

# La suite des $\lambda$ -dérivés d'un polygone

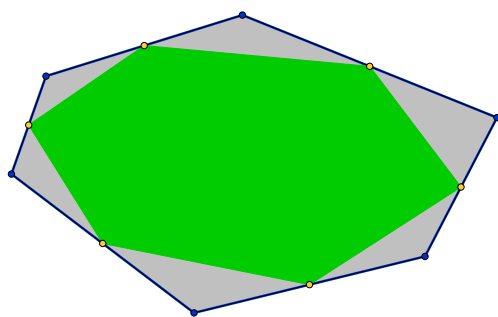
Écrit par Aziz EL KACIMI & Abdellatif ZEGGAR

Publié le 19 octobre 2025

DOI : 10.60868/g10j-2w88 — CC BY-NC-ND 4.0

🔪🔪🔪 — ⌚ ≥ 30 min — 📄

📐 Géométrie



Soient  $\mathcal{P} = [A_1, \dots, A_n]$  un polygone pair<sup>1</sup> et  $\lambda \in ]0, 1[$ . Le polygone  $\mathcal{P}_1 = [A_{1,1}, \dots, A_{1,n}]$  dont les sommets  $A_{1,k}$  sont les barycentres  $(1 - \lambda)A_k + \lambda A_{k+1}$  (avec  $A_{n+1} = A_1$ ) est appelé  $\lambda$ -dérivé de  $\mathcal{P}$ . En itérant cette opération, on obtient la suite  $\mathcal{P}_q$  des polygones  $\lambda$ -dérivés successifs de  $\mathcal{P}$ . Sur un dessin soigneusement exécuté, on peut voir que plus  $q$  est grand, plus  $\mathcal{P}_q$  est proche d'un multiparallélogramme (polygone pair ayant un centre de symétrie). C'est ce que nous démontrons dans cet article en donnant d'abord un sens précis à cette notion de proximité. Le sujet étant assez proche du théorème de Varignon, nous en rappelons l'énoncé classique et en donnons une généralisation aux polygones pairs. Ce travail est parti d'une revisite de l'article [3] et d'une remarque sur l'itération des polygones des milieux d'un polygone pair qu'on trouve [5, p. 57].

1. Un polygone est *pair* si son nombre de sommets  $n$  est pair.

## Plan de l'article

|                                                    |     |
|----------------------------------------------------|-----|
| Le théorème de Varignon et ses variantes . . . . . | 101 |
| Les $\epsilon$ -multiparallélogrammes . . . . .    | 103 |
| Références . . . . .                               | 114 |

Une « figure du plan » est une partie de celui-ci ayant une certaine particularité. Nous pouvons la voir dans son intégralité, ou nous en comprenons globalement la forme, même lorsqu'elle échappe à notre regard. Un polygone en est probablement l'exemple le plus simple : il est bordé par un nombre fini de segments. Généralement, quand les gens disent « ceci est une figure géométrique », ils dénomment ainsi tout ce qui est délimité par des segments, des arcs de cercle ou des morceaux d'une courbe familière (ellipse, parabole, hyperbole...).

Les polygones du plan abondent ; ils sont de toutes formes et de tailles arbitraires. Bien évidemment, la meilleure façon d'en donner une image est d'en dessiner ; en voici deux exemples (figure 1).

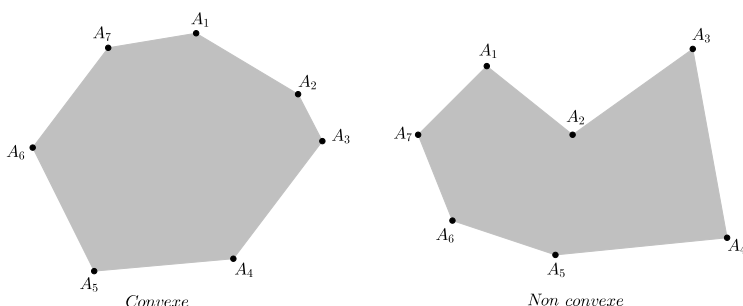


FIGURE 1

Nous commençons par rappeler les définitions essentielles et précises qui nous serviront dans ce que nous nous proposons d'exposer. Ce sera un peu abondant, mais cela est indispensable.

Soit  $E$  un plan vectoriel euclidien, c'est-à-dire un espace vectoriel réel de dimension 2, muni d'un produit scalaire  $\langle \cdot, \cdot \rangle$  auquel est associée la norme  $\|\vec{u}\| = \sqrt{\langle \vec{u}, \vec{u} \rangle}$ . Sa structure *affine canonique* permet de voir tout élément  $M$  de  $E$  aussi bien comme un *point* que comme le vecteur  $\overrightarrow{OM}$  (où  $O$  est l'origine de  $E$ ). L'écriture  $\lambda M + \mu N$  a donc parfaitement un sens, parce qu'elle désigne le vecteur  $\lambda \overrightarrow{OM} + \mu \overrightarrow{ON}$ .

Si  $A$  et  $B$  sont deux points de  $\mathbf{E}$ ,  $[A, B]$  sera le segment d'extrémités  $A$  et  $B$  donné par  $[A, B] = \{(1-t)A + tB : t \in [0, 1]\}$ . Sa longueur  $AB$  est la norme  $\|\overrightarrow{AB}\|$  du vecteur  $\overrightarrow{AB}$ .

Si  $A_1, \dots, A_n$  sont des points et  $\lambda_1, \dots, \lambda_n \in \mathbb{R}$  de somme  $\tau \neq 0$ , le point  $\frac{1}{\tau}(\lambda_1 A_1 + \dots + \lambda_n A_n)$  est appelé *barycentre* des points  $A_1, \dots, A_n$  affectés respectivement des *coefficients* (ou *poids*)  $\lambda_1, \dots, \lambda_n$ . On le note  $\text{Bar}\{(A_1, \lambda_1), \dots, (A_n, \lambda_n)\}$ . Il reste inchangé si l'on multiplie tous les coefficients  $\lambda_1, \dots, \lambda_n$  par un même réel non nul, ce qui amène souvent à supposer  $\lambda_1 + \dots + \lambda_n = 1$ .

### → p. 108 La notion d'espace affine

Pour tout entier  $n \geq 3$ , on appelle *n-polygone* (on dit aussi *n-gone* ou encore polygone d'ordre  $n$ ) de  $\mathbf{E}$  tout ensemble  $\{[A_1, A_2], \dots, [A_n, A_1]\}$  (ligne brisée décrite dans cet ordre) constitué de segments  $[A_1, A_2], \dots, [A_n, A_1]$  de  $\mathbf{E}$ . Les points  $A_1, \dots, A_n$  sont supposés distincts deux à deux.

