

Université Paris-Dauphine
DU MI2E, 2ème année

Algèbre linéaire 3 :
produits scalaires, espaces euclidiens,
formes quadratiques.

Cours 2010/2011

Olivier Glass

Le polycopié qui suit peut avoir des différences notables avec le cours dispensé en amphithéâtre (qui seul fixe le programme de l'examen). Il comporte des passages qui ne seront pas traités en amphithéâtre, et a contrario dans ce dernier pourront être donnés des compléments ne figurant pas dans ces notes, les preuves être plus détaillées, etc.

Certains points plus délicats (en particulier certaines démonstrations) sont signalés par un symbole de virages dangereux. Ils peuvent être passés en première lecture (mais pas nécessairement en seconde...)

Table des matières

1	Espaces vectoriels	1
1.1	Premières définitions	1
1.2	Famille génératrice, libre, bases	4
1.3	Matrice d'application linéaire	6
1.4	Dual d'un espace vectoriel	8
1.5	Espaces vectoriels normés	10
2	Espaces euclidiens	11
2.1	Définitions	11
2.2	Exemples	15
2.2.1	Produit scalaire canonique de $\mathbb{R}^n$	15
2.2.2	Produit scalaire canonique de $\mathbb{C}^n$	15
2.2.3	Autres exemples d'espaces munis d'un produit scalaire : les fonctions continues	16
2.2.4	Un autre exemple d'espace euclidien	16
2.3	Premières propriétés des espaces euclidiens	17
2.3.1	Identité du parallélogramme	17
2.3.2	Angle non-orienté des deux vecteurs	18
2.3.3	Théorème de Pythagore	18
2.3.4	Orthogonalité et liberté	19
2.4	Bases orthonormées, procédé d'orthonormalisation de Schmidt	19
2.4.1	Motivation	19
2.4.2	Bases orthonormées	19
2.4.3	Procédé d'orthonormalisation de Schmidt	20
2.4.4	Structure des espaces euclidiens par les bases orthonormées. Equivalence avec $\mathbb{R}^n$	21

2.5	Projections	23
2.5.1	Définition	23
2.5.2	Application au problème des moindres carrés	25
2.6	Structure du dual d'un espace euclidien	27
2.7	Changement de base orthonormée et groupe orthogonal	28
2.7.1	Rappels sur matrices et les changements de coordonnées	28
2.7.2	Groupe orthogonal	29
2.7.3	Endomorphismes orthogonaux	30
3	Formes bilinéaires et quadratiques	31
3.1	Motivation	31
3.2	Formes bilinéaires	32
3.3	Formes quadratiques	34
3.4	Méthode de Gauss	36
3.5	Signature d'une forme quadratique	39
4	Théorème de diagonalisation	41
4.1	Motivation	41
4.2	Changement de base pour les matrices de forme bilinéaire	42
4.3	Vecteurs et valeurs propres	42
4.4	Théorème de diagonalisation	44
4.5	Critères de positivité des formes quadratiques	45
4.6	Deux applications de la diagonalisation	47
5	Applications	49
5.1	Optimisation de fonctions (conditions nécessaires et condition suffisante) . .	49
5.1.1	Outils en œuvre	49
5.1.2	Conditions d'optimalité	51
5.2	Coniques	51
5.3	Interprétation géométrique des groupes $O(2)$ et $O(3)$	53
5.3.1	Le groupe $O(2)$	53
5.3.2	Le groupe $O(3)$	54
6	Annexe : structures algébriques	55
6.1	Lois de composition	55
6.2	Groupes	56
6.3	Anneaux	57

6.4	Corps	58
-----	-----------------	----

1 Espaces vectoriels

Les trois premiers paragraphes de ce chapitre sont composés de rappels. Les preuves y seront donc en général omises.

1.1 Premières définitions

Définition 1. Soit $\mathbb{K} = \mathbb{R}$ ou $\mathbb{K} = \mathbb{C}$, que l'on appellera le “corps de base”. On appelle espace vectoriel (souvent abrégé e.v.) sur $\mathbb{K}$ un ensemble E doté des lois de composition suivantes :

- une loi interne, c'est-à-dire une opération $E \times E \rightarrow E$ qui à $(x, y) \in E \times E$ associe $x + y$, dite addition, telle que $(E, +)$ soit un groupe commutatif, c'est-à-dire que (voir aussi §6.2) :
 - elle est associative : $\forall x, y, z \in E, (x + y) + z = x + (y + z)$,
 - elle est commutative : $\forall x, y \in E, x + y = y + x$,
 - elle admet un élément neutre (noté 0_E ou plus simplement 0) : $\forall x \in E, x + 0_E = 0_E + x = x$,
 - tout élément admet un inverse (dit opposé) : $\forall x \in E, \exists y \in E, x + y = y + x = 0_E$; cet élément est noté $-x$,
- une loi externe sur E par les éléments de $\mathbb{K}$, c'est-à-dire une opération $\mathbb{K} \times E \rightarrow E$ qui à $(\lambda, x) \in \mathbb{K} \times E$ associe $\lambda \cdot x$, dite multiplication par un scalaire de $\mathbb{K}$, telle que :

$$\begin{aligned}\forall \lambda, \mu \in \mathbb{K}, \forall x \in E, \quad \lambda \cdot (\mu \cdot x) &= (\lambda\mu) \cdot x, \\ \forall \lambda, \mu \in \mathbb{K}, \forall x \in E, \quad (\lambda + \mu) \cdot x &= \lambda \cdot x + \mu \cdot x, \\ \forall x, y \in E, \quad \forall \lambda \in \mathbb{K}, \quad \lambda \cdot (x + y) &= \lambda \cdot x + \lambda \cdot y, \\ \forall x \in E, \quad 1 \cdot x &= x \text{ (où 1 est l'élément neutre pour la multiplication dans } \mathbb{K}).\end{aligned}$$

On notera dans toute la suite “ λx ” plutôt que “ $\lambda \cdot x$ ”. On appellera les éléments de $\mathbb{K}$ des **scalaires** et les éléments de E des **vecteurs**.

Exemples.

- $\mathbb{R}^n$ muni de l'addition et de la multiplication par un scalaire “coordonnée par coordonnée”

$$\begin{aligned}(a_1, \dots, a_n) + (b_1, \dots, b_n) &= (a_1 + b_1, \dots, a_n + b_n), \\ \lambda(a_1, \dots, a_n) &= (\lambda a_1, \dots, \lambda a_n).\end{aligned}$$

est un espace vectoriel sur $\mathbb{R}$.

- Le corps des complexes $E = \mathbb{C}$ est un espace vectoriel réel (sur $\mathbb{K} = \mathbb{R}$) ou complexe (sur $\mathbb{K} = \mathbb{C}$).
- L'espace $\mathbb{C}^n$ avec les mêmes opérations d'addition et de multiplication que ci-dessus (voir $\mathbb{R}^n$) est un espace vectoriel sur $\mathbb{R}$ ou sur $\mathbb{C}$.

- Soit X un ensemble quelconque (non vide), soit E un espace vectoriel sur $\mathbb{K}$. On introduit alors $\mathcal{A}(X; E)$ (noté encore X^E) l'ensemble des applications de X dans E . Muni des opérations standard d'addition et multiplication par un scalaire de fonctions (on fait les opérations "point par point"), ceci est aussi un espace vectoriel sur $\mathbb{K}$.
- L'espace des matrices à n lignes et p colonnes $M_{n,p}(\mathbb{K})$ à entrées dans $\mathbb{K}$ est un $\mathbb{K}$ -espace vectoriel.
- L'espace $\mathbb{K}[X]$ des polynômes à coefficients dans $\mathbb{K}$ est un $\mathbb{K}$ -espace vectoriel.

Proposition 1. *Pour tout $\lambda \in \mathbb{K}$ et $x \in E$, on a :*

- $\lambda 0 = 0$ (0 est ici l'élément neutre de E pour l'addition i.e. "le vecteur nul"),
- $0x = 0$ (0 est ici l'élément neutre de $\mathbb{K}$ pour l'addition i.e. le 0 de $\mathbb{R}$ ou $\mathbb{C}$),
- si $\lambda x = 0$ alors soit $\lambda = 0$, soit $x = 0$,
- $(-\lambda)x = \lambda(-x) = -(\lambda x)$.

Démonstration. Pour la première utiliser $0 + 0 = 0$ (dans E), pour la deuxième utiliser $0 + 0 = 0$ (dans $\mathbb{K}$), pour la troisième multiplier par λ^{-1} et pour la dernière vérifier les propriétés de l'inverse. $\square$

Définition 2. *Soient E un espace vectoriel sur $\mathbb{K}$ et $F \subset E$ un sous ensemble de E . On dit que F est un **sous-espace vectoriel** (souvent abrégé : s.e.v.) si et seulement si c'est une partie non vide et stable par les deux lois de E . Autrement dit F est un s.e.v. si et seulement si :*

- $F \neq \emptyset$
- $\forall x, y \in F, x + y \in F$
- $\forall x \in F, \lambda \in \mathbb{K}, \lambda x \in F$

On sait qu'on peut aussi abréger les conditions précédentes de la façon qui suit.

Proposition 2. *Soient E un espace vectoriel et F une partie de E . Alors F est un sous-espace vectoriel de E si et seulement si :*

- $0 \in F$,
- $\forall x, y \in F, \forall \lambda, \mu \in \mathbb{K}, \lambda x + \mu y \in F$.

Un des grands intérêts de la notion de sous-espace vectoriel est qu'elle permet de montrer facilement qu'un certain ensemble est muni d'une structure d'espace vectoriel. En effet on a le résultat suivant bien connu.

Proposition 3. *Soient E un espace vectoriel et $F \subset E$ un sous-espace vectoriel de E . Alors F muni des lois interne et externe de E devient lui-même un espace vectoriel sur $\mathbb{K}$.*

Exemples.

- Une droite passant par l'origine d'un espace vectoriel E : pour tout $x \in E \setminus \{0\}$: $\{\lambda x, \lambda \in \mathbb{K}\}$ est un sous espace vectoriel de E .

- Les sous-espaces engendrés (voir plus bas) sont des s.e.v.
- L'espace des fonctions continues de $\mathbb{R}$ dans $\mathbb{R}$ est un sous-espace vectoriel de l'espace $\mathcal{A}(\mathbb{R}; \mathbb{R})$. C'est donc lui-même un espace vectoriel.
- L'espace $\mathbb{K}_n[X]$ des polynômes à coefficients dans $\mathbb{K}$ et de degré au plus n est un sous-espace vectoriel de $\mathbb{K}[X]$.
- En revanche le groupe $GL(n; \mathbb{K})$ des matrices carrées de taille n inversibles, n'est pas un sous-espace vectoriel de l'espace $M_n(\mathbb{K})$ des matrices carrées de taille n .

Proposition 4. Soient F et G deux s.e.v de E . Alors :

- $F \cap G$ est un s.e.v. de E ,
- $F + G := \{x + y \text{ pour } x \in F \text{ et } y \in G\}$ est un s.e.v. de E ,
- mais $F \cup G$ n'est pas en général un s.e.v de E (faire un dessin dans $\mathbb{R}^2$!)

Définition 3. On dit que la somme de deux sous-espaces F et G d'un espace vectoriel E est **directe** et l'on note $F + G = F \oplus G$ lorsque tout x de $F + G$ s'écrit de manière unique comme somme d'un élément de F et d'un élément de G . Il est équivalent que $F \cap G = \{0\}$. Deux sous-espaces F et G sont dits **supplémentaires** lorsque leur somme est directe et que $E = F + G$

Définition 4. Soient E et F deux $\mathbb{K}$ -espaces vectoriels. Une application $T : E \rightarrow F$ est dite **linéaire** (ou $\mathbb{K}$ -linéaire) si elle satisfait aux relations :

- $\forall x, y \in E, T(x + y) = T(x) + T(y)$,
- $\forall \lambda \in \mathbb{K}, \forall x \in E, T(\lambda x) = \lambda T(x)$.

En abrégé, ces deux conditions peuvent s'écrire sous la forme :

$$\forall \lambda, \mu \in \mathbb{K}, \forall x, y \in E, T(\lambda x + \mu y) = \lambda T(x) + \mu T(y). \quad (1)$$

On notera $L(E, F)$ l'ensemble des applications linéaires de E dans F . C'est un sous-espace vectoriel de $\mathcal{A}(E, F)$. Lorsque $E = F$, on parle d'**endomorphisme** de E et on note $End(E) := L(E, E)$.

Exemples. L'application de $\mathbb{R}^3$ dans $\mathbb{R}^2$ qui à (x, y, z) associe $(x, y + 2z)$ est linéaire. Pas l'application de $\mathbb{R}^2$ dans $\mathbb{R}$ qui à (x, y) associe xy .

L'application de $\mathbb{C}^2$ dans $\mathbb{C}$ qui à (z, ζ) associe $\bar{z}$ est $\mathbb{R}$ -linéaire (linéaire entre $\mathbb{C}^2$ et $\mathbb{C}$ considérés comme $\mathbb{R}$ -espaces vectoriels), mais pas $\mathbb{C}$ -linéaire (linéaire entre $\mathbb{C}^2$ et $\mathbb{C}$ considérés comme $\mathbb{C}$ -espaces vectoriels). On observe en effet que dans ce cas (1) est satisfait pour $\lambda, \mu \in \mathbb{R}$, mais pas en général pour $\lambda, \mu \in \mathbb{C}$.

Proposition 5. Soient T une application linéaire de E dans F . Alors les parties

$$\begin{aligned} Ker(T) &:= \{x \in E / T(x) = 0\}, \\ Im(T) &:= \{y \in F / \exists x \in E, y = T(x)\}, \end{aligned}$$

appelées respectivement **noyau** et **image** de T , sont des sous-espaces vectoriels de E et F respectivement.

1.2 Famille génératrice, libre, bases

Définition 5. Soit E un e.v. sur $\mathbb{K}$, et soient p vecteurs de $E : x_1, \dots, x_p \in E$ et p scalaires $\lambda_1, \dots, \lambda_p \in \mathbb{K}$. On appelle **combinaison linéaire** des vecteurs x_i de coefficients λ_i ($i = 1, \dots, p$) le vecteur

$$\lambda_1 x_1 + \dots + \lambda_p x_p.$$

Définition 6. Soient E un e.v. sur $\mathbb{K}$, et $x_1, \dots, x_p \in E$ des vecteurs de E . On appelle **espace vectoriel engendré** par cette famille et on note $\text{Vect}\{x_1, \dots, x_p\}$, l'ensemble de toutes les combinaisons linéaires des vecteurs $x_1, \dots, x_p$:

$$\text{Vect}\{x_1, \dots, x_p\} = \left\{ \sum_{i=1}^p \lambda_i x_i, (\lambda_i)_{i=1, \dots, p} \in \mathbb{K}^p \right\}.$$

Proposition 6. L'espace vectoriel engendré par la famille $x_1, \dots, x_p \in E$ est le plus petit s.e.v. de E contenant les vecteurs $x_1, \dots, x_p$.

Démonstration. On procède en deux étapes : d'abord on observe que l'espace vectoriel engendré par la famille $x_1, \dots, x_p$ tel que défini plus haut est bien un s.e.v. de E ; ensuite qu'il est contenu dans tout s.e.v. de E qui contient $x_1, \dots, x_p$. $\square$

Définition 7. On dit que $\{x_1, \dots, x_p\}$ est un **système de générateurs** ou encore une **famille génératrice** de E si tout élément de E peut se mettre sous la forme d'une combinaison linéaire de $x_1, \dots, x_p$, i.e. si

$$\text{Vect}\{x_1, \dots, x_p\} = E.$$

On dira également que $x_1, \dots, x_p$ engendrent E .

Interprétation en termes de système linéaire dans $\mathbb{R}^n$. Si $x_1, \dots, x_p$ sont des éléments de $\mathbb{R}^n$, dire qu'ils forment une famille génératrice signifie que le système linéaire de n équations à p inconnues $\lambda_1, \dots, \lambda_p$:

$$\lambda_1 x_1 + \dots + \lambda_p x_p = u,$$

admet toujours **au moins** une solution quel que soit u de $\mathbb{R}^n$.

Exemple. La famille $\{(1, 0, 0), (0, 1, 0), (0, 0, 1)\}$ est un système de générateurs de $\mathbb{R}^3$ car tout élément $(a, b, c) \in \mathbb{R}^3$ est combinaison linéaire des éléments de la famille :

$$(a, b, c) = a(1, 0, 0) + b(0, 1, 0) + c(0, 0, 1).$$

Définition 8. Soient E un e.v. sur $\mathbb{K}$, et $x_1, \dots, x_p \in E$ des vecteurs de E . On dit que la famille $\{x_1, \dots, x_p\}$ est **liée** (ou que les vecteurs $x_1, \dots, x_p$ sont **linéairement dépendants**), s'il existe une combinaison linéaire de ces vecteurs, à coefficients non tous nuls qui soit égale à zéro, i.e.

$$\exists \lambda_1, \dots, \lambda_p \in \mathbb{K} \text{ tels que } \sum_{i=1}^p \lambda_i x_i = 0, \text{ et } \sum_{i=1}^p |\lambda_i| > 0.$$

Remarque 1. Dans la définition précédente on peut avoir **certain**s λ_i nuls. La définition stipule seulement qu'il existe au moins un non nul.

Exemple. Soient E un e.v. et $x_1, x_2, x_3 \in E$. Alors la famille $\{x_1, x_2, x_3, 7x_1 + 2x_2\}$ est linéairement dépendante.

Définition 9. Soient E un e.v. sur $\mathbb{K}$, et $x_1, \dots, x_p \in E$ des vecteurs de E . On dit que la famille $\{x_1, \dots, x_p\}$ est **libre** (ou encore que les vecteurs $x_1, \dots, x_p$ sont **linéairement indépendants**), si cette famille n'est pas liée, autrement dit, si la seule combinaison linéaire des vecteurs $x_1, \dots, x_p$ qui soit nulle, est celle dont les coefficients sont tous nuls. Ce que l'on peut encore exprimer sous la forme : si $\sum_{i=1}^p \lambda_i x_i = 0$ alors $\lambda_1 = \dots = \lambda_p = 0$.

Interprétation en termes de système linéaire dans $\mathbb{R}^n$. Si $x_1, \dots, x_p$ sont des éléments de $\mathbb{R}^n$, dire qu'ils forment une famille libre signifie que le système linéaire de n équations à p inconnues $\lambda_1, \dots, \lambda_p$:

$$\lambda_1 x_1 + \dots + \lambda_p x_p = u,$$

admet toujours **au plus** une solution quel que soit u de $\mathbb{R}^n$. En effet, s'il y en a deux différentes, en faisant la différence entre les deux équations, on trouve une combinaison nulle des vecteurs $x_1, \dots, x_p$, à coefficients non tous nuls.

Définition 10. Soient E un e.v. sur $\mathbb{K}$ et $X \subset E$ un ensemble de vecteurs de E . On dit que l'ensemble X est une **base** de E s'il forme à la fois une famille libre et génératrice.

Théorème 1. Soient E un e.v. sur $\mathbb{K}$ et $X \subset E$ un ensemble de vecteurs de E . L'ensemble X est une base de E si et seulement si tout vecteur de E s'écrit de façon unique comme combinaison linéaire d'éléments de X .

Démonstration. Condition nécessaire. Si X est une base et s'il y a deux façons d'écrire un vecteur y comme combinaisons linéaires des éléments de X , on fait la différence entre ces deux écritures et on déduit que X n'est pas libre. Par ailleurs, comme X est génératrice, on déduit qu'il y a au moins une écriture.

Condition suffisante. Si la propriété de droite est satisfaite, alors X est de toute évidence d'une famille génératrice; l'unicité de l'écriture donne le fait que la famille est libre (car l'unique écriture de 0 comme combinaison des éléments de X est la combinaison à coefficients nuls). $\square$

Définition 11. Soient E un e.v. sur $\mathbb{K}$, et $X = \{x_1, \dots, x_n\}$ une base de E . Pour tout $x \in E$ les coefficients λ_i de l'écriture

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

s'appellent les coordonnées de x dans la base X .

Théorème 2. Soit E un e.v. sur $\mathbb{K}$. On suppose que E est engendré par un nombre fini de vecteurs. Alors toutes les bases de E ont le même nombre d'éléments. Cet nombre s'appelle **la dimension** de l'espace vectoriel E .

Nous ne refaisons pas la preuve de ce théorème fondamental, mais rappelons qu'un lemme central (parfois appelé lemme de Steinitz) est le suivant.

Lemme 1. Soient E un espace vectoriel sur $\mathbb{K}$, et $\{x_1, \dots, x_p\}$ une famille génératrice de E . Alors toute famille d'au moins $p + 1$ vecteurs est liée.

Exemples. L'espace $\mathbb{R}^n$ est de dimension n . L'espace $\mathbb{R}_n[X]$ est de dimension $n + 1$ puisque $1, X, X^2, \dots, X^n$ en forme une base. L'espace $\mathbb{R}[X]$ est de dimension infinie, puisqu'il contient une famille infinie de vecteurs libres (les X^k).

Terminons ce paragraphe par un théorème très utile en algèbre linéaire et deux de ses conséquences.

Théorème 3 (de la base incomplète). Soient E un espace vectoriel sur $\mathbb{K}$ de dimension finie, et $A = \{x_1, \dots, x_p\} \subset B = \{x_1, \dots, x_n\}$ (où $p \leq n$) deux familles de vecteurs de E . On suppose que A est une famille libre et que B est une famille génératrice. Alors il existe une base de E contenant A et incluse dans B . Autrement dit, on peut piocher parmi $\{x_{p+1}, \dots, x_n\}$ pour compléter $A = \{x_1, \dots, x_p\}$ en une base.

