

Chapitre 18 : Applications linéaires

PTSI B Lycée Eiffel

20 mai 2021

J'ai simplement pensé à l'idée d'une projection, d'une quatrième dimension invisible, autrement dit que tout objet de trois dimensions, que nous voyons froidement, est une projection d'une chose à quatre dimensions, que nous ne connaissons pas.

MARCEL DUCHAMP

*Un mathématicien et un ingénieur assistent à une conférence sur les processus physiques intervenant dans les espaces de dimension 9. Le mathématicien est assis et apprécie beaucoup la conférence, pendant que l'ingénieur fronce les sourcils et semble complètement embrouillé. À la fin, l'ingénieur demande au matheux :
« Comment fais-tu pour comprendre tout cela ? » « C'est simple ! D'abord tu visualises le processus en dimension n , et ensuite il suffit de prendre $n = 9$. »*

Ce deuxième chapitre d'algèbre linéaire sera un simple complément du premier, permettant de présenter une notion complètement fondamentale, celle d'application linéaire, qui va éclairer d'un jour nouveau tous les termes vus depuis le début de l'année et faisant intervenir ce fameux mot « linéaire ». En gros, les applications linéaires sont des applications (des fonctions, quoi) « naturelles » dans les espaces vectoriels, qui apparaissent dans tous les domaines des mathématiques, et pour lesquels une étude tout à fait générale et théorique est possible, ce qui permet d'appréhender un peu mieux la puissance de l'algèbre linéaire pour résoudre des problèmes de maths très divers. Ce petit chapitre sera essentiellement constitué de vocabulaire, les quelques calculs à savoir faire se résumant à nouveau la plupart du temps à des résolutions de petits systèmes (linéaires, cela va de soit !).

Objectifs du chapitre :

- maîtriser tout le vocabulaire introduit dans ce chapitre.
- comprendre l'intérêt du théorème du rang.
- savoir manipuler les projections et symétries vectorielles.

1 Vocabulaire.

1.1 Morphismes.

Définition 1. Soient E et F deux espaces vectoriels, une **application linéaire** de E dans F est une application $f : E \rightarrow F$ vérifiant les conditions suivantes :

- $\forall (u, v) \in E^2, f(u + v) = f(u) + f(v)$

- $\forall \lambda \in \mathbb{R}, \forall u \in E, f(\lambda u) = \lambda f(u)$

Remarque 1. Autrement dit, une application linéaire est une application compatible avec les deux opérations définissant la structure d'espace vectoriel. Pour les plus curieux, voir le développement de ce principe dans le paragraphe suivant (celui qui est complètement hors-programme). Une application linéaire « transforme les sommes en sommes » et les produits par des réels en produits par des réels, même si les espaces vectoriels E et F sont de nature complètement différente, et même si les opérations de somme dans E et dans F ne sont techniquement pas « les mêmes ».

Remarque 2. Si $f : E \rightarrow F$ est une application linéaire, on a nécessairement $f(0_E) = 0_F$ (exceptionnellement, j'ai précisé ici dans quel espace vectoriel se situait le vecteur nul car c'est très important). En effet, $f(0 + 0) = f(0) + f(0)$, donc $f(0) = 2f(0)$, ce qui implique manifestement $f(0) = 0$.

Proposition 1. Une application $f : E \rightarrow F$ est linéaire si et seulement si $\forall (\lambda, \mu) \in \mathbb{R}^2, \forall (u, v) \in E^2, f(\lambda u + \mu v) = \lambda f(u) + \mu f(v)$.

Démonstration. Si f vérifie les conditions de la définition, alors $f(\lambda u + \mu v) = f(\lambda u) + f(\mu v) = \lambda f(u) + \mu f(v)$ en utilisant successivement les deux propriétés. Réciproquement, en prenant $\lambda = \mu = 1$, on retrouve la première condition ; et en prenant $v = 0$, on retrouve la deuxième (en exploitant la remarque faite juste avant l'énoncé de la propriété). $\square$

Remarque 3. Autrement dit, une application est linéaire si et seulement si elle est compatible avec les combinaisons linéaires. On a d'ailleurs plus généralement, pour une application linéaire, $f\left(\sum_{i=1}^n \lambda_i e_i\right) = \sum_{i=1}^n \lambda_i f(e_i)$. C'est d'ailleurs cette propriété (avec seulement deux vecteurs, pas besoin de se compliquer la vie) qu'on vérifiera lorsqu'il faudra prouver dans les exercices qu'une application est linéaire : on vérifie bien que l'application est définie sur E et à valeurs dans F , et on calcule $f(\lambda u + \mu v)$ en essayant de prouver que c'est égal à $\lambda f(u) + \mu f(v)$. Attention bien sûr à faire attention au fait qu'ici u et v sont des vecteurs de l'espace E , pas de simples coordonnées.

Exemple : Dans l'espace vectoriel $E = \mathcal{M}_2(\mathbb{R})$, on pose $A = \begin{pmatrix} 2 & 1 \\ -1 & 1 \end{pmatrix}$, et on définit l'application $f : E \rightarrow E$ par $f(M) = AM$. Pour prouver que cette application est linéaire, même pas besoin de calculer les coordonnées explicites de $f(M)$, on se contente d'écrire que, si M et N sont deux matrices quelconques de E , et $(\lambda, \mu) \in \mathbb{R}^2$, alors $f(\lambda M + \mu N) = A(\lambda M + \mu N) = \lambda AM + \mu AN = \lambda f(M) + \mu f(N)$. On a simplement utilisé les règles de base du calcul matriciel (si on veut faire savant, on a en fait simplement utilisé la linéarité du produit matriciel à gauche, ce qui signifie en fait exactement que l'application f est linéaire).

Exemples : On pourrait penser que les conditions définissant une application linéaire sont restrictives, et qu'il va donc y avoir « peu » d'applications linéaires. C'est en partie vrai : si on se restreint à des espaces très simples (et de très petite dimension), il y a effectivement très peu d'applications qui seront linéaires. Mais dans des espaces plus gros, il y en aura largement assez pour qu'on s'y intéresse de très près, d'autant plus que l'ensemble de ces applications linéaires va être doté d'une structure forte, justement à cause des conditions très restrictives qui les définissent.

- Si on considère toutes les applications $f : \mathbb{R} \rightarrow \mathbb{R}$ (autrement dit toutes les fonctions réelles définies sur $\mathbb{R}$ tout entier), les seules à être linéaires sont celles de la forme $x \mapsto ax$, avec $a \in \mathbb{R}$. Ce qui tombe d'ailleurs très bien puisque ce sont justement les fonctions qu'on appelle traditionnellement fonctions linéaires !