- Un tel ensemble sera noté  $[A_1, \dots, A_n]$ . On dira que  $[A_1, \dots, A_n]$  est *pair* (resp. *impair*) si  $n$  est pair (resp. impair).
- Pour tout entier naturel  $p$ , on posera  $A_p = A_k$ , où  $k$  est le représentant dans  $\{1, \dots, n\}$  de la classe de congruence de  $p$  modulo  $n$ . En particulier,  $A_{n+1} = A_1$ .
- Il est bien connu qu'un 3-polygone, un 4-polygone, un 5-polygone, un 6-polygone..., sont appelés respectivement *triangle*, *quadrilatère*, *pentagone*, *hexagone*...
- Les points  $A_k$  et les segments  $[A_k, A_{k+1}]$  sont respectivement les *sommets* et les *côtés* (ou *arêtes*) du  $n$ -polygone  $[A_1, \dots, A_n]$ . Si  $A_k$  et  $A_\ell$  sont deux sommets non consécutifs, on dit que le segment  $[A_k, A_\ell]$  est une *diagonale* du polygone.
- Le polygone  $[A_1, \dots, A_n]$  est *convexe* si pour chaque côté  $[A_k, A_{k+1}]$ , l'un des deux demi-plans fermés de frontière la droite  $(A_k A_{k+1})$  contient tous les sommets du polygone.
- Soit  $[A_1, \dots, A_n]$  un  $n$ -polygone pair avec  $n = 2r$ . Les sommets  $A_k$  et  $A_{k+r}$  seront dits *sommets opposés*. De même, les segments  $[A_k, A_\ell]$  et  $[A_{k+r}, A_{\ell+r}]$  seront dits *segments opposés*. Une diagonale du type  $[A_k, A_{k+r}]$  sera appelée *diagonale principale*. Un quadrilatère du type  $[A_k, A_{k+1}, A_{k+r}, A_{k+r+1}]$ , avec  $1 \leq k \leq r$ , sera appelé *quadrilatère principal* (figure 2 page suivante). Un *parallélogramme* est un quadrilatère  $[A, B, C, D]$  dont les diagonales  $[A, C]$  et  $[B, D]$  se coupent

en leur milieu. Cette condition est équivalente à l'égalité vectorielle  $\overrightarrow{AB} + \overrightarrow{CD} = \vec{0}$ .

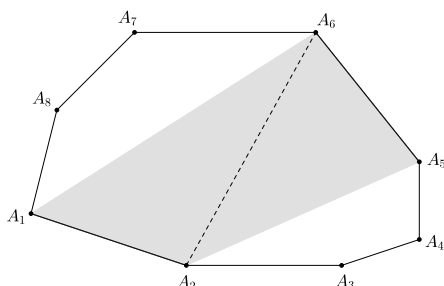


FIGURE 2 – Dans cet octogone,  $[A_2, A_6]$  est une diagonale principale et  $[A_1, A_2, A_5, A_6]$  est un quadrilatère principal.

- On appellera *multipartalléogramme* tout polygone pair  $\mathcal{P}$  dont tous les quadrilatères principaux sont des parallélogrammes (figure 3). Si ces quadrilatères sont des rectangles, on dira que  $\mathcal{P}$  est un *multirectangle*. Par exemple tout polygone régulier pair est un *multirectangle*. L'affirmation qui suit est presque immédiate à établir :

**Proposition 1**

Un polygone pair  $[A_1, \dots, A_{2r}]$  est un multipartalléogramme si, et seulement si, il existe une symétrie centrale qui envoie  $A_k$  sur  $A_{k+r}$  pour  $k = 1, \dots, r$ . Un multipartalléogramme est un multirectangle si, et seulement si, il est inscrit dans un cercle.

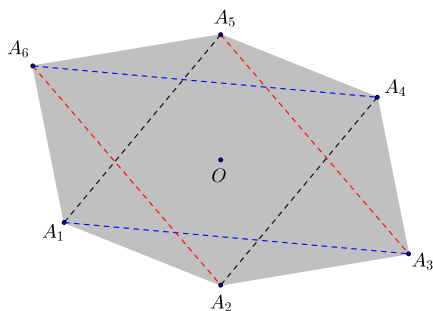


FIGURE 3

## Le théorème de Varignon et ses variantes

Voici la version classique connue pour un quadrilatère par tout le monde. Sa démonstration est presque immédiate, mais ceci n'enlève rien à la beauté et à l'intérêt géométrique de ce théorème.

### Théorème 2

Soient  $[A_1, A_2, A_3, A_4]$  un quadrilatère quelconque et  $A'_1, A'_2, A'_3$  et  $A'_4$  les milieux respectifs des côtés  $[A_1, A_2]$ ,  $[A_2, A_3]$ ,  $[A_3, A_4]$  et  $[A_4, A_1]$ . Alors  $[A'_1, A'_2, A'_3, A'_4]$  est un parallélogramme (figure 4).

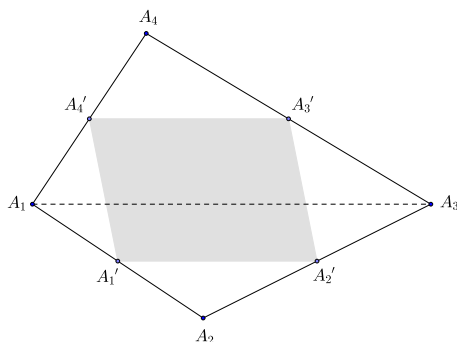


FIGURE 4

On peut généraliser cet énoncé de la façon suivante avec une démonstration similaire. Mais nous en donnerons une autre qui se transpose plus facilement au cas d'un polygone pair.

### Théorème 3

Soient  $[A_1, A_2, A_3, A_4]$  un quadrilatère et  $\lambda$  un nombre réel dans  $]0, 1[$ . On pose :

$$\begin{cases} A'_1 &= (1 - \lambda)A_1 + \lambda A_2 \\ A'_3 &= (1 - \lambda)A_3 + \lambda A_4 \\ A'_2 &= \lambda A_2 + (1 - \lambda)A_3 \\ A'_4 &= \lambda A_4 + (1 - \lambda)A_1 \end{cases}$$

Alors le quadrilatère  $[A'_1, A'_2, A'_3, A'_4]$  est un parallélogramme (figure 5 page suivante).

*Preuve.* Soit  $O = \text{Bar}\{(A_1, (1 - \lambda)), (A_2, \lambda), (A_3, (1 - \lambda)), (A_4, \lambda)\}$  le barycentre des quatre points  $A_1, A_2, A_3$  et  $A_4$  affectés respectivement des coefficients  $(1 - \lambda), \lambda, (1 - \lambda)$  et  $\lambda$ . Alors, par l'associativité des barycentres,

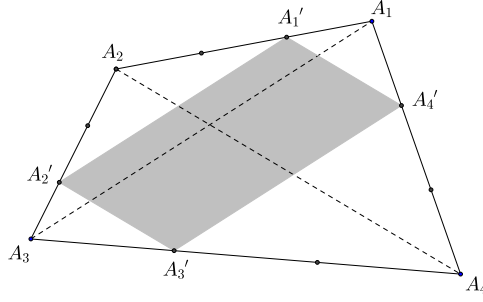


FIGURE 5

on a immédiatement :

$$O = \text{Bar}\{(A'_1, 1), (A'_3, 1)\} \text{ et } O = \text{Bar}\{(A'_2, 1), (A'_4, 1)\}$$

Le point  $O$  est donc le milieu commun des deux diagonales  $[A'_1, A'_3]$  et  $[A'_2, A'_4]$  du quadrilatère  $[A'_1, A'_2, A'_3, A'_4]$ . Par suite, ce dernier est un parallélogramme.  $\square$

Voici maintenant une généralisation du théorème de Varignon à un polygone d'ordre pair tout à fait quelconque. Elle a la forme qui suit.

