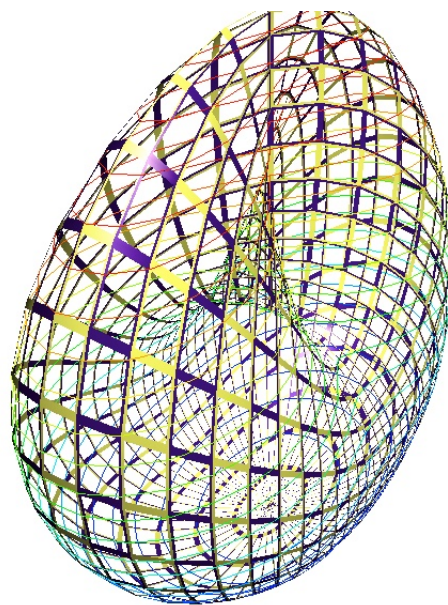




Introduction à la Géométrie Différentielle

Marc Troyanov



Avant-Propos

La géométrie différentielle étudie les courbes et les surfaces dans le plan et l'espace, et plus généralement les variétés différentiables. Dans ce domaine, nous appliquons les techniques du calcul différentiel et intégral à divers objets géométriques, nous permettant ainsi d'explorer leurs propriétés par des méthodes analytiques.

Ce polycopié accompagne le cours de géométrie différentielle 1 du programme de 2ème année du bachelor en mathématiques à l'EPFL. Tout au long de l'année, des exercices hebdomadaires viendront compléter ce document. Ces exercices sont une partie intégrante du cours : ils constituent une part fondamentale des compétences et connaissances que vous devrez maîtriser.

Votre participation active est essentielle ; je vous encourage vivement à partager vos remarques, questions ou corrections éventuelles sur le forum dédié. Vos retours sont précieux pour améliorer la qualité du polycopié et du cours en général. De plus, d'autres ressources et documents seront mis à disposition sur le site Moodle du cours.

Je vous souhaite une enrichissante découverte de la géométrie différentielle.

Marc Troyanov,
septembre 2024

Voici quelques références récentes, parmi d'autres possibles, sur le sujet de ce cours. Il existe aussi d'excellentes références plus classiques, à commencer par le traité en 3 volumes de Gaston-Darboux publiés entre 1887 et 1896.

1. Kobayashi, Shoshichi *Differential geometry of curves and surfaces*. Springer Undergraduate Mathematics Series. Springer, Singapore, 2021.
2. Needham, Tristan *Visual differential geometry and forms, a mathematical drama in five acts*. Princeton University Press, Princeton, NJ, 2021.
3. Toponogov, Victor Andreevich *Differential geometry of curves and surfaces*. Birkhäuser Boston, Inc., Boston, MA, 2006.
4. Umehara, Masaaki ; Yamada, Kotaro *Differential geometry of curves and surfaces*. World Scientific Publishing Co. Pte. Ltd., Hackensack, NJ, 2017.
5. Marc Troyanov *Cours de Géométrie*. Presses Polytechniques et Universitaires Romandes (PPUR), 2009.

Table des matières

1	Rappels sur les espaces euclidiens	2
1.1	Définitions de bases	2
1.2	Orientation d'un espace vectoriel réel de dimension finie	4
1.3	Similitudes et isométries d'un espace vectoriel euclidien.	5
1.4	Le groupe orthogonal	7
1.5	Un théorème d'Euler	8
1.6	Géométrie vectorielle dans l'espace euclidien orienté $\mathbb{E}^3$	10
1.7	Géométrie vectorielle dans le plan euclidien orienté $\mathbb{E}^2$	11
2	Courbes dans l'espace et le plan euclidien	14
2.1	Qu'est ce qu'une courbe?	14
2.2	Notions fondamentales	14
2.3	Champs de vecteurs le long d'une courbe	18
2.4	Longueur et abscisse curviligne	20
2.5	Changement de paramétrisation d'une courbe	22
2.6	Quantités géométriques et quantités cinématiques	24
2.7	Paramétrisation naturelle d'une courbe régulière	25
2.8	Courbure d'une courbe de $\mathbb{R}^n$	27
2.9	Contact entre deux courbes	29
2.10	Le repère de Frenet d'une courbes dans $\mathbb{R}^3$	31
	2.10.1 Variation angulaire du plan osculateur	34
	2.10.2 Courbes de pente constante	34
	2.10.3 Le théorème fondamental de la théorie des courbes de $\mathbb{R}^3$	35
2.11	Courbes dans un plan orienté	37
2.12	Le théorème des quatre sommets	41
3	Calcul différentiel et sous-variétés	43
3.1	Rappels de calcul différentiel	43
	3.1.1 Dérivées directionnelles et dérivées partielles	43
	3.1.2 Applications de classe C^k sur un ouvert de $\mathbb{R}^m$	44
	3.1.3 Applications Différentiables au sens de Fréchet	45
	3.1.4 Une autre interprétation de la différentielle	48
	3.1.5 Le théorème d'inversion locale	49
	3.1.6 Le théorème du rang constant	50
3.2	Sous-Variétés de $\mathbb{R}^n$	53
3.3	L'espace tangent à une sous-variété	55
3.4	Applications différentiables entre deux sous-variétés	58
3.5	Le fibré tangent à une sous-variété	60

4	Géométrie des sous-variétés	61
4.1	Distances extrinsèque et intrinsèque sur une sous-variété	62
4.2	Le tenseur métrique associé à une paramétrisation locale	63
4.3	Signification géométrique du tenseur métrique	67
4.4	Sur les isométries entre sous-variétés paramétrées	69
4.5	Intégration sur une sous-variété	71
4.6	Domaines riemanniens	72
5	Les surfaces et leur courbure	74
5.1	Co-orientation d'une surface et application de Gauss	74
5.2	Courbure d'une courbe tracée sur une surface	75
5.2.1	Géodésiques	75
5.2.2	Repère de Darboux, courbures normale et géodésique	75
5.2.3	Le théorème de Meusnier	78
5.3	L'application de Weingarten et la deuxième forme fondamentale	78
5.4	Les différentes courbures d'une surface	81
5.4.1	La courbure normale	81
5.4.2	Courbures principales, moyenne et de Gauss	81
5.4.3	Interprétation locale des courbures principales.	83
5.4.4	Courbure des surfaces de révolution	84
5.5	Quelques théorèmes classiques de la théorie des surfaces	87
A	Notions de topologie et espaces vectoriels normés	89
A.1	Rappels de topologie	89
A.2	Rappels sur la notion de norme	91
B	Sur les notations classiques	93
C	Les Symboles de Christoffel	96
C.1	Les symboles de Christoffel	96
C.2	Accélération des courbes tracées sur une surface	98
C.3	Preuve du Théorème Egregium	99
C.4	Les équations de Codazzi-Mainardi	100
D	Formulaire	102

Chapitre 1

Rappels sur les espaces vectoriels euclidiens et pseudo-euclidiens

1.1 Définitions de bases

Définitions. (i) Un *espace vectoriel euclidien* est un espace vectoriel de dimension finie sur le corps des réel muni d'un produit scalaire. On notera génériquement un tel espace par $(\mathbb{E}^n, \langle \cdot, \cdot \rangle)$, où $n \in \mathbb{N}$ est la dimension de l'espace vectoriel et $\langle \cdot, \cdot \rangle$ est le produit scalaire. Rappelons qu'un produit scalaire est une forme bilinéaire, symétrique et définie-positive sur l'espace vectoriel $\mathbb{E}^n$.

En particulier, le *produit scalaire standard* sur $\mathbb{R}^n$ est défini par

$$\langle x, y \rangle = x_1 y_1 + \cdots + x_n y_n = \sum_{i=1}^n x_i y_i.$$

(ii) La *norme* d'un vecteur $x \in \mathbb{E}^n$ est le nombre réel défini par

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

La norme est bien définie car $\langle x, x \rangle \geq 0$ pour tout x .

(iii) Une base $\{e_1, \dots, e_n\}$ de l'espace vectoriel euclidien $x \in \mathbb{E}^n$ est dite *orthonormée* si $\langle e_i, e_j \rangle = \delta_{ij}$ (le symbole de Kronecker). Cela signifie que les vecteurs de cette base sont de norme 1 et qu'ils sont deux à deux orthogonaux.

Il est facile de démontrer l'existence de bases orthonormées

Le produit scalaire peut se retrouver à partir de la norme en utilisant la *formule de polarisation* :

$$\langle x, y \rangle = \frac{1}{4} (\|x + y\|^2 - \|x - y\|^2).$$

Les deux identités suivantes sont également utiles :

$$\begin{aligned} \langle x, y \rangle &= \frac{1}{2} (\|x + y\|^2 - \|x\|^2 - \|y\|^2) \\ &= \frac{1}{2} (\|x\|^2 + \|y\|^2 - \|x - y\|^2). \end{aligned}$$

Le résultat suivant est une propriété fondamentale des produits scalaires.

Proposition 1.1. (*Inégalité de Cauchy-Schwartz.*) Pour tous vecteurs x, y de l'espace euclidien $\mathbb{E}^n$ on a

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

De plus on a égalité si et seulement si x et y sont colinéaires.

Preuve. On pose $a = \langle x, y \rangle$, que l'on suppose non nul (sinon le résultat est trivial), et $p(t) = \|tax + y\|^2$. On calcule en utilisant les propriétés du produit scalaire :

$$p(t) = \|tax + y\|^2 = \langle tax + y, tax + y \rangle = \|x\|^2 a^2 t^2 + 2a^2 t + \|y\|^2,$$

Ainsi $p(t)$ est un polynôme à coefficients réel du second degré qui est ≥ 0 pour tout $t \in \mathbb{R}$. Par conséquent le discriminant $\Delta = 4|a|^2 (|a|^2 - \|x\|^2 \|y\|^2)$ doit être négatif ou nul, c'est-à-dire $|a| \leq \|x\| \|y\|$. De plus on a égalité si et seulement s'il existe $t \in \mathbb{R}$ tel que $y = -tax$. □

Proposition 1.2. *La norme vérifie les propriétés suivantes pour tous $x, y \in \mathbb{E}^n$ et $\lambda \in \mathbb{R}$:*

(a) $\|x\| \geq 0$ et $\|x\| = 0$ si et seulement si $x = 0$.

(b) $\|\lambda x\| = |\lambda| \|x\|$.

