

Chapitre 24 : Isométries.

PTSI B Lycée Eiffel

26 juillet 2021

*Si quelqu'un, en l'éveil de son intelligence, n'a pas
été capable de s'enthousiasmer pour une telle architecture,
alors jamais il ne pourra réellement s'initier à la recherche théorique.*

ALBERT EINSTEIN, à propos de la géométrie euclidienne.

Pour ce dernier chapitre de l'année (mais oui, déjà), nous allons en quelque sorte boucler la boucle puisque nous ferons le lien entre la géométrie du plan et de l'espace, et le formidable outil que constitue l'algèbre linéaire (notamment les applications linéaires et les matrices). On peut en fait faire de la géométrie dans à peu près tous les espaces vectoriels comme on le fait dans ces espaces usuels. Pourquoi presque ? Parce qu'il nous manque tout de même une notion fondamentale pour cela, la notion d'orthogonalité de vecteurs dont découle tout l'aspect métrique du travail géométrique, c'est-à-dire les calculs de distance notamment. C'est l'objet de notre début de chapitre, que vous approfondirez largement l'an prochain. Nous appliquerons simplement ici ces notions à la classification des isométries du plan et de l'espace, qui vous rappelleront de bons souvenirs datant du chapitre sur les nombres complexes.

Objectifs du chapitre :

- comprendre la notion de distance dans un espace vectoriel autre que $\mathbb{R}^2$ et $\mathbb{R}^3$.
- connaître les différents types d'isométries dans le plan et dans l'espace.

1 Produit scalaire dans $\mathbb{R}^n$.

1.1 Norme et orthogonalité.

Certaines notions géométriques comme le parallélisme sont intrinsèques à la structure d'espace vectoriel. D'autres, au contraire, nécessitent d'ajouter une structure supplémentaire pour être définies. C'est le cas de la notion fondamentale en géométrie vectorielle de distance, qui découle de celle de produit scalaire (autrement dit, les notions d'orthogonalité et les calculs de distance sont liés, on ne peut pas faire l'un sans l'autre).

Définition 1. Le **produit scalaire** de deux vecteurs $u = (x_1, x_2, \dots, x_n)$ et $v = (y_1, y_2, \dots, y_n)$ appartenant à $\mathbb{R}^n$ est le nombre réel $u.v = \sum_{i=1}^n x_i y_i$.

On suppose bien sûr dans cette définition que les coordonnées sont prises dans la base canonique de $\mathbb{R}^n$. C'est donc une généralisation extrêmement naturelle du produit scalaire que nous connaissons bien dans $\mathbb{R}^2$ et dans $\mathbb{R}^3$.

Remarque 1. Le produit scalaire dans $\mathbb{R}^n$ est une application bilinéaire, symétrique et définie positive.

Remarque 2. Il existe plusieurs notations courantes pour les produits scalaires. Dans ce cours, nous prendrons toujours la plus économique en le notant $u.v$, mais on croise aussi régulièrement (u, v) , $\langle u, v \rangle$ ou encore $\langle u|v \rangle$ (surtout en physique pour cette dernière).

Remarque 3. On peut en fait plus généralement définir des produits scalaires (le pluriel est ici volontaire) sur n'importe quel espace vectoriel. Toute application vérifiant les trois propriétés fondamentales (bilinéaire, symétrique et définie positive) convient et permet de définir des notions d'orthogonalité et de distance (cf ci-après) permettant de faire de la géométrie dans l'espace considéré. Encore mieux, comme il existe en fait énormément de choix possibles pour l'application qui servira de produit scalaire, cela signifie qu'on peut définir par exemple plusieurs notions de distance très différentes sur un même espace vectoriel. Nous reviendrons sur quelques exemples dans la dernière partie de ce cours.

Définition 2. Soit $u \in \mathbb{R}^n$, la **norme de u** est le réel positif $\|u\| = \sqrt{u.u}$.

Remarque 4. La notion de **distance** découle immédiatement de celle de norme : la distance entre deux vecteurs u et v sera simplement égale à $\|u - v\|$, ce qui est cohérent avec la façon dont on calcule les distances dans $\mathbb{R}^2$ et dans $\mathbb{R}^3$. On garde par ailleurs la même interprétation du produit scalaire comme outil de détermination d'angles. En particulier, deux vecteurs u et v dans $\mathbb{R}^n$ sont **orthogonaux** si $u.v = 0$.

Proposition 1. Identité de polarisation :

$$\forall (u, v) \in (\mathbb{R}^n)^2, u.v = \frac{1}{4}(\|u + v\|^2 - \|u - v\|^2)$$

Démonstration. C'est une conséquence de la bilinéarité du produit scalaire : $\|u + v\|^2 - \|u - v\|^2 = u.u + 2u.v + v.v - u.u + 2u.v - v.v = 4u.v$. $\square$

Remarque 5. Cette identité (valable pour toute application vérifiant les propriétés théoriques du produit scalaire) est très importante d'un point de vue théorique, puisqu'elle signifie qu'on peut reconstituer le produit scalaire à partir de la connaissance de la norme. Autrement dit, la notion de distance ne peut pas être définie en-dehors du contexte du produit scalaire et de l'orthogonalité. Si on connaît une distance sur un espace vectoriel, on a forcément un produit scalaire caché derrière.