**Théorème 4 (théorème général de Varignon)**

Soient  $[A_1, \dots, A_{2r}]$  un polygone pair et  $(\lambda_1, \dots, \lambda_r)$  un  $r$ -uplet d'éléments de l'intervalle  $]0, 1[$  tels que  $\lambda_1 + \dots + \lambda_r = 1$ . On pose  $\lambda_{j+r} = \lambda_j$  pour  $j = 1, \dots, r$ , et, pour tout  $k = 1, \dots, 2r$ , on note  $G_k = \text{Bar}\{(A_k, \lambda_k), \dots, (A_{k+r-1}, \lambda_{k+r-1})\}$ . Alors le polygone  $[G_1, \dots, G_{2r}]$  est un multipartallélogramme (figure 6).

*Preuve.* Nous utiliserons à cet effet l'associativité des barycentres (comme la deuxième preuve que nous avons donnée dans le cas d'un quadrilatère). On doit prouver que toutes les diagonales principales  $[G_k, G_{k+r}]$  du polygone  $[G_1, \dots, G_{2r}]$  ont le même milieu.

Notons  $G$  le barycentre  $\text{Bar}\{(A_1, \lambda_1), \dots, (A_{2r}, \lambda_{2r})\}$ . Les deux ensembles de points pondérés (où l'on a posé  $A_{k+2r-1} = A_{k-1}$ ) :

$$\begin{cases} \Gamma_k &= \{(A_k, \lambda_k), \dots, (A_{k+r-1}, \lambda_{k+r-1})\} \\ \Gamma'_k &= \{(A_{k+r}, \lambda_{k+r}), \dots, (A_{k+2r-1}, \lambda_{k+2r-1})\} \end{cases}$$

forment une partition de l'ensemble des sommets pondérés  $\Gamma = \{(A_1, \lambda_1), \dots, (A_{2r}, \lambda_{2r})\}$ . De plus,  $\Gamma_k$ , ayant un poids total égal à 1, a pour barycentre le point  $G_k$ , et  $\Gamma'_k$ , ayant également un poids total égal à 1, admet pour barycentre le point  $G_{k+r}$ . La règle d'associativité des barycentres permet

donc d'avoir la relation  $G = \text{Bar}\{ \langle G_k, 1 \rangle, \langle G_{k+r}, 1 \rangle \}$ . Toutes les diagonales principales  $[G_k, G_{k+r}]$  ont donc le point  $G$  comme milieu commun. Ce qui prouve que le polygone  $[G_1, \dots, G_{2r}]$  est bien un multiparallélogramme.  $\square$

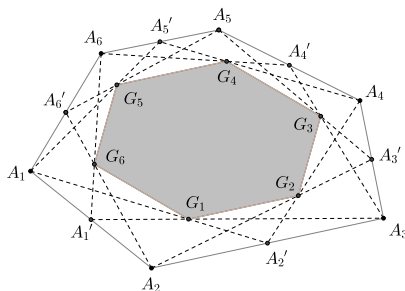


FIGURE 6 – Le polygone  $[A_1, \dots, A_6]$  est un hexagone. Pour  $k = 1, \dots, 6$ ,  $A'_k$  est le milieu de  $[A_k, A_{k+1}]$ . Donc  $G_k$  est l'isobarycentre de  $A_k, A_{k+1}$  et  $A_{k+2}$ . On obtient ainsi l'hexagone  $[G_1, \dots, G_6]$ , dont on voit clairement sur le dessin que c'est un multiparallélogramme.

## Les $\epsilon$ -multiparallélogrammes

**Définition 5.** Soit  $\epsilon$  un nombre réel strictement positif. On dira qu'un quadrilatère  $[A, B, C, D]$  est un  $\epsilon$ -parallélogramme si  $\|\overrightarrow{AB} + \overrightarrow{CD}\| \leq \epsilon$ . De même, un polygone pair  $[A_1, \dots, A_{2r}]$  sera appelé  $\epsilon$ -multiparallélogramme si tous ses quadrilatères principaux  $[A_k, A_{k+1}, A_{k+r}, A_{k+r+1}]$ , avec  $k = 1, \dots, r$ , sont des  $\epsilon$ -parallélogrammes.

**Exemple 6.** Si l'on munit  $E$  d'un repère cartésien orthonormé, les points  $A(2, 0)$ ,  $B\left(\frac{\epsilon}{\sqrt{2}}, 1 + \frac{\epsilon}{\sqrt{2}}\right)$ ,  $C(-2, 0)$  et  $D(0, -1)$  sont tels que  $\|\overrightarrow{AB} + \overrightarrow{CD}\| = \epsilon$ . Le quadrilatère  $[A, B, C, D]$  est donc bien un  $\epsilon$ -parallélogramme (figure 7 page suivante).

Considérons un polygone pair quelconque  $\mathcal{P} = [A_1, \dots, A_{2r}]$ . Pour un nombre réel  $\lambda$  fixé dans l'intervalle  $]0, 1[$ , on peut considérer le polygone  $\mathcal{P}_1 = [A_{1,1}, \dots, A_{1,2r}]$  tel que, pour chaque indice  $k$ , le sommet  $A_{1,k}$  est le barycentre  $(1-\lambda)A_k + \lambda A_{k+1}$ ; c'est le premier polygone  $\lambda$ -dérivé de  $\mathcal{P}$ . En itérant cette opération, on définit une suite de polygones  $\mathcal{P}_q = [A_{q,1}, \dots, A_{q,2r}]$  par  $\mathcal{P}_0 = \mathcal{P}$ , et  $\mathcal{P}_{q+1}$  est le premier dérivé de  $\mathcal{P}_q$  pour tout  $q \in \mathbb{N}$ . C'est la

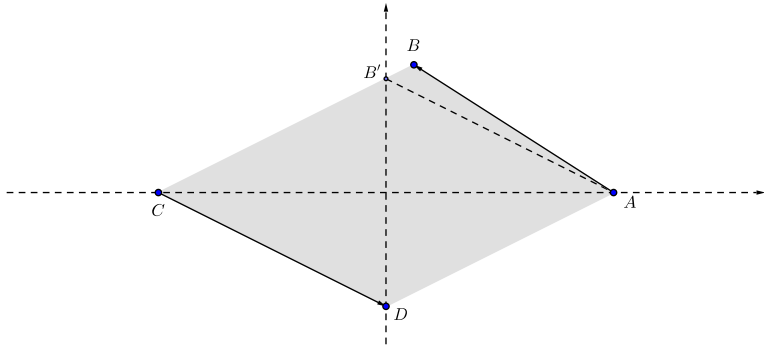


FIGURE 7

suite des  $\lambda$ -dérivés de  $\mathcal{P}$ . La figure 8 en montre une image avec  $\lambda = \frac{1}{2}$ .