(c) $\|x \pm y\| \leq \|x\| + \|y\|$.

Preuve. Les deux premières propriétés suivent facilement des définitions. La troisième propriété est une conséquence de l'inégalité de Cauchy-Schwartz :

$$\|x \pm y\|^2 = \|x\|^2 \pm 2\langle x, y \rangle + \|y\|^2 \leq \|x\|^2 + 2\|x\| \|y\| + \|y\|^2 = (\|x\| + \|y\|)^2.$$

Comme les normes de x , y et $x + y$ sont positives ou nulles, on peut prendre la racine carrée dans l'inégalité ci-dessus, ce qui nous donne $\|x \pm y\| \leq \|x\| + \|y\|$. □

Définitions. Dans un espace vectoriel euclidien :

(1.) La *distance* entre deux éléments x et y de $\mathbb{E}^n$ est définie par

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

(2.) L'*angle* $\alpha \in [0, \pi]$ entre deux vecteurs non nuls $x, y \in \mathbb{E}^n$ est défini par

$$\cos(\alpha) = \frac{\langle x, y \rangle}{\|x\| \|y\|}.$$

Cette notion est bien définie car d'une part $\|x\| \|y\| \neq 0$ lorsque x et y sont non nuls et d'autre part on a

$$-1 \leq \frac{\langle x, y \rangle}{\|x\| \|y\|} \leq +1$$

par l'inégalité de Cauchy-Schwartz. Notons que le produit scalaire est parfois défini géométriquement à partir de la notion d'angle via la formule

$$\langle x, y \rangle = \|x\| \|y\| \cos(\alpha),$$

mais du point de vue de l'algèbre linéaire, c'est le produit scalaire qui est la notion de base et l'angle est une notion dérivée.

(3.) L'*aire* du parallélogramme $\mathcal{P}(x, y)$ construit sur les vecteurs x et y est définie par

$$\text{Aire}(\mathcal{P}(x, y)) = \sqrt{\|x\|^2 \|y\|^2 - \langle x, y \rangle^2}.$$

A nouveau, l'inégalité de Cauchy-Schwartz justifie aussi que $\text{Aire}(\mathcal{P}(x, y))$ est bien définie. On vérifie d'autre part facilement que

$$\text{Aire}(\mathcal{P}(x, y)) = \|x\| \|y\| \sin(\alpha),$$

ce qui correspond à la définition de l'aire d'un parallélogramme comme le produit de la "base" par la "hauteur".

(4.) On dit que deux vecteurs $x, y \in \mathbb{E}^n$ sont *orthogonaux*, et on note $x \perp y$, si $\langle x, y \rangle = 0$.

Proposition 1.3. *Tout espace vectoriel euclidien $\mathbb{E}^n$ est un espace métrique pour la distance définie ci-dessus.*

Preuve. Nous devons vérifier que la distance $d(x, y) = \|y - x\|$ vérifie les trois propriétés suivantes pour tous $x, y, z \in \mathbb{E}^n$:

- (i.) $d(x, y) \geq 0$ et $d(x, y) = 0$ si et seulement si $x = y$.
- (ii.) $d(x, y) = d(y, x)$.
- (iii.) $d(x, z) \leq d(x, y) + d(y, z)$ (inégalité du triangle).

Ces propriétés se déduisent très facilement de la proposition 1.2. Vérifions par exemple l'inégalité du triangle :

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

□

Proposition 1.4. *Les conditions suivantes sont équivalentes pour deux vecteurs non nuls $x, y \in \mathbb{E}^n$:*

- (i) $x \perp y$, i.e. $\langle x, y \rangle = 0$.
- (ii) L'angle θ entre x et y est égal à $\frac{\pi}{2}$.
- (iii) On a $\|x + y\| = \|x - y\|$.
- (iv) On a $\|x + y\|^2 = \|x\|^2 + \|y\|^2$ (théorème de Pythagore).

Preuve. L'équivalence entre (i) et (ii) vient de

$$\theta = \frac{\pi}{2} \Leftrightarrow \cos(\theta) = 0 \Leftrightarrow \langle x, y \rangle = 0.$$

L'équivalence entre (i) et (iii) vient de

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

et celle entre (i) et (iv) de

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

□

1.2 Orientation d'un espace vectoriel réel de dimension finie

Dans ce bref paragraphe nous définissons la notion d'orientation d'un espace vectoriel de dimension finie sur le corps des réels. Rappelons que si $\{u_1, \dots, u_n\}$ et $\{v_1, \dots, v_n\}$ sont deux bases d'un espace vectoriel V , alors on appelle *matrice de changement de base* de la base $\{u_i\}$ vers la base $\{v_j\}$ la matrice $P = (p_{ij})$ définie par

$$v_j = \sum_{i=1}^n p_{ij} u_i$$

Définition. On dit que deux bases $\{u_1, \dots, u_n\}$ et $\{v_1, \dots, v_n\}$ d'un espace vectoriel réel ont la *même orientation* si le déterminant de la matrice de changement de base P est positif. Sinon on dit que les bases ont des *orientations opposées*.

Il n'est pas difficile de vérifier que "avoir la même orientation" est une relation d'équivalence sur l'ensemble des bases de V . De plus il existe exactement deux classes d'équivalences.

Définition. On appelle *orientation* de V le choix d'une classe d'équivalence pour cette relation. Un *espace vectoriel réel orienté* est un espace vectoriel muni du choix d'une orientation.

Une orientation de V est donc définie dès qu'on a choisi une base $\mathcal{B} = \{v_1, \dots, v_n\}$ et qu'on l'a déclarée d'orientation positive. Toute autre base est dite d'*orientation positive* si elle a la même orientation que $\mathcal{B}$; on dit aussi que c'est une *base directe*. Une base est dite d'*orientation négative* si elle a l'orientation opposée à la base $\mathcal{B}$.

Finalement, on dit qu'une application linéaire $f : V \rightarrow V$ *préserve l'orientation* si son déterminant est positif et qu'elle *inverse l'orientation* si son déterminant est négatif. Cette notion est indépendante du choix d'une orientation sur V . On note

$$\mathrm{GL}_+(V) = \{f \in \mathrm{GL}(V) \mid \det(f) > 0\},$$

c'est un sous-groupe du groupe linéaire général de V .

1.3 Similitudes et isométries d'un espace vectoriel euclidien.

Définition. Un *similitude de rapport* $\lambda > 0$ d'un espace vectoriel euclidien $\mathbb{E}^n$ est une application bijective $f : \mathbb{E}^n \rightarrow \mathbb{E}^n$ telle que

$$d(f(x), f(y)) = \lambda d(x, y), \quad \forall x, y \in \mathbb{E}^n.$$

Une *isométrie* de $\mathbb{E}^n$ est une similitude de rapport 1. C'est donc une bijection qui respecte les distances.

Il est facile de vérifier à partir de cette définition que les similitudes de $\mathbb{E}^n$ forment un groupe et que les isométries forment un sous-groupe normal de ce groupe. Pour décrire le groupe des isométries, nous commençons par décrire les isométries qui fixent l'origine.

Lemme 1.5. *Toute isométrie $g : E \rightarrow E$ d'un espace euclidien $\mathbb{E}^n$ qui fixe l'origine vérifie*

- (a) *g préserve le produit scalaire, i.e. $\langle g(x), g(y) \rangle = \langle x, y \rangle$ pour tous $x, y \in \mathbb{E}^n$.*
- (b) *L'application g est linéaire.*

Preuve. (a) Puisque g est une isométrie, on a $\|g(y) - g(x)\| = \|y - x\|$ pour tous $x, y \in \mathbb{E}^n$. Notons aussi que $\|g(x)\| = \|x\|$ pour tout x car

$$\|g(x)\| = \|g(x) - 0\| = \|g(x) - g(0)\| = \|x - 0\| = \|x\|,$$

puisque $g(0) = 0$. On a donc

$$\begin{aligned} \langle g(x), g(y) \rangle &= \frac{1}{2} (\|g(y)\|^2 + \|g(x)\|^2 - \|g(y) - g(x)\|^2) \\ &= \frac{1}{2} (\|y\|^2 + \|x\|^2 - \|y - x\|^2) \\ &= \langle x, y \rangle. \end{aligned}$$

(b) Nous pouvons maintenant montrer la linéarité de g . Soient $x \in E$ un vecteur quelconque et $\alpha \in \mathbb{R}$, alors

$$\begin{aligned}\|g(\alpha x) - \alpha g(x)\|^2 &= \|g(\alpha x)\|^2 - 2\langle g(\alpha x), \alpha g(x) \rangle + \alpha^2 \|g(x)\|^2 \\ &= \|g(\alpha x)\|^2 - 2\alpha \langle g(\alpha x), g(x) \rangle + \alpha^2 \|g(x)\|^2 \\ &= \|\alpha x\|^2 - 2\alpha \langle \alpha x, x \rangle + \alpha^2 \|x\|^2 \\ &= 0,\end{aligned}$$

ce qui prouve que $g(\alpha x) = \alpha g(x)$.

D'autre part, si $x, y \in E$ sont deux vecteurs, alors

$$\begin{aligned}\|g(x) + g(y) - g(x+y)\|^2 &= \langle g(x) + g(y) - g(x+y), g(x) + g(y) - g(x+y) \rangle \\ &= \|g(x)\|^2 + \|g(y)\|^2 + \|g(x+y)\|^2 + 2\langle g(x), g(y) \rangle - 2\langle g(x), g(x+y) \rangle - 2\langle g(x+y), g(y) \rangle \\ &= \|x\|^2 + \|y\|^2 + \|x+y\|^2 + 2\langle x, y \rangle - 2\langle x, x+y \rangle - 2\langle x+y, y \rangle \\ &= \langle x+y - (x+y), x+y - (x+y) \rangle = 0,\end{aligned}$$

ce qui prouve que $g(x+y) = g(x) + g(y)$. On a donc démontré qu'une isométrie de E qui fixe l'origine est une application linéaire. □

Théorème 1.6. *L'application $f : E \rightarrow E$ est une similitude de rapport λ si et seulement s'il existe un vecteur $b \in \mathbb{E}^n$ et une isométrie linéaire $g : \mathbb{E}^n \rightarrow \mathbb{E}^n$ tels que $f(x) = \lambda g(x) + b$ pour tout $x \in \mathbb{E}^n$.*

On dit que λg est la *partie linéaire* de l'isométrie f et b est le *vecteur de translation* de f . Remarquons que ce vecteur est donné par $b = f(0)$.