Une conséquence bien connue et importante du théorème de la base incomplète est le suivant.

Théorème 4 (du rang). Soit $T \in L(E, F)$, où E est de dimension finie. Alors on a :

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

Corollaire 1. Soit E un $\mathbb{K}$ -espace vectoriel de dimension finie. Alors pour $T \in \text{End}(E)$, on a :

$$T \text{ injectif} \iff T \text{ surjectif} \iff T \text{ bijectif}.$$

1.3 Matrice d'application linéaire

Définition 12. Soient E et F deux espaces vectoriels de dimension finie, de base respective $\{e_1, \dots, e_n\}$ et $\{f_1, \dots, f_d\}$. Soit $T : E \rightarrow F$ une application linéaire. Pour chaque k entre 1 et n , le vecteur $T(e_k)$ de F se décompose dans la base $\{f_1, \dots, f_d\}$ selon

$$T(e_k) = \sum_{j=1}^d a_{jk} f_j.$$

On appelle **matrice de T dans les bases $\{e_1, \dots, e_n\}$ et $\{f_1, \dots, f_d\}$** , le tableau à n colonnes et d lignes, dont l'entrée (c'est-à-dire l'élément) à la k -ème colonne et j -ème ligne est a_{jk} . On la note aussi parfois $(a_{jk})_{j=1\dots d, k=1\dots n}$.

Convention. Quand on désigne l'entrée a_{ij} d'une matrice, le premier indice i correspond à la ligne, le second j à la colonne. On note $M_{d,n}(\mathbb{K})$ l'espace des matrices à d lignes et n colonnes, dont les entrées sont des éléments de $\mathbb{K}$.

Autrement dit, dans la k -ième colonne (k allant de 1 à n), on range les d coordonnées de $T(e_k)$ dans la base $\{f_1, \dots, f_d\}$.

Remarque 2. Si $E = F$ et que T est donc un endomorphisme, on choisit souvent la même base $\{e_1, \dots, e_n\}$ "au départ" et à l'arrivée. Dans ce cas, on parle de matrice de l'endomorphisme T dans la base $\{e_1, \dots, e_n\}$. Il s'agit bien sûr, dans ce cas, d'une matrice « carrée », c'est-à-dire ayant le même nombre de lignes et de colonnes.

Exemple 1. Soit $E := \mathbb{R}_1[X]$ l'espace des polynômes de degré inférieur ou égal à 1. On vérifie sans peine que $\{X, X - 1\}$ est une base de E . On considère alors l'endomorphisme T de E qui à P associe le polynôme $Q(X) := P(X + 1)$. On voit sans peine que $T(X) = X + 1 = 2X - 1 \cdot (X - 1)$ tandis que $T(X - 1) = X = 1 \cdot X + 0 \cdot (X - 1)$. La matrice de T dans la base $\{X, X - 1\}$ est donc

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

Définition 13. Soit $A = (a_{ij}) \in M_{p,q}(\mathbb{K})$ et $B = (b_{jk}) \in M_{q,r}(\mathbb{K})$. On définit le **produit matriciel** de A et B , dans cet ordre, comme la matrice $C \in M_{p,r}(\mathbb{K})$, dont les coefficients c_{ik} pour $i = 1 \dots p$ et $k = 1 \dots r$ sont

$$c_{ik} = \sum_{j=1}^q a_{ij} b_{jk}.$$

Proposition 7. Soient E , F et G trois espaces vectoriels de base respective $\{e_1, \dots, e_l\}$, $\{f_1, \dots, f_m\}$ et $\{g_1, \dots, g_n\}$. Soient $S \in L(E, F)$ et $T \in L(F, G)$, de matrice respective M et N dans ces bases. Alors la matrice de $T \circ S \in L(E, G)$ dans ces bases est la matrice produit $N \cdot M$.

Définition 14. Soit E un espace vectoriel de dimension finie, et soient $\{e_1, \dots, e_n\}$ et $\{f_1, \dots, f_n\}$ deux bases de E . On appelle **matrice de passage de la base $\{e_1, \dots, e_n\}$ vers la base $\{f_1, \dots, f_n\}$** , la matrice carrée de taille n , dont la k -ième colonne contient les coordonnées du vecteur f_k dans la base $\{e_1, \dots, e_n\}$.

Remarque 3. On retiendra qu'on exprime les « nouveaux » vecteurs dans « l'ancienne base ».

Exemple 2. Soit $E := \mathbb{R}_1[X]$ l'espace des polynômes de degré inférieur ou égal à 1. On vérifie que $\{1, X\}$ et $\{2X, 3X - 1\}$ sont deux bases de E . On voit que la matrice de passage de $\{1, X\}$ vers $\{2X, 3X - 1\}$ est

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

Proposition 8. Soit E un espace vectoriel de dimension finie, et soient $\{e_1, \dots, e_n\}$ et $\{f_1, \dots, f_n\}$ deux bases de E . Soit x un vecteur E , de coordonnées $(x_1, \dots, x_n)$ dans la base $\{e_1, \dots, e_n\}$ et $(y_1, \dots, y_n)$ dans la base $\{f_1, \dots, f_n\}$. Soit P la matrice de passage de $\{e_1, \dots, e_n\}$ vers $\{f_1, \dots, f_n\}$. Alors on a :

$$\begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} = P \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix}.$$

Remarque 4. Le produit matriciel naturel donne donc les **anciennes coordonnées** en fonction des nouvelles.

Proposition 9. Soit E un espace vectoriel de dimension finie, soient $\{e_1, \dots, e_n\}$ et $\{f_1, \dots, f_n\}$ deux bases de E , et soit P la matrice de passage de $\{e_1, \dots, e_n\}$ vers $\{f_1, \dots, f_n\}$. Soit T un endomorphisme de E , de matrice A dans la base $\{e_1, \dots, e_n\}$. Alors sa matrice $\tilde{A}$ dans la base $\{f_1, \dots, f_n\}$ est

$$\tilde{A} = P^{-1}AP.$$

1.4 Dual d'un espace vectoriel

Définition 15. Soit E un e.v. sur $\mathbb{K}$. On appelle **dual** de E et on note par E^* ou par E' l'ensemble des applications linéaires sur E à valeurs dans $\mathbb{K}$. Autrement dit, $E^* := L(E, \mathbb{K})$.

Remarque 5. On appelle **forme linéaire** une application linéaire dont l'espace d'arrivée est $\mathbb{K}$. On fera attention à distinguer cette expression du cas général où l'on parle « d'application » linéaire.

Exemples. Soit $E = \mathbb{R}^3$; alors $\mu(x, y, z) = x - y$ est une forme linéaire sur E donc $\mu \in E^*$. De même pour $\nu(x, y, z) = x + y - 2z$. En revanche, $\ell(x, y, z) = x + 1$ et $m(x, y, z) = x + |z|$ ne sont pas des formes linéaires.

Remarque 6. Pour $f \in E^*$ et $x \in E$ on notera également $f(x)$ par $\langle f, x \rangle_{E^*, E}$ ou encore par $\langle x, f \rangle_{E, E^*}$.

Théorème 5. Soit E un e.v. sur $\mathbb{K}$. Alors :

- le dual E^* est aussi un e.v. sur $\mathbb{K}$.
- si de plus E est **de dimension finie** alors E^* est aussi de dimension finie; mieux, $\dim(E^*) = \dim(E)$.

Démonstration.

Pour le premier point, il suffit d'observer que pour tous $f, g \in E^*$, $\alpha, \beta \in \mathbb{K}$ la fonction $\alpha f + \beta g$ définie sur E à valeurs en $\mathbb{K}$ est encore linéaire et à valeurs dans $\mathbb{K}$; donc $\alpha f + \beta g \in E^*$.

Pour le second point, soit $\{e_i\}_{i=1}^n$ une base de E ; on définit les applications $\mu_i : E \rightarrow \mathbb{K}$ de la manière suivante :

$$\mu_i \left(\sum_{k=1}^n x_k e_k \right) = x_i.$$

On vérifie qu'il s'agit bien de formes linéaires. De plus, on a que $\mu_i(e_j) = \delta_{ij}$, pour tous $i, j = 1, \dots, n$. (Rappelons que symbole de Kronecker δ_{ij} vaut 0 si $i \neq j$ et 1 sinon.)

Il suffit maintenant de montrer que la famille $\{\mu_i\}_{i=1}^n$ forme une base de E^* . D'une part, du fait des valeurs des μ_j sur les vecteurs e_j , elles forment une famille libre. D'autre part n'importe quelle forme linéaire $f \in E^*$ s'écrit comme combinaison linéaire des μ_i :

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

En effet, on sait en effet que les valeurs d'une application linéaire est déterminée par ses valeurs sur une base ; les deux formes linéaires coïncident ici sur les (e_i) . $\square$

Définition 16. La base $(\mu_1, \dots, \mu_n)$ de E^* ainsi déterminée est appelée **base duale** de la base (e_j) . Elle est parfois notée $(e_1^*, \dots, e_n^*)$.

Définition 17. Soient E un e.v. de dimension finie sur $\mathbb{K}$ et F un s.e.v. de E . L'**annulateur** F^0 de F est l'ensemble des formes linéaires $f \in E^*$ s'annulant sur F :

$$F^0 = \{f \in E^* \mid \forall x \in F, f(x) = 0\}.$$

Remarque 7. L'annulateur F^0 de F est un s.e.v. de E^* : en effet, toute combinaison linéaire $\alpha f + \beta g$ avec $f, g \in F^0$ s'annule sur F .

Exemples. Soit $E = \mathbb{R}^3$; $\mu(x, y, z) = x - y$ et $\nu(x, y, z) = x + y - 2z$ sont dans l'annulateur de $F = \text{Vect} \{(1, 1, 1)\}$ vu comme s.e.v. de $E = \mathbb{R}^3$.

Théorème 6. Soient E un e.v. de dimension finie sur $\mathbb{K}$ et F un s.e.v. de E . Alors

$$\dim(F^0) = \dim(E) - \dim(F).$$

Démonstration. On introduit une base de F , mettons $(e_1, \dots, e_k)$ (avec $k = \dim(F)$). On complète cette base de F à une base de E , mettons que nous obtenons $(e_1, \dots, e_n)$ (avec $n = \dim(E)$). Maintenant, on introduit la base duale $(e_1^*, \dots, e_n^*)$ comme décrit précédemment. On voit alors que

$$F^0 = \text{Vect} \{e_{k+1}^*, \dots, e_n^*\}.$$

En effet, il est clair que $e_{k+1}^*, \dots, e_n^*$ sont dans F^0 . Par ailleurs, un élément ℓ de F^0 ne peut avoir de composante selon $e_1^*, \dots, e_k^*$, comme on le voit en regardant $\ell(e_1) = \dots = \ell(e_k) = 0$. Donc $\dim(F^0) = n - k$. $\square$

1.5 Espaces vectoriels normés

Soit E un espace vectoriel sur $\mathbb{K} = \mathbb{R}$ ou $\mathbb{K} = \mathbb{C}$. Il est fréquemment utile de pouvoir mesurer la “taille” d’un vecteur donné ou la distance entre deux vecteurs de l’espace E . Pour cela on introduit la notion de **norme**.

Définition 18. Une **norme** est une application de E dans $\mathbb{R}_+$ qui à tout vecteur $x \in E$ associe le nombre noté $\|x\| \geq 0$ vérifiant :

$$\|x\| = 0 \iff x = 0, \quad (2)$$

$$\forall \alpha \in \mathbb{K}, \forall x \in E : \|\alpha x\| = |\alpha| \|x\|, \quad (3)$$

$$\forall x, y \in E : \|x + y\| \leq \|x\| + \|y\| \quad (\text{inégalité triangulaire}). \quad (4)$$

On appelle **espace vectoriel normé** (en abrégé, *e.v.n.*) un espace vectoriel muni d’une norme.

Un espace vectoriel normé est un cas particulier d’espace métrique avec la **distance** entre deux éléments donnée par :

$$d(x, y) = \|x - y\|.$$

On appelle également $\|x\|$ la longueur du vecteur x .

Interprétation géométrique. Lorsqu’on est dans $\mathbb{R}^2$ avec la norme

$$\|(x, y)\| = \sqrt{x^2 + y^2},$$

$\|x + y\|$, $\|x\|$ et $\|y\|$ sont les longueurs des cotés d’un triangle dans $\mathbb{R}^2$ (mettons $x = \overrightarrow{AB}$, $y = \overrightarrow{BC}$, et donc $x + y = \overrightarrow{AC}$), d’où l’appellation d’inégalité triangulaire.

Exemples. On peut munir $\mathbb{R}^n$ (ou $\mathbb{C}^n$) de plusieurs normes différentes comme :

$$\begin{aligned} \|(x_1, \dots, x_n)\| &:= |x_1| + \dots + |x_n|, \\ \|(x_1, \dots, x_n)\| &:= \sqrt{|x_1|^2 + \dots + |x_n|^2}, \\ \|(x_1, \dots, x_n)\| &:= \max\{|x_1|, \dots, |x_n|\}. \end{aligned}$$

De même, on peut munir l’espace $C^0([0, 1]; \mathbb{R})$ des fonctions continues sur $[0, 1]$ à valeurs réelles des normes suivantes :

$$\begin{aligned} \|f\|_1 &:= \int_0^1 |f(x)| dx, \\ \|f\|_\infty &:= \sup_{x \in [0, 1]} |f(x)|. \end{aligned}$$

Le fait que ces deux dernières soient bien définies vient du théorème affirmant qu’une fonction continue sur un segment est bornée.

2 Espaces euclidiens

2.1 Définitions

Nous avons vu que dans des espaces normés, on peut obtenir des informations sur la façon dont la norme se comporte par rapport aux deux opérations de l'espace : addition et multiplication par des scalaires. Mais en ce qui concerne l'addition, on dispose seulement d'une inégalité. Dans ce chapitre, nous allons voir que l'on peut donner plus d'informations sur ce sujet pour certaines normes particulières : celles qui dérivent d'un “produit scalaire”, qui est le concept principal de ce chapitre. En outre, le produit scalaire permettra d'introduire d'autres notions importantes comme l'orthogonalité.

Nous allons donc commencer par définir la notion de produit scalaire, puis voir comment à tout produit scalaire on associe une norme particulière. Il sera plus simple pour cela de distinguer les cas réel et complexe.

Définition 19. Soit E un espace vectoriel sur $\mathbb{K}$. Une application $a : E \times E \rightarrow \mathbb{K}$, $(x, y) \mapsto a(x, y) \in \mathbb{K}$ est dite **bilinéaire**, lorsqu'elle est linéaire par rapport à chacune de ses variables, c'est-à-dire :

- pour $y \in E$ fixé, l'application $x \mapsto a(x, y)$ est linéaire,
- pour $x \in E$ fixé, l'application $y \mapsto a(x, y)$ est linéaire.

Autrement dit, pour tout $x, y, z \in E$ et $\lambda, \mu \in \mathbb{K}$:

$$\begin{aligned}a(\lambda x + \mu y, z) &= \lambda a(x, z) + \mu a(y, z), \\a(x, \lambda y + \mu z) &= \lambda a(x, y) + \mu a(x, z).\end{aligned}$$

Définition 20. Soit E un espace vectoriel **réel**. Une application $\langle \cdot, \cdot \rangle : E \times E \rightarrow \mathbb{R}$, $(x, y) \mapsto \langle x, y \rangle$ est appelée un **produit scalaire** si elle possède les propriétés suivantes :

- elle est bilinéaire,
- elle est symétrique :

$$\forall x, y \in E, \langle x, y \rangle = \langle y, x \rangle,$$

- elle est positive : $\forall x \in E, \langle x, x \rangle \geq 0$,
- elle est “définie” : $\langle x, x \rangle = 0 \Rightarrow x = 0$.

Remarque 8. On dit aussi “**non dégénérée positive**” à la place de “définie positive”. Voir la remarque 31 beaucoup plus loin !

Il est aussi usuel de noter le produit scalaire, lorsqu'il n'y a pas d'ambiguïté, $\langle x, y \rangle = x \cdot y$.

Exemple. Dans $\mathbb{R}^2$, l'application

$$\langle (x_1, x_2), (y_1, y_2) \rangle = x_1 y_1 + x_2 y_2,$$

définit un produit scalaire. Un autre exemple (exercice) :

$$\langle (x_1, x_2), (y_1, y_2) \rangle = (x_1 + x_2)(y_1 + y_2) + 2x_2y_2.$$

Nous allons maintenant présenter le cas complexe, dont le cas réel est en fait un cas particulier, mais qu'il est plus simple d'introduire à part.

Définition 21. Soit E un espace vectoriel **complexe**. Une application $T : E \rightarrow \mathbb{C}$, $x \mapsto T(x)$ est dite **antilinéaire** (ou parfois semi-linéaire), lorsqu'elle satisfait : pour tous $x, y \in E$ et $\lambda, \mu \in \mathbb{C}$:

$$T(\lambda x + \mu y) = \bar{\lambda}T(x) + \bar{\mu}T(y).$$

Rappel. Rappelons la définition du conjugué d'un nombre complexe (en espérant que cela est en fait inutile) : si un complexe z s'écrit sous la forme $z = a + ib$ avec a et b réels, alors $\bar{z} = a - ib$.

Remarque 9. La différence avec la linéarité ne réside donc que dans les conjugaisons des scalaires. Mais cela implique en particulier qu'une application antilinéaire n'est pas linéaire ! (Soulignons que nous sommes dans le contexte d'espaces vectoriels **complexes**.)

Définition 22. Soit E un espace vectoriel sur $\mathbb{C}$. Une application $a : E \times E \rightarrow \mathbb{C}$, $(x, y) \mapsto a(x, y)$ est dite **sesquilinéaire**, lorsqu'elle est linéaire par rapport à la première variable, et antilinéaire par rapport à la seconde, c'est-à-dire :

- pour $y \in E$ fixé, l'application $x \mapsto a(x, y)$ est linéaire,
- pour $x \in E$ fixé, l'application $y \mapsto a(x, y)$ est antilinéaire.

Autrement dit, pour tout $x, y, z \in E$ et $\lambda, \mu \in \mathbb{K}$:

$$\begin{aligned} a(\lambda x + \mu y, z) &= \lambda a(x, z) + \mu a(y, z), \\ a(x, \lambda y + \mu z) &= \bar{\lambda}a(x, y) + \bar{\mu}a(x, z). \end{aligned}$$

Remarque 10. Dans certains ouvrages (minoritaires à ce qu'il semble), les auteurs préfèrent que l'application soit antilinéaire par rapport à la première variable, et linéaire par rapport à la seconde : ce n'est qu'une histoire de convention...

Définition 23. Soit E un espace vectoriel **complexe**. Une application $\langle \cdot, \cdot \rangle : E \times E \rightarrow \mathbb{C}$ $(x, y) \mapsto \langle x, y \rangle$ est appelée un **produit scalaire** (parfois produit scalaire **hermitien**) si elle possède les propriétés suivantes :

- elle est sesquilinéaire,
- elle est à symétrie hermitienne :

$$\forall x, y \in E, \langle x, y \rangle = \overline{\langle y, x \rangle},$$

- elle est positive : $\forall x \in E, \langle x, x \rangle \in \mathbb{R}^+$,

- elle est définie : $\langle x, x \rangle = 0 \Rightarrow x = 0$.

On dit aussi de la même façon “non dégénérée positive” à la place de “définie positive”.

Remarque 11. La seule différence avec le cas réel, consiste en ces conjugaisons auxquelles il faut faire un peu attention. Bien sûr, dans le cas d’un espace réel, les produits scalaires sont réels, et donc les propriétés ci-dessus sont encore valables. Notons enfin que le caractère réel de $\langle x, x \rangle$ provient de la symétrie hermitienne avant même de supposer l’application positive.

Exemple. Dans $\mathbb{C}^2$, l’application

$$\langle (z_1, z_2), (\zeta_1, \zeta_2) \rangle = z_1 \overline{\zeta_1} + z_2 \overline{\zeta_2},$$

définit un produit scalaire. Sur cet exemple, on comprend peut-être mieux la raison des conjugaisons complexes sur la deuxième variable : comment pourrait-on s’assurer d’obtenir une valeur positive pour $\langle (z_1, z_2), (z_1, z_2) \rangle$ sans elles ? Et comme nous le verrons bientôt, la propriété de positivité est indispensable pour définir la norme associée au produit scalaire.

Ce qui suit est valable dans les deux cas réel et complexe.

Notation. On notera

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

qui est bien défini grâce à la propriété de positivité !

Lemme 2 (Inégalité de Cauchy-Schwarz). Dans un espace réel ou complexe E muni d’un produit scalaire, on a pour tous $x, y \in E$,

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

avec égalité si et seulement si x et y sont linéairement dépendants.

Démonstration. On commence d’abord par le cas réel. On considère l’application $\lambda \mapsto \|x + \lambda y\|^2$ qui doit être positive pour tout $\lambda \in \mathbb{R}$:

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

Du fait de cette positivité sur $\mathbb{R}$, comme il s’agit d’un polynôme de second ordre, son discriminant doit être négatif, i.e.

$$4|\langle x, y \rangle|^2 \leq 4\|x\|^2 \cdot \|y\|^2,$$

d’où l’inégalité dans le cas réel.

En outre, s’il y a égalité dans l’inégalité, cela signifie que $|\langle x, y \rangle| = \|x\| \|y\|$, alors le discriminant précédent est nul, et donc le polynôme a une racine réelle μ . On voit bien que cela indique que $x + \mu y = 0$.