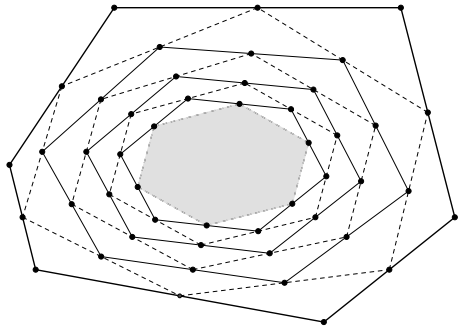


FIGURE 8

Le dessin de la figure 8 suggère que lorsque  $q$  est suffisamment grand, le  $q$ -ième dérivé de  $\mathcal{P}$  a l'aspect d'un multiparallélogramme. C'est effectivement le cas et c'est ce qu'exprime le théorème qui suit.

**Théorème 7 (théorème principal)**

Soit  $\lambda$  un nombre réel dans l'intervalle  $]0, 1[$ . Pour tout nombre réel  $\epsilon > 0$ , il existe un rang  $q_\epsilon$  tel que, pour tout  $q \geq q_\epsilon$ , le  $q$ -ième dérivé  $\mathcal{P}_q$  de  $\mathcal{P}$  est un  $\epsilon$ -multiparallélogramme.

**Remarque 8.** Certains auteurs ont étudié indépendamment le cas particulier  $\lambda = \frac{1}{2}$  par des méthodes légèrement différentes. On les trouve sur le net en faisant une recherche par mots-clés, par exemple [1].

*Preuve.* Elle est simple sur le plan des idées. Elle consiste essentiellement

en un processus de calcul matriciel du niveau de la licence, mais qu'il faut mener quand même.

Pour  $q \in \mathbf{N}$  et  $k \in \{1, \dots, r\}$ , considérons le quadrilatère principal  $[A_{q,k}, A_{q,k+1}, A_{q,k+r}, A_{q,k+r+1}]$  et posons :

$$v_{q,k} = \overrightarrow{A_{q,k}A_{q,k+1}} + \overrightarrow{A_{q,k+r}A_{q,k+r+1}}$$

Notons  $\mathcal{V}_q$  le  $r$ -uplet  $(v_{q,1}, \dots, v_{q,r})$  ; c'est un vecteur de  $\mathbf{E}^r$  dont les composantes sont les vecteurs  $\overrightarrow{A_{q,k}A_{q,k+1}} + \overrightarrow{A_{q,k+r}A_{q,k+r+1}}$  des quadrilatères principaux  $[A_{q,k}, A_{q,k+1}, A_{q,k+r}, A_{q,k+r+1}]$  du polygone  $\mathcal{P}_q$ . Nous avons :

$$\sum_{k=1}^r v_{q,k} = \vec{0} \quad \text{pour tout } q \in \mathbf{N}$$

En effet, pour chaque rang  $q \in \mathbf{N}$ , avec la relation de Chasles, nous avons :

$$\sum_{k=1}^r v_{q,k} = \sum_{k=1}^r \overrightarrow{A_{q,k}A_{q,k+1}} + \sum_{k=1}^r \overrightarrow{A_{q,k+r}A_{q,k+r+1}} = \overrightarrow{A_{q,1}A_{q,r+1}} + \overrightarrow{A_{q,r+1}A_{q,1}} = \vec{0}$$

D'autre part, pour  $q \in \mathbf{N}$  et  $k \in \{1, \dots, r\}$ , nous avons, par définition :

$$A_{q+1,k} = (1 - \lambda)A_{q,k} + \lambda A_{q,k+1} \quad \text{ou encore} \quad \overrightarrow{A_{q,k}A_{q+1,k}} = \lambda \overrightarrow{A_{q,k}A_{q,k+1}}$$

On en déduit que :

$$\begin{aligned} \overrightarrow{A_{q+1,k}A_{q+1,k+1}} &= \overrightarrow{A_{q+1,k}A_{q,k}} + \overrightarrow{A_{q,k}A_{q,k+1}} + \overrightarrow{A_{q,k+1}A_{q+1,k+1}} \\ &= (1 - \lambda)\overrightarrow{A_{q,k}A_{q,k+1}} + \lambda\overrightarrow{A_{q,k+1}A_{q,k+2}} \end{aligned}$$

On a donc :

$$\begin{aligned} v_{q+1,k} &= \overrightarrow{A_{q+1,k}A_{q+1,k+1}} + \overrightarrow{A_{q+1,k+r}A_{q+1,k+r+1}} \\ &= (1 - \lambda)\overrightarrow{A_{q,k}A_{q,k+1}} + \lambda\overrightarrow{A_{q,k+1}A_{q,k+2}} \\ &\quad + (1 - \lambda)\overrightarrow{A_{q,k+r}A_{q,k+r+1}} + \lambda\overrightarrow{A_{q,k+r+1}A_{q,k+r+2}} \\ &= (1 - \lambda)\left[\overrightarrow{A_{q,k}A_{q,k+1}} + \overrightarrow{A_{q,k+r}A_{q,k+r+1}}\right] \\ &\quad + \lambda\left[\overrightarrow{A_{q,k+1}A_{q,k+2}} + \overrightarrow{A_{q,k+r+1}A_{q,k+r+2}}\right] \\ &= (1 - \lambda)v_{q,k} + \lambda v_{q,k+1} \end{aligned}$$

Il en résulte que  $\mathcal{V}_{q+1} = \mathcal{L}(\mathcal{V}_q)$ , où  $\mathcal{L}$  est l'endomorphisme  $\mathcal{L} : \mathbf{E}^r \rightarrow \mathbf{E}^r$  qui à  $Z = (z_1, \dots, z_r)$  associe :

$$\mathcal{L}(Z) = ((1 - \lambda)z_1 + \lambda z_2, \dots, (1 - \lambda)z_{r-1} + \lambda z_r, (1 - \lambda)z_r + \lambda z_1)$$

On a donc, pour tout  $q \in \mathbf{N}^*$  :

$$\mathcal{V}_q = \mathcal{L}(\mathcal{V}_{q-1}) = \mathcal{L}(\mathcal{L}(\mathcal{V}_{q-2})) = \dots = \mathcal{L}^q(\mathcal{V}_0)$$

En identifiant  $\mathbf{E}$  au plan complexe  $\mathbf{C}$ , on peut représenter l'endomorphisme  $\mathcal{L}$  par sa matrice  $\mathcal{M}$  dans la base canonique de l'espace vectoriel complexe  $\mathbf{C}^r$ . Cette matrice s'écrit :

$$\mathcal{M} = \begin{pmatrix} 1-\lambda & \lambda & 0 & \dots & 0 \\ 0 & 1-\lambda & \lambda & \ddots & \vdots \\ \vdots & \ddots & \ddots & \ddots & 0 \\ 0 & \dots & 0 & 1-\lambda & \lambda \\ \lambda & 0 & \dots & 0 & 1-\lambda \end{pmatrix}$$