- L'application $f : \mathbb{R}^2 \rightarrow \mathbb{R}^3$ définie par $f(x, y) = (2x - 3y, 4x + y, -x + 2y)$ est une application linéaire. On le vérifie comme ci-dessus en développant $f(\lambda(x, y, z) + \mu(x', y', z'))$ et en constatant que c'est égal à $\lambda f(x, y, z) + \mu f(x', y', z')$ (ici, il faut vraiment faire le calcul avec les coordonnées, c'est un peu pénible mais essentiellement trivial).
- L'application $f : \mathbb{R}^2 \rightarrow \mathbb{R}^3$ définie par $f(x, y) = (2x - 3, 4 + y, -x + 2y)$ n'est pas une application linéaire (on peut constater par exemple qu'en général $f(2x, 2y) \neq 2f(x, y)$). En fait, en réfléchissant un peu, on constate assez facilement que les seules applications linéaires de $\mathbb{R}^2$ dans $\mathbb{R}^3$ sont celles pour lesquelles les coordonnées de $f(x, y)$ sont obtenues en faisant des combinaisons linéaires de x et de y . Ce constat s'étend en fait aux applications linéaires entre n'importe quels espaces de dimension finie. Il faut et il suffit que chaque coordonnée de $f(u)$ soit une combinaison linéaire des coordonnées de u pour que f soit une application linéaire.
- L'application $f : \mathcal{M}_3(\mathbb{R}) \rightarrow \mathcal{M}_3(\mathbb{R})$ définie par $f(M) = A^2 M$ est une application linéaire, quelle que soit la matrice $A \in \mathcal{M}_3(\mathbb{R})$ (même principe que dans l'exemple ci-dessus).
- L'application $f : \mathcal{M}_3(\mathbb{R}) \rightarrow \mathcal{M}_3(\mathbb{R})$ définie par $f(M) = M^2$ n'est pas une application linéaire (en général, $(M + N)^2 \neq M^2 + N^2$, ce qu'on peut résumer en disant justement que l'élévation au carré « n'est pas linéaire »).
- Soit E l'ensemble des suites réelles. L'application $f : E \rightarrow \mathbb{R}^3$ définie par $f(u_n) = (u_0, u_8, u_{35})$ est une application linéaire.
- Soit E l'ensemble des fonctions C^∞ de $\mathbb{R}$ dans $\mathbb{R}$. L'application $f : E \rightarrow E$ définie par $f(g) = g'$ est une application linéaire (dont la variable est une fonction!). J'ai d'ailleurs déjà du dire au moins dix fois depuis le début de l'année la phrase « La dérivation est linéaire », ce qui signifie exactement cela.
- Soit E l'ensemble des fonctions continues sur l'intervalle $[0; 1]$. L'application $f : E \rightarrow \mathbb{R}$ définie par $f(g) = \int_0^1 g(t) dt$ est une application linéaire. Là encore, on a parlé de linéarité de l'intégrale largement avant d'avoir une définition rigoureuse de ce qu'est une application linéaire.
- L'application $f : \mathbb{R}_n[X] \rightarrow \mathbb{R}_n[x]$ définie par $f(P) = 2X^2 P'' - X P'$ est une application linéaire (ici, il faut bien faire attention quand on vérifie la linéarité à vérifier que $f(\lambda P + \mu Q) = \lambda f(P) + \mu f(Q)$, il n'y absolument aucune raison de mettre des coefficients λ et μ sur l'indéterminée X à l'intérieur du polynôme, ce qui fait que le X^2 qui est simplement en facteur de P'' ne perturbe pas du tout la linéarité de l'application).

Définition 2. Une application linéaire $f : E \rightarrow F$ est aussi appelée **morphisme** de E dans F . On note $\mathcal{L}(E, F)$ l'ensemble de toutes les applications linéaires de E dans F .

Une application linéaire $f : E \rightarrow E$ est appelée **endomorphisme** de l'espace vectoriel E . On note plus simplement $\mathcal{L}(E)$ l'ensemble des endomorphismes de E .

Une application linéaire bijective est appelée **isomorphisme**.

Un endomorphisme bijectif est appelé **automorphisme**. L'ensemble des automorphismes d'un espace vectoriel E est noté $GL(E)$.

Remarque 4. Tout ce vocabulaire est assez logique si on a quelques rudiments de grec (mais oui). Le terme morphisme découle du grec $\mu\omicron\rho\phi\eta$ qui signifie forme. Un morphisme est donc une application qui « respecte la forme » de l'ensemble, c'est-à-dire ici qui respecte la structure d'espace vectoriel des deux ensembles (au départ et à l'arrivée). Il existe en fait des morphismes pour tous types de structures algébriques sur les ensembles, qui seront définis en lien avec la structure. On devrait donc parler non pas simplement de morphismes dans ce chapitre, mais bien de « morphismes d'espaces vectoriels » puisqu'il existe aussi des morphismes de groupes, des morphismes d'algèbres et autres joyeusetés que vous ne devez absolument pas maîtriser (cf plus bas pour les curieux).

Les préfixes ajoutés devant le terme morphisme sont tout autant dérivés du grec. Le préfixe endo signifie intérieur, un endomorphisme est donc un morphisme où on reste dans le même espace vectoriel (tout comme un endogame est quelqu'un qui se marie avec une personne de la même catégorie

(typiquement socio-professionnelle) que lui). Vous devez déjà avoir croisé le préfixe iso qui signifie « même », surtout si vous avez déjà fait de la thermodynamique avec M.Raimi (isobares, isochores et autres trucs sympathiques ; sinon le terme isotherme est d'un emploi nettement plus fréquent). Ici, un isomorphisme est une application qui permet de comprendre que les espaces vectoriels E et F ont la « même » structure. Bien sûr les ensembles ne sont pas les mêmes, et les opérations non plus, mais le fait d'avoir un isomorphisme entre E et F permet de les voir comme des espaces qui « se ressemblent » énormément. Ainsi, ils auront forcément la même dimension (s'ils sont de dimension finie), l'image de n'importe quelle base de E par l'isomorphisme f sera une base de F , etc.

Proposition 2. Si E et F sont deux espaces vectoriels réels, c'est aussi le cas de $\mathcal{L}(E, F)$.

Démonstration. La somme de deux applications linéaires est linéaire, l'élément neutre étant l'application nulle, et l'opposé d'une application linéaire étant toujours défini. De plus, les produits par des constantes d'applications linéaires sont linéaires, et les relations de distributivité sont immédiates. Bref, tout est trivial. $\square$

Remarque 5. C'est cette structure d'espace vectoriel sur l'ensemble des applications linéaires de E dans F (attention, il s'agit d'un espace complètement différent de E et de F) qui nous permettra dans le dernier chapitre d'algèbre linéaire que nous ferons cette année de représenter les applications linéaires entre espaces vectoriels de dimension finie par des matrices, et (encore mieux !) d'utiliser toute la puissance du calcul matriciel pour étudier plus efficacement les applications linéaires.

1.2 Compléments hors-programme.

Ce paragraphe n'est à lire que si vous voulez comprendre un peu mieux la notion de morphisme et surtout d'isomorphisme qui est omniprésente en algèbre (et pas seulement en algèbre linéaire puisque, comme expliqué plus haut, il n'existe pas que des morphismes entre espaces vectoriels). Comme j'ai essayé de vous le faire comprendre plus haut, l'idée derrière la notion de morphisme est de créer des applications qui préservent la structure (quelle qu'elle soit), et la notion d'isomorphisme permet en fait d'identifier deux ensembles ayant une structure identique. On peut même aller plus loin en essayant de comprendre quelles sont les différentes façons de structurer un ensemble « à partir de rien ». Pour cela il est plus facile de travailler avec des ensembles « petits » (ici on prendra des ensembles finis à peu d'éléments) et une structure plus basique que celle d'espace vectoriel : celle de **groupe**.

C'est en fait assez simple : définir une structure de groupe sur un ensemble E consiste simplement à créer une opération dans l'ensemble E (qu'on notera ici $\star$ car il ne s'agit pas forcément d'une addition) qui vérifie les deux propriétés suivantes :

- il existe un élément de E qui joue le rôle d'élément neutre pour l'opération $\star$, c'est-à-dire un élément $e \in E$ tel que, $\forall x \in E$, $e \star x = x$.
- tout élément de E admet un **inverse** pour l'opération $\star$, c'est-à-dire que, si $x \in E$, il existe $y \in E$ tel que $x \star y = e$, où e est l'élément neutre de ce qu'on appellera désormais le groupe E .