Définition 3. Un vecteur $u \in \mathbb{R}^n$ est **unitaire** (ou **normé**) si $\|u\| = 1$.

Les notions de bases orthogonales et orthonormales découlent naturellement de cette dernière définition.

1.2 Endomorphismes orthogonaux.

Définition 4. Un endomorphisme $f \in \mathcal{L}(\mathbb{R}^n)$ est appelé **endomorphisme orthogonal** ou **isométrie** si $\forall (u, v) \in E^2, f(u).f(v) = u.v$.

Remarque 6. Un endomorphisme est donc orthogonal s'il conserve le produit scalaire, ce qui est une condition naturelle. Attention tout de même au vocabulaire, une symétrie orthogonale est un endomorphisme orthogonal, mais pas une projection orthogonale !

Proposition 2. Un endomorphisme est orthogonal si et seulement s'il conserve la norme : $\forall u \in \mathbb{R}^n, \|f(u)\| = \|u\|$.

Démonstration. En effet, si f est orthogonale, on aura toujours $f(u).f(u) = u.u$, donc $\|f(u)\|^2 = \|u\|^2$. La réciproque découle de l'identité de polarisation. Cette propriété explique le terme d'isométrie, ce sont des applications qui conservent les longueurs. On retrouve bien entendu la définition classique des isométries dans le plan par exemple, que nous avons eu le bonheur d'étudier dans le chapitre consacré aux nombres complexes. Il n'existe pas d'isométries qui ne soient pas des applications linéaires dans $\mathbb{R}^n$ (exercice laissé au lecteur très motivé, ce n'est pas trivial). $\square$

Proposition 3. Un endomorphisme orthogonal est nécessairement bijectif.

Démonstration. Pour un endomorphisme d'un espace vectoriel de dimension finie, être bijectif ou injectif est équivalent. Or, si $f(u) = 0$, d'après la propriété précédente, on aura $\|u\| = 0$, donc $u = 0$, ce qui prouve l'injectivité et donc la bijectivité de f . $\square$

Proposition 4. Un endomorphisme f est orthogonal si et seulement s'il vérifie, au choix, l'une des deux conditions suivantes :

- L'image par f de toute base orthonormale de $\mathbb{R}^n$ est une base orthonormale.
- L'image par f d'une base orthonormale fixée de $\mathbb{R}^n$ est une base orthonormale.

Démonstration. Le fait que ces conditions soient nécessaire est évident : puisqu'une isométrie conserve à la fois les normes et les produits scalaires, l'image d'une base orthonormale par une isométrie sera toujours une base orthonormale. La réciproque est évidemment plus compliquée, elle peut en fait être démontrée comme conséquence de la propriété matricielle énoncée ci-dessous, mais à condition de ne pas démontrer celle-ci à l'aide de la propriété précédente, comme on va le faire dans quelques instants ! $\square$

Définition 5. Une matrice $A \in \mathcal{M}_n(\mathbb{R})$ est **orthogonale** si elle est la matrice représentative d'un endomorphisme orthogonal f dans la base canonique (ou dans toute base orthonormale) de $\mathbb{R}^n$. On note $\mathcal{O}_n(\mathbb{R})$ l'ensemble de toutes les matrices orthogonales.

Proposition 5. $A \in \mathcal{O}_n(\mathbb{R}) \Leftrightarrow {}^tAA = I$.

Démonstration. Supposons que ${}^tAA = I$, et notons f l'endomorphisme représenté par A dans une base $(e_1, e_2, \dots, e_n)$ orthonormale (la base canonique convient très bien). Dire que ${}^tAA = I$ signifie deux choses : pour chaque colonne de A , la somme des carrés des éléments de la colonne est égale à 1 (c'est le calcul qu'on fait pour obtenir $({}^tAA)_{ii}$), ce qui revient à dire que $f(e_i)$ est un vecteur de norme 1 ; et si on effectue le calcul $\sum_{k=1}^n a_{ik}a_{jk}$ pour $i \neq j$, on obtient 0, ce qui revient cette fois-ci à dire que $f(e_i)$ et $f(e_j)$ sont orthogonaux. Autrement dit, l'image de notre base orthonormale est une base orthonormale, donc f est orthogonal. La réciproque se fait exactement de la même façon. $\square$

Remarque 7. Attention dans la définition des matrices orthogonales à ne pas oublier qu'on doit se placer dans une base orthonormale. Dans une base quelconque, la matrice d'une application orthogonale peut ressembler à n'importe quoi (d'inversible tout de même puisque f est bijective). Au passage, notre dernière propriété confirme que A est une matrice inversible, et même que son inverse est sa transposée.