Il s'agit d'une matrice circulante de la forme :

$$\mathcal{C} = \begin{pmatrix} a_0 & a_1 & \dots & a_{r-1} \\ a_{r-1} & a_0 & \dots & a_{r-2} \\ \vdots & \vdots & \ddots & \vdots \\ a_1 & a_2 & \dots & a_0 \end{pmatrix}$$

Le déterminant d'une telle matrice est le nombre complexe ([\[4\]](#)) :

$$\det(\mathcal{C}) = P(1)P(\omega)P(\omega^2) \dots P(\omega^{r-1})$$

où  $\omega = e^{\frac{2i\pi}{r}}$  et  $P$  est le polynôme défini par  $P(X) = a_0 + a_1X + \dots + a_{r-1}X^{r-1}$ . On en déduit que si  $I_r$  est la matrice identité d'ordre  $r$  et  $\mu$  un nombre complexe quelconque, nous avons :

$$\det(\mathcal{C} - \mu I_r) = P_\mu(1)P_\mu(\omega)P_\mu(\omega^2) \dots P_\mu(\omega^{r-1})$$

où  $P_\mu(X) = (a_0 - \mu) + a_1X + \dots + a_{r-1}X^{r-1}$ , ce qui implique que les valeurs propres de  $\mathcal{C}$  sont les nombres complexes  $\mu$  vérifiant la relation :

$$P_\mu(1)P_\mu(\omega)P_\mu(\omega^2) \dots P_\mu(\omega^{r-1}) = 0$$

Dans le cas de la matrice circulante  $\mathcal{M}$  où  $a_0 = 1 - \lambda$ ,  $a_1 = \lambda$  et  $a_j = 0$  pour  $j \in \{2, \dots, r-1\}$ , les valeurs propres sont les nombres complexes  $\mu$  tels que  $P_\mu(\omega^j) = (1 - \lambda - \mu) + \lambda\omega^j = 0$  pour  $j \in \{0, 1, \dots, r-1\}$ , c'est-à-dire les nombres complexes  $\mu_j = 1 - \lambda + \lambda\omega^j$  pour  $j \in \{0, 1, \dots, r-1\}$ . Ces valeurs propres étant deux à deux distinctes, la matrice  $\mathcal{M}$  est diagonalisable sur  $\mathbf{C}$  et l'on peut donc la décomposer sous la forme :

$$\mathcal{M} = \mathcal{N}^{-1} \times \text{diag}(1, \mu_1, \dots, \mu_{r-1}) \times \mathcal{N} \text{ où } \mathcal{N} \in \text{GL}_r(\mathbf{C})$$

Ce qui implique que :

$$\forall q \in \mathbf{N}, \mathcal{M}^q = \mathcal{N}^{-1} \times \text{diag}(1, \mu_1^q, \dots, \mu_{r-1}^q) \times \mathcal{N}$$

avec  $|\mu_j| = \left| (1 - \lambda) + \lambda e^{\frac{2i\pi j}{r}} \right| < 1$  pour  $1 \leq j \leq r - 1$ . En effet, comme  $e^{\frac{2i\pi j}{r}} = \cos \frac{2\pi j}{r} + i \sin \frac{2\pi j}{r}$  :

$$\begin{aligned} |\mu_j|^2 &= \left[ (1 - \lambda) + \lambda \cos \frac{2\pi j}{r} \right]^2 + \lambda^2 \left( \sin \frac{2\pi j}{r} \right)^2 \\ &= (1 - \lambda)^2 + 2(1 - \lambda)\lambda \cos \frac{2\pi j}{r} + \lambda^2 \\ &= 1 - 2(1 - \lambda)\lambda \left( 1 - \cos \frac{2\pi j}{r} \right) \\ &< 1 \end{aligned}$$

La suite  $(\mathcal{M}^q)_{q \in \mathbf{N}}$  converge donc vers la matrice :

$$\mathcal{M}^* = \mathcal{N}^{-1} \times \text{diag}(1, 0, \dots, 0) \times \mathcal{N}$$

et la suite vectorielle  $(\mathcal{V}_q) = (\mathcal{M}^q(\mathcal{V}_0))$  converge vers  $\mathcal{V}_* = \mathcal{M}^*(\mathcal{V}_0)$ . D'autre part, par passage aux limites dans la relation de récurrence  $\mathcal{M}(\mathcal{V}_q) = \mathcal{V}_{q+1}$ , on obtient  $\mathcal{M}(\mathcal{V}_*) = \mathcal{V}_*$ . Ce qui implique que  $\mathcal{V}_*$  est un élément de la droite vectorielle  $\text{Ker}(\mathcal{M} - I_r) = \text{Vect}\{(1, \dots, 1)\}$ . Il existe donc un nombre complexe  $\alpha$  tel que  $\mathcal{V}_* = (\alpha, \dots, \alpha)$ . De plus, si  $\varphi : \mathbf{C}^r \rightarrow \mathbf{C}$  est la forme linéaire (continue bien sûr) définie par :

$$\varphi(z_1, \dots, z_r) = z_1 + \dots + z_r$$

on a  $\forall q \in \mathbf{N}, \varphi(\mathcal{V}_q) = 0$ . Donc, en passant à la limite, on obtient  $\varphi(\mathcal{V}_*) = 0$ , soit  $r\alpha = 0$ . D'où  $\alpha = 0$  et  $\mathcal{V}_* = (0, \dots, 0)$ .

La suite  $(\mathcal{V}_q)$  étant convergente vers  $\mathcal{V}_* = (0, \dots, 0)$ , on en déduit que pour tout  $\epsilon > 0$ , il existe  $q_\epsilon$  tel que  $\|v_{q,k}\| < \epsilon$  pour  $q \geq q_\epsilon$  et  $k \in \{1, \dots, r\}$ . Ce qui prouve que  $\mathcal{P}_q$  est un  $\epsilon$ -multiparallélogramme dès que  $q \geq q_\epsilon$ .  $\square$

Dans [4, p. 57], il est aussi affirmé (sans démonstration) que même si le polygone de départ n'est pas contenu dans un plan, ses  $\lambda$ -dérivés successifs se rapprochent d'un multiparallélogramme plan. Cela semble vrai, et nous pensons que notre démonstration pourrait être transposée presque telle quelle après avoir formulé le bon énoncé. C'est un joli problème qui intéresserait probablement certains lecteurs d'*Images des mathématiques*. Bon vent aux amateurs !