On ajoutera en plus le fait que l'opération $\star$ doit être associative et commutative (et encore, la commutativité n'est même pas obligatoire), et c'est tout. La question est maintenant, si on prend un ensemble fini contenant un certain nombre d'éléments, de savoir le nombre d'opérations différentes qu'on peut créer (et donc le nombre de structures de groupe différentes qu'on peut mettre sur notre ensemble) en respectant ces conditions. Le tout **à isomorphisme près**, ce qui revient à dire que, si on définit une nouvelle opération qui est en fait la même qu'une opération précédemment définie à permutation des éléments près, ça ne compte pas. Donnons des exemples concrets pour éclairer tout ça. Pour représenter une opération sur un ensemble fini, le plus simple est de présenter sa **table**, similaire à une bête table de multiplication, où on met en lignes et en colonnes les éléments de l'ensemble et on indique dans les cases du tableau le résultat de l'opération associant deux éléments de l'ensemble. Par exemple :

	e	a	b	c
e	e	a	b	c
a	a	e	c	b
b	b	c	e	a
c	c	b	a	e

On dispose ici d'un ensemble contenant quatre éléments notés e (qui jouera le rôle d'élément neutre), a , b et c . Et on a décidé que l'opération $\star$ vérifierait que $a \star a = e$, $a \star b = c$ etc (on lit les résultats dans la table ci-dessus). Je vous laisse constater si vous avez du temps à perdre que l'opération ainsi définie est bien associative et commutative, que e en est un élément neutre et que chaque élément admet un inverse (en l'occurrence, ici, chaque élément est son propre inverse). La question qu'on se pose maintenant est : aurait-on pu décider autrement et créer d'autres opérations qui fonctionnent ? Si on décide par exemple de permuter les rôles des quatre éléments (on décide que a devient l'élément neutre, donc $a \star b = b$, etc, et par contre $e \star c = b$ et ainsi de suite), la structure sera la même, ça ne compte pas (on peut prouver qu'il existe un isomorphisme de groupe entre notre ensemble muni de la première opération et le même ensemble muni de la deuxième opération). En fait il existe bel et bien une deuxième solution qui **ne peut pas** être obtenue à partir de la première par simple permutation des variables :

	e	a	b	c
e	e	a	b	c
a	a	b	c	e
b	b	c	e	a
c	c	e	a	b

Il ne s'agit pas de la même structure car ici, contrairement au tableau précédent, il existe des éléments qui ne sont pas leur propre inverse (a et c sont inverses l'un de l'autre, e et b sont leur propre inverse). On peut prouver (mais ce n'est pas évident !) qu'il n'y a en fait que deux structures de groupes différentes sur un ensemble contenant quatre éléments (les deux que nous venons de citer). Si notre ensemble contient un nombre d'éléments qui est un nombre **premier**, c'est encore pire, il n'y a qu'une seule façon de structurer l'ensemble pour en faire un groupe ! D'ailleurs, vous la connaissez déjà, cette structure, sans même le savoir : c'est la structure de l'ensemble $\mathbb{U}_n$ des racines n -èmes de l'unité, muni de l'opération de multiplication (je vous laisse y réfléchir). Pour un nombre d'éléments non premier, il existe en général plusieurs structures possibles, et bien sûr il peut y en avoir beaucoup si le nombre d'éléments est élevé, et encore plus si on accepte les opérations non commutatives. On connaît de toute façon bien les structures de groupes sur tous les ensembles finis mais les démonstrations de ces résultats sont assez monstrueuses (elles prennent plusieurs **milliers** de pages, les plus curieux iront par exemple consulter la page Wikipédia « Liste des groupes finis simples », mais n'y comprendront probablement pas grand chose). Plus accessible, la « Liste des petits groupes » toujours disponible sur Wikipédia vous donne le nombre de structures possibles pour des nombres d'éléments inférieurs ou égaux à 20. On a par exemple pas moins de 14 opérations différentes (dont neuf ne sont pas commutatives) si notre ensemble contient 16 éléments.

Un dernier exemple quand même un peu plus tordu que les précédents, puisqu'il s'agit d'une structure de groupe non commutatif, qui a pourtant une origine simple et géométrique. Prenez un beau triangle équilatéral (dessinez-le sur votre feuille si vous voulez) ABC . On s'intéresse aux isométries du plan laissant stable le triangle équilatéral (autrement dit, on va déplacer ou symétriser notre triangle, mais à la fin on doit toujours avoir un triangle équilatéral à la même place). Avec un peu de motivation, on peut prouver qu'il n'existe que six transformations géométriques laissant effectivement stable notre triangle équilatéral :

- l'application identité, notée i , qui ne fait rien bouger.
- la rotation par rapport au centre O du triangle d'angle $\frac{2\pi}{3}$ (qui permute les trois sommets, et donc les trois côtés, du triangle), qu'on notera r_1 .
- la rotation par rapport au centre O du triangle d'angle $\frac{4\pi}{3}$, qu'on notera r_2 .
- la réflexion par rapport à la droite (AI) , où I est le milieu du côté $[BC]$ opposé à A . On notera cette réflexion s_1 .
- de même, la réflexion par rapport à la droite (BJ) , où J est le milieu du côté $[AC]$. On notera cette réflexion s_2 .
- la réflexion par rapport à la droite (CK) , où K est le milieu du côté $[AB]$. On notera cette réflexion s_3 .

On dispose donc d'un ensemble à six éléments, sur lequel une opération très naturelle va créer une structure de groupe : l'opération de composition (celle qu'on note $\circ$). Et c'est logique : la composée de deux applications laissant stable notre triangle continuera à le laisser stable, on a un élément neutre évident qui est l'identité i , et chaque application a une réciproque (c'est elle qui joue le rôle d'inverse pour l'opération de composition) qui est dans la liste puisqu'elle va elle-même laisser le triangle stable. Si on écrit la table complète de l'opération $\circ$ sur cet ensemble, on obtient :

	i	r_1	r_2	s_1	s_2	s_3
i	i	r_1	r_2	s_1	s_2	s_3
r_1	r_1	r_2	i	s_2	s_3	s_1
r_2	r_2	i	r_1	s_3	s_1	s_2
s_1	s_1	s_3	s_2	i	r_2	r_1
s_2	s_2	s_1	s_3	r_1	i	r_2
s_3	s_3	s_2	s_1	r_2	i	r_1

Attention, la lecture de ce tableau est plus compliquée que pour les précédents puisque l'opération n'est pas commutative : le résultat donné est celui obtenu en composant l'élément de la ligne (à gauche) par l'élément de la colonne (à droite). Ainsi, on a par exemple $r_2 \circ s_1 = s_3$ mais $s_1 \circ r_2 = s_2$. Cette structure est en fait la première structure de groupe non commutatif sur un ensemble fini (on ne peut pas en créer sur des ensembles à moins de six éléments). Les mathématiciens qui font de l'algèbre tous les jours et qui aiment le vocabulaire compliqué l'appellent **groupe diédral** d'ordre 6, les autres l'appellent plus simplement groupe des symétries du triangle. On pourrait bien sûr faire de même avec un carré, un hexagone ou même n'importe quel polygone régulier à n cotés. On obtient toujours un groupe non commutatif à $2n$ éléments (il n'y a que des rotations et des réflexions, comme pour le triangle).

1.3 Noyau et image.

Définition 3. Le **noyau** d'une application linéaire $f : E \rightarrow F$ est l'ensemble $\ker(f) = \{u \in E \mid f(u) = 0\}$.

Remarque 6. Les lettres Ker sont les premières du mot allemand *Kernel* qui signifie, comme vous auriez pu le deviner, noyau. Il correspond en fait aux antécédents d'un élément très particulier de l'ensemble d'arrivée, son vecteur nul. On va voir ci-dessous que, lorsqu'une application est linéaire, le fait de connaître les antécédents de 0 suffit en fait à en déduire des choses sur les antécédents de **tous** les éléments de l'espace d'arrivée. Première remarque : on a dit en début de ce chapitre qu'une application linéaire vérifiait toujours $f(0) = 0$, le noyau de f contient donc toujours au moins un élément, il ne peut jamais être vide. On peut en fait être beaucoup plus précis que ça :

Proposition 3. Si $f \in \mathcal{L}(E, F)$ est une application linéaire, alors $\ker(f)$ est un sous-espace vectoriel de E .