Preuve. On définit une application $g : E \rightarrow E$ par $g(x) = \frac{1}{\lambda} (f(x) - f(0))$. Alors il est clair que $g(0) = 0$ et g est une isométrie car

$$\begin{aligned}d(g(x), g(y)) &= \|g(x) - g(y)\| \\ &= \left\| \frac{1}{\lambda} (f(x) - f(0)) - \frac{1}{\lambda} (f(y) - f(0)) \right\| \\ &= \frac{1}{\lambda} \|f(x) - f(y)\| \\ &= \|x - y\|.\end{aligned}$$

Par le lemme précédent, g est linéaire. On a donc montré que l'application f s'écrit $f(x) = \lambda g(x) + b$, où $b = f(0) \in E$ est un vecteur constant et g est une isométrie linéaire. □

Corollaire 1.7. *Une application $f : \mathbb{R}^n \rightarrow \mathbb{R}^n$ est une isométrie pour le produit scalaire standard de $\mathbb{R}^n$ si et seulement si on a*

$$f(x) = Ax + b,$$

où $b = f(0) \in \mathbb{R}^n$ et $A \in GL_n(\mathbb{R})$ est une matrice vérifiant $A^\top A = I_n$.

Preuve. Par définition du produit scalaire standard de $\mathbb{R}^n$, on a

$$\langle e_i, e_j \rangle = \delta_{ij},$$

où $\{e_1, \dots, e_n\}$ est la base canonique de $\mathbb{R}^n$ (cette relation exprime que la base canonique est une base orthonormée).

D'autre part on a $Ae_r = \sum_{i=1}^n a_{ir}e_i$ et $Ae_s = \sum_{j=1}^n a_{js}e_j$, par conséquent :

$$\delta_{rs} = \langle e_r, e_s \rangle = \langle Ae_r, Ae_s \rangle = \left\langle \sum_{i=1}^n a_{ir}e_i, \sum_{j=1}^n a_{js}e_j \right\rangle = \sum_{i=1}^n \sum_{j=1}^n a_{ir}a_{js}\delta_{ij} = \sum_{i=1}^n a_{ir}a_{is} = \left(A^\top A \right)_{rs},$$

ce qui prouve que $A^\top A = I_n$.

□

1.4 Le groupe orthogonal

Le résultat précédent justifie la définition importante suivante :

Définition 1.8. Une matrice $A \in M_n(\mathbb{R})$ est *orthogonale* si $A^\top A = I_n$. L'ensemble des $n \times n$ matrices orthogonales se note

$$O(n) = \{A \in M_n(\mathbb{R}) \mid A^\top A = I_n\}$$

Proposition 1.9. Pour toute matrice $A \in M_n(\mathbb{R})$ les propriétés suivantes sont équivalentes :

- (i) $A \in O(n)$, c'est-à-dire $A^\top A = I_n$.
- (ii) A est inversible et $A^{-1} = A^\top$.
- (iii) $\|Ax\| = \|x\|$ pour tout $x \in \mathbb{R}^n$.
- (iv) $\langle Ax, Ay \rangle = \langle x, y \rangle$ pour tous $x, y \in \mathbb{R}^n$.
- (v) Les colonnes de A forment une base orthonormée de $\mathbb{R}^n$.
- (vi) Les lignes de A forment une base orthonormée de $\mathbb{R}^n$.
- (vii) Pour tout vecteur $b \in \mathbb{R}^n$, l'application affine $f : \mathbb{R}^n \rightarrow \mathbb{R}^n$ définie par $f(x) = Ax + b$ est une isométrie.

De plus $O(n)$ est un sous-groupe de $GL_n(\mathbb{R})$ et pour tout $A \in O(n)$ on a $\det(A) = \pm 1$.

Dans cette proposition, le produit scalaire est le produit scalaire standard de $\mathbb{R}^n$ et la norme et la distance sont associées à ce produit scalaire. Nous laissons la preuve de cette proposition en exercice.

Remarquons que l'application déterminant définit un homomorphisme de groupes

$$\det : O(n) \rightarrow \{\pm 1\},$$

le noyau de cet homomorphisme est le *groupe spécial orthogonal* :

$$SO(n) = O(n) \cap SL_n(\mathbb{R}) = \{A \in M_n(\mathbb{R}) \mid A^\top A = I_n \text{ et } \det(A) = +1\}.$$

La proposition suivante décrit les matrices orthogonales de taille 2×2 .

Proposition 1.10. Pour toute matrice $A \in O(2)$, il existe $\theta \in \mathbb{R}$ tel que

$$A = R_\theta = \begin{pmatrix} \cos(\theta) & -\sin(\theta) \\ \sin(\theta) & \cos(\theta) \end{pmatrix}, \text{ si } \det(A) = +1,$$

et

$$A = S_{\theta/2} = \begin{pmatrix} \cos(\theta) & \sin(\theta) \\ \sin(\theta) & -\cos(\theta) \end{pmatrix}, \text{ si } \det(A) = -1.$$

La matrice R_θ représente une rotation d'angle θ et $S_{\theta/2}$ représente la réflexion à travers la droite vectorielle formant un angle $\theta/2$ avec le premier vecteur e_1 de la base canonique.

Preuve. Les colonnes d'une matrice orthogonale $A \in O(2)$ doivent former une base orthonormée de $\mathbb{R}^2$. Il existe donc $\theta \in (-\pi, \pi]$ tel que la première colonne s'écrive $\begin{pmatrix} \cos(\theta) \\ \sin(\theta) \end{pmatrix}$.

La deuxième colonne de A doit être un vecteur de norme 1 orthogonal à la première colonne, par conséquent $\pm \begin{pmatrix} -\sin(\theta) \\ \cos(\theta) \end{pmatrix}$. Ceci démontre que ou bien $A = R_\theta$ ou bien $A = S_{\theta/2}$.

Pour voir que R_θ est une matrice de rotation, on peut vérifier que l'angle orienté entre tout vecteur non nul $\mathbf{x}$ et $R_\theta(\mathbf{x})$ est égal à θ .

Finalement, $S_{\theta/2}$ est une symétrie car cette matrice possède deux vecteurs propres orthogonaux de valeurs propre $+1$ et -1 respectivement. Ces vecteurs propres sont (au signe près)

$$\begin{pmatrix} \cos(\theta/2) \\ \sin(\theta/2) \end{pmatrix} \quad \text{et} \quad \begin{pmatrix} -\sin(\theta/2) \\ \cos(\theta/2) \end{pmatrix}.$$

Nous laissons la vérification de ces deux dernières affirmations en exercice. □

1.5 Un théorème d'Euler

Le théorème d'Euler décrit les isométries directes fixant un point dans l'espace à trois dimensions.

Théorème 1.11 (Théorème d'Euler). *Toute isométrie directe $f : \mathbb{E}^3 \rightarrow \mathbb{E}^3$ fixant un point est ou bien l'identité ou bien une rotation autour d'un axe passant par ce point.*

Preuve. On peut supposer que f fixe l'origine O . Alors f est une transformation linéaire. On a $f(\mathbf{x}) = A\mathbf{x}$. On sait également que $A \in SO(3)$ (c'est-à-dire $A^\top A = \mathbf{I}$ et $\det(A) = +1$). Pour montrer qu'il existe un axe, il suffit de montrer qu'il existe un vecteur propre de valeur propre $\lambda = 1$. En effet, s'il existe un vecteur non nul $\mathbf{a}$ tel que $A\mathbf{a} = \mathbf{a}$, alors la droite $\mathbb{R}\mathbf{a}$ est fixe pour la transformation f (c'est donc un axe pour f) car

$$f(t\mathbf{a}) = A(t\mathbf{a}) = tA(\mathbf{a}) = t\mathbf{a}$$

pour tout $t \in \mathbb{R}$. Pour montrer que 1 est une valeur propre de A , il faut montrer que

$$\det(A - \mathbf{I}) = 0.$$

On a

$$\begin{aligned} \det(A - \mathbf{I}) &= \underbrace{\det(A^\top)}_{=1} \det(A - \mathbf{I}) = \det(A^\top (A - \mathbf{I})) \\ &= \det(A^\top A - A^\top) = \det(\mathbf{I} - A^\top) = \det(\mathbf{I} - A). \end{aligned}$$

Or, comme $(A - \mathbf{I})$ est une matrice 3×3 , on a $\det(A - \mathbf{I}) = -\det(\mathbf{I} - A)$, donc

$$\det(\mathbf{I} - A) = -\det(\mathbf{I} - A).$$

Il en résulte que $\det(\mathbf{I} - A) = 0$.

Il faut encore prouver que f est bien une rotation autour de l'axe $\mathbb{R}\mathbf{a}$; considérons pour cela un vecteur $\mathbf{u}_1$ de longueur 1 et perpendiculaire à $\mathbf{a}$ et notons $\mathbf{u}_2 := \mathbf{a} \times \mathbf{u}_1$. Observons que $A\mathbf{u}_1$ et $A\mathbf{u}_2$ sont aussi orthogonaux à l'axe car

$$\langle A\mathbf{u}_i, \mathbf{a} \rangle = \langle A\mathbf{u}_i, A\mathbf{a} \rangle = \langle \mathbf{u}_i, \mathbf{a} \rangle = 0.$$

Ceci implique que les vecteurs $A\mathbf{u}_1$ et $A\mathbf{u}_2$ sont des combinaisons linéaires de $\mathbf{u}_1$ et $\mathbf{u}_2$, et comme ces vecteurs sont aussi de longueur 1 et orthogonaux, on a

$$\begin{aligned} A\mathbf{u}_1 &= \cos(\theta)\mathbf{u}_1 + \sin(\theta)\mathbf{u}_2, \\ A\mathbf{u}_2 &= -\sin(\theta)\mathbf{u}_1 + \cos(\theta)\mathbf{u}_2 \end{aligned}$$

où θ est l'angle entre $\mathbf{u}_1$ et $A\mathbf{u}_1$.

La matrice de la transformation linéaire A dans la base orthonormée $\mathbf{u}_1, \mathbf{u}_2, \mathbf{a}$ est donc la matrice

$$\begin{pmatrix} \cos(\theta) & -\sin(\theta) & 0 \\ \sin(\theta) & \cos(\theta) & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

et il s'agit bien d'une rotation autour de l'axe $\mathbb{R}\mathbf{a}$.

