P_1 = 1$  le polynôme constant égal à 1. On pose  $P = X - \lambda P_1$  et on cherche  $\lambda \in \mathbb{R}$  de sorte que  $\langle P_1, P \rangle = 0$ . On a

$$\langle P_1, P \rangle = \int_0^1 (x - \lambda) dx = \frac{1}{2} - \lambda .$$

On veut donc  $\lambda = \frac{1}{2}$ . Et alors

$$\begin{aligned} \|P\|^2 &= \int_0^1 \left(x - \frac{1}{2}\right)^2 dx \\ &= \int_0^1 \left(x^2 - x + \frac{1}{4}\right) dx \\ &= \frac{1}{3} - \frac{1}{2} + \frac{1}{4} \\ &= \frac{1}{12} . \end{aligned}$$

On pose alors

$$P_2 = 2\sqrt{3} \left( X - \frac{1}{2} \right) .$$

On pose ensuite

$$P = X^2 - \lambda P_2 - \mu P_1$$

et on cherche  $\lambda \in \mathbb{R}$  et  $\mu \in \mathbb{R}$  pour que  $\langle P_2, P \rangle = 0$  et  $\langle P_1, P \rangle = 0$ . On a

$$\begin{aligned} \langle P_1, P \rangle &= \int_0^1 (x^2 - \lambda P_2(x) - \mu P_1(x)) dx \\ &= \int_0^1 x^2 dx - \mu \\ &= \frac{1}{3} - \mu \end{aligned}$$

tandis que

$$\begin{aligned} \langle P_2, P \rangle &= \int_0^1 (x^2 - \lambda P_2(x) - \mu P_1(x)) P_2(x) dx \\ &= \int_0^1 x^2 P_2(x) dx - \lambda \\ &= 2\sqrt{3} \int_0^1 \left(x - \frac{1}{2}\right) x^2 dx - \lambda \\ &= 2\sqrt{3} \left(\frac{1}{4} - \frac{1}{6}\right) - \lambda \\ &= \frac{\sqrt{3}}{6} - \lambda . \end{aligned}$$

On veut donc

$$\lambda = \frac{\sqrt{3}}{6} \text{ et } \mu = \frac{1}{3} .$$

On a en particulier calculé

$$\int_0^1 x^2 P_2(x) dx = \frac{\sqrt{3}}{6} .$$

On a aussi calculé

$$\int_0^1 x^2 P_1(x) dx = \frac{1}{3} .$$

On a alors

$$\begin{aligned} \|P\|^2 &= \int_0^1 \left( x^2 - \frac{\sqrt{3}}{6} P_2(x) - \frac{1}{3} P_1(x) \right)^2 dx \\ &= \int_0^1 x^4 dx - \frac{\sqrt{3}}{3} \int_0^1 x^2 P_2(x) dx - \frac{2}{3} \int_0^1 x^2 P_1(x) dx \\ &\quad + \frac{1}{12} \int_0^1 P_2(x)^2 dx + \frac{\sqrt{3}}{9} \int_0^1 P_2(x) P_1(x) dx + \frac{1}{9} \int_0^1 P_1(x)^2 dx \\ &= \frac{1}{5} - \frac{\sqrt{3}}{3} \times \frac{\sqrt{3}}{6} - \frac{2}{3} \times \frac{1}{3} + \frac{1}{12} + 0 + \frac{1}{9} \end{aligned}$$

de sorte que

$$\begin{aligned}\|P\|^2 &= \frac{1}{5} - \frac{1}{12} - \frac{1}{9} \\ &= \frac{9 \times 4 - 3 \times 5 - 4 \times 5}{4 \times 5 \times 9} \\ &= \frac{1}{4 \times 5 \times 9} \\ &= \frac{1}{180}\end{aligned}$$

On pose

$$\begin{aligned}P_3 &= 3\sqrt{20} \left( X^2 - \left( X - \frac{1}{2} \right) - \frac{1}{3} \right) \\ &= 3\sqrt{20} \left( X^2 - X + \frac{1}{6} \right) .\end{aligned}$$

La famille  $(P_1, P_2, P_3)$  est alors une base orthonormée de  $\mathbb{R}_2[X]$ .  $\square$

## 15. ORTHOGONALITÉ ET SOUS ESPACES VECTORIELS

Les idées contenues dans le procédé d'orthogonalisation de Gram-Schmidt s'adaptent facilement pour démontrer le résultat suivant.

**Théorème 15.1.** *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie muni d'un produit scalaire. Toute famille orthonormale de vecteurs de  $E$  se complète en une base orthonormale de  $E$ .*

*Démonstration.* Notons  $n$  la dimension de  $E$ , et  $\langle \cdot, \cdot \rangle$  le produit scalaire placé sur  $E$ . Soit  $(e_1, \dots, e_k)$  une famille orthonormale. Elle est libre (car orthonormale) donc  $k \leq n$ . Si  $k = n$ , c'est une base et il n'y a rien à ajouter. Sinon  $k < n$  et la famille n'est pas génératrice. Du coup, c'est une idée que l'on utilise souvent, il existe un vecteur  $x$  de  $E$  tel que  $x \notin \text{Vect}(e_1, \dots, e_k)$ . En particulier,  $x \neq 0$ . On pose

$$y = x - \sum_{i=1}^k \lambda_i e_i ,$$

où les  $\lambda_i \in \mathbb{R}$  sont choisis de sorte que  $\langle e_i, y \rangle = 0$  pour tout  $i = 1, \dots, k$ . Donc,  $\lambda_i = \langle e_i, x \rangle$  puisque

$$\langle e_i, y \rangle = \langle e_i, x \rangle - \lambda_i .$$

Comme  $x \notin \text{Vect}(e_1, \dots, e_k)$ , on a  $y \neq 0$ . On pose

$$e_{k+1} = \frac{1}{\|y\|} y ,$$

où  $\|\cdot\|$  désigne la norme associée au produit scalaire  $\langle \cdot, \cdot \rangle$ . Alors  $(e_1, \dots, e_{k+1})$  est une famille orthonormale composée de  $k+1$  vecteurs. Dès lors, soit  $k+1 = n$ , et on a en fait une base de  $E$ , soit  $k+1 < n$  et on recommence jusqu'à obtenir (il s'agit d'un processus fini) une famille orthonormale qui soit composée de  $n$  vecteurs, et donc une base orthonormale de  $E$ . D'où le résultat.  $\square$

Une conséquence importante de ce théorème est la suivante.

**Théorème 15.2.** Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie muni d'un produit scalaire. Pour tout sous espace vectoriel  $F$  de  $E$ , on a  $E = F \oplus F^\perp$ .

*Démonstration.* Soit  $\langle \cdot, \cdot \rangle$  le produit scalaire placé sur  $E$ . Il est déjà clair que la somme de  $F$  et  $F^\perp$  est nécessairement directe. En effet,  $F \cap F^\perp = \{0\}$  car si  $x \in F$  et  $x \in F^\perp$ , alors  $\langle x, x \rangle = 0$  de sorte que  $x = 0$ . Par ailleurs,  $F$  est un espace vectoriel de dimension finie. Il possède donc une base orthonormale  $(e_1, \dots, e_k)$ ,  $k$  la dimension de  $F$ . On complète cette base en une base orthonormale  $(e_1, \dots, e_n)$  de  $E$ ,  $n$  la dimension de  $E$ . Alors, clairement,

$$F^\perp = \text{Vect}(e_{k+1}, \dots, e_n) .$$

Les  $e_{k+1}, \dots, e_n$  sont en effet dans  $F^\perp$ . Comme  $(e_{k+1}, \dots, e_n)$  est une famille libre (toute sous famille d'une famille libre étant une famille libre) on en déduit que  $\dim F^\perp \geq n - k$ . Or  $\dim(F \oplus F^\perp) \leq \dim E$  et  $\dim(F \oplus F^\perp) = \dim F + \dim F^\perp$ . Donc  $\dim F^\perp \leq n - k$  puis  $\dim F^\perp = n - k$  et on récupère bien que  $E = F \oplus F^\perp$ . Le théorème est démontré.  $\square$

**Exercice:** Soit  $\mathcal{M}_n(\mathbb{R})$  l'espace vectoriel des matrices réelles carrées  $n \times n$  muni du produit scalaire  $\langle M, N \rangle = \text{tr}({}^t M N)$  vu précédemment à la Section 10. On considère la base canonique de  $\mathcal{M}_n(\mathbb{R})$  constituée des matrices  $A_{ij} = (a(ij)_{kl})$  pour  $i, j = 1, \dots, n$ , où  $a(ij)_{kl} = 1$  si  $(k, l) = (i, j)$ , et  $a(ij)_{kl} = 0$  sinon. Montrer que la base formée des  $A_{ij}$  est une base orthonormale pour  $\langle \cdot, \cdot \rangle$ .

**Solution:** Notons  ${}^t A_{ij} A_{ij} = (b_{kl})$ . On a alors

$$b_{ll} = \sum_{k=1}^n a(ij)_{kl}^2$$

pour tout  $l = 1, \dots, n$ . Par suite,  $b_{ll} = 0$  si  $l \neq j$  et  $b_{jj} = 1$ . Donc,  $\text{tr}({}^t A_{ij} A_{ij}) = 1$  et ainsi,  $\|A_{ij}\| = 1$ . Soient maintenant  $(i, j)$  et  $(\alpha, \beta)$  avec  $(i, j) \neq (\alpha, \beta)$ . Notons  ${}^t A_{ij} A_{\alpha\beta} = (b_{kl})$ . On a alors

$$b_{ll} = \sum_{k=1}^n a(ij)_{kl} a(\alpha\beta)_{kl}$$

pour tout  $l = 1, \dots, n$ . Par suite, pour que  $b_{ll} \neq 0$  il faut qu'il y ait au moins un  $k$  pour lequel  $a(ij)_{kl} = 1$  et  $a(\alpha\beta)_{kl} = 1$ . Or  $a(ij)_{kl} = 1$  si et seulement si  $(k, l) = (i, j)$  et  $a(\alpha\beta)_{kl} = 1$  si et seulement si  $(k, l) = (\alpha, \beta)$ . Ces deux conditions ne peuvent être réalisées en même temps car elles impliquent que  $(i, j) = (k, l) = (\alpha, \beta)$  alors que  $(i, j) \neq (\alpha, \beta)$ . Donc  $b_{ll} = 0$  pour tout  $l$  dans ce cas, et ainsi  $\langle A_{ij}, A_{\alpha\beta} \rangle = 0$  si  $(i, j) \neq (\alpha, \beta)$ . La base canonique de  $\mathcal{M}_n(\mathbb{R})$  est bien une base orthonormale pour le produit scalaire  $\langle M, N \rangle = \text{tr}({}^t M N)$ .

## 16. PROJECTEURS ORTHOGONAUX

On utilise ici la décomposition  $E = F \oplus F^\perp$  du Théorème 15.2 pour définir la notion de projection orthogonale sur  $F$ . La définition qui suit fonctionne pour  $F = \{0\}$  (alors  $F^\perp = E$ ) et pour  $F = E$  (alors  $F^\perp = \{0\}$ ), mais évidemment elle renvoie plutôt à des situations où  $F$  n'est ni réduit au seul vecteur nul ni l'espace  $E$  tout entier.

**Définition 16.1.** Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie muni d'un produit scalaire. Soit  $F$  un sous espace vectoriel de  $E$ . Comme  $E = F \oplus F^\perp$ , tout vecteur  $x$  de  $E$  se décompose de façon unique en  $x = x_1 + x_2$  avec  $x_1 \in F$  et  $x_2 \in F^\perp$ . On définit la projection orthogonale sur  $F$  comme étant l'application  $p_F : E \rightarrow F$  donnée par  $p_F(x) = x_1$ .

On montre que  $p_F$  est linéaire et on caractérise son noyau et son image.

**Proposition 16.1.** Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie muni d'un produit scalaire. Soit  $F$  un sous espace vectoriel de  $E$ . La projection orthogonale  $p_F$  est une application linéaire de  $E$  dans  $F$  qui vérifie que  $p_F \circ p_F = p_F$ . De plus  $\text{Im}(p_F) = F$  et  $\text{Ker}(p_F) = F^\perp$ .

*Démonstration.* Soient  $x, y \in E$  et  $x = x_1 + x_2$ ,  $y = y_1 + y_2$  les décompositions de  $F$  suivant la décomposition  $E = F \oplus F^\perp$ . Soit  $\lambda \in \mathbb{R}$ . Alors

$$x + \lambda y = (x_1 + \lambda y_1) + (x_2 + \lambda y_2)$$

et on a écrit la décomposition de  $x + \lambda y$ . Donc

$$\begin{aligned} p_F(x + \lambda y) &= x_1 + \lambda y_1 \\ &= p_F(x) + \lambda p_F(y) \end{aligned}$$

et  $p_F$  est donc bien linéaire. Si  $x \in F$ ,  $x = x + 0$  est la décomposition de  $x$  suivant la décomposition  $E = F \oplus F^\perp$ . Donc  $p_F(x) = x$  pour tout  $x \in F$ . On en déduit que  $p_F \circ p_F = p_F$  et que  $F \subset \text{Im}(p_F)$ . Puisque  $\text{Im}(p_F) \subset F$  on obtient que  $\text{Im}(p_F) = F$ . Si  $x \in F^\perp$ , alors  $x = 0 + x$  est la décomposition de  $x$  suivant la décomposition  $E = F \oplus F^\perp$ . Donc  $p_F(x) = 0$  pour tout  $x \in F^\perp$ . Soit  $F^\perp \subset \text{Ker}(p_F)$ . Réciproquement, partons de l'écriture  $x = x_1 + x_2$ . Si  $p_F(x) = 0$  alors  $x_1 = 0$  et donc  $x = 0 + x_2$  (et automatiquement donc  $x_2 = x$ ). En particulier,  $x \in F^\perp$ . Donc  $\text{Ker}(p_F) \subset F^\perp$  et, par double inclusion,  $\text{Ker}(p_F) = F^\perp$ .  $\square$

L'expression explicite suivante a lieu. On notera que l'existence d'une base orthonormée  $(e_1, \dots, e_n)$  de  $E$  pour laquelle la sous famille  $(e_1, \dots, e_k)$  soit une base orthonormée de  $F$  ( $n = \dim(E)$ ,  $k = \dim(F)$ ) suit du Théorème 15.1.

**Proposition 16.2.** Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie muni d'un produit scalaire  $\langle \cdot, \cdot \rangle$ . Soit  $F$  un sous espace vectoriel de  $E$ . Soit  $(e_1, \dots, e_n)$  une base orthonormée de  $E$  choisie de sorte que la sous famille  $(e_1, \dots, e_k)$  soit une base orthonormée de  $F$ . Alors  $p_F(x) = \sum_{i=1}^k \langle x, e_i \rangle e_i$  pour tout  $x \in E$ .

*Démonstration.* Tout  $x$  de  $E$  s'écrit

$$\begin{aligned} x &= \sum_{i=1}^n \langle x, e_i \rangle e_i \\ &= \sum_{i=1}^k \langle x, e_i \rangle e_i + \sum_{i=k+1}^n \langle x, e_i \rangle e_i \\ &= x_1 + x_2 \end{aligned}$$

avec  $x_1 = \sum_{i=1}^k \langle x, e_i \rangle e_i \in F$  et  $x_2 = \sum_{i=k+1}^n \langle x, e_i \rangle e_i \in F^\perp$ . D'où le résultat.  $\square$

Dans la pratique il suffit de connaître les  $(e_1, \dots, e_k)$  pour obtenir  $p_F(x)$  suivant la formule de la Proposition 16.2. On montre pour finir que la projection orthogonale de  $x$  sur  $F$  réalise la minimum de la distance de  $x$  aux vecteurs de  $F$ .

**Théorème 16.1.** *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie muni d'un produit scalaire  $\langle \cdot, \cdot \rangle$ . Soit  $F$  un sous espace vectoriel de  $E$ . Pour tout  $x \in E$ , et tout  $y \in F$ ,*

$$\|x - p_F(x)\| \leq \|x - y\|$$

*où  $p_F$  est la projection orthogonale sur  $F$  et  $\|\cdot\|$  est la norme associée au produit scalaire. De plus, si  $y \in F$  est tel que  $\|x - p_F(x)\| = \|x - y\|$ , alors forcément  $y = p_F(x)$ .*

*Démonstration.* Soit  $x = x_1 + x_2$  la décomposition de  $x$  suivant la décomposition  $E = F \oplus F^\perp$ . Donc  $x_1 = p_F(x)$ . Pour tout  $y \in F$ ,

$$\begin{aligned} \|x - y\|^2 &= \|(x_1 - y) + x_2\|^2 \\ &= \|x_1 - y\|^2 + \|x_2\|^2 \end{aligned}$$

puisque  $x_1 - y \in F$  et  $x_2 \in F^\perp$  sont orthogonaux. En particulier,  $\|x - p_F(x)\|^2 = \|x_2\|^2$ . Donc, clairement,  $\|x - p_F(x)\| \leq \|x - y\|$  pour tout  $y \in F$ , et si maintenant on a que  $\|x - p_F(x)\| = \|x - y\|$  pour un  $y \in F$ , alors forcément  $\|x_1 - y\| = 0$  et donc  $y = p_F(x)$ .  $\square$

## 17. SYMÉTRIES ORTHOGONALES

On utilise la décomposition  $E = F \oplus F^\perp$  du Théorème 15.2 pour définir la notion de symétrie orthogonale par rapport à  $F$ .

**Définition 17.1.** *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie muni d'un produit scalaire. Soit  $F$  un sous espace vectoriel de  $E$ . Comme  $E = F \oplus F^\perp$ , tout vecteur  $x$  de  $E$  se décompose de façon unique en  $x = x_1 + x_2$  avec  $x_1 \in F$  et  $x_2 \in F^\perp$ . On définit la symétrie orthogonale par rapport à  $F$  comme étant l'application  $s_F : E \rightarrow E$  donnée par  $s_F(x) = x_1 - x_2$ .*

La proposition suivante a lieu.

**Proposition 17.1.** *On a  $s_F = 2p_F - \text{Id}_E$ , où  $\text{Id}_E$  est l'application linéaire identité de  $E$  et  $p_F$  est la projection orthogonale sur  $F$  (Définition 16.1). En particulier  $s_F$  est un endomorphisme de  $E$ . Il est involutif au sens où  $s_F \circ s_F = \text{Id}_E$ , et  $s_F$  est donc bijective de réciproque elle-même. On a enfin que  $s_F = \text{Id}_E$  sur  $F$ .*

*Démonstration.* La preuve de cette proposition est simple. Si  $x = x_1 + x_2$  est la décomposition de  $x$  suivant la décomposition  $E = F \oplus F^\perp$ , alors

$$\begin{aligned} (2p_F - \text{Id}_E)(x) &= 2x_1 - x \\ &= x_1 - x_2 \\ &= s_F(x) \end{aligned}$$

pour tout  $x \in E$ . Donc  $s_F = 2p_F - \text{Id}_E$ . Comme  $p_F$  et  $\text{Id}_E$  sont linéaire,  $s_F$  l'est aussi. On a clairement que

$$(s_F \circ s_F)(x) = s_F(s_F(x)) = s_F(x_1 - x_2) = x_1 - (-x_2) = x_1 + x_2 = x$$

pour tout  $x$ , et donc  $s_F \circ s_F = \text{Id}_E$ . Enfin il est clair que  $s_F(x) = x$  pour tout  $x \in F$  puisqu'alors  $x_2 = 0$ . D'où la proposition.  $\square$

On démontre pour finir le résultat suivant. En présence d'un produit scalaire on dit qu'une application  $f$  est une isométrie si elle préserve le produit scalaire au sens où  $\langle f(x), f(y) \rangle = \langle x, y \rangle$  pour tous  $x, y$ . En présence d'une norme on dit

qu'une application  $f$  est une isométrie si elle préserve la norme au sens où  $\|f(x)\| = \|x\|$  pour tout  $x$ . En présence d'une distance on dit qu'une application  $f$  est une isométrie si elle préserve la distance au sens où  $d(f(x), f(y)) = d(x, y)$  pour tous  $x, y$ . Qui préserve un produit scalaire préserve automatiquement la norme associée.

**Théorème 17.1.** *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie muni d'un produit scalaire  $\langle \cdot, \cdot \rangle$ . Soit  $F$  un sous espace vectoriel de  $E$ . La symétrie orthogonale  $s_F$  par rapport à  $F$  est une isométrie au sens où  $\langle s_F(x), s_F(y) \rangle = \langle x, y \rangle$  pour tous  $x, y \in E$ . En particulier  $\|s_F(x)\| = \|x\|$  pour tout  $x \in E$ , où  $\|\cdot\|$  est la norme associée au produit scalaire  $\langle \cdot, \cdot \rangle$ .*

*Démonstration.* Soient  $x, y \in E$  et  $x = x_1 + x_2$ ,  $y = y_1 + y_2$  les décompositions de  $F$  suivant la décomposition  $E = F \oplus F^\perp$ . On a

$$\begin{aligned} \langle s_F(x), s_F(y) \rangle &= \langle x_1 - x_2, y_1 - y_2 \rangle \\ &= \langle x_1, y_1 \rangle + \langle x_2, y_2 \rangle - \langle x_1, y_2 \rangle - \langle x_2, y_1 \rangle \\ &= \langle x_1, y_1 \rangle + \langle x_2, y_2 \rangle \end{aligned}$$

puisque  $x_1, y_1 \in F$  et  $x_2, y_2 \in F^\perp$  de sorte que  $\langle x_1, y_2 \rangle = 0$  et  $\langle x_2, y_1 \rangle = 0$ . De même,

$$\begin{aligned} \langle x, y \rangle &= \langle x_1 + x_2, y_1 + y_2 \rangle \\ &= \langle x_1, y_1 \rangle + \langle x_2, y_2 \rangle + \langle x_1, y_2 \rangle + \langle x_2, y_1 \rangle \\ &= \langle x_1, y_1 \rangle + \langle x_2, y_2 \rangle \end{aligned}$$

puisque, comme ci-dessus,  $\langle x_1, y_2 \rangle = 0$  et  $\langle x_2, y_1 \rangle = 0$ . Donc  $\langle s_F(x), s_F(y) \rangle = \langle x, y \rangle$  pour tous  $x, y \in E$ . En prenant  $y = x$  on obtient que  $\|s_F(x)\| = \|x\|$  pour tout  $x \in E$ . Le théorème est démontré.  $\square$

## CHAPITRE 3

### LE THÉORÈME SPECTRAL

On revient ici à la problématique de diagonalisation des endomorphismes mais dans le cadre des endomorphismes symétriques. Ce que l'on commence par définir.

#### 18. ENDOMORPHISME ADJOINT.

La notion d'endomorphisme adjoint est liée au produit scalaire. Elle est définie dans le théorème qui suit.

**Théorème 18.1.** *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie muni d'un produit scalaire  $\langle \cdot, \cdot \rangle$ . Soit  $f \in \text{End}(E)$  un endomorphisme de  $E$ . Il existe alors un unique endomorphisme  $f^* \in \text{End}(E)$ , appelé adjoint de  $f$  (sous entendu pour  $\langle \cdot, \cdot \rangle$ ), tel que pour tous  $x, y \in E$ ,*

$$\langle f(x), y \rangle = \langle x, f^*(y) \rangle .$$

*Si  $\mathcal{B}$  est une base orthonormale de  $E$ ,  $f^*$  est caractérisé par le fait que sa matrice dans  $\mathcal{B}$  est la transposée de la matrice de  $f$  dans  $\mathcal{B}$ . En d'autres termes,  $M_{\mathcal{B}\mathcal{B}}(f^*) = {}^t M_{\mathcal{B}\mathcal{B}}(f)$  pour toute base orthonormale  $\mathcal{B}$  de  $E$ .*

*Démonstration.* Notons  $\mathcal{B} = (e_1, \dots, e_n)$  une base orthonormale de  $E$ , et notons  $a_{ij}$  les coefficients de la matrice de représentation de  $f$  dans  $\mathcal{B}$  de sorte que  $M_{\mathcal{B}\mathcal{B}}(f) = (a_{ij})$ . Donc

$$M_{\mathcal{B}\mathcal{B}}(f) = \begin{pmatrix} a_{11} & \dots & a_{1n} \\ \vdots & & \vdots \\ a_{n1} & \dots & a_{nn} \end{pmatrix} .$$

On construit l'endomorphisme  $f^* \in \text{End}(E)$  de  $E$  par

$$M_{\mathcal{B}\mathcal{B}}(f^*) = {}^t M_{\mathcal{B}\mathcal{B}}(f)$$

de sorte que  $M_{\mathcal{B}\mathcal{B}}(f^*) = (a_{ji})$ . Donc

$$M_{\mathcal{B}\mathcal{B}}(f^*) = \begin{pmatrix} a_{11} & \dots & a_{n1} \\ \vdots & & \vdots \\ a_{1n} & \dots & a_{nn} \end{pmatrix} .$$

Soient maintenant  $x$  et  $y$  deux vecteurs de  $E$  de coordonnées respectives  $x_i$  et  $y_i$  dans  $\mathcal{B}$ . Alors

$$\begin{aligned} f(x) &= \sum_{j=1}^n x_j f(e_j) \\ &= \sum_{i,j=1}^n a_{ij} x_j e_i , \end{aligned}$$

et donc, puisque  $\mathcal{B}$  est une base orthonormale,

$$\langle f(x), y \rangle = \sum_{i,j=1}^n a_{ij} y_i x_j .$$

Par ailleurs,

$$\begin{aligned} f^*(y) &= \sum_{j=1}^n y_j f^*(e_j) \\ &= \sum_{i,j=1}^n a_{ji} y_j e_i, \end{aligned}$$

et donc, toujours puisque  $\mathcal{B}$  est une base orthonormale,

$$\langle x, f^*(y) \rangle = \sum_{i,j=1}^n a_{ji} y_j x_i.$$

On trouve donc bien que pour tous  $x, y \in E$ ,

$$\langle f(x), y \rangle = \langle x, f^*(y) \rangle.$$

Il reste à ce stade à montrer que  $f^*$  est unique (et cela suffira à démontrer le théorème). Or

$$\langle f(x), y \rangle = \langle x, f^*(y) \rangle$$

entraîne en particulier que

$$\langle f(e_i), e_j \rangle = \langle e_i, f^*(e_j) \rangle$$

pour tous  $i, j$ . Mais si  $X$  est un vecteur de  $E$ , ses coordonnées dans  $\mathcal{B}$  sont exactement les  $X_i = \langle X, e_i \rangle$ . Il suit que

$$\langle f(e_i), e_j \rangle = a_{ji}$$

pour tous  $i, j$ , puis que pour tout  $j$ ,

$$f^*(e_j) = \sum_{i=1}^n a_{ji} e_i.$$

Par suite  $f^*$  est uniquement déterminée, et donc unique. Le théorème est démontré.  $\square$

On rappelle que si

$$M_{\mathcal{B}\mathcal{B}}(f) = \begin{pmatrix} a_{11} & \dots & a_{1n} \\ \vdots & & \vdots \\ a_{n1} & \dots & a_{nn} \end{pmatrix}$$

dans une base  $\mathcal{B} = (e_1, \dots, e_n)$ , alors pour tout  $j$ ,  $f(e_j) = \sum_{i=1}^n a_{ij} e_i$ .

**Proposition 18.1.** *Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie muni d'un produit scalaire  $\langle \cdot, \cdot \rangle$ . Soit  $f \in \text{End}(E)$  un endomorphisme de  $E$ . Soit  $f^*$  l'endomorphisme adjoint de  $f$ . On a  $\text{Ker}(f^*) = \text{Im}(f)^\perp$  et  $\text{Im}(f^*) = \text{Ker}(f)^\perp$ .*

*Démonstration.* On a  $y \in \text{Ker}(f^*)$  si et seulement si  $f^*(y) = 0$ . Et  $f^*(y) = 0$  si et seulement si pour tout  $x$ ,  $\langle f(x), y \rangle = 0$ . Donc  $f^*(y) = 0$  si et seulement si  $y \in \text{Im}(f)^\perp$ . Soit maintenant  $z \in \text{Im}(f^*)$ . Alors  $z = f^*(x)$ . Soit  $y \in \text{Ker}(f)$  quelconque. On a  $\langle y, z \rangle = \langle f(y), x \rangle = 0$  puisque  $f(y) = 0$ . Donc  $\text{Im}(f^*) \subset \text{Ker}(f)^\perp$ . Avec le théorème du rang,

$$\dim \text{Ker}(f^*) + \dim \text{Im}(f^*) = n$$

si  $n$  est la dimension de  $E$  (donc de  $E^*$ ). Le Théorème 15.2 et l'égalité  $\text{Ker}(f^*) = \text{Im}(f)^\perp$  donnent que  $\dim \text{Ker}(f^*) = n - p$  si  $p = \dim \text{Im}(f) = \text{Rg}(f)$ . Donc

$\dim \text{Im}(f^*) = p$ . Encore avec le Théorème 15.2 on peut écrire que  $\dim \text{Ker}(f)^\perp = n - \dim \text{Ker}(f)$  tandis que le théorème du rang donne que  $\dim \text{Ker}(f) = n - \text{Rg}(f) = n - p$ . Donc  $\dim \text{Ker}(f)^\perp = p$ . Les deux espaces vectoriels  $\text{Im}(f^*)$  et  $\text{Ker}(f)^\perp$  ont donc même dimension et  $\text{Im}(f^*) \subset \text{Ker}(f)^\perp$ . Donc  $\text{Im}(f^*) = \text{Ker}(f)^\perp$ . La proposition est démontrée.  $\square$

**Exercice:** Soit  $E$  un  $\mathbb{R}$ -espace vectoriel muni d'un produit scalaire  $\langle \cdot, \cdot \rangle$ . Soient  $u, v \in \text{End}(E)$  deux endomorphismes de  $E$ . Montrer que  $(v \circ u)^* = u^* \circ v^*$ . On suppose maintenant que  $u$  et  $v$  sont symétriques au sens de la Définition 19.1, à savoir  $u^* = u$  et  $v^* = v$ . Montrer que  $v \circ u$  est symétrique, à savoir  $(v \circ u)^* = v \circ u$ , si et seulement si  $v \circ u = u \circ v$ .