Démonstration. On peut en fait prouver un résultat plus général que celui-ci, puisque l'image réciproque de tout sous-espace vectoriel de l'espace d'arrivée F est un sous-espace vectoriel de E (comme $\{0\}$ est un sous-espace de F , notre proposition est un cas particulier de ce résultat). Soit H un sous-espace vectoriel de F , et $(u, v) \in f^{-1}(H)$, alors $f(u) = z \in H$ et $f(v) = w \in H$, donc $f(\lambda u + \mu v) = \lambda z + \mu w \in H$, et $\lambda u + \mu v \in f^{-1}(H)$. $\square$

Proposition 4. Une application linéaire $f : E \rightarrow F$ est injective si et seulement si $\ker(f) = \{0_E\}$.

Démonstration. Cette propriété revient à dire que, pour démontrer qu'une application linéaire est injective, il suffit de démontrer qu'un seul élément bien particulier (le vecteur nul) n'a pas plus d'un antécédent par f . Supposons d'abord le noyau réduit au vecteur nul et montrons que f est injective : soient $(u, v) \in E^2$ tels que $f(u) = f(v)$, alors, par linéarité de f , on a $f(u - v) = f(u) - f(v) = 0$, donc $u - v \in \ker(f)$. Comme le noyau est réduit au vecteur nul, on a donc $u - v = 0$, c'est-à-dire $u = v$, ce qui prouve bien l'injectivité. Réciproquement, supposons f injective, alors 0 a (au plus) un seul antécédent par f . Or, le vecteur nul est toujours un antécédent de 0 par une application linéaire. Ceci prouve bien qu'il est le seul élément de E à appartenir à $\ker(f)$. $\square$

Exemple : En pratique, un noyau se calcule très simplement en résolvant un système. Déterminons par exemple le noyau de l'endomorphisme f de $\mathbb{R}^3$ défini par $f : (x, y, z) \mapsto (x - y + z, 3x - 2y + 5z, -x - 3z)$ (je vous dispense de la vérification pénible du fait que f est bien une application linéaire). Les éléments du noyau sont les triplets de réels (x, y, z) solutions du système

$$\begin{cases} x - y + z = 0 \\ 3x - 2y + 5z = 0 \\ -x - 3z = 0 \end{cases} . \text{ Le système n'est pas de Cramer } (2L_1 - L_2 = L_3), \text{ les solutions sont}$$

les triplets de la forme $(-3z, -2z, z)$, avec $z \in \mathbb{R}$. Autrement dit, $\ker(f) = \text{Vect}((-3, -2, 1))$, ce qui prouve que l'application f n'est pas injective. On essaiera de systématiquement présenter les noyaux d'applications linéaires « sous forme de Vect » pour en déduire rapidement leur dimension.

Exemple 2 : L'application $f : \mathbb{R}_n[X] \rightarrow \mathbb{R}_n[X]$ définie par $f(P) = P'$ a pour noyau l'ensemble des polynômes constants, ce qu'on peut écrire $\ker(f) = \text{Vect}(1)$. Là encore, l'application f n'est pas injective.

Définition 4. L'image d'une application linéaire $f : E \rightarrow F$ est l'ensemble $\text{Im}(f) = \{v \in F \mid \exists u \in E, f(u) = v\}$.

Remarque 7. Il ne s'agit en fait pas d'une nouvelle définition puisque nous avons déjà défini il y a un certain temps l'image d'une application quelconque entre deux ensembles de la même façon. La première propriété énoncée ci-dessus découle de façon immédiate de la définition de l'image et ne mérite donc aucune démonstration :

Proposition 5. Une application linéaire $f : E \rightarrow F$ est surjective si et seulement si $\text{Im}(f) = F$.

Proposition 6. Soit $f \in \mathcal{L}(E, F)$ et $(e_1, \dots, e_n)$ une base de E , alors $\text{Im}(f) = \text{Vect}(f(e_1), \dots, f(e_n))$.

Démonstration. Comme les vecteurs $f(e_1), \dots, f(e_n)$ appartiennent évidemment à $\text{Im}(f)$, on a nécessairement $\text{Vect}(f(e_1), \dots, f(e_n)) \subset \text{Im}(f)$. De plus, si $v \in \text{Im}(f)$, alors $v = f(u)$ avec $u \in E$, et comme $(e_1, \dots, e_n)$ est une base de E , on peut écrire $u = \sum_{i=1}^n \lambda_i e_i$. Alors $v = f(u) = f\left(\sum_{i=1}^n \lambda_i e_i\right) = \sum_{i=1}^n \lambda_i f(e_i)$, donc $v \in \text{Vect}(f(e_1), \dots, f(e_n))$, et les deux ensembles sont bien égaux. $\square$

Remarque 8. Attention, en général, $(f(e_1), \dots, f(e_n))$ n'est pas une **base** de $\text{Im}(f)$, mais seulement une famille génératrice (elle n'a aucune raison d'être libre, ce ne sera d'ailleurs en pratique le cas que si f est un isomorphisme). Si on veut connaître la dimension de l'image, il sera donc nécessaire de supprimer dans la famille $(f(e_1), \dots, f(e_n))$ les vecteurs inutiles pour la transformer en famille libre (en pratique, on va voir que le théorème du rang permet souvent de procéder en sens inverse : on connaît la dimension de l'image avant de la calculer concrètement, et on sait donc à l'avance le nombre de vecteurs à supprimer dans la famille génératrice obtenue).

Exemple 1 : La méthode élémentaire pour calculer une image est d'utiliser la définition, mais elle n'est pas pratique du tout. Prenons par exemple l'application linéaire de $\mathbb{R}^2$ dans $\mathbb{R}^3$ définie par $f(x, y) = (2x - y, x + 2y, -2x + y)$. Un triplet (a, b, c) appartient à $\text{Im}(f)$ si et seulement si le système

$$\begin{cases} 2x - y = a \\ x + 2y = b \\ -2x + y = c \end{cases}$$

admet une solution. Les membres de gauche des deux équations extrêmes étant opposés, il faut nécessairement avoir $a = -c$, et on vérifie facilement que cette condition est suffisante. On a donc $\text{Im}(f) = \{(a, b, -a) \mid a, b \in \mathbb{R}\} = \text{Vect}((1, 0, -1), (0, 1, 0))$.

Exemple 2 : En pratique, on utilisera plutôt notre dernière proposition, car c'est beaucoup plus rapide ! Reprenons le même exemple. La base canonique de $\mathbb{R}^2$ est constituée des deux vecteurs $(1, 0)$ et $(0, 1)$, donc l'image est engendrée par $f(1, 0) = (2, 1, -2)$ et $f(0, 1) = (-1, 2, 1)$. On a donc $\text{Im}(f) = \text{Vect}((2, 1, -2), (-1, 2, 1))$ (ce ne sont pas les mêmes vecteurs que tout à l'heure mais on peut vérifier qu'ils engendrent le même espace vectoriel).

Proposition 7. Si f est un isomorphisme de E dans F , alors sa réciproque f^{-1} est un isomorphisme de F dans E .

Démonstration. La seule chose à vérifier est la linéarité de la réciproque (qui existe nécessairement puisque f est supposée bijective). Soit donc $(z, w) \in F^2$ et notons $u = f^{-1}(z)$ et $v = f^{-1}(w)$. Par linéarité de f , on peut dire que $f(\lambda u + \mu v) = \lambda f(u) + \mu f(v) = \lambda z + \mu w$, donc $f^{-1}(\lambda z + \mu w) = \lambda u + \mu v = \lambda f^{-1}(z) + \mu f^{-1}(w)$, ce qui prouve la linéarité de f^{-1} . $\square$

Remarque 9. L'ensemble $\mathcal{L}(E)$ muni des deux opérations $+$ et $\circ$ est ce qu'on appelle un anneau non commutatif. L'addition joue son rôle usuel et la composition joue à peu de choses près le rôle de la multiplication dans les ensembles de nombres usuels ($\mathbb{R}$ ou $\mathbb{C}$ par exemple). En effet, la composition admet un élément neutre qui est l'application identité, mais toutes les applications linéaires ne sont pas inversibles (seuls les automorphismes le sont), et surtout la composition est distributive par

rapport à l'addition, tout comme le produit dans les ensembles de nombres (c'est-à-dire qu'on a toujours $f \circ (g + h) = f \circ g + f \circ h$ et $(g + h) \circ f = g \circ f + h \circ f$, ce qu'on ne se privera pas d'exploiter dans certains calculs). En fait, nous verrons plus tard que la composition d'applications linéaires s'identifie effectivement à un produit, celui des matrices. Pour l'instant, nous utiliserons déjà cette analogie pour justifier l'énorme abus de notation suivant :

Définition 5. Si f est un endomorphisme d'un espace vectoriel E , on notera f^2 la composée $f \circ f$, et plus généralement $f^n = \underbrace{f \circ f \circ \dots \circ f}_{n \text{ fois}}$.