**Remarque 9.** Le polygone  $\mathcal{P}_1$ ,  $\lambda$ -dérivé de  $\mathcal{P} = \mathcal{P}_0$ , a le même centre de gravité  $G$  que  $\mathcal{P}$ , car nous avons la relation vectorielle :

$$\begin{aligned} \sum_{k=1}^n \overrightarrow{GA_{1,k}} &= \sum_{k=1}^n \left( (1-\lambda) \overrightarrow{GA_k} + \lambda \overrightarrow{GA_{k+1}} \right) \\ &= (1-\lambda) \underbrace{\sum_{k=1}^n \overrightarrow{GA_k}}_{=\vec{0}} + \lambda \underbrace{\sum_{k=1}^n \overrightarrow{GA_{k+1}}}_{=\vec{0}} = \vec{0} \end{aligned}$$

Ce qui implique (par récurrence) que le polygone  $\mathcal{P}$  et tous ses  $\lambda$ -dérivés successifs  $\mathcal{P}_q$  ont le même isobarycentre  $G$  et qu'on a ainsi pour chaque polygone  $\mathcal{P}_q$  la relation vectorielle :  $\sum_{k=1}^n \overrightarrow{GA_{q,k}} = \vec{0}$ . Ceci permet de montrer, exactement comme dans la preuve du théorème principal, que la suite  $(\mathcal{P}_q)$  des  $\lambda$ -dérivés converge vers le  $n$ -polygone trivial  $(G, \dots, G)$  réduit à  $G$ .

## Complément

### La notion d'espace affine

Le contenu de cette section est constitué de fragments en partie extraits du livre [2]. Ceux qui souhaitent en savoir plus y trouveront un développement assez fourni sur le sujet.

### Préliminaires

On se donne un espace vectoriel réel, c'est-à-dire un ensemble  $E$  non vide muni :

- d'une *addition*  $\vec{u} + \vec{v}$  vérifiant les propriétés :
  1.  $\vec{u} + \vec{v} = \vec{v} + \vec{u}$  (*commutativité*),
  2.  $(\vec{u} + \vec{v}) + \vec{w} = \vec{u} + (\vec{v} + \vec{w})$  (*associativité*),
  3. il existe un élément  $\vec{0}$  (*élément neutre*) tel que  $\vec{u} + \vec{0} = \vec{u}$  pour tout  $\vec{u} \in E$ ,
  4. pour tout  $\vec{u} \in E$ , il existe un élément  $\vec{u}'$  tel que  $\vec{u} + \vec{u}' = \vec{0}$ ; on le note  $-\vec{u}$  et on l'appelle *opposé de  $\vec{u}$* ;
- d'une *multiplication par les réels*  $\lambda \vec{u}$  vérifiant les propriétés :
  1.  $1 \cdot \vec{u} = \vec{u}$  pour tout  $\vec{u} \in E$ ,
  2.  $\lambda(\mu \vec{u}) = (\lambda\mu) \vec{u}$  pour tous  $\lambda, \mu \in \mathbf{R}$  et tout  $\vec{u} \in E$ ,

- 3.  $(\lambda + \mu)\vec{u} = \lambda\vec{u} + \mu\vec{u}$  pour tous  $\lambda, \mu \in \mathbf{R}$  et tout  $\vec{u} \in \mathbf{E}$ ,
- 4.  $\lambda(\vec{u} + \vec{v}) = \lambda\vec{u} + \lambda\vec{v}$  pour tout  $\lambda \in \mathbf{R}$  et tous  $\vec{u}, \vec{v} \in \mathbf{E}$ .

Un élément  $\vec{u}$  de  $\mathbf{E}$  est appelé *vecteur*. Dans le bon vieux temps, on nous le définissait comme un *segment orienté* : une *force* s'exerçant sur les éléments  $M$  (appelés *points*) d'un ensemble  $\mathcal{E}$  dans une direction, un sens, et avec une certaine intensité. Ceci en respectant certaines règles qui décrivent de façon géométrique les axiomes définissant l'espace  $\mathbf{E}$ . Nous nous contentons d'en montrer quelques-unes sur les dessins ci-dessous (figures 9 à 11 pages 109 et 110).

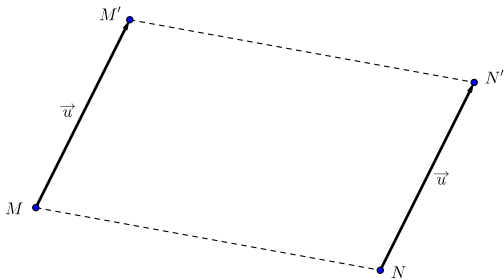


FIGURE 9 – Le vecteur  $\vec{u}$  agit de la même manière sur tous les éléments de  $\mathcal{E}$ .

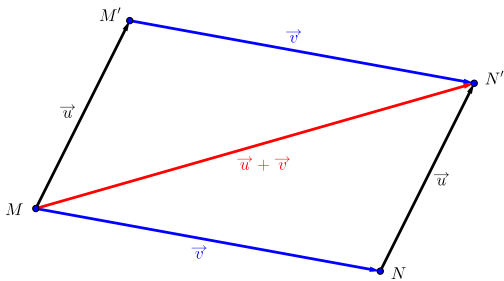


FIGURE 10 – On voit bien la commutativité de l'addition  $\vec{u} + \vec{v} = \vec{v} + \vec{u}$ .

Le vecteur  $\vec{0}$  représente une force nulle et n'a donc aucun effet sur les points  $M \in \mathcal{E}$ . Et c'est le seul ayant cette propriété.

On a donc une *action* de l'espace vectoriel  $\mathbf{E}$  sur l'ensemble  $\mathcal{E}$ , c'est-à-dire une application  $\Psi : \mathbf{E} \times \mathcal{E} \rightarrow \mathcal{E}$  qui au couple  $(\vec{u}, M)$  associe son déplacé

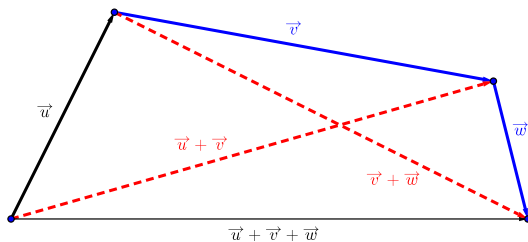


FIGURE 11 – On voit l’associativité de l’addition  $(\vec{u} + \vec{v}) + \vec{w} = \vec{u} + (\vec{v} + \vec{w})$ . Ce qui autorise l’écriture  $\vec{u} + \vec{v} + \vec{w}$  de la somme de trois vecteurs.

$M' = \Psi(\vec{u}, M)$  (qu’on note aussi  $M + \vec{u}$ ) tel qu’on le voit par exemple sur le premier dessin. Cette action possède les propriétés suivantes :

1.  $M + \vec{0} = M$  pour tout  $M \in \mathcal{E}$  ;
2.  $M + (\vec{u} + \vec{v}) = (M + \vec{u}) + \vec{v}$  pour tous  $\vec{u}, \vec{v} \in \mathbf{E}$  et tout  $M \in \mathcal{E}$  ;
3.  $M + \vec{u} = M$  implique  $\vec{u} = \vec{0}$  (simplicité de l’action) ;
4. pour tous  $M, M'$  de  $\mathcal{E}$ , il existe  $\vec{u} \in \mathbf{E}$  tel que  $M + \vec{u} = M'$  (transitivité de l’action).