**Solution:** Soient  $x, y \in E$  quelconques. On a

$$\begin{aligned} \langle (v \circ u)(x), y \rangle &= \langle v(u(x)), y \rangle \\ &= \langle u(x), v^*(y) \rangle \\ &= \langle x, u^*(v^*(y)) \rangle \\ &= \langle x, (u^* \circ v^*)(y) \rangle. \end{aligned}$$

Par suite, comme  $x$  et  $y$  sont quelconques,  $(v \circ u)^* = u^* \circ v^*$ . On suppose maintenant que  $u$  et  $v$  sont symétriques, donc que  $u^* = u$  et  $v^* = v$ . Supposons que  $v \circ u$  est symétrique. Alors  $(v \circ u)^* = v \circ u$ . Comme par ailleurs on a aussi  $u^* \circ v^* = u \circ v$ , et comme  $(v \circ u)^* = u^* \circ v^*$ , on en déduit que  $v \circ u = u \circ v$ . Réciproquement supposons que  $v \circ u = u \circ v$ . Comme  $u^* \circ v^* = u \circ v$  on a écrit que  $v \circ u = u^* \circ v^*$ . Par suite, comme  $(v \circ u)^* = u^* \circ v^*$ , on a écrit que  $(v \circ u)^* = v \circ u$ . Donc que  $v \circ u$  est symétrique.  $\square$

## 19. LE THÉORÈME SPECTRAL

La notion d'endomorphisme symétrique est associée à un produit scalaire.

**Définition 19.1.** Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie muni d'un produit scalaire. On dit qu'un endomorphisme  $f \in \text{End}(E)$  est symétrique (ou encore autoadjoint) si  $f = f^*$ , où  $f^*$  est l'adjoint de  $f$  (pour le produit scalaire).

Dire que  $f$  est symétrique revient encore à dire que pour toute base orthonormale  $\mathcal{B}$  de  $E$ , la matrice de  $f$  dans  $\mathcal{B}$  est symétrique. L'objectif de cette section est de démontrer le théorème spectral suivant.

**Théorème 19.1** (Théorème spectral). Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie muni d'un produit scalaire, et soit  $f \in \text{End}(E)$  un endomorphisme symétrique de  $E$ , donc tel que  $f = f^*$ . Alors  $f$  est diagonalisable, qui plus est dans une base orthonormale de  $E$ .

En d'autres termes, pour tout endomorphisme symétrique  $f \in \text{End}(E)$  de  $E$ , il existe une base orthonormale de  $E$  qui soit composée de vecteurs propres de  $f$ .

On propose tout d'abord une preuve de ce théorème dans le cas  $n = 2$ . On montre successivement les trois points suivants:

- 1 - Les valeurs propres de  $f$  sont réelles,
- 2 - Les espaces propres de  $f$  sont deux à deux orthogonaux,
- 3 -  $f$  est diagonalisable dans une base orthonormale,

de sorte que, au total,  $f$  est diagonalisable, qui plus est dans une base orthonormale de  $E$ , ce qui démontre le théorème lorsque  $n = 2$ . Une preuve lorsque  $n = 3$  est proposée à la section suivante. La preuve dans le cas général  $n \geq 3$  est proposée à la dernière section de ce chapitre. On suppose donc dans la suite que  $n = 2$ .

*Démonstration de l'étape 1.* On affirme que si  $f \in \text{End}(E)$  est un endomorphisme symétrique de  $E$ , alors les valeurs propres de  $f$  sont toutes réelles. Pour le voir, dans ce cas “simple” où  $n = 2$ , on considère une base orthonormale  $\mathcal{B}$  de  $E$ . Alors la matrice de représentation de  $f$  dans cette base est symétrique, et donc du type

$$M_{\mathcal{B}\mathcal{B}}(f) = \begin{pmatrix} a & b \\ b & c \end{pmatrix} .$$

Par suite, le polynôme caractéristique de  $f$  vaut

$$\begin{aligned} P(X) &= (a - X)(c - X) - b^2 \\ &= X^2 - (a + c)X + (ac - b^2) . \end{aligned}$$

Son discriminant vaut  $\Delta = (a + c)^2 - 4(ac - b^2) = (a - c)^2 + 4b^2$  qui est toujours positif ou nul. Donc  $P$  a toutes ses racines dans  $\mathbb{R}$ . En d'autres termes, les valeurs propres de  $f$  sont toutes réelles.  $\square$

*Démonstration de l'étape 2.* On affirme que si  $f \in \text{End}(E)$  est un endomorphisme symétrique de  $E$ , alors les espaces propres de  $f$  sont deux à deux orthogonaux. Dans ce cas particulier où  $n = 2$ , soit  $f$  a une seule valeur propre, et il n'y a rien à démontrer, soit  $f$  a deux valeurs propres distinctes. On note  $\lambda_1$  et  $\lambda_2$  les deux valeurs propres distinctes de  $f$ , et  $x$  et  $y$  deux vecteurs propres de  $f$  respectivement associés aux valeurs propres  $\lambda_1$  et  $\lambda_2$ . Puisque  $f$  est symétrique,

$$\langle f(x), y \rangle = \langle x, f(y) \rangle ,$$

et donc

$$(\lambda_1 - \lambda_2)\langle x, y \rangle = 0$$

puisque

$$f(x) = \lambda_1 x \text{ et } f(y) = \lambda_2 y .$$

Comme  $\lambda_1 \neq \lambda_2$ , il suit que

$$\langle x, y \rangle = 0 .$$

D'où l'affirmation.  $\square$

*Démonstration de l'étape 3.* On affirme que si  $f \in \text{End}(E)$  est un endomorphisme symétrique de  $E$ , alors  $E$  est la somme directe des espaces propres de  $f$ . Si  $f$  a deux valeurs propres distinctes  $\lambda_1$  et  $\lambda_2$ , alors on sait déjà que  $E$  est somme (directe) des espaces propres de  $f$ . Si  $e_1$  est un vecteur de norme 1 de  $E_{\lambda_1}$  et  $e_2$  est un vecteur de norme 1 de  $E_{\lambda_2}$ , alors d'après l'étape 2,  $(e_1, e_2)$  est une base orthonormale de  $E$ . Dans cette base  $M_{\mathcal{B}\mathcal{B}}(f)$  est diagonale. Le théorème est démontré dans ce cas. Si maintenant  $f$  a une seule valeur propre  $\lambda_1$ , alors le discriminant  $\Delta$  du polynôme caractéristique de  $f$  doit être nul. Etant donnée une base orthonormale  $\mathcal{B}$ , on a vu que dans cette base

$$\Delta = (a - c)^2 + 4b^2 .$$

Donc  $a = c$  et  $b = 0$ . En d'autres termes,

$$M_{\mathcal{B}\mathcal{B}}(f) = \begin{pmatrix} a & 0 \\ 0 & a \end{pmatrix} ,$$

et  $f$  est diagonalisable dans  $\mathcal{B}$  qui est orthonormale. Le théorème est démontré.  $\square$

**Exercice:** Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie muni d'un produit scalaire  $\langle \cdot, \cdot \rangle$ . Soit  $u \in \text{End}(E)$  un endomorphisme de  $E$ . On suppose que  $u$  est symétrique et que  $\langle u(x), x \rangle = 0$  pour tout  $x \in E$ . Montrer que  $u$  est l'endomorphisme nul.

**Solution:** Comme  $u$  est symétrique,  $u$  est diagonalisable, et même  $u$  est diagonalisable dans une base orthonormée de  $E$ . En particulier il existe  $\tilde{\mathcal{B}} = (\tilde{e}_1, \dots, \tilde{e}_n)$  une base (orthonormée) qui diagonalise  $u$ . Si  $\lambda_i$  est la valeur propre associée à  $e_i$  alors

$$\langle u(e_i), e_i \rangle = \lambda_i \|e_i\|^2.$$

Par hypothèse  $\langle u(e_i), e_i \rangle = 0$ , et puisque  $\|e_i\| \neq 0$ , on récupère  $\lambda_i = 0$ . Comme  $i$  est quelconque,  $\lambda_i = 0$  pour tout  $i$ , et il suit que  $M_{\tilde{\mathcal{B}}\tilde{\mathcal{B}}}(u)$  est la matrice nulle. En particulier  $u$  est l'endomorphisme nul.  $\square$

Une conséquence directe de la preuve est donnée par le résultat suivant (cf. l'étape 2 ci-dessus, ou l'étape 2 de la section 24 en dimension quelconque).

**Lemme 19.1.** *Les espaces propres d'un endomorphisme symétrique sont deux à deux orthogonaux.*

## 20. PREUVE LORSQUE $n = 3$

On suppose ici que  $(E, \langle \cdot, \cdot \rangle)$  est un espace préhilbertien de dimension 3 et que  $f \in \text{End}(E)$  est un endomorphisme symétrique de  $E$ . On se propose de montrer, qu'alors,  $f$  est diagonalisable dans une base orthonormée de  $(E, \langle \cdot, \cdot \rangle)$ .

Le polynôme caractéristique  $P$  de  $f$  est de degré 3 puisque  $E$  est de dimension 3. Son terme de plus haut degré est  $-X^3$ . On a donc que

$$\lim_{x \rightarrow -\infty} P(x) = +\infty \text{ et } \lim_{x \rightarrow +\infty} P(x) = -\infty.$$

Par suite, et par théorème des valeurs intermédiaires, puisqu'un polynôme est toujours une fonction continue, il existe  $\lambda_1 \in \mathbb{R}$  une racine de  $P$ . En d'autres termes,  $f$  a au moins une valeur propre réelle  $\lambda_1$ .

On prétend que  $(f - \lambda_1 \text{Id}_E)^* = f^* - \lambda_1 \text{Id}_E$  pour tout endomorphisme  $f$  de  $E$ . En effet, pour tous  $x, y \in E$ ,

$$\begin{aligned} \langle f(x) - \lambda_1 x, y \rangle &= \langle f(x), y \rangle - \lambda_1 \langle x, y \rangle \\ &= \langle x, f^*(y) \rangle - \langle x, \lambda_1 y \rangle \\ &= \langle x, (f^* - \lambda_1 \text{Id}_E)(y) \rangle, \end{aligned}$$

ce qui prouve notre assertion. Dans notre cas, puisque  $f$  est symétrique,  $(f - \lambda_1 \text{Id}_E)^* = f - \lambda_1 \text{Id}_E$ .

On prétend maintenant que  $E_{\lambda_1}$ , l'espace propre associé à  $\lambda_1$ , et  $E_{\lambda_1}^\perp$ , son orthogonal, sont tous deux stables par  $f$ , à savoir que  $f(E_{\lambda_1}) \subset E_{\lambda_1}$  et  $f(E_{\lambda_1}^\perp) \subset E_{\lambda_1}^\perp$ . Pour  $E_{\lambda_1}$  c'est immédiat. En effet, si  $x \in E_{\lambda_1}$  alors  $f(x) = \lambda_1 x$  et donc  $f(x) \in E_{\lambda_1}$  puisque  $E_{\lambda_1}$  est un espace vectoriel, et l'on a donc bien que  $f(E_{\lambda_1}) \subset E_{\lambda_1}$ . Par ailleurs, avec la Proposition 18.1, et ce qui a été dit ci-dessus,

$$E_{\lambda_1}^\perp = \text{Ker}(f - \lambda_1 \text{Id}_E)^\perp = \text{Im}(f - \lambda_1 \text{Id}_E).$$

Par suite  $y \in E_{\lambda_1}^\perp$  si et seulement si il existe  $x \in E$  tel que  $y = f(x) - \lambda_1 x$ . On a que

$$f(y) = f(f(x) - \lambda_1 x) = (f - \lambda_1 \text{Id}_E)(f(x))$$

si  $y = f(x) - \lambda_1 x$ , et donc  $f(y) \in E_{\lambda_1}^\perp$  si  $y \in E_{\lambda_1}^\perp$ . Donc  $f(E_{\lambda_1}^\perp) \subset E_{\lambda_1}^\perp$  comme annoncé.

Si  $E_{\lambda_1} = E$  il n'y a rien à démontrer puisqu'alors  $f = \lambda_1 \text{Id}_E$ . On peut donc supposer dans ce qui suit que  $\dim(E_{\lambda_1}) = 1$  ou  $\dim(E_{\lambda_1}) = 2$ . De plus, cf. le chapitre précédent, on a toujours que  $E = E_{\lambda_1} \oplus E_{\lambda_1}^\perp$ . En particulier, on obtient que  $\dim(E_{\lambda_1}^\perp) = 3 - \dim(E_{\lambda_1})$ .

Si  $\dim(E_{\lambda_1}) = 1$  alors  $\dim(E_{\lambda_1}^\perp) = 2$ . Soit  $g = f|_{E_{\lambda_1}^\perp}$  la restriction de  $f$  à  $E_{\lambda_1}^\perp$ . En raison de ce qui a été dit plus haut,  $g \in \text{End}(E_{\lambda_1}^\perp)$ . De plus  $g$  est symétrique (pour le produit scalaire restreint à  $E_{\lambda_1}^\perp$ ) puisque  $f$  est symétrique:  $\forall x, y \in E_{\lambda_1}^\perp$ ,

$$\langle g(x), y \rangle = \langle x, g(y) \rangle.$$

On a démontré le théorème spectral en dimension 2. En vertu de ce théorème il existe  $(\tilde{e}_2, \tilde{e}_3)$  une base orthonormale de  $E_{\lambda_1}^\perp$  qui diagonalise  $g$ . Et donc, en particulier,  $f(\tilde{e}_2) = \lambda \tilde{e}_2$  et  $f(\tilde{e}_3) = \mu \tilde{e}_3$  pour des  $\lambda, \mu \in \mathbb{R}$ . Soit  $\tilde{e}_1$  un vecteur directeur de norme 1 de  $E_{\lambda_1}$ . On a  $f(\tilde{e}_1) = \lambda_1 \tilde{e}_1$  et  $(\tilde{e}_1, \tilde{e}_2, \tilde{e}_3)$  est une famille orthonormée. Cette famille a 3 vecteurs et  $E$  est de dimension 3. C'est donc une base orthonormée de  $E$  et elle diagonalise  $f$ . C'est ce qu'il fallait démontrer.

Si  $\dim(E_{\lambda_1}) = 2$  alors  $\dim(E_{\lambda_1}^\perp) = 1$ . Soit  $g = f|_{E_{\lambda_1}}$  la restriction de  $f$  à  $E_{\lambda_1}$ . Pour tout  $x \in E_{\lambda_1}$ ,  $g(x) = f(x) = \lambda_1 x$ . Donc  $g = \lambda_1 \text{Id}_{E_{\lambda_1}}$  sur  $E_{\lambda_1}$ . Si  $(\tilde{e}_1, \tilde{e}_2)$  est une base orthonormale de  $E_{\lambda_1}$ , alors  $f(\tilde{e}_1) = \lambda_1 \tilde{e}_1$  et  $f(\tilde{e}_2) = \lambda_1 \tilde{e}_2$ . Soit  $\tilde{e}_3$  un vecteur directeur de norme 1 de  $E_{\lambda_1}^\perp$ . On a  $f(\tilde{e}_3) = \lambda \tilde{e}_3$  pour un certain  $\lambda \in \mathbb{R}$  puisque  $f(E_{\lambda_1}^\perp) \subset E_{\lambda_1}^\perp$ . La famille  $(\tilde{e}_1, \tilde{e}_2, \tilde{e}_3)$  est une famille orthonormée. Cette famille a 3 vecteurs et  $E$  est de dimension 3. C'est donc une base orthonormée de  $E$  et elle diagonalise  $f$ . C'est ce qu'il fallait démontrer.

Au total, le théorème spectral est démontré en dimension 3.

## 21. LE POINT DE VUE DES MATRICES

On commence par définir la notion de matrice orthogonale.

**Définition 21.1.** Soit  $P$  une matrice  $n \times n$ . On dit que  $P$  est orthogonale si  $P$  est inversible et  $P^{-1} = {}^t P$ .

Les matrices orthogonales sont les matrices de passage d'une base orthonormale à une base orthonormale. C'est l'objet du résultat suivant.

**Lemme 21.1.** Soit  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie muni d'un produit scalaire. Si  $\mathcal{B}$  et  $\tilde{\mathcal{B}}$  sont deux bases orthonormales de  $E$ , alors la matrice  $M_{\mathcal{B} \rightarrow \tilde{\mathcal{B}}}$  de passage de  $\mathcal{B}$  à  $\tilde{\mathcal{B}}$  est telle que  $M_{\mathcal{B} \rightarrow \tilde{\mathcal{B}}}^{-1} = {}^t M_{\mathcal{B} \rightarrow \tilde{\mathcal{B}}}$ . En d'autres termes, la matrice de passage  $M_{\mathcal{B} \rightarrow \tilde{\mathcal{B}}}$  d'une base orthonormale  $\mathcal{B}$  à une base orthonormale  $\tilde{\mathcal{B}}$  est une matrice orthogonale.

*Démonstration.* Notons  $\mathcal{B} = (e_1, \dots, e_n)$ ,  $\tilde{\mathcal{B}} = (\tilde{e}_1, \dots, \tilde{e}_n)$ , et  $\langle \cdot, \cdot \rangle$  le produit scalaire de  $E$ . Pour  $i, j = 1, \dots, n$ , on note de même

$$a_{ij} = \langle e_i, \tilde{e}_j \rangle$$

et on note  $\Phi$  et  $\Psi$  les endomorphismes de  $E$  définis par

$$\Phi(e_i) = \tilde{e}_i \quad \text{et} \quad \Psi(\tilde{e}_i) = e_i$$

pour tout  $i$ . Alors

$$M_{\mathcal{B} \rightarrow \tilde{\mathcal{B}}} = M_{\mathcal{B}\mathcal{B}}(\Phi) \quad \text{tandis que} \quad M_{\mathcal{B} \rightarrow \tilde{\mathcal{B}}}^{-1} = M_{\mathcal{B}\mathcal{B}}(\Psi) .$$

Sachant que les coordonnées d'un vecteur  $X$  dans une base orthonormale  $(e_1, \dots, e_n)$  sont exactement les  $X_i = \langle X, e_i \rangle$ , on voit que  $M_{\mathcal{B}\mathcal{B}}(\Phi) = (a_{ij})$ . Par ailleurs,

$$\begin{aligned} e_j &= \sum_{i=1}^n \langle e_j, \tilde{e}_i \rangle \tilde{e}_i \\ &= \sum_{i=1}^n a_{ji} \tilde{e}_i \end{aligned}$$

de sorte que

$$\Psi(e_j) = \sum_{i=1}^n a_{ji} e_i .$$

On en déduit que  $M_{\mathcal{B}\mathcal{B}}(\Psi) = (a_{ji})$ . Donc

$$M_{\mathcal{B}\mathcal{B}}(\Psi) = {}^t M_{\mathcal{B}\mathcal{B}}(\Phi) ,$$

et le lemme est démontré.  $\square$

Une autre justification de la terminologie “matrice orthogonale” est la suivante. Notons  $\langle \cdot, \cdot \rangle$  le produit scalaire euclidien de  $\mathbb{R}^n$  donné par

$$\langle x, y \rangle = \sum_{i=1}^n x_i y_i$$

où les  $x_i$  et  $y_i$  sont les coordonnées de  $x$  et  $y$  dans la base canonique de  $\mathbb{R}^n$ . Si on note  $X$  la matrice vecteur colonne constituée des  $x_i$  et  $Y$  la matrice vecteur colonne constituée des  $y_i$ , on a aussi

$$\langle x, y \rangle = {}^t X Y .$$

On constate alors très facilement que les matrices orthogonales préservent le produit scalaire, et donc en particulier aussi l'orthogonalité (d'où une autre explication possible du nom qui leur est donné). Plus précisément, on montre très facilement que si  $P$  est orthogonale, alors pour tous  $X, Y$ ,

$${}^t(PX)(PY) = {}^tXY$$

et ce tout simplement puisque  ${}^t(PX)(PY) = {}^tX({}^tPP)Y$  et  ${}^tPP = \text{Id}_n$  lorsque  $P$  est orthogonale. En particulier si  $X$  et  $Y$  sont orthogonaux, alors  $PX$  et  $PY$  le sont aussi.

On passe maintenant à la transcription du théorème spectral du langage des endomorphismes symétriques au langage des matrices, ce qui donne lieu au théorème suivant.

**Théorème 21.1.** *Soit  $A$  une matrice réelle symétrique. Il existe alors une matrice orthogonale  $P$ , donc telle que  $P^{-1} = {}^tP$ , pour laquelle la matrice  $P^{-1}AP$  est diagonale. En particulier, une matrice symétrique est diagonalisable.*

*Démonstration.* Soit  $A$  une matrice symétrique  $n \times n$ . On considère  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension  $n$ ,  $\langle \cdot, \cdot \rangle$  un produit scalaire sur  $E$ , et  $\mathcal{B} = (e_1, \dots, e_n)$  une base orthonormale de  $E$ . Par exemple,  $E = \mathbb{R}^n$ ,  $\langle \cdot, \cdot \rangle$  le produit scalaire euclidien, et  $\mathcal{B}$  la base canonique de  $\mathbb{R}^n$ . On construit l'endomorphisme  $f \in \text{End}(E)$  de  $E$  par

$$M_{\mathcal{B}\mathcal{B}}(f) = A .$$

Puisque  $A$  est symétrique, et  $\mathcal{B}$  est orthonormale,  $f$  est aussi symétrique. Il existe donc, en vertu du théorème vu plus haut, une base orthonormale  $\tilde{\mathcal{B}}$  de  $E$  qui diagonalise  $f$ , et donc pour laquelle  $M_{\tilde{\mathcal{B}}\tilde{\mathcal{B}}}(f)$  est diagonale. En vertu de la formule du changement de base vue dans les chapitres précédents,

$$M_{\tilde{\mathcal{B}}\tilde{\mathcal{B}}}(f) = M_{\mathcal{B} \rightarrow \tilde{\mathcal{B}}}^{-1} M_{\mathcal{B}\mathcal{B}}(f) M_{\mathcal{B} \rightarrow \tilde{\mathcal{B}}} .$$

Si  $P = M_{\mathcal{B} \rightarrow \tilde{\mathcal{B}}}$ , alors  $P$  est orthogonale, à savoir  $P^{-1} = {}^tP$ , et la formule précédente nous dit que  $P^{-1}AP$  est diagonale. D'où le théorème.  $\square$

**Exercice:** Soit  $A$  une matrice réelle symétrique. On suppose qu'il existe  $p \in \mathbb{N}^*$  tel que  $A^p = 0$  (on dit que  $A$  est nilpotente). Montrer que  $A$  est la matrice nulle.

**Solution:** Comme  $A$  est symétrique elle est diagonalisable. En particulier il existe une matrice orthogonale  $P$  (i.e.  $P^{-1} = {}^tP$ ) et une matrice diagonale  $D$  telles que

$${}^tPAP = D .$$

On en déduit que

$${}^tPA^pP = D^p ,$$

et si  $A^p = 0$  alors  $D^p = 0$ . Si la diagonale de  $D$  est constituée des  $\lambda_i$ , alors  $D^p$  est encore diagonale et elle est constituée des  $\lambda_i^p$ . On a donc  $\lambda_i^p = 0$  pour tout  $i$ , donc  $\lambda_i = 0$  pour tout  $i$ . Par suite  $D = 0$ . Comme  $P$  est inversible et  $P^{-1} = {}^tP$ , on a que  $A = PD^tP$ . D'où  $A = 0$ .  $\square$

A titre de remarque, il faut faire attention au langage des matrices. En général on ne retient que la dernière affirmation du théorème (une matrice symétrique est diagonalisable) en oubliant l'information supplémentaire importante que la matrice qui "la diagonalise" peut-être choisie de sorte qu'elle soit orthogonale. Cela correspond au fait qu'un endomorphisme symétrique est diagonalisable *dans une base qui peut être choisie orthonormale*.

**Lemme 21.2.** Soient  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie  $n$  muni d'un produit scalaire. Si  $\mathcal{B}$  et  $\tilde{\mathcal{B}}$  sont deux bases de  $E$ , si  $\mathcal{B}$  est orthonormale, et si  $M_{\mathcal{B} \rightarrow \tilde{\mathcal{B}}}$  est une matrice orthogonale, alors  $\tilde{\mathcal{B}}$  est aussi une base orthonormale.

*Démonstration.* On note  $\langle \cdot, \cdot \rangle$  le produit scalaire placé sur  $E$ , et on note  $\mathcal{B} = (e_1, \dots, e_n)$ ,  $\tilde{\mathcal{B}} = (\tilde{e}_1, \dots, \tilde{e}_n)$ . Soit de plus  $\Phi$  et  $\Psi$  les endomorphismes de  $E$  définis par

$$\Phi(e_i) = \tilde{e}_i \quad \text{et} \quad \Psi(\tilde{e}_i) = e_i$$

pour tout  $i$ . Alors  $\Psi = \Phi^{-1}$ , et  $M_{\mathcal{B}\mathcal{B}}(\Phi) = M_{\mathcal{B} \rightarrow \tilde{\mathcal{B}}}$ . En particulier,

$$\begin{aligned} M_{\mathcal{B}\mathcal{B}}(\Psi) &= M_{\mathcal{B}\mathcal{B}}(\Phi)^{-1} && (\text{puisque } \Psi = \Phi^{-1}) \\ &= {}^tM_{\mathcal{B}\mathcal{B}}(\Phi) && (\text{puisque } M_{\mathcal{B} \rightarrow \tilde{\mathcal{B}}} \text{ est orthogonale}) \end{aligned}$$

et il suit que l'on a aussi  $\Psi = \Phi^*$ , l'adjoint de  $\Phi$ . Donc, pour tous  $x, y \in E$ ,

$$\langle \Phi(x), y \rangle = \langle x, \Psi(y) \rangle .$$

En particulier, pour tous  $i, j = 1, \dots, n$ ,

$$\begin{aligned}\langle \tilde{e}_i, \tilde{e}_j \rangle &= \langle \Phi(e_i), \tilde{e}_j \rangle \\ &= \langle e_i, \Psi(\tilde{e}_j) \rangle \\ &= \langle e_i, e_j \rangle ,\end{aligned}$$

et puisque  $\mathcal{B}$  est orthonormale,  $\tilde{\mathcal{B}}$  l'est aussi. D'où la preuve du lemme.  $\square$

**Exercice:** Soit  $A$  la matrice

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

Montrer que  $A$  est diagonalisable. Trouver une matrice orthogonale  $P$  pour laquelle  ${}^tPAP$  est diagonale.