Il est essentiel de garder constamment à l'esprit que, malgré la notation, f^n n'a rien à voir avec un calcul de puissances de coordonnées.

2 Rang.

Définition 6. Soit $\mathcal{F} = (e_1, \dots, e_n)$ une famille de vecteurs d'un espace vectoriel E , on appelle **rang de la famille** $\mathcal{F}$ la dimension de $\text{Vect}(e_1, \dots, e_n)$.

Remarque 10. La famille est libre si et seulement si son rang est égal au nombre de vecteurs qu'elle contient. Dans le cas où la famille n'est pas libre, le rang donne immédiatement le nombre de vecteurs « inutiles » de la famille : si une famille de cinq vecteurs a un rang égal à 3, cela signifie qu'il faudra supprimer deux vecteurs pour obtenir une base de l'espace vectoriel engendré. Cette notion de rang peut très bien s'appliquer à une famille de vecteurs dans un espace vectoriel de dimension infinie (puisque la famille elle-même est finie, et l'espace vectoriel qu'elle engendre est donc de dimension finie).

Définition 7. Soit $f \in \mathcal{L}(E, F)$, le **rang de** f , s'il existe, est la dimension de l'image de f . On le note $\text{rg}(f)$.

Remarque 11. Pour faire le lien avec la définition précédente, $\text{rg}(f) = \text{rg}(f(e_1), f(e_2), \dots, f(e_n))$, où $(e_1, \dots, e_n)$ est une base quelconque de E . Pour être très rigoureux, cette égalité ne peut avoir de sens que si E est de dimension finie, alors que le rang de f peut éventuellement être défini même si E n'est pas de dimension finie.

Théorème 1. Théorème du rang.

Soit $f \in \mathcal{L}(E, F)$, où E est un espace vectoriel de dimension finie, alors $\text{rg}(f) + \dim(\ker(f)) = \dim(E)$.

Remarque 12. Le théorème du rang n'affirme absolument pas que le noyau et l'image de f sont supplémentaires (ce qui ne pourrait de toute façon avoir de sens que pour un endomorphisme, sinon le noyau et l'image ne sont pas des sous-espaces vectoriels du même espace), c'est faux en général (voir les exemples suivant la démonstration), mais simplement que les dimensions du noyau et de l'image sont « complémentaires ». En gros, dans le cas où f est un endomorphisme, le théorème affirme qu'une application linéaire est « d'autant moins injective qu'elle n'est pas surjective ». Imaginons par exemple $f : E \rightarrow E$, avec E un espace vectoriel de dimension 3. Si l'application n'est pas surjective, cela signifie que $\text{rg}(f) < 3$. Si ce rang est égal à 2, alors la dimension du noyau sera égale à 1, si ce rang est égal à 1, le noyau sera de dimension 2. Plus l'image est réduite, plus le noyau est au contraire important.

Démonstration. L'idée est de démontrer qu'à défaut d'être supplémentaire de $\ker(f)$, $\text{Im}(f)$ est isomorphe à tout supplémentaire de $\ker(f)$ (et donc a la même dimension, ce qui prouve immédiatement le théorème). Soit donc G un supplémentaire de $\ker(f)$ (cela existe forcément dans un espace

vectorel de dimension finie, c'est une conséquence du théorème de la base incomplète : on considère une base de $\ker(f)$ et on la complète en base de E . Tous les vecteurs ajoutés lors de cette complétion forment une base d'un sous-espace de E qui est par construction supplémentaire de $\ker(f)$. Montrons que $f|_G$ (la restriction de f au sous-espace vectoriel G) est un isomorphisme de G sur $\text{Im}(f)$. Soit $u \in \ker(f|_G)$, on a donc $f(u) = 0$ mais aussi $u \in G$ (puisque $f|_G$ n'est définie que sur G), donc $u \in G \cap \ker(f)$. Comme G et $\ker(f)$ sont supplémentaires, leur intersection est réduite au vecteur nul, donc $u = 0$, ce qui prouve l'injectivité de $f|_G$. Soit maintenant $v \in \text{Im}(f)$, donc $v = f(u)$, avec $u \in E$. Ce vecteur u peut se décomposer en $u_G + u_K$ (toujours en exploitant le fait que G et $\ker(f)$ sont supplémentaires), avec $u_G \in G$ et $u_K \in \ker(f)$. Par définition, $f(u_K) = 0$, donc $v = f(u) = f(u_G + u_K) = f(u_G)$, et $v \in \text{Im}(f|_G)$ (le vecteur v est bien l'image par f d'un vecteur u_G appartenant à G), ce qui prouve que $\text{Im}(f) = \text{Im}(f|_G)$. L'application f restreinte à G est donc également surjective, c'est bien un isomorphisme, le théorème en découle. $\square$

Exemple 1 : On pose $A = \begin{pmatrix} 2 & 1 \\ -1 & 1 \end{pmatrix}$ et on définit un endomorphisme f de l'espace vectoriel $\mathcal{M}_2(\mathbb{R})$ par $f(M) = AM - MA$. Le noyau de cet endomorphisme est constitué de toutes les matrices commutant avec A , qu'on obtient par un calcul brutal : en posant $M = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$,

on a $AM = \begin{pmatrix} 2a+c & 2b+d \\ c-a & d-b \end{pmatrix}$ et $MA = \begin{pmatrix} 2a-b & a+b \\ 2c-d & c+d \end{pmatrix}$, donc $M \in \ker(f)$ si et seulement si
$$\begin{cases} b + c = 0 \\ -a + b + d = 0 \\ -a - c + d = 0 \\ -b - c = 0 \end{cases}.$$
 Les deux équations extrêmes donnent $c = -b$, et les

deux autres sont alors équivalentes : $a = b + d$, donc $\ker(f) = \left\{ \begin{pmatrix} b+d & b \\ -b & d \end{pmatrix} \mid (b, d) \in \mathbb{R}^2 \right\} = \text{Vect} \left(\begin{pmatrix} 1 & 1 \\ -1 & 0 \end{pmatrix}, \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \right)$. Ce noyau étant manifestement de dimension 2, le théorème du rang assure que l'image de f sera également de dimension 2. On peut donc tricher un peu et se contenter par exemple de calculer les images par f des deux premières matrices de la base canonique de $\mathcal{M}_2(\mathbb{R})$ (à condition qu'elles ne soient pas proportionnelles, bien entendu) pour obtenir $\text{Im}(f) = \text{Vect} \left(\begin{pmatrix} 0 & -1 \\ -1 & 0 \end{pmatrix}, \begin{pmatrix} 1 & 1 \\ 0 & -1 \end{pmatrix} \right)$. Il n'est pas difficile ici de prouver que le noyau et l'image de f sont supplémentaires.