On considère maintenant le cas complexe et on rappelle la décomposition polaire de $\langle x, y \rangle : \langle x, y \rangle = \rho e^{i\theta}$, $\rho = |\langle x, y \rangle| > 0$, $\theta \in \mathbb{R}$. On calcule $\|x' + \lambda y\|^2$ pour $\lambda \in \mathbb{R}$ et $x' = e^{-i\theta}x$:

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

On évalue un par un les trois premiers termes :

$$\begin{aligned} \|x'\|^2 &= \|e^{-i\theta}x\|^2 = |e^{-i\theta}|^2 \|x\|^2 = 1^2 \|x\|^2 = \|x\|^2, \\ \langle \lambda y, x' \rangle &= \lambda \langle y, e^{-i\theta}x \rangle = \lambda \overline{\langle e^{-i\theta}x, y \rangle} = \lambda \overline{e^{-i\theta} \rho e^{i\theta}} = \lambda \rho, \\ \langle x', \lambda y \rangle &= \langle e^{-i\theta}x, \lambda y \rangle = \lambda e^{-i\theta} \langle x, y \rangle = \lambda e^{-i\theta} \rho e^{i\theta} = \lambda \rho. \end{aligned}$$

Donc $\|x' + \lambda y\|^2 = \|x'\|^2 + 2\lambda\rho + \lambda^2\|y\|^2$ et cette quantité est positive pour tout $\lambda \in \mathbb{R}$. On conclut donc comme dans le cas réel. $\square$

Comme annoncé, à chaque produit scalaire on peut associer une norme particulière.

Théorème 7. *L'application $x \mapsto \|x\|$ est une norme, appelée norme **induite par le produit scalaire** $\langle \cdot, \cdot \rangle$ et aussi appelée norme **euclidienne** dans le cas réel et norme **hermitienne** dans le cas complexe.*

Démonstration. Il s'agit de montrer que l'application ainsi définie a toutes les propriétés d'une norme : la positivité et les propriétés (2)-(4). Le seul point non évident est l'inégalité triangulaire (4). Pour cela on écrit :

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

puis par utilisation de Cauchy-Schwarz on obtient

$$\begin{aligned} \|x + y\|^2 &\leq \|x\|^2 + 2\|x\|\|y\| + \|y\|^2 \\ &= (\|x\| + \|y\|)^2. \end{aligned}$$

$\square$

Définition 24. *Deux vecteurs $x, y \in E$ tels que $\langle x, y \rangle = 0$ sont dits **orthogonaux**, ce que l'on note $x \perp y$. Un vecteur x est dit orthogonal à un s.e.v. F et l'on note $x \perp F$ lorsque pour tout $y \in F$, on a $x \perp y$. Deux s.e.v. F et G de E sont dits orthogonaux, et l'on note $F \perp G$ lorsque tout vecteur de F est orthogonal à tout vecteur de G .*

Définition 25. *On appelle espace euclidien tout espace vectoriel **réel**, de dimension finie, muni d'un produit scalaire.*

Note. Tout espace euclidien est en particulier un espace normé.

2.2 Exemples

2.2.1 Produit scalaire canonique de $\mathbb{R}^n$

Dans l'espace $\mathbb{R}^n$, on introduit le produit scalaire **canonique**, c'est-à-dire celui qu'on utilisera de manière standard dans $\mathbb{R}^n$ (sauf précision contraire bien entendu). Il s'agit du produit scalaire suivant :

$$\langle (x_1, \dots, x_n), (y_1, \dots, y_n) \rangle_{\mathbb{R}^n} = x_1 y_1 + \dots + x_n y_n.$$

Il arrive parfois qu'on note ce produit scalaire simplement $\langle \cdot, \cdot \rangle$ lorsque le contexte n'est pas ambigu.

Remarque 12. Avec ce choix du produit scalaire la longueur d'un vecteur et l'orthogonalité prennent les sens classiques :

- la norme de (x_1, x_2, x_3) (ou distance de 0 à (x_1, x_2, x_3)) dans $\mathbb{R}^3$ est $\sqrt{x_1^2 + x_2^2 + x_3^2}$,
- $x \perp y$ est équivalent $\langle x, y \rangle = 0$ ce qui est équivalent à $\|x + y\|^2 = \|x\|^2 + \|y\|^2$.

Pour ajouter une remarque, rappelons une notation.

Notation. Dans toute la suite, si $M \in M_{n,p}(\mathbb{K})$ est une matrice, nous désignerons par ${}^t M \in M_{p,n}(\mathbb{K})$ sa **transposée** (obtenue par écriture de gauche à droite à la j -ième ligne des entrées de la j -ième colonne de M lues de haut en bas). Rappelons la formule ${}^t(AB) = {}^t B {}^t A$, valable dès que le produit matriciel AB a un sens.

Avec cette notation, notons que si U et V sont des vecteurs colonnes de $\mathbb{R}^n$, on a :

$$\langle U, V \rangle_{\mathbb{R}^n} = {}^t U V, \tag{5}$$

en identifiant (ce que nous ferons systématiquement) les matrices scalaires (à une ligne et une colonne) et le scalaire qu'elles contiennent.

2.2.2 Produit scalaire canonique de $\mathbb{C}^n$

Dans l'espace $\mathbb{C}^n$, on utilise de la même façon le produit scalaire **canonique** complexe :

$$\langle (z_1, \dots, z_n), (\zeta_1, \dots, \zeta_n) \rangle_{\mathbb{C}^n} = z_1 \bar{\zeta}_1 + \dots + z_n \bar{\zeta}_n.$$

Soulignons que dans le cas particulier où les n -uplets $(z_1, \dots, z_n)$ et $(\zeta_1, \dots, \zeta_n)$ sont des vecteurs réels, alors $\langle (z_1, \dots, z_n), (\zeta_1, \dots, \zeta_n) \rangle_{\mathbb{C}^n} = \langle (z_1, \dots, z_n), (\zeta_1, \dots, \zeta_n) \rangle_{\mathbb{R}^n}$.

Ici, si U et V sont des vecteurs colonnes de $\mathbb{C}^n$, on a :

$$\langle U, V \rangle_{\mathbb{C}^n} = {}^t U \bar{V}. \tag{6}$$

2.2.3 Autres exemples d'espaces munis d'un produit scalaire : les fonctions continues

Soit $E = C([a, b]; \mathbb{R})$ le $\mathbb{R}$ -e.v. des applications continues de $[a, b]$ dans $\mathbb{R}$. On définit le produit scalaire sur E par

$$\langle f, g \rangle = \int_a^b f(x)g(x)dx.$$

De même on peut aussi définir le produit scalaire sur l'e.v. complexe $F = C([a, b], \mathbb{C})$ des applications continues de $[a, b]$ dans $\mathbb{C}$ par :

$$\langle f, g \rangle = \int_a^b f(x)\overline{g(x)}dx.$$

Remarque 13. *Approximation numérique. Approchons les fonctions f de $C([a, b], \mathbb{R})$ par des fonctions constantes par morceaux $T_n[f]$ de la manière suivante.*

La fonction $T_n[f]$ est constante sur chaque intervalle $[a + \frac{k}{n}(b-a), a + \frac{k+1}{n}(b-a)[$: on lui donne comme valeur sur cet intervalle celle de f au point le plus à gauche de celui-ci, à savoir $f(a + \frac{k}{n}(b-a))$ (voir la Figure 1).

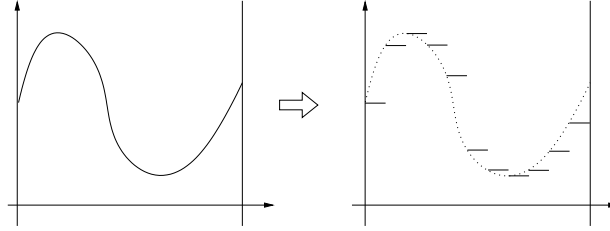


FIGURE 1 – Approximation constante par morceaux

Le produit scalaire obtenu sur ces approximations constantes par morceaux :

$$\int_a^b T_n[f](x)T_n[g](x)dx$$

est le même (à un facteur multiplicatif près) que le produit euclidien sur les n -uplets $(f(a), f(a + \frac{b-a}{n}), \dots, f(a + (n-1)\frac{b-a}{n}))$ et $(g(a), g(a + \frac{b-a}{n}), \dots, g(a + (n-1)\frac{b-a}{n}))$.

2.2.4 Un autre exemple d'espace euclidien

On peut munir l'espace des matrices carrées à entrées réelles d'une structure euclidienne, comme suit.

Soit $E = M_n(\mathbb{R})$ l'espace des matrices $n \times n$. On peut définir un produit scalaire sur E par la formule :

$$\langle A, B \rangle = \text{Tr}({}^tAB).$$

2.3 Premières propriétés des espaces euclidiens

Dans cette section, nous donnons quelques propriétés importantes des espaces euclidiens.

2.3.1 Identité du parallélogramme

L'identité suivante est en fait valable dans tout espace muni d'un produit scalaire (même s'il est complexe ou de dimension infinie).

Proposition 10 (Identité du parallélogramme). *Dans tout espace euclidien E on a, $\forall x, y \in E$:*

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

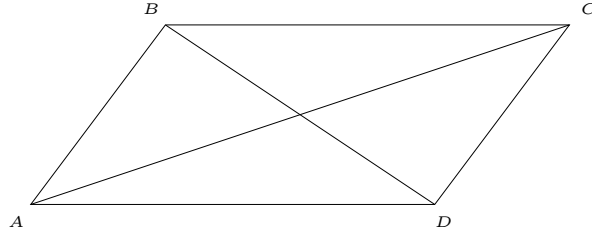


FIGURE 2 – Identité du parallélogramme

Cette égalité est dite *identité du parallélogramme*, car elle permet de voir que dans un parallélogramme, la somme des carrés des longueurs des diagonales est égale à la somme des carrés des longueurs des côtés. Ainsi, dans la Figure 2, on a la relation

$$\|\overrightarrow{AC}\|^2 + \|\overrightarrow{BD}\|^2 = \|\overrightarrow{AB}\|^2 + \|\overrightarrow{BC}\|^2 + \|\overrightarrow{CD}\|^2 + \|\overrightarrow{DA}\|^2 = 2(\|\overrightarrow{AB}\|^2 + \|\overrightarrow{BC}\|^2).$$

Démonstration. Il suffit de développer $\|x + y\|^2 + \|x - y\|^2 = \langle x + y, x + y \rangle + \langle x - y, x - y \rangle \dots$

□

Remarque 14. *Certaines normes ne provenant pas d'un produit scalaire ne satisfont pas cette identité. C'est même un moyen simple de déterminer qu'une norme ne provient pas d'un produit scalaire. Exemple : montrer que la norme sur $\mathbb{R}^2$: $\|(x, y)\|_1 := |x| + |y|$ ne provient pas d'un produit scalaire.*

Une autre identité utile des espaces euclidiens : on sait déjà déterminer la norme à partir du produit scalaire, mais **pour une norme dont on sait qu'elle provient d'un produit scalaire**, on sait faire le trajet inverse :

Proposition 11. Dans un espace euclidien, on a, $\forall x, y \in E$:

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

Démonstration. Il suffit de développer $\|x + y\|^2$ à l'aide du produit scalaire. $\square$

Remarque 15. La Proposition 10 est valable dans le cas d'un $\mathbb{C}$ -espace vectoriel muni d'un produit scalaire complexe. Pas la Proposition 11 (pourquoi ?).

2.3.2 Angle non-orienté des deux vecteurs

Soient u et v deux vecteurs non nuls d'un espace euclidien E . On sait d'après l'inégalité de Cauchy-Schwarz que $\left| \frac{\langle u, v \rangle}{\|u\| \cdot \|v\|} \right| \leq 1$. On peut donc trouver un angle $\alpha \in [0, \pi]$ tel que

$$\cos(\alpha) = \frac{\langle u, v \rangle}{\|u\| \cdot \|v\|}.$$

Définition 26. Cet α est appelé **angle non orienté** entre u et v .

Note. On retrouve bien la notion usuelle d'angle non orienté du plan euclidien.

2.3.3 Théorème de Pythagore

Théorème 8. Soit E un espace euclidien. Alors :

- on a l'équivalence

$$\|x + y\|^2 = \|x\|^2 + \|y\|^2 \Leftrightarrow \langle x, y \rangle = 0,$$

- si $x_1, \dots, x_n$ sont des vecteurs orthogonaux deux à deux, on a :

$$\|x_1 + \dots + x_n\|^2 = \sum_{i=1}^n \|x_i\|^2. \quad (7)$$

Démonstration. Développer les carrés à l'aide du produit scalaire ! $\square$

Note. À partir de trois vecteurs, la réciproque n'est pas toujours vraie (l'identité (7) n'implique pas que les vecteurs soient deux à deux orthogonaux). Par exemple, pour $(1, 0)$, $(0, 1)$, $(1, -1)$:

$$\begin{aligned} \|(1, 0) + (0, 1) + (1, -1)\|^2 &= \|(2, 0)\|^2 = 4, \\ \|(1, 0)\|^2 + \|(0, 1)\|^2 + \|(1, -1)\|^2 &= 1 + 1 + 2 = 4, \end{aligned}$$

donc les deux quantités sont égales, mais les vecteurs ne sont pas deux à deux orthogonaux :

$$\langle (0, 1), (1, -1) \rangle = 0 \cdot 1 + 1 \cdot (-1) = -1 \neq 0.$$

2.3.4 Orthogonalité et liberté

Proposition 12. *Soit E un espace euclidien. Soit $X = \{x_1, \dots, x_n\}$ une famille de vecteurs, deux à deux orthogonaux et non nuls. Alors X est une famille libre de E .*

Démonstration. Si une combinaison linéaire des x_i est nulle, on en prend le produit scalaire par x_k (quel que soit k). En utilisant les relations d'orthogonalité (et le fait que x_k soit non nul), on voit que le coefficient de x_k de cette combinaison est nul. $\square$

2.4 Bases orthonormées, procédé d'orthonormalisation de Schmidt

2.4.1 Motivation

Soient E un espace euclidien et $X = \{e_1, \dots, e_n\}$ une base de E . Pour $x = \sum_{k=1}^n \alpha_k e_k$ et $y = \sum_{j=1}^n \beta_j e_j$, le calcul de $\langle x, y \rangle$ donne :

$$\langle x, y \rangle = \left\langle \sum_{k=1}^n \alpha_k e_k, \sum_{j=1}^n \beta_j e_j \right\rangle = \sum_{j,k=1}^n \alpha_k \beta_j \langle e_k, e_j \rangle. \quad (8)$$

Ce calcul n'est en général pas aisé car il faut prendre en compte tous les produits $\langle e_k, e_j \rangle$. Mais toutes les bases ne sont donc pas identiques de ce point de vue. En particulier, celles qui ont beaucoup de facteurs $\langle e_k, e_j \rangle$ nuls vont permettre un calcul plus rapide. Il faut néanmoins remarquer que quelle que soit la base, on aura $\langle e_k, e_k \rangle = \|e_k\|^2 \neq 0$.

2.4.2 Bases orthonormées

Définition 27. *Soient E un espace euclidien et $X = \{e_1, \dots, e_n\}$ une base de E . La base X est dite **orthogonale** lorsque*

$$\langle e_i, e_j \rangle = 0, \text{ pour tout } i \neq j.$$

*La base X est dite **orthonormée** lorsque*

$$\langle e_i, e_j \rangle = \delta_{ij} \text{ pour tout } i, j \leq n.$$

Remarque 16. *Dans une base orthonormée, le calcul (8) devient très simple :*

$$\langle x, y \rangle = \left\langle \sum_{k=1}^n \alpha_k e_k, \sum_{j=1}^n \beta_j e_j \right\rangle = \sum_{j,k=1}^n \alpha_k \beta_j \langle e_k, e_j \rangle = \sum_{k=1}^n \alpha_k \beta_k. \quad (9)$$

Ainsi, si on raisonne sur les coordonnées dans cette base, on est ramené au cas de $\mathbb{R}^n$ avec le produit euclidien usuel. Un énoncé précis donnant cette équivalence d'un espace euclidien à $\mathbb{R}^n$ muni du produit scalaire canonique est donné plus bas (voir Théorème 10).

Proposition 13. Soient E un espace euclidien et $\{e_1, \dots, e_n\}$ une base orthonormée de E . Alors

$$\forall x \in E, \quad x = \sum_{i=1}^n \langle x, e_i \rangle e_i \quad \text{et} \quad \|x\|^2 = \sum_{i=1}^n |\langle x, e_i \rangle|^2.$$

Démonstration. Il suffit d'écrire x comme combinaison linéaire des e_i , et de considérer les produits scalaires avec les e_i pour identifier les coefficients. Ensuite, c'est le théorème de Pythagore. $\square$

Remarque 17. Donnons un exemple de construction d'une base orthonormée dans $\mathbb{R}^2$, comme une "déformation" d'une base donnée.

À partir de $v_1 = (2, 0)$ et $v_2 = (2, 1)$ on peut chercher une base $\{e_1, e_2\}$ orthonormée en "normalisant" le premier : $e_1 = \frac{v_1}{\|v_1\|} = (1, 0)$ et en cherchant le second comme une combinaison linéaire des deux : $e_2 = \alpha e_1 + \beta v_2$. On obtient $e_2 = (0, 1)$.

2.4.3 Procédé d'orthonormalisation de Schmidt

Un procédé pour construire une base orthonormée à partir d'une base quelconque et qui généralise la méthode de la Remarque 17 est le procédé d'orthonormalisation de Schmidt (aussi appelé procédé de Gram-Schmidt) décrit dans le théorème et la preuve suivants.

Théorème 9. Soient E un espace euclidien de dimension n et $X = \{x_1, \dots, x_n\}$ une base de E . Alors il existe une unique base orthonormée $X_S = \{e_1, \dots, e_n\}$ telle que pour $k = 1, \dots, n$,

$$\begin{aligned} \text{Vect}\{e_1, \dots, e_k\} &= \text{Vect}\{x_1, \dots, x_k\}, \\ \langle x_k, e_k \rangle &> 0. \end{aligned}$$

Corollaire 2. Tout espace euclidien admet une base orthonormée.

Démonstration. On définit par récurrence les vecteurs suivants :

$$\begin{aligned} e_1 &= \frac{x_1}{\|x_1\|}, \\ e_2 &= \frac{x_2 - \langle x_2, e_1 \rangle e_1}{\|x_2 - \langle x_2, e_1 \rangle e_1\|}, \\ &\vdots \\ e_k &= \frac{x_k - \sum_{j=1}^{k-1} \langle x_k, e_j \rangle e_j}{\|x_k - \sum_{j=1}^{k-1} \langle x_k, e_j \rangle e_j\|}, \\ &\vdots \\ e_n &= \frac{x_n - \sum_{j=1}^{n-1} \langle x_n, e_j \rangle e_j}{\|x_n - \sum_{j=1}^{n-1} \langle x_n, e_j \rangle e_j\|}. \end{aligned}$$

On montre maintenant différentes propriétés.

- Chaque e_k appartient à $\text{Vect}\{x_1, \dots, x_k\}$: cela est clair dans les formules ci-dessus, par récurrence sur k .
- Pour tout k , $\|x_k - \sum_{j=1}^{k-1} \langle x_k, e_j \rangle e_j\| \neq 0$: cela vient du fait que X est une base et du fait que nous venons de prouver que pour $i = 1, \dots, k-1$, $e_i \in \text{Vect}\{x_1, \dots, x_{k-1}\}$. Cela permet en particulier de voir que les définitions précédentes ne comprennent pas de division par 0.
- Chaque x_k est inclus dans $\text{Vect}\{e_1, \dots, e_k\}$: cela vient encore des formules ci-dessus.
- Les e_j , $j \leq n$ sont orthogonaux deux à deux : il suffit de montrer que e_k est orthogonal à tous les e_j pour $j = 1, \dots, k-1$; pour s'en convaincre on multiplie la formule précédente de e_k par e_j avec $j < k$.
- Enfin, pour tout k , on a $\langle x_k, e_k \rangle > 0$ car

$$0 < \|e_k\|^2 = \langle e_k, e_k \rangle = \left\langle \frac{x_k - \sum_{j=1}^{k-1} \langle x_k, e_j \rangle e_j}{\|x_k - \sum_{j=1}^{k-1} \langle x_k, e_j \rangle e_j\|}, e_k \right\rangle = \frac{\langle x_k, e_k \rangle}{\|x_k - \sum_{j=1}^{k-1} \langle x_k, e_j \rangle e_j\|}.$$

Pour ce qui est de l'unicité : supposons avoir une base orthonormée $(e_1, \dots, e_n)$ satisfaisant les conditions précédentes. On commence par voir que e_1 est nécessairement défini comme précédemment : il appartient à $\text{Vect}\{x_1\}$ et est donc de la forme λx_1 ; ensuite $\|e_1\| = 1$ et $\langle e_1, x_1 \rangle > 0$ montrent que nécessairement $\lambda = 1/\|x_1\|$. Supposons ensuite que nous avons prouvé que $e_1, \dots, e_k$ avaient nécessairement la forme précédente ; prouvons que e_{k+1} a lui aussi cette forme. D'après les hypothèses, on sait que $e_{k+1} \in \text{Vect}\{x_1, \dots, x_{k+1}\} = \text{Vect}\{e_1, \dots, e_k, x_{k+1}\}$. On écrit donc e_{k+1} sous la forme :

$$e_{k+1} = \sum_{i=1}^k \alpha_i e_i + \lambda x_{k+1}.$$

On prend le produit scalaire avec e_i pour $i = 1, \dots, k$ et on trouve $\alpha_i = -\lambda \langle x_{k+1}, e_i \rangle$, donc

$$e_{k+1} = \lambda \left(x_{k+1} - \sum_{i=1}^k \langle x_{k+1}, e_i \rangle e_i \right).$$

Enfin les conditions $\|e_{k+1}\| = 1$ et $\langle e_{k+1}, x_{k+1} \rangle > 0$ déterminent λ et montrent que e_{k+1} a bien la forme précédente. $\square$

2.4.4 Structure des espaces euclidiens par les bases orthonormées. Equivalence avec $\mathbb{R}^n$.

Définition 28. Soient E et F deux espaces euclidiens. On appelle **isométrie vectorielle** entre E et F (ou encore, **isomorphisme d'espaces euclidiens**), une bijection linéaire $f : E \rightarrow F$ qui préserve le produit scalaire, c'est-à-dire telle que :

$$\forall u, v \in E, \langle f(u), f(v) \rangle_F = \langle u, v \rangle_E.$$

Théorème 10. Soit E un espace euclidien de dimension n . Alors E est en isométrie avec $\mathbb{R}^n$, c'est-à-dire qu'il existe une isométrie vectorielle entre E et $\mathbb{R}^n$.