**Solution:** La matrice  $A$  est symétrique, donc diagonalisable et il existe  $P$  orthogonale telle que  ${}^tPAP$  est diagonale. Soient  $E = \mathbb{R}^3$ ,  $\langle \cdot, \cdot \rangle$  le produit scalaire euclidien de  $\mathbb{R}^3$ ,  $\mathbb{B} = (e_1, e_2, e_3)$  la base canonique de  $\mathbb{R}^3$  (elle est orthonormée) et  $f \in \text{End}(\mathbb{R}^3)$  l'endomorphisme défini par

$$M_{\mathcal{B}\mathcal{B}}(f) = A .$$

Alors  $f$  est diagonalisable, qui plus est dans une base orthonormée  $\tilde{\mathcal{B}}$  pour  $\langle \cdot, \cdot \rangle$ . La matrice orthogonale  $P$  recherchée est alors donnée par  $P = M_{\mathcal{B} \rightarrow \tilde{\mathcal{B}}}$  (cours). Le polynôme caractéristique de  $f$  est donné par

$$P(X) = -X(X-3)^2 + 24X + 8 = -X^3 + 6X^2 + 15X + 8 .$$

On remarque que  $P(-1) = 0$ . Par suite  $P$  se factorise en

$$P(X) = -(X+1)(X^2 + \alpha X - 8)$$

avec  $\alpha \in \mathbb{R}$ . On trouve  $\alpha + 1 = -6$  et  $\alpha - 8 = -15$  de sorte que  $\alpha = -7$ . Par suite

$$\begin{aligned}P(X) &= -(X+1)(X^2 - 7X - 8) \\ &= -(X+1)^2(X-8) .\end{aligned}$$

Les valeurs propres de  $f$  sont  $-1$  et  $8$ . Soient  $E_{-1}$  et  $E_8$  les espaces propres associés. On sait déjà que  $E_8$  est de dimension 1, et comme  $f$  est diagonalisable  $E_{-1}$  est forcément de dimension 2. On trouve que

$$\begin{aligned}E_{-1} &= \{(x, y, z) / 2x + y + 2z = 0\} \\ &= \{x(1, -2, 0) + z(0, -2, 1) , x, z \in \mathbb{R}\} .\end{aligned}$$

Soient  $u = (1, -2, 0)$  et  $v = (0, -2, 1)$ . Alors  $(u, v)$  est génératrice pour  $E_{-1}$ . On vérifie facilement que  $(u, v)$  est aussi libre. Donc  $(u, v)$  est une base pour  $E_{-1}$ . On orthonormalise maintenant  $(u, v)$  en appliquant Gram-Schmidt. On a  $\|u\|^2 = 5$ . On pose

$$\tilde{e}_1 = \frac{1}{\sqrt{5}}u = \left( \frac{1}{\sqrt{5}}, -\frac{2}{\sqrt{5}}, 0 \right) .$$

On calcule

$$\langle v, \tilde{e}_1 \rangle = \frac{4}{\sqrt{5}} .$$

On pose

$$\tilde{v} = v - \langle v, \tilde{e}_1 \rangle \tilde{e}_1 = \left( -\frac{4}{5}, -\frac{2}{5}, 1 \right) .$$

On a

$$\|\tilde{v}\|^2 = \frac{45}{25} = \frac{9}{5} .$$

On pose

$$\tilde{e}_2 = \frac{\sqrt{5}}{3} \tilde{v} = \left( -\frac{4}{3\sqrt{5}}, -\frac{2}{3\sqrt{5}}, \frac{\sqrt{5}}{3} \right) .$$

Alors  $(\tilde{e}_1, \tilde{e}_2)$  est une base orthonormée de  $E_{-1}$ . Indépendamment, on trouve que

$$\begin{aligned} E_8 &= \{(x, y, z) / x = z = 2y\} \\ &= \{y(2, 1, 2) , y \in \mathbb{R}\} . \end{aligned}$$

Donc  $E_8$  est la droite vectorielle de vecteur de base  $w = (2, 1, 2)$ . On a  $\|w\| = 3$ .

On pose

$$\tilde{e}_3 = \frac{1}{3}w = \left( \frac{2}{3}, \frac{1}{3}, \frac{2}{3} \right) .$$

Alors (cours) la famille  $\tilde{\mathcal{B}} = (\tilde{e}_1, \tilde{e}_2, \tilde{e}_3)$  est une base orthonormée et si

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

on a que

$${}^tPAP = \begin{pmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 8 \end{pmatrix}$$

□

## 22. DÉCOMPOSITIONS D'IWASAWA ET DE CHOLESKY

On aborde ici deux décompositions (écritures) particulières des matrices carrées (invertible pour l'une, symétrique définie positive pour l'autre). La première des deux décompositions, dite d'Iwasawa, ou encore décomposition  $QR$ , concerne les matrices inversibles (il y a plus général pour les matrices  $n \times p$  avec  $n \geq p$ ).

**Théorème 22.1** (Décomposition d'Iwasawa d'une matrice inversible). *Soit  $M$  une matrice carrée  $n \times n$  inversible. Il existe alors  $Q$  orthogonale  $n \times n$  et  $R$  triangulaire supérieure  $n \times n$  à diagonale strictement positive telles que  $M = QR$ .*

*Démonstration.* Soit  $\mathbb{R}^n$  muni du produit scalaire euclidien  $\langle \cdot, \cdot \rangle$ . Soit  $\mathcal{B}$  la base canonique de  $\mathbb{R}^n$ . Elle est orthonormale pour  $\langle \cdot, \cdot \rangle$ . Comme  $M$  est inversible il existe une (et une seule) base  $\mathcal{B}'$  de  $\mathbb{R}^n$  telle que

$$M_{\mathcal{B} \rightarrow \mathcal{B}'} = M .$$

On note  $\mathcal{B}''$  la base orthonormale obtenue à partir de  $\mathcal{B}'$  par le procédé d'orthonormalisation de Gram-Schmidt. On note

$$Q = M_{\mathcal{B} \rightarrow \mathcal{B}''} \text{ et } R = M_{\mathcal{B}'' \rightarrow \mathcal{B}'} .$$

Comme  $\mathcal{B}$  et  $\mathcal{B}''$  sont orthonormales,  $Q$  est orthogonale. Soit  $\text{Id}$  l'application linéaire identité de  $\mathbb{R}^n$ . On écrit que

$$\begin{aligned} M_{\mathcal{B} \rightarrow \mathcal{B}'} &= M_{\mathcal{B}'\mathcal{B}}(\text{Id}) = M \\ M_{\mathcal{B} \rightarrow \mathcal{B}''} &= M_{\mathcal{B}''\mathcal{B}}(\text{Id}) = Q \\ M_{\mathcal{B}'' \rightarrow \mathcal{B}'} &= M_{\mathcal{B}'\mathcal{B}''}(\text{Id}) = R . \end{aligned}$$

On rappelle la formule de composition: si

$$(E_1, \mathcal{B}_1) \xrightarrow{f} (E_2, \mathcal{B}_2) \xrightarrow{g} (E_3, \mathcal{B}_3)$$

alors

$$M_{\mathcal{B}_1\mathcal{B}_3}(g \circ f) = M_{\mathcal{B}_2\mathcal{B}_3}(g)M_{\mathcal{B}_1\mathcal{B}_2}(f) .$$

Clairement,  $\text{Id} = \text{Id} \circ \text{Id}$  et donc

$$M_{\mathcal{B}'\mathcal{B}}(\text{Id}) = M_{\mathcal{B}''\mathcal{B}}(\text{Id})M_{\mathcal{B}'\mathcal{B}''}(\text{Id}) ,$$

soit encore  $M = QR$ . Il reste à prouver que  $R$  est triangulaire supérieure à diagonale strictement positive. Cela tient au fonctionnement même du procédé d'orthonormalisation de Gram-Schmidt. Notons

$$\mathcal{B}' = (e'_1, \dots, e'_n) \text{ et } \mathcal{B}'' = (e''_1, \dots, e''_n) .$$

Le procédé d'orthonormalisation de Gram-Schmidt fonctionne de tel façon que

$$(i) \quad \forall k = 1, \dots, n, \text{Vect}(e''_1, \dots, e''_k) = \text{Vect}(e'_1, \dots, e'_k)$$

$$(ii) \quad \forall k = 1, \dots, n, e''_k = \lambda_k e'_k + \sum_{i=1}^{k-1} \alpha_{ik} e'_i$$

où  $\lambda_k > 0$  pour tout  $k$  puisqu'il s'agit en fait de l'inverse d'une norme d'un vecteur non nul. On déduit de (i) et (ii) que  $R = M_{\mathcal{B}'' \rightarrow \mathcal{B}'}$  est triangulaire supérieure à diagonale les  $\lambda_k^{-1}$ , donc à diagonale strictement positive. D'où le théorème.  $\square$

On aborde maintenant la notion de matrice symétrique définie positive.

**Définition 22.1.** Soit  $A$  une matrice symétrique  $n \times n$ . On dit que  $A$  est définie positive si pour tout vecteur colonne non nul  $X$  à  $n$  lignes,  ${}^tXAX > 0$ .

Une matrice symétrique est diagonalisable. On montre le théorème suivant.

**Théorème 22.2.** Soit  $A$  une matrice symétrique. Alors  $A$  est définie positive si et seulement si ses valeurs propres sont toutes strictement positives.

*Démonstration.* Soit  $A$  une matrice symétrique  $n \times n$ . En vertu du théorème spectral il existe  $P$  orthogonale  $n \times n$  et  $D$  une matrice diagonale  $n \times n$ , de diagonale les valeurs propres de  $A$  répétées avec leur multiplicité, telles que  ${}^tPAP = D$ . Pour  $X$  un vecteur colonne non nul à  $n$  lignes, on pose  $Y = {}^tPX$ . On a

$$\begin{aligned} {}^tYDY &= {}^t({}^tPX)D({}^tPX) \\ &= {}^tX(PD{}^tP)X \\ &= {}^tXAX . \end{aligned}$$

Comme  ${}^tP$  est inversible, il est clair que

$$\begin{aligned} \forall X \neq 0, \quad {}^tXAX &> 0 \\ \Leftrightarrow \forall Y \neq 0, \quad {}^tYDY &> 0 . \end{aligned}$$

Notons  $\lambda_1, \dots, \lambda_n$  les éléments qui composent la diagonale de  $D$  (donc les valeurs propres de  $A$  répétées avec leur multiplicité). Si on note  $y_1, \dots, y_n$  les éléments qui composent le vecteur colonne  $Y$ ,

$${}^tYDY = \sum_{i=1}^n \lambda_i y_i^2.$$

On voit donc que  ${}^tYDY > 0$  pour tout  $Y \neq 0$  si et seulement si  $\lambda_i > 0$  pour tout  $i = 1, \dots, n$ . En d'autres termes, en raison de l'équivalence ci-dessus,  $A$  est définie positive si et seulement si ses valeurs propres sont toutes strictement positives. Le théorème est démontré.  $\square$

On aborde maintenant la décomposition de Cholesky des matrices symétriques définies positives.

**Théorème 22.3** (Décomposition de Cholesky des matrices symétriques définies positives). *Soit  $A$  une matrice symétrique  $n \times n$  définie positive. Il existe alors une matrice  $R$  triangulaire supérieure à diagonale strictement positive telle que  $A = {}^tRR$ .*

*Démonstration.* Soit  $A$  symétrique  $n \times n$  définie positive. En vertu du théorème spectral et du théorème précédent il existe  $P$  orthogonale  $n \times n$ , et  $D$  diagonale  $n \times n$  à diagonale strictement positive constituée des valeurs propres de  $A$ , telles que  ${}^tPAP = D$ . On commence par montrer qu'il existe  $M$  inversible  $n \times n$  telle que  $A = {}^tMM$ . Notons  $\lambda_1, \dots, \lambda_n$  les éléments qui composent la diagonale de  $D$  (donc les valeurs propres de  $A$  répétées avec leur multiplicité). On définit la matrice diagonale  $D'$  dont la diagonale est constituée des  $\sqrt{\lambda_i}$ ,  $i = 1, \dots, n$ . On a alors  $D = D'D'$ . Puisque  $P$  est orthogonale,

$$A = PD{}^tP.$$

On pose  $\hat{P} = {}^tP$ . Comme  ${}^tD' = D'$  puisque  $D'$  est diagonale,

$$\begin{aligned} A &= {}^t\hat{P}D'D'\hat{P} \\ &= {}^t\hat{P}{}^tD'D'\hat{P} \\ &= {}^t(D'\hat{P})(D'\hat{P}) \\ &= {}^tMM \end{aligned}$$

si on pose  $M = D'\hat{P}$ . Les  $\sqrt{\lambda_i}$  sont tous strictement positifs, donc non nuls, et donc  $D'$  est inversible. Comme  $\hat{P}$  est aussi inversible,  $M$  est inversible. D'où l'affirmation ci-dessus. Maintenant, d'après la décomposition d'Iwasawa, Il existe  $Q$  orthogonale et  $R$  triangulaire supérieure à diagonale strictement positive telles que  $M = QR$ . Mais alors,

$$\begin{aligned} A &= {}^t(QR)(QR) \\ &= {}^tR{}^tQQR \\ &= {}^tRR \end{aligned}$$

puisque  $Q$  étant orthogonale on a que  ${}^tQQ = \text{Id}_n$  la matrice identité  $n \times n$ . D'où le théorème.  $\square$

On peut montrer que la matrice  $R$  triangulaire supérieure à diagonale strictement positive dans la décomposition de Cholesky est unique.

## 23. DÉCOMPOSITION SVD

On aborde ici la décomposition SVD (pour “singular value decomposition”) dans le cadre spécifique des matrices carrées inversibles (il y a plus général pour les matrices  $n \times p$ ). Le premier résultat que l’on démontre est le suivant.

**Lemme 23.1.** *Soit  $A \in \mathcal{M}_n(\mathbb{R})$  une matrice carrée inversible  $n \times n$ . La matrice  $M = {}^tAA$  est alors symétrique définie positive.*

*Proof.* Il est clair que  $M$  est symétrique puisque

$${}^tM = {}^tA({}^tA) = {}^tAA = M .$$

Par ailleurs, si  $X$  est un vecteur colonne de  $\mathbb{R}^n$ , alors

$$\begin{aligned} {}^tXMX &= {}^tX{}^tAAX \\ &= {}^t(AX)(AX) \\ &= \langle AX, AX \rangle , \end{aligned}$$

où  $\langle \cdot, \cdot \rangle$  est le produit scalaire euclidien de  $\mathbb{R}^n$ . Donc,  ${}^tXMX \geq 0$  et il y a égalité à zéro si et seulement si  $AX = 0$ , soit si et seulement si  $X = 0$  puisque  $A$  est inversible. Par suite,  ${}^tXMX > 0$  pour tout vecteur colonne non nul  $X$  de  $\mathbb{R}^n$ . Le lemme est démontré.  $\square$

Les valeurs singulières d’une matrice  $A$  sont les  $\sigma_i = \sqrt{\lambda_i}$ , où les  $\lambda_i$  sont les valeurs propres de la matrice symétrique  ${}^tAA$ . Dans notre cas les  $\sigma_i$ , répétés avec la multiplicité des  $\lambda_i$ , sont tous strictement positifs.

**Théorème 23.1** (Décomposition SVD des matrices carrées inversibles). *Soit  $A \in \mathcal{M}_n(\mathbb{R})$  une matrice carrée inversible  $n \times n$ . Il existe alors  $U, V \in \mathcal{M}_n(\mathbb{R})$  deux matrices carrées orthogonales  $n \times n$  pour lesquelles*

$$A = V\Sigma^tU ,$$

où  $\Sigma = \text{Diag}(\sigma_1, \dots, \sigma_n)$  est la matrice diagonale composée des valeurs singulières de  $A$  répétées avec la multiplicité des valeurs propres correspondantes.

*Proof.* On note  $(\mathbb{R}^n, \langle \cdot, \cdot \rangle)$  l’espace euclidien standard et  $\mathcal{B}$  la base canonique de  $\mathbb{R}^n$ . Cette base est orthonormée pour  $\langle \cdot, \cdot \rangle$ . La matrice  $M = {}^tAA$  est symétrique. Par théorème spectral elle est donc diagonalisable et il existe  $U$  orthogonale pour laquelle

$${}^tUMU = \text{Diag}(\sigma_1^2, \dots, \sigma_n^2) ,$$

la matrice diagonale composée des carrés des valeurs singulières de  $A$  répétées avec la multiplicité des valeurs propres correspondantes (on rappelle que  $\sigma_i^2 = \lambda_i$  pour tout  $i$ , où les  $\lambda_1, \dots, \lambda_n$  sont les valeurs propres de  $M$ ). Soit  $\mathcal{B}' = (u_1, \dots, u_n)$  la base de  $\mathbb{R}^n$  pour laquelle la matrice de passage  $M_{\mathcal{B} \rightarrow \mathcal{B}'}$  est donnée par  $M_{\mathcal{B} \rightarrow \mathcal{B}'} = U$ . Comme  $U$  est orthogonale,  $\mathcal{B}'$  est orthonormée. On note  $U_i$  le vecteur colonne de  $\mathbb{R}^n$  constitué des coordonnées des  $u_i$  dans la base canonique  $\mathcal{B}$ . On a  $MU_i = \lambda_i U_i$  pour tout  $i$  (les  $u_i$  sont des vecteurs propres associés aux valeurs propres  $\lambda_i$ ). On pose

$$V_i = \frac{1}{\sigma_i} AU_i ,$$

ce qui fait sens car les  $\sigma_i$  sont strictement positifs en vertu du Théorème 22.2 et du Lemme 23.1. Soit  $V$  la matrice dont les colonnes sont les  $V_1, \dots, V_n$ . Pour tous  $i, j$ ,

$$\begin{aligned} \langle V_i, V_j \rangle &= \frac{1}{\sigma_i \sigma_j} \langle AU_i, AU_j \rangle \\ &= \frac{1}{\sigma_i \sigma_j} {}^t(AU_i)(AU_j) \\ &= {}^tU_i({}^tAA)U_j \\ &= \frac{1}{\sigma_i \sigma_j} \langle U_i, {}^tAAU_j \rangle \\ &= \frac{\sigma_j}{\sigma_i} \langle U_i, U_j \rangle \end{aligned}$$

et donc, puisque  $\mathcal{B}'$  est orthonormée,  $\langle V_i, V_j \rangle = \delta_{ij}$  pour tous  $i, j$ , où les  $\delta_{ij} = 1$  si  $i = j$  et 0 sinon sont les symboles de Kroenecker. Donc  $\hat{\mathcal{B}} = (V_1, \dots, V_n)$  est orthonormée. On a  $V = M_{\mathcal{B} \rightarrow \hat{\mathcal{B}}}$  (matrice de passage de  $\mathcal{B}$  à  $\hat{\mathcal{B}}$ ) et donc  $V$  est elle-aussi une matrice orthogonale. Or  $AU_i = \sigma_i V_i$  pour tout  $i$ , soit encore

$$AU = V\Sigma$$

en termes de matrices. Puisque  $U^{-1} = {}^tU$ , il s'ensuit que  $A = V\Sigma{}^tU$ . Le théorème est démontré.  $\square$

#### 24. PREUVE DU THÉORÈME SPECTRAL - CAS GÉNÉRAL

Une des principales difficultés dans le cas général, où la dimension  $n$  est quelconque, est de montrer que les valeurs propres d'un endomorphisme symétrique  $f$  sont toutes réelles. On donne une preuve "à la main" de ce fait lorsque  $n = 3$ .

Si  $n = 3$ , le polynôme caractéristique  $P$  de  $f$  est de degré 3. Donc de degré impair, de sorte que  $P(x) \rightarrow \pm\infty$  lorsque  $x \rightarrow \pm\infty$ . De plus  $P$  est une fonction continue sur  $\mathbb{R}$ . On en déduit qu'il existe forcément  $\lambda_1 \in \mathbb{R}$  telle que  $P(\lambda_1) = 0$ . En particulier,  $f$  a au moins une valeur propre réelle. Soit  $e_1$  un vecteur propre (non nul) associé à cette valeur propre  $\lambda_1$ . On complète  $e_1$  par des vecteurs  $e_2, e_3$  pour obtenir une base orthonormale  $\mathcal{B} = (e_1, e_2, e_3)$  de  $E$ . Si  $D$  est la droite vectorielle de base  $e_1$ , et si  $F = Vect(e_2, e_3)$ , alors

$$E = D \oplus F$$

Par ailleurs, puisque  $f$  est symétrique, et puisque  $\mathcal{B}$  est orthonormale,

$$\begin{aligned} \langle f(e_2), e_1 \rangle &= \langle e_2, f(e_1) \rangle \\ &= \lambda_1 \langle e_2, e_1 \rangle \\ &= 0 \end{aligned}$$

De même

$$\langle f(e_3), e_1 \rangle = 0$$

et on en déduit que la restriction  $f|_F$  de  $f$  à  $F$  définit un endomorphisme de  $F$ . Notons  $\tilde{\mathcal{B}} = (e_2, e_3)$ . Alors  $\tilde{\mathcal{B}}$  est une base orthonormale de  $F$ , et on vérifie facilement que

$$M_{\mathcal{B}\mathcal{B}}(f) = \begin{pmatrix} \lambda_1 & 0 & 0 \\ 0 & M_{\tilde{\mathcal{B}}\tilde{\mathcal{B}}}(f|_F) & \\ 0 & & \end{pmatrix}$$

Comme  $f$  est symétrique,  $M_{\mathcal{B}\mathcal{B}}(f)$  est symétrique. On en déduit que  $M_{\tilde{\mathcal{B}}\tilde{\mathcal{B}}}(f|_F)$  est aussi symétrique, et donc que  $f|_F$  est un endomorphisme symétrique (car  $\tilde{\mathcal{B}}$  est une base orthonormale). Par ailleurs on voit que si  $P_E$  désigne le polynôme caractéristique de  $f$ , et si  $P_F$  désigne le polynôme caractéristique de  $f|_F$ , alors

$$P_E(X) = -(X - \lambda_1)P_F(X)$$

On ramène ainsi le cas de la dimension 3 à celui de la dimension 2, déjà traité ci-dessus. Puisque  $M_{\tilde{\mathcal{B}}\tilde{\mathcal{B}}}(f|_F)$  est symétrique,  $P_F$  a toutes ses racines réelles. Par suite,  $P_E$  a aussi toutes ses racines réelles.

On revient au cas  $n$  quelconque et on montre successivement que:

**Etape 1** - Les valeurs propres de  $f$  sont réelles

**Etape 2** - Les espaces propres de  $f$  sont 2 à 2 orthogonaux

**Etape 3** -  $E$  est somme (directe) des espaces propres de  $f$

Par suite  $f$  est diagonalisable, qui plus est dans une base orthonormale de  $E$ .

**Etape 1:** On affirme que si  $f \in \text{End}(E)$  est un endomorphisme symétrique de  $E$ , alors les valeurs propres de  $f$  sont toutes réelles au sens où il existe des réels  $\lambda_1, \dots, \lambda_p$  et des entiers  $k_1, \dots, k_p$  pour lesquels le polynôme caractéristique  $P$  de  $f$  se décompose en

$$P(X) = (-1)^n \prod_{i=1}^p (X - \lambda_i)^{k_i}$$

Pour démontrer cette affirmation, on fait un petit détour par les nombres complexes. Le tout est assez naturel dans la mesure où  $\mathbb{C}$  est un corps algébriquement clos, à savoir tel que les polynômes sur  $\mathbb{C}$  ont toutes leurs racines dans  $\mathbb{C}$ . Le principe de preuve est alors le suivant: par le choix d'une base orthonormale  $\mathcal{B}$  de  $E$ ,

1 - on définit un espace vectoriel complexe  $E^c$  qui étend  $E$  des réels aux complexes,

2 - on définit un endomorphisme  $f^c$  de  $E^c$  qui étend  $f$ ,

3 - on construit le tout pour que  $f$  et  $f^c$  aient même polynôme caractéristique, et

4 - on montre que les valeurs propres de  $f^c$ , à priori complexes, sont en fait réelles.

Soit donc  $\mathcal{B} = (e_1, \dots, e_n)$  une base orthonormale de  $E$ . On note  $E^c$  le "compléxifié" de  $E$  défini par

$$E^c = \left\{ \sum_{i=1}^n z_i e_i, z_i \in \mathbb{C} \right\}$$

que l'on munit des opérations

$$\begin{aligned} \left( \sum_{i=1}^n z_i e_i \right) + \left( \sum_{i=1}^n \hat{z}_i e_i \right) &= \sum_{i=1}^n (z_i + \hat{z}_i) e_i \\ \lambda \left( \sum_{i=1}^n z_i e_i \right) &= \sum_{i=1}^n (\lambda z_i) e_i \end{aligned}$$

pour  $\lambda \in \mathbb{C}$ . Alors  $E^c$  muni de ces deux opérations est un espace vectoriel complexe (même définition que les espaces vectoriels réels, mais on remplace  $\mathbb{R}$  par  $\mathbb{C}$ ). Clairement,  $E$  est un sous ensemble de  $E^c$ . On étend alors  $f$  de  $E$  à  $E^c$  en

définissant un endomorphisme  $f^c$  de  $E^c$  (là encore, même définition que pour le cas réel, mais on remplace  $\mathbb{R}$  par  $\mathbb{C}$ ). On pose

$$f^c \left( \sum_{i=1}^n z_i e_i \right) = \sum_{i=1}^n z_i f(e_i)$$

Si  $M_{\mathcal{B}\mathcal{B}}(f) = (a_{ij})$ , de sorte que  $f(e_j) = \sum_i a_{ij} e_i$ , on a encore que

$$f^c \left( \sum_{i=1}^n z_i e_i \right) = \sum_{i=1}^n \left( \sum_{j=1}^n a_{ij} z_j \right) e_i$$

La restriction de  $f^c$  à  $E$  vaut  $f$ . On définit maintenant la forme hermitienne  $\langle \cdot, \cdot \rangle_c$  sur  $E^c$  par

$$\left\langle \sum_{i=1}^n z_i e_i, \sum_{i=1}^n z_i^2 e_i \right\rangle_c = \sum_{i=1}^n z_i^1 \overline{z_i^2}$$

où si  $z \in \mathbb{C}$ ,  $\bar{z}$  est le conjugué de  $z$ . La restriction de  $\langle \cdot, \cdot \rangle_c$  à  $E \times E$  coïncide avec  $\langle \cdot, \cdot \rangle$ . Comme  $f$  est symétrique, les  $a_{ij}$  sont symétriques au sens où  $a_{ij} = a_{ji}$  (cf. section 1). Par définition même de tous ces objets, on a alors que

$$\left\langle f^c \left( \sum_{i=1}^n z_i^1 e_i \right), \sum_{i=1}^n z_i^2 e_i \right\rangle_c = \left\langle \sum_{i=1}^n z_i^1 e_i, f^c \left( \sum_{i=1}^n z_i^2 e_i \right) \right\rangle_c$$

pour tous  $\sum_{i=1}^n z_i^1 e_i, \sum_{i=1}^n z_i^2 e_i \in E^c$ . On trouve en effet que

$$\begin{aligned} \left\langle f^c \left( \sum_{i=1}^n z_i^1 e_i \right), \sum_{i=1}^n z_i^2 e_i \right\rangle_c &= \sum_{i,j=1}^n a_{ij} z_i^1 \overline{z_j^2} \\ &= \sum_{i,j=1}^n z_i^1 \overline{a_{ij} z_j^2} \\ &= \left\langle \sum_{i=1}^n z_i^1 e_i, f^c \left( \sum_{i=1}^n z_i^2 e_i \right) \right\rangle_c \end{aligned}$$

Pour en finir avec ces généralités,  $\mathcal{B}$  est encore une base de  $E^c$  (c'est fait pour), et on a que

$$M_{\mathcal{B}\mathcal{B}}(f^c) = M_{\mathcal{B}\mathcal{B}}(f)$$

Comme dans le cas réel, si  $E^c$  est un espace vectoriel complexe de dimension finie  $n$ , si  $\mathcal{B}$  est une base de  $E^c$ , et si  $f^c$  est un endomorphisme de  $E^c$ , un nombre complexe  $\lambda$  est une valeur propre de  $f^c$ , au sens où il existe  $x \in E^c \setminus \{0\}$  pour lequel  $f^c(x) = \lambda x$ , si et seulement si

$$\det(M_{\mathcal{B}\mathcal{B}}(f^c) - \lambda Id_n) = 0$$

où  $Id_n$  est la matrice identité  $n \times n$ . En d'autres termes,  $\lambda \in \mathbb{C}$  est valeur propre de  $f^c$  si et seulement si elle annule le polynôme caractéristique de  $f^c$ . Dans notre cas,  $M_{\mathcal{B}\mathcal{B}}(f^c) = M_{\mathcal{B}\mathcal{B}}(f)$ , et le polynôme caractéristique de  $f^c$  vaut donc le polynôme caractéristique de  $f$ . Un polynôme complexe a toutes ses racines dans  $\mathbb{C}$  car  $\mathbb{C}$  est algébriquement clos. Notons  $\lambda_1, \dots, \lambda_p$  les racines dans  $\mathbb{C}$  du polynôme caractéristique de  $f^c$ , ou encore de  $f$  (ce qui revient au même). Démontrer l'étape 1 consiste à démontrer que les  $\lambda_1, \dots, \lambda_p$  sont en fait des réels. Pour  $j = 1, \dots, p$ ,

notons  $x_j = \sum_{i=1}^n z_i^j e_i \in E^c \setminus \{0\}$  un vecteur propre de  $f^c$  associé à la valeur propre  $\lambda_j$ . On a vu que

$$\langle f^c(\sum_{i=1}^n z_i^j e_i), \sum_{i=1}^n z_i^j e_i \rangle_c = \langle \sum_{i=1}^n z_i^j e_i, f^c(\sum_{i=1}^n z_i^j e_i) \rangle_c$$

Or, par définition même de  $\langle \cdot, \cdot \rangle_c$ , et puisque  $f(x_j) = \lambda_j x_j$ , on a que

$$\begin{aligned} \langle f^c(\sum_{i=1}^n z_i^j e_i), \sum_{i=1}^n z_i^j e_i \rangle_c &= \lambda_j \sum_{i=1}^n |z_i^j|^2 \\ \langle \sum_{i=1}^n z_i^j e_i, f^c(\sum_{i=1}^n z_i^j e_i) \rangle_c &= \overline{\lambda_j} \sum_{i=1}^n |z_i^j|^2 \end{aligned}$$

Sachant que  $\sum_{i=1}^n |z_i^j|^2 \neq 0$  puisque  $x_j \neq 0$ , on obtient que  $\lambda_j = \overline{\lambda_j}$ , et donc que  $\lambda_j$  est réel. Comme  $j$  est quelconque dans  $\{1, \dots, p\}$ , il suit que les valeurs propres de  $f$  sont toutes réelles.  $\square$

**Etape 2:** On affirme que les espaces propres d'un endomorphisme symétrique sont deux à deux orthogonaux. Il s'agit là d'une affirmation simple à démontrer. Pour le voir, on note  $\lambda$  et  $\mu$  deux valeurs propres distinctes de  $f$ , et  $x$  et  $y$  deux vecteurs propres de  $f$  respectivement associés aux valeurs propres  $\lambda$  et  $\mu$ . Alors, puisque  $f$  est symétrique,

$$\langle f(x), y \rangle = \langle x, f(y) \rangle$$

et donc

$$(\lambda - \mu) \langle x, y \rangle = 0$$

puisque  $f(x) = \lambda x$  et  $f(y) = \mu y$ . Comme  $\lambda \neq \mu$ , c'est que  $\langle x, y \rangle = 0$ . L'affirmation suit.  $\square$

**Etape 3:** On affirme que si  $f \in \text{End}(E)$  est un endomorphisme symétrique de  $E$ , alors  $E$  est la somme directe des espaces propres de  $f$ . On démontre cette affirmation par récurrence totale sur la dimension  $n$  de  $E$ . On sait déjà (cf. le cours d'algèbre linéaire) que la somme des espaces propres est toujours directe. La seule chose qu'il y ait à démontrer est que cette somme vaut l'espace tout entier.