Exemple 2 : Soit $f \in \mathcal{L}(\mathbb{R}^4)$ l'application linéaire définie par $f(x, y, z, t) = (x - y, x - y, t - z, z - t)$. On constate aisément que $\text{Im}(f) = \text{Vect}((1, 1, 0, 0), (0, 0, 1, -1))$ (les images des deux premiers vecteurs de la base canonique sont opposées, celles des deux derniers également), donc f est de rang 2. La détermination du noyau amène de même à un système constitué de deux paires d'équations identiques, et $\ker(f) = \{(x, x, z, z) \mid (x, z) \in \mathbb{R}^2\} = \text{Vect}((1, 1, 0, 0), (0, 0, 1, 1))$. Le noyau est également de dimension 2 (encore heureux, sinon le théorème du rang serait mis en défaut), mais noyau et image ne sont pas supplémentaires (puisque'ils contiennent tous deux le vecteur $(1, 1, 0, 0)$, l'intersection est en fait de dimension 1).

On peut aller ici un peu plus loin et s'amuser à calculer $f^2(x, y, z, t)$ (comme toujours, ce qu'on note f^2 c'est $f \circ f$). En notant $X = x - y$ et $Z = z - t$, on a $f(x, y, z, t) = (X, X, Z, -Z)$, donc $f^2(x, y, z, t) = f(X, X, Z, -Z) = (X - X, X - X, Z + Z, -Z - Z) = (0, 0, 2z - 2t, 2t - 2z)$. On en déduit que $\text{Im}(f^2) = \text{Vect}((0, 0, 1, -1))$ et $\ker(f^2) = \text{Vect}((1, 0, 0, 0), (0, 1, 0, 0), (0, 0, 1, 1))$ (calculs vraiment débiles laissés au lecteur). Cette fois-ci, les deux sous-espaces sont supplémentaires. D'ailleurs, les composées suivantes de f ont les mêmes noyaux et images que f^2 . On peut prouver plus généralement que, pour tout endomorphisme f d'un espace vectoriel de dimension finie, les noyaux et images de f^k finissent par se stabiliser sur des sous-espaces supplémentaires (c'est l'objet d'un exercice de la feuille d'exercices).

Corollaire 1. Soit $f \in \mathcal{L}(E)$, où E est un espace vectoriel de dimension finie, f est bijectif si et seulement si il est injectif ou surjectif.

Démonstration. En effet, si par exemple f est injectif, $\dim(\ker(f)) = 0$, donc en appliquant le théorème du rang $\text{rg}(f) = \dim(E)$, ce qui assure que $\text{Im}(f) = E$, et donc que f est également surjectif. C'est à peu près la même chose si on suppose f surjectif. $\square$

Exemple : La remarque précédente ne s'applique absolument pas en dimension infinie. Si on note f l'application définie sur l'ensemble de toutes les suites réelles par $f(u_n) = v_n$, où $v_n = u_{n+1}$ (décalage de tous les termes de la suite vers la gauche, en supprimant le premier). L'application f est surjective mais pas injective. En fait, en notant $g(u_n) = w_n$, avec $w_0 = 0$ et $\forall n \geq 1, w_n = u_{n-1}$, on a $f \circ g = \text{id}$, mais $g \circ f \neq \text{id}$ (on transforme le premier terme de la suite en 0). Autrement dit, g est une « réciproque à droite » de f mais pas à gauche. Ce genre de choses n'arrive absolument jamais en dimension finie, si on a deux applications linéaires f et g qui vérifient $f \circ g = \text{id}$, on aura automatiquement $g \circ f = \text{id}$, et f et g seront donc bijectives et réciproques l'une de l'autre. Sur l'exemple donné en dimension infinie, on peut vérifier que g est injective mais pas surjective.

Exemple : Soient $x_0, x_1, \dots, x_n$ des réels deux à deux distincts et $f \in \mathcal{L}(\mathbb{R}_n[X])$ l'endomorphisme défini par $f(P) = (P(x_0), P(x_1), \dots, P(x_n))$. Le noyau de f est réduit au polynôme nul car c'est le seul polynôme de degré inférieur ou égal à n pouvant avoir $n+1$ racines distinctes. Par conséquent, f est un isomorphisme. Cela prouve que, pour tous réels $a_0, a_1, \dots, a_n$, il existe un unique polynôme de degré n (au plus) tel que $\forall i \in \{0, \dots, n\}, P(x_i) = a_i$ (P est l'unique antécédent de $(a_0, \dots, a_n)$ par l'application f). Vous le saviez en fait déjà, il s'agit des polynômes interpolateurs de Lagrange (mais le théorème du rang permet donc de prouver leur existence sans avoir besoin le moins du monde de les calculer).

Définition 8. Une **forme linéaire** sur un espace vectoriel E est une application linéaire $f \in \mathcal{L}(E, \mathbb{R})$.

Exemple : Si on note E l'ensemble de toutes les fonctions continues sur $[0, 1]$, l'application $\varphi : f \mapsto \int_0^1 f(t) dt$ est une forme linéaire sur E (application linéaire donnant des images qui sont simplement des nombres réels).

Proposition 8. Si E est de dimension finie, le noyau d'une forme linéaire non nulle est un hyperplan de E .

Démonstration. Comme, par définition, $\text{Im}(f) \subset \mathbb{R}$, il n'y a pas beaucoup de choix pour le rang d'une forme linéaire : soit il est nul, et f est alors l'application nulle ; soit il vaut 1, et d'après le théorème du rang on a alors $\dim(\ker(f)) = \dim(E) - 1$. $\square$

Exemple : La trace est une forme linéaire sur $\mathcal{M}_n(\mathbb{R})$ (rappelons que la trace d'une matrice carrée est définie comme la somme de ses coefficients diagonaux). L'ensemble des matrices de trace nulle est donc un hyperplan (de dimension $n^2 - 1$ dans ce cas) de $\mathcal{M}_n(\mathbb{R})$.

3 Applications linéaires et géométrie.

Nous allons retrouver dans ce paragraphe un premier lien vraiment concret entre algèbre linéaire et géométrie, en étudiant quelques types d'applications linéaires bien particulières, que vous connaissez déjà en géométrie plane depuis longtemps (et qui représentent vraiment le même type de transformations géométriques que dans le plan ou dans l'espace, même si c'est bien sûr plus délicat à visualiser dans des espaces vectoriels de dimension plus grande).

Définition 9. Soit E un espace vectoriel réel. L'**homothétie de rapport** λ est l'endomorphisme de E défini par $f(x) = \lambda x$. Autrement dit, $f = \lambda \text{id}$, où $\lambda \in \mathbb{R}$.

Remarque 13. Cela correspond bien à la notion usuelle d'homothétie de rapport λ , mais dans le cadre des espaces vectoriels, le centre de l'homothétie sera toujours en 0 (c'est indispensable si on veut qu'il puisse s'agir d'une application linéaire, puisque le vecteur nul doit rester invariant par notre application).

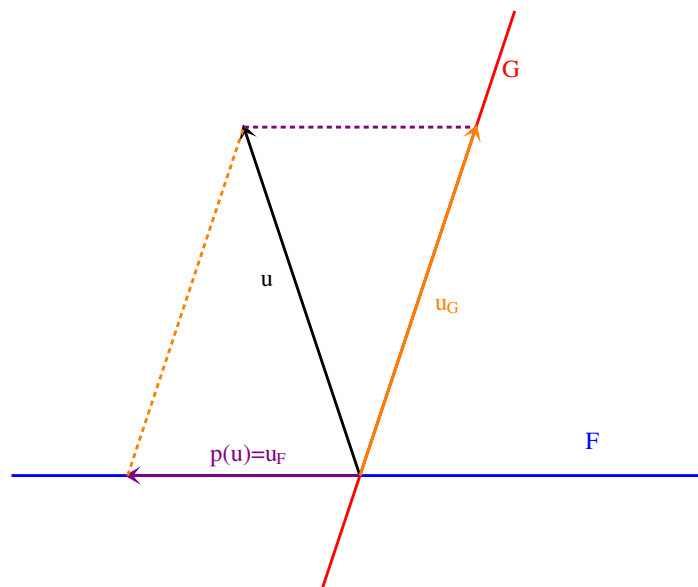
Proposition 9. Si $\lambda \neq 0$, l'homothétie de rapport λ est un automorphisme de E , son automorphisme réciproque est l'homothétie de rapport $\frac{1}{\lambda}$.