Démonstration. On dénote par $\{e_1, \dots, e_n\}$ une base orthonormée de E . On sait que tout $u \in E$ s'écrit d'une manière unique

$$u = \sum_{i=1}^n U^i e_i, \quad (10)$$

puisque E est une base. On définit l'application

$$f : \begin{cases} \mathbb{R}^n & \longrightarrow E \\ U = (U^i)_{i=1}^n & \longmapsto f(U) = \sum_{i=1}^n U^i e_i. \end{cases}$$

Dire que tout u s'écrit de manière unique comme dans (10), c'est exactement dire que f est bijective.

Montrons que f est linéaire et préserve le produit scalaire. Soient $U = (U^i)_{i=1}^n \in \mathbb{R}^n$ et $V = (V^i)_{i=1}^n \in \mathbb{R}^n$. La linéarité se montre de manière standard :

$$\begin{aligned} f(U + V) &= f((U^i)_{i=1}^n + (V^i)_{i=1}^n) = f((U^i + V^i)_{i=1}^n) \\ &= \sum_{i=1}^n (U^i + V^i) e_i = \sum_{i=1}^n U^i e_i + \sum_{i=1}^n V^i e_i = f(U) + f(V), \\ f(\lambda U) &= f((\lambda U^i)_{i=1}^n) = \sum_{i=1}^n (\lambda U^i) e_i = \lambda \sum_{i=1}^n U^i e_i = \lambda f(U). \end{aligned}$$

Le caractère isométrique est donné par le calcul suivant :

$$\langle U, V \rangle_{\mathbb{R}^n} = \sum_{i=1}^n U^i V^i = \left\langle \sum_{i=1}^n U^i e_i, \sum_{i=1}^n V^i e_i \right\rangle = \langle f(U), f(V) \rangle_E.$$

Par conséquent l'application f qui à $(U^i)_{i=1}^n \in \mathbb{R}^n$ associe $u \in E$ est un isomorphisme d'espaces euclidiens. $\square$

Remarque 18. L'isomorphisme est **différent** pour chaque base orthonormée $\{e_1, \dots, e_n\}$ choisie.

Remarque 19. Cette équivalence s'étend pour les espaces complexes de dimension finie dotés de produits scalaires complexes : il sont de la même façon en isométrie vectorielle (complexe) avec $\mathbb{C}^n$.

2.5 Projections

2.5.1 Définition

Définition 29. Soient E un espace euclidien et $A \subset E$ une partie non vide de E . On note

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

et on l'appelle l'**orthogonal** de A .

Proposition 14. L'orthogonal $A^\perp$ de tout partie non vide A est un s.e.v de E . En outre on a :

$$\{0\}^\perp = E, \quad E^\perp = \{0\}.$$

Démonstration. Les propriétés de s.e.v. se vérifient trivialement. De même, tout vecteur est orthogonal à zéro. Pour voir que $E^\perp = \{0\}$, il suffit de remarquer que si $v \in E^\perp$, alors $\langle v, v \rangle = 0$. $\square$

Proposition 15. Soit F un s.e.v. d'un espace euclidien E . Alors :

- $E = F \oplus F^\perp$,
- $F^{\perp\perp} = F$.

Démonstration. Pour la première partie : on considère une base de F ; on étend ensuite celle-ci à une base de E (on utilise ici le Théorème 3 de la base incomplète). Ensuite on orthonormalise (par le procédé de Schmidt) et on obtient donc une base orthonormée de E qui contient une base orthonormée de F , mettons $\{e_1, \dots, e_n\}$ où les p premiers vecteurs forment une base de F . Il n'est alors pas difficile de voir que

$$F^\perp = \text{Vect} \{e_{p+1}, \dots, e_n\},$$

en utilisant la formule (9). On conclut ainsi que $E = F \oplus F^\perp$ et que $\dim(E) = \dim(F) + \dim(F^\perp)$.

Pour la deuxième partie on a tout d'abord $\dim(F^{\perp\perp}) = \dim(E) - \dim(F^\perp) = \dim(F)$ et par ailleurs on vérifie facilement que $F \subset F^{\perp\perp}$. On a donc l'égalité. $\square$

Définition 30. Le théorème précédent nous montre que tout $x \in E$ s'écrit d'une manière unique comme $x = u + v$ avec $u \in F$ et $v \in F^\perp$. On appelle u la **projection orthogonale** de x sur F et on note par p_F l'application $x \mapsto p_F(x) = u$ ainsi construite.

Exemple. Dans $E = \mathbb{R}^3$ la projection de $(3, 2, 4)$ sur le plan $x = 0$ est $(0, 2, 4)$ alors que sa projection sur $x + y + z = 0$ est $(0, -1, 1)$.

Proposition 16. Soit F un s.e.v. d'un espace euclidien E . Alors

1. **Orthogonalité :** pour tout $x \in E$ on a : $\forall y \in F, x - p_F(x) \perp y$ et $x = p_F(x) + p_{F^\perp}(x)$.
De plus $\|p_F(x)\| \leq \|x\|$.

2. si $\{e_1, \dots, e_k\}$ est une base orthonormée de F on a

$$\forall x \in E, p_F(x) = \sum_{i=1}^k \langle x, e_i \rangle e_i. \quad (11)$$

3. Si p est une projection orthogonale $p \circ p = p$.

4. Si p est un endomorphisme vérifiant $p \circ p = p$ et si en plus $\text{Im}(p) \perp \text{Ker}(p)$ alors p est la projection orthogonale sur $\text{Im}(p)$.

5. Caractérisation par la propriété de minimalité :

$$\forall x \in E, \|x - p_F(x)\| = \inf_{w \in F} \|x - w\|.$$

Démonstration. Le 1. est une conséquence simple de l'écriture $x = u + v$ et de $F = F^{\perp\perp}$. L'inégalité $\|p_F(x)\| \leq \|x\|$ découle de Pythagore.

Pour le 2., on écrit la projection comme une combinaison linéaire dans cette base, ensuite les coefficients se déterminent en utilisant $x - p_F(x) \perp F$, donc pour $i = 1 \dots k$, $\langle x, e_i \rangle = \langle p_F(x), e_i \rangle$.

Pour le 3., on utilise que pour $x \in F$, $p_F(x) = x$ (ce qui est clair puisque l'écriture $x = x + 0$ avec $x \in F$ et $0 \in F^\perp$ est unique).

Pour le 4., on démontre d'abord que $\text{Im}(p) \oplus \text{Ker}(p) = E$: tout $x \in E$ s'écrit $x = p(x) + x - p(x)$ et on voit que $p(x - p(x)) = 0$ donc $x - p(x) \in \text{Ker}(p)$. De plus l'écriture est unique puisque si $x = a + b$ avec $b \in \text{Ker}(p)$ alors $p(x) = p(a) = a$ (une autre manière de faire était de voir que les espaces $\text{Im}(p)$ et $\text{Ker}(p)$ sont en somme directe (du fait de leur orthogonalité) et d'utiliser le théorème du rang). On conclut par la définition de la projection orthogonale.

Pour le 5., on écrit $\|x - w\|^2 = \|x - p_F(x) + p_F(x) - w\|^2 = \|x - p_F(x)\|^2 + \|p_F(x) - w\|^2$, la dernière égalité venant de la propriété d'orthogonalité. $\square$

Remarque 20. On reconnaît dans la formule (11) une écriture que nous avons vue dans le procédé d'orthonormalisation de Schmidt. On voit que la base obtenue par orthonormalisation peut également s'écrire :

$$e_k = \frac{x_k - P_{F_{k-1}}(x_k)}{\|x_k - P_{F_{k-1}}(x_k)\|}, \text{ où } F_{k-1} := \text{Vect}\{e_1, \dots, e_{k-1}\} = \text{Vect}\{x_1, \dots, x_{k-1}\}.$$

Proposition 17. Soit E un espace euclidien muni d'une base orthonormée $(v_1, \dots, v_n)$. Soient F un s.e.v. de E et $(e_1, \dots, e_k)$ une base orthonormée de F . Notons U la matrice des coordonnées de $e_1, \dots, e_k$ en colonne dans la base $(v_1, \dots, v_n)$, (qui est donc une matrice à n lignes et k colonnes). Soit $x \in E$, et soient X et Y les matrices coordonnées respectives de x et $p_F(x)$ dans la base $(v_1, \dots, v_n)$. Alors on a :

$$Y = U^t U X.$$

Démonstration. Ce n'est en fait rien d'autre que la formule (11) : tUX est le vecteur colonne (à k lignes) :

$$\begin{pmatrix} \langle x, e_1 \rangle \\ \vdots \\ \langle x, e_k \rangle \end{pmatrix},$$

et on conclut aisément. □

2.5.2 Application au problème des moindres carrés

Soient m points de $\mathbb{R}$ notés x_i , $i = 1, \dots, m$ et c_i des valeurs associées à ces points. On s'intéresse au problème d'approcher cet ensemble de points par le graphe d'une fonction élémentaire, c'est-à-dire de trouver une fonction simple telle que $U(x_i) \simeq c_i$. Un bon choix peut être constitué par les fonctions affines de la forme $U(x) = ax + b$. Autrement dit, nous allons chercher à approcher un “nuage de points” par une droite. C'est le classique problème de la “régression linéaire”, tel que décrit sur la Figure 3.

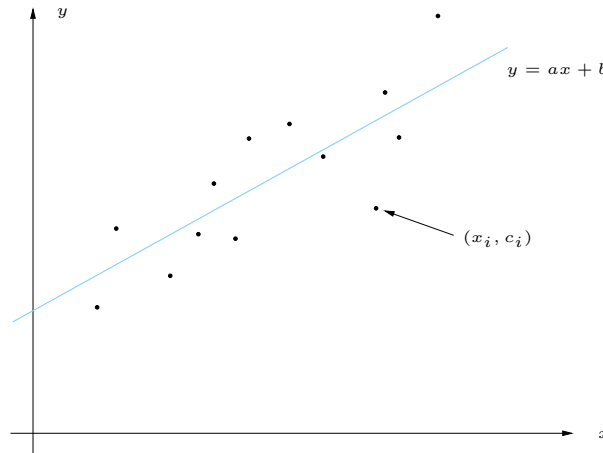


FIGURE 3 – Droite de régression linéaire

On peut formuler notre problème comme la **minimisation** de

$$S(a, b) = \sum_{i=1}^m (ax_i + b - c_i)^2.$$

Soit c le vecteur colonne d'entrées c_i , autrement dit, $c = {}^t(c_1, \dots, c_m)$. On considère aussi $x = {}^t(x_1, \dots, x_m)$, $u = {}^t(1, \dots, 1) \in \mathbb{R}^m$. On observe alors que

$$S(a, b) = \|ax + bu - c\|^2.$$

On peut donc reformuler le problème comme suit : trouver la combinaison linéaire des vecteurs x et u (de coefficients a et b) à distance minimale du vecteur c . Soit donc l'espace

vecteuriel V engendré par x et u dans $\mathbb{R}^m$. D'après ce qui précède, nous cherchons donc la **projection orthogonale** de c sur V :

$$ax + bu = p_V(c).$$

Pour effectuer le calcul on utilise la propriété d'orthogonalité :

$$(ax + bu - c) \perp x \quad \text{et} \quad (ax + bu - c) \perp u.$$

Cela donne un système de deux équations à deux inconnues à résoudre.

Ceci se généralise de plusieurs façons :

- On peut chercher la meilleure approximation parmi des fonctions non affines. Par exemple, soient $x \mapsto g_1(x)$ et $x \mapsto g_2(x)$. On peut chercher la meilleure approximation comme combinaison de ces deux fonctions

$$ag_1(x_i) + bg_2(x_i) \simeq c_i,$$

au sens des moindres carrés, c'est-à-dire qui minimise

$$\sum_{i=1}^m |ag_1(x_i) + bg_2(x_i) - c_i|^2.$$

Le cas précédent correspond à $g_1 : x \mapsto x$ et à g_2 la fonction constante à 1. On introduit alors les vecteurs

$$G_1 = {}^t(g_1(x_1), \dots, g_1(x_m)) \quad \text{et} \quad G_2 = {}^t(g_2(x_1), \dots, g_2(x_m)).$$

On cherche donc le vecteur au sein de $\text{Vect} \{G_1, G_2\}$ qui soit à distance minimale du vecteur c . Il s'agit donc encore du projeté orthogonal de c sur le sous-espace $\text{Vect} \{G_1, G_2\}$.

- On peut utiliser plus de deux fonctions, par exemple $g_1, \dots, g_k$, et chercher la meilleure approximation sous la forme

$$\alpha_1 g_1(x_i) + \dots + \alpha_k g_k(x_i) \simeq c_i.$$

On introduira alors comme précédemment les vecteurs colonnes $G_1, \dots, G_k$ donnés par

$$G_j = \begin{pmatrix} g_j(x_1) \\ \vdots \\ g_j(x_m) \end{pmatrix} \quad \text{pour } j = 1, \dots, k,$$

et la meilleure approximation au sens des moindres carrés sera le projeté orthogonal de c sur le sous-espace $\text{Vect} \{G_1, \dots, G_k\}$.

- On peut considérer le cas d'une dépendance de plusieurs variables, par exemple quand on a un nuage tridimensionnel $(x_i, y_i, c_i)_{i=1\dots m}$ et qu'on cherche à approcher c par une fonction des deux variables x, y . Dans le cas linéaire pour simplifier, on cherche donc la meilleure approximation

$$\alpha x_i + \beta y_i + b \simeq c_i,$$

ce qui donne au sens des moindres carrés

$$\sum_{i=1}^m |\alpha x_i + \beta y_i + b - c_i|^2.$$

On cherche encore le projeté orthogonal de c sur le sous-espace de $\mathbb{R}^m$:

$$\text{Vect} \left\{ \begin{pmatrix} x_1 \\ \vdots \\ x_m \end{pmatrix}, \begin{pmatrix} y_1 \\ \vdots \\ y_m \end{pmatrix}, \begin{pmatrix} 1 \\ \vdots \\ 1 \end{pmatrix} \right\}.$$

On obtient dans ce cas précis trois équations à trois inconnues.

2.6 Structure du dual d'un espace euclidien

Soit E un espace euclidien et soit E^* le dual de E (i.e., rappelons-nous, l'espace des formes linéaires sur E).

Pour $E = \mathbb{R}^3$, nous avons vu comme exemples d'éléments de E^* les applications $\mu(x, y, z) = x + 2y - z$, et $\nu(x, y, z) = x - z$. Plus généralement, toute application $\lambda_{\alpha, \beta, \gamma}(x, y, z) = \alpha x + \beta y + \gamma z$ ($\alpha, \beta, \gamma \in \mathbb{R}$) est linéaire donc $\lambda_{\alpha, \beta, \gamma} \in E^*$. Remarquons que l'on peut écrire $\lambda_{\alpha, \beta, \gamma}(x, y, z) = \langle (\alpha, \beta, \gamma), (x, y, z) \rangle$.

Plus généralement, pour tout $e \in E$, l'application $d_e : E \rightarrow \mathbb{R}$ définie par $d_e(v) = \langle v, e \rangle$ est une forme linéaire sur E donc $d_e \in E^*$. Nous sommes ainsi conduits à nous demander si la réciproque est également vraie, c'est-à-dire, si toutes les applications $f \in E^*$ sont de la forme précédente. Autrement dit, pour tout f de E^* , existe-t-il w dans E tel que pour tout v de E on ait $f(v) = \langle v, w \rangle$? La réponse est donnée par le résultat suivant.

Théorème 11. *L'application $d : E \rightarrow E^*$ qui à tout v de E associe l'élément $d(v)$ de E^* (qu'on notera d_v) défini par*

$$\forall u \in E, d_v(u) = \langle u, v \rangle,$$

est un isomorphisme (i.e. une bijection linéaire) de E sur E^ . De plus si F est un s.e.v. de E et F^0 est l'annulateur de F (voir la Définition 17 page 9) on a*

$$d(F^\perp) = F^0.$$

Démonstration. On sait d'après le Théorème 5 (page 8) que E et E^* ont la même dimension. Comme la linéarité de d est évidente, il suffit juste de montrer l'injectivité. Mais si jamais $d(v) = 0$ il suit qu'en particulier $d_v(v) = \langle v, v \rangle = 0$ d'où l'injectivité.

Une autre façon de faire est de considérer une base $\{e_1, \dots, e_n\}$ de E et pour $f \in E^*$, $v = \sum_{i=1}^n v_i e_i$ constater que $f(v) = \sum_{i=1}^n v_i f(e_i)$ donc $f(v) = \langle v, w \rangle$ avec $w = \sum_{i=1}^n f(e_i) e_i \in E$ (un peu comme dans la preuve du Théorème 5).

Pour le dernier point, on écrit

$$\begin{aligned} d^{-1}(F^0) &= \{y \in E \mid d_y \in F^0\} \\ &= \{y \in E \mid d_y(v) = 0, \forall v \in F\}, \\ &= \{y \in E \mid \langle v, y \rangle = 0, \forall v \in F\} \\ &= F^\perp. \end{aligned}$$

Comme d est une bijection il suit que $d(F^\perp) = F^0$. □

2.7 Changement de base orthonormée et groupe orthogonal

2.7.1 Rappels sur matrices et les changements de coordonnées

Rappelons comment former la matrice de passage de la base $\{e_1, \dots, e_n\}$ vers une base $\{v_1, \dots, v_n\}$: il s'agit de former une matrice carrée de taille n , dont la k -ième colonne est formée des coordonnées du vecteur v_k dans la base $\{e_1, \dots, e_n\}$. C'est donc la matrice de l'unique endomorphisme de E envoyant e_k sur v_k pour $k = 1, \dots, n$. Ou encore, c'est la matrice $P \in M_n(\mathbb{R})$ dont les entrées satisfont :

$$v_i = \sum_{j=1}^n P_{ji} e_j.$$

On sait ensuite que si la matrice de passage de $\{e_1, \dots, e_n\}$ vers $\{v_1, \dots, v_n\}$ est P , alors on obtient les coordonnées $(y_1, \dots, y_n)$ dans $\{v_1, \dots, v_n\}$ à partir des coordonnées $(x_1, \dots, x_n)$ dans $\{e_1, \dots, e_n\}$ de la manière suivante :

$$\begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} = P \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix}.$$

Rappelons également comment on forme la matrice M d'un endomorphisme $f \in \text{End}(E)$ dans la base $\{e_1, \dots, e_n\}$. Il s'agit d'une matrice carrée de taille n , dont la k -ième colonne est formée des coordonnées du vecteur $v_i = f(e_k)$ dans la base $\{e_1, \dots, e_n\}$. Autrement dit, c'est la matrice $M \in M_n(\mathbb{R})$ dont les entrées satisfont :

$$f(e_i) = \sum_{j=1}^n M_{ji} e_j.$$

Alors, si A est la matrice d'un endomorphisme f dans la base $\{e_1, \dots, e_n\}$, la matrice A' de f dans la base $\{v_1, \dots, v_n\}$ est donnée par

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

2.7.2 Groupe orthogonal

Soit $\{e_1, \dots, e_n\}$ une base orthonormée d'un espace euclidien E ; soit $f \in \text{End}(E)$ un endomorphisme de E . On pose $v_i = f(e_i)$.

La question qui nous intéresse ici est la suivante : quelles sont les propriétés de M qui assurent que $\{f(e_1), \dots, f(e_n)\}$ soit une base orthonormée de E ?

Avant d'aborder cette question, commençons par un lemme qui nous sera utile ici, mais également plus tard dans le cours.

Lemme 3 (Produit scalaire et transposée). *Soit $M \in M_n(\mathbb{R})$. Pour tous $U, V \in \mathbb{R}^n$ des vecteurs colonnes,*

$$\langle MU, V \rangle_{\mathbb{R}^n} = \langle U, {}^t M V \rangle_{\mathbb{R}^n}.$$

Dans le cas complexe, pour $M \in M_n(\mathbb{C})$ et $U, V \in \mathbb{C}^n$,

$$\langle MU, V \rangle_{\mathbb{C}^n} = \langle U, {}^t \overline{M} V \rangle_{\mathbb{C}^n}.$$

Démonstration. C'est un simple calcul. Dans le cas réel, en utilisant (5), on voit que

$$\langle MU, V \rangle_{\mathbb{R}^n} = {}^t(MU)V = {}^tU {}^t M V = \langle U, {}^t M V \rangle_{\mathbb{R}^n}.$$

Dans le cas complexe, c'est le même calcul, en faisant attention aux conjugaisons (voir (6)) :

$$\langle MU, V \rangle_{\mathbb{C}^n} = {}^t(MU)\overline{V} = {}^tU {}^t M \overline{V} = {}^tU \overline{{}^t M V} = \langle U, {}^t \overline{M} V \rangle_{\mathbb{C}^n}.$$

□

Revenons à la question initiale (qui, rappelons-le se place dans un espace euclidien donc réel). Il faut donc calculer $\langle v_i, v_j \rangle$ et voir à quelle condition on obtient δ_{ij} . Faisons le calcul :

$$\langle v_i, v_j \rangle = \langle M e_i, M e_j \rangle = \langle {}^t M M e_i, e_j \rangle.$$

Or il est facile de voir que pour une matrice carrée $A = (A_{ij})_{i,j=1}^n$, on a $\langle A e_i, e_j \rangle = A_{ji}$. Nous avons donc démontré le résultat suivant.