□

Rappelons que la *trace* d'une matrice est la somme de ses éléments diagonaux. On prouve dans le cours d'algèbre linéaire que deux matrices semblables ont la même trace. Donc la trace de A dans la base originale coïncide avec la trace de A dans la base $\mathbf{u}_1, \mathbf{u}_2, \mathbf{a}$. Cette trace vaut donc $1 + 2\cos(\theta)$; on a prouvé le résultat suivant.

Proposition 1.12. *L'angle θ d'une rotation $A \in SO(3)$ est donné par l'équation*

$$\text{Trace}(A) = 1 + 2\cos(\theta).$$

Certaines matrices de rotation sont très simples. Par exemple la rotation d'angle θ autour de l'axe Ox est donnée par la matrice

$$R_x(\theta) = \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos(\theta) & -\sin(\theta) \\ 0 & \sin(\theta) & \cos(\theta) \end{pmatrix}$$

et la rotation d'angle θ autour de l'axe Oy est donnée par la matrice

$$R_y(\theta) = \begin{pmatrix} \cos(\theta) & 0 & \sin(\theta) \\ 0 & 1 & 0 \\ -\sin(\theta) & 0 & \cos(\theta) \end{pmatrix}$$

(observer la place du signe $-$ dans cette matrice!).

Finalement, la rotation d'angle θ autour de l'axe Oz est donnée par la matrice

$$R_z(\theta) = \begin{pmatrix} \cos(\theta) & -\sin(\theta) & 0 \\ \sin(\theta) & \cos(\theta) & 0 \\ 0 & 0 & 1 \end{pmatrix}$$

Si on effectue une rotation d'angle φ autour de l'axe Oz , puis une rotation d'angle θ autour de l'axe Oy et enfin une rotation d'angle ψ de nouveau autour de l'axe Oz , on obtient une matrice

$$\begin{aligned} A &= R_z(\psi) \circ R_y(\theta) \circ R_z(\varphi) = \\ &\begin{pmatrix} \cos(\psi)\cos(\theta)\cos(\varphi) - \sin(\psi)\sin(\varphi) & -\cos(\psi)\cos(\theta)\sin(\varphi) - \sin(\psi)\cos(\varphi) & \cos(\psi)\sin(\theta) \\ \sin(\psi)\cos(\theta)\cos(\varphi) + \cos(\psi)\sin(\varphi) & -\sin(\psi)\cos(\theta)\sin(\varphi) + \cos(\psi)\cos(\varphi) & \sin(\psi)\sin(\theta) \\ -\sin(\theta)\cos(\varphi) & \sin(\theta)\sin(\varphi) & \cos(\theta) \end{pmatrix} \end{aligned}$$

Toute matrice de rotation dans $\mathbb{R}^3$ s'obtient de cette manière (avec $0 \leq \psi < 2\pi$, $0 \leq \theta < \pi$ et $0 \leq \varphi < 2\pi$). Les angles ψ, θ, φ s'appellent les *angles d'Euler* de la rotation A .

1.6 Géométrie vectorielle dans l'espace euclidien orienté $\mathbb{E}^3$

Soit $\mathbb{E}^3$ un espace vectoriel euclidien orienté de dimension 3. On appelle *produit vectoriel* de deux vecteurs $\mathbf{x}, \mathbf{y} \in \mathbb{E}^3$ le vecteur $\mathbf{x} \times \mathbf{y} \in \mathbb{E}^3$, vérifiant les conditions géométriques suivantes :

- (i) $(\mathbf{x} \times \mathbf{y}) \perp \mathbf{x}$ et $(\mathbf{x} \times \mathbf{y}) \perp \mathbf{y}$.
- (ii) $\|\mathbf{x} \times \mathbf{y}\| = \text{aire}(\mathcal{P}(\mathbf{x}, \mathbf{y}))$, où $\mathcal{P}(\mathbf{x}, \mathbf{y})$ est le parallélogramme construit sur les vecteurs $\mathbf{x}$ et $\mathbf{y}$.
- (iii) Si $\mathbf{x}$ et $\mathbf{y}$ sont linéairement indépendants, alors $\{\mathbf{x}, \mathbf{y}, \mathbf{x} \times \mathbf{y}\}$ est une base directe de $\mathbb{E}^3$.

La proposition suivante justifie cette définition :

Proposition 1.13. *Le produit vectoriel est uniquement défini par les trois conditions ci-dessus. De plus, si $\{\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3\}$ est une base orthonormée directe de $\mathbb{E}^3$, alors le produit vectoriel de $\mathbf{x} = x_1\mathbf{e}_1 + x_2\mathbf{e}_2 + x_3\mathbf{e}_3$ et $\mathbf{y} = y_1\mathbf{e}_1 + y_2\mathbf{e}_2 + y_3\mathbf{e}_3$ se calcule par la formule suivante :*

$$\mathbf{x} \times \mathbf{y} = (x_2y_3 - x_3y_2)\mathbf{e}_1 + (x_3y_1 - x_1y_3)\mathbf{e}_2 + (x_1y_2 - x_2y_1)\mathbf{e}_3. \quad (1.1)$$

Preuve. Si $\mathbf{x}$ et $\mathbf{y}$ sont linéairement dépendants, alors $\text{aire}(\mathcal{P}(\mathbf{x}, \mathbf{y})) = 0$, par conséquent $\mathbf{x} \times \mathbf{y}$ doit être le vecteur nul, et on vérifie facilement que dans ce cas le membre de droite de (1.1) est en effet nul. Lorsque $\mathbf{x}$ et $\mathbf{y}$ sont linéairement indépendants, l'ensemble des vecteurs qui sont à la fois orthogonaux à $\mathbf{x}$ et à $\mathbf{y}$ est un sous-espace vectoriel de dimension 1. Ce sous-espace contient exactement deux vecteurs dont la norme est égale à $\text{aire}(\mathcal{P}(\mathbf{x}, \mathbf{y}))$, et pour un seul de ces deux vecteurs, que l'on notera $\mathbf{x} \times \mathbf{y}$, la base $\{\mathbf{x}, \mathbf{y}, \mathbf{x} \times \mathbf{y}\}$ est d'orientation positive.

Nous devons maintenant prouver que le produit vectoriel est donné par la formule (1.1). Cela demande un peu de calcul. Notons

$$\mathbf{z} = (x_2y_3 - x_3y_2)\mathbf{e}_1 + (x_3y_1 - x_1y_3)\mathbf{e}_2 + (x_1y_2 - x_2y_1)\mathbf{e}_3,$$

et observons pour commencer que

$$\langle \mathbf{z}, \mathbf{x} \rangle = (x_2y_3 - x_3y_2)x_1 + (x_3y_1 - x_1y_3)x_2 + (x_1y_2 - x_2y_1)x_3 = \det \begin{pmatrix} x_1 & y_1 & x_1 \\ x_2 & y_2 & x_2 \\ x_3 & y_3 & x_3 \end{pmatrix} = 0.$$

De même $\langle \mathbf{z}, \mathbf{y} \rangle = 0$, ce qui montre que $\mathbf{z}$ est orthogonal à $\mathbf{x}$ et $\mathbf{y}$. Pour montrer la propriété (ii), on calcule le carré de la norme¹ de $\mathbf{z}$, et on en réorganise les termes :

$$\begin{aligned} \|\mathbf{z}\|^2 &= (x_2y_3 - x_3y_2)^2 + (x_3y_1 - x_1y_3)^2 + (x_1y_2 - x_2y_1)^2 \\ &= \sum_{i \neq j} x_i^2 y_j^2 - 2 \sum_{i < j} x_i x_j y_i y_j \\ &= \sum_{i,j} x_i^2 y_j^2 - \left(\sum_i x_i^2 y_i^2 + 2 \sum_{i < j} x_i y_i x_j y_j \right) \\ &= \left(\sum_i x_i^2 \right) \left(\sum_j y_j^2 \right) - \left(\sum_i x_i y_i \right)^2 \\ &= \|x\|^2 \|y\|^2 - \langle \mathbf{x}, \mathbf{y} \rangle^2 \\ &= \text{aire}(\mathcal{P}(\mathbf{x}, \mathbf{y}))^2. \end{aligned}$$

1. Il est souvent plus commode de calculer le carré d'une norme que la norme elle-même.

Finalement, pour prouver (iii) on remarque que si $\mathbf{x}$ et $\mathbf{y}$ sont linéairement indépendants, alors $\mathbf{z}$ est non nul et donc

$$\begin{aligned} \det \begin{pmatrix} x_1 & y_1 & z_1 \\ x_2 & y_2 & z_2 \\ x_3 & y_3 & z_3 \end{pmatrix} &= (x_2 y_3 - x_3 y_2) z_1 + (x_3 y_1 - x_1 y_3) z_2 + (x_1 y_2 - x_2 y_1) z_3 \\ &= z_1^2 + z_2^2 + z_3^2 = \|\mathbf{z}\|^2 > 0, \end{aligned}$$

ce qui implique que $\{\mathbf{x}, \mathbf{y}, \mathbf{z} = \mathbf{x} \times \mathbf{y}\}$ est une base directe. $\square$

Remarque. Le produit vectoriel peut aussi s'écrire

$$\mathbf{x} \times \mathbf{y} = \begin{vmatrix} x_2 & y_2 \\ x_3 & y_3 \end{vmatrix} \mathbf{e}_1 - \begin{vmatrix} x_1 & y_1 \\ x_3 & y_3 \end{vmatrix} \mathbf{e}_2 + \begin{vmatrix} x_1 & y_1 \\ x_2 & y_2 \end{vmatrix} \mathbf{e}_3,$$

que l'on écrit aussi parfois sous la forme d'un "déterminant formel"

$$\mathbf{x} \times \mathbf{y} = \begin{vmatrix} x_1 & y_1 & \mathbf{e}_1 \\ x_2 & y_2 & \mathbf{e}_2 \\ x_3 & y_3 & \mathbf{e}_3 \end{vmatrix}.$$

Observons aussi que le produit vectoriel définit une application bilinéaire antisymétrique $\times : \mathbb{E}^3 \times \mathbb{E}^3 \rightarrow \mathbb{E}^3$.