Démonstration. Une démonstration à la portée de tous : $(\lambda \text{id}) \circ \left(\frac{1}{\lambda} \text{id}\right) = \text{id}$. □

Remarque 14. En tant que multiples de l'identité, les homothéties commutent avec tous les autres endomorphismes de E (ici, le terme « commutent » se rapporte à l'opération de composition, deux endomorphismes f et g commutent si $f \circ g = g \circ f$). On peut d'ailleurs prouver que ce sont les seules applications linéaires dans ce cas.

Définition 10. Soient F et G deux sous-espaces vectoriels supplémentaires d'un même espace vectoriel E . On peut donc décomposer tout vecteur $u \in E$ sous la forme $u = u_F + u_G$, avec $u_F \in F$ et $u_G \in G$ (et cette décomposition est unique). La **projection sur F parallèlement à G** est alors l'application linéaire $p : u \mapsto u_F$.

Remarque 15. Cette application est bien une application linéaire, car la décomposition de $\lambda u + \mu v$ dans $F \oplus G$ est simplement $(\lambda u_F + \mu v_F) + (\lambda u_G + \mu v_G)$. Surtout, ce type d'application correspond vraiment à une projection comme on en fait dans le plan. Attention par contre, la notion de projection **orthogonale** n'a absolument aucun sens dans un espace vectoriel quelconque (en fait, la notion d'orthogonalité n'existe pas dans les espaces vectoriels tant qu'on n'a pas ajouté à la structure d'espace vectoriel un outil supplémentaire, en l'occurrence ce qu'on appellera plus tard un produit scalaire ; au contraire, la notion de parallélisme est intrinsèque à la structure d'espace vectoriel, elle existe naturellement dès que l'ensemble E est muni des deux opérations permettant d'en faire un espace vectoriel). Pour illustrer, je vais me contenter d'un schéma dans $\mathbb{R}^2$ (la droite en bleu est ici F , la droite en rouge est G), où on projette donc sur une droite parallèlement à une autre (droites passant nécessairement par l'origine, ce sont des sous-espaces vectoriels de $\mathbb{R}^2$). C'est le seul cas intéressant qu'on puisse faire en dimension 2, puisque les sous-espaces F et G doivent être supplémentaires. Mais le principe est le même dans un espace de dimension quelconque : on « découpe » le vecteur u en deux morceaux appartenant respectivement à F et à G , et on ne garde que le morceau appartenant à F , ce qui revient bien à projeter sur F parallèlement à G .



Les algébristes, qui sont comme chacun sait des gens bizarres, utilisent également le terme **projecteur** comme synonyme de projection.

Remarque 16. Si p est le projecteur sur F parallèlement à G , $q = \text{id} - p$ est le projecteur sur G parallèlement à F . En effet, avec les notations introduites ci-dessus, $q(u) = u_G = u - u_F = (\text{id} - p)(u)$.

Théorème 2. Un endomorphisme $f \in \mathcal{L}(E)$ est un projecteur si et seulement si $f \circ f = f$.

Démonstration. Il y a un sens évident : si f est une projection, alors, pour tout vecteur $u \in E$, on a $f(f(u)) = f(u_F) = u_F$ (en gardant les notations précédentes) puisque la décomposition du vecteur u_F dans $F \oplus G$ est simplement $u_F = u_F + 0$, donc $f^2 = f$ (c'est géométriquement cohérent, la première projection a envoyé tous les vecteurs sur F , projeter une deuxième fois sur ce même espace F ne peut rien changer). La réciproque est nettement moins évidente. Supposons donc que f vérifie $f^2 = f$, et notons $F = \text{Im}(f)$ et $G = \text{ker}(f)$. Vérifions pour commencer que ces deux sous-espaces vectoriels sont supplémentaires dans E (on s'interdit ici d'utiliser des arguments de dimension, car ce théorème et sa démonstration restent valables dans un espace de dimension infinie) :

- si $u \in F \cap G$, alors $f(u) = 0$ et $\exists v \in E, u = f(v)$. Mais on a alors $f(f(v)) = 0$, donc $f(v) = 0$ puisque par hypothèse $f^2 = f$. Or, $u = f(v)$, ce qui prouve que $u = 0$ et donc que $F \cap G = \{0\}$.
- tout vecteur $u \in E$ peut s'écrire sous la forme $u = u - f(u) + f(u)$. Le vecteur $f(u)$ appartient clairement à F , et le vecteur $u - f(u)$ appartient bien à G car $f(u - f(u)) = f(u) - f^2(u) = 0$ en exploitant à nouveau l'hypothèse $f^2 = f$. On peut donc décomposer tout vecteur de E dans $F + G$, ce qui prouve que $E = F + G$.

On a bien prouvé que $E = F \oplus G$. On a même en passant écrit explicitement la décomposition d'un vecteur u dans cette somme directe : en reprenant les notations précédentes, $u_F = f(u)$ et $u_G = u - f(u)$. Ceci prouve que la projection p sur F parallèlement à G est définie par $p(u) = u_F = f(u)$, autrement dit que f est une projection, sur F parallèlement à G . $\square$

Remarque 17. Quand on étudie un endomorphisme f et qu'on souhaite prouver qu'il s'agit d'un projecteur, on calcule donc tout simplement f^2 . Si on obtient l'égalité $f^2 = f$, il est ensuite très

facilement de déterminer les sous-espaces correspondant à la projection (qu'on appelle aussi éléments caractéristiques de la projection) : le sous-espace sur lequel on projette est l'image de f , celui parallèlement auquel on projette est le noyau de f .

Proposition 10. Soit $p \in \mathcal{L}(E)$ un projecteur. Alors $\text{Im}(p) = \ker(p - \text{id})$. Autrement dit, $\text{Im}(p) = \{u \in E \mid p(u) = u\}$.

Démonstration. C'est en fait évident : avec les notations habituelles, $p(u) = u \Leftrightarrow u_F = u \Leftrightarrow u \in F$. $\square$

Remarque 18. Cette propriété permet de calculer l'image d'un projecteur sous forme de noyau (donc en résolvant un système plutôt qu'en calculant les images des vecteurs d'une base). Attention, cette méthode ne fonctionne que pour les projecteurs et pas du tout pour une application linéaire quelconque.

Exemple : L'application définie sur $\mathbb{R}^2$ par $f(x, y) = \left(\frac{x+y}{2}, \frac{x+y}{2}\right)$ est une projection. Le plus simple pour le prouver est de constater que $f \circ f = f$ (c'est ici un calcul immédiat). Le noyau de f est constitué des vecteurs pour lesquels $x + y = 0$, autrement dit $F = \ker(f) = \text{Vect}((1, -1))$, et l'image de ceux vérifiant $f(x, y) = (x, y)$, soit $\frac{x+y}{2} = x = y$, donc $x = y$. Autrement dit, $\text{Im}(f) = \text{Vect}((1, 1))$.

Définition 11. Avec les mêmes hypothèses et notations que dans la définition des projections, la **symétrie par rapport à F parallèlement à G** est l'application linéaire $s : x \mapsto u_F - u_G$.

Théorème 3. Un endomorphisme $f \in \mathcal{L}(E)$ est une symétrie si et seulement si $f \circ f = \text{id}$.