Proposition 18. *La matrice M définit un changement de bases orthonormées si et seulement si*

$${}^t M M = I_n, \tag{12}$$

ou, d'une manière équivalente,

$$M^{-1} = {}^t M.$$

Définition 31. Une matrice $M \in M_n(\mathbb{R})$ qui satisfait (12) est dite **orthogonale**. L'ensemble des matrices orthogonales de dimension n est noté $O(n)$ et est appelé **groupe orthogonal**.

Remarque 21. Une matrice M est orthogonale si et seulement si elle transforme la base canonique de $\mathbb{R}^n$ (qui est orthonormée pour le produit scalaire canonique) en une base orthonormée de $\mathbb{R}^n$, donc si et seulement si ses vecteurs colonnes sont normés et deux à deux orthogonaux pour le produit scalaire usuel.

Remarque 22. Insistons sur le fait que l'analyse qui précède est valable dans les espaces euclidiens, donc dans des espaces **réels**. Dans le cas complexe, une matrice qui transforme une base orthonormée en une autre base orthonormée est dite **unitaire**, et satisfait l'équation :

$$\overline{^t M} M = I_n.$$

Proposition 19. $O(n)$ est un groupe, c'est-à-dire :

- $\forall A, B \in O(n), AB \in O(n)$,
- $I_n \in O(n)$,
- $\forall A \in O(n), A^{-1} \in O(n)$.

Démonstration. On peut le montrer soit par un calcul direct en utilisant (12), soit en utilisant les propriétés des applications orthogonales (l'isométrie vectorielle), que nous définissons dans le paragraphe suivant. $\square$

Définition 32. L'ensemble des matrices orthogonales de déterminant égal à 1 est appelé **groupe spécial orthogonal** et est noté $SO(n)$. Autrement dit :

$$SO(n) := \{M \in O(n) \mid \det(M) = 1\}.$$

2.7.3 Endomorphismes orthogonaux

Soit E un espace euclidien.

Théorème 12. Soit $f \in \text{End}(E)$. Les propriétés suivantes sont équivalentes :

1. f conserve le produit scalaire, i.e.

$$\forall x, y \in E, \quad \langle f(x), f(y) \rangle = \langle x, y \rangle. \quad (13)$$

2. f conserve la norme, i.e.

$$\forall x \in E, \quad \|f(x)\| = \|x\|.$$

3. si $\{e_1, \dots, e_n\}$ est une base orthonormée de E alors la matrice M de f dans cette base est orthogonale.

Si l'une de ces conditions est satisfaite, on dit que l'endomorphisme f est **orthogonal** (ou, rappelons-nous, une isométrie vectorielle). De plus si f est orthogonal alors $\det(M) = \pm 1$ et f est une bijection.

Remarque 23. Notons à l'aide de ce théorème que si une application envoie une base orthonormée particulière sur une base orthonormée, il envoie toute base orthonormée sur une base orthonormée (puisqu'il satisfait (13)).

Démonstration. $1 \Rightarrow 2$: ceci est immédiat en prenant $y = x$.

$2 \Rightarrow 1$: ceci s'obtient de la relation de calcul de $\langle f(x), f(y) \rangle$:

$$\langle f(x), f(y) \rangle = \frac{\|f(x+y)\|^2 - \|f(x)\|^2 - \|f(y)\|^2}{2}$$

$1 \Leftrightarrow 3$: si X et Y sont les vecteurs colonnes des coordonnées de x et y de $\mathbb{R}^n$ dans la base $(e_1, \dots, e_n)$ alors $f(x)$ (respectivement $f(y)$) a comme coordonnées MX (resp. MY) dans cette base. Comme la base $(e_1, \dots, e_n)$ est orthonormée, il suit que :

$$\langle f(x), f(y) \rangle = {}^t(MX)MY = {}^tX({}^tMM)Y.$$

$$\langle x, y \rangle = {}^tXY,$$

donc (13) est équivalent à ${}^tMM = I_n$, c'est-à-dire $M \in O(n)$.

Si ${}^tMM = I_n$ il résulte que $1 = \det I_n = \det({}^tMM) = \det({}^tM)\det(M) = \det(M)^2$ donc $\det(M) = \pm 1$. En particulier $\det(M) \neq 0$ donc f est bijective (ce qu'on savait puisque ${}^tMM = I_n$). $\square$

3 Formes bilinéaires et quadratiques

3.1 Motivation

Lorsqu'on optimise une fonctionnelle $F : E \rightarrow \mathbb{R}$ (c'est-à-dire qu'on en recherche selon les cas un maximum ou un minimum), il s'agit d'habitude d'applications plus complexes que des applications linéaires ou affines (au demeurant celles-ci n'ont pas de minimum en général ; par exemple $f : \mathbb{R} \rightarrow \mathbb{R}$, $f(x) = ax + b$ a $-\infty$ comme infimum).

L'exemple le plus simple de fonctions après les applications affines est alors donnée par les polynômes de degré au plus 2 (mais de plusieurs variables). De plus, celui-ci se rencontre souvent dans les applications, voir par exemple le paragraphe 2.5.2, et cet exemple a également des conséquences dans un cadre plus général, voir le paragraphe 5.1 plus loin. On considère donc $E = \mathbb{R}^n$ et la fonction $f : \mathbb{R}^n \rightarrow \mathbb{R}$: pour $X = (x_1, \dots, x_n) \in \mathbb{R}^n$,

$$f(x_1, \dots, x_n) = \alpha + \sum_{i=1}^n \alpha_i x_i + \sum_{i,j=1}^n \beta_{ij} x_i x_j.$$

Une question naturelle est de déterminer s'il existe un minimum et d'expliquer comment le trouver. On sait lorsqu'il n'y a qu'une variable que si X_0 est un minimum de f alors la dérivée de f est nulle au point X_0 . Cela reste vrai à plusieurs variables : les dérivées partielles de f obtenues en gelant $(n-1)$ coordonnées et en dérivant par rapport à la variable restante sont nulles.

En changeant la variable X par $Y = X - X_0$, le développement de Taylor à l'ordre 2 (à plusieurs variables) montre que :

$$\begin{aligned} f(X) &= f(X_0) + f'(X_0)(X - X_0) + Q(X - X_0, X - X_0) \\ &= f(X_0) + Q(Y, Y), \end{aligned}$$

où Q est une "forme quadratique", c'est-à-dire un polynôme homogène de degré 2 :

$$Q(Y, Y) = \sum_{i,j=1}^n \beta_{ij} y_i y_j.$$

(Nous verrons cela avec un peu plus de détail à la Section 5.1.)

La question de la minimalité est donc résolue par l'analyse de la forme quadratique Q . En particulier, des questions importantes associées à cette analyse sont les suivantes :

- la positivité de Q : (si elle a lieu alors X_0 est un minimum dans l'exemple ci-dessus),
- l'existence d'une base où cette forme s'écrit plus simplement.

Dans une autre ligne d'idées, l'étude des formes bilinéaires permettra également de mieux comprendre à quoi ressemble un produit scalaire général.

3.2 Formes bilinéaires

Rappelons la définition d'une forme bilinéaire (voir page 11).

Définition 33. Soit E un espace vectoriel sur $\mathbb{K} = \mathbb{R}$ ou $\mathbb{K} = \mathbb{C}$. Une **forme bilinéaire** sur E est une application $f : E \times E \rightarrow \mathbb{K}$, linéaire par rapport à chacune de ses variables.

Exemple. Le produit scalaire dans un espace euclidien est une forme bilinéaire. Mais c'est loin d'être le seul exemple, on peut ainsi considérer dans $\mathbb{R}^2$ la forme bilinéaire $f((x_1, x_2), (y_1, y_2)) = -x_1 y_1 - 2x_1 y_2 + x_2 y_2$ (mais pas $g((x_1, x_2), (y_1, y_2)) = -x_1^2 - 2y_1 y_2 + x_1 y_2$!)

Remarque 24. On gardera en tête qu'une forme séquilinéaire **n'est pas** bilinéaire sur $\mathbb{C}$.

Soient maintenant un espace vectoriel réel E et $\{e_1, \dots, e_n\}$ une base de E . Alors pour $u = \sum_{i=1}^n U^i e_i$ et $v = \sum_{i=1}^n V^i e_i$ on aura

$$f(u, v) = f\left(\sum_{i=1}^n U^i e_i, \sum_{i=1}^n V^i e_i\right) = \sum_{i,j=1}^n U^i V^j f(e_i, e_j). \quad (14)$$

Définition 34. On appelle **matrice de la forme bilinéaire** f dans la base $\{e_1, \dots, e_n\}$, la matrice carrée A de taille n , d'entrées $A_{ij} = f(e_i, e_j)$ (c'est-à-dire la matrice ayant $f(e_i, e_j)$ en ligne i et colonne j).

Exemples. Le produit scalaire exprimé dans une base orthonormée a comme matrice la matrice identité. La matrice dans la base canonique de $\mathbb{R}^2$ de la forme bilinéaire donnée par $f((x_1, x_2), (y_1, y_2)) = -x_1y_1 - 2x_1y_2 + x_2y_2$ est la suivante

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

Proposition 20. Pour $u = \sum_{i=1}^n U^i e_i$ et $v = \sum_{i=1}^n V^i e_i$, on note par U le vecteur colonne $(U^i)_{i=1}^n$ et V le vecteur colonne $(V^i)_{i=1}^n$. On a alors :

$$f(u, v) = {}^t U A V. \quad (15)$$

Corollaire 3. Dans le cas réel, on a donc :

$$f(u, v) = \langle U, AV \rangle_{\mathbb{R}^n} = \langle {}^t A U, V \rangle_{\mathbb{R}^n}.$$

Remarque 25. Ci-dessus, la matrice ${}^t U$ est une matrice ligne, la matrice V est une matrice colonne, la matrice A est une matrice carrée. Donc le produit matriciel ${}^t U A V = {}^t U (A V)$ est le produit d'une matrice ligne et d'une matrice colonne (dans cet ordre), donc c'est une matrice scalaire (réduite à une entrée). Cela explique pourquoi on peut avoir l'égalité avec $f(u, v)$ qui est un élément de $\mathbb{K}$. En outre, le produit scalaire $\langle U, AV \rangle_{\mathbb{R}^n}$ est bien défini, puisqu'il s'opère entre deux vecteurs de $\mathbb{R}^n$ (écrits sous forme colonne).

Démonstration. C'est un simple calcul :

$$f(u, v) = \sum_{i,j=1}^n U^i V^j f(e_i, e_j) = \sum_{i=1}^n U^i \sum_{j=1}^n V^j A_{ij} = \sum_{i=1}^n U^i (A V)_i = {}^t U A V.$$

Le corollaire suit immédiatement puisque dans le cas réel, ${}^t U A V = \langle U, AV \rangle_{\mathbb{R}^n} = \langle {}^t A U, V \rangle_{\mathbb{R}^n}$. □

On rappelle également la définition de la symétrie d'une forme bilinéaire.

Définition 35. La forme bilinéaire $f : E \times E \rightarrow \mathbb{R}$ est dite **symétrique** lorsqu'elle satisfait

$$\forall u, v \in E, \quad f(u, v) = f(v, u).$$

Exemples. La forme bilinéaire sur $\mathbb{R}^2$ donnée par $g((x_1, x_2), (y_1, y_2)) = -x_1y_1 - 2x_1y_2 - 2x_2y_1 + x_2y_2$ est symétrique. La forme bilinéaire $f((x_1, x_2), (y_1, y_2)) = -x_1y_1 - 2x_1y_2 + x_2y_2$ ne l'est pas.

Définition 36. Une matrice carrée $A \in M_n(\mathbb{K})$ est dite **symétrique** lorsque ${}^tA = A$, c'est-à-dire lorsque $A_{ij} = A_{ji}$, pour tous $i, j = 1, \dots, n$.

Exemples. La matrice $\begin{pmatrix} 1 & 6 \\ 6 & 0 \end{pmatrix}$ est symétrique, pas la matrice $\begin{pmatrix} 1 & 6 \\ 7 & 0 \end{pmatrix}$.

On a alors l'équivalence suivante.

Proposition 21. Soit f une forme bilinéaire sur un espace euclidien E et soit A sa matrice dans une base $(e_i)_{i=1}^n$ de E . Alors f est symétrique si et seulement si A l'est.

Démonstration. Soit f symétrique. Alors $A_{ij} = f(e_i, e_j) = f(e_j, e_i) = A_{ji}$. Pour la réciproque on écrit u, v dans une base, on utilise (14) pour calculer $f(u, v)$ et $f(v, u)$ et on compare! $\square$

Proposition 22. Soit f une forme bilinéaire symétrique sur E un espace réel, on a pour $u = \sum_{i=1}^n U^i e_i$ et $v = \sum_{i=1}^n V^i e_i$:

$$f(u, v) = \langle AU, V \rangle_{\mathbb{R}^n} = \langle U, AV \rangle_{\mathbb{R}^n}.$$

Démonstration. Cela vient de (15). $\square$

3.3 Formes quadratiques

Définition 37. Soit $f : E \times E \rightarrow \mathbb{K}$ une forme bilinéaire sur un espace vectoriel E . On appelle **forme quadratique** associée à f l'application $s : E \rightarrow \mathbb{K}$ donnée par

$$s(u) = f(u, u). \quad (16)$$

On appelle forme quadratique (tout court), une application $s : E \rightarrow \mathbb{K}$ telle qu'il existe f une forme bilinéaire pour laquelle la formule (16) est valable.

Proposition 23. Soit $(e_i)_{i=1}^n$ une base de E , et soit A la matrice de la forme bilinéaire f dans cette base. Pour $u = \sum_{i=1}^n U^i e_i$, on note par U le vecteur colonne $(U^i)_{i=1}^n$; on a alors :

$$s(u) = {}^tU A U.$$

Corollaire 4. Dans le cas réel, on a :

$$s(u) = \langle U, AU \rangle_{\mathbb{R}^n}.$$

Démonstration. C'est une conséquence immédiate de la Proposition 20 (et de son corollaire). $\square$

En particulier s est un polynôme homogène de degré 2 en les U^i :

$$s(u) = \sum_{i,j=1}^n U^i U^j A_{ij},$$

d'où le nom de “quadratique”.

Remarque 26. Des formes bilinéaires différentes peuvent donner la même forme quadratique ! Par exemple, la forme bilinéaire $g((x_1, x_2), (y_1, y_2)) = x_1y_1 - 2x_1y_2 + 2x_2y_1 + x_2y_2$ et la forme bilinéaire $h((x_1, x_2), (y_1, y_2)) = x_1y_1 + x_2y_2$ donnent la même forme quadratique $s(x_1, x_2) = x_1^2 + x_2^2$.

Proposition 24. Soit s une forme quadratique. Parmi toutes les formes bilinéaires dont la forme quadratique associée est s , il en existe une unique qui soit symétrique. Celle-ci est donnée par

$$f(u, v) = \frac{s(u+v) - s(u) - s(v)}{2}. \quad (17)$$

Démonstration. Si s est associée à une forme bilinéaire symétrique f , alors on a

$$s(u+v) = f(u+v, u+v) = f(u, u) + f(u, v) + f(v, u) + f(v, v) = s(u) + 2f(u, v) + s(v).$$

Donc f est nécessairement donnée par (17), il y a donc au plus une forme bilinéaire symétrique qui fonctionne.

Réciproquement, soit s une forme quadratique, vérifions que f donnée par (17) est bien bilinéaire symétrique et que s est sa forme quadratique associée. Comme s est une forme quadratique, par définition, elle est associée à une certaine forme bilinéaire g (non nécessairement symétrique). Le même calcul que précédemment appliqué à g montre que :

$$\frac{g(u, v) + g(v, u)}{2} = \frac{s(u+v) - s(u) - s(v)}{2} = f(u, v).$$

Il est facile de voir sur cette formule que f est bien bilinéaire symétrique et que s est sa forme quadratique associée. $\square$

Définition 38. On appellera **matrice d'une forme quadratique** Q , la matrice de la forme bilinéaire symétrique qui lui est associée.

Proposition 25. Au niveau matriciel, si une forme quadratique q donnée dans une base (e_i) par

$$q\left(\sum_i X_i e_i\right) = \sum_{\substack{i, k=1, \dots, n \\ i \leq k}} q_{ik} X_i X_k,$$

alors la matrice A de la forme bilinéaire symétrique associée est $A_{ik} = q_{ii}$ si $i = k$ et $A_{ik} = q_{ik}/2$ sinon.

Démonstration. Il suffit d'observer que la matrice A en question est symétrique, et qu'elle donne bien la forme linéaire q : le calcul est direct (le " $q_{ik}/2$ " vient de ce que nous avons regroupé le terme en $X_i X_k$ avec le terme en $X_k X_i$ dans l'expression ci-dessus). $\square$

Exemple. Si $q(X, Y) = X^2 + XY - Y^2$. La matrice A doit satisfaire $q(X, Y) = (X, Y)A^t(X, Y)$ donc $A = \begin{pmatrix} 1 & 1/2 \\ 1/2 & -1 \end{pmatrix}$.

3.4 Méthode de Gauss

Commençons ce paragraphe par une définition concernant les formes quadratiques et bilinéaires **réelles**.

Définition 39. Soit un espace vectoriel **réel** E . Une forme quadratique $s : E \rightarrow \mathbb{R}$, est dite **positive** si pour tout $u \in E$, $s(u) \geq 0$. Une forme bilinéaire $f : E \times E \rightarrow \mathbb{R}$, est dite **positive** si la forme quadratique associée est positive, c'est-à-dire si $f(u, u) \geq 0$ pour tout $u \in E$.

La question que nous soulevons dans ce paragraphe est la suivante. Étant donnée une forme quadratique, peut-on l'exprimer sous une forme plus simple, qui en particulier permette de déterminer si elle est positive ou non ?

Rappel. Un petit rappel avant de commencer : la **forme canonique** d'un polynôme de degré 2 à une variable est la suivante :

$$aX^2 + bX + c = a \left(X + \frac{b}{2a} \right)^2 + \left(c - \frac{b^2}{4a} \right).$$

Premiers exemples. Pour introduire la méthode de Gauss, commençons par quelques exemples de formes quadratiques à deux variables X et Y :

- $Q_1(X, Y) = X^2 + 3Y^2$ est clairement positive.
- $Q_2(X, Y) = 3X^2 - Y^2$ n'est clairement pas positive (prendre $(X, Y) = (0, 1)$).
- $Q_3(X, Y) = (X + Y)^2 + 3(X - Y)^2$ est positive.
- $Q_4(X, Y) = 3(X + Y)^2 - (X - Y)^2$ n'est pas positive (prendre $(X, Y) = (-1, 1)$, ou, de manière plus générale, un couple (X, Y) qui annule $X + Y$ sans annuler $X - Y$).

Remarquons maintenant qu'en développant on obtient

$$Q_3(X, Y) = 4X^2 - 4XY + 4Y^2 \text{ et } Q_4(X, Y) = 2X^2 + 8XY + 2Y^2. \quad (18)$$

Sous ces formes-là, il est beaucoup moins facile de voir que Q_3 est positive et pas Q_4 !

D'où la question naturelle : trouver une méthode qui permette de remonter de la forme (18) à une forme telle que décrite plus haut. Un algorithme général permettant de faire ce travail est appelé *méthode (ou réduction) de Gauss*.

Décrivons-là sur les exemples du dessus : partons de la formule (18). On fait les opérations suivantes.

- On regroupe les termes en X . S'il n'y a pas de terme en X^2 , on travaillera plutôt avec Y en premier. (S'il n'y a pas non plus de terme en Y^2 , alors c'est plus compliqué ; voir plus bas).

Dans le cas de (18), c'est déjà fait, on écrit juste

$$Q_3(X, Y) = [4X^2 - 4XY] + 4Y^2 \quad \text{et} \quad Q_4(X, Y) = [2X^2 + 8XY] + 2Y^2.$$

- Ensuite, on met *sous forme canonique* en X (comme pour les polynômes de degré 2 en une variable!), on obtient :

$$Q_3(X, Y) = [4(X - Y/2)^2 - Y^2] + 4Y^2 \quad \text{et} \quad Q_4(X, Y) = [2(X + 2Y)^2 - 8Y^2] + 2Y^2.$$

Nous avons donc trouvé des formes réduites à une somme de carrés :

$$Q_3(X, Y) = 4(X - Y/2)^2 + 3Y^2 \quad \text{et} \quad Q_4(X, Y) = 2(X + 2Y)^2 - 6Y^2.$$

Mais nous voyons qu'il **n'y a pas unicité** de ce type d'écriture !