Définition. On appelle *produit mixte* de trois vecteurs $\mathbf{x}, \mathbf{y}, \mathbf{w} \in \mathbb{E}^3$, le produit scalaire

$$[\mathbf{x}, \mathbf{y}, \mathbf{w}] = \langle \mathbf{x} \times \mathbf{y}, \mathbf{w} \rangle.$$

Il est clair à partir de la formule (1.1) que dans une base orthonormée directe, la produit mixte est donné par le déterminant suivant :

$$[\mathbf{x}, \mathbf{y}, \mathbf{w}] = \begin{vmatrix} x_1 & y_1 & w_1 \\ x_2 & y_2 & w_2 \\ x_3 & y_3 & w_3 \end{vmatrix}.$$

Cette quantité représente le *volume orienté* du parallélépipède $\mathcal{P}(\mathbf{x}, \mathbf{y}, \mathbf{w})$ construit sur les trois vecteurs.

1.7 Géométrie vectorielle dans le plan euclidien orienté $\mathbb{E}^2$

Dans ce paragraphe et le suivant nous travaillons dans un espace vectoriel euclidien $\mathbb{E}^2$ muni d'une orientation ; on se donne également une base orthonormée directe (i.e. d'orientation positive) $\{\mathbf{e}_1, \mathbf{e}_2\}$.

Par définition les vecteurs $\mathbf{a} = a_1 \mathbf{e}_1 + a_2 \mathbf{e}_2$ et $\mathbf{b} = b_1 \mathbf{e}_1 + b_2 \mathbf{e}_2$ forment une autre base directe de $\mathbb{E}^2$ si et seulement si $a_1 b_2 - a_2 b_1 > 0$. Ils forment une base d'orientation négative si $a_1 b_2 - a_2 b_1 < 0$.

L'opérateur $\mathbf{J}$

On note $\mathbf{J} : \mathbb{E}^2 \rightarrow \mathbb{E}^2$ l'application linéaire qui est donnée dans une base orthonormée directe par $\mathbf{J}(v_1 \mathbf{e}_1 + v_2 \mathbf{e}_2) = -v_2 \mathbf{e}_1 + v_1 \mathbf{e}_2$. Sa matrice dans une telle base est donc

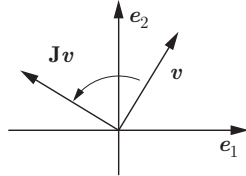
$$\mathbf{J} = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}.$$

L'opérateur $\mathbf{J}$ est caractérisée par les propriétés géométriques suivantes :

- (i) $\|\mathbf{J}\mathbf{v}\| = \|\mathbf{v}\|$ (en particulier $\mathbf{J}(0) = 0$),
- (ii) $\mathbf{J}\mathbf{v} \perp \mathbf{v}$,
- (iii) Si $\mathbf{v} \neq 0$, alors $\{\mathbf{v}, \mathbf{J}\mathbf{v}\}$ est une base d'orientation positive.

En particulier, $\mathbf{J}$ ne dépend pas de la base orthonormée directe choisie (mais cet opérateur dépend de l'orientation de $\mathbb{E}^2$.)

Géométriquement, l'opérateur $\mathbf{J}$ est la rotation qui fait tourner le vecteur $\mathbf{v}$ d'un quart de tour dans le sens positif.



Définition. Le *produit extérieur* de deux vecteurs $\mathbf{a}, \mathbf{b} \in \mathbb{E}^2$ est le scalaire $\mathbf{a} \wedge \mathbf{b} \in \mathbb{R}$ défini par

$$\mathbf{a} \wedge \mathbf{b} = \langle \mathbf{J}(\mathbf{a}), \mathbf{b} \rangle = -\langle \mathbf{a}, \mathbf{J}(\mathbf{b}) \rangle.$$

Dans une base orthonormée directe $\{\mathbf{e}_1, \mathbf{e}_2\}$, le produit extérieur de $\mathbf{a} = a_1\mathbf{e}_1 + a_2\mathbf{e}_2$ et $\mathbf{b} = b_1\mathbf{e}_1 + b_2\mathbf{e}_2$ est donné par

$$\mathbf{a} \wedge \mathbf{b} = \langle -a_2\mathbf{e}_1 + a_1\mathbf{e}_2, b_1\mathbf{e}_1 + b_2\mathbf{e}_2 \rangle = a_1b_2 - a_2b_1,$$

c'est-à-dire :

$$\mathbf{a} \wedge \mathbf{b} = \det \begin{pmatrix} a_1 & b_1 \\ a_2 & b_2 \end{pmatrix}.$$

Noter aussi que

$$(\mathbf{a} \wedge \mathbf{b})^2 = (a_1^2 + a_2^2)(b_1^2 + b_2^2) - (a_1b_1 + a_2b_2)^2$$

Les propriétés suivantes découlent immédiatement de ces formules :

Proposition 1.14. *Le produit extérieur vérifie les propriétés suivantes :*

- (i) *Le produit extérieur est bilinéaire et antisymétrique.*
- (ii) *$\mathbf{a} \wedge \mathbf{b} = 0$ si et seulement si $\mathbf{a}$ et $\mathbf{b}$ sont colinéaires.*
- (iii) *$|\mathbf{a} \wedge \mathbf{b}| = \text{aire}(\mathcal{P}(\mathbf{a}, \mathbf{b}))$.*
- (iv) *$|\mathbf{a} \wedge \mathbf{b}| \leq \|\mathbf{a}\| \|\mathbf{b}\|$ et on a égalité si et seulement si $\mathbf{a} \perp \mathbf{b}$*
- (v) *$\mathbf{a} \wedge \mathbf{b} > 0$ si et seulement si $\{\mathbf{a}, \mathbf{b}\}$ est une base directe de $\mathbb{E}^2$.*

On définit alors l'*angle orienté* $\theta_{\text{or}} \in (-\pi, \pi]$ entre deux vecteurs non nuls $\{\mathbf{a}_1, \mathbf{b}_2\}$ de $\mathbb{E}^2$ par

$$\theta_{\text{or}} = \begin{cases} \angle(\mathbf{a}_1, \mathbf{b}_2), & \text{si } \mathbf{a} \wedge \mathbf{b} \geq 0, \\ -\angle(\mathbf{a}_1, \mathbf{b}_2), & \text{si } \mathbf{a} \wedge \mathbf{b} < 0. \end{cases}$$

où $\angle(\mathbf{a}_1, \mathbf{b}_2) \in [0, \pi]$ est l'angle non orienté. Ainsi l'angle orienté entre $\mathbf{a}_1$ et $\mathbf{b}_2$ est négatif si et seulement si ces deux vecteurs forment une base d'orientation négative (et dans ce cas le signe du sinus de l'angle orienté est négatif). L'angle orienté est complètement déterminé par les formules :

$$\cos(\theta_{\text{or}}) = \frac{\langle \mathbf{a}, \mathbf{b} \rangle}{\|\mathbf{a}\| \|\mathbf{b}\|}, \quad \sin(\theta_{\text{or}}) = \frac{\mathbf{a} \wedge \mathbf{b}}{\|\mathbf{a}\| \|\mathbf{b}\|}.$$

De la même manière, on définit l'*aire orientée* du parallélogramme construit sur les vecteurs $\{\mathbf{a}_1, \mathbf{b}_2\} \in \mathbb{E}^2$ par

$$\text{aire}_{\text{or}}(\mathcal{P}(\mathbf{a}, \mathbf{b})) = \begin{cases} \text{aire}(\mathcal{P}(\mathbf{a}, \mathbf{b})), & \text{si } \mathbf{a} \wedge \mathbf{b} \geq 0, \\ -\text{aire}(\mathcal{P}(\mathbf{a}, \mathbf{b})), & \text{si } \mathbf{a} \wedge \mathbf{b} < 0. \end{cases}$$

On voit donc que

$$\text{aire}_{\text{or}}(\mathcal{P}(\mathbf{a}, \mathbf{b})) = \|\mathbf{a}\| \|\mathbf{b}\| \sin(\theta_{\text{or}}) = \mathbf{a} \wedge \mathbf{b}.$$

Chapitre 2

Courbes dans l'espace et le plan euclidien

2.1 Qu'est ce qu'une courbe ?

La notion mathématique de *courbe* ou de *ligne* formalise l'idée intuitive d'un objet du plan ou de l'espace qui est continu et n'a qu'une dimension. Euclide en donne la définition suivante dans le livre I des *Eléments* : *une ligne est une longueur sans largeur*. Les droites, les cercles et les ellipses sont des exemples familiers de courbes. Dans la vie courante, un fil de fer ou la trajectoire d'un projectile sont des exemples concrets de courbes.

La formalisation de la notion de courbe conduit à plusieurs concepts qu'il faudra distinguer. Le premier est celui de « lieu géométrique » des points satisfaisant certaines propriétés : cette idée nous conduit à la notion *implicite* d'une courbe comme ensemble des points satisfaisant une équation (dans le plan) ou deux équations (dans l'espace de dimension 3). Le second concept est celui de courbe comme « trajectoire » : on ne conçoit plus la courbe comme un ensemble de points, mais comme un « point mobile », c'est-à-dire une fonction d'un paramètre à valeurs dans le plan ou dans l'espace : c'est le point de vue *paramétrique* ou *cinématique* en théorie des courbes. L'acte de tracer une courbe au crayon noir sur une feuille blanche se décrit par le point de vue paramétrique, le résultat de cette action, la courbe qu'on a tracée, correspond au point de vue implicite. Dans ce chapitre, nous privilégions le point de vue paramétrique.

2.2 Notions fondamentales

Dans ce chapitre, on suppose que l'espace est muni d'un système de coordonnées fixe. On l'identifie donc à $\mathbb{R}^n$ et on admet que n est un entier quelconque. On supposera, sauf mention du contraire, que le système de coordonnées est orthonormé. La norme d'un vecteur $\mathbf{v} = (v_1, v_2, \dots, v_n)$ est alors donnée par

$$\|\mathbf{v}\| = \sqrt{v_1^2 + v_2^2 + \dots + v_n^2},$$

et si $\mathbf{w} = (w_1, w_2, \dots, w_n)$ est un second vecteur alors leur produit scalaire est donné par

$$\langle \mathbf{v}, \mathbf{w} \rangle = \sum_{i=1}^n v_i w_i.$$

Définitions. Une *courbe paramétrée* dans $\mathbb{R}^n$ est une application continue $\alpha : I \rightarrow \mathbb{R}^n$:

$$\alpha : u \mapsto (\alpha_1(u), \alpha_2(u), \dots, \alpha_n(u)) \in \mathbb{R}^n, \quad u \in I$$

où $I \subset \mathbb{R}$ est un intervalle appelé l'*intervalle de paramétrisation* de la courbe.