Exemple : La matrice $A = \frac{1}{9} \begin{pmatrix} 8 & -1 & -4 \\ -1 & 8 & -4 \\ 4 & 4 & 7 \end{pmatrix}$ est une matrice orthogonale. Pour le vérifier, on peut évidemment calculer tAA , mais on peut aussi utiliser le petit truc suivant : on vérifie que les colonnes « sont de norme 1 » et « orthogonales entre elles » comme expliqué dans la démonstration ci-dessus. Ici, $\frac{1}{9}\|(8, -1, 4)\| = \frac{1}{9}\sqrt{64 + 1 + 16} = 1$ et de même pour les deux autres colonnes ; et $(8, -1, 4) \cdot (-1, 8, 4) = -8 - 8 + 16 = 0$, et de même pour les deux autres produits scalaires de colonnes.

Proposition 6. Une matrice A est orthogonale si et seulement si ses vecteurs-colonnes forment une base orthonormale de E . De même pour les vecteurs-lignes.

Démonstration. On vient de voir que c'est une autre façon de dire exactement la même chose que dans la proposition précédente. Si ça marche pour les colonnes, ça marche pour les lignes, car A est orthogonale si et seulement si tA l'est (cela découle de l'égalité ${}^tAA = I$). $\square$

Proposition 7. La matrice de passage entre deux bases orthonormales est une matrice orthogonale.

Proposition 8. Si $A \in \mathcal{O}_n(\mathbb{R})$, alors $\det(A) = \pm 1$.

Démonstration. On sait que ${}^tAA = I$, et que $\det({}^tA)\det(A)$, donc $\det(A)^2 = \det(I) = 1$, la propriété en découle. $\square$

Définition 6. L'ensemble des matrices orthogonales de déterminant 1 est noté $\mathcal{SO}_n(\mathbb{R})$ (groupe spécial orthogonal) ou $\mathcal{O}_n^+(\mathbb{R})$, l'ensemble de celles de déterminant -1 est noté $\mathcal{O}_n^-(\mathbb{R})$ (non, n'insistez pas, il n'y a pas d'équivalent à la notation $\mathcal{SO}_n(\mathbb{R})$ dans ce cas-là).

Définition 7. Une isométrie est **directe** si sa matrice dans une base orthonormale appartient à $\mathcal{SO}_n(\mathbb{R})$, elle est **indirecte** sinon.

2 Isométries du plan.

Théorème 1. Toute matrice $A \in \mathcal{SO}_2(\mathbb{R})$ peut s'écrire sous la forme $A = \begin{pmatrix} \cos(\theta) & -\sin(\theta) \\ \sin(\theta) & \cos(\theta) \end{pmatrix}$, où $\theta \in \mathbb{R}$. L'application f ayant pour matrice A dans n'importe quelle base orthonormale est appelée **rotation d'angle θ** .

Démonstration. Soit $A = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \in \mathcal{SO}_2(\mathbb{R})$. On peut traduire l'égalité ${}^tAA = I$ par le système d'équations suivant :
$$\begin{cases} a^2 + b^2 = 1 \\ ac + bd = 0 \\ c^2 + d^2 = 1 \end{cases}$$
 (deux des équations obtenues sont identiques). On a de plus la condition $\det(A) = ad - bc = 1$. La dernière équation signifie que le point de coordonnées (c, d) dans le plan appartient au cercle trigonométrique (il est à une distance 1 de l'origine), ce qui permet de poser $c = \cos(\theta)$ et $d = \sin(\theta)$, pour un certain réel θ . De même, la première équation permet de poser $a = \cos(\alpha)$ et $b = \sin(\alpha)$. La deuxième condition devient alors $\cos(\alpha)\cos(\theta) + \sin(\alpha)\sin(\theta) = 0$, soit $\cos(\alpha - \theta) = 0$, et celle sur le déterminant donne de même $\sin(\theta - \alpha) = 1$. Autrement dit, on aura $\alpha = \theta + \frac{\pi}{2}[2\pi]$, et donc $\cos(\alpha) = \cos(\theta)$ et $\sin(\alpha) = -\sin(\theta)$, ce qui donne bien la matrice annoncée. $\square$

Remarque 8. L'application n'est évidemment pas appelée rotation par hasard : il s'agit bel et bien d'une rotation d'angle θ autour de l'origine. On vient en fait de prouver que les seules isométries directes dans $\mathbb{R}^2$ sont les rotations (on parle ici d'isométries vectorielles, qui doivent laisser fixe l'origine O , ce qui explique l'absence des translations). La matrice d'une rotation dans le plan est indépendante de la base orthonormale choisie. Bien évidemment, dans une base qui n'est pas orthonormale, la matrice peut changer.

Définition 8. Un **retournement** est une rotation plane d'angle $\theta = \pi$.