Cependant, sous les formes que nous venons d'obtenir, la positivité ou non de la forme quadratique se reconnaît facilement : soit tous les coefficients des "carrés" sont positifs ou nuls et la réponse est positive, soit ce n'est pas le cas, et on montre que la réponse est négative en prenant des vecteurs qui annulent les termes positifs. (On utilise pour cela le fait que les formes linéaires intervenant au sein des carrés sont indépendantes). Au passage, sur nos exemples, le résultat que nous trouvons est cohérent avec les formes dont nous sommes partis...

Principe général. De manière plus générale, l'algorithme est le suivant pour une forme quadratique générale à n variables $X_1, \dots, X_n$:

1. Regrouper tous les termes en X_1 . S'il n'y a pas de terme en X_1^2 , passer à X_2 et ainsi de suite. S'il n'y a aucun terme en X_k^2 , passer à l'étape 3, sinon à l'étape 2.
2. Mettre sous forme canonique en la variable sélectionnée. Aller à l'étape 4.
3. S'il n'y a pas de terme en X_k^2 , on isole les termes en X_1 et X_2 (s'il y en a, sinon on fait la même chose avec des variables pour lesquelles il y a effectivement des termes!). Nous obtenons une partie de la forme quadratique (on garde les autres termes pour plus tard) de la forme :



$$\tilde{Q} := \alpha X_1 X_2 + X_1 L_1 + X_2 L_2,$$

où L_1 et L_2 sont des expressions linéaires des variables restantes X_3, X_4 , etc. (L_1 et L_2 sont éventuellement nulles, par exemple s'il n'y a plus d'autres variables). On écrit alors

$$\tilde{Q} := \alpha \left[\left(X_1 + \frac{L_2}{\alpha} \right) \left(X_2 + \frac{L_1}{\alpha} \right) - \frac{L_1 L_2}{\alpha^2} \right].$$

On utilise enfin l'identité remarquable $ab = \frac{1}{4}[(a+b)^2 - (a-b)^2]$ pour obtenir

$$\tilde{Q} := \alpha \left[\left(X_1 + X_2 + \frac{L_1 + L_2}{\alpha} \right)^2 - \left(X_1 - X_2 + \frac{L_1 - L_2}{\alpha} \right)^2 \right] - \frac{L_1 L_2}{\alpha}.$$

4. On itère avec les termes “restants” (ceux que nous n'avons pas mis sous forme de carré ou de somme de carrés) : ceux-ci ont strictement moins de variables !

Cela permet d'obtenir le théorème suivant.

Théorème 13. *Soit Q une forme quadratique sur $\mathbb{K}^n$. Alors il existe un entier $r \leq n$, r formes linéaires $l_1, \dots, l_r$ sur $\mathbb{K}^n$ **linéairement indépendantes** et r coefficients $\alpha_1, \dots, \alpha_r$ de $\mathbb{K}$ tous non nuls, tels que pour tout x de $\mathbb{K}^n$ on ait*

$$Q(x) = \sum_{i=1}^r \alpha_i l_i(x)^2. \quad (19)$$

Une telle écriture d'une forme quadratique est appelée **forme réduite**.

Remarque 27. *Il faut bien retenir que les formes linéaires sont linéairement indépendantes. En particulier, il y en a moins que n ($r \leq n$). Cela fait partie de la définition d'une forme réduite.*

On peut alors répondre à notre question initiale, qui rappelons-le ne concerne que le cas de formes quadratiques réelles. Nous nous placerons dans ce cas jusqu'à la fin du Chapitre 3.

Proposition 26. *Une forme quadratique Q sur $\mathbb{R}^n$ est positive si et seulement si tous les α_i apparaissant dans la forme réduite (19) sont positifs ou nuls.*

Démonstration. Ceci est une conséquence du fait que les (l_i) sont linéairement indépendants et du lemme suivant.

Lemme 4. *Soit E un espace de dimension finie, et soit $(f_i)_{i=1}^k$ une famille libre de E^* . Il existe une base $(e_j)_{j=1 \dots n}$ de E telle que pour tout $i = 1 \dots k$ et $j = 1 \dots n$ on ait $f_i(e_j) = \delta_{ij}$.*

Preuve. Soit n la dimension de E (et donc de E^*). On complète les $(f_i)_{i=1}^k$ en une base de E^* , mettons $(f_i)_{i=1}^n$. Alors l'application $\Phi : E \rightarrow \mathbb{R}^n$ donnée par $x \mapsto (f_1(x), \dots, f_n(x))$ est un isomorphisme. En effet, la linéarité est claire, et il suffit donc de montrer Φ est injectif (pour des raisons de dimension). Maintenant, pour tout vecteur non nul v , il existe une forme linéaire ne s'annulant pas en ce point (par exemple $\langle v, \cdot \rangle$). Donc si $v \neq 0$, il ne peut pas appartenir à $\text{Ker}(\Phi)$, puisque les $(f_i)_{i=1}^n$ forment une base de E^* . Par suite, l'image réciproque de la base canonique de $\mathbb{R}^n$ par Φ répond à la question, ce qui achève la preuve du lemme. $\square$

Revenons à la preuve de la Proposition 26. Si les α_i sont tous positifs ou nuls, alors clairement Q est positive. Réciproquement, si les α_i ne sont pas tous positifs ou nuls, soit k un indice tel que $\alpha_k < 0$. On introduit la base $(e_j)_{j=1 \dots n}$ comme dans le Lemme 4. On vérifie alors que $Q(e_k) < 0$. $\square$

Remarque 28. Une bonne nouvelle : on peut voir en suivant l'algorithme précédent que si à un moment ou un autre on atteint l'étape 3 décrite précédemment ("S'il n'y a pas de carré"), cela veut dire que dans la réduction de Gauss que nous obtiendrons à la fin, il y aura des coefficients α_i négatifs. On peut donc s'arrêter là si la seule question qui se posait était la positivité.

Remarque 29. On peut donner une forme un peu plus compacte à (19). En effet, réordonnons les $\alpha_i l_i(x)^2$ pour mettre ceux qui correspondent à $\alpha_i > 0$ en premier (mettons de 1 à p) puis ceux qui correspondent à $\alpha_i < 0$ (de $p+1$ à r). Alors on pose $\tilde{l}_i = \sqrt{|\alpha_i|} l_i$ pour $1 \leq i \leq r$ et on obtient la forme

$$Q(x) = \sum_{i=1}^p \tilde{l}_i(x)^2 - \sum_{i=p+1}^r \tilde{l}_i(x)^2. \quad (20)$$

3.5 Signature d'une forme quadratique

Dans ce paragraphe, on ne considère encore que des formes quadratiques réelles.

Théorème 14 (Signature). Soit E un espace euclidien de dimension n . Soit $q : E \rightarrow \mathbb{R}$ une forme quadratique. Alors il existe une base $(e_1, \dots, e_n)$ de E , pas nécessairement orthonormée, et des entiers p et r avec $p \leq r \leq n$, tels que pour $u = \sum_{i=1}^n x_i e_i$, la forme q s'écrit :

$$q(u) = \sum_{i=1}^p x_i^2 - \sum_{i=p+1}^r x_i^2.$$

De plus les entiers p et r sont indépendants de la base $(e_i)_{i=1, \dots, n}$ qui met q sous cette forme. Le couple d'entiers $(p, r-p)$ sera appelé la **signature** de la forme quadratique q et sera noté $\text{sign}(q) = (p, r-p)$. L'entier r est appelé **rang** de la forme quadratique.

Remarque 30. L'écriture précédente sous-entend que la première somme (respectivement la seconde) est nulle (car l'ensemble d'indices est vide) si $p = 0$ (resp. si $p = r$).

Démonstration. L'existence d'une base (e_i) vient de la forme réduite (20) : on introduit les $\tilde{l}_i$ comme dans (20), puis des (e_i) qui leur sont associés par le Lemme 4. Les r premières coordonnées de x dans la base (e_i) sont alors précisément les $\tilde{l}_i(x)$.

Ensuite, pour l'unicité du couple (p, r) : soient $(e_i)_{i=1 \dots n}$ et $(e'_i)_{i=1 \dots n}$ deux bases et $u = \sum_{i=1}^n x_i e_i = \sum_{j=1}^n y_j e'_j$; on a

$$q(u) = \sum_{i=1}^p x_i^2 - \sum_{i=p+1}^r x_i^2 = \sum_{j=1}^{p'} y_j^2 - \sum_{j=p'+1}^{r'} y_j^2. \quad (21)$$

On introduit les notations :

$$\begin{aligned} F &= \text{Vect} \{e_1, \dots, e_p\}, & F' &= \text{Vect} \{e'_1, \dots, e'_{p'}\}, \\ G &= \text{Vect} \{e_{p+1}, \dots, e_r\}, & G' &= \text{Vect} \{e'_{p'+1}, \dots, e'_{r'}\}, \\ H &= \text{Vect} \{e_{r+1}, \dots, e_n\}, & H' &= \text{Vect} \{e'_{r'+1}, \dots, e'_n\}. \end{aligned}$$

Montrons que seules les paires (F, F') , (G, G') et (H, H') peuvent avoir une intersection non triviale i.e., $F \cap G' = \{0\} = F \cap H' = G \cap F' = \dots$, etc. Si $u \in F \cap G$ et $u \neq 0$, alors d'après (21), $u \in F$ implique $q(u) > 0$ et $u \in G'$ implique $q(u) < 0$. Ceci est une contradiction, donc $F \cap G' = \{0\}$.

Par conséquent, on a par exemple que les espaces F , G' et H' sont en somme directe : $F \oplus G' \oplus H' \subset E$. Il suit que $\dim(F) + \dim(G') + \dim(H') \leq \dim(E)$ c'est-à-dire $p + r' - p' + n - r' \leq n$ et donc, $p \leq p'$. Mais en changeant les rôles des espaces nous obtenons aussi $p' \leq p$, donc $p = p'$. On applique ensuite le même raisonnement pour montrer $r = r'$. $\square$

Exemples. La signature de la forme quadratique à trois variables $X^2 + (Y + Z)^2 - (Y - Z)^2$ est $(2, 1)$. Celle de $X^2 - (Y - Z)^2$ est $(1, 1)$.

On peut caractériser la positivité d'une forme quadratique en fonction de sa signature. Plus généralement que la Définition 39, on a les définitions suivantes.

Définition 40. Une forme quadratique $q(u) : E \rightarrow \mathbb{R}$ est dite *positive* si $q(x) \geq 0, \forall x \in E$. Elle est dite **définie positive** si $q(x) > 0, \forall x \neq 0$. De même pour une forme **négative** ou **définie négative**. Une matrice symétrique A est dite *positive* ou *définie positive* (ce que l'on notera $A \geq 0$ et $A > 0$ respectivement) si la forme quadratique $q(x) = \langle Ax, x \rangle$ l'est.

Exemples. Dans $\mathbb{R}^2$, la forme quadratique $Q_1(X, Y) = X^2 + (X + Y)^2$ est définie positive. La forme quadratique $Q_2(X, Y) = (X + Y)^2$ est positive, mais pas définie positive.

Introduisons également la définition de la non dégénérescence.

Définition 41. Une forme bilinéaire $f : E \times E \rightarrow \mathbb{R}$ est dite **non dégénérée** lorsque pour tout $u \in E \setminus \{0\}$, il existe $v \in E$ tel que $f(u, v) \neq 0$.

Corollaire 5. Soit $q(u) : E \rightarrow \mathbb{R}$ une forme quadratique, $n = \dim(E)$ et $\text{sign}(q) = (n_+, n_-)$. Alors :

- q est définie positive si et seulement si $n_+ = n$ (et donc $n_- = 0$); q est positive si et seulement si $n_- = 0$ (et donc $n_+ \leq n$),
- la forme bilinéaire symétrique f associée à q est non dégénérée si et seulement si $n_+ + n_- = n$.

Démonstration. Introduisons une base $(e_1, \dots, e_n)$ comme dans le Théorème 14. Si $n_+ = n$, il est clair que q est définie positive. Si $n_+ < n$, alors pour $n_+ < j \leq n$, on voit que $q(e_j) \leq 0$, et donc q n'est pas définie positive. De même dans le cas positif : il s'agit qu'il n'y ait pas de vecteur de base tel que $q(e_j) < 0$.

Prouvons maintenant la deuxième assertion. Notons que la forme symétrique f associée à q a la forme suivante : pour $u = \sum_{i=1}^p x_i e_i$ et $v = \sum_{i=1}^r y_i e_i$:

$$f(u, v) = \sum_{i=1}^p x_i y_i - \sum_{i=p+1}^r x_i y_i.$$

En effet, cette forme est clairement symétrique et satisfait $f(u, u) = q(u)$. Notons que sous cette forme et par définition de la signature, on a $r = n_- + n_+$. Il est clair que si $r < n$, alors on a $f(e_n, \cdot) = 0$ et la forme est clairement dégénérée. Mais si $r = n$, alors quel que soit $u \neq 0$, il existe $i \in \{1, \dots, n\}$ tel que la i -ème coordonnée de v dans la base $(e_1, \dots, e_n)$ soit non nulle. Alors on voit que $f(u, e_i) \neq 0$. $\square$

Remarque 31. On voit facilement à partir du Corollaire 5 qu'il y a équivalence pour une forme bilinéaire symétrique réelle **positive** entre “définie positive” et “non dégénérée positive”. Mais on peut le montrer plus facilement en utilisant pour le sens non trivial (“non dégénérée positive” $\Rightarrow$ “définie positive”) :

$$\forall u, v \in E, \forall t \in \mathbb{R}, f(u + tv, u + tv) = f(u, u) + 2tf(u, v) + t^2f(v, v) \geq 0,$$

donc le discriminant de ce polynôme en t est négatif ou nul :

$$f(u, v)^2 \leq f(u, u)f(v, v).$$

(C'est Cauchy-Schwarz pour une forme bilinéaire dont on ne sait pas encore que c'est un produit scalaire.) Pour $u \in E \setminus \{0\}$, on introduit v venant de la non dégénérescence de f , et on applique !

Remarque 32. Pour une forme bilinéaire non nécessairement positive, l'équivalence entre “non dégénérée” et “définie” est fausse en général.

4 Théorème de diagonalisation

Dans ce chapitre, nous utiliserons de nouveau les deux corps $\mathbb{R}$ et $\mathbb{C}$.

4.1 Motivation

Un des grands enjeux de l'algèbre linéaire est de “réduire” les matrices d'endomorphismes, par exemple de les diagonaliser (c'est-à-dire de trouver une base où la matrice de l'endomorphisme est diagonal) lorsque cela est possible. Cela permet de simplifier considérablement certains calculs, comme par exemple les puissances n -ièmes de ces matrices.

La même question peut se poser pour une forme bilinéaire : existe-t-il des bases où la matrice de la forme bilinéaire est plus simple à écrire, par exemple diagonale ? On voit en particulier que si la matrice A de la forme bilinéaire f est diagonale dans la base (e_i) , notons $A = \text{diag}(\alpha_i)$, on peut écrire facilement la forme bilinéaire $f(u, v) = \sum_i \alpha_i U^i V^i$, où U et V sont les vecteurs colonnes des coordonnées de u et v dans la base (e_i) .

Nous allons décrire dans la suite un théorème qui montre que certaines matrices sont diagonalisables (à la fois en tant que matrice d'endomorphisme et en tant que matrice de forme quadratique).

4.2 Changement de base pour les matrices de forme bilinéaire

Soit f une forme bilinéaire sur un espace réel de dimension finie E , et soit $(e_i)_{i=1}^n$ une base de E . On note A la matrice d'une forme bilinéaire f dans cette base. La question dans ce paragraphe est la suivante : comment s'écrit la matrice de f dans une autre base $(\tilde{e}_i)_{i=1}^n$?

Proposition 27. *Soit $P_{\tilde{e} \leftarrow e}$ la matrice de passage de $(e_i)_{i=1}^n$ vers $(\tilde{e}_i)_{i=1}^n$. Alors la matrice de f dans la base $(\tilde{e}_i)_{i=1}^n$ est donnée par*

$$\tilde{A} = {}^t P_{\tilde{e} \leftarrow e} A P_{\tilde{e} \leftarrow e}.$$

Démonstration. Soient u et v dans E ; notons U et V les vecteurs colonnes des coordonnées de u et v dans la base $(e_i)_{i=1}^n$ et $\tilde{U}$ et $\tilde{V}$ les vecteurs colonnes des coordonnées de u et v dans la base $(\tilde{e}_i)_{i=1}^n$. On sait alors que

$$\tilde{U} = P_{\tilde{e} \leftarrow e}^{-1} U \text{ et } \tilde{V} = P_{\tilde{e} \leftarrow e}^{-1} V$$

Mais en utilisant la Proposition 20, on voit que

$$f(u, v) = {}^t U A V = {}^t (P_{\tilde{e} \leftarrow e} \tilde{U}) A P_{\tilde{e} \leftarrow e} \tilde{V},$$

d'où le résultat. □

4.3 Vecteurs et valeurs propres

Définition 42. *Soit M une matrice $n \times n$ à entrées dans $\mathbb{K} = \mathbb{R}$ ou $\mathbb{K} = \mathbb{C}$. On dit que λ est une **valeur propre** de M lorsqu'il existe $v \neq 0$ de $\mathbb{R}^n$ (qu'on appelle **vecteur propre associé**) tel que*

$$Mv = \lambda v.$$

Remarque 33. *On parle de la même façon de **valeurs et vecteurs propres de l'endomorphisme** f représenté par M : l'équation devient*

$$f(v) = \lambda v.$$

Exemples.

- $M = \begin{pmatrix} 2 & 0 \\ 0 & 3 \end{pmatrix}$. Alors $\lambda = 3$ est une valeur propre et $v = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$ un vecteur propre associé car

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

- $M = \begin{pmatrix} 0 & 1 \\ -1 & 0 \end{pmatrix}$. Alors $\lambda = i$ est une valeur propre et $v = \begin{pmatrix} 1 \\ i \end{pmatrix}$ un vecteur propre associé car

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

Remarque 34. Si v est vecteur propre, tout αv avec $\alpha \in \mathbb{K}$ **non nul**, l'est également (pour la même valeur propre). Si v et w sont deux vecteurs propres non opposés associés à une même valeur propre, alors $v + w$ l'est également.

Définition 43. Pour λ une valeur propre de la matrice M , l'ensemble des vecteurs propres associés à cette valeur propre, réuni avec $\{0\}$, forme un s.e.v. de $\mathbb{R}^n$. Ce sous espace V_λ est appelé **sous-espace propre** associé à λ .

Remarque 35. On voit que

$$V_\lambda = \{v / Mv = \lambda v\} = \text{Ker}(M - \lambda I).$$

Une conséquence importante du fait que $\mathbb{C}$ soit algébriquement clos, c'est-à-dire que les polynômes de degré n aient toujours n racines comptées avec multiplicité (Théorème de d'Alembert-Gauss) est la suivante.

Lemme 5. Soit A une matrice carrée sur $\mathbb{K}$. Toute racine du polynôme caractéristique $P(\lambda) = \det(A - \lambda I)$ est valeur propre de A . Par conséquent, A a toujours au moins une valeur propre complexe.

Preuve. Ce lemme est classique : on sait que $\det(A - \lambda I) = 0 \Leftrightarrow \text{Ker}(A - \lambda I) \neq \{0\}$, ce qui donne le résultat. Rappelons au passage que le polynôme caractéristique d'une matrice carrée de taille n est de degré n (ce qui vient des propriétés de multilinéarité du déterminant). $\square$

On remarque que même si M a toutes ses entrées réelles il n'est pas nécessaire que les valeurs propres ou vecteurs propres complexes soient réels. Cependant ceci est vrai pour les matrices symétriques (c'est-à-dire, rappelons-nous, qui satisfont ${}^tM = M$).

Lemme 6. Soit A une matrice symétrique réelle de taille $n \times n$.

1. Soit $\lambda \in \mathbb{C}$ une valeur propre de A . Alors $\lambda \in \mathbb{R}$ et il existe un vecteur propre réel $v \in \mathbb{R}^n$ associé à λ .
2. Des vecteurs propres v_i et v_j de A qui correspondent à des valeurs propres différentes λ_i et λ_j sont orthogonaux deux à deux.

Démonstration. Notons tout d'abord que le fait que $A = (A_{ij})_{i,j=1,\dots,n}$ soit symétrique réelle permet d'écrire, avec (6) :

$$\forall u, v \in \mathbb{C}^n, \langle Au, v \rangle_{\mathbb{C}^n} = \langle u, {}^t\overline{A}v \rangle_{\mathbb{C}^n} = \langle u, Av \rangle_{\mathbb{C}^n}. \quad (22)$$

Pour le premier point, nous introduisons w un vecteur propre associé à une valeur propre λ (a priori les deux sont complexes). Alors on a :

$$\langle Aw, w \rangle_{\mathbb{C}^n} = \lambda \langle w, w \rangle_{\mathbb{C}^n} = \lambda \|w\|^2,$$

et par ailleurs, grâce à (22), on a

$$\langle Aw, w \rangle_{\mathbb{C}^n} = \langle w, Aw \rangle_{\mathbb{C}^n} = \bar{\lambda} \|w\|^2.$$

Il suit que λ est réel. Donc il existe un vecteur propre réel, car la matrice réelle $A - \lambda I$ n'est pas inversible.

Pour le deuxième point, on écrit

$$\lambda_i \langle v_i, v_j \rangle_{\mathbb{C}^n} = \langle Av_i, v_j \rangle_{\mathbb{C}^n} = \langle v_i, Av_j \rangle_{\mathbb{C}^n} = \lambda_j \langle v_i, v_j \rangle_{\mathbb{C}^n},$$

puisque λ_j est réel. Donc $\langle v_i, v_j \rangle_{\mathbb{C}^n} = 0$. □

4.4 Théorème de diagonalisation

Le théorème suivant est l'un des points les plus importants de ce cours.