## Espaces affines

Voici la première définition d’un espace affine que suggère tout ce qu’on vient de voir (les dessins montrent qu’on ne peut pas faire autrement!).

**Définition 10.** Un espace affine réel est la donnée d’un couple  $(\mathcal{E}, \Psi)$  où  $\mathcal{E}$  est un ensemble non vide et  $\Psi : (\vec{u}, M) \in \mathbf{E} \times \mathcal{E} \mapsto M + \vec{u} \in \mathcal{E}$  une action de  $\mathbf{E}$  sur  $\mathcal{E}$  possédant les propriétés (1), (2), (3) et (4) énumérées ci-dessus. L’espace vectoriel  $\mathbf{E}$  est appelé direction de  $\mathcal{E}$ .

Il y a aussi une deuxième définition dont nous laisserons le soin au lecteur de montrer qu’elle est équivalente à la première.

**Définition 11.** Un espace affine réel dirigé par  $\mathbf{E}$  est la donnée d’un couple  $(\mathcal{E}, \Phi)$  où  $\mathcal{E}$  est un ensemble non vide et  $\Phi : (M, N) \in \mathcal{E} \times \mathcal{E} \mapsto \overrightarrow{MN} \in \mathbf{E}$  une application possédant les propriétés suivantes :

1.  $\overrightarrow{ML} + \overrightarrow{LN} = \overrightarrow{MN}$  pour tous  $M, N, L \in \mathcal{E}$  ; c’est la relation de Chasles ;
2. pour tout  $\omega \in \mathcal{E}$  fixé, l’application partielle  $\Phi_\omega : M \in \mathcal{E} \mapsto \overrightarrow{\omega M} \in \mathbf{E}$  est une bijection.

À toute application  $f : \mathcal{E} \rightarrow \mathcal{E}$  est associée de façon naturelle l'application  $\vec{f} : \mathbf{E} \rightarrow \mathbf{E}$  définie par  $\vec{f}(\overrightarrow{MN}) = \overrightarrow{f(M)f(N)}$ . On dira que  $f$  est *affine* si  $\vec{f}$  est linéaire ; dans ce cas  $\vec{f}$  est appelée *direction* de  $f$ . Si  $f$  est affine et bijective, on dira que  $f$  est un *automorphisme affine* de  $\mathcal{E}$ . L'ensemble  $\text{Aff}(\mathcal{E})$  des automorphismes affines de  $\mathcal{E}$  est un groupe pour la composition des applications.

Si l'on fixe  $\vec{u} \in \mathbf{E}$ , l'application partielle  $\Psi_{\vec{u}} : M \in \mathcal{E} \mapsto M + \vec{u} \in \mathcal{E}$  est une bijection. L'application  $\vec{u} \mapsto \Psi_{\vec{u}}$  est un morphisme injectif du groupe additif  $(\mathbf{E}, +)$  dans le groupe des bijections de  $\mathcal{E}$ . La transformation  $\Psi_{\vec{u}}$  est appelée *translation* de vecteur  $\vec{u}$ . Les translations de  $\mathcal{E}$  forment un sous-groupe du groupe  $\text{Aff}(\mathcal{E})$ .

**Remarque 12.** — Pour tout  $M \in \mathcal{E}$ , on a  $\overrightarrow{MM} + \overrightarrow{MM} = \overrightarrow{MM}$ , donc  $\overrightarrow{MM} = \vec{0}$ .

- Pour tous  $M, N \in \mathcal{E}$ , on a  $\overrightarrow{MN} + \overrightarrow{NM} = \overrightarrow{MM} = \vec{0}$ , donc  $\overrightarrow{MN} = -\overrightarrow{NM}$ .
- $\overrightarrow{MN} = \vec{0}$  si et seulement si  $M = N$ .
- $\overrightarrow{MN} = \overrightarrow{M'N'}$  si et seulement si  $\overrightarrow{MM'} = \overrightarrow{NN'}$  : c'est la *règle du parallélogramme*.

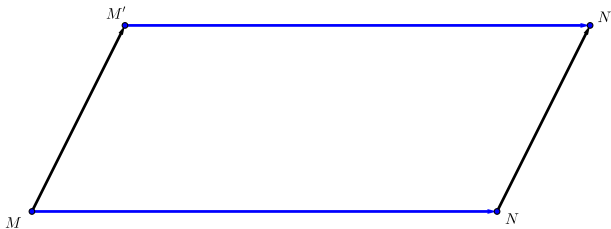


FIGURE 12

Dans toute la suite, ce qu'on entendra par espace affine sera la donnée de l'espace vectoriel  $\mathbf{E}$ , d'un ensemble  $\mathcal{E}$  et d'une application  $\Phi : \mathcal{E} \times \mathcal{E} \rightarrow \mathbf{E}$  ou d'une action  $\Psi : (\vec{u}, M) \in \mathbf{E} \times \mathcal{E} \mapsto M + \vec{u} \in \mathcal{E}$  possédant les propriétés que nous avons décrites pour chacune d'elles. Reste à savoir ce qu'on peut prendre pour ensemble  $\mathcal{E}$  ; voici quelques exemples.

**Exemple 13.** 1. On peut prendre  $\mathcal{E} = \mathbf{E}$  et  $\Phi(M, N) = N - M$  comme application  $\Phi : \mathcal{E} \times \mathcal{E} \rightarrow \mathbf{E}$ . On munit ainsi  $\mathcal{E} = \mathbf{E}$  d'une structure d'espace affine appelée *structure affine canonique* sur l'espace vectoriel  $\mathbf{E}$ .

2. Soient  $\mathcal{E}$  un ensemble quelconque et  $f : \mathcal{E} \rightarrow \mathbf{E}$  une bijection. On définit  $\Phi : \mathcal{E} \times \mathcal{E} \rightarrow \mathbf{E}$  par  $\Phi(M, N) = f(N) - f(M)$ . Il est alors facile de voir que  $\Phi$  munit  $\mathcal{E}$  d'une structure d'espace affine. Si  $g : \mathcal{E} \rightarrow \mathbf{E}$  est une autre bijection, alors elle définit la même structure affine si, et seulement si,  $g \circ f^{-1}$  est une translation (la démonstration est laissée en exercice au lecteur).

## Barycentres

On se situe dans un espace affine  $(\mathcal{E}, \mathbf{E})$ . Soient  $M_1, \dots, M_n$  des points de  $\mathcal{E}$ ,  $\lambda_1, \dots, \lambda_n$  des réels de somme  $\tau$  et  $\omega \in \mathcal{E}$ . On suppose  $\tau \neq 0$ .

### Proposition 14

Le point  $\omega + \frac{1}{\tau}(\lambda_1 \overrightarrow{\omega M_1} + \dots + \lambda_n \overrightarrow{\omega M_n})$  ne dépend pas de  $\omega$ . On l'appelle barycentre de  $M_1, \dots, M_n$  affectés respectivement des coefficients  $\lambda_1, \dots, \lambda_n$ . On le note  $\frac{1}{\tau}(\lambda_1 M_1 + \dots + \lambda_n M_n)$  ou  $\text{Bar}\{(M_1, \lambda_1), \dots, (M_n, \lambda_n)\}$ .