Remarque 9. C'est ce que vous avez appelé pendant des années une symétrie centrale, terme que nous n'utiliserons plus jamais même si l'application f est de fait dans ce cas une symétrie (on obtient l'identité si on la compose avec elle-même).

Théorème 2. Toute matrice $A \in \mathcal{O}_2^-(\mathbb{R})$ peut s'écrire sous la forme $A = \begin{pmatrix} \cos(\theta) & \sin(\theta) \\ \sin(\theta) & -\cos(\theta) \end{pmatrix}$. L'application u ayant pour matrice A dans une base orthonormale est une réflexion.

Démonstration. La preuve que la matrice peut se mettre sous cette forme est identique à celle vue pour les isométries directes, à un changement de signe près à un endroit, nous nous épargnerons les calculs. Il est facile de constater que u est une symétrie (nécessairement orthogonale) en calculant $A^2 = \begin{pmatrix} \cos^2(\theta) + \sin^2(\theta) & 0 \\ 0 & \sin^2(\theta) + \cos^2(\theta) \end{pmatrix} = I$. C'est nécessairement une symétrie par rapport à une droite vectorielle, donc une réflexion. En effet, si le sous-espace par rapport auquel on symétrise n'est pas une droite, il s'agit soit de $\{0\}$ et alors $u = -\text{id}$, soit de $\mathbb{R}^2$ et $u = \text{id}$. Dans les deux cas, u ne serait pas une isométrie indirecte. $\square$

Remarque 10. Dans le cas des réflexions, la matrice de u n'est pas du tout la même dans toutes les bases orthonormales. Par ailleurs, l'angle θ est beaucoup moins facile à interpréter géométriquement que dans le cas d'une rotation.

Remarque 11. Toute rotation dans le plan peut s'écrire comme composée de deux réflexions. En effet,
$$\begin{pmatrix} \cos(\theta) & -\sin(\theta) \\ \sin(\theta) & \cos(\theta) \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \times \begin{pmatrix} \cos(-\theta) & \sin(-\theta) \\ \sin(-\theta) & -\cos(-\theta) \end{pmatrix}.$$

3 Isométries de l'espace.

Définition 9. Tout élément de $\mathcal{SO}_3(\mathbb{R})$ est appelé **matrice de rotation** dans l'espace.

Théorème 3. Soit f une isométrie directe de l'espace, alors $F = \ker(f - \text{id})$ est de dimension 1. Si (v, w) est une base orthonormale du plan orthogonal à F , alors la matrice de f dans la base orthonormale $(v, w, v \wedge w)$ est $A = \begin{pmatrix} \cos(\theta) & -\sin(\theta) & 0 \\ \sin(\theta) & \cos(\theta) & 0 \\ 0 & 0 & 1 \end{pmatrix}$.

Définition 10. Une isométrie directe ayant pour matrice A dans une base orthonormale est appelée **rotation d'axe dirigé par $v \wedge w$ et d'angle θ** .

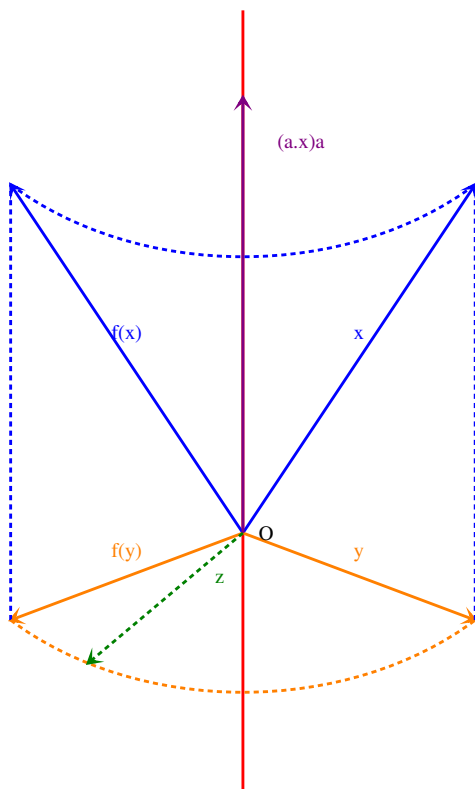
Remarque 12. Une rotation spatiale reste assez facile à visualiser : l'axe ne bouge pas, et le plan orthogonal à l'axe subit une rotation d'angle θ (autrement dit, on tourne d'un angle θ autour de l'axe). Voir le schéma plus bas. Notons quand même que la définition est un peu douteuse, car l'angle de la rotation dépend du vecteur choisi pour diriger l'axe. En effet, si on prend le vecteur opposé, l'angle va être également changé en son opposé.

Proposition 9. La composée de deux rotations dans l'espace est une rotation.