Théorème 15 (Diagonalisation des matrices symétriques réelles). *Une matrice symétrique réelle $A \in M_n(\mathbb{R})$ est diagonalisable en base orthonormée, i.e. il existe $P \in O(n)$ tel que*

$$\Delta = P^{-1}AP \text{ soit une matrice diagonale.} \quad (23)$$

*De plus, les colonnes de P sont **des** vecteurs propres de A qui forment une base orthonormée de $\mathbb{R}^n$ (en particulier elles sont deux à deux orthogonales). Les éléments diagonaux de Δ sont **les** valeurs propres de A .*

Remarque 36. *Lorsque l'on peut obtenir la relation (23) avec P inversible seulement (mais pas nécessairement orthogonale), on dit que A est diagonalisable (tout court). On retiendra que **diagonaliser** une matrice, c'est déterminer les matrices A et P (et éventuellement P^{-1}).*

Le théorème précédent a pour corollaire le théorème suivant.

Théorème 16 (Diagonalisation des formes bilinéaires). *Soit f une forme bilinéaire symétrique sur l'espace euclidien E de dimension n . Alors il existe une base orthonormée $(e_i)_{i=1}^n$ de E telle que la matrice de f par rapport à cette base soit diagonale, c'est-à-dire*

$$f(e_i, e_j) = 0 \text{ si } i \neq j$$

L'action de f se calcule alors par

$$f(u, v) = \sum_{i=1}^n \alpha_i \langle u, e_i \rangle \langle v, e_i \rangle.$$

Démonstration du Théorème 15. Soit $A \in M_n(\mathbb{R})$ une matrice symétrique réelle (notons f l'endomorphisme associé qui a pour matrice A dans la base canonique). Notons $\{\lambda_1, \dots, \lambda_k\}$ l'ensemble de ses valeurs propres, c'est-à-dire des racines complexes de son polynôme caractéristique (comptées avec multiplicité, il y en aurait n). D'après le Lemme 6, les λ_i sont réels. Notons $V_i = \text{Ker}(A - \lambda_i I)$ l'espace propre dans $\mathbb{R}^n$ associé à λ_i . D'après le Lemme 6, on voit que pour $i \neq j$, $V_i \perp V_j$.

Montrons que

$$E = V_1 + \dots + V_k. \quad (24)$$

Si ce n'est pas le cas, notons $F = V_1 + \dots + V_k$, on a donc $F^\perp$ (son orthogonal dans $\mathbb{R}^n$) non réduit à $\{0\}$. Alors en utilisant (22) on voit que pour $x \in F^\perp$, on a pour tout $v_i \in V_i$,

$$\langle v_i, Ax \rangle = \langle Av_i, x \rangle = \lambda_i \langle v_i, x \rangle = 0.$$

Il suit que pour $x \in F^\perp$, $Ax \in F^\perp$. Donc la restriction $f|_{F^\perp}$ de f à $F^\perp$ peut être vue comme un endomorphisme de $F^\perp$. Donc il admet une valeur propre complexe. Celle-ci est bien entendu aussi une valeur propre de f , elle est donc réelle. Donc $f|_{F^\perp}$ admet un vecteur propre réel dans $F^\perp$! Ce qui contredit le fait que nous avons considéré toutes les valeurs et vecteurs propres précédemment.

Une fois (24) démontré, on choisit une base orthonormée dans chaque V_i . Comme $V_i \perp V_j$ pour $i \neq j$, leur réunion fait une base orthonormée de E . Dans cette base composée de vecteurs propres, la matrice de l'endomorphisme f est donc diagonale avec les éléments diagonaux les valeurs propres de f .

Enfin on sait que la matrice de passage de la base canonique de $\mathbb{R}^n$ vers une autre base orthonormée est une matrice orthogonale (voir la Proposition 18). $\square$

Démonstration du Théorème 16. On construit la matrice A de f dans une base orthonormée de E . Cette matrice est diagonalisable en base orthonormée d'après ce qui précède : il existe $P \in O(n)$ telle que (23) soit satisfait. Mais alors la matrice de f dans la nouvelle base est

$${}^tPAP = P^{-1}AP = \Delta,$$

c'est-à-dire que dans cette base, la matrice de la forme bilinéaire f est diagonale (voir la Proposition 27). $\square$

Remarque 37. *La diagonalisation en base orthonormée d'une matrice symétrique réelle n'est pas unique, dans le sens où P et Δ dans l'écriture (23) ne le sont pas. Dans la preuve du Théorème 15, on a vu que l'on avait le choix de l'ordre des valeurs propres et sous-espaces propres, mais aussi le choix d'une base orthonormée dans chaque sous-espace propre.*

4.5 Critères de positivité des formes quadratiques

Un corollaire du Théorème 16 est le suivant.

Corollaire 6. Soit $A \in M_n(\mathbb{R})$ une matrice symétrique réelle. Soit (n_+, n_-) la signature de la forme quadratique $q(x) = \langle Ax, x \rangle$ associée. Alors n_+ (respectivement n_-) est le nombre de valeurs propres strictement positives (resp. négatives) de A .

Par conséquent, A est positive (respectivement définie positive) si et seulement si les valeurs propres de A sont positives (resp. strictement positives).

Démonstration. Nous pouvons utiliser le Théorème 16 pour donner une nouvelle démonstration de la partie existence du Théorème 14 : il suffit de prendre la base (e_i) qui diagonalise A , puis de “normaliser” les e_i (les remplacer par $e_i/\sqrt{|\alpha_i|}$ si $\alpha_i \neq 0$ et ne pas transformer les autres) pour ramener les α_i à ± 1 . Au passage, on voit bien que n_+ (respectivement n_-) est le nombre de valeurs propres strictement positives (resp. négatives) de A . $\square$

Un autre critère classique de positivité des formes quadratiques est le suivant.

Théorème 17 (Sylvester). Soit A une matrice symétrique réelle. Alors A est définie positive si et seulement si tous les “mineurs principaux” $\det \left((A_{ij})_{i,j=1}^l \right)$ sont strictement positifs pour $l = 1, \dots, n$.



Démonstration. On démontre le théorème par récurrence.

Pour $n = 1$, A est un nombre réel $A = (A_{11})$, et $\langle Ax, x \rangle = A_{11}x^2$ est positif si et seulement si $A_{11} > 0$; par ailleurs A_{11} est le (seul) mineur principal.

Supposons maintenant le résultat vrai pour des matrices de taille $n - 1$ et prouvons-le pour des matrices de taille n . Soit A une telle matrice. Notons par (λ_i, v_i) des couples valeurs propres/vecteurs propres de la matrice A , les v_i formant une base orthonormée de $\mathbb{R}^n$.

Nous observons que si l’on note par $A_{n-1} \in M_{n-1}(\mathbb{R})$ la matrice carrée “nord-ouest” des $n - 1$ premières lignes et colonnes de A , alors :

$$\left\langle A \begin{pmatrix} x_1 \\ \vdots \\ x_{n-1} \\ 0 \end{pmatrix}, \begin{pmatrix} x_1 \\ \vdots \\ x_{n-1} \\ 0 \end{pmatrix} \right\rangle_{\mathbb{R}^n} = \left\langle A_{n-1} \begin{pmatrix} x_1 \\ \vdots \\ x_{n-1} \end{pmatrix}, \begin{pmatrix} x_1 \\ \vdots \\ x_{n-1} \end{pmatrix} \right\rangle_{\mathbb{R}^n}.$$

Sens direct. Nous savons que $\langle Ax, x \rangle > 0$ pour tout $x \in \mathbb{R}^n$, $x \neq 0$, donc aussi en particulier pour tous les x de la forme

$$x = \begin{pmatrix} x_1 \\ \vdots \\ x_{n-1} \\ 0 \end{pmatrix}.$$

Par conséquent, A_{n-1} est une matrice définie positive donc tous les mineurs de A_{n-1} sont strictement positifs, donc les $n - 1$ premiers mineurs de A sont strictement positifs. Il ne reste donc plus qu’à montrer que $\det(A) > 0$. Mais le fait que A soit définie positive implique que $\lambda_i > 0$ pour tous les i et comme $\det(A) = \prod_{i=1}^n \lambda_i$, il suit que $\det(A) > 0$.

Sens réciproque. Si tous les mineurs sont strictement positifs on applique l'hypothèse pour déduire que A_{n-1} est une matrice définie positive. Comme $\det(A) > 0$ on a l'alternative suivante : soit toutes les valeurs propres sont strictement positives (et donc A est définie positive), soit au moins deux valeurs propres λ_i et λ_j sont strictement négatives. Dans ce cas, il existe au moins une combinaison linéaire $\alpha v_i + \beta v_j$ avec $\alpha^2 + \beta^2 \neq 0$ telle que $\alpha v_i + \beta v_j$ ait zéro sur la dernière coordonnée. Puisque nous avons démontré que A_{n-1} était définie positive il suit que $\langle A(\alpha v_i + \beta v_j), \alpha v_i + \beta v_j \rangle > 0$. Mais d'un autre coté

$$\langle A(\alpha v_i + \beta v_j), \alpha v_i + \beta v_j \rangle = \alpha^2 \lambda_i + \beta^2 \lambda_j < 0,$$

d'où une contradiction. □

4.6 Deux applications de la diagonalisation

Donnons deux applications immédiates de la diagonalisation. Soit $A \in M_n(\mathbb{R})$ une matrice diagonalisable.

1. Puissance d'une matrice, récurrence linéaire. Si l'on a la relation (23), alors il suit immédiatement que

$$A^k = P \Delta^k P^{-1}.$$

Or Δ^k est très simple à calculer :

$$\text{si } \Delta = \begin{pmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \dots & 0 & \lambda_n \end{pmatrix}, \text{ alors } \Delta^k = \begin{pmatrix} \lambda_1^k & 0 & \dots & 0 \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \dots & 0 & \lambda_n^k \end{pmatrix}.$$

Application : récurrence linéaire. La diagonalisation est très utilisée pour calculer le résultat d'une suite récurrente vectorielle : si la suite $(x_k)_{k \in \mathbb{N}}$ d'éléments de $\mathbb{R}^n$ est donnée par la relation de récurrence

$$x_{k+1} = Ax_k,$$

alors on montre facilement que

$$x_k = A^k x_0.$$

Cas non homogène. On peut aussi considérer les suites satisfaisant des récurrences du type :

$$x_{k+1} = Ax_k + b, \tag{25}$$

où b est un vecteur fixe de $\mathbb{R}^n$. La méthode suivante fonctionne *si 1 n'est pas valeur propre de A* (au moins dans ce cas devrions-nous dire). Il s'agit de déterminer $v \in \mathbb{R}^n$ un "point fixe" de la relation de récurrence, i.e. un point satisfaisant :

$$v = Av + b.$$

Si 1 n'est pas valeur propre de A , cela est toujours possible car dans ce cas $A - I$ est injectif donc inversible, et on prend $v = -(A - I)^{-1}(b)$. Et dans ce cas, on observe qu'une suite (x_n) satisfait (25) si et seulement si la suite de terme général $y_n = x_n - v$ satisfait

$$y_{k+1} = Ay_k,$$

ce que l'on sait calculer par la puissance de matrice !

2. Exponentielle de matrice. On se propose ici de donner un sens à l'exponentielle d'une matrice carrée, qu'on notera $\exp(A)$, dans le cas où A est symétrique réelle. Comme on sait que pour x réel, on a :

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

on a envie de poser

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

Grâce au calcul précédent, on peut voir que pour les matrices symétriques réelles A diagonalisées comme dans (23), on a

$$\sum_{k=0}^N \frac{A^k}{k!} = P \left(\sum_{k=0}^N \frac{\Delta^k}{k!} \right) P^{-1}.$$

Cela permet de voir qu'on a bien convergence (élément par élément de la matrice) vers

$$\sum_{k=0}^{\infty} \frac{A^k}{k!} = P \begin{pmatrix} \exp(\lambda_1) & 0 & \dots & 0 \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \dots & 0 & \exp(\lambda_n) \end{pmatrix} P^{-1}.$$

Remarque 38. *En fait, on peut définir l'exponentielle de matrice de toute matrice carrée sur $\mathbb{R}$ ou $\mathbb{C}$ (non nécessairement diagonalisable), mais cela est un peu plus délicat dans le cas général (et c'est beaucoup plus difficile à calculer). Et on peut étendre le procédé à d'autres fonctions que la fonction exponentielle, en particulier aux fonctions entières (c'est-à-dire développables en série entières avec rayon de convergence infini).*

L'exponentielle de matrice a d'importantes conséquences en théorie des équations différentielles vectorielles (souvent appelées *systèmes différentiels*). On montre que A étant une matrice fixée, l'unique solution y (dépendant du temps et à valeurs dans $\mathbb{R}^n$), de l'équation différentielle vectorielle :

$$y' = Ay,$$

avec condition initiale

$$y(0) = y_0,$$

est donnée par

$$y(t) := \exp(tA)y_0.$$

Ce n'est pas sans rappeler le cas où y est à valeurs réelles !

5 Applications

5.1 Optimisation de fonctions (conditions nécessaires et condition suffisante)

Dans ce paragraphe, on s'intéresse au problème de déterminer les maxima et minima d'une fonction de plusieurs variables. On donne des conditions pour que le minimum ou maximum *local* d'une fonction soit réalisé en un point. Cela est intimement relié à la théorie des formes quadratiques et à leur positivité. Les démonstrations précises de ce qui suit relèvent plutôt d'un cours d'analyse ou de calcul différentiel que d'un cours d'algèbre ; elles seront omises.

Soit f une fonction définie sur $\mathbb{R}^n$ aux valeurs dans $\mathbb{R}$. Remarquons qu'on passe facilement de la notion de minimum à celle de maximum : on vérifie en effet immédiatement que :

$$\min_{x \in \Omega} f(x) = -\max_{x \in \Omega} (-f(x)).$$

Nous donnons à présent les principaux outils intervenant dans l'analyse.

5.1.1 Outils en œuvre

• **Dérivées partielles.** Pour une fonction f de plusieurs variables, mettons $x = (x_1, \dots, x_n)$ et $f(x) = f(x_1, \dots, x_n)$, la **dérivée partielle** de f au point $\bar{x} := (\bar{x}_1, \dots, \bar{x}_n)$ par rapport à sa k -ième variable x_k est la dérivée au point $x_k = \bar{x}_k$ de la fonction d'une variable :

$$x_k \mapsto (\bar{x}_1, \dots, \bar{x}_{k-1}, x_k, \bar{x}_{k+1}, \dots, \bar{x}_n)$$

Cette dérivée partielle est notée $\frac{\partial f}{\partial x_k}|_{x=\bar{x}}$.

Exemple. Soit $f(x, y) = xy + 2x + y$. Alors

$$\frac{\partial f}{\partial x}|_{(x,y)=(\bar{x},\bar{y})} = \bar{y} + 2 \quad \text{et} \quad \frac{\partial f}{\partial y}|_{(x,y)=(\bar{x},\bar{y})} = \bar{x} + 1.$$

• **Dérivées partielles secondes.** On peut ensuite définir les dérivées partielles secondes, comme les dérivées partielles des dérivées partielles (premières). Ainsi $\frac{\partial^2 f}{\partial x_j \partial x_k}|_{x=\bar{x}}$ est obtenu en dérivant f d'abord par rapport à la j -ième variable, puis par rapport à la k -ième.

Exemple. Soit encore $f(x, y) = xy + 2x + y$. Alors

$$\frac{\partial^2 f}{\partial x^2}|_{(x,y)=(\bar{x},\bar{y})} = \frac{\partial^2 f}{\partial y^2}|_{(x,y)=(\bar{x},\bar{y})} = 0 \quad \text{et} \quad \frac{\partial^2 f}{\partial x \partial y}|_{(x,y)=(\bar{x},\bar{y})} = \frac{\partial^2 f}{\partial y \partial x}|_{(x,y)=(\bar{x},\bar{y})} = 1.$$

• **Gradient.** On peut regarder la différentielle première de f au point $x^0 = (x_1^0, \dots, x_n^0)$, $Df(x^0)$, comme le “**gradient**” de f , c’est-à-dire le vecteur des dérivées partielles, et en déduire $Df(x^0)(u)$, la dérivée de f au point x^0 dans la direction u :

$$\nabla f(x^0) = \left(\frac{\partial f}{\partial x_i} \Big|_{x=x^0} \right)_{i=1}^n \quad \text{et} \quad Df(x^0)(u) = \langle \nabla f(x^0), u \rangle_{\mathbb{R}^n}.$$

• **Matrice hessienne.** Nous pouvons voir la dérivée seconde au point x^0 , $D^2f(x^0)$, comme la forme bilinéaire obtenue à partir de la **matrice hessienne** des dérivées partielles d’ordre 2 au point x^0 :

$$H_{f,x^0} = \left(\frac{\partial^2 f}{\partial x_i \partial x_j} \Big|_{x=x^0} \right)_{i,j=1}^n \quad \text{et} \quad D^2f(x^0)(u, v) = \langle H_{f,x^0} u, v \rangle_{\mathbb{R}^n}$$

Le lemme de Schwarz montre que la matrice hessienne est symétrique.

Exemple. Pour $f(x, y) = x^2 + e^y + 3xy$, on a :

$$Df(\bar{x}, \bar{y})(h_1, h_2) = \left\langle \begin{pmatrix} 2\bar{x} + 3\bar{y} \\ e^{\bar{y}} + 3\bar{x} \end{pmatrix}, \begin{pmatrix} h_1 \\ h_2 \end{pmatrix} \right\rangle,$$

$$D^2f(\bar{x}, \bar{y})((h_1, h_2), (h'_1, h'_2)) = \left\langle \begin{pmatrix} 2 & 3 \\ 3 & e^{\bar{y}} \end{pmatrix} \begin{pmatrix} h_1 \\ h_2 \end{pmatrix}, \begin{pmatrix} h'_1 \\ h'_2 \end{pmatrix} \right\rangle.$$

• **Minima et maxima locaux.** On dit qu’une fonction f a un **minimum** (resp. **maximum**) **local** au point x^0 si $f(x^0)$ minimise (resp. maximise) les valeurs de f , parmi les points d’une certaine boule centrée en x^0 , c’est-à-dire parmi les points suffisamment proches de x^0 . Autrement dit

$$f \text{ a un minimum (respectivement maximum) local en } x^0 \iff$$

$$\exists r > 0 \text{ tel que } \forall x \in \mathbb{R}^n, \text{ si } \|x - x^0\| \leq r, \text{ alors } f(x^0) \leq f(x) \text{ (resp. } f(x^0) \geq f(x)).$$

• **Formule de Taylor-Young à plusieurs variables.**

Théorème 18. Si $f : \mathbb{R}^n \rightarrow \mathbb{R}$ est deux fois dérivable par rapport à chacune de ses variables et de dérivées partielles continues, on a :

$$f(x) = f(x_0) + Df(x_0)(x - x_0) + \frac{1}{2}D^2f(x_0)(x - x_0, x - x_0) + o(\|x - x_0\|^2).$$

5.1.2 Conditions d'optimalité

Les résultats suivants sont des conséquences du Théorème 18. Ils concernent des fonctions $f : \mathbb{R}^n \rightarrow \mathbb{R}$, deux fois continûment dérivables par rapport à toutes leurs variables.

Proposition 28 (Condition nécessaire d'optimalité). *Si $x^0 \in \mathbb{R}^n$ est un extremum local (minimum ou maximum) de f alors $Df(x^0) = 0$.*

Remarque 39. *Cette équation servira pour trouver les “candidats” possibles pour être un optimum. Dans les applications, on a souvent seulement un nombre fini de tels points. Mais on notera bien que la réciproque est en général fausse : on peut trouver des fonctions f et des points x^0 tels que $Df(x^0) = 0$ (qu'on appelle des points “critiques”), et qui ne sont ni un minimum, ni un maximum local (pensons à la fonction xy au point 0 par exemple).*

Proposition 29 (Condition nécessaire d'optimalité, suite). *Si $x^0 \in \mathbb{R}^n$ est un minimum local (respectivement un maximum) de f alors $D^2f(x^0)$ est une forme bilinéaire positive (resp. négative).*

Remarque 40. *Voilà qui permettra d'être encore plus sélectif, mais ce n'est toujours pas suffisant (pensons à la fonction x^3y^3 au point 0 par exemple).*

Certaines conditions assurent toutefois qu'on a bien un extremum local.

Proposition 30 (Condition suffisante d'optimalité). *Si $x^0 \in \mathbb{R}^n$ est tel que $Df(x^0) = 0$ et que $D^2f(x^0)$ soit une forme bilinéaire **définie** positive (resp. définie négative), alors x^0 est un minimum (resp. maximum) local.*

Remarque 41. *Cette fois-ci la condition est suffisante, mais elle n'est pas nécessaire (pensons à la fonction x^2y^2 au point 0 par exemple). Insistons également sur le fait que le minimum (resp. maximum) que nous obtenons est **local**.*

5.2 Coniques



On considère les courbes du plan définies par une équation polynomiale du deuxième degré, c'est-à-dire pour $a, b, c, d, e, f \in \mathbb{R}$ l'ensemble des points

$$\mathcal{C} = \{(X, Y) \in \mathbb{R}^2 / aX^2 + bXY + cY^2 + dX + eY + f = 0\}. \quad (26)$$

Les courbes qui satisfont cette équation sont appelées “coniques”.