**3.1.** On démontre l'affirmation pour  $n = 2$  (on l'a déjà fait, mais répétons tout de même l'argument). Soit donc  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension 2 muni d'un produit scalaire, et  $f \in \text{End}(E)$  un endomorphisme symétrique de  $E$ . Si  $f$  a deux valeurs propres distinctes  $\lambda_1$  et  $\lambda_2$ , alors on sait déjà, cf. encore le chapitre 3, que  $E$  est somme (directe) des espaces propres de  $f$ . Si maintenant  $f$  a une seule valeur propre  $\lambda_1$ , et si  $E_1$  est le sous espace propre associé à  $\lambda_1$ , on veut montrer que  $E = E_1$ . Notons  $e_1$  un vecteur propre de  $f$  (pour  $\lambda_1$ ). Sans perdre en généralité, quitte à diviser  $e_1$  par sa norme, on peut supposer que  $\|e_1\| = 1$  (où  $\|\cdot\|$  est la norme associée au produit scalaire). On complète  $e_1$  par un vecteur  $e_2$  pour obtenir une base orthonormale  $(e_1, e_2)$  de  $E$ . Alors

$$\begin{aligned} \langle f(e_2), e_1 \rangle &= \langle e_2, f(e_1) \rangle \\ &= \lambda_1 \langle e_2, e_1 \rangle \end{aligned}$$

puisque d'une part  $f$  est symétrique, et d'autre part  $e_1$  est un vecteur propre de  $f$  pour  $\lambda_1$ . Par suite

$$\langle f(e_2), e_1 \rangle = 0$$

et donc la première coordonnée de  $f(e_2)$  dans  $(e_1, e_2)$  est nulle. On en déduit que

$$f(e_2) = \lambda e_2$$

pour un certain  $\lambda$  réel. Il suit que  $\lambda$  est valeur propre de  $f$ , et que  $e_2$  est un vecteur propre de  $f$ . Donc  $\lambda = \lambda_1$ , puisque nous avons supposé que  $f$  a une seule valeur propre, et on a tout à la fois que  $e_1 \in E_1$  et que  $e_2 \in E_1$ , de sorte que  $E = E_1$ . L'affirmation est vraie lorsque  $n = 2$ .

**3.2.** On suppose que l'affirmation est vraie pour tout espace vectoriel  $E$  de dimension  $k \leq n$  et tout endomorphisme symétrique de  $E$ , et on démontre qu'elle est encore vraie pour tout espace vectoriel  $E$  de dimension  $n+1$  et tout endomorphisme symétrique de  $E$ . On note  $E$  un  $\mathbb{R}$ -espace vectoriel de dimension finie  $n+1$ ,  $\langle \cdot, \cdot \rangle$  le produit scalaire sur  $E$ ,  $f \in \text{End}(E)$  un endomorphisme symétrique de  $E$ , et  $\lambda_1, \dots, \lambda_p$  les valeurs propres de  $f$ . Soit  $E_1$  le sous espace propre associé à la valeur propre  $\lambda_1$ . On note  $F = E_1^\perp$  l'orthogonal de  $E_1$ . Alors

$$E = E_1 \oplus F$$

comme on l'a démontré au chapitre précédent. En particulier,  $\dim F \leq n$  (car  $\dim E_1 \geq 1$ ). Une affirmation simple est que la restriction  $f|_F$  de  $f$  à  $F$  définit un endomorphisme de  $F$ . En d'autres termes que si  $x \in F$  alors  $f(x) \in F$ . Pour le voir on considère une base orthonormale  $(e_1, \dots, e_{n+1})$  de  $E$  ayant la propriété suivante: les premiers vecteurs de la base, disons les  $(e_1, \dots, e_k)$ , forment une base de  $E_1$ . Alors, on l'a déjà vu au chapitre précédent, les  $(e_{k+1}, \dots, e_n)$  forment une base de  $F$ . Dès lors, pour  $i = 1, \dots, k$  et pour  $j = k+1, \dots, n$ ,

$$\begin{aligned} \langle f(e_j), e_i \rangle &= \langle e_j, f(e_i) \rangle \\ &= \lambda_1 \langle e_j, e_i \rangle \\ &= 0 \end{aligned}$$

puisque  $f$  est symétrique,  $f(e_i) = \lambda_1 e_i$ , et la base  $(e_1, \dots, e_n)$  est orthonormale. Une autre façon de dire les choses est que pour tout  $j = k+1, \dots, n$ ,

$$f(e_j) \in \text{Vect}(e_{k+1}, \dots, e_n)$$

et donc que  $f|_F$  réalise un endomorphisme de  $F$ . Par hypothèse de récurrence, puisque  $\dim F \leq n$ ,  $F$  est somme directe des espaces propres de  $f|_F$ . Afin de finir de démontrer l'affirmation de l'étape 3, puisque  $E = E_1 \oplus F$ , il nous reste donc à montrer que

- (1) Valeurs propres de  $f|_F = \{\lambda_2, \dots, \lambda_p\}$ , à savoir les valeurs propres de  $f$  hors  $\lambda_1$ ,  
et que pour tout  $i = 2, \dots, p$ ,

- (2) Espace propre de  $f|_F$  associé à  $\lambda_i$  = Espace propre de  $f$  associé à  $\lambda_i$ .

Pour démontrer (1) et (2) on procède comme suit. Tout d'abord on remarque que les inclusions

- (i) Valeurs propres de  $f|_F \subset \{\lambda_2, \dots, \lambda_p\}$   
(ii) Espace propre de  $f|_F$  associé à  $\lambda_i \subset$  Espace propre de  $f$  associé à  $\lambda_i$

sont immédiates. Réciproquement, notons  $\lambda$  un des  $\lambda_i$  pour  $i = 2, \dots, p$ . Soit de plus  $x$  un vecteur propre de  $f$  associé à la valeur propre  $\lambda$ . On a vu à l'étape précédente que les espaces propres de  $f$  sont deux à deux orthogonaux. Donc  $x$  est orthogonal à tous les vecteurs de  $E_1$ , ce qui signifie que  $x \in F$ . Donc  $x$  est aussi un

vecteur propre de  $f|_F$ , et  $\lambda$  est valeur propre de  $f|_F$ . On en déduit, puisque  $x$  est un vecteur propre quelconque de  $f$  associé à l'une quelconque des valeurs propres  $\lambda_2, \dots, \lambda_p$  que les inclusions réciproques de (i) et (ii) sont vraies. Donc (1) et (2) sont vraies, et l'étape 3 est démontrée.  $\square$

Reste pour démontrer le théorème 2.1 à collecter les informations des étapes 1 à 3. D'après l'étape 1, si  $f$  est un endomorphisme symétrique sur un  $\mathbb{R}$ -espace vectoriel de dimension finie, alors  $f$  a toutes ses valeurs propres réelles. D'après l'étape 2, les espaces propres de  $f$  sont deux à deux orthogonaux. D'après l'étape 3,  $E$  est somme directe des espaces propres de  $f$ . On en déduit que  $f$  est diagonalisable. De plus, si les  $E_i$ ,  $i = 1, \dots, p$ , sont les espaces propres de  $f$ , en notant  $\mathcal{B}_i$  une base orthonormale de  $E_i$ , on obtient que

$$\mathcal{B} = (\mathcal{B}_1, \dots, \mathcal{B}_p)$$

est une base orthonormale de  $E$ . Cette base diagonalise  $f$ : en d'autres termes, la matrice  $M_{\mathcal{B}\mathcal{B}}(f)$  de représentation de  $f$  dans  $\mathcal{B}$  est diagonale. Le théorème spectral en dimension quelconque est démontré.

## 25. INÉGALITÉ DE KANTOROVITCH

On démontre ici l'inégalité de Kantorovitch qui est souvent utilisée pour mesurer la vitesse de convergence dans les méthodes de gradient à pas optimal. Les vecteurs  $x$  sont ici compris comme des vecteurs colonnes.

**Théorème 25.1** (Inégalité de Kantorovitch). *Soit  $A \in \mathcal{M}_n(\mathbb{R})$  une matrice  $n \times n$  symétrique définie positive. On note  $\lambda_1 \geq \dots \geq \lambda_n$  les valeurs propres de  $A$  répétées avec leur multiplicité et rangées par ordre décroissant. Alors*

$$\|x\|^4 \leq \langle Ax, x \rangle \langle A^{-1}x, x \rangle \leq \frac{1}{4} \left( \sqrt{\frac{\lambda_1}{\lambda_n}} + \sqrt{\frac{\lambda_n}{\lambda_1}} \right)^2 \|x\|^4 \quad (25.1)$$

pour tout  $x \in \mathbb{R}^n$ .

*Démonstration.* Comme  $A$  est définie positive les  $\lambda_i$  sont tous strictement positifs (cf. Théorème 22.2). Pour démontrer la double inégalité (25.1), par homogénéité d'ordre 4 en  $x$ , il suffit de la démontrer pour les  $x$  de norme 1. Soit  $P$  orthogonale donnée par le théorème spectral, donc telle que

$${}^t P A P = D$$

où  $D = \text{Diag}(\lambda_1, \dots, \lambda_n)$  est la matrice diagonale dont la diagonale est constituée des  $\lambda_1, \dots, \lambda_n$ . Comme  $P$  est orthogonale, si  $y = Px$  et  $\|x\| = 1$ , alors  $\|y\| = 1$ . De plus,

$$\begin{aligned} \langle Ax, x \rangle &= \langle APy, Py \rangle \\ &= \langle {}^t P A P y, y \rangle \\ &= \langle Dy, y \rangle \end{aligned}$$

et, de même,  $\langle A^{-1}x, x \rangle = \langle D^{-1}y, y \rangle$  (puisque  $P^{-1} = {}^t P$ ). Démontrer (25.1) revient donc à montrer que pour tout  $y$  de norme 1,

$$1 \leq \langle Dy, y \rangle \langle D^{-1}y, y \rangle \leq \frac{1}{4} \left( \sqrt{\frac{\lambda_1}{\lambda_n}} + \sqrt{\frac{\lambda_n}{\lambda_1}} \right)^2. \quad (25.2)$$

Posons  $\alpha_i = y_i^2$ . Alors  $\sum_{i=1}^n \alpha_i = 1$  et

$$\begin{aligned}\langle Dy, y \rangle &= \sum_{i=1}^n \alpha_i \lambda_i , \\ \langle D^{-1}y, y \rangle &= \sum_{i=1}^n \frac{\alpha_i}{\lambda_i} .\end{aligned}$$

Soit  $\Lambda > 0$  fixé. On note  $f_\Lambda$  la fonction définie pour  $x > 0$  par

$$f_\Lambda(x) = x + \frac{\Lambda}{x} .$$

Un tableau de variation facile montre que la fonction  $f_\Lambda$  atteint un minimum en  $x = \sqrt{\Lambda}$ . En particulier  $f_1$  atteint un minimum en 1 qui vaut 2. Donc  $x + \frac{1}{x} \geq 2$  pour tout  $x > 0$ . On en déduit, puisque les  $\alpha_i$  sont positifs et les  $\lambda_i$  strictement positifs, que

$$\begin{aligned}\langle Dy, y \rangle \langle D^{-1}y, y \rangle &= \sum_{i,j=1}^n \alpha_i \alpha_j \frac{\lambda_i}{\lambda_j} \\ &= \sum_{i=1}^n \alpha_i^2 + \sum_{i < j} \left( \frac{\lambda_i}{\lambda_j} + \frac{\lambda_j}{\lambda_i} \right) \alpha_i \alpha_j \\ &\geq \sum_{i=1}^n \alpha_i^2 + 2 \sum_{i < j} \alpha_i \alpha_j \\ &= \left( \sum_{i=1}^n \alpha_i \right)^2 = 1 .\end{aligned}$$

L'inégalité de gauche dans (25.2) est démontrée. Pour ce qui est de l'inégalité de droite, la plus importante, on remarque aussi à partir du tableau de variation de  $f_\Lambda$  (en considérant les trois cas  $\lambda_n \leq \lambda_1 \leq \sqrt{\Lambda}$ ,  $\lambda_n \leq \sqrt{\Lambda} \leq \lambda_1$  et  $\sqrt{\Lambda} \leq \lambda_n \leq \lambda_1$ ) que pour tout  $x \in [\lambda_n, \lambda_1]$ ,

$$f_\Lambda(x) \leq \max(f_\Lambda(\lambda_1), f_\Lambda(\lambda_n)) .$$

On pose  $\Lambda = \lambda_1 \lambda_n$ . Alors

$$f_\Lambda(\lambda_1) = f_\Lambda(\lambda_n) = \lambda_1 + \lambda_n$$

et ainsi

$$\begin{aligned}\langle Dy, y \rangle + \lambda_1 \lambda_n \langle D^{-1}y, y \rangle &= \sum_{i=1}^n \alpha_i f_\Lambda(\lambda_i) \\ &\leq (\lambda_1 + \lambda_n) \sum_{i=1}^n \alpha_i = \lambda_1 + \lambda_n .\end{aligned}$$

Par suite

$$\langle Dy, y \rangle \langle D^{-1}y, y \rangle \leq \frac{\Theta(\lambda_1 + \lambda_n - \Theta)}{\lambda_1 \lambda_n} ,$$

où  $\Theta = \langle Dy, y \rangle$ . La fonction  $\Phi$  définie pour  $\Theta > 0$  par

$$\Phi(\Theta) = \frac{\Theta(\lambda_1 + \lambda_n - \Theta)}{\lambda_1 \lambda_n}$$

atteint un maximum en  $\Theta = \frac{\lambda_1 + \lambda_n}{2}$  et

$$\begin{aligned}\Phi\left(\frac{\lambda_1 + \lambda_n}{2}\right) &= \frac{(\lambda_1 + \lambda_n)^2}{4\lambda_1\lambda_n} \\ &= \frac{1}{4}\left(\sqrt{\frac{\lambda_1}{\lambda_n}} + \sqrt{\frac{\lambda_n}{\lambda_1}}\right)^2.\end{aligned}$$

D'où l'inégalité de droite dans (25.2). Le théorème est démontré.  $\square$

## 26. CARACTÉRISATION DES VALEURS PROPRES

Soit  $A$  une matrice symétrique carrée  $n \times n$ . On définit le quotient de Rayleigh de  $A$  par

$$\mathcal{R}(X) = \frac{{}^t X A X}{{}^t X X}$$

pour tout vecteur colonne à  $n$  lignes  $X \neq 0$ . Si  $f \in \text{End}(\mathbb{R}^n)$  est l'endomorphisme de  $\mathbb{R}^n$  dont la matrice de représentation dans la base canonique  $\mathcal{B}$  de  $\mathbb{R}^n$  est  $A$ , on a encore que

$$\mathcal{R}(X) = \frac{\langle f(x), x \rangle}{\|x\|^2},$$

où  $\langle \cdot, \cdot \rangle$  est le produit scalaire euclidien,  $\|\cdot\|$  est la norme euclidienne et  $x \in \mathbb{R}^n$  est le vecteur dont les coordonnées dans  $\mathcal{B}$  sont données par les composantes de  $X$ . On notera dans la suite  $\mathcal{R}(x)$  ou  $\mathcal{R}(X)$  selon le contexte.

**Théorème 26.1.** *Soit  $A$  une matrice symétrique carrée  $n \times n$ . On note maintenant  $\lambda_1 > \dots > \lambda_p$  les valeurs propres distinctes de  $A$  rangées par ordre décroissant. Alors*

$$\lambda_1 = \max_{x \in \mathbb{R}^n \setminus \{0\}} \mathcal{R}(x)$$

*et le maximum est réalisé par les  $x \in E_1$ ,  $x \neq 0$ , où  $E_1$  est l'espace propre associé à la valeur propre  $\lambda_1$ . Ensuite,*

$$\lambda_2 = \max_{x \in E_1^\perp \setminus \{0\}} \mathcal{R}(x)$$

*et le maximum est réalisé par les  $x \in E_2$ ,  $x \neq 0$ , où  $E_2$  est l'espace propre associé à la valeur propre  $\lambda_2$ . Et ainsi de suite de sorte que pour  $k \in \{0, \dots, p-1\}$ ,*

$$\lambda_{k+1} = \max_{x \in (E_1 \oplus \dots \oplus E_k)^\perp \setminus \{0\}} \mathcal{R}(x)$$

*et le maximum est réalisé par les  $x \in E_{k+1}$ ,  $x \neq 0$ , où les  $E_i$  sont les espaces propres associés aux  $\lambda_i$ .*

*Démonstration.* On note  $f \in \text{End}(\mathbb{R}^n)$  l'endomorphisme de  $\mathbb{R}^n$  dont la matrice de représentation dans la base canonique  $\mathcal{B}$  de  $\mathbb{R}^n$  est  $A$ . La base  $\mathcal{B}$  est alors orthonormée (pour le produit scalaire euclidien) et  $A$  est symétrique. Par théorème spectral il existe  $\tilde{\mathcal{B}}$  une base orthonormée qui diagonalise  $f$ . On note  $\tilde{\lambda}_1 \geq \tilde{\lambda}_2 \geq \dots \geq \tilde{\lambda}_n$  les valeurs propres de  $A$  (donc aussi de  $f$ ) répétées avec leur multiplicité et rangées en ordre décroissant. Si on réécrit le quotient de Rayleigh dans  $\tilde{\mathcal{B}}$ , si on note  $\tilde{x}_i$  les coordonnées dans  $\tilde{\mathcal{B}}$  d'un vecteur  $x$  de  $\mathbb{R}^n$ , alors

$$\mathcal{R}(x) = \frac{\sum_{i=1}^n \tilde{\lambda}_i \tilde{x}_i^2}{\sum_{i=1}^n \tilde{x}_i^2}.$$

Clairement  $\tilde{\lambda}_1 = \lambda_1$  et donc, puisque  $\tilde{\lambda}_i \leq \tilde{\lambda}_1$  pour tout  $i$ ,  $\mathcal{R}(x) \leq \lambda_1$  pour tout  $x$ . Par ailleurs, si  $x \in E_1$  et  $x \neq 0$ , alors  $\mathcal{R}(x) = \lambda_1$  puisqu'alors  $f(x) = \lambda_1 x$ . D'où

$$\lambda_1 = \max_{x \in \mathbb{R}^n \setminus \{0\}} \mathcal{R}(x)$$

et le fait que le maximum est réalisé par les  $x \in E_1$ ,  $x \neq 0$ . Soit  $k$  la dimension de  $E_1$ . Pour  $x \in E_1^\perp$ ,

$$f(x) = \sum_{i=k+1}^n \tilde{x}_i f(\tilde{e}_i) = \sum_{i=1}^n \tilde{\lambda}_i \tilde{x}_i \tilde{e}_i$$

et donc

$$\mathcal{R}(x) = \frac{\sum_{i=k+1}^n \tilde{\lambda}_i \tilde{x}_i^2}{\sum_{i=k+1}^n \tilde{x}_i^2}.$$

On a  $\tilde{\lambda}_{k+1} = \lambda_2$  par notre choix de notation, et ainsi  $\mathcal{R}(x) \leq \lambda_2$  pour  $x \in E_1^\perp$ ,  $x \neq 0$ . De plus, si  $x \in E_2$  alors  $f(x) = \lambda_2 x$  et donc  $\mathcal{R}(x) = \lambda_2$ . On a  $E_2 \subset E_1^\perp$  et on a donc bien là encore la relation voulue. La relation générale se démontre de la même façon.  $\square$

On peut aussi partir de la plus petite valeur propre  $\lambda_p$  puis remonter. On a alors

$$\lambda_p = \min_{x \in \mathbb{R}^n \setminus \{0\}} \mathcal{R}(x)$$

et le minimum est réalisé par les  $x \in E_p$ ,  $x \neq 0$ , où  $E_p$  est l'espace propre associé à la valeur propre  $\lambda_p$ . Ensuite,

$$\lambda_{p-1} = \min_{x \in E_p^\perp \setminus \{0\}} \mathcal{R}(x)$$

et le minimum est réalisé par les  $x \in E_{p-1}$ ,  $x \neq 0$ , où  $E_{p-1}$  est l'espace propre associé à la valeur propre  $\lambda_{p-1}$ , et ainsi de suite.

## 27. CARACTÉRISATION DE COURANT-FISCHER

On présente ici la caractérisation de Courant-Fischer des valeurs propres d'une matrice symétrique. Dans cette caractérisation, on note toujours les valeurs propres par ordre décroissant mais cette fois-ci répétées avec leur multiplicité.

**Théorème 27.1.** *Soit  $A$  une matrice symétrique carrée  $n \times n$  et  $\lambda_1 \geq \dots \geq \lambda_n$  les valeurs propres de  $A$  répétées avec leur multiplicité et rangées par ordre décroissant. Pour  $k \in \{1, \dots, n\}$  on note  $\mathcal{F}_k$  l'ensemble des sous espaces vectoriels de  $\mathbb{R}^n$  qui sont de dimension  $k$  et  $\mathcal{S}$  la sphère unité de  $\mathbb{R}^n$  qui est donc constituée des vecteurs de  $\mathbb{R}^n$  qui sont de norme 1. On note  $f \in \text{End}(\mathbb{R}^n)$  l'endomorphisme de  $\mathbb{R}^n$  dont la matrice de représentation dans la base canonique  $\mathcal{B}$  de  $\mathbb{R}^n$  est  $A$ . Alors*

$$\begin{aligned} \lambda_k &= \inf_{F \in \mathcal{F}_{n+1-k}} \sup_{x \in F \cap \mathcal{S}} \langle f(x), x \rangle \\ &= \sup_{F \in \mathcal{F}_k} \inf_{x \in F \cap \mathcal{S}} \langle f(x), x \rangle \end{aligned}$$

pour tout  $k \in \{1, \dots, n\}$ , où  $\langle \cdot, \cdot \rangle$  est le produit scalaire euclidien.

*Démonstration.* La base  $\mathcal{B}$  est alors orthonormée et  $A$  est symétrique. Donc  $f$  est auto adjoint et, par théorème spectral, il existe  $\tilde{\mathcal{B}} = (\tilde{e}_1, \dots, \tilde{e}_n)$  une base orthonormée qui diagonalise  $f$ . Donc  $f(\tilde{e}_i) = \lambda_i \tilde{e}_i$  pour tout  $i = 1, \dots, n$  et les

valeurs propres de  $f$  répétées avec leur multiplicité sont bien sûr les  $\lambda_1, \dots, \lambda_n$ . Pour  $k \in \{1, \dots, n\}$  fixé, on note

$$\begin{cases} \Lambda_-^k(f) = \inf_{F \in \mathcal{F}_{n+1-k}} \sup_{x \in F \cap \mathcal{S}} \langle f(x), x \rangle \\ \Lambda_+^k(f) = \sup_{F \in \mathcal{F}_k} \inf_{x \in F \cap \mathcal{S}} \langle f(x), x \rangle \end{cases}$$

et on note encore  $E_k = \text{Vect}(\tilde{e}_1, \dots, \tilde{e}_k)$  et  $F_k = \text{Vect}(\tilde{e}_k, \dots, \tilde{e}_n)$ . Si  $x \in E_k$  et  $y \in F_k$  alors

$$\langle f(x), x \rangle = \sum_{i=1}^k \lambda_i x_i^2 \geq \lambda_k \|x\|^2 \quad \text{et} \quad \langle f(y), y \rangle = \sum_{i=k}^n \lambda_i y_i^2 \leq \lambda_k \|y\|^2$$

où les  $x_i$  et  $y_i$  sont les coordonnées de  $x$  et de  $y$  dans la base  $\tilde{B}$  et  $\|\cdot\|$  est la norme euclidienne. Dans les deux cas il y a égalité lorsque  $x = \tilde{e}_k$  ou  $y = \tilde{e}_k$ . On en déduit que

$$\inf_{x \in E_k \cap \mathcal{S}} \langle f(x), x \rangle = \lambda_k \leq \Lambda_+^k(f) \quad \text{et} \quad \sup_{x \in F_k \cap \mathcal{S}} \langle f(x), x \rangle = \lambda_k \geq \Lambda_-^k(f).$$

Par suite

$$\Lambda_-^k(f) \leq \lambda_k \leq \Lambda_+^k(f).$$

Soit maintenant  $F \in \mathcal{F}_{n+1-k}$ . Automatiquement  $F \cap E_k \neq \{0\}$  car  $\dim(F) = n+1-k$ ,  $\dim(E_k) = k$  et  $k + (n+1-k) > n$ . Soit  $x_0 \in F \cap E_k$ . Quitte à renormaliser on peut supposer que  $x_0 \in \mathcal{S}$ . On a alors

$$\sup_{x \in F \cap \mathcal{S}} \langle f(x), x \rangle \geq \langle f(x_0), x_0 \rangle \geq \inf_{x \in E_k \cap \mathcal{S}} \langle f(x), x \rangle$$

et il suit facilement en passant à l'infimum sur  $F \in \mathcal{F}_{n+1-k}$  dans le membre de gauche de cette double inégalité que  $\Lambda_-^k(f) \geq \lambda_k$ . Donc  $\Lambda_-^k(f) = \lambda_k$ . De même, si  $F \in \mathcal{F}_k$  alors  $F \cap F_k \neq \{0\}$ . On prend  $x_0 \in F \cap F_k$  avec  $x_0 \in \mathcal{S}$ . On a alors que

$$\sup_{x \in F_k \cap \mathcal{S}} \langle f(x), x \rangle \geq \langle f(x_0), x_0 \rangle \geq \inf_{x \in F \cap \mathcal{S}} \langle f(x), x \rangle$$

et en prenant le supremum sur  $F \in \mathcal{F}_k$  dans le membre de droite de cette double inégalité on obtient que  $\lambda_k \geq \Lambda_+^k(f)$ . Au final,  $\Lambda_-^k(f) = \Lambda_+^k(f) = \lambda_k$  et le théorème est démontré.  $\square$

## 28. PRINCIPE DU MAXIMUM DE KY FAN

Le principe du maximum de Ky Fan présente une autre vision de la caractérisation des valeurs propres qui généralise la caractérisation du Théorème 26.1. Si  $k = 1$  on retrouve exactement la caractérisation classique du  $\lambda_1$  du Théorème 26.1.

**Théorème 28.1** (Principe du maximum de Ky-Fan). *Soit  $A \in \mathcal{M}_n(\mathbb{R})$  une matrice réelle symétrique  $n \times n$ . On note  $\lambda_1(A) \geq \lambda_2(A) \geq \dots \geq \lambda_n(A)$  les valeurs propres de  $A$  répétées avec leur multiplicité et rangées en ordre décroissant. Pour tout  $k \in \{1, \dots, n\}$ ,*

$$\sum_{i=1}^k \lambda_i(A) = \max_{(X_1, \dots, X_k) \in \mathcal{O}_k} \sum_{i=1}^k \langle AX_i, X_i \rangle,$$

où  $\mathcal{O}_k$  est l'ensemble des  $k$ -uplets  $(X_1, \dots, X_k)$  de vecteurs colonnes qui forment une famille orthonormée de l'espace euclidien  $(\mathbb{R}^n, \langle \cdot, \cdot \rangle)$ .