La variable u parcourant l'intervalle I s'appelle le *paramètre* (elle est aussi parfois notée par les lettres s, t, φ ou θ) et l'ensemble

$$\alpha(I) = \{\alpha(u) \mid u \in I\} \subset \mathbb{R}^n$$

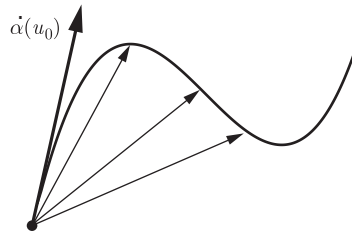
s'appelle la *trace* ou le *support* de la courbe paramétrée α .

On dit que la courbe α est *différentiable* en $u_0 \in I$ si la limite

$$\frac{d\alpha}{du}(u_0) := \lim_{u \rightarrow u_0} \frac{\alpha(u) - \alpha(u_0)}{u - u_0}$$

existe. Cette limite s'appelle alors le *vecteur vitesse* de la courbe α en u_0 et on le note $\dot{\alpha}(u_0)$ ou $\alpha'(u_0)$.

Remarquons que la direction du vecteur vitesse est tangente à la courbe en $\alpha(u_0)$ car cette direction est la limite des directions prises par une suite de cordes reliant le point $p = \alpha(u_0)$ à un point de la courbe se rapprochant du point p .



La *vitesse* de α en u_0 est la norme du vecteur vitesse, on la note

$$V_\alpha(u_0) = \|\dot{\alpha}(u_0)\|.$$

Lemme 2.1. La courbe $\alpha(u) = (\alpha_1(u), \dots, \alpha_n(u))$ est différentiable en u_0 si et seulement si les fonctions $\alpha_i(u)$ sont dérivables en u_0 . De plus,

$$\dot{\alpha}(u_0) = \left(\frac{d\alpha_1}{du}(u_0), \dots, \frac{d\alpha_n}{du}(u_0) \right).$$

Si le système de coordonnées est orthonormé, alors on a aussi :

$$V_\alpha(u_0) = \sqrt{\left(\frac{d\alpha_1}{du}(u_0) \right)^2 + \dots + \left(\frac{d\alpha_n}{du}(u_0) \right)^2}.$$

Nous laissons la vérification de ce lemme en exercice. □

Définitions. Voyons quelques définitions supplémentaires.

- (a) La courbe $\alpha : I \rightarrow \mathbb{R}^n$ est dite de *classe* C^1 si elle est différentiable en tout point de I et si les dérivées

$$\dot{\alpha}_j = \frac{d\alpha_j}{du}$$

sont continues sur l'intervalle I pour tout $j = 1, 2, \dots, n$.

- (b) La courbe est dite de *classe* C^k (où k est un entier) si les dérivées d'ordre m

$$\frac{d^m \alpha_j}{du^m}(u)$$

existent et sont continues pour tout $j = 1, 2, \dots, n$ et tout $m = 1, 2, \dots, k$.

Si une courbe est de classe C^k pour tout entier k , on dit qu'elle est de *classe* C^∞ . Si la courbe est simplement continue, on dit qu'elle est de *classe* C^0 .

- (c) Si $\alpha : I \rightarrow \mathbb{R}^n$ est une courbe de classe C^2 , alors son *accélération* est le vecteur défini par

$$\ddot{\alpha}(u) = \frac{d^2 \alpha}{du^2}(u).$$

- (d) Soit α une courbe de classe C^1 et u_0 une valeur du paramètre. On dit que le point $p = \alpha(u_0)$ est *singulier* si $\dot{\alpha}(u_0) = \mathbf{0}$ (de façon équivalente, p est singulier si et seulement si $V_\alpha(u_0) = 0$). Le point $p = \alpha(u_0)$ est *régulier* s'il n'est pas singulier.

Une courbe est *régulière* si elle est de classe C^1 et si tous ses points sont réguliers.

- (e) Le point $p = \alpha(u_0)$ sur une courbe de classe C^2 est *birégulier* si $\dot{\alpha}(u_0)$ et $\ddot{\alpha}(u_0)$ sont linéairement indépendants.

Une courbe est *birégulière* si elle est de classe C^2 et si tous ses points sont biréguliers.

- (f) Le *plan osculateur* à la courbe α au point $p = \alpha(u_0)$ est le plan passant par p et qui est parallèle aux vecteurs $\dot{\alpha}(u_0)$ et $\ddot{\alpha}(u_0)$. Ce plan n'est défini que si p est un point birégulier.

- (g) Un point p sur une courbe $\alpha : I \rightarrow \mathbb{R}^n$ est un *point double* s'il existe deux valeurs distinctes du paramètre ($u_1, u_2 \in I$, $u_1 \neq u_2$) telles que

$$p = \alpha(u_1) = \alpha(u_2).$$

- (h) Si $\alpha : I \rightarrow \mathbb{R}^n$ est une courbe et si $J \subset I$ est un intervalle, alors on dit que la restriction de α à J est un *arc de la courbe* α (un arc de courbe n'est donc rien d'autre qu'un « morceau de courbe »).

- (i) On dit qu'un arc de courbe est *simple* s'il ne contient pas de point double.

- (j) La *droite tangente* à la courbe γ au point régulier $\gamma(u_0)$ est la droite $T_{u_0}\gamma$ parcourue à vitesse constante, passant par $\gamma(u_0)$ dans la direction du vecteur vitesse $\dot{\gamma}(u_0)$:

$$T_{u_0}\gamma : \lambda \mapsto \gamma(u_0) + \lambda \dot{\gamma}(u_0), \quad \lambda \in \mathbb{R}$$

Premiers exemples

- 1) La *cubique* dans $\mathbb{R}^3$ est la courbe $\alpha : \mathbb{R} \rightarrow \mathbb{R}^3$ définie par

$$\alpha(u) = (au, bu^2, cu^3),$$

où a, b, c sont des constantes non nulles. Cette courbe est de classe C^∞ , son vecteur vitesse est

$$\dot{\alpha}(u) = (a, 2bu, 3cu^2)$$

et son accélération est

$$\ddot{\alpha}(u) = (0, 2b, 6cu).$$

La cubique est donc birégulière et sa vitesse est

$$V_\alpha(u) = \|\dot{\alpha}(u)\| = \sqrt{a^2 + 4b^2u^2 + 9c^2u^4}.$$

2) La courbe $\beta : \mathbb{R} \rightarrow \mathbb{R}^n$ définie par

$$\beta(u) = (u^2, u^3, \dots, u^{n+1})$$

est de classe C^∞ . Son vecteur vitesse est

$$\dot{\beta}(u) = (2u, 3u^2, \dots, (n+1)u^n),$$

et sa vitesse est $V_\beta(u) = \|\dot{\beta}(u)\| = \sqrt{4u^2 + \dots + ((n+1)u^n)^2}$. Cette courbe a un unique point singulier en $\beta(0) = (0, 0, \dots, 0)$.

3) La droite passant par les points distincts $p = (p_1, p_2, \dots, p_n)$ et $q = (q_1, q_2, \dots, q_n)$ admet la paramétrisation affine $\delta : \mathbb{R} \rightarrow \mathbb{R}^n$ suivante :

$$\delta(t) = p + t\vec{pq} = (p_1 + t(q_1 - p_1), p_2 + t(q_2 - p_2), \dots, p_n + t(q_n - p_n)).$$

En posant $\mathbf{w} = \vec{pq} = (w_1, w_2, \dots, w_n)$, on a $\delta(t) = (p_1 + tw_1, \dots, p_n + tw_n)$. Le vecteur vitesse et la vitesse sont donnés pour tout t par

$$\dot{\delta}(t) = \mathbf{w} \quad \text{et} \quad V_\delta(t) = \|\mathbf{w}\|,$$

et l'accélération est nulle. La courbe est donc régulière, de classe C^∞ et sa vitesse est constante. Son accélération est nulle et la droite n'est donc pas birégulière.

4) La même droite admet de nombreux autres paramétrisations, par exemple :

$$\varepsilon(t) = p + t^3\mathbf{w} = (p_1 + t^3w_1, \dots, p_n + t^3w_n) \quad (t \in \mathbb{R}).$$

Dans ce cas,

$$\dot{\varepsilon}(t) = 3t^2\mathbf{w} \quad \text{et} \quad V_\varepsilon(t) = 3t^2\|\mathbf{w}\|.$$

Cette courbe est de classe C^∞ et elle possède un unique point singulier en $\varepsilon(0) = p$.

5) Ou encore

$$\eta(t) = p + \sqrt[3]{t}\mathbf{w} = (p_1 + \sqrt[3]{t}w_1, \dots, p_n + \sqrt[3]{t}w_n) \quad (t \in \mathbb{R}).$$

Cette courbe n'est pas de classe C^1 , elle n'est en effet pas différentiable en $t = 0$. Nous avons pour $t \neq 0$:

$$\dot{\eta}(t) = \frac{1}{3}t^{-2/3}\mathbf{w} \quad \text{et} \quad V_\eta(t) = \frac{1}{3}t^{-2/3}\|\mathbf{w}\|,$$

et donc $V_\eta(t) \rightarrow \infty$ lorsque $t \rightarrow 0$.

6) Le cercle de centre p et rayon r dans le plan $\Pi \subset \mathbb{R}^n$ admet la paramétrisation

$$c(t) = p + r \cos(\omega t)\mathbf{b}_1 + r \sin(\omega t)\mathbf{b}_2 \quad (0 \leq t \leq \frac{2\pi}{\omega})$$

où $p, \mathbf{b}_1, \mathbf{b}_2$ est un repère orthonormé dans le plan Π et $\omega > 0$ est une constante appelée la *vitesse angulaire* (on vérifie en effet facilement que $\|c(t) - p\| = r$). La vitesse de cette courbe est

$$\dot{c}(t) = -\omega r \sin(\omega t)\mathbf{b}_1 + \omega r \cos(\omega t)\mathbf{b}_2 \quad \text{et} \quad V_c(t) = \omega r.$$

Son accélération est

$$\ddot{c}(t) = -\omega^2 r \cos(\omega t)\mathbf{b}_1 - \omega^2 r \sin(\omega t)\mathbf{b}_2.$$

Cette courbe est birégulière, elle est de classe C^∞ , et sa vitesse est constante. Elle admet un point double puisque $c(0) = c(\frac{2\pi}{\omega})$.