Avec les notations

$$A = \begin{pmatrix} a & b/2 \\ b/2 & c \end{pmatrix}, \quad B = \begin{pmatrix} d \\ e \end{pmatrix}, \quad u = \begin{pmatrix} X \\ Y \end{pmatrix},$$

l'équation de (26) admet l'écriture matricielle

$$\mathcal{C} = \{u \in \mathbb{R}^2 / \langle Au, u \rangle + \langle B, u \rangle + f = 0\}.$$

Nous distinguerons plusieurs cas.

Premier cas. Si $\text{rang}(A) = 0$. Dans ce cas, on a donc $A = 0$; l'équation sera $\langle B, u \rangle + f = 0$ et donc la courbe est une droite.

Deuxième cas. Si $\text{rang}(A) = 2$. On peut alors résoudre de manière unique l'équation $2Aw + B = 0$ d'inconnue w (car A est inversible). Ensuite, par le changement d'inconnue $u = w + v$ on obtient l'équation suivante, en posant $\alpha = -\langle Aw, w \rangle + \langle B, w \rangle + f$:

$$\mathcal{C} = \{u = w + v \in \mathbb{R}^2 / \langle Av, v \rangle = \alpha\}.$$

On écrit alors la décomposition

$$A = {}^tU \begin{pmatrix} \lambda & 0 \\ 0 & \mu \end{pmatrix} U,$$

(avec λ et μ non nuls pour des raisons de rang). Cherchons ce que devient l'équation dans la base des vecteurs colonnes de U . On note $V = Uv$ et avec la notation $V = (X, Y)$ l'équation devient

$$\lambda X^2 + \mu Y^2 = \alpha.$$

La courbe sera donc

- une ellipse si λ et μ sont tous les deux positifs ou tous les deux négatifs (strictement pour des raisons de rang), cette ellipse étant éventuellement vide si α a le mauvais signe.
- une hyperbole sinon.

Exemple. $(X/2)^2 + Y^2 = 2$ donne une ellipse, $Y^2 - X^2 = 1$ une hyperbole.

Troisième cas. Si $\text{rang}(A) = 1$. Dans ce cas, on peut écrire

$$A = {}^tU \begin{pmatrix} \lambda & 0 \\ 0 & 0 \end{pmatrix} U,$$

(avec $\lambda \neq 0$). Après changement de repère (comme précédemment) on obtient l'équation $\lambda X^2 + \alpha X + \beta Y + \gamma = 0$. Alors nous avons deux sous-cas différents.

- Si $\beta = 0$, alors

$$\mathcal{C} = \{u = (X, Y) \in \mathbb{R}^2 / \lambda X^2 + \alpha X + \gamma = 0\}.$$

Selon le nombre de solutions de l'équation en X , on obtient l'ensemble vide, une ou deux droites parallèles à l'axe des ordonnées dans les coordonnées (X, Y) , d'équation $X = x_1$ et $X = x_2$.

- Si $\beta \neq 0$, on obtient la parabole

$$\mathcal{C} = \left\{ u = (X, Y) \in \mathbb{R}^2 / Y = -\frac{\lambda X^2 + \alpha X + \gamma}{\beta} \right\}.$$

5.3 Interprétation géométrique des groupes $O(2)$ et $O(3)$

Dans cette partie, nous décrivons géométriquement les isométries vectorielles dans le plan et l'espace tridimensionnel euclidiens.

5.3.1 Le groupe $O(2)$

Soit $A \in O(2)$, $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$. Les vecteurs $\begin{pmatrix} a \\ c \end{pmatrix}$ et $\begin{pmatrix} b \\ d \end{pmatrix}$ forment donc une base orthonormée de $\mathbb{R}^2$. On en déduit que :

- $a^2 + c^2 = 1$ donc il existe θ avec $a = \cos \theta$, $c = \sin \theta$,
- $b^2 + d^2 = 1$ donc il existe ϕ avec $b = \cos \phi$, $d = \sin \phi$,
- $ab + cd = 0$ donc $\cos \theta \cos \phi + \sin \theta \sin \phi = 0$ donc $\cos(\phi - \theta) = 0$; en particulier $\sin(\phi - \theta) = \pm 1$.

Donc $d = \sin(\phi - \theta + \theta) = \sin(\phi - \theta) \cos \theta + \cos(\phi - \theta) \sin \theta = \sin(\phi - \theta) \cos \theta$. Par conséquent, nous sommes dans l'une des deux situations suivantes.

Première situation. Si $\sin(\phi - \theta) = 1$, alors $d = \cos \theta$ et $b = -\sin \theta$. Donc

$$A = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}.$$

Dans ce cas $\det A = 1$ donc $A \in SO(2)$.

Interprétation géométrique. Soit le vecteur $\begin{pmatrix} x \\ y \end{pmatrix}$ dont l'écriture en coordonnées polaires est $\begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} \rho \cos \alpha \\ \rho \sin \alpha \end{pmatrix}$. Alors on a :

$$A \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix} \begin{pmatrix} \rho \cos \alpha \\ \rho \sin \alpha \end{pmatrix} = \begin{pmatrix} \rho \cos(\alpha + \theta) \\ \rho \sin(\alpha + \theta) \end{pmatrix}$$

donc A est une rotation d'angle θ dans le sens trigonométrique.

Deuxième situation. Si $\sin(\phi - \theta) = -1$, alors $d = -\cos \theta$ et $b = \sin \theta$. Donc

$$A = \begin{pmatrix} \cos \theta & \sin \theta \\ \sin \theta & -\cos \theta \end{pmatrix}.$$

Dans ce cas $\det A = -1$ donc $A \notin SO(2)$.

Interprétation géométrique. Soit le vecteur $\begin{pmatrix} x \\ y \end{pmatrix}$ dont l'écriture en coordonnées polaires est $\begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} \rho \cos \alpha \\ \rho \sin \alpha \end{pmatrix}$. Alors

$$A \begin{pmatrix} x \\ y \end{pmatrix} = \begin{pmatrix} \cos \theta & \sin \theta \\ \sin \theta & -\cos \theta \end{pmatrix} \begin{pmatrix} \rho \cos \alpha \\ \rho \sin \alpha \end{pmatrix} = \begin{pmatrix} \rho \cos(\theta - \alpha) \\ \rho \sin(\theta - \alpha) \end{pmatrix}$$

donc A est une symétrie (sur le cercle de rayon ρ) par rapport à la droite qui fait un angle $\theta/2$ avec Ox dans le sens trigonométrique.

5.3.2 Le groupe $O(3)$

Théorème 19. Soit $A \in O(3)$ et f un endomorphisme de $\mathbb{R}^3$ dont A est la matrice dans la base canonique notée (e_1, e_2, e_3) . Alors il existe une base orthonormée (v_1, v_2, v_3) telle que la matrice M de f dans cette nouvelle base soit

$$M = \begin{pmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & \det(A) \end{pmatrix}.$$



Démonstration. Remarquons d'abord que si λ est valeur propre réelle de A alors $\lambda = \pm 1$; en effet $Av = \lambda v$ mais $\|Av\| = \|v\|$ donc $|\lambda| = 1$.

Prouvons ensuite que $\det(A)$ est valeur propre de A . En effet, les valeurs propres de A sont des racines du polynôme caractéristique $\det(A - \lambda I)$ qui est à coefficients réels et de degré 3. Il a donc soit une, soit trois racines réelles. Examinons ces deux cas.

- S'il a une racine réelle λ et deux complexes λ_2 et λ_3 . Il suit que $\lambda_3 = \overline{\lambda_2}$. Alors $\det(A) = \lambda \lambda_2 \lambda_3 = \lambda |\lambda_2|^2$ donc $\det(A)$ et λ ont le même signe, et comme les deux sont de module 1, elles sont égales.
- S'il a trois valeurs propres réelles λ_1, λ_2 et λ_3 , elles sont nécessairement toutes ± 1 . Comme $\det(A) = \lambda_1 \lambda_2 \lambda_3$ au moins une aura le signe du $\det(A)$ et on obtient donc la conclusion.

Soit donc maintenant v un vecteur propre (de norme 1) correspondant à la valeur propre $\lambda = \det(A)$. On complète v en une base orthonormée de $\mathbb{R}^3$ notée (v_1, v_2, v_3) avec $v = v_3$. Les entrées de la matrice M de A dans cette base sont alors $M_{ij} = \langle f(v_j), v_i \rangle_{\mathbb{R}^3}$. Par exemple $M_{i3} = \langle f(v_3), v_i \rangle_{\mathbb{R}^3} = \langle \lambda v_3, v_i \rangle_{\mathbb{R}^3} = \lambda \delta_{3i}$. Donc

$$M = \begin{pmatrix} M_{11} & M_{12} & M_{13} \\ M_{21} & M_{22} & M_{23} \\ 0 & 0 & \det(A) \end{pmatrix}.$$

Mais, comme f est orthogonale, M l'est aussi. En particulier les lignes 1 et 2 de M sont des vecteurs orthogonaux à la ligne 3, donc $M_{13} = M_{23} = 0$.

Notons maintenant

$$M_2 = \begin{pmatrix} M_{11} & M_{12} \\ M_{21} & M_{22} \end{pmatrix},$$

ce qui nous permet d'écrire

$$M = \begin{pmatrix} M_2 & 0 \\ 0 & 0 & \det(A) \end{pmatrix}.$$

Par conséquent

$$I_3 = {}^t M = \begin{pmatrix} {}^t M_2 M_2 & 0 \\ 0 & 0 & 1 \end{pmatrix},$$

donc $M_2 \in O(2)$. De plus $\det(A) = \det(M) = \det(M_2) \det(A)$ donc $\det(M_2) = 1$. □

Par conséquent, la matrice $A \in O(3)$ est dans l'une des situations suivantes :

- si $\det(A) = 1$ alors $A \in SO(3)$ donc c'est une rotation dans le plan $\{v\}^\perp$ orthogonal au vecteur propre de la valeur propre 1,
- si $\det(A) = -1$ alors $A \notin SO(3)$ et c'est une rotation dans le plan $\{v\}^\perp$ orthogonal au vecteur propre de la valeur propre -1 , suivie d'une symétrie par rapport au même plan.

En pratique, pour déterminer la transformation il faut déterminer v , ensuite considérer un vecteur u orthogonal à v et calculer l'angle θ entre u et $f(u)$. La transformation sera soit une rotation, soit une "rotation + symétrie" selon le signe du $\det(A)$. Remarquons enfin que l'on peut calculer l'angle θ en remarquant que $Tr(A) = Tr(M) = \det(A) + 2 \cos(\theta)$.

6 Annexe : structures algébriques



Dans ce chapitre, on rappelle les principales structures algébriques. On pourra ainsi donner un sens plus précis aux expressions "corps de base" ou "groupe orthogonal".

6.1 Lois de composition

Définition 44. Une *loi de composition interne* sur un ensemble E est une application

$$\circ : E \times E \rightarrow E$$

Il est usuel de noter l'image $\circ(x, y)$ du couple (x, y) par la loi de composition $\circ$ par $x \circ y$ ou, s'il n'y a pas d'ambiguïté, par xy .

Exemples. Voici quelques exemples classiques de lois de composition interne :

- l'addition “+” sur $\mathbb{N}, \mathbb{Z}, \mathbb{Q}, \mathbb{R}, \mathbb{C}$,
- la multiplication “.” sur $\mathbb{N}, \mathbb{Z}, \mathbb{Q}, \mathbb{R}, \mathbb{C}$,
- la soustraction “−” sur $\mathbb{Z}, \mathbb{Q}, \mathbb{R}, \mathbb{C}$ mais pas sur $\mathbb{N}$,
- la division sur $\mathbb{Q}^* (= \mathbb{Q} \setminus \{0\}), \mathbb{R}^*, \mathbb{C}^*$, mais pas sur $\mathbb{N}^*, \mathbb{N}, \mathbb{Z}, \mathbb{Q}, \mathbb{R}, \mathbb{C} \dots$

Définition 45. Une loi de composition (ou multiplication) externe sur un ensemble E par les éléments de $\mathbb{K}$ une application

$$\cdot : \mathbb{K} \times E \rightarrow E$$

Il est usuel de noter l'image $\cdot(\lambda, x)$ du couple (λ, x) par la loi de composition $\cdot$ par $\lambda \cdot x$ ou encore λx .

Définition 46. Une loi de composition interne $\circ$ est dite :

- **associative** si

$$(x \circ y) \circ z = x \circ (y \circ z), \quad \forall x, y, z \in E.$$

- **commutative** si

$$(x \circ y) = (y \circ x), \quad \forall x, y \in E.$$

Un élément $e \in E$ est dit **élément neutre** si

$$e \circ x = x \circ e = x, \quad \forall x \in E.$$

Si l'élément neutre existe, un élément $x \in E$ est dit **inversible** s'il existe $x' \in E$ tel que

$$x \circ x' = x' \circ x = e.$$

Dans ce cas, l'élément x' est appelé **l'inverse** de x noté x^{-1} . Pour une loi notée “+” l'inverse sera noté “− x ” et sera appelé l'opposé de x .

Remarque 42. Il est immédiat de montrer que l'élément neutre est unique.

Exemples. L'addition et la multiplication sur $\mathbb{N}, \mathbb{Z}, \mathbb{Q}, \mathbb{R}, \mathbb{C}$ sont associatives et commutatives, et y admettent un élément neutre. En revanche dans $\mathbb{N}$ les éléments ne sont pas inversibles pour l'addition (sauf 0). Et de même dans tous ces ensembles l'élément 0 n'est pas inversible pour la multiplication ! Enfin, la soustraction sur les mêmes ensembles n'est ni associative, ni commutative, et elle n'admet pas d'élément neutre.

6.2 Groupes

Définition 47. Un **groupe** est un ensemble E muni d'une loi de composition interne qui est associative, admet un élément neutre et est telle que tout élément est inversible. Si la loi est commutative le groupe est dit **abélien** ou **commutatif**.

Exemples.

- $(\mathbb{Z}, +)$, $(\mathbb{Q}, +)$, $(\mathbb{R}, +)$, $(\mathbb{C}, +)$, $(\mathbb{R}_+^*, \cdot)$ mais pas $(\mathbb{R}_+, +)$,
- $(\mathbb{Q}^*, \cdot)$, $(\mathbb{R}^*, \cdot)$, $(\mathbb{C}^*, \cdot)$ mais pas $(\mathbb{Z}^*, \cdot)$,
- $SO(2)$ qui est l'ensemble des rotations du plan $\mathbb{R}^2$ avec la composition des fonctions comme loi interne.

Proposition 31. *L'ensemble $GL(n; \mathbb{R})$ des matrices $n \times n$ inversibles est un groupe avec la multiplication comme loi de composition. De même pour $GL(n; \mathbb{C})$. Aucun des deux n'est groupe abélien (sauf si $n = 1$).*

6.3 Anneaux

Définition 48. *On appelle **anneau** tout ensemble A doté de deux lois de composition interne, notées $+$ et $\cdot$, telles que :*

- $(A, +)$ est un groupe commutatif,
- la loi $\cdot$ est associative,
- la loi $\cdot$ est distributive par rapport à la loi $+$, i.e., pour tous $x, y, z \in A$,

$$\begin{aligned}x \cdot (y + z) &= x \cdot y + x \cdot z, \\(x + y) \cdot z &= x \cdot z + y \cdot z,\end{aligned}$$

- la loi $\cdot$ possède un élément neutre.

L'anneau est dit commutatif si la loi $\cdot$ l'est.

Remarque 43. *On peut encore trouver des vieux grimoires (des années 1980 par exemple), où l'existence de l'élément neutre pour la loi $\cdot$ n'était pas requise. L'anneau était dit "unitaire" dans ce cas.*

Exemples.

- $(\mathbb{Z}, +, \cdot)$, $(\mathbb{Q}, +, \cdot)$, $(\mathbb{R}, +, \cdot)$, $(\mathbb{C}, +, \cdot)$.
- L'anneau des entiers modulo n : $\mathbb{Z}_n = \{0, \dots, n-1\}$ avec les lois $+$ et $\cdot$ suivantes :
 - ▷ $x + y$ est par définition le reste de la division euclidienne de $x + y$ par n ;
 - ▷ $x \cdot y$ est par définition le reste de la division euclidienne de $x \cdot y$ par n .
- L'ensemble de polynômes d'une variable (ou plusieurs) sur $\mathbb{R}$ (ou encore $\mathbb{Q}, \mathbb{C}, \mathbb{Z}$) : $\mathbb{R}[X]$ (ou $\mathbb{R}[X, Y]$, $\mathbb{R}[X, Y, Z]$, ...).
- L'ensemble $M_n(\mathbb{R})$ (ou $M_n(\mathbb{C})$) des matrices $n \times n$ à entrées réelles ou complexes muni de l'addition et de la multiplication (non commutative) des matrices.
- $(\mathbb{R}_+^*, \cdot, +)$ n'est pas un anneau car la distributivité n'est pas vérifiée. De même pour $(\mathbb{R}_+^*, \cdot, \cdot)$.

Notations. Dans un anneau $(A, +, \cdot)$ il est d'usage de noter par 0 l'élément neutre de la loi " $+$ " et par 1 l'élément neutre de la loi " $\cdot$ ". On obtient en particulier $a \cdot 0 = 0$ pour tout $a \in A$.

6.4 Corps

Définition 49. *Un ensemble $\mathbb{K}$ est appelé **corps** s'il est doté de deux lois de composition interne, notées $+$ et $\cdot$, telles que :*

- $(\mathbb{K}, +, \cdot)$ est un anneau commutatif,
- $(\mathbb{K} \setminus \{0\}, \cdot)$ est un groupe, où 0 est l'élément neutre de la loi $+$.

Remarque 44. *Dans certains livres du XXème siècle, un corps n'est pas nécessairement commutatif...*

Exemples. $(\mathbb{Q}, +, \cdot)$, $(\mathbb{R}, +, \cdot)$, $(\mathbb{C}, +, \cdot)$. Mais pas $(\mathbb{Z}, +, \cdot)$.

Références

- [1] Joseph Grifone. *Algèbre linéaire*. Éditions Cépaduès, Toulouse, 1990.
- [2] Bernard Guerrien. *Algèbre linéaire pour économistes*. Economica, 1997.
- [3] David Lay. *Algèbre linéaire. Théorie, exercices et applications*. De Boeck, 2004.
- [4] François Liret. *Mathématiques pour le DEUG. Algèbre et Géométrie ; 2ème année*. Dunod, 1999.

Index

- Angle non orienté, 18
- Anneau, 57
- Annulateur, 9
- Application antilinéaire, 12
- Application bilinéaire, 11
- Application linéaire, 3
- Application semilinéaire, 12
- Application sesquilinéaire, 12

- Base, 5
- Base duale, 9
- Base orthogonale, 19
- Base orthonormée, 19

- Changement de base, 28, 42
- Combinaison linéaire, 4
- Conditions d'optimalité, 51
- Conique, 51
- Coordonnées, 5
- Corps, 1, 58

- Dépendance linéaire, 4
- Dérivée partielle, 49
- Dimension, 6
- Distance, 10
- Droite de régression linéaire, 25
- Dual, 8

- Élément inversible, 56
- Élément neutre, 56
- Endomorphisme, 3
- Endomorphisme orthogonal, 30
- Espace euclidien, 14
- Espace vectoriel, 1
- Espace vectoriel normé, 10
- Exponentielle de matrice, 48

- Famille génératrice, 4
- Famille libre, 5
- Forme bilinéaire, 32
- Forme définie positive, 40
- Forme linéaire, 8
- Forme non dégénérée, 11, 40
- Forme positive, 11, 12, 36, 45
- Forme quadratique, 34
- Forme réduite, 38
- Forme symétrique, 11, 33
- Formule de Taylor-Young, 50

- Gradient, 50
- Groupe, 56
- Groupe orthogonal, 30, 53
- Groupe spécial orthogonal, 30

- Identité du parallélogramme, 17
- Image, 3
- Inégalité de Cauchy-Schwarz, 13
- Inverse, 56
- Isométrie vectorielle, 21
- Isomorphisme, 27
- Isomorphisme d'espaces euclidiens, 21

- Loi associative, 56
- Loi commutative, 56
- Loi de composition externe, 56
- Loi de composition interne, 55

- Matrice d'une application linéaire, 6
- Matrice d'un endomorphisme, 28
- Matrice d'une forme bilinéaire, 33
- Matrice d'une forme quadratique, 35
- Matrice de passage, 7, 28
- Matrice hessienne, 50
- Matrice orthogonale, 30
- Matrice symétrique, 34, 43
- Matrice unitaire, 30
- Mineurs principaux, 46
- Minimum, maximum, 49
- Minimum, maximum local, 50

- Norme, 10
- Norme euclidienne, 14

Norme hermitienne, 14
 Norme induite par le produit scalaire, 14
 Noyau, 3

 Orthogonal d'une partie, 23
 Orthogonalité, 14

 Problème des moindres carrés, 25
 Procédé d'orthonormalisation de Schmidt, 20, 24
 Produit matriciel, 7
 Produit scalaire, 11, 12
 Produit scalaire canonique, 15
 Produit scalaire hermitien, 12
 Projection orthogonale, 23
 Puissance d'une matrice, 47

 Récurrence linéaire, 47
 Réduction de Gauss, 36

 Signature, 39
 Somme directe $\oplus$, 3
 Sous-espace engendré, 4
 Sous-espace propre, 43
 Sous-espace vectoriel, 2
 Sous-espaces supplémentaires, 3
 Symbole de Kronecker δ_{ij} , 9
 Symétrie hermitienne, 12
 Système différentiel, 48

 Théorème de diagonalisation, 44
 Théorème de la base incomplète, 6
 Théorème de Pythagore, 18
 Théorème de Sylvester, 46
 Théorème du rang, 6
 Transposée, 15

 Valeur propre, 42
 Vecteur propre, 42