*Démonstration.* Par théorème spectral, il existe  $(U_1, \dots, U_n) \in \mathcal{O}_n$  qui est telle que  $AU_i = \lambda_i(A)U_i$  pour tout  $i$ . On fixe  $k \in \{1, \dots, n\}$ . Clairement,

$$\begin{aligned} \max_{(X_1, \dots, X_k) \in \mathcal{O}_k} \sum_{i=1}^k \langle AX_i, X_i \rangle &\geq \sum_{i=1}^k \langle AU_i, U_i \rangle \\ &= \sum_{i=1}^k \lambda_i(A) . \end{aligned}$$

Réciproquement,

$$\max_{(X_1, \dots, X_k) \in \mathcal{O}_k} \sum_{i=1}^k \langle AX_i, X_i \rangle = \max_{(X_1, \dots, X_k) \in \mathcal{O}_k} \text{Trace}({}^t X A X)$$

où  $X$  est la matrice  $n \times k$  dont les colonnes sont les  $X_1, \dots, X_k$ . Par suite,

$$\max_{(X_1, \dots, X_k) \in \mathcal{O}_k} \sum_{i=1}^k \langle AX_i, X_i \rangle = \max_{(X_1, \dots, X_k) \in \mathcal{O}_k} \text{Trace}(A X {}^t X)$$

en raison du Lemme 29.1. Pour tout  $(X_1, \dots, X_k) \in \mathcal{O}_k$ ,

$${}^t X X = I_k ,$$

où  $I_k$  est la matrice identité  $k \times k$ . Pour  $k = 1, \dots, n$ , si  $E_i$  est le  $i$ ème vecteur colonne de la base canonique de  $\mathbb{R}^k$ , alors

$${}^t X X E_i = E_i$$

pour tout  $i = 1, \dots, k$ . On a  $X_i = X E_i$ . C'est donc un vecteur colonne non nul de  $\mathbb{R}^n$  et en appliquant  $X$  aux deux cotés de l'égalité ci-dessus on obtient que

$$(X {}^t X) X_i = X_i$$

Pour  $V \in \mathbb{R}^n$  un vecteur colonne de  $\mathbb{R}^n$  on a aussi que

$$X {}^t X V = 0 \Rightarrow ({}^t X X)({}^t X V) = 0 \Rightarrow {}^t X V = 0$$

et donc, au final,

$$\begin{aligned} X {}^t X V = 0 &\Leftrightarrow {}^t X V = 0 \\ &\Leftrightarrow \langle X_i, V \rangle = 0 \text{ pour tout } i = 1, \dots, k . \end{aligned}$$

Soient  $V_{k+1}, \dots, V_n$  des vecteurs colonnes de  $\mathbb{R}^n$  qui complètent la famille des  $X_i$  en une base orthonormée  $(X_1, \dots, X_k, V_{k+1}, \dots, V_n)$  de  $\mathbb{R}^n$ . Si  $P$  est la matrice de passage de la base canonique de  $\mathbb{R}^n$  à cette nouvelle base, alors  $P$  est orthogonale et

$${}^t P (X {}^t X) P = \begin{pmatrix} I_k & 0 \\ 0 & 0 \end{pmatrix}$$

où  $X {}^t X$  est carrée  $n \times n$ . On note  $\chi_k$  cette matrice en blocs. Alors

$$\begin{aligned} \text{Trace}(A X {}^t X) &= \text{Trace}({}^t P (A X {}^t X) P) \\ &= \text{Trace}({}^t P A P \times {}^t P (X {}^t X) P) \\ &= \text{Trace}({}^t P A P \chi_k) . \end{aligned}$$

Si  ${}^t P A P = (a_{ij})$  alors

$$\text{Trace}({}^t P A P \chi_k) = \sum_{i=1}^k a_{ii} .$$

Soit  $B = {}^tPAP$ . Les valeurs propres de  ${}^tPAP$  sont les mêmes (avec la même multiplicité) que celles de  $A$ . On note  $Q$  une matrice orthogonale telle que

$${}^tQBQ = D$$

où  $D$  est la matrice diagonale  $D = \text{diag}(\lambda_1(A), \dots, \lambda_n(A))$ . En vertu de la majoration de Schur du Lemme 30.1,  $\sum_{i=1}^k a_{ii} \leq \sum_{i=1}^k \lambda_i(A)$  et le théorème est démontré.  $\square$

## 29. TRACE D'UN PRODUIT DE MATRICES

Le résultat suivant est classique. Il est souvent énoncé dans le cas de matrices carrées mais il fonctionne tout aussi bien dans le cas rectangulaire dont nous avons eu besoin dans la preuve du Théorème 28.1.

**Lemme 29.1.** *Soit  $A$  une matrice  $k \times n$  et  $B$  une matrice  $n \times k$ . Alors  $AB$  est une matrice  $k \times k$ ,  $BA$  est une matrice  $n \times n$  et  $\text{Trace}(AB) = \text{Trace}(BA)$ .*

*Démonstration.* Soit  $A = (a_{ij})$  une matrice  $k \times n$  et  $B = (b_{ij})$  une matrice  $n \times k$ . Alors  $AB = (c_{ij})$  est une matrice carrée  $k \times k$  et  $BA = (d_{ij})$  est une matrice carrée  $n \times n$  avec

$$c_{ij} = \sum_{\alpha=1}^n a_{i\alpha} b_{\alpha j} \quad \text{et} \quad d_{ij} = \sum_{\beta=1}^k b_{i\beta} a_{\beta j}.$$

On a alors

$$\text{Trace}(AB) = \sum_{i=1}^k \sum_{\alpha=1}^n a_{i\alpha} b_{\alpha i}$$

tandis que

$$\begin{aligned} \text{Trace}(BA) &= \sum_{i=1}^n \sum_{\beta=1}^k b_{i\beta} a_{\beta i} \\ &= \sum_{i=1}^k \sum_{\alpha=1}^n b_{\alpha i} a_{i\alpha} \quad (i = \alpha, \beta = i) \end{aligned}$$

et donc, même dans le cas rectangulaire,  $\text{Trace}(AB) = \text{Trace}(BA)$ . Le lemme est démontré.  $\square$

## 30. MAJORATION DE SCHUR

**Lemme 30.1** (Majoration de Schur). *Soit  $A = (a_{ij})$  une matrice  $n \times n$ . On suppose qu'il existe  $Q$  une matrice orthogonale telle que  ${}^tQAQ = D$  où  $D = \text{diag}(\lambda_1, \dots, \lambda_n)$  avec  $\lambda_1 \geq \dots \geq \lambda_n$ . Alors pour tout  $1 \leq k \leq n$ ,  $\sum_{i=1}^k a_{ii} \leq \sum_{i=1}^k \lambda_i$ .*

*Démonstration.* L'existence de  $Q$  impose que  $A$  est symétrique. Soit  $Q = (q_{ij})$ . On a  $A = QD^tQ$  et donc

$$a_{ii} = \sum_{j=1}^n q_{ij}^2 \lambda_j$$

pour tous  $i, j$ . Comme  ${}^tPP = P^tP = I_n$  (matrice identité  $n \times n$ ), la matrice  $(p_{ij}^2)$  est doublement stochastique au sens où

$$\sum_{j=1}^n q_{ij}^2 = 1 \text{ pour tous } i \text{ et } \sum_{i=1}^n q_{ij}^2 = 1 \text{ pour tous } j.$$

Soit  $1 \leq k \leq n$  donné. Pour  $1 \leq j \leq n$  on pose  $t_j = \sum_{i=1}^k q_{ij}^2$ . Alors  $0 \leq t_j \leq 1$  et  $\sum_{j=1}^n t_j = k$ . On écrit maintenant que

$$\begin{aligned} \sum_{i=1}^k a_{ii} - \sum_{i=1}^k \lambda_i &= \sum_{j=1}^n t_j \lambda_j - \sum_{i=1}^k \lambda_i \\ &= \sum_{i=1}^n t_i \lambda_i - \sum_{i=1}^k \lambda_i + \left(k - \sum_{i=1}^n t_i\right) \lambda_k \\ &= \sum_{i=1}^k (t_i - 1)(\lambda_i - \lambda_k) + \sum_{i=k+1}^n t_i (\lambda_i - \lambda_k) \\ &\leq 0 \end{aligned}$$

puisque  $\lambda_i \geq \lambda_k$  pour  $1 \leq i \leq k$  (et  $t_i - 1 \leq 0$ ) et puisque  $\lambda_i \leq \lambda_k$  pour  $k+1 \leq i \leq n$  (et  $t_i \geq 0$ ). Le lemme est démontré.  $\square$

### 31. INÉGALITÉS DE KY FAN ET DE VON NEUMANN

Une première inégalité de Ky Fan est donnée par le Théorème 31.1. Si  $A$  est une matrice symétrique réelle  $n \times n$  et si  $1 \leq k \leq n$  on note

$$\sum_{i=1}^k \lambda_i^\downarrow(A) = \lambda_1(A) + \cdots + \lambda_k(A)$$

la somme ordonnée des  $k$  premières valeurs propres de  $A$ , les valeurs propres  $\lambda_1(A) \geq \lambda_2(A) \geq \cdots \geq \lambda_n(A)$  de  $A$  étant répétées avec leur multiplicité et rangées en ordre décroissant. L'inégalité de Ky Fan ci-dessous découle facilement du principe du maximum de Ky Fan donné dans le Théorème 28.1.

**Théorème 31.1** (Inégalité de Ky Fan). *Soient  $A, B$  sont deux matrices réelles symétriques  $n \times n$  alors*

$$\sum_{i=1}^k \lambda_i^\downarrow(A+B) \leq \sum_{i=1}^k \lambda_i^\downarrow(A) + \sum_{i=1}^k \lambda_i^\downarrow(B)$$

pour tout  $1 \leq k \leq n$ .

*Démonstration.* Soit  $1 \leq k \leq n$  fixé. On utilise le Théorème 28.1 pour écrire que

$$\begin{aligned} \sum_{i=1}^k \lambda_i^\downarrow(A+B) &= \max_{(X_1, \dots, X_k) \in \mathcal{O}_k} \sum_{i=1}^k \langle (A+B)X_i, X_i \rangle \\ &\leq \max_{(X_1, \dots, X_k) \in \mathcal{O}_k} \sum_{i=1}^k \langle AX_i, X_i \rangle + \max_{(X_1, \dots, X_k) \in \mathcal{O}_k} \sum_{i=1}^k \langle BX_i, X_i \rangle \\ &= \sum_{i=1}^k \lambda_i^\downarrow(A) + \sum_{i=1}^k \lambda_i^\downarrow(B). \end{aligned}$$

Le théorème est démontré.  $\square$

Une seconde inégalité, aussi répertoriée comme inégalité de Ky Fan, ou encore de Von Neumann, est donnée par le théorème suivant.

**Théorème 31.2** (Inégalité de Von Neumann). *Soient  $A, B$  deux matrices réelles symétriques  $n \times n$ . On note  $\lambda_1(A) \geq \lambda_2(A) \geq \dots \geq \lambda_n(A)$  et  $\lambda_1(B) \geq \lambda_2(B) \geq \dots \geq \lambda_n(B)$  les valeurs propres de  $A$  et  $B$  répétées avec leur multiplicité et rangées en ordre décroissant. Alors*

$$\text{Trace}(AB) \leq \sum_{i=1}^n \lambda_i(A) \lambda_i(B) .$$

En particulier,  $\langle A, B \rangle \leq {}^t\lambda(A)\lambda(B)$  où  $\langle \cdot, \cdot \rangle$  est le produit scalaire canonique sur  $\mathcal{M}_n(\mathbb{R})$ ,  $\lambda(A)$  est le vecteur colonne constitué des  $\lambda_i(A)$  et  $\lambda(B)$  est le vecteur colonne constitué des  $\lambda_i(B)$ .

L'inégalité  $\langle A, B \rangle \leq {}^t\lambda(A)\lambda(B)$  est un raffinement de l'inégalité de Cauchy-Schwarz classique puisque  ${}^t\lambda(A)\lambda(B) \leq \|A\| \times \|B\|$  une fois que l'on remarque que

$$\|A\|^2 = \text{Trace}({}^tAA) = \text{Trace}(A^2) = \text{Trace}(PD^{2t}P) = \text{Trace}(D^2)$$

(et pareil pour  $B$ ) où  $P$  orthogonale est telle que  $A = PD^tP$  et  $D$  est la matrice diagonale des  $\lambda_i(A)$  répétées avec leur multiplicité.

*Démonstration.* On démontre l'inégalité du théorème par récurrence sur la dimension  $n$ . Si  $n = 1$  il n'y a rien à démontrer. Le résultat est vrai. On suppose le résultat vrai pour toutes matrices symétriques  $A, B$  de tailles  $(n-1) \times (n-1)$ . On considère alors  $A$  et  $B$  symétriques de tailles  $n \times n$ . Soit  $D$  la matrice diagonale des  $\lambda_i(A)$  répétées avec leur multiplicité. Soit  $P$  orthogonale donnée par le théorème spectral et qui est telle que  ${}^tPAP = D$ . On a

$$\text{Trace}(AB) = \text{Trace}({}^tPABP) = \text{Trace}(D^tPBP) .$$

On note  $\hat{B} = {}^tPBP$ . Alors  $\hat{B}$  est aussi symétrique et  $B$  et  $\hat{B}$  ont mêmes valeurs propres avec mêmes multiplicités, l'isomorphisme associé à  $P$  réalisant le passage des espaces propres de  $\hat{B}$  aux espaces propres de  $B$  puisque  ${}^tPBU = \lambda U \Leftrightarrow B(PU) = \lambda(PU)$ . On note  $\hat{B}_{ii}$  la diagonale de  $\hat{B}$ . Alors

$$\begin{aligned} \text{Trace}(AB) &= \sum_{i=1}^n \lambda_i(A) \hat{B}_{ii} \\ &= \sum_{i=1}^{n-1} (\lambda_i(A) - \lambda_n(A)) \hat{B}_{ii} + \lambda_n(A) \sum_{i=1}^n \hat{B}_{ii} \\ &= \text{Trace}(\Lambda \tilde{B}) + \lambda_n(A) \text{Trace}(\hat{B}) , \end{aligned}$$

où  $\Lambda$  est la matrice diagonale  $(n-1) \times (n-1)$  dont la diagonale est constituée des  $\lambda_i(A) - \lambda_n(A)$ ,  $i = 1, \dots, n-1$ , et où  $\tilde{B}$  est la matrice  $(n-1) \times (n-1)$  donnée par

$$\tilde{B} = \begin{pmatrix} \tilde{B} & C \\ {}^tC & d \end{pmatrix}$$

avec  $C$  un vecteur colonne de taille  $n-1$  et  $d \in \mathbb{R}$  un réel. Clairement  $\tilde{B}$  est symétrique. Le Lemme 32.1 d'entrelacement de Cauchy donne que si on note

$$\lambda_1(\tilde{B}) \geq \dots \geq \lambda_{n-1}(\tilde{B})$$

les valeurs propres de  $\tilde{B}$  répétées avec leur multiplicité et rangées en ordre décroissant et

$$\lambda_1(\hat{B}) \geq \dots \geq \lambda_n(\hat{B})$$

les valeurs propres de  $\hat{B}$  répétées avec leur multiplicité et rangées en ordre décroissant (ce sont en fait les  $\lambda_i(B)$ , comme déjà dit), alors

$$\begin{aligned} \lambda_1(\hat{B}) &\geq \lambda_1(\tilde{B}) \geq \lambda_2(\hat{B}) \geq \lambda_2(\tilde{B}) \geq \cdots \geq \lambda_i(\hat{B}) \\ &\geq \lambda_i(\tilde{B}) \geq \lambda_{i+1}(\hat{B}) \geq \cdots \geq \lambda_{n-1}(\tilde{B}) \geq \lambda_n(\hat{B}) . \end{aligned}$$

Par hypothèse de récurrence sur  $\Lambda\tilde{B}$ , sachant que  $\lambda_i(A) \geq \lambda_n(A)$  et en raison de l'entrelacement ci-dessus,

$$\begin{aligned} \text{Trace}(AB) &\leq \sum_{i=1}^{n-1} (\lambda_i(A) - \lambda_n(A)) \lambda_i(\tilde{B}) + \lambda_n(A) \sum_{i=1}^n \lambda_i(\hat{B}) \\ &\leq \sum_{i=1}^{n-1} (\lambda_i(A) - \lambda_n(A)) \lambda_i(\hat{B}) + \lambda_n(A) \sum_{i=1}^n \lambda_i(\hat{B}) \\ &= \sum_{i=1}^n \lambda_i(A) \lambda_i(\hat{B}) \end{aligned}$$

et on trouve donc que

$$\text{Trace}(AB) \leq \sum_{i=1}^n \lambda_i(A) \lambda_i(B)$$

puisque  $\lambda_i(\hat{B}) = \lambda_i(B)$  pour tout  $i = 1, \dots, n$ . Le théorème est démontré.  $\square$

### 32. LE LEMME D'ENTRELACEMENT DE CAUCHY

Soient  $P_1, P_2$  deux polynômes réels, respectivement de degrés  $n$  et  $n-1$ ,  $n \geq 2$ , qui ont toutes leurs racines réelles  $\lambda_1 \geq \dots \geq \lambda_n$  et  $\mu_1 \geq \dots \geq \mu_{n-1}$ . On dit que  $P_1$  et  $P_2$  s'entrelacent si

$$\lambda_1 \geq \mu_1 \geq \lambda_2 \geq \dots \geq \lambda_i \geq \mu_i \geq \lambda_{i+1} \geq \dots \geq \mu_{n-1} \geq \lambda_n .$$

Le lemme d'entrelacement de Cauchy, utilisé dans la preuve de l'inégalité de Von-Neumann, est donné par le lemme suivant.

**Lemme 32.1** (Lemme d'entrelacement de Cauchy). *Soit  $A$  une matrice symétrique réelle de taille  $n \times n$  et soit  $B$  une matrice symétrique réelle de taille  $(n-1) \times (n-1)$ , les deux matrices étant liées par*

$$A = \begin{pmatrix} B & C \\ {}^t C & d \end{pmatrix} ,$$

où  $C$  un vecteur colonne de taille  $n-1$  et  $d \in \mathbb{R}$  un réel. Les polynômes caractéristiques de  $A$  et  $B$  sont alors entrelacés.

*Démonstration.* On utilise la caractérisation de Courant-Fischer du Théorème 27.1. On note  $\lambda_1 \geq \dots \geq \lambda_n$  les valeurs propres de  $A$  (donc les racines du polynôme caractéristique de  $A$ ) et  $\mu_1 \geq \dots \geq \mu_{n-1}$  les valeurs propres de  $B$  (donc les racines du polynôme caractéristique de  $B$ ). Soit  $X$  un vecteur colonne de  $\mathbb{R}^{n-1}$ . On note  $Y$  le vecteur colonne de  $\mathbb{R}^n$  donné par

$$Y = \begin{pmatrix} X \\ 0 \end{pmatrix} .$$

On remarque que si  $f$  est l'endomorphisme de  $\mathbb{R}^n$  dont la matrice de représentation dans la base canonique de  $\mathbb{R}^n$  est  $A$ , et si  $g$  est l'endomorphisme de  $\mathbb{R}^{n-1}$  dont la matrice de représentation dans la base canonique de  $\mathbb{R}^{n-1}$  est  $B$ , alors

$$\begin{aligned} {}^tYAY &= \langle f(Y), Y \rangle \\ &= {}^tXBX \\ &= \langle g(X), X \rangle . \end{aligned}$$

Soit  $k \in \{1, \dots, n-1\}$ . Etant donné  $F$  un sous espace vectoriel de dimension  $n-k$  de  $\mathbb{R}^{n-1}$ , on note  $\tilde{F} = F \times \{0\}$  le sous espace de  $\mathbb{R}^n$  de dimension  $n-k$  construit par assimilation de  $\mathbb{R}^{n-1}$  à  $\mathbb{R}^{n-1} \times \{0\} \subset \mathbb{R}^n$ . En vertu du Théorème 27.1,

$$\begin{aligned} \lambda_{k+1} &\leq \sup_{X \in \mathcal{S}_n \cap \tilde{F}} ({}^tXAX) \\ &\leq \sup_{Y \in \mathcal{S}_{n-1} \cap F} ({}^tYBY) \end{aligned}$$

où  $\mathcal{S}_n$  et  $\mathcal{S}_{n-1}$  sont les sphères unités de  $\mathbb{R}^n$  et  $\mathbb{R}^{n-1}$ . En passant à l'infimum sur les  $F$  sous espaces vectoriels de dimension  $n-k$  de  $\mathbb{R}^{n-1}$ , on obtient, toujours avec le Théorème 27.1, que

$$\lambda_{k+1} \leq \mu_k .$$

Les valeurs propres de  $-A$  rangées par ordre décroissant sont  $-\lambda_n \geq \dots \geq -\lambda_1$  et les valeurs propres de  $-B$  rangées par ordre décroissant sont  $-\mu_{n-1} \geq \dots \geq -\mu_1$ . La  $(k+1)$ -ème valeur propre de  $-A$  est donc  $-\lambda_{n-k}$  tandis que la  $k$ -ème valeur propre de  $-B$  est  $-\mu_{n-k}$ . En appliquant ce que l'on vient de démontrer à  $-A$  (et donc aussi  $-B$ ) on trouve que  $-\lambda_{n-k} \leq -\mu_{n-k}$  pour tout  $k \in \{1, \dots, n-1\}$  et donc que  $\lambda_{n-k} \geq \mu_{n-k}$  soit encore que

$$\lambda_k \geq \mu_k$$

pour tout  $k \in \{1, \dots, n-1\}$ . Le lemme est démontré.  $\square$

### 33. ENTRELACEMENTS D'HERMITE ET DE CAUCHY

Une preuve plus simple du Lemme 32.1 d'entrelacement de Cauchy (il faut ici comprendre n'utilisant pas la caractérisation de Courant-Fischer du Théorème 27.1) est donnée par S. Fisk (arXiv:math/0502408v1). Elle est basée sur la condition d'entrelacement de deux polynômes due à Hermite (que l'on admettra). Le résultat est démontré dans Q. I. Rahman and G. Schmeisser, *Analytic theory of polynomials*, London Mathematical Society Monographs. New Series, vol. 26, The Clarendon Press Oxford University Press, Oxford, 2002, ISBN 0-19-853493-0. MR 2004b:30015

**Lemme 33.1** (Condition d'entrelacement d'Hermite). *Deux polynômes réels  $P_1$  et  $P_2$ , avec  $\deg P_2 = \deg P_1 - 1$ , et  $P_1$  et  $P_2$  qui ont toutes leurs racines réelles, s'entrelacent si et seulement si pour tout  $\alpha \in \mathbb{R}$ , les polynômes  $P_1 + \alpha P_2$  ont toutes leurs racines réelles.*

On peut maintenant donner la preuve alternative que Fisk donne du lemme d'entrelacement de Cauchy.

*Démonstration alternative du lemme 32.1.* Soit  $\alpha \in \mathbb{R}$  fixé. Soient  $P_1$  le polynôme caractéristique de  $A$  et  $P_2$  le polynôme caractéristique de  $B$ . Par linéarité (pour

être exact par multilinéarité) du déterminant par rapport aux colonnes,

$$\begin{aligned} & \det \begin{pmatrix} B - xI & C \\ {}^tC & d - x + \alpha \end{pmatrix} \\ &= \det \begin{pmatrix} B - xI & C \\ {}^tC & d - x \end{pmatrix} + \det \begin{pmatrix} B - xI & 0 \\ {}^tC & \alpha \end{pmatrix} \\ &= P_1(x) + \alpha P_2(x) . \end{aligned}$$

Or

$$M_\alpha = \begin{pmatrix} B - xI & C \\ {}^tC & d - x + \alpha \end{pmatrix}$$

est symétrique. Elle a donc toutes ses racines réelles. On peut alors appliquer le Lemme 33.1 et on obtient le résultat voulu.  $\square$

## CHAPITRE 4

### ESPACES HERMITIENS

On considère dans tout ce qui suit des espaces vectoriels sur le corps  $\mathbb{C}$  des complexes. La notion d'espace vectoriel sur  $\mathbb{C}$  suit mot pour mot la définition des espaces vectoriels sur  $\mathbb{R}$ , mais maintenant avec une loi externe  $\mathbb{C} \times E \rightarrow E$  (et on peut donc multiplier les vecteurs par des complexes). Les définitions classiques d'applications linéaires, de familles libres, génératrices et bases suivent facilement en copiant les définitions du cas réel. Les choses changent un peu lorsqu'on parle de produit scalaire. Le modèle du produit scalaire dans  $\mathbb{R}$  est  $\langle x, y \rangle = xy$  qui donne bien  $\|x\|^2 = x^2$ . Dans  $\mathbb{C}$  la situation est plus subtile et on doit récupérer la norme de  $\mathbb{R}^2$ . On a alors comme modèle

$$\langle z, z' \rangle = z \overline{z'}$$

qui donne bien que  $\|z\|^2 = \langle z, z \rangle = z \overline{z} = x^2 + y^2$  si  $z = x + iy$ . On perd donc la notion d'application bilinéaire en vertu du passage au conjugué. Il faut introduire d'autres notion et on parlera dès lors de produit scalaire hermitien.

#### 34. PREMIÈRES DÉFINITIONS

On commence avec la définition des applications semi-linéaires.

**Définition 34.1.** Soient  $E, F$  deux espaces vectoriels sur  $\mathbb{C}$  et  $f : E \rightarrow F$  une application. L'application  $f$  est dite *linéaire* si  $\forall \lambda, \mu \in \mathbb{C}$  et  $\forall x, y \in E$ ,  $f(\lambda x + \mu y) = \lambda f(x) + \mu f(y)$ . L'application  $f$  est dite *semi-linéaire* si  $\forall \lambda, \mu \in \mathbb{C}$  et  $\forall x, y \in E$ ,  $f(\lambda x + \mu y) = \overline{\lambda} f(x) + \overline{\mu} f(y)$ , où  $\overline{\lambda}$  et  $\overline{\mu}$  sont les conjugués de  $\lambda$  et  $\mu$ .

On poursuit avec la définition des formes sesquilinéaires et hermitiennes.

**Définition 34.2.** Soit  $E$  un espace vectoriel sur  $\mathbb{C}$ . Une forme sesquilinéaire sur  $E$  est une application  $B : E \times E \rightarrow \mathbb{C}$  qui vérifie les deux points suivants :

- (i)  $\forall y \in E, x \rightarrow B(x, y)$  est linéaire sur  $E$ ,
- (ii)  $\forall x \in E, y \rightarrow B(x, y)$  est semi-linéaire sur  $E$ .

L'application  $B : E \times E \rightarrow \mathbb{C}$  est appelée une *forme hermitienne* si  $B$  est une forme sesquilinéaire et si  $B(y, x) = \overline{B(x, y)}$  pour tous  $x, y \in E$ , où  $\overline{B(x, y)}$  est le conjugué de  $B(x, y)$ .

Une conséquence simple de la définition est que si  $B : E \times E \rightarrow \mathbb{C}$  est hermitienne, alors forcément  $B(x, x) \in \mathbb{R}$  pour tout  $x \in E$  (puisque  $B(x, x) = \overline{B(x, x)}$  pour tout  $x \in E$ ). On aborde maintenant la définition des produits scalaires hermitiens, et donc aussi des espaces hermitiens.

**Définition 34.3.** Soit  $E$  un espace vectoriel sur  $\mathbb{C}$ . Un produit scalaire hermitien  $\langle \cdot, \cdot \rangle$  sur  $E$  est une forme hermitienne sur  $E$  qui est de plus définie positive au sens où  $\langle x, x \rangle \geq 0$  pour tout  $x$  avec la propriété que  $\langle x, x \rangle = 0$  si et seulement si  $x = 0$ . Un espace hermitien est un  $\mathbb{C}$ -espace vectoriel de dimension finie qui est muni d'un produit scalaire hermitien.

Un espace hermitien  $(E, \langle \cdot, \cdot \rangle)$  possède une norme naturelle associée donnée par  $\|x\| = \sqrt{\langle x, x \rangle}$  pour tout  $x \in E$ . La preuve est identique au cas réel et passe par l'identité de Cauchy-Schwarz:  $\forall x, y \in E$ ,

$$|\langle x, y \rangle| \leq \|x\| \times \|y\| .$$

Pour démontrer cette inégalité, on écrit que  $\forall x, y \in E, \forall \lambda \in \mathbb{C}$ ,

$$\begin{aligned} 0 &\leq \|\lambda x + y\|^2 \\ &= \langle \lambda x + y, \lambda x + y \rangle \\ &= \lambda \langle x, \lambda x + y \rangle + \langle y, \lambda x + y \rangle \\ &= |\lambda|^2 \|x\|^2 + \lambda \langle x, y \rangle + \bar{\lambda} \langle y, x \rangle + \|y\|^2 \\ &= |\lambda|^2 \|x\|^2 + \lambda \langle x, y \rangle + \overline{\lambda \langle x, y \rangle} + \|y\|^2 . \end{aligned}$$

On écrit que  $\langle x, y \rangle = |\langle x, y \rangle| e^{i\theta}$ . On pose  $\lambda = t e^{-i\theta}$  pour  $t \in \mathbb{R}$ . On a alors que:  $\forall t \in \mathbb{R}$ ,

$$t^2 \|x\|^2 + 2t |\langle x, y \rangle| + \|y\|^2 \geq 0 .$$

Comme dans le cas réel le discriminant du polynôme du second degré du membre de gauche de l'inégalité doit donc forcément être négatif ou nul, et on en déduit l'inégalité de Cauchy-Schwarz. Comme dans le cas réel, l'égalité dans l'inégalité de Cauchy-Schwarz a lieu si et seulement si  $x$  et  $y$  sont colinéaires.