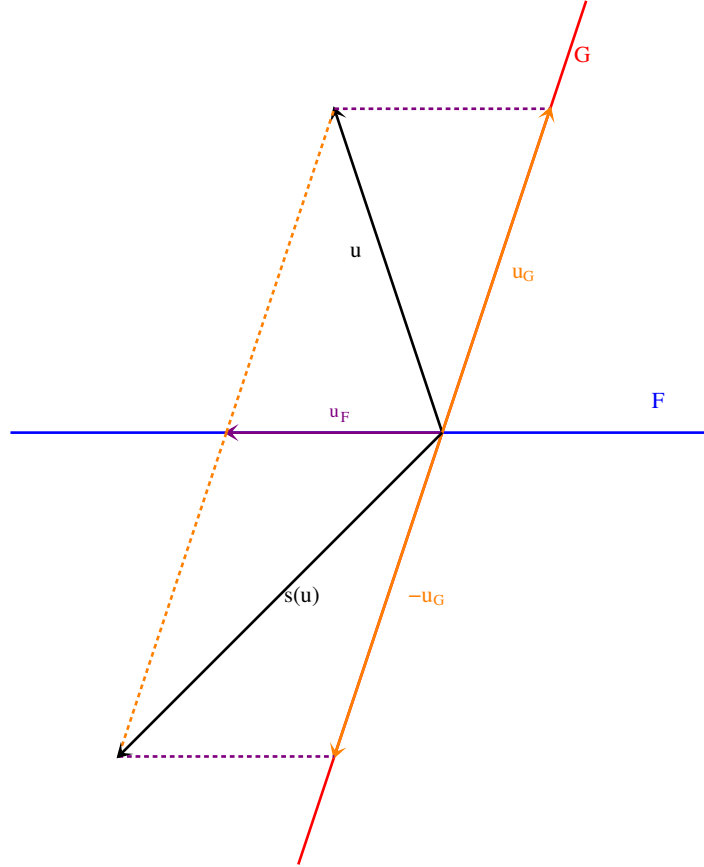
Démonstration. C'est exactement le même principe que pour les projections : il y a un sens évident ($f(f(u)) = u_F - (-u_G) = u_F + u_G = u$, donc $f^2 = \text{id}$), et la réciproque l'est nettement moins. Il faut surtout « deviner » à quoi correspondent les sous-espaces F et G dans le cas d'une symétrie. Posons donc $F = \ker(f - \text{id})$, ce qui correspond bien à l'espace par rapport auquel on symétrise (F est constitué des vecteurs u vérifiant $f(u) = u$, autrement dit des vecteurs invariants par f), et $G = \ker(f + \text{id})$ (u appartient à l'espace parallèlement auquel on symétrise si $f(u) = -u$). Comme dans le cas des projecteurs, on prouve que ces deux sous-espaces sont supplémentaires : un vecteur u non nul ne peut pas vérifier à la fois $f(u) = u$ et $f(u) = -u$, donc $F \cap G = \{0\}$. De plus, on peut écrire tout vecteur appartenant à E sous la forme $u = \frac{u + s(u)}{2} + \frac{u - s(u)}{2}$. Le premier élément de cette décomposition vérifie $f\left(\frac{u + f(u)}{2}\right) = \frac{f(u) + u}{2}$ (en exploitant l'hypothèse $f^2 = \text{id}$), donc il appartient à F . De même, le deuxième appartient à G . Ces deux éléments correspondent donc aux vecteurs notés u_F et u_G dans nos définitions. Enfin, $u_F - u_G = \frac{u + f(u)}{2} - \frac{u - f(u)}{2} = f(u)$, donc f est bien une symétrie, plus précisément la symétrie par rapport à F parallèlement à G . $\square$

Remarque 19. Dans le cas d'une symétrie, calculer l'image et le noyau de l'application n'a strictement aucun intérêt, puisqu'une symétrie est toujours bijective (sa réciproque étant la symétrie elle-même puisque $f \circ f = \text{id}$).

Proposition 11. Soit $s \in \mathcal{L}(E)$ une symétrie. En notant $F = \ker(s - \text{id})$ et $G = \ker(s + \text{id})$, s est la symétrie par rapport à F parallèlement à G . De plus, en notant p la projection sur ce même espace F parallèlement à G , $s = 2p - \text{id}$.

Démonstration. Avec les notations habituelles, $2p(u) - u = 2u_F - (u_F + u_G) = u_F - u_G = s(u)$. $\square$

Remarque 20. Comme dans les cas des projections, la notion de symétrie orthogonale n'a aucun sens.



Exemple : Dans $\mathbb{R}^3$, on cherche à déterminer une expression analytique de la symétrie par rapport à $F = \text{Vect}((1, 0, 1))$ et parallèlement à $G = \text{Vect}((1, 2, 3), (1, 0, 0))$. Il faudrait déjà commencer par prouver que $F \oplus G = E$. Comme nous avons de toute façon besoin de connaître la décomposition d'un vecteur dans $F \oplus G$ pour calculer son image par s , le calcul ne peut pas faire de mal. Considérons donc $u = (x, y, z) \in \mathbb{R}^3$, et cherchons trois réels a, b et c tels que $(x, y, z) = a(1, 0, 1) + b(1, 2, 3) + c(1, 0, 0)$. Il n'est même pas indispensable d'écrire entièrement le système : la deuxième coordonnée donne immédiatement $y = 2b$, soit $b = \frac{y}{2}$, puis la troisième donne $a + 3b = z$, soit $a = z - 3b = z - \frac{3}{2}y$, et enfin via la première coordonnée $x = a + b + c$, donc $c = x - a - b = x + y - z$. Finalement, le système admet toujours une solution unique, ce qui prouve la suppléantarité de F et de G . Par ailleurs, avec les notations usuelles, $u_F = (a, 0, a)$ et $u_G = (b + c, 2b, 3b)$, donc $s(u) = a(1, 0, 1) - b(1, 2, 3) - c(1, 0, 0) = \left(z - \frac{3}{2}y, 0, z - \frac{3}{2}y\right) - \left(\frac{1}{2}y, y, \frac{3}{2}y\right) - (x + y - z, 0, 0) = (-x - 3y + 2z, -y, z - 3y)$. On vérifie facilement que $s \circ s = \text{id}$ avec cette expression.

Définition 12. Toujours avec les mêmes notations que dans les définitions précédentes, l'**affinité de rapport λ par rapport à F et parallèlement à G** est l'application linéaire $a_\lambda : u \mapsto u_F + \lambda u_G$.

Exemple : Je ne vais pas vraiment donner un exemple concret d'affinité (ça n'a pas grand intérêt), mais simplement signaler que les affinités dans le plan ont l'intéressante particularité de transformer les cercles (centrés en l'origine) en ellipses. Les affinités sont en quelque sorte des hybrides entre l'identité et les homothéties.

Remarque 21. Les projections et symétries sont des cas particuliers d'affinités, lorsque $\lambda = 0$ et $\lambda = -1$ respectivement. Si $\lambda \neq 0$, l'affinité est un automorphisme, dont la réciproque est l'affinité de rapport $\frac{1}{\lambda}$. Notons aussi que, pour une affinité a , on a toujours $E = \ker(a - \text{id}) \oplus \ker(a - \lambda \text{id})$. Poussons d'ailleurs cette remarque un peu plus loin pour les plus motivés d'entre vous : les projecteurs sont des applications linéaires vérifiant $p^2 = p$, ou si on préfère $p^2 - p = 0$, soit encore $p \circ (p - \text{id}) = 0$. Or, on a vu que, pour une projection, $E = \ker(p) \oplus \ker(p - \text{id})$ (puisque $\text{Im}(p) = \ker(p - \text{id})$). De même, une symétrie vérifie $s^2 - \text{id} = 0$, soit $(s + \text{id}) \circ (s - \text{id}) = 0$, et on a par ailleurs $E = \ker(s + \text{id}) \oplus \ker(s - \text{id})$. On peut constater qu'une affinité vérifie quand à elle l'égalité $a_\lambda^2 = (\lambda + 1)a_\lambda - \lambda \text{id}$, soit $(a_\lambda - \text{id}) \circ (a_\lambda - \lambda \text{id}) = 0$, ce qui est cohérent avec le fait que $E = \ker(a_\lambda - \text{id}) \oplus \ker(a_\lambda - \lambda \text{id})$. Il y a sûrement des choses très profondes à généraliser à partir de ces constatations, mais ça nous dépasse un peu pour l'instant (déjà, pour être honnête, la notion d'affinité est hors programme).