Démonstration. En effet, comme dans le plan, la composée reste une isométrie directe, donc une rotation. Mais pour le coup, ça n'a rien de géométriquement évident (et en particulier, l'angle de la rotation composée n'est pas évident à déterminer à partir de ceux des deux rotations). $\square$

Proposition 10. Soit f la rotation d'axe $F = \text{Vect}(a)$, avec a unitaire, et d'angle θ , alors $\forall u \in \mathbb{R}^3, f(u) = \cos(\theta)u + \sin(\theta)a \wedge u + (1 - \cos(\theta))(a.u)a$.

Remarque 13. Dans le cas particulier où u appartient au plan orthogonal à F , on trouve plus simplement $f(u) = \cos(\theta)u + \sin(\theta)(a \wedge u)$, on peut donc retrouver facilement l'angle de la rotation, quasiment comme dans le cas d'une rotation plane, à l'aide des relations $u.f(u) = \cos(\theta)$ et $\det(u, f(u), a) = \sin(\theta)$ (en choisissant un vecteur u unitaire). Sur la figure ci-dessous, l'axe de la rotation est en rouge, le plan orthogonal en marron, le vecteur noté z est $y \wedge a$.



Démonstration. La projection de u sur l'axe de la rotation est simplement donnée par $(a.u)a$. Posons $y = u - (a.u)a$, alors $(y, y \wedge a)$ est une base orthogonale directe du plan orthogonal à F constituée de deux vecteurs de même norme. De plus, $y \wedge a = (u - (a.u)a) \wedge a = u \wedge a$ puisque $a \wedge a = 0$. On en déduit que $f(y) = \cos(\theta)y - \sin(\theta)(u \wedge a)$, donc $f(u) = f(y + (a.u)a) = \cos(\theta)y + \sin(\theta)(a \wedge u) + (a.u)a$ (puisque a est laissé fixe par la rotation). Autrement dit, $f(u) = \cos(\theta)u - \cos(\theta)(a.u)a + \sin(\theta)(a \wedge u) + (a.u)a$, ce qui est bien la formule donnée. $\square$

Exemple : Soit f l'endomorphisme de $\mathbb{R}^3$ représenté dans la base canonique par la matrice $M = \frac{1}{3} \begin{pmatrix} 2 & 2 & 1 \\ -2 & 1 & 2 \\ 1 & -2 & 2 \end{pmatrix}$. La matrice est orthogonale puisque $\frac{1}{3}\|(2, 2, 1)\| = \frac{1}{3}\sqrt{4+4+1} = 1$; $\frac{1}{3}\|(-2, 1, 2)\| = 1$; $\frac{1}{3}\|(1, -2, 2)\| = 1$; $(2, 2, 1) \cdot (-2, 1, 2) = -4+2+2 = 0$; $(2, 2, 1) \cdot (1, -2, 2) = 0$ et $(-2, 1, 2) \cdot (1, -2, 2) = 0$. Il s'agit donc de la matrice d'une isométrie.

On calcule ensuite le déterminant pour déterminer si l'isométrie est directe (dans le calcul qui suit, on effectue l'opération $L_3 \leftarrow L_3 - L_2$ puis on développe par rapport à la dernière ligne) :

$$\begin{vmatrix} 2 & 2 & 1 \\ -2 & 1 & 2 \\ 1 & -2 & 2 \end{vmatrix} = \begin{vmatrix} 2 & 2 & 1 \\ -2 & 1 & 2 \\ 3 & -3 & 0 \end{vmatrix} = 3 \times \begin{vmatrix} 2 & 1 \\ 1 & 2 \end{vmatrix} + 3 \times \begin{vmatrix} 2 & 1 \\ -2 & 2 \end{vmatrix} = 3 \times 3 + 3 \times 6 = 27, \text{ donc}$$

$\det(A) = \frac{1}{3^3} \times 27 = 1$. Il s'agit d'une isométrie directe, donc d'une rotation.

On cherche ensuite l'axe de la rotation, en déterminant simplement $\ker(f - \text{id})$. Quitte à tout multiplier par 3, on doit résoudre le système

$$\begin{cases} 2x + 2y + z = 3x \\ -2x + y + 2z = 3y \\ x - 2y + 2z = 3z \end{cases} \text{ . La deuxième équation}$$

donne $z = x + y$, et la dernière donne $2y = x - z$, soit $2y = -y$. Manifestement, cela implique $y = 0$, puis $z = x$. La première équation donne alors $3x = 3x$, donc $\ker(f - \text{id}) = \text{Vect}((1, 0, 1))$, qui est bien une droite F .