### 35. BASES ORTHONORMÉES

Soit  $(E, \langle \cdot, \cdot \rangle)$  un espace hermitien. Une base  $(e_1, \dots, e_n)$  de  $E$  est dite orthogonale si  $\langle e_i, e_j \rangle = \delta_{ij}$  pour tous  $i, j$ , où les  $\delta_{ij}$  sont les symboles de Kroenecker.

**Théorème 35.1.** *Tout espace hermitien possède une base orthonormale.*

*Démonstration.* Soit  $(E, \langle \cdot, \cdot \rangle)$  un espace hermitien de dimension  $n$ . On applique le procédé d'orthonormalisation de Gram-Schmidt. On part d'une base  $(e_1, \dots, e_n)$  de  $E$ . On pose

$$\tilde{e}_1 = \frac{1}{\|e_1\|} e_1 .$$

Alors  $\|\tilde{e}_1\| = 1$ . On pose maintenant  $u = e_2 + \lambda \tilde{e}_1$ . On veut  $\lambda \in \mathbb{C}$  tel que  $\langle \tilde{e}_1, u \rangle = 0$ . On a

$$\langle \tilde{e}_1, u \rangle = \langle \tilde{e}_1, e_2 \rangle + \bar{\lambda} ,$$

et on pose ainsi  $\lambda = -\overline{\langle \tilde{e}_1, e_2 \rangle} = -\langle e_2, \tilde{e}_1 \rangle$ . On est assuré que  $u \neq 0$  car  $(e_1, e_2)$  est une famille libre. On pose ensuite

$$\tilde{e}_2 = \frac{1}{\|u\|} u .$$

Alors  $(\tilde{e}_1, \tilde{e}_2)$  est une famille orthonormale. Si  $n = 2$  on a fini (une famille orthonormale est automatiquement libre, et ayant le même nombre de vecteurs que la dimension il s'agit d'une base). Sinon,  $n \geq 3$ , et on recommence. On pose  $u = e_3 + \lambda \tilde{e}_1 + \mu \tilde{e}_2$ . On veut  $\lambda, \mu \in \mathbb{C}$  tels que  $\langle \tilde{e}_1, u \rangle = 0$  et  $\langle \tilde{e}_2, u \rangle = 0$ . On a

$$\langle \tilde{e}_1, u \rangle = \langle \tilde{e}_1, e_3 \rangle + \bar{\lambda} \text{ et } \langle \tilde{e}_2, u \rangle = \langle \tilde{e}_2, e_3 \rangle + \bar{\mu} .$$

On pose  $\lambda = -\overline{\langle \tilde{e}_1, e_3 \rangle}$  et  $\mu = -\overline{\langle \tilde{e}_2, e_3 \rangle}$ . Alors  $\langle \tilde{e}_1, u \rangle = 0$  et  $\langle \tilde{e}_2, u \rangle = 0$ . On est assuré que  $u \neq 0$  car  $(e_1, e_2, e_3)$  est une famille libre. On pose

$$\tilde{e}_3 = \frac{1}{\|u\|} u .$$

Alors  $(\tilde{e}_1, \tilde{e}_2, \tilde{e}_3)$  est une famille orthonormale. Si  $n = 3$  on a fini (une famille orthonormale est automatiquement libre, et ayant le même nombre de vecteurs que la dimension il s'agit d'une base). Sinon,  $n \geq 4$ , et on recommence. Etc.  $\square$

On notera que le procédé de Gram-Schmidt permet de construire une base orthonormée avec une partie imposée comme base d'un sous espace vectoriel  $F$ . Plus précisément, étant donné  $F \subset E$  un sous espace vectoriel de  $E$  de dimension  $k < n$ ,  $n$  la dimension de  $E$ , il existe une base orthonormée  $(e_1, \dots, e_n)$  de  $E$  avec la propriété que  $(e_1, \dots, e_k)$  est une base (orthonormée) de  $F$ .

Les coordonnées de  $x$  dans une base orthonormales sont les  $x_i = \langle x, e_i \rangle = \overline{\langle e_i, x \rangle}$ . L'expression du produit scalaire hermitien dans une base orthonormale est donnée par  $\langle x, y \rangle = \sum_{i,j} x_i \bar{y}_j$ .

### 36. RÉDUCTION DES ENDOMORPHISMES NORMAUX

On commence cette section avec la définition d'endomorphisme adjoint.

**Définition 36.1.** Soit  $(E, \langle \cdot, \cdot \rangle)$  un espace hermitien. Soit  $u : E \rightarrow E$  un endomorphisme de  $E$  (application linéaire de  $E$  dans lui-même). Il existe un unique endomorphisme  $u^* : E \rightarrow E$  de  $E$  tel que:  $\forall x, y \in E$ ,

$$\langle u(x), y \rangle = \langle x, u^*(y) \rangle .$$

On dit que  $u^*$  est l'endomorphisme adjoint de  $u$ .

Cette définition relève plus du théorème que de la définition. On doit montrer l'existence et l'unicité de  $u^*$ .

*Démonstration.* On considère  $(e_1, \dots, e_n)$  une base orthonormée de  $E$ . Si  $x, y \in E$  on note  $x_i$  et  $y_i$  les coordonnées de  $x$  et de  $y$  dans cette base. On a alors

$$\langle x, y \rangle = \sum_{i=1}^n x_i \bar{y}_i .$$

On considère la matrice complexe  $a_{ij}$  qui représente  $u$  dans la base. Donc:  $\forall j = 1, \dots, n$ ,

$$u(e_j) = \sum_{i=1}^n a_{ij} e_i .$$

On définit  $u^* : E \rightarrow E$  l'endomorphisme de  $E$  donné par:  $\forall i = 1, \dots, n$ ,

$$u^*(e_i) = \sum_{j=1}^n \bar{a}_{ij} e_j .$$

Alors, pour tous  $x, y \in E$ ,

$$\begin{aligned} \langle u(x), y \rangle &= \left\langle \sum_{i,j} a_{ij} x_j e_i, \sum_i y_i e_i \right\rangle \\ &= \sum_{i,j} a_{ij} \bar{y}_i x_j \end{aligned}$$

tandis que

$$\begin{aligned} \langle x, u^*(y) \rangle &= \left\langle \sum_i x_i e_i, \sum_{i,j} a_{ij} \bar{y}_i e_j \right\rangle \\ &= \sum_{i,j} a_{ij} y_i x_j . \end{aligned}$$

On a donc bien l'égalité  $\langle u(x), y \rangle = \langle x, u^*(y) \rangle$  pour tous  $x, y$ , ce qui prouve l'existence de  $u^*$ . En ce qui concerne l'unicité, supposons qu'il existe deux endomorphismes  $v, w$  de  $E$  tels que:  $\forall x, y \in E$ ,

$$\langle u(x), y \rangle = \langle x, v(y) \rangle \text{ et } \langle u(x), y \rangle = \langle x, w(y) \rangle .$$

En posant  $\hat{u} = v - w$  on obtient alors que

$$\langle x, \hat{u}(y) \rangle = 0$$

pour tous  $x, y \in E$ . En particulier:  $\forall i, j, \langle e_i, \hat{u}(e_j) \rangle = 0$ . On en déduit (voir l'expression des coordonnées dans une base orthonormée) que pour tout  $j, \hat{u}(e_j) = 0$ , puis que  $\hat{u}$  est l'endomorphisme nul. D'où l'unicité.  $\square$

A titre de remarque on montre facilement que  $u^{**} = u$ . On aborde maintenant la réduction proprement dite.

**Définition 36.2.** Soit  $(E, \langle \cdot, \cdot \rangle)$  un espace hermitien. Soit  $u : E \rightarrow E$  un endomorphisme de  $E$ . On dit que  $u$  est normal s'il commute avec son adjoint, et donc si  $u \circ u^* = u^* \circ u$ . On dira par ailleurs que:

- (1)  $u$  est unitaire si  $\forall x, y \in E, \langle u(x), u(y) \rangle = \langle x, y \rangle$ ,
- (2)  $u$  est hermitien si  $u = u^*$ ,
- (3) antihermitien si  $u = -u^*$ .

Un endomorphisme hermitien ou antihermitien est clairement normal. Il en va de même d'un endomorphisme unitaire. En effet, d'une part

$$\langle x, y \rangle = \langle u(x), u(y) \rangle = \langle x, u^* \circ u(y) \rangle$$

pour tous  $x$  et  $y$ . Par ailleurs,  $\text{Ker}(u) = \{0\}$  si  $u$  est unitaire. Donc  $u$  est un isomorphisme. On écrit alors que

$$\langle x, u \circ u^*(y) \rangle = \langle u \circ u^{-1}(x), u \circ u^*(y) \rangle = \langle u^{-1}(x), u^*(y) \rangle = \langle x, y \rangle$$

pour tous  $x$  et  $y$ . Donc  $\langle x, u^* \circ u(y) \rangle = \langle x, u \circ u^*(y) \rangle$  pour tous  $x$  et  $y$ . On en déduit que  $u \circ u^* = u^* \circ u$ . On dit maintenant d'une base  $(e_1, \dots, e_n)$  qu'elle diagonalise  $u$  si  $u(e_i) = \lambda_i e_i$  pour tout  $i = 1, \dots, n$ . Les  $\lambda_i \in \mathbb{C}$  sont les valeurs propres de  $u$  (racines du polynôme caractéristique). Le théorème de réduction concerne les endomorphismes normaux (ce qui étend la condition  $u^* = u$  du cas réel). Il s'énonce de la façon suivante.

**Théorème 36.1.** Soit  $(E, \langle \cdot, \cdot \rangle)$  un espace hermitien. Pour tout endomorphisme normal  $u$  de  $E$  il existe une base orthonormée de  $E$  qui diagonalise  $u$ . De plus, si  $u$  est unitaire, les valeurs propres de  $u$  sont de normes 1. Si  $u$  est hermitien les valeurs propres de  $u$  sont réelles. Si  $u$  est antihermitien, les valeurs propres de  $u$  sont imaginaires pures.

*Démonstration.* C'est en fait la même preuve que celle que l'on a faite dans le cas réel, mais simplifiée ici par le fait que l'on accepte de rester dans les complexes. On montre la réduction par récurrence totale sur la dimension  $n$  de  $E$ . Si  $n = 1$  la réduction est trivialement vraie. Il n'y a rien à démontrer. On suppose maintenant que pour tout  $\mathbb{C}$ -espace hermitien  $(E, \langle \cdot, \cdot \rangle)$  de dimension inférieure ou égale à  $n$  et tout endomorphisme normal  $u$  de  $E$ , il existe une base orthonormée de  $E$  qui diagonalise  $u$ . On considère  $(E, \langle \cdot, \cdot \rangle)$  un espace hermitien de dimension  $n + 1$ . Soit aussi  $u$  un endomorphisme normal de  $E$ . Soit  $\lambda_1 \in \mathbb{C}$  une racine du polynôme caractéristique de  $u$ , et donc une valeur propre de  $u$ . Soit  $E_1$  l'espace propre

associé à  $\lambda_1$ . On peut supposer que  $E_1 \neq E$  (sinon il n'y a rien à démontrer). Avec le théorème 35.1, voir aussi la remarque suivant le Théorème 35.1, on obtient une base orthonormée  $(e_1, \dots, e_k, e_{k+1}, \dots, e_{n+1})$  de  $E$  telle que  $(e_1, \dots, e_k)$  est une base orthonormée de  $E_1$  (ici  $k = \dim(E)$ , dimension complexe). Soit  $F$  l'espace vectoriel engendré par  $(e_{k+1}, \dots, e_n)$ . Alors  $F = E_1^\perp$ ,  $E_1 = F^\perp$  (i.e  $F$  est l'orthogonal de  $E_1$  et réciproquement) et on a que  $u(F) \subset F$ . En effet,  $\forall x \in E_1$ , puisque  $u$  est normal,

$$u(u^*(x)) = u^*(u(x)) = \lambda_1 u^*(x)$$

et ainsi,  $\forall x \in E_1$ ,  $u^*(x) \in E_1$  de sorte que  $u^*(E_1) \subset E_1$ . On écrit maintenant que  $\forall x \in E_1$ ,  $\forall y \in F$ ,

$$\begin{aligned} \langle x, u(y) \rangle &= \overline{\langle u(y), x \rangle} \\ &= \overline{\langle y, u^*(x) \rangle} \\ &= \langle u^*(x), y \rangle \\ &= 0 \end{aligned}$$

puisque  $u^*(x) \in E_1$  et  $F = E_1^\perp$ . Donc, comme  $x \in E_1$  est quelconque, on a montré que  $u(y) \in E_1^\perp$  pour tout  $y \in F$ , et ainsi  $u(F) \subset F$  comme annoncé. Par hypothèse de récurrence, il existe une base orthonormée  $(\tilde{e}_{k+1}, \dots, \tilde{e}_{n+1})$  de  $F$  qui diagonalise  $u$  restreint à  $F$ . Clairement,  $(e_1, \dots, e_k, \tilde{e}_{k+1}, \dots, \tilde{e}_{n+1})$  est une base orthonormée de  $E$  et, tout aussi clairement, elle diagonalise  $u$ . D'où la partie réduction du théorème. Soit maintenant  $\lambda \in \mathbb{C}$  une valeur propre d'un endomorphisme  $u$  d'un espace hermitien  $(E, \langle \cdot, \cdot \rangle)$ , et soit  $x \neq 0$  un vecteur propre associé. Si  $u$  est unitaire, alors

$$\langle u(x), u(x) \rangle = |\lambda|^2 \|x\|^2 = \langle x, x \rangle = \|x\|^2$$

et ainsi, comme annoncé, on a bien que  $|\lambda| = 1$ . Si  $u$  est hermitien, alors

$$\langle u(x), x \rangle = \lambda \|x\|^2 = \langle x, u(x) \rangle = \bar{\lambda} \|x\|^2$$

et donc  $\bar{\lambda} = \lambda$  de sorte que  $\lambda \in \mathbb{R}$ . Enfin si  $u$  est antihermitien, alors

$$\langle u(x), x \rangle = \lambda \|x\|^2 = -\langle x, u(x) \rangle = -\bar{\lambda} \|x\|^2$$

et donc  $\bar{\lambda} = -\lambda$  de sorte que  $\lambda$  est imaginaire pur. Le théorème est démontré.  $\square$

## CHAPITRE 5

### LA DIMENSION INFINIE

On commence par quelques notions de topologie. Soit  $X$  un ensemble quelconque. On appelle distance sur  $X$  toute application  $d : X \times X \rightarrow \mathbb{R}^+$  vérifiant les propriétés suivantes:

- (1)  $\forall x, y \in X, d(x, y) = 0 \Leftrightarrow x = y,$
- (2)  $\forall x, y \in X, d(x, y) = d(y, x),$
- (3) (inégalité triangulaire)  $\forall x, y, z, d(x, z) \leq d(x, y) + d(y, z) .$

Le couple  $(X, d)$  constitué d'un ensemble et d'une distance est appelé espace métrique.

**Lemme 36.1.** *Tout espace vectoriel normé  $(E, \|\cdot\|)$  possède une structure naturelle d'espace métrique  $(E, d)$ , la distance  $d$  associée à la norme étant définie par*

$$d(x, y) = \|y - x\|$$

*pour tous  $x, y \in E$ .*

Les propriétés (1), (2) et (3) caractéristiques des distances suivent très facilement des propriétés (1), (2) et (3) caractéristiques des normes. On a ainsi produit scalaire  $\Rightarrow$  norme  $\Rightarrow$  distance.

Dès lors que  $(E, d)$  est un espace métrique, on récupère plusieurs notions sur  $E$  comme celles de boules ouvertes et fermées, d'ouverts, de fermés, de suites convergentes et de Cauchy, d'application continue, etc. Si  $a$  est un point de  $E$ , et si  $r > 0$  est un réel strictement positif, on définit les boules ouvertes et fermées de centre  $a$  et de rayon  $r$  par  $B_a(r) = \{x \in E \text{ t.q. } d(a, x) < r\}$  et  $\bar{B}_a(r) = \{x \in E \text{ t.q. } d(a, x) \leq r\}$ . On a alors la définition suivante des ouverts et fermés.

**Définition 36.3.** *Soit  $(E, \|\cdot\|)$  un espace vectoriel normé, et  $d$  la distance associée à  $\|\cdot\|$ . On dit qu'un sous ensemble  $\Omega \subset E$  de  $E$  est un ouvert de  $E$  si pour tout point  $a \in \Omega$ , il existe  $r = r_a$  strictement positif, tel que  $B_a(r) \subset \Omega$ , où  $B_a(r)$  est la boule ouverte de centre  $a$  et de rayon  $r$  définie ci-dessus. On dit qu'un sous ensemble  $F \subset E$  est un fermé de  $E$  si  $E \setminus F$  est un ouvert de  $E$ .*

Par convention  $\emptyset$  et  $E$  sont à la fois ouverts et fermés. En langage intuitif, qui peut être rendu précis: (i) un ouvert est un ensemble qui ne contient aucun des points de sa frontière, et (ii) un fermé est un ensemble qui contient tous les points de sa frontière. Un ensemble peut donc n'être ni ouvert ni fermé. On montre sans trop de difficultés que les boules ouvertes sont des ouverts et les boules fermées des fermés.

**Définition 36.4.** *Soit  $(E, \|\cdot\|)$  un espace vectoriel normé, et  $d$  la distance associée à  $\|\cdot\|$ . Soit  $(x_n)_n$  une suite de points de  $E$ . On dit que  $(x_n)_n$  converge vers un point  $x$  de  $E$ , sous entendu pour la distance  $d$  ou pour la norme  $\|\cdot\|$ , si*

$$\forall \varepsilon > 0, \exists N \in \mathbb{N} \text{ t.q. } \forall n \geq N, d(x_n, x) < \varepsilon .$$

*On dit que  $x$  est la limite de la suite  $(x_n)_n$  (elle est forcément unique). Lorsqu'il n'y a pas de confusion possible avec le choix d'une norme, on note  $x_n \rightarrow x$  ou  $\lim_{n \rightarrow +\infty} x_n = x$ .*

Les fermés sont caractérisés par la propriété suivante:

**Lemme 36.2.** *Un sous ensemble  $F$  de  $E$  est un fermé de  $E$  si et seulement si pour toute suite  $(x_n)_n$  de points de  $F$ , si  $(x_n)_n$  converge dans  $E$ , alors sa limite  $x$  est forcément dans  $F$ .*

En remarquant (l'argument est intuitif mais il peut être rendu précis) que tout point de la frontière d'un patatoïde peut s'écrire comme limite d'une suite de points intérieurs, ce lemme signifie juste que l'ensemble possède bien tous les points de sa frontière.

*Démonstration.* (i) Supposons tout d'abord que  $F$  est un fermé de  $E$ , et soit  $(x_n)_n$  une suite de points de  $F$  qui converge dans  $E$  vers un point  $x$ . On veut montrer que  $x \in F$ . Par l'absurde, si  $x \notin F$ , alors  $x \in E \setminus F$ , et comme  $F$  est un fermé,  $E \setminus F$  est un ouvert. Par définition d'un ouvert, il existe alors  $r > 0$  tel que  $B_x(r) \subset E \setminus F$ . Or par définition de la convergence de  $(x_n)_n$  vers  $x$ , il doit exister  $N$  tel que  $x_n \in B_x(r)$  pour  $n \geq N$ . Mais comme  $x_n \in F$  et  $B_x(r) \subset E \setminus F$ , on obtient une contradiction.

(ii) A l'inverse, supposons que pour toute suite  $(x_n)_n$  de points de  $F$ , si  $(x_n)_n$  converge vers un point  $x$  dans  $E$ , alors  $x \in F$ . On veut montrer qu'alors  $F$  est un fermé de  $E$ . La encore, on raisonne par l'absurde. Si  $E \setminus F$  n'est pas un ouvert de  $E$ , alors il existe  $x \in E \setminus F$  avec la propriété que pour tout  $r > 0$ ,  $B_x(r) \cap F \neq \emptyset$ . On pose  $r = 1/n$  pour  $n \in \mathbb{N}$ ,  $n \geq 1$ , et on fait varier  $n$ . On obtient alors une suite  $(x_n)_n$  de points de  $F$  telle que  $d(x_n, x) \leq 1/n$  pour tout  $n \geq 1$ . Donc, par définition même,  $(x_n)_n$  converge vers  $x$  et, là encore, on récupère une contradiction.  $\square$

La continuité des applications est définie comme suit.

**Définition 36.5.** *Soient  $(E, \|\cdot\|_E)$  et  $(F, \|\cdot\|_F)$  deux espaces vectoriels normés,  $d_E$  la distance associée à  $\|\cdot\|_E$ , et  $d_F$  la distance associée à  $\|\cdot\|_F$ . Soient aussi  $X \subset E$  un sous ensemble de  $E$ ,  $a$  un point de  $X$ , et  $f : X \rightarrow F$  une application définie sur  $X$  et à valeurs dans  $F$ . On dit que  $f$  est continue au point  $a$ , sous entendu pour les distances  $d_E$  et  $d_F$ , ou pour les normes  $\|\cdot\|_E$  et  $\|\cdot\|_F$ , si l'une des trois propriétés équivalentes suivantes est vérifiée:*

- (i)  $\lim_{x \rightarrow a} f(x) = f(a)$ ,
- (ii)  $\forall \varepsilon > 0, \exists \delta > 0 / \forall x \in X, d_E(a, x) < \delta \Rightarrow d_F(f(a), f(x)) < \varepsilon$ ,
- (iii) pour toute suite  $(x_n)_n$  de points de  $X$  vérifiant que  $\lim_{n \rightarrow +\infty} x_n = a$ , on a forcément que  $\lim_{n \rightarrow +\infty} f(x_n) = f(a)$ .

Par extension, on dit que  $f$  est continue sur  $X$  si  $f$  est continue en tout point de  $X$ .

L'équivalence (i)  $\Leftrightarrow$  (ii) suit du fait que (ii) n'est quasiment rien d'autre que la définition mathématique de (i). L'équivalence (i)  $\Leftrightarrow$  (iii) se démontre comme dans  $\mathbb{R}$  en remplaçant les valeurs absolues de  $\mathbb{R}$  par les normes des espaces  $E$  et  $F$ .

Une remarque simple mais importante est la suivante: si  $(E, \|\cdot\|)$  est un  $\mathbb{R}$ -espace vectoriel normé, alors l'application norme de  $E$  dans  $\mathbb{R}$ , qui à  $x$  associe  $\|x\|$ , est continue sur  $E$ . On le voit en remarquant que  $|\|y\| - \|x\|| \leq \|y - x\|$ , soit encore que  $|\|y\| - \|x\|| \leq d(x, y)$  où  $d$  est la distance associée à  $\|\cdot\|$  (l'application est donc même lipschitzienne).

**Lemme 36.3.** Soient  $(E, \|\cdot\|_E)$ ,  $(F, \|\cdot\|_F)$  et  $(G, \|\cdot\|_G)$  trois espaces vectoriels normés,  $X$  est un sous ensemble de  $E$ ,  $Y$  est un sous ensemble de  $F$  tel que  $f(X) \subset Y$  et  $a \in X$  un point de  $X$ . Si  $f : X \rightarrow F$  est continue en  $a$ , et si  $g : Y \rightarrow G$  est continue en  $f(a)$ , alors  $g \circ f : X \rightarrow G$  est continue en  $a$ .

### 37. COMPACTITÉ ET DIMENSION

Soit  $(E, \|\cdot\|)$  un  $\mathbb{R}$ -espace vectoriel normé. Soit  $K \subset E$  un sous ensemble de  $E$ . On appelle recouvrement ouvert de  $K$  toute famille  $(\Omega_i)_{i \in I}$  d'ouverts de  $E$  telle que  $K \subset \bigcup_{i \in I} \Omega_i$ .

**Définition 37.1.** Soit  $(E, \|\cdot\|)$  un  $\mathbb{R}$ -espace vectoriel normé, et soit  $K \subset E$  un sous ensemble de  $E$ . Par définition,  $K$  est un compact de  $E$  si de tout recouvrement ouvert de  $K$  on peut extraire un sous recouvrement qui soit fini.

En d'autres termes,  $K \subset E$  est un compact de  $E$  si pour tout recouvrement ouvert  $(\Omega_i)_{i \in I}$  de  $K$ , il existe un sous ensemble fini  $\{i_1, \dots, i_p\} \subset I$  tel que  $K \subset \bigcup_{j=1}^p \Omega_{i_j}$ . Un théorème important en topologie est le suivant, dit de Bolzano-Weierstrass (ici énoncé dans le cas des espaces vectoriels normés, mais il possède une version topologique).

**Théorème 37.1** (Théorème de Bolzano-Weierstrass). Soit  $(E, \|\cdot\|)$  un  $\mathbb{R}$ -espace vectoriel normé, et soit  $K \subset E$  un sous ensemble de  $E$ . Alors  $K$  est un compact de  $E$  si et seulement si toute suite de points de  $K$  possède une sous suite qui converge dans  $K$ .

En d'autres termes,  $K \subset E$  est un compact de  $E$  si pour toute suite  $(x_n)_n$  de points de  $K$ , il existe  $\varphi : \mathbb{N} \nearrow \mathbb{N}$ , et il existe  $x \in K$ , tels que  $x_{\varphi(n)} \rightarrow x$  lorsque  $n \rightarrow +\infty$  (la notation  $\varphi : \mathbb{N} \nearrow \mathbb{N}$  signifie "application strictement croissante de  $\mathbb{N}$  dans  $\mathbb{N}$ "). Plusieurs propriétés importantes se rattachent à la notion de compacité. Parmi les résultats célèbres on citera le théorème de Weierstrass (l'image d'un compact par une application continue est un compact), et le théorème de Tychonoff (tout produit d'espaces compacts est compact). On pourra encore montrer le théorème suivant:

**Théorème 37.2.** Soit  $(E, \|\cdot\|)$  un  $\mathbb{R}$ -espace vectoriel normé, et soit  $K \subset E$  un sous ensemble compact de  $E$ . Soit  $f : K \rightarrow \mathbb{R}$  une fonction continue sur  $K$ . Alors  $f$  est bornée et elle atteint ses bornes.

*Démonstration.* Que  $f$  soit bornée, à savoir qu'il existe  $M > 0$  tel que  $|f(x)| \leq M$  pour tout  $x \in K$ , suit facilement du théorème de Weierstrass (un compact de  $\mathbb{R}$  est forcément borné). Montrons maintenant que  $f$  atteint son minimum (la démonstration pour le maximum est identique, ou on peut changer  $f$  en  $-f$ ). Soit

$$m = \inf_{x \in K} f(x)$$

la borne inférieure de l'ensemble  $\{f(x), x \in K\}$ . Par définition,  $m$  est le plus grand des minorants de cet ensemble et donc:

- (i)  $\forall x \in K, m \leq f(x)$ , et
- (ii)  $\forall \varepsilon > 0, \exists x \in K / m \leq f(x) \leq m + \varepsilon$ .

On pose  $\varepsilon = 1/n$  pour  $n \in \mathbb{N} \setminus \{0\}$ . On obtient ainsi l'existence d'un  $x_n \in K$  tel que  $\forall n \in \mathbb{N} \setminus \{0\}$ ,

$$m \leq f(x_n) \leq m + \frac{1}{n}.$$

Donc  $f(x_n) \rightarrow m$  lorsque  $n \rightarrow +\infty$ . Comme  $K$  est un compact, il existe  $\varphi : \mathbb{N} \nearrow \mathbb{N}$  telle que  $(x_{\varphi(n)})_n$  converge. Notons  $x$  sa limite. Une sous suite d'une suite convergente étant convergente et de même limite, on a encore que  $f(x_{\varphi(n)}) \rightarrow m$  lorsque  $n \rightarrow +\infty$ . Par continuité de  $f$  en  $x$  il s'ensuit que  $f(x) = m$ .  $\square$

Un autre résultat fréquemment associé à la compacité est que toute application continue sur un compact  $Y$  est uniformément continue.

### 38. LE SECOND THÉORÈME DE RIESZ

On discute ici du second théorème de Riesz. Pour rappel, en dimension finie, la boule unité fermée est compact (les compacts de  $\mathbb{R}^n$  sont les fermés bornés). Le second théorème de Riesz marque une rupture spectaculaire entre dimension finie et dimension infinie.