7) Le *graphe* d'une fonction $f : I \rightarrow \mathbb{R}$ de classe C^1 est la courbe $\gamma_f : I \rightarrow \mathbb{R}^2$ définie par

$$\gamma_f(x) = (x, f(x)).$$

Remarquons que dans cet exemple, la variable x est à la fois une coordonnée du plan et le paramètre de la courbe. Si f est continûment dérivable, alors la courbe est de classe C^1 et on a

$$\dot{\gamma}_f(x) = (1, f'(x)) \quad \text{et} \quad V_{\gamma}(x) = \sqrt{1 + (f'(x))^2}.$$

Cette courbe est toujours régulière puisqu'en tout point $V_{\gamma}(x) \geq 1$.

8) L'*hélice circulaire* est la courbe $\gamma : \mathbb{R} \rightarrow \mathbb{R}^3$ définie par

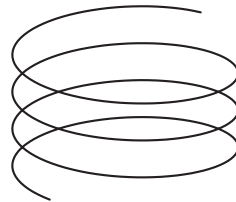
$$\gamma(u) = (a \cos(u), a \sin(u), bu),$$

où a et b sont des réels non nuls. Son vecteur vitesse et son accélération sont donnés par

$$\begin{aligned} \dot{\gamma}(u) &= (-a \sin(u), a \cos(u), b) \\ \ddot{\gamma}(u) &= a(-\cos(u), -\sin(u), 0). \end{aligned}$$

L'hélice circulaire est donc une courbe birégulière et la vitesse est constante :

$$\|\dot{\gamma}\| = \sqrt{a^2 + b^2}.$$



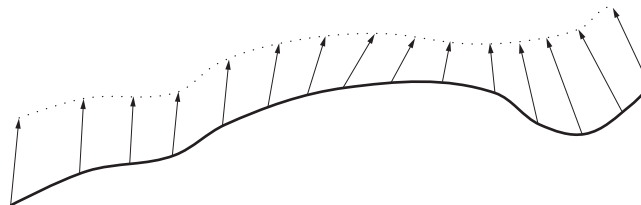
Hélice circulaire

2.3 Champs de vecteurs le long d'une courbe

Définition. Un *champ de vecteurs* le long d'une courbe $\gamma : I \rightarrow \mathbb{R}^n$ est la donnée d'un vecteur

$$\mathbf{W}(u) = w_1(u)\mathbf{e}_1 + w_2(u)\mathbf{e}_2 + \cdots + w_n(u)\mathbf{e}_n$$

pour toute valeur du paramètre $u \in I$. Le vecteur $\mathbf{W}(u)$ est en général considéré comme un vecteur fixe d'origine $\gamma(u)$, mais on peut aussi le voir comme un vecteur libre.



Champ de vecteurs le long d'une courbe.

Le champ de vecteurs $\mathbf{W}(u)$ est dit *de classe C^k* si les dérivées de $w_1, w_2, \dots, w_n$ existent et sont continues jusqu'à l'ordre k .

Exemples de champs de vecteurs

- (1) Si γ est de classe C^1 , alors son vecteur vitesse définit un champ $u \mapsto \dot{\gamma}(u)$.
- (2) Si γ est de classe C^2 , alors son accélération définit un champ $u \mapsto \ddot{\gamma}(u)$.
- (3) Si $\mathbf{W}(u)$ et $\mathbf{Z}(u)$ sont deux champs de vecteurs le long de la courbe $\gamma : I \rightarrow \mathbb{R}^n$ et $f, g : I \rightarrow \mathbb{R}$ sont deux fonctions, alors

$$u \mapsto f(u)\mathbf{W}(u) + g(u)\mathbf{Z}(u)$$

est un nouveau champ de vecteurs le long de la courbe.

- (4) Si $\mathbf{W}(u)$ est un champ de vecteurs de classe C^k , alors sa dérivée $\dot{\mathbf{W}}(u)$ est un champ de vecteurs de classe C^{k-1} et $\ddot{\mathbf{W}}(u)$ est un champ de classe C^{k-2} .
- (5) En dimension 3, un autre champ est donné par $u \mapsto \mathbf{W}(u) \times \mathbf{Z}(u)$.
- (6) Si $\gamma, \beta : I \rightarrow \mathbb{R}^n$ sont deux courbes ayant même intervalle de paramétrisation, alors on peut définir un champ de vecteurs le long de γ par

$$\mathbf{W}(u) = \beta(u) - \gamma(u).$$

Ce champ s'appelle le *champ de poursuite* de la courbe β depuis la courbe γ .

- (7) Un champ important est le *vecteur tangent* d'une courbe régulière $\gamma : I \rightarrow \mathbb{R}^n$ de classe C^1 . C'est le champ de vecteurs le long de la courbe obtenu en normalisant le vecteur vitesse :

$$\mathbf{T}_\gamma(u) := \frac{\dot{\gamma}(u)}{\|\dot{\gamma}(u)\|} = \frac{\dot{\gamma}(u)}{V_\gamma(u)}.$$

- (8) Le *vecteur normal principal* d'une courbe birégulière γ de classe C^2 est le champ de vecteurs le long de la courbe défini par

$$\mathbf{N}_\gamma(u) := \frac{\ddot{\gamma}(u) - \langle \ddot{\gamma}(u), \mathbf{T}_\gamma(u) \rangle \cdot \mathbf{T}_\gamma(u)}{\|\ddot{\gamma}(u) - \langle \ddot{\gamma}(u), \mathbf{T}_\gamma(u) \rangle \cdot \mathbf{T}_\gamma(u)\|}.$$

Exercice. Vérifier que, en chaque point d'une courbe birégulière, les vecteurs $\mathbf{T}_\gamma(u)$ et $\mathbf{N}_\gamma(u)$ forment un repère orthonormé du plan osculateur.

Lemme 2.2 (Règle de Leibniz). Soient $\mathbf{W}(u)$ et $\mathbf{Z}(u)$ deux champs de vecteurs de classe C^1 le long de la courbe $\gamma : I \rightarrow \mathbb{R}^n$, alors

$$\frac{d}{du} \langle \mathbf{W}(u), \mathbf{Z}(u) \rangle = \langle \dot{\mathbf{W}}(u), \mathbf{Z}(u) \rangle + \langle \mathbf{W}(u), \dot{\mathbf{Z}}(u) \rangle.$$

Si $n = 3$, alors on a de même

$$\frac{d}{du} (\mathbf{W}(u) \times \mathbf{Z}(u)) = \dot{\mathbf{W}}(u) \times \mathbf{Z}(u) + \mathbf{W}(u) \times \dot{\mathbf{Z}}(u),$$

et si $n = 2$,

$$\frac{d}{du} (\mathbf{W}(u) \wedge \mathbf{Z}(u)) = \dot{\mathbf{W}}(u) \wedge \mathbf{Z}(u) + \mathbf{W}(u) \wedge \dot{\mathbf{Z}}(u).$$

Preuve. Démontrons la première formule. Pour simplifier on écrit

$$\mathbf{W}(u) \cdot \mathbf{Z}(u) = \langle \mathbf{W}(u), \mathbf{Z}(u) \rangle.$$

On a alors par bilinéarité

$$\begin{aligned} & \mathbf{W}(u + \varepsilon) \cdot \mathbf{Z}(u + \varepsilon) - \mathbf{W}(u) \cdot \mathbf{Z}(u) \\ &= \left(\mathbf{W}(u + \varepsilon) \cdot \mathbf{Z}(u + \varepsilon) - \mathbf{W}(u) \cdot \mathbf{Z}(u + \varepsilon) \right) + \left(\mathbf{W}(u) \cdot \mathbf{Z}(u + \varepsilon) - \mathbf{W}(u) \cdot \mathbf{Z}(u) \right) \\ &= \left(\mathbf{W}(u + \varepsilon) - \mathbf{W}(u) \right) \cdot \mathbf{Z}(u + \varepsilon) + \mathbf{W}(u) \cdot \left(\mathbf{Z}(u + \varepsilon) - \mathbf{Z}(u) \right). \end{aligned}$$

Il suffit de diviser cette identité par ε et faire tendre ε vers 0 pour obtenir le lemme. Les autres formules se vérifient de la même manière. $\square$

Le corollaire suivant est important et sera fréquemment utilisé dans la suite :

Corollaire 2.1 (a) Si $\mathbf{W}_1(u)$ et $\mathbf{W}_2(u)$ sont deux champs de vecteurs de classe C^1 le long de γ tel que $\langle \mathbf{W}_1(u), \mathbf{W}_2(u) \rangle$ est constant, alors on a

$$\langle \mathbf{W}_1(u), \dot{\mathbf{W}}_2(u) \rangle = -\langle \dot{\mathbf{W}}_1(u), \mathbf{W}_2(u) \rangle$$

pour tout $u \in I$.

(b) Si $\mathbf{W}(u)$ est un champ de vecteurs de classe C^1 le long de γ tel que $\|\mathbf{W}\|$ est constant, alors $\dot{\mathbf{W}}(u)$ est orthogonal à $\mathbf{W}(u)$ pour tout $u \in I$.