- Remarque 15.**
1. Si  $\lambda_i = 0$ , le point  $M_i$  n'intervient pas. On peut donc supposer que chacun des  $\lambda_i$  est non nul. On dira alors que le couple  $(M_i, \lambda_i)$  est un *point massique*.
  2. Le point  $G = \text{Bar}\{(M_1, \lambda_1), \dots, (M_n, \lambda_n)\}$  ne change pas si l'on multiplie tous les  $\lambda_i$  par un même réel non nul  $s$ . On peut donc supposer que  $\tau = \lambda_1 + \dots + \lambda_n = 1$ .
  3. Si les  $\lambda_i$  sont tous égaux (et non nuls), on dira que  $G$  est l'*isobarycentre* ou le *centre de gravité* des points  $M_1, \dots, M_n$ .

Détermination pratique des barycentres. On se donne toujours des points  $M_1, \dots, M_n$  et  $\lambda_1, \dots, \lambda_n$  des réels. On pose  $\tau = \lambda_1 + \dots + \lambda_n$ , qu'on suppose non nul. Le lecteur est invité à établir les assertions qui suivent.

1. Origine quelconque  $\omega : G = \frac{1}{\tau} \sum_{i=1}^n \lambda_i M_i \iff \overrightarrow{\omega G} = \frac{1}{\tau} \sum_{i=1}^n \lambda_i \overrightarrow{\omega M_i}$ .
2. Origine en  $G : G = \frac{1}{\tau} \sum_{i=1}^n \lambda_i M_i \iff \sum_{i=1}^n \lambda_i \overrightarrow{GM_i} = \vec{0}$ .
3. Origine en  $M_j : G = \frac{1}{\tau} \sum_{i=1}^n \lambda_i M_i \iff \overrightarrow{M_j G} = \frac{1}{\tau} \sum_{i \neq j} \lambda_i \overrightarrow{M_j M_i}$ .

La proposition qui suit dit comment on détermine le barycentre d'un ensemble de points massiques en utilisant une partition de celui-ci.

### Proposition 16

Soient  $(M_1, \lambda_1), \dots, (M_n, \lambda_n)$   $n$  points massiques. On se donne une partition de  $\{1, \dots, n\}$  par des parties  $J_1 = \{1, \dots, n_1\}, \dots, J_k = \{n_{k-1} + 1, \dots, n_k\}$  où  $n_k = n$ . On suppose que les nombres  $\lambda_1, \dots, \lambda_n$  sont tels que  $\lambda_1 + \dots + \lambda_n \neq 0$  et  $\tau_\ell = \sum_{j \in J_\ell} \lambda_j \neq 0$  pour tout  $\ell \in \{1, \dots, k\}$ . On note  $G_\ell$  le barycentre des points massiques  $\{(M_j, \lambda_j) : j \in J_\ell\}$ .

Alors le barycentre  $G$  des points massiques  $(M_1, \lambda_1), \dots, (M_n, \lambda_n)$  est égal à celui des points massiques  $\{(G_\ell, \tau_\ell) : \ell \in \{1, \dots, k\}\}$ . C'est l'associativité du calcul du barycentre.

**Exemple 17.** Détermination du centre de gravité  $G$  d'un triangle (figure 13), d'un quadrilatère (figure 14) et d'un pentagone (figure 15 page suivante).

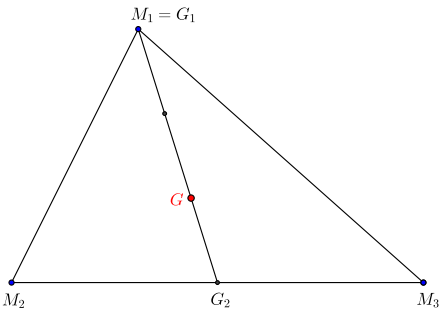


FIGURE 13 –  $G = \text{Bar}\{(M_1, 1), (M_2, 1), (M_3, 1)\} = \text{Bar}\{(G_1, 1), (G_2, 2)\}$ . La partition est  $\{M_1\} \cup \{M_2, M_3\}$ ;  $G_1 = \text{Bar}\{(M_1, 1)\} = M_1$ .

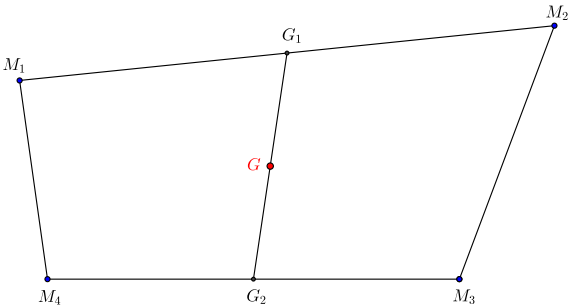


FIGURE 14 –  $G = \text{Bar}\{(M_1, 1), (M_2, 1), (M_3, 1), (M_4, 1)\}$  ou tout aussi bien  $G = \text{Bar}\{(G_1, 2), (G_2, 2)\}$ .

↑ page 99

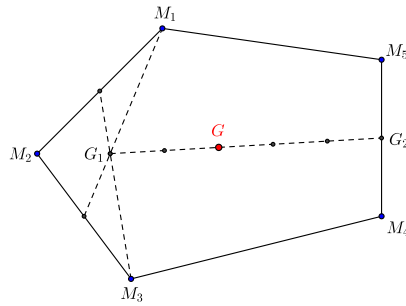


FIGURE 15 –  $G = \text{Bar}\{(M_1, 1), (M_2, 1), (M_3, 1), (M_4, 1), (M_5, 1)\}$  ou aussi bien  $G = \text{Bar}\{(G_1, 3), (G_2, 2)\}$ .

## Références

- [1] Baptiste CALMÈS. « Itération du polygone des milieux ». *Prépublication LML, Université d'Artois* (2022).
- [2] Aziz EL KACIMI. *Géométrie euclidienne élémentaire*. Éditions Ellipses, 2012.
- [3] Aziz EL KACIMI. « Le problème de Sin Pan ». *Images des mathématiques* (2014). DOI : [10.60868/8a1t-vg55](https://doi.org/10.60868/8a1t-vg55).
- [4] Jean FRESNEL et Michel MATIGNON. *Algèbre et géométrie*. Éditions Ellipses, 2017.
- [5] David WELLS. *Le dictionnaire Penguin des curiosités géométriques*. Éditions Eyrolles, 1996.

## Remerciements

Nous remercions Samuel Sendera pour sa relecture détaillée et pour nous avoir signalé des points à corriger ; cela nous a permis d'améliorer la rédaction de l'article. Merci aussi à Jonathan Chappelon (qui en est l'éditeur) et Sébastien Perrono pour leurs relectures soignées.



**Aziz EL KACIMI**

Professeur émérite – Université  
polytechnique Hauts-de-France.

**Abdellatif ZEGGAR**  
Maître de conférences – Université  
polytechnique Hauts-de-France.



Article édité par Jonathan Chappelon.