On choisit maintenant un vecteur unitaire v orthogonal à F . Ce n'est pas bien compliqué ici, il suffit de prendre $v = (0, 1, 0)$. On peut alors calculer, en notant θ l'angle de la rotation, $\cos(\theta) = v \cdot f(v)$. Puisque $f(0, 1, 0) = \frac{1}{3}(2, 1, -2)$, on trouve $\cos(\theta) = \frac{1}{3}$. En posant $a = \frac{1}{\sqrt{2}}(1, 0, 1)$ (vecteur directeur

unitaire de l'axe), on peut ensuite calculer $\det(v, f(v), a) = \frac{1}{3\sqrt{2}} \begin{vmatrix} 0 & 1 & 0 \\ 2 & 1 & -2 \\ 1 & 0 & 1 \end{vmatrix} = -\frac{1}{3\sqrt{2}} \begin{vmatrix} 2 & -2 \\ 1 & 1 \end{vmatrix} = -\frac{2\sqrt{2}}{3}$. On en déduit que $\sin(\theta) = -\frac{2\sqrt{2}}{3}$, donc $\theta = -\arccos\left(\frac{1}{3}\right)$. Seul le signe de $\sin(\theta)$ était nécessaire pour conclure, mais l'avantage de l'avoir calculé explicitement est de pouvoir vérifier que $\cos^2(\theta) + \sin^2(\theta) = \frac{1}{9} + \frac{8}{9} = 1$. On peut en tout cas désormais affirmer que f est la rotation d'axe $\text{Vect}((1, 0, 1))$ et d'angle $\arccos\left(\frac{1}{3}\right)$.

Remarque 14. Un autre moyen de vérifier la cohérence des calculs est d'utiliser la trace de la matrice. En effet, celle-ci est invariante par changement de repère, donc est égale à $2\cos(\theta) + 1$ dans n'importe quel repère. Ici, on pouvait donc calculer dès le départ $\text{Tr}(A) = \frac{5}{3} = 2\cos(\theta) + 1$, et en déduire que $\cos(\theta) = \frac{1}{3}$.

Théorème 4. (Hors-programme). Si f est une réflexion dans l'espace, elle a pour matrice dans une base orthonormale bien choisie $A = \begin{pmatrix} \cos(\theta) & \sin(\theta) & 0 \\ \sin(\theta) & -\cos(\theta) & 0 \\ 0 & 0 & 1 \end{pmatrix}$. De plus, toute isométrie indirecte de l'espace est la composée d'une rotation (éventuellement égale à id) et d'une réflexion.

Remarque 15. On peut toujours écrire dans l'espace une rotation comme composée de deux réflexions (c'est exactement le même calcul que dans le plan, en ajoutant un 1 en bas à droite de la matrice et des 0 ailleurs sur la dernière ligne et la dernière colonne). Il existe donc trois types d'isométries vectorielles dans l'espace :

- les réflexions, qui laissent tout un plan fixe et sont indirectes.
- les composées de deux réflexions, qui sont des rotations, laissent une droite fixe et sont directes.
- les composées de trois réflexions ne laissent que le vecteur nul invariant et sont indirectes.

Ainsi, dans l'espace, $-\text{id}$ est une isométrie indirecte qui est une composée de trois réflexions (par exemple par rapport aux trois axes du repère canonique). Ce n'est absolument pas la rotation d'angle π autour de l'origine, cette application n'étant pas une rotation avec la définition que nous avons prise. Pour les plus curieux, le théorème précédent se généralise en dimension n , où les isométries peuvent toujours être écrites comme composées d'au plus n réflexions par rapport à des hyperplans (et laissant stable un sous-espace vectoriel de dimension $n - k$, où k est le nombre de réflexions composées pour obtenir l'isométrie).

4 Espaces vectoriels préhilbertiens.

4.1 Produits scalaires dans un espace vectoriel quelconque.

Comme annoncé plus haut, on peut en fait étendre la notion de produit scalaire à n'importe quel espace vectoriel réel (et même complexe), y compris de dimension infinie. Il n'est bien sûr plus question de définir directement le produit scalaire à l'aide d'une jolie formule comme on l'a fait dans $\mathbb{R}^n$, mais de le décrire par ses propriétés théoriques :

Définition 11. Un **produit scalaire** dans un espace vectoriel E est une application $\varphi : E \times E \rightarrow \mathbb{R}$ vérifiant les trois propriétés fondamentales suivantes :

- linéarité à gauche : $\forall (u, v, w) \in E^2, \forall (\lambda, \mu) \in \mathbb{R}^2, \varphi(\lambda u + \mu v, w) = \lambda \varphi(u, w) + \mu \varphi(v, w)$.
- symétrie : $\forall (u, v) \in E^2, \varphi(u, v) = \varphi(v, u)$.
- définie positivité : $\forall u \in E, \varphi(u, u) \geq 0$, et $\varphi(u, u) = 0 \Leftrightarrow u = 0$.

Remarque 16. La symétrie et la linéarité à gauche permettent de prouver que φ est aussi linéaire à droite, et donc bilinéaire. Il peut bien sûr sembler curieux de noter le produit scalaire comme une sorte de fonction à deux variables, mais on procède en fait de la façon suivante : on définit l'application qui va jouer le rôle de produit scalaire, on prouve qu'elle vérifie les propriétés fondamentales, et ensuite on change sa notation pour la noter comme un produit scalaire habituel.