**Théorème 38.1** (Second théorème de Riesz). *Soit  $(E, \|\cdot\|)$  un espace vectoriel normé, et soit  $B' = \overline{B}_0(1)$  la boule fermée de  $E$  de centre 0 et rayon 1. Alors  $E$  est de dimension finie si et seulement si  $B'$  est un compact de  $E$ .*

Bien sûr le résultat reste valable si on remplace  $B'$  par n'importe quelle boule fermée  $\overline{B}_a(r)$ . En particulier tout compact dans un espace vectoriel normé de dimension infinie est d'intérieur vide (i.e ne contient aucune boule ouverte du type  $B_a(r)$ ,  $r > 0$ ).

*Démonstration.* On montre que les compacts d'un espace vectoriel normé de dimension finie sont les fermés bornés de cet espace. En dimension finie  $n$  l'espace est homéomorphe (et même isomorphe) à  $\mathbb{R}^n$  et, dans  $\mathbb{R}^n$ , les compacts sont précisément les fermés bornés (ce qui se démontre par extractions successives de sous suites en utilisant le fait que les suites réelles bornées possèdent des sous suites convergentes). On utilise ici (cf. plus loin), que toutes les normes sur  $\mathbb{R}^n$  sont équivalentes. En particulier, si  $E$  est de dimension finie, alors  $B'$  est compacte. Réciproquement, supposons que  $B'$  est un compact de  $E$ . On veut montrer que  $E$  est de dimension finie. Pour cela, on commence par remarquer que  $(B_{a_i}(\frac{1}{2}))_{a_i \in B'}$  est un recouvrement ouvert de  $B'$ . Par suite, et puisque  $B'$  est supposé compact, il existe des  $a_1, \dots, a_n \in B'$  tels que

$$B' \subset \bigcup_{i=1}^n B_{a_i}(\frac{1}{2}) .$$

Notons  $F$  le sous espace vectoriel de  $E$  engendré par les  $a_i$ ,  $F = \text{Vect}(a_1, \dots, a_n)$ . Clairement,  $F$  est de dimension finie, puisque par construction,  $(a_1, \dots, a_n)$  est une famille génératrice de  $F$ . On va montrer que  $E = F$ . On raisonne par l'absurde et on suppose que  $F \neq E$ . Il existe alors  $x \in E \setminus F$ . Notons

$$\alpha = \inf_{y \in F} \|x - y\| .$$

Puisque  $F$  est de dimension finie, il existe  $\tilde{x} \in F$  tel que  $\alpha = \|x - \tilde{x}\|$  (les fermés bornés dans  $F$  sont compacts). En particulier,  $\alpha > 0$ . Soit maintenant  $y \in F$  tel que

$$\alpha \leq \|x - y\| \leq \frac{3\alpha}{2}$$

et soit  $z = \frac{1}{\|x-y\|}(x-y)$ . Alors  $z \in B'$ , et il existe ainsi  $i \in \{1, \dots, n\}$  pour lequel  $z \in B_{a_i}(\frac{1}{2})$ . On remarque maintenant que

$$y' = y + \|x-y\|a_i$$

appartient à  $F$ , et que

$$x = y + \|x-y\|z.$$

Ainsi

$$\begin{aligned} \alpha &\leq \|x-y'\| \\ &= \|x-y\| \cdot \|z-a_i\| \\ &\leq \frac{3\alpha}{2} \cdot \frac{1}{2} = \frac{3\alpha}{4} < \alpha \end{aligned}$$

et on obtient donc une contradiction. Il s'ensuit que  $E = F$ . En particulier,  $E$  est de dimension finie. D'où le théorème.  $\square$

### 39. ESPACES DE BANACH ET DE HILBERT

**Définition 39.1.** Soit  $(E, \|\cdot\|)$  un espace vectoriel normé, et  $d$  la distance associée à  $\|\cdot\|$ . Soit aussi  $(x_n)_n$  une suite de points de  $E$ . On dit que  $(x_n)_n$  est une suite de Cauchy, sous entendu pour la distance  $d$  ou pour la norme  $\|\cdot\|$ , si

$$\forall \varepsilon > 0, \exists N \in \mathbb{N} \text{ t.q. } \forall p, q \geq N, d(x_p, x_q) < \varepsilon.$$

En d'autres termes une suite  $(x_n)_n$  est de Cauchy si les  $x_n$  s'accumulent les uns sur les autres quand  $n \rightarrow +\infty$ .

Une remarque simple est qu'une suite convergente est toujours de Cauchy. Il suffit pour le voir d'écrire que

$$d(x_p, x_q) \leq d(x_p, x) + d(x, x_q),$$

où  $x$  est la limite de la suite. Si les  $x_n$  s'accumulent sur un point  $x$ , alors ils s'accumulent les uns sur les autres.

**Définition 39.2.** Soit  $(E, \|\cdot\|)$  un espace vectoriel normé. On dit que  $(E, \|\cdot\|)$  est un espace vectoriel normé complet, ou encore un espace de Banach, si toute suite de Cauchy dans  $E$  pour  $\|\cdot\|$  converge dans  $E$  pour  $\|\cdot\|$ . Si la norme provient d'un produit scalaire  $\langle \cdot, \cdot \rangle$ , on dit que  $(E, \langle \cdot, \cdot \rangle)$  est un espace de Hilbert.

Des exemples d'espaces de Banach sont:

Ex.1: L'espace euclidien  $\mathbb{R}^n$  muni du produit scalaire euclidien  $\langle x, y \rangle = \sum_i x_i y_i$  est un espace de Hilbert.

Ex.2: L'espace  $\mathcal{L}^2(\mathbb{R})$  constitué des suites réelles  $(x_n)_n$  pour lesquelles la série  $\sum x_n^2$  converge, muni du produit scalaire  $\langle (x_n)_n, (y_n)_n \rangle = \sum_{n=0}^{+\infty} x_n y_n$ , est un espace de Hilbert.

Ex.3: Si  $(X, d)$  est un espace métrique, l'espace  $C_B^0(X)$  des fonctions réelles continues et bornées sur  $X$ , muni de la norme de la convergence uniforme  $\|f\|_{L^\infty} = \sup_{x \in X} |f(x)|$ , est un espace de Banach.

Il est bien sûr des espaces qui ne sont pas de Banach. Ces espaces sont forcément de dimension infinie (comme on le verra plus tard).

**Définition 39.3.** Deux normes  $\|\cdot\|_1$  et  $\|\cdot\|_2$  sur un espace vectoriel  $E$  sont équivalentes s'il existe des constantes  $C_1, C_2 > 0$  telles que

$$C_1\|x\|_1 \leq \|x\|_2 \leq C_2\|x\|_1$$

pour tout  $x \in E$ .

On vérifie que l'équivalence des normes préservent les notions de

- (1) suites convergentes et limites,
- (2) suites de Cauchy.

En particulier, si  $\|\cdot\|_1$  et  $\|\cdot\|_2$  sont équivalentes, alors  $(E, \|\cdot\|_1)$  est un Banach si et seulement si  $(E, \|\cdot\|_2)$  est un Banach. Plus généralement, deux normes équivalentes préservent la topologie.

**Théorème 39.1.** Sur un espace vectoriel réel de dimension finie, deux normes sont toujours équivalentes.

*Démonstration.* Pour simplifier, on va supposer que  $E = \mathbb{R}^n$ , même si la démonstration procède en fait de façon analogue pour  $E$  un espace vectoriel quelconque de dimension finie. On note  $\|\cdot\|_0$  la norme standard de  $\mathbb{R}^n$ , définie par

$$\|x\|_0 = \sqrt{\sum_{i=1}^n x_i^2}.$$

Comme on s'en convaincra facilement, démontrer le théorème revient à montrer que toute norme  $\|\cdot\|$  sur  $\mathbb{R}^n$  est équivalente à  $\|\cdot\|_0$ . Étant donnée  $\|\cdot\|$  une norme sur  $\mathbb{R}^n$ , et  $(e_1, \dots, e_n)$  la base canonique de  $\mathbb{R}^n$ , on écrit que

$$\|x_1 e_1 + \dots + x_n e_n\| \leq \sum_{i=1}^n |x_i| \|e_i\|.$$

Par suite, et si

$$M = \sum_{i=1}^n \|e_i\|$$

on obtient que pour tout  $x \in E$ ,

$$\|x\| \leq M \left( \max_{i=1, \dots, n} |x_i| \right) \leq M \|x\|_0.$$

Il s'ensuit que l'application  $f$  de  $\mathbb{R}^n$  dans  $\mathbb{R}$  qui à  $x$  associe  $\|x\|$  est continue pour  $\|\cdot\|_0$ . Notons  $S$  la sphère unité de  $\mathbb{R}^n$  pour  $\|\cdot\|_0$ :

$$S = \{x \in \mathbb{R}^n / \|x\|_0 = 1\}.$$

Clairement,  $S$  est un fermé et borné de  $\mathbb{R}^n$ , donc un compact de  $\mathbb{R}^n$ . Il s'ensuit que la restriction de  $f$  à  $S$  est bornée et qu'elle atteint ses bornes. Sachant que pour  $x \in S$ ,  $f(x) \neq 0$ , puisque sinon on devrait avoir  $0 \in S$ , on obtient qu'il existe un réel  $m > 0$  tel que pour tout  $x \in S$ ,  $f(x) \geq m$ . En d'autres termes, il existe  $m > 0$  tel que pour tout  $x \in S$ ,  $\|x\| \geq m$ . Si maintenant  $x$  est un vecteur non nul quelconque de  $\mathbb{R}^n$ , en remarquant que  $\frac{1}{\|x\|_0} x \in S$ , on obtient que pour tout  $x \in \mathbb{R}^n$ ,  $\|x\| \geq m \|x\|_0$ . Posons pour finir  $\lambda = \max(\frac{1}{m}, M)$ . On voit alors que pour tout  $x \in \mathbb{R}^n$ ,

$$\frac{1}{\lambda} \|x\|_0 \leq \|x\| \leq \lambda \|x\|_0.$$

D'où le théorème. □

On a alors aussi le résultat suivant.

**Théorème 39.2.** *Tout espace vectoriel normé de dimension finie est un espace de Banach.*

*Démonstration.* Soient  $\|\cdot\|$  la norme de  $E$  et  $\mathcal{B}$  une base de  $E$ . On note  $\|\cdot\|_0$  la norme euclidienne de  $E$  définie par  $\|x\|_0 = \sqrt{\sum_{i=0}^n x_i^2}$ , où les  $x_i$  sont les coordonnées de  $x$  dans  $\mathcal{B}$  et  $n = \dim(E)$ . Alors  $(E, \|\cdot\|_0)$  est un Banach (pour les mêmes raisons que  $\mathbb{R}^n$  muni de la norme euclidienne est un Banach). Comme  $\|\cdot\|$  est équivalente à  $\|\cdot\|_0$  d'après le théorème précédent,  $(E, \|\cdot\|)$  est aussi un Banach. □

#### 40. APPLICATIONS LINÉAIRES CONTINUES

On aborde la notion importante d'application linéaire continue. Pour mémoire, une application linéaire d'un espace vectoriel  $E$  dans un espace vectoriel  $F$  est une application  $f : E \rightarrow F$  qui vérifie que  $f(x + y) = f(x) + f(y)$  et  $f(\lambda x) = \lambda f(x)$  pour tous  $x, y \in E$ , et tout  $\lambda \in \mathbb{R}$ .

**Théorème 40.1.** *Soient  $(E, \|\cdot\|_E)$  et  $(F, \|\cdot\|_F)$  deux espaces vectoriels normés, et soit  $f : E \rightarrow F$  une application linéaire de  $E$  dans  $F$ . Les trois propriétés suivantes sont équivalentes*

- 1)  $f$  est continue sur  $E$
- 2)  $f$  est continue en 0
- 3)  $\exists M > 0 / \forall x \in E, \|f(x)\|_F \leq M\|x\|_E$

où  $f$  est dite continue si elle l'est en tant qu'application de l'espace vectoriel normé  $(E, \|\cdot\|_E)$  dans l'espace vectoriel normé  $(F, \|\cdot\|_F)$ .

*Démonstration.* L'implication 1)  $\Rightarrow$  2) est immédiate. Supposons maintenant que 2) a lieu. Alors

$$\forall \varepsilon > 0, \exists \eta > 0 / \forall x \in E, \|x\|_E < \eta \Rightarrow \|f(x)\|_F < \varepsilon.$$

En prenant  $\varepsilon = 1$  dans cette relation, on obtient ainsi qu'il existe  $\eta > 0$  tel que pour tout  $x \in E$ ,  $\|x\|_E < \eta$  entraîne  $\|f(x)\|_F \leq 1$ . Etant donné  $x$  un point quelconque de  $E$ , avec  $x \neq 0$ , on pose

$$y = \frac{\eta}{2\|x\|_E} x.$$

En remarquant que  $\|y\|_E < \eta$ , on obtient que  $\|f(y)\|_F \leq 1$ . Par linéarité de  $f$ , cela entraîne que

$$\|f(x)\|_F \leq \frac{2}{\eta} \|x\|_E.$$

On a ainsi montré que 2)  $\Rightarrow$  3). Supposons pour finir que 3) a lieu. Etant donné  $x \in E$ , soit  $y \in E$  quelconque. Alors

$$\|f(y) - f(x)\|_F = \|f(y - x)\|_F.$$

de sorte qu'avec 3) on obtient que pour tout  $y \in E$ ,

$$\|f(y) - f(x)\|_F \leq M\|y - x\|_E.$$

Donc  $f$  est lipschitzienne, et il s'ensuit que  $f$  est continue au point  $x$ . Ainsi, 3)  $\Rightarrow$  1), et le théorème est démontré. □

Le théorème suivant a lieu.

**Théorème 40.2.** *Soient  $(E, \|\cdot\|_E)$  et  $(F, \|\cdot\|_F)$  deux espaces vectoriels normés. Si  $E$  est de dimension finie, toute application linéaire de  $E$  dans  $F$  est continue.*

*Démonstration.* Soit  $(e_1, \dots, e_n)$  une base de  $E$ . On note  $\|\cdot\|'_E$  la norme de  $E$  définie pour tout  $x = x_1e_1 + \dots + x_ne_n$  de  $E$  par

$$\|x\|'_E = \sum_{i=1}^n |x_i| .$$

Sachant que  $E$  est de dimension finie, et d'après ce qui a été dit plus haut sur l'équivalence des normes, il existe une constante réelle  $C > 0$  telle que pour tout  $x \in E$ ,  $\|x\|'_E \leq C\|x\|_E$ . Etant donné  $f$  une application linéaire de  $E$  dans  $F$ , on écrit que pour tout  $x = x_1e_1 + \dots + x_ne_n$  de  $E$

$$\begin{aligned} \|f(x)\|_F &= \left\| \sum_{i=1}^n x_i f(e_i) \right\|_F \\ &\leq \sum_{i=1}^n |x_i| \cdot \|f(e_i)\|_F \\ &\leq \left( \max_{i=1, \dots, n} \|f(e_i)\|_F \right) \|x\|'_E \\ &\leq C \left( \max_{i=1, \dots, n} \|f(e_i)\|_F \right) \|x\|_E . \end{aligned}$$

Il s'ensuit qu'il existe un réel  $M > 0$ ,

$$M = C \left( \max_{i=1, \dots, n} \|f(e_i)\|_F \right) ,$$

tel que pour tout  $x \in E$ ,  $\|f(x)\|_F \leq M\|x\|_E$ . D'après le théorème précédent, cela entraîne que  $f$  est continue sur  $E$ .  $\square$

L'espace des linéaires continues devient très vite un Banach.

**Théorème 40.3.** *Soient  $(E, \|\cdot\|_E)$  et  $(F, \|\cdot\|_F)$  deux espaces vectoriels normés, et soit  $L_c(E, F)$  l'espace vectoriel des applications linéaires continues de  $E$  dans  $F$ . Soit  $\|\cdot\|_{L_c(E, F)}$  la norme sur  $L_c(E, F)$  définie par*

$$\|f\|_{L_c(E, F)} = \sup_{x \in E \setminus \{0\}} \frac{\|f(x)\|_F}{\|x\|_E} .$$

*Alors non seulement  $\|\cdot\|_{L_c(E, F)}$  est bien une norme sur  $L_c(E, F)$ , mais en plus, l'espace  $(L_c(E, F), \|\cdot\|_{L_c(E, F)})$  est un espace de Banach dès que  $(F, \|\cdot\|_F)$  est un espace de Banach.*

*Démonstration.* On vérifie facilement que  $\|\cdot\|_{L_c(E, F)}$  est bien une norme. On montre que  $(L_c(E, F), \|\cdot\|_{L_c(E, F)})$  est un espace de Banach dès que  $(F, \|\cdot\|_F)$  est un espace de Banach. On considère  $(f_n)_n$  une suite de Cauchy dans  $L_c(E, F)$ . Alors

$$\forall \varepsilon > 0 , \exists N \in \mathbb{N} / \forall p, q \geq N , \|f_p - f_q\|_{L_c(E, F)} < \varepsilon . \quad (1)$$

En particulier, par définition même de  $\|\cdot\|_{L_c(E, F)}$ , on obtient avec (1) que

$$\begin{aligned} \forall \varepsilon > 0 , \exists N \in \mathbb{N} / \forall p, q \geq N , \forall x \in E , \\ \|f_p(x) - f_q(x)\|_F \leq \varepsilon \|x\|_E \quad (2) \end{aligned}$$

La suite  $(f_n(x))_n$  est donc de Cauchy dans  $(F, \|\cdot\|_F)$ . Comme  $(F, \|\cdot\|_F)$  est un espace de Banach, elle converge dans  $(F, \|\cdot\|_F)$ . On note  $f(x)$  sa limite, et comme  $x$  est quelconque dans  $E$ , on récupère une application  $f : E \rightarrow F$ . Pour  $x, y \in E$  et  $\lambda \in \mathbb{R}$ ,

$$\begin{aligned} & \|f(x + \lambda y) - f(x) - \lambda f(y)\|_F \\ &= \|f(x + \lambda y) - f_n(x + \lambda y) + f_n(x) + \lambda f_n(y) - f(x) - \lambda f(y)\|_F \\ &\leq \|f(x + \lambda y) - f_n(x + \lambda y)\|_F + \|f_n(x) - f(x)\|_F + |\lambda| \|f_n(y) - f(y)\|_F \end{aligned}$$

pour tout  $n$ , puisque  $f_n$  est linéaire. En faisant tendre  $n \rightarrow +\infty$ , on voit que  $f$  est bien linéaire, à savoir que

$$f(x + \lambda y) = f(x) + \lambda f(y)$$

pour tous  $x, y \in E$  et tout  $\lambda \in \mathbb{R}$ . En faisant tendre  $p \rightarrow +\infty$  dans (2), avec  $q$  et  $\varepsilon$  fixés, par exemple  $\varepsilon = 1$  et  $q = N$ , on voit aussi que pour tout  $x \in E$ ,

$$\begin{aligned} \|f(x)\|_F &= \|f(x) - f_q(x) + f_q(x)\|_F \\ &\leq \varepsilon \|x\|_E + \|f_q(x)\|_F \\ &\leq (\|f_q\|_{L_c(E,F)} + \varepsilon) \|x\|_E . \end{aligned}$$

En particulier,  $f$  est continue sur  $E$ . Reste à montrer que  $(f_n)_n$  converge vers  $f$  dans  $(L_c(E, F), \|\cdot\|_{L_c(E,F)})$ . Pour cela on revient à (2), on fait tendre  $q$  vers l'infini, et on prend ensuite le supremum sur les  $x \in E \setminus \{0\}$  dans la conclusion de (2). Alors

$$\forall \varepsilon > 0, \exists N \in \mathbb{N} / \forall p \geq N, \sup_{x \in E \setminus \{0\}} \frac{\|f_p(x) - f(x)\|_F}{\|x\|_E} \leq \varepsilon \quad (3)$$

et puisque  $\varepsilon > 0$  est quelconque, (3) n'exprime rien d'autre que la convergence de la suite  $(f_n)_n$  vers  $f$  dans  $(L_c(E, F), \|\cdot\|_{L_c(E,F)})$ , à savoir une relation du type

$$\forall \varepsilon > 0, \exists N \in \mathbb{N} / \forall p \geq N, \forall x \in E, \|f_p - f\|_{L_c(E,F)} < \varepsilon .$$

L'espace  $(L_c(E, F), \|\cdot\|_{L_c(E,F)})$  est un espace de Banach.  $\square$

A titre de remarque, on a aussi

$$\|f\|_{L_c(E,F)} = \sup_{\{x \in E / \|x\|_E \leq 1\}} \|f(x)\|_F$$

ou encore

$$\|f\|_{L_c(E,F)} = \sup_{\{x \in E / \|x\|_E = 1\}} \|f(x)\|_F .$$

En fait  $\|f\|_{L_c(E,F)}$  est le plus petit des réels  $M$  pour lesquels  $\|f(x)\|_F \leq M\|x\|_E$  pour tout  $x$ .

Une autre remarque importante mais simple à démontrer est que la composée d'applications linéaires continues est encore une application linéaire continue. Le résultat se démontre en écrivant que si  $f : E \rightarrow F$  et  $g : F \rightarrow G$  sont linéaires, alors pour tout  $x \in E$ ,

$$\begin{aligned} \|g \circ f(x)\|_G &\leq \|g\|_{L_c(F,G)} \|f(x)\|_F \\ &\leq \|g\|_{L_c(F,G)} \|f\|_{L_c(E,F)} \|x\|_E . \end{aligned}$$

En particulier, on voit que

$$\|g \circ f\|_{L_c(E,G)} \leq \|g\|_{L_c(F,G)} \|f\|_{L_c(E,F)} .$$

## 41. PROJECTIONS ORTHOGONALES

Un premier résultat dont on aura besoin est le suivant.

**Lemme 41.1.** *Soit  $(E, \|\cdot\|)$  un espace de Banach. Soit  $F$  un sous espace vectoriel de  $E$ . Si  $F$  est un fermé de  $E$ , alors  $(F, \|\cdot\|)$  est aussi un espace de Banach.*

*Démonstration.* Soit  $(x_n)_n$  une suite de Cauchy dans  $F$ . La suite  $(x_n)_n$  est naturellement de Cauchy dans  $E$ . Elle converge donc dans  $E$ . En vertu du Lemme 36.2 sa limite est dans  $F$  dès que  $F$  est fermé. On en déduit que  $(x_n)_n$  converge dans  $F$  et  $(F, \|\cdot\|)$  est ainsi un espace de Banach. Le lemme est démontré.  $\square$

On aborde maintenant le théorème de projection orthogonale qui généralise à la dimension infinie ce qui a été dit à la Section 16. Etant donnée un espace vectoriel normé  $(E, \|\cdot\|)$ ,  $X \subset E$  un sous ensemble de  $E$  et  $x \in E$ , on définit la distance  $d(x, X)$  de  $x$  à  $X$  par

$$d(x, X) = \inf_{y \in X} d(x, y) ,$$

où  $d$  est la distance associée à la norme de  $E$ .

**Théorème 41.1** (Théorème de projection orthogonale). *Soit  $(E, \langle \cdot, \cdot \rangle)$  un espace de Hilbert et soit  $F \subset E$  un sous espace vectoriel fermé de  $E$ . Il existe une unique application  $P_F : E \rightarrow F$ , dite projection orthogonale sur  $F$ , qui est telle que pour tout  $x \in E$ ,*

$$d(x, F) = d(x, P_F(x)) ,$$

où  $d$  est la distance associée au produit scalaire de  $E$ . De plus,  $P_F(x)$  est entièrement caractérisé par la condition que  $P_F(x) \in F$  et la condition que  $x - P_F(x) \in F^\perp$ , où  $F^\perp$  est l'orthogonal de  $F$ .

*Démonstration.* Soit  $\delta = d(x, F)$ . Par définition de  $d(x, F)$  il existe une suite  $(y_n)_n$  de points dans  $F$  telle que

$$\delta \leq d(x, y_n) \leq \delta + \frac{1}{n} \quad (41.1)$$

pour tout  $n \geq 1$ . Soit  $\|\cdot\|$  la norme associée au produit scalaire. En écrivant que  $y_m + y_n - 2x = (y_m - x) + (y_n - x)$  et que  $y_m - y_n = (y_m - x) - (y_n - x)$ , puis en développant  $\|u + v\|^2 = \|u\|^2 + \|v\|^2 + 2\langle u, v \rangle$ , on remarque que pour tous  $m, n \geq 1$ ,

$$2\|y_n - x\|^2 + 2\|y_m - x\|^2 = \|y_m + y_n - 2x\|^2 + \|y_m - y_n\|^2 . \quad (41.2)$$

Or  $z_{m,n} = \frac{y_m + y_n}{2} \in F$  puisque  $F$  est un sous espace vectoriel et donc

$$\|y_m + y_n - 2x\|^2 = 4d(x, z_{m,n})^2 \geq 4\delta^2 .$$

Avec (41.1) et (41.2) on obtient donc que

$$2(\delta + \frac{1}{m})^2 + 2(\delta + \frac{1}{n})^2 \geq 4\delta^2 + \|y_m - y_n\|^2$$

pour tous  $m, n \geq 1$ , et donc  $(y_n)_n$  est de Cauchy dans  $F$ . Comme  $F$  est un fermé dans un Hilbert (donc aussi un Banach),  $(F, \|\cdot\|)$  est un Banach en vertu du Lemme 41.1. Donc  $y_n \rightarrow y$  pour un  $y \in F$  et, clairement,  $d(x, y) = d(x, F)$  par (41.1). Supposons maintenant que  $y, z \in F$  sont tels que  $d(x, y) = d(x, z) = d(x, F)$ . Comme pour l'établissement de (41.2),

$$2\|y - x\|^2 + 2\|z - x\|^2 = \|y + z - 2x\|^2 + \|z - y\|^2 . \quad (41.3)$$

Et comme  $\frac{y+z}{2} \in F$ , on récupère que  $4\delta^2 \geq 4\delta^2 + \|z - y\|^2$ , donc que  $z = y$ . Ainsi il existe un et un seul  $y \in F$  qui est tel que  $d(x, y) = d(x, F)$ . On pose alors  $P_F(x) = y$ . On a donc,

$$\|x - P_F(x)\| \leq \|x - z\| \quad (41.4)$$

pour tout  $z \in F$ . Soit  $\tilde{z} \in F$  quelconque. Comme  $F$  est un sous espace vectoriel,  $z = P_F(x) + \tilde{z}$  est encore dans  $F$ . Donc

$$\|x - P_F(x) - \tilde{z}\|^2 \geq \|x - P_F(x)\|^2$$

et en développant on trouve que

$$\|\tilde{z}\|^2 \geq 2\langle x - P_F(x), \tilde{z} \rangle$$

pour tout  $\tilde{z} \in F$ . En particulier, étant donné  $\tilde{z} \in F$ , on obtient que pour tout  $t \in \mathbb{R}$ ,

$$t^2 \|\tilde{z}\|^2 \geq 2t \langle x - P_F(x), \tilde{z} \rangle.$$

Une telle inégalité n'est possible pour tout  $t$  que si  $\langle x - P_F(x), \tilde{z} \rangle = 0$ . Donc  $x - P_F(x) \in F^\perp$ . Réciproquement, si  $z \in F$  est tel que  $x - z \in F^\perp$  alors pour tout  $y \in F$ ,

$$\begin{aligned} \|x - y\|^2 &= \|x - z + z - y\|^2 \\ &= \|x - z\|^2 + \|z - y\|^2 + 2\langle x - z, z - y \rangle \\ &= \|x - z\|^2 + \|z - y\|^2 \quad (\text{car } z - y \in F) \\ &\geq \|x - z\|^2 \end{aligned}$$

et on obtient que  $\|x - z\| \leq \|x - y\|$  pour tout  $y \in F$  et donc que  $d(x, z) = d(x, F)$  puisque  $z \in F$ . Le théorème est démontré.  $\square$

Du théorème de projection dans le cas des sous espaces vectoriels on déduit facilement le théorème du supplémentaire orthogonal dans les Hilbert.