Preuve. L'affirmation (a) est une conséquence immédiate de la règle de Leibniz et (b) découle de (a). $\square$

2.4 Longueur et abscisse curviligne

Définition. La longueur d'un arc de courbe $\gamma : [a, b] \rightarrow \mathbb{R}^n$ de classe C^1 par morceaux est l'intégrale de sa vitesse :

$$\ell(\gamma) = \int_a^b V_\gamma(t) dt, \quad \text{où } V_\gamma = \|\dot{\gamma}(t)\|.$$

Exemple 2.3. 1) Il est clair que si la vitesse est constante : $V_\gamma(t) \equiv v$, alors on a

$$\ell(\gamma) = v \cdot (b - a).$$

Ainsi, la longueur d'un chemin parcouru à vitesse constante est égale à la vitesse multipliée par le temps de parcours :

$$\text{longueur} = \text{vitesse} \times \text{temps}.$$

2) Comme cas particulier, nous avons le segment $[p, q]$ paramétré par

$$\delta(t) = (p_1 + t(q_1 - p_1), \dots, p_n + t(q_n - p_n)),$$

avec $t \in [0, M]$. On a vu que $V_\delta(t) = \|\vec{pq}\| = \|q - p\|$, et donc

$$\ell(\delta) = M \|q - p\|.$$

3) L'arc de cercle de centre p et rayon r dans $\mathbb{R}^2$ est paramétré par

$$c(\theta) = (p_1 + r \cos(\theta), p_2 + r \sin(\theta)),$$

où θ varie de θ_0 à θ_1 . La vitesse de cette courbe est constante : $V_c(\theta) = r$, et donc

$$\ell(c) = \int_{\theta_0}^{\theta_1} V_c d\theta = r(\theta_1 - \theta_0).$$

On a donc montré que *la longueur d'un arc de cercle est égale au produit du rayon par l'angle qui sous-tend l'arc.*

4) La longueur du graphe $\gamma_f : [a, b] \rightarrow \mathbb{R}^2$ de la fonction $f : [a, b] \rightarrow \mathbb{R}$ est donnée par

$$\ell(\gamma_f) = \int_a^b V_\gamma(x) dx = \int_a^b \sqrt{1 + (f'(x))^2} dx.$$

Voyons à présent quelques propriétés importantes de la longueur.

Proposition 2.4. *Si $g : \mathbb{R}^n \rightarrow \mathbb{R}^n$ est une similitude de rapport $\lambda > 0$ et $\gamma : [a, b] \rightarrow \mathbb{R}^n$ est un arc de courbe de classe C^1 , alors $\tilde{\gamma} := g \circ \gamma : [a, b] \rightarrow \mathbb{R}^n$ est aussi de classe C^1 et*

$$\ell(\tilde{\gamma}) = \lambda \ell(\gamma).$$

En particulier la longueur d'une courbe est invariante par isométrie.

Preuve. On sait que toute similitude g est de la forme $g(x) = \lambda Ax + \mathbf{b}$, où $\mathbf{b}$ est un vecteur et A une matrice orthogonale. On a donc $\tilde{\gamma}(u) = \lambda A\gamma(u) + \mathbf{b}$, et, par la règle de Leibniz,

$$\dot{\tilde{\gamma}}(u) = \lambda \dot{A}\gamma(u) + \lambda A\dot{\gamma}(u) + \dot{\mathbf{b}} = \lambda A\dot{\gamma}(u)$$

puisque A et $\mathbf{b}$ sont constantes. Comme A est une matrice orthogonale, on a

$$V_{\tilde{\gamma}}(u) = \|\dot{\tilde{\gamma}}(u)\| = \|\lambda A\dot{\gamma}(u)\| = \lambda \|\dot{\gamma}(u)\| = \lambda V_\gamma(u),$$

et donc

$$\ell(\tilde{\gamma}) = \int_a^b V_{\tilde{\gamma}}(u) du = \lambda \int_a^b V_\gamma(u) du = \lambda \ell(\gamma).$$

□

Proposition 2.5 (additivité de la longueur). *Soit $\alpha : [a, b] \rightarrow \mathbb{R}^n$ une courbe de classe C^1 et $c \in [a, b]$. Notons $\beta := \alpha|_{[a, c]} : [a, c] \rightarrow \mathbb{R}^n$ et $\gamma := \alpha|_{[c, b]} : [c, b] \rightarrow \mathbb{R}^n$ les restrictions de α aux intervalles $[a, c]$ et $[c, b]$. Alors*

$$\ell(\alpha) = \ell(\beta) + \ell(\gamma).$$

Preuve. Cette proposition découle de la propriété correspondante de l'intégrale, nous laissons le lecteur compléter les détails.

□

Proposition 2.6. *Pour tout arc de courbe $\alpha : [a, b] \rightarrow \mathbb{R}^n$ de classe C^1 on a*

$$d(\alpha(a), \alpha(b)) \leq \ell(\alpha).$$

Cette proposition dit que *le plus court chemin reliant deux points est le segment de droite reliant ces deux points*.

Preuve. Quitte à composer α par une isométrie, on peut supposer que $\alpha(a) = \mathbf{0} = (0, 0, \dots, 0)$ et $\alpha(b) = d\mathbf{e}_1 = (d, 0, \dots, 0)$, où $d = d(\alpha(a), \alpha(b))$. Écrivons $\alpha(t) = (x_1(t), x_2(t), \dots, x_n(t))$ dans un système de coordonnées orthonormé, alors

$$\ell(\alpha) = \int_a^b \|\dot{\alpha}(t)\| dt = \int_a^b \left(\sum_{i=1}^n \dot{x}_i^2(t) \right)^{1/2} dt \geq \int_a^b |\dot{x}_1(t)| dt \geq \int_a^b \dot{x}_1(t) dt = x_1(b) - x_1(a) = d.$$

□

Définition 2.7. Soit $\alpha : I \rightarrow \mathbb{R}^n$ une courbe paramétrée de classe C^1 et $u_0 \in I$ une valeur du paramètre. L'*abscisse curviligne* (aussi appelé le *paramètre naturel*) sur α correspondant au point initial $p_0 = \alpha(u_0)$ est la fonction $s_\alpha : I \rightarrow \mathbb{R}$ définie par

$$s_\alpha(u) = \int_{u_0}^u V_\alpha(\tau) d\tau.$$

L'abscisse curviligne mesure donc la longueur du chemin parcouru sur la courbe depuis le point initial, elle est négative avant le point initial et positive après :

$$s_\alpha(u) = \begin{cases} \ell(\alpha|_{[u_0, u]}) & \text{si } u \geq u_0 \\ -\ell(\alpha|_{[u, u_0]}) & \text{si } u \leq u_0. \end{cases}$$

Lorsqu'il n'y a pas de risque de confusion, nous noterons l'abscisse curviligne par $s(u)$ au lieu de $s_\alpha(u)$.

2.5 Changement de paramétrisation d'une courbe

La notion de courbe que nous avons introduite plus haut est une notion *cinématique*¹, i.e. fondée sur la notion de paramétrisation. Il est naturel, d'un point de vue géométrique, d'admettre qu'une « même » courbe puisse avoir plusieurs paramétrisations distincts.

Définition. Soit $\alpha(t)$ ($t \in I$) une courbe paramétrée. On dit qu'une courbe $\beta(u)$ ($u \in J$) est une *reparamétrisation directe* de α s'il existe une bijection

$$h : I \rightarrow J$$

transformant le paramètre t en $u = h(t)$ et telle que

- a) h est continûment différentiable ;
- b) $h'(t) > 0$ quel que soit $t \in I$;
- c) $\alpha = \beta \circ h$.

Le diagramme triangulaire suivant représente la situation schématiquement :

$$\begin{array}{ccc} I & \xrightarrow{h} & J \\ & \searrow \alpha & \downarrow \beta \\ & & \mathbb{R}^n \end{array}$$

1. Le mot *cinématique* vient du grec *κίνησις*, qui signifie « mouvement ».

Observons que les deux courbes ont alors la même trace, i.e. $\alpha(I) = \beta(J)$. Les vecteurs vitesses sont reliés par

$$\frac{d\alpha}{dt} = \frac{d\beta}{du} \frac{du}{dt} = h'(t) \frac{d\beta}{du} \quad (2.1)$$

En particulier, comme $h'(t) \neq 0$, on voit que les courbes α et β ont les mêmes points singuliers.

Les formules ci-dessus montrent en particulier que lorsqu'on reparamétrise une courbe, celle-ci ne change pas de sens de parcours (car les vecteurs vitesses des deux courbes ont même direction et même sens). On peut toutefois inverser le sens de parcours d'une courbe par une procédure similaire à une reparamétrisation.

Définition. On dit qu'une courbe $\beta(u)$ ($u \in J$) est une *reparamétrisation indirecte*, ou une *inversion* de la courbe $\alpha(t)$ ($t \in I$) s'il existe une bijection

$$h : I \rightarrow J$$

transformant le paramètre t en $u = h(t)$ et telle que

- a) h est continûment différentiable ;
- b) $h'(t) < 0$ quel que soit $t \in I$;
- c) $\alpha = \beta \circ h$.

Si $\beta(u)$ est une reparamétrisation directe ou indirecte de la courbe $\alpha(t)$, alors les vitesses de ces deux courbes sont reliées par

$$V_\alpha(t) = |h'(t)| V_\beta(u). \quad (2.2)$$

Voici un exemple simple : considérons les courbes du plan $\mathbb{R}^2$

$$\alpha(\theta) = (\cos(\theta), \sin(\theta)) \quad (0 < \theta < \pi)$$

et

$$\gamma(x) = (x, \sqrt{1-x^2}), \quad (-1 < x < 1).$$

Ces deux courbes ont la même trace, qui est le demi-cercle unité :

$$\{(x, y) \mid x^2 + y^2 = 1, y > 0\}.$$

La fonction $h : (0, \pi) \rightarrow (-1, 1)$ définie par $h(\theta) = x = \cos(\theta)$ fait le lien entre les deux paramétrisations car

$$\gamma(h(\theta)) = (x, \sqrt{1-x^2}) = (\cos(\theta), \sin(\theta)) = \alpha(\theta).$$

Comme $h'(\theta) = \frac{dx}{d\theta} = -\sin(\theta) < 0$, on voit que la courbe α est une inversion de γ .

Remarque. Observons que si $\theta = 0$ ou $\theta = \pi$, alors $h'(\theta) = 0$. Le reparamétrage h cesse d'être admissible aux extrémités de l'intervalle. Cela correspond au fait que la vitesse de γ

$$V_\gamma(x) = \left\| \frac{d\gamma}{dx} \right\| = \frac{1}{\sqrt{1-x^2}}$$

tend vers l'infini lorsque $x \rightarrow \pm 1$.

2.6 Quantités géométriques et quantités cinématiques

Définition. Une quantité ou une notion attachée à une courbe est dite *géométrique* si elle est invariante par rapport aux changements de paramètres, et elle est dite *cinématique* dans le cas contraire.

Par exemple, la vitesse et l'accélération sont des notions cinématiques alors que la notion de point singulier, de point régulier et de direction tangente sont des notions géométriques.

Lemme 2.8. *Le vecteur tangent $\mathbf{T}_\alpha(t)$ est une quantité géométrique.*

Cette affirmation est géométriquement évidente, puisque $\mathbf{T}$ est le champ de vecteurs unitaire indiquant la direction de