Exemple : plaçons-nous dans l'espace vectoriel $E = \mathbb{R}_3[X]$ des polynômes de degré inférieur ou égal à 3. On souhaite définir un produit scalaire sur cet espace vectoriel. Pourquoi, me direz-vous ? Tout simplement pour pouvoir faire de la géométrie dans E , et notamment pouvoir calculer des **distances** entre éléments de E . La notion de distance entre deux polynômes peut tout à fait s'interpréter intuitivement : on peut par exemple avoir envie de dire que deux polynômes sont « proches » s'ils ont des coefficients proches, ou bien s'ils prennent des valeurs proches à certains endroits, ou encore si leurs courbes ne s'éloignent pas trop l'une de l'autre (tout ça restant bien sûr assez flou pour l'instant). Eh bien, ces différentes notions de proximité entre polynômes peuvent en fait être matérialisées par différents choix de produits scalaires sur l'espace E , qui amèneront à des définitions différentes des distances.

On peut ainsi prouver que l'application φ définie par $\varphi(P, Q) = \sum_{i=0}^3 a_i b_i$ (où a_i et b_i désignent les coefficients des polynômes P et Q) est un produit scalaire. C'est un exercice facile, et c'est en fait tout à fait logique puisque cette définition « correspond » à celle du produit scalaire usuel sur $\mathbb{R}^4$ si on considère que les coefficients d'un polynôme peuvent être identifiés à des coordonnées de vecteurs dans $\mathbb{R}^4$. À partir de ce produit scalaire, la distance entre deux polynômes P et Q sera définie par

la formule $\|P - Q\| = \sqrt{\sum_{i=0}^3 (a_i - b_i)^2}$. On retrouve ici une notion de distance où des polynômes sont à une faible distance l'un de l'autre si leurs coefficients sont proches. Par exemple, la distance entre $P = 1$ et $Q = 2X^2 - X$ sera égale à $\sqrt{1+1+4} = \sqrt{6}$, mais la distance entre $P = X$ et $Q = 2X + X^2$ sera égale à $\sqrt{1+1} = \sqrt{2}$ et donc plus petite que la précédente (ce qui devrait vous paraître vaguement logique). Pour ce produit scalaire, la base canonique $(1, X, X^2, X^3)$ est une base orthonormale.

Définissons maintenant sur ce même espace vectoriel un deuxième produit scalaire : on pose $\psi(P, Q) = \int_0^1 P(t)Q(t) dt$. Je vous laisse vérifier que cette définition est bien celle d'un produit scalaire, à partir duquel on peut donc définir une notion de distance qui n'est plus du tout la même que la précédente.

Cette fois-ci, la distance entre deux polynômes P et Q sera égale à $\|P - Q\| = \sqrt{\int_0^1 (P(t) - Q(t))^2 dt}$.

Déjà, la base canonique $(1, X, X^2, X^3)$ n'est plus du tout une base orthonormale : ses vecteurs ne sont pas tout normés (on a bien $\|1\| = \sqrt{\psi(1, 1)} = \sqrt{\int_0^1 1 dt} = 1$, mais par exemple $\|X^2\| =$

$\sqrt{\psi(X^2, X^2)} = \sqrt{\int_0^1 t^4 dt} = \frac{1}{\sqrt{5}}$), et ils ne sont même plus orthogonaux, puisque par exemple

$\psi(1, X) = \int_0^1 t dt = \frac{1}{2}$. Vous apprendrez l'an prochain comment déterminer des bases orthonormales pour de tels produits scalaires (il en existe toujours). En attendant, calculons quelques distances

pour voir ce que ça donne. Par exemple, la distance entre $P = 1$ et $Q = 2X^2 - X$ sera égale à $\sqrt{\int_0^1 (2t^2 - t - 1)^2 dt} = \frac{2}{\sqrt{5}}$, et la distance entre $P = X$ et $Q = 2X + X^2$ vaut $\sqrt{\int_0^1 (t^2 + t)^2 dt} = \sqrt{\frac{31}{30}}$. Cette fois-ci c'est la deuxième distance qui est plus grande, ce qui peut se traduire par le fait que les courbes des polynômes P et Q sont « plus proches l'une de l'autre sur l'intervalle $[0, 1]$ » dans le deuxième cas. Bien entendu, ce genre de calcul peut sembler sans intérêt, mais toute la puissance de la géométrie euclidienne appliquée aux espaces vectoriels réside dans le fait qu'on peut réellement exploiter toutes les notions et tous les calculs classiques de géométrie dans ce cadre. Un exemple (pour lequel vous saurez très bien faire les calculs l'an prochain) : on cherche une courbe qui approche le plus possible celle d'une fonction (allez, disons celle de l'exponentielle pour fixer les idées) sur l'intervalle $[0, 1]$, en imposant que la fonction recherchée soit un polynôme de degré 3. Les développements limités fournissent une réponse si on cherche la meilleure approximation possible en un point donné, mais pas sur un intervalle. En fait, le produit scalaire dont nous venons de parler fournit une façon de faire ce calcul : il « suffit » d'effectuer une projection orthogonale de la fonction exponentielle sur le sous-espace vectoriel $\mathbb{R}_3[X]$, et ce calcul est réellement un calcul de projection orthogonale tout à fait classique (au détail près que la notion d'orthogonalité utilisée est bien entendu celle qui découle de notre produit scalaire).