**Théorème 41.2.** *Soit  $(E, \langle \cdot, \cdot \rangle)$  un espace de Hilbert. Soit  $F \subset E$  un sous espace vectoriel fermé de  $E$ . Alors  $E = F \oplus F^\perp$ , où  $F^\perp$  est l'orthogonal de  $F$ .*

*Démonstration.* On a toujours  $F \cap F^\perp = \{0\}$ . Il reste donc "uniquement" à montrer que  $E = F + F^\perp$ . Comme  $E$  est un Hilbert et  $F$  est fermé dans  $E$ , la projection orthogonale  $P_F : E \rightarrow F$  existe en vertu du Théorème 41.1. Pour tout  $x \in E$ ,

$$x = P_F(x) + (x - P_F(x)),$$

et comme  $P_F(x) \in F$  et  $x - P_F(x) \in F^\perp$ , on obtient que  $E = F + F^\perp$ . D'où le théorème.  $\square$

## 42. LE PREMIER THÉORÈME DE RIESZ

Soit  $(E, \|\cdot\|, \langle \cdot, \cdot \rangle)$  un espace préhilbertien. Soit aussi  $f \in L(E, \mathbb{R})$  une forme linéaire sur  $E$ . Si  $E$  est de dimension finie il est clair qu'il existe  $a \in E$ , unique, qui est tel que

$$f(x) = \langle a, x \rangle$$

pour tout  $x \in E$ . Pour le voir il suffit de considérer  $(e_1, \dots, e_n)$  une base orthonormée de  $(E, \|\cdot\|, \langle \cdot, \cdot \rangle)$  puis de remarquer que pour tout  $x \in E$ ,  $f(x) = \langle a, x \rangle$  si  $a$  est le vecteur de coordonnées les  $f(e_i)$  dans la base  $(e_1, \dots, e_n)$ . Cela démontre l'existence d'un tel  $a$ . Pour l'unicité on remarque que si  $\langle a, x \rangle = \langle a', x \rangle$  pour tout  $x \in E$ , alors  $\langle a' - a, x \rangle = 0$  pour tout  $x \in E$  et donc, en particulier, pour  $x = a' - a$ ,

ce qui implique que  $a' = a$ . En dimension infinie cette preuve s'effondre mais le résultat perdure dans le cas hilbertien des formes linéaires continues. On sait donc caractériser toutes les formes linéaires continues dans le cadre hilbertien.

**Théorème 42.1** (Premier théorème de Riesz). *Soit  $(E, \langle \cdot, \cdot \rangle)$  un espace de Hilbert et soit  $f : E \rightarrow \mathbb{R}$  une forme linéaire continue sur  $E$ . Il existe alors  $a \in E$  tel que  $f(x) = \langle a, x \rangle$  pour tout  $x \in E$ . De plus  $a$  est unique.*

*Démonstration.* Sans perdre en généralité on peut supposer que  $f$  n'est pas identiquement nulle. Soit  $F = \text{Ker}(f)$ . Comme  $f$  est continue,  $F$  est un sous espace vectoriel fermé de  $E$  (ce que l'on vérifie facilement avec le Lemme 36.2). Donc, en vertu du Théorème 41.2,  $E = F \oplus F^\perp$ . Soit  $b \in F^\perp$  tel que  $\|b\| = 1$ , où  $\|\cdot\|$  est la norme associée au produit scalaire. Un tel  $b$  existe car  $F \neq E$  dès que  $f$  n'est pas identiquement nulle. Comme  $b \neq 0$  on a que  $f(b) \neq 0$  (car sinon  $b \in F \cap F^\perp$ ). Pour tout  $x \in E$  il existe alors  $\lambda \in \mathbb{R}$  tel que  $f(x) = \lambda f(b)$  qui est une égalité dans  $\mathbb{R}$ . En particulier, pour tout  $z \in F^\perp$  il existe  $\lambda \in \mathbb{R}$  tel que  $f(z) = \lambda f(b)$ . Mais alors  $z - \lambda b \in F \cap F^\perp$  et ainsi  $z = \lambda b$ . Donc  $F^\perp$  est de dimension un. Soit  $\lambda_0 = f(b)$  et  $a = \lambda_0 b$ . Alors  $a$  est un vecteur directeur (une base) de  $F^\perp$  et  $f(a) = \|a\|^2$ . Tout  $x \in E$  s'écrit  $x = y + z$  avec  $y \in F$  et  $z \in F^\perp$ . Donc tout  $x \in E$  s'écrit  $x = y + \lambda a$  avec  $y \in F$  et  $\lambda \in \mathbb{R}$ . On a alors  $f(x) = \lambda f(a)$  puisque  $f(y) = 0$  tandis que

$$\langle x, a \rangle = \langle y, a \rangle + \lambda \|a\|^2 = \lambda f(a)$$

puisque  $y \in F$ ,  $a \in F^\perp$  et  $f(a) = \|a\|^2$ . D'où  $f(x) = \langle a, x \rangle$  pour tout  $x \in E$ . Pour montrer l'unicité supposons que  $a, a' \in E$  sont tels que  $f(x) = \langle a, x \rangle$  et  $f(x) = \langle a', x \rangle$  pour tout  $x \in E$ . Alors  $\langle a' - a, x \rangle = 0$  pour tout  $x \in E$  et donc, en particulier,  $\|a' - a\| = 0$ . D'où  $a' = a$  et on a bien existence et unicité d'un  $a \in E$  pour lequel  $f(x) = \langle a, x \rangle$  pour tout  $x \in E$ . Le théorème est démontré.  $\square$

#### 43. APPLICATIONS MULTILINÉAIRES CONTINUES

On considère  $E_1, \dots, E_n, F$  des espaces vectoriels. On considère aussi

$$f : E_1 \times \dots \times E_n \rightarrow F$$

une application de  $E_1 \times \dots \times E_n$  dans  $F$ . L'application  $f$  est dite multilinéaire si elle est linéaire par rapport à chacune de ses variables. En d'autres termes,  $f$  est multilinéaire si pour tout point  $(a_1, \dots, a_n) \in E_1 \times \dots \times E_n$ , et tout  $i = 1, \dots, n$ , les applications partielles  $f_i = E_i \rightarrow F$  définies par

$$f_i(x) = f(a_1, \dots, a_{i-1}, x, a_{i+1}, \dots, a_n)$$

sont linéaires en tant qu'applications de  $E_i$  dans  $F$ . Si  $(E_1, \|\cdot\|_{E_1}), \dots, (E_n, \|\cdot\|_{E_n})$  sont des espaces vectoriels normés, on place sur  $E = E_1 \times \dots \times E_n$  l'une des deux normes équivalentes

$$\|X\| = \sqrt{\sum_{i=1}^n \|x_i\|_{E_i}^2} \quad \text{ou} \quad \|X\|' = \sum_{i=1}^n \|x_i\|_{E_i} \quad (\star)$$

où  $X = (x_1, \dots, x_n)$ . Elles sont équivalentes, comme on le voit par exemple en montrant la double inégalité

$$\sum_{i=1}^n \|x_i\|_{E_i} = \langle (\|x_i\|_{E_i})_i, (1)_i \rangle_{\mathbb{R}^n} \leq \sqrt{n} \sqrt{\sum_{i=1}^n \|x_i\|_{E_i}^2}, \text{ et}$$

$$\sqrt{\sum_{i=1}^n \|x_i\|_{E_i}^2} \leq \sqrt{n} \max_i \|x_i\|_{E_i} \leq \sqrt{n} \sum_{i=1}^n \|x_i\|_{E_i}$$

de sorte que

$$\frac{1}{\sqrt{n}} \|X\| \leq \|X\|' \leq \sqrt{n} \|X\|$$

pour tout  $X = (x_1, \dots, x_n)$  dans  $E_1 \times \dots \times E_n$ .

**Théorème 43.1.** Soient  $(E_1, \|\cdot\|_1), \dots, (E_n, \|\cdot\|_n)$ , et  $(F, \|\cdot\|_F)$  des espaces vectoriels normés. Soit aussi  $f : E_1 \times \dots \times E_n \rightarrow F$  une application multilinéaire. Les affirmations suivantes sont équivalentes:

- (i)  $f$  est continue en tout point de  $E_1 \times \dots \times E_n$ ,
- (ii)  $f$  est continue à l'origine  $(0, \dots, 0)$  de  $E_1 \times \dots \times E_n$ ,
- (iii)  $\exists M > 0$  tel que  $\|f(x_1, \dots, x_n)\|_F \leq M \|x_1\|_1 \times \dots \times \|x_n\|_n$  pour tout  $(x_1, \dots, x_n) \in E_1 \times \dots \times E_n$ ,

où la norme placée sur  $E_1 \times \dots \times E_n$  qui sert à définir la continuité est l'une des normes équivalentes  $(\star)$ .

*Démonstration.* Comme dans le cas linéaire on montre que (i)  $\Rightarrow$  (ii)  $\Rightarrow$  (iii)  $\Rightarrow$  (i). Déjà, il est clair que (i)  $\Rightarrow$  (ii). Supposons maintenant (ii). Alors

$$\forall \varepsilon > 0, \exists \eta > 0 / \forall (x_1, \dots, x_n), \|(x_1, \dots, x_n)\| < \eta$$

$$\Rightarrow \|f(x_1, \dots, x_n)\|_F < \varepsilon$$

où, par exemple,  $\|(x_1, \dots, x_n)\| = \sum_{i=1}^n \|x_i\|_i$ . On fixe  $\varepsilon = 1$ . Soit  $(x_1, \dots, x_n)$  un point quelconque de  $E_1 \times \dots \times E_n$ . Si l'un des  $x_i$  est nul, alors  $f(x_1, \dots, x_n) = 0$  (car  $f(0) = 0$  pour une application linéaire), et (iii) est trivialement vérifié par  $(x_1, \dots, x_n)$ . On suppose maintenant que  $x_i \neq 0$  pour tout  $i$ . Le point  $(y_1, \dots, y_n)$  défini par  $y_i = \frac{\eta}{2n\|x_i\|_i} x_i$  est alors tel que  $\|(y_1, \dots, y_n)\| < \eta$ . En particulier,  $\|f(y_1, \dots, y_n)\|_F \leq 1$ , et puisque  $f$  est multilinéaire,

$$f(y_1, \dots, y_n) = \left(\frac{\eta}{2n}\right)^n \frac{1}{\prod_{i=1}^n \|x_i\|_i} f(x_1, \dots, x_n).$$

On a ainsi obtenu que

$$\|f(x_1, \dots, x_n)\|_F \leq \left(\frac{2n}{\eta}\right)^n \prod_{i=1}^n \|x_i\|_i$$

pour tout point  $(x_1, \dots, x_n) \in E_1 \times \dots \times E_n$ . Donc (ii)  $\Rightarrow$  (iii). Pour finir on suppose (iii). Soient aussi  $X = (x_1, \dots, x_n)$  et  $Y = (y_1, \dots, y_n)$  deux points de  $E_1 \times \dots \times E_n$ . On écrit que

$$\begin{aligned} & f(y_1, \dots, y_n) - f(x_1, \dots, x_n) \\ &= f(y_1 - x_1, y_2, \dots, y_n) + f(x_1, y_2 - x_2, y_3, \dots, y_n) \\ &+ \dots + f(x_1, \dots, x_{n-1}, y_n - x_n). \end{aligned}$$

Par exemple, lorsque  $n = 2$ ,

$$f(y_1, y_2) - f(x_1, x_2) = f(y_1 - x_1, y_2) + f(x_1, y_2 - x_2) .$$

On fixe  $X = (x_1, \dots, x_n)$ . Si  $\|Y - X\| \leq 1$ , alors  $\|y_i\|_i \leq 1 + \|x_i\|_i$  pour tout  $i$ , et il existe ainsi  $K > 0$ ,  $K = 1 + \max_i \|x_i\|_i$ , tel que  $\|x_i\|_i \leq K$  et  $\|y_i\|_i \leq K$  pour tout  $i$ . Par suite, par inégalité triangulaire, et d'après (iii),

$$\begin{aligned} \|f(y_1, \dots, y_n) - f(x_1, \dots, x_n)\|_F &\leq MK^{n-1} \sum_{i=1}^n \|y_i - x_i\|_i \\ &= MK^{n-1} \|Y - X\| \end{aligned}$$

Il s'ensuit facilement que

$$\forall \varepsilon > 0, \exists \eta > 0 / \forall Y, \|Y - X\| < \eta \Rightarrow \|f(Y) - f(X)\|_F < \varepsilon .$$

Il suffit de choisir  $\eta$  tel que  $\eta \leq 1$  et  $\eta < \varepsilon / MK^{n-1}$ . Donc (iii)  $\Rightarrow$  (i). Le théorème est démontré.  $\square$

Etant donnés  $(E_1, \|\cdot\|_1), \dots, (E_n, \|\cdot\|_n)$ , et  $(F, \|\cdot\|_F)$  des espaces vectoriels normés, on note  $L_c(E_1, \dots, E_n; F)$  l'espace des applications multilinéaires continues de  $E_1, \dots, E_n$  dans  $F$ . L'espace hérite d'une norme définie par

$$\|f\|_{L_c(E_1, \dots, E_n; F)} = \sup_{x_i \in E_i, \|x_i\|_i \leq 1} \|f(x_1, \dots, x_n)\|_F .$$

On a alors que

$$\|f(x_1, \dots, x_n)\|_F \leq \|f\|_{L_c(E_1, \dots, E_n; F)} \|x_1\|_1 \times \dots \times \|x_n\|_n$$

pour tous  $(x_1, \dots, x_n) \in E_1 \times \dots \times E_n$ .

Une définition équivalente de  $\|f\|_{L_c(E_1, \dots, E_n; F)}$  est qu'il s'agit du plus petit des réels positifs  $M$  pour lesquels

$$\|f(x_1, \dots, x_n)\|_F \leq M \|x_1\|_1 \times \dots \times \|x_n\|_n$$

pour tous  $(x_1, \dots, x_n) \in E_1 \times \dots \times E_n$ . On pourra vérifier que le résultat suivant a lieu. Il se démontre comme dans le cas des applications linéaires.

**Théorème 43.2.** *Soient  $(E_1, \|\cdot\|_1), \dots, (E_n, \|\cdot\|_n)$ , et  $(F, \|\cdot\|_F)$  des espaces vectoriels normés. L'espace  $(L_c(E_1, \dots, E_n; F), \|\cdot\|_{L_c(E_1, \dots, E_n; F)})$  est un espace de Banach dès que  $(F, \|\cdot\|_F)$  est un espace de Banach.*

En dimension finie, les multilinéaires, comme les linéaires, sont toujours continues.

**Théorème 43.3.** *Soient  $(E_1, \|\cdot\|_1), \dots, (E_n, \|\cdot\|_n)$  et  $(F, \|\cdot\|_F)$  des espaces vectoriels normés. On suppose que les  $E_i$  sont de dimensions finies. Toute application multilinéaire  $f : E_1 \times \dots \times E_n \rightarrow F$  est alors continue.*

*Démonstration.* Il suffit de montrer qu'il existe  $M > 0$  tel que

$$\|f(x_1, \dots, x_n)\|_F \leq M \|x_1\|_1 \times \dots \times \|x_n\|_n$$

pour tout  $(x_1, \dots, x_n) \in E_1 \times \dots \times E_n$ . Notons  $m_k$  les dimensions des  $E_k$  et  $(e_1^k, \dots, e_{m_k}^k)$  des bases de  $E_k$ ,  $k = 1, \dots, n$ . Pour tout  $x_k \in E_k$  on note  $x_1^k, \dots, x_{m_k}^k$

les coordonnées de  $x_k$  dans la base  $(e_1^k, \dots, e_{m_k}^k)$ ,  $k = 1, \dots, n$ . Par équivalence des normes en dimension finie, pour tout  $k = 1, \dots, n$ , il existe  $C_k > 0$  tel que

$$\sqrt{\sum_{i=1}^{m_k} (x_i^k)^2} \leq C_k \|x_k\|_k$$

pour tout  $x_k \in E_k$ . Pour  $i_1 \in \{1, \dots, m_1\}, \dots, i_n \in \{1, \dots, m_n\}$ , on note maintenant

$$a_{i_1 \dots i_n} = f(e_{i_1}^1, \dots, e_{i_n}^n) .$$

Pour tout  $(x_1, \dots, x_n) \in E_1 \times \dots \times E_n$ , on a alors, par multilinéarité de  $f$ ,

$$f(x_1, \dots, x_n) = \sum_{i_1, \dots, i_n} x_{i_1}^1 \dots x_{i_n}^n a_{i_1 \dots i_n} ,$$

et donc, puisqu'on a toujours  $|x_i| \leq \sqrt{\sum_j x_j^2}$ , on obtient avec ce qui a été dit ci-dessus que

$$\begin{aligned} \|f(x_1, \dots, x_n)\|_F &= \left\| \sum_{i_1, \dots, i_n} x_{i_1}^1 \dots x_{i_n}^n a_{i_1 \dots i_n} \right\|_F \\ &\leq C_1 \dots C_n \sum_{i_1, \dots, i_n} \|a_{i_1 \dots i_n}\|_F \|x_1\|_1 \dots \|x_n\|_n \end{aligned}$$

En posant  $M = C_1 \dots C_n \sum \|a_{i_1 \dots i_n}\|_F$ , on obtient l'inégalité voulue. Le théorème est démontré.  $\square$

#### 44. L'ISOMÉTRIE NATURELLE $L_c(E, L_c(F, G)) \sim L_c(E, F; G)$ .

On commence par définir ce qu'est une isométrie vectorielle entre espaces vectoriels normés. Soient  $(E, \|\cdot\|_E)$  et  $(F, \|\cdot\|_F)$  deux espaces vectoriels normés et  $f$  une application linéaire de  $E$  dans  $F$ . On dit que  $f$  est une isométrie vectorielle de  $(E, \|\cdot\|_E)$  sur  $(F, \|\cdot\|_F)$  si:

- (i)  $f$  est un isomorphisme de  $E$  sur  $F$ , et
- (ii)  $\|f(x)\|_F = \|x\|_E$  pour tout  $x \in E$ .

En d'autres termes, une isométrie vectorielle est un isomorphisme qui préserve les normes.

Bien sûr  $f$  est continue (car, en particulier,  $\|f(x)\|_F \leq \|x\|_E$  pour tout  $x$ ), et  $f^{-1}$  préserve aussi les normes (i.e  $\|f^{-1}(y)\|_E = \|y\|_F$  pour tout  $y$ ). En particulier,  $f^{-1}$  est continue, et  $f$  est un isomorphisme bi-continu (qui préserve les normes).

**Théorème 44.1.** *Soient  $(E, \|\cdot\|_E)$ ,  $(F, \|\cdot\|_F)$ , et  $(G, \|\cdot\|_G)$  trois espaces vectoriels normés. L'application*

$$\Phi : L_c(E, F; G) \rightarrow L_c(E, L_c(F, G))$$

*qui à  $f \in L_c(E, F; G)$  associe  $\Phi(f) \in L_c(E, L_c(F, G))$ , définie par*

$$(\Phi(f)(x))(y) = f(x, y)$$

*pour tout  $x \in E$  et tout  $y \in F$ , est une isométrie vectorielle de  $(L_c(E, F; G), \|\cdot\|_{L_c(E, F; G)})$  sur  $(L_c(E, L_c(F, G)), \|\cdot\|_{L_c(E, L_c(F, G))})$ . En particulier, via  $\Phi$ , l'espace  $L_c(E, F; G)$  s'assimile naturellement à l'espace  $L_c(E, L_c(F, G))$ .*

On vérifie facilement que  $\Phi$  est bien bijective et que l'application réciproque de  $\Phi$  est l'application

$$\Phi^{-1} : L_c(E, L_c(F, G)) \rightarrow L_c(E, F; G)$$

qui à  $f \in L_c(E, L_c(F, G))$  associe  $\Phi^{-1}(f) \in L_c(E, F; G)$  définie par

$$\Phi^{-1}(f)(x, y) = (f(x))(y)$$

pour tout  $x \in E$  et tout  $y \in F$ .

On démontre maintenant le théorème.

*Démonstration.* On commence par montrer que  $\Phi$  et  $\Phi^{-1}$  sont bien définies. Soit  $f$  fixée,  $f \in L_c(E, F; G)$ . Clairement, puisque  $f$  est bilinéaire,  $\Phi(f)(x) \in L(F, G)$  (i.e.  $\Phi(f)(x)$  est linéaire de  $F$  dans  $G$ ) où  $(\Phi(f)(x))(y) = f(x, y)$ . De même, on vérifie facilement que  $\Phi(f) \in L_c(E, L(F, G))$ . Comme  $f$  est continue,

$$\|f(x, y)\|_G \leq \|f\|_{L_c(E, F; G)} \|x\|_E \|y\|_F .$$

On en déduit que

$$\|(\Phi(f)(x))(y)\|_G \leq (\|f\|_{L_c(E, F; G)} \|x\|_E) \|y\|_F .$$

et donc que  $\Phi(f)(x) \in L_c(F, G)$ , avec

$$\|\Phi(f)(x)\|_{L_c(F, G)} \leq \|f\|_{L_c(E, F; G)} \|x\|_E .$$

Ensuite, on tire de cette inégalité que  $\Phi(f) \in L_c((E, L_c(F, G)))$  et que

$$\|\Phi(f)\|_{L_c((E, L_c(F, G)))} \leq \|f\|_{L_c(E, F; G)} . \quad (\star)$$

Donc  $\Phi$  est bien une application de  $L_c(E, F; G)$  dans  $L_c(E, L_c(F, G))$ . Clairement,  $\Phi$  est linéaire en  $f$ . Avec  $(\star)$  on récupère alors que  $\Phi$  est continue et que  $\|\Phi\| \leq 1$ . De la même façon on montre que  $\Phi^{-1}$  est bien définie, qu'elle est linéaire continue, et que  $\|\Phi^{-1}\| \leq 1$ . Comme par ailleurs  $\Phi \circ \Phi^{-1} = Id$ , en passant aux normes, on trouve que

$$1 \leq \|\Phi^{-1}\| \times \|\Phi\| .$$

Du coup,  $\|\Phi\| = 1$  et  $\|\Phi^{-1}\| = 1$ . En écrivant que

$$\|f\|_{L_c(E, F; G)} = \|\Phi^{-1}(\Phi(f))\|_{L_c(E, F; G)} \leq \|\Phi^{-1}\| \times \|\Phi(f)\|_{L_c(E, L_c(F, G))} ,$$

on en déduit que  $\|\Phi(f)\|_{L_c(E, L_c(F, G))} = \|f\|_{L_c(E, F; G)}$  et donc que  $\Phi$  est bien une isométrie vectorielle.  $\square$

#### 45. PRINCIPE DE CONTRACTION DE BANACH-PICARD

On traite ici du principe de contraction de Banach-Picard. C'est ce principe de contraction qui est à la base de l'existence de solutions mild des équations de Navier-Stokes. Dire qu'une application bilinéaire  $B : E \times E \rightarrow E$ ,  $(E, \|\cdot\|_E)$  un espace de Banach, est continue (cf. Section 43) c'est dire qu'il existe  $M > 0$  tel que pour tout  $x, y \in E$ ,  $\|B(x, y)\|_E \leq M \|x\|_E \|y\|_E$ . Le plus petit de tels  $M$  est, par définition, la norme  $\|B\|_{L_c(E, E; E)}$  de  $B$ .

**Théorème 45.1.** *Soit  $(E, \|\cdot\|_E)$  un espace de Banach et soit  $B : E \times E \rightarrow E$  une application bilinéaire continue. On note  $\Theta = \|B\|_{L_c(E, E; E)}$ . Soit  $e \in E$  tel que  $\|e\|_E < 1/4\Theta$ . Il existe alors une et une seule solution  $x_e$  à l'équation*

$$x = e - B(x, x)$$

vérifiant que  $\|x_e\| \leq 2\|e\|_E$ . De plus, pour tout  $0 < \theta < 1/4\Theta$ , si  $e_1, e_2 \in E$  sont tels que  $\|e_1\|_E \leq \theta$  et  $\|e_2\|_E \leq \theta$ , alors  $\|x_{e_2} - x_{e_1}\|_E \leq M_\theta \|e_2 - e_1\|_E$ , où  $M_\theta = 1/(1 - 4\Theta\theta)$ .

*Démonstration.* On construit  $x_e$  comme limite d'une suite construite itérativement (ce qui est classique dans les résultats de type point fixe). On construit la suite  $(x_n)$  par  $x_0 = e$  puis, pour tout  $n \in \mathbb{N}$ , par l'itération

$$x_{n+1} = e - B(x_n, x_n) .$$

On montre par récurrence que  $\|x_n\|_E \leq 2\|e\|_E$  pour tout  $n$ . L'amorce est clairement vérifiée. Pour ce qui est de l'hérédité, si  $\|x_n\|_E \leq 2\|e\|_E$ , alors

$$\begin{aligned} \|x_{n+1}\|_E &\leq \|e\|_E + \|B(x_n, x_n)\|_E \\ &\leq \|e\|_E + \Theta\|x_n\|_E^2 \\ &\leq \|e\|_E + 4\Theta\|e\|_E^2 \\ &\leq 2\|e\|_E \end{aligned}$$

puisque  $\|e\|_E < 1/4\Theta$ . On écrit maintenant que

$$\begin{aligned} x_{n+1} - x_n &= e - B(x_n, x_n) - (e - B(x_{n-1}, x_{n-1})) \\ &= B(x_{n-1}, x_{n-1}) - B(x_n, x_n) \\ &= B(x_{n-1} - x_n, x_n) + B(x_{n-1}, x_{n-1} - x_n) \end{aligned}$$

et donc, en vertu de ce que l'on vient de démontrer,

$$\begin{aligned} \|x_{n+1} - x_n\|_E &= \|B(x_{n-1} - x_n, x_n) + B(x_{n-1}, x_{n-1} - x_n)\|_E \\ &\leq \|B(x_{n-1} - x_n, x_n)\|_E + \|B(x_{n-1}, x_{n-1} - x_n)\|_E \\ &\leq \Theta(\|x_{n-1}\|_E + \|x_n\|_E)\|x_n - x_{n-1}\|_E \\ &\leq 2\Theta\|e\|_E\|x_n - x_{n-1}\|_E . \end{aligned}$$

Par suite, si  $A = 2\Theta\|e\|_E$ , alors  $A < 1$  et

$$\|x_{n+1} - x_n\|_E \leq A^n \|x_1 - x_0\|_E$$

pour tout  $n$ . On en déduit par télescopage et inégalité triangulaire que pour tous  $p < q$  dans  $\mathbb{N}$ ,

$$\begin{aligned} \|x_q - x_p\|_E &\leq \left( \sum_{k=p}^{q-1} A^k \right) \|x_1 - x_0\|_E \\ &\leq \left( \frac{1 - A^q}{1 - A} - \frac{1 - A^p}{1 - A} \right) \|x_1 - x_0\|_E \\ &\leq \frac{A^p}{1 - A} \|x_1 - x_0\|_E \end{aligned}$$

et donc,  $(x_n)$  est une suite de Cauchy puisque  $A^p \rightarrow 0$  lorsque  $p \rightarrow +\infty$  (car  $0 \leq A < 1$ ). Par suite  $x_n$  converge puisque l'espace est de Banach. On note  $x_e$  sa limite. Par continuité de  $B$ , et passage à la limite dans l'équation de récurrence pour les  $x_n$ ,  $x_e = e - B(x_e, x_e)$ . On a donc bien une solution à notre équation. Reste à montrer l'unicité de cette solution et la dernière estimée de dépendance continue par rapport aux données initiales. Les deux découlent d'un même type

de calcul. Supposons que  $x, y$  sont deux solutions à notre équation, toutes deux vérifiant  $\|x\| \leq 2\|e\|_E$  et  $\|y\| \leq 2\|e\|_E$ . Alors

$$\begin{aligned} \|y - x\|_E &= \|B(y, y) - B(x, x)\|_E \\ &= \|B(y - x, x) + B(y, y - x)\|_E \\ &\leq \|B(y - x, x)\|_E + \|B(y, y - x)\|_E \\ &\leq \Theta(\|x\|_E + \|y\|_E)\|y - x\|_E \\ &\leq 4\Theta\|e\|_E\|y - x\|_E. \end{aligned}$$

Or  $4\Theta\|e\|_E < 1$ . Donc  $x = y$ . De même, si  $e_1, e_2 \in E$  sont tels que  $\|e_1\|_E \leq \theta$  et  $\|e_2\|_E \leq \theta$ , où  $0 < \theta < 1/4\Theta$ , alors

$$\begin{aligned} \|x_{e_2} - x_{e_1}\|_E &= \|e_2 - e_1 + B(x_{e_1}, x_{e_1}) - B(x_{e_2}, x_{e_2})\|_E \\ &= \|e_2 - e_1\|_E + \|B(x_{e_2} - x_{e_1}, x_{e_1}) + B(x_{e_2}, x_{e_2} - x_{e_1})\|_E \\ &\leq \|e_2 - e_1\|_E + \|B(x_{e_2} - x_{e_1}, x_{e_1})\|_E + \|B(x_{e_2}, x_{e_2} - x_{e_1})\|_E \\ &\leq \|e_2 - e_1\|_E + \Theta(\|x_{e_1}\|_E + \|x_{e_2}\|_E)\|x_{e_2} - x_{e_1}\|_E \\ &\leq \|e_2 - e_1\|_E + 4\Theta\theta\|x_{e_2} - x_{e_1}\|_E. \end{aligned}$$

D'où l'inégalité de dépendance continue voulue. Le théorème est démontré.  $\square$

EMMANUEL HEBEY, UNIVERSITÉ DE CERGY-PONTOISE, DÉPARTEMENT DE MATHÉMATIQUES,  
SITE DE SAINT-MARTIN, 2 AVENUE ADOLPHE CHAUVIN, 95302 CERGY-PONTOISE CEDEX, FRANCE  
Email address: [Emmanuel.Hebey@cyu.fr](mailto:Emmanuel.Hebey@cyu.fr)