Définition 12. Un **espace préhilbertien** est un espace vectoriel E muni d'un produit scalaire.

4.2 Application aux séries de Fourier.

Un domaine qui fournit une application inattendue des notions de géométrie rapidement développées ci-dessus est celui des séries de Fourier. Comme vous devez déjà le savoir, le principe de la décomposition en séries de Fourier consiste à partir d'une fonction périodique et à essayer de l'écrire comme une somme (techniquement infinie, d'où le terme de **séries** de Fourier) de fonctions cosinus et sinus, dont les fréquences sont toutes des multiples entiers de celle de la fonction dont on est parti. Tout cela semble avoir plus de lien avec l'analyse qu'avec l'algèbre ou la géométrie, et pourtant, le principe même de la décomposition en séries de Fourier est, comme au paragraphe précédent, une projection orthogonale pour un produit scalaire adapté.

Proposition 11. On note E l'espace vectoriel constitué de toutes les fonctions définies et continues sur $\mathbb{R}$, et 2π -périodiques. L'application $\varphi : E \times E \rightarrow \mathbb{R}$ définie par $\varphi(f, g) = \frac{1}{2\pi} \int_0^{2\pi} f(t)g(t) dt$ définit un produit scalaire sur l'espace E .

On remarquera une énorme similitude entre la définition de ce produit scalaire et de l'un de ceux étudiés au paragraphe précédent.

Proposition 12. Les fonctions $u_n : t \mapsto \cos(nt)$ et $v_n : t \mapsto \sin(nt)$, pour $n \geq 1$, sont constituées de fonctions de norme 1 pour le produit scalaire précédemment défini. De plus, la famille $\mathcal{F}_n = \text{Vect}(1, u_1, v_1, \dots, u_n, v_n)$ est une famille libre et orthogonale dans E quel que soit $n \geq 1$.

Définition 13. Soit $f \in E$, la projection orthogonale de f sur le sous-espace vectoriel de E engendré par la famille $\mathcal{F}_n$ est appelée **somme partielle d'ordre n** de la série de Fourier associée à la fonction f .

Il s'agit en fait d'approcher au mieux (au sens de la distance induite par notre produit scalaire) la fonction f par une fonction f_n de la forme $f_n(x) = a_0 + \sum_{k=1}^n a_k \cos(kx) + \sum_{k=1}^n b_k \sin(kx)$ (ce qu'on appelle aussi des polynômes trigonométriques). Il n'est bien sûr pas nécessaire de comprendre cette histoire de projection orthogonale pour calculer en pratique des séries de Fourier, il suffit d'avoir à disposition des formules permettant de calculer les coefficients a_n et b_n , ce qui se fait très bien en calculant des intégrales bien choisies.

Théorème 5. Les sommes partielles de la série de Fourier associée à f convergent simplement vers la fonction f , c'est-à-dire que :

$$\forall x \in \mathbb{R}, f(x) = a_0 + \sum_{k=1}^{+\infty} a_k \cos(kx) + \sum_{k=1}^{+\infty} b_k \sin(kx).$$

Il existe diverses notions de convergence pour les suites de fonctions, que vous n'étudierez d'ailleurs même pas l'an prochain. Ici, le résultat du théorème signifie en gros la chose suivante : si on considère la famille de fonctions **infinie** (et qui n'est donc plus techniquement une famille de fonctions) $(1, u_1, v_1, \dots, u_n, v_n, \dots)$, cette famille se comporte en quelque sorte comme une base de l'espace vectoriel E (qui est bien sûr de dimension infinie). En effet, on peut décomposer toute fonction appartenant à E (et ce de manière unique) comme une « combinaison linéaire infinie » des fonctions de la famille. Attention tout de même, ça marche moins bien si on ne suppose pas les fonctions continues. Une telle famille porte dans les espaces vectoriels préhilbertiens le nom de **base hilbertienne**. Cette définition me permet d'ailleurs de conclure le cours de cette année sur l'une des plus inénarrables blagues de matheux qui soient (personnellement je l'adore) : monsieur et madame Bertienne ont un fils, comment s'appelle-t-il ? Basile ! Le simple fait de pouvoir comprendre cette blague hilarante vous placera pour le restant de vos jours dans une classe supérieure à laquelle peu d'être humains (même pas Battista, c'est dire) peuvent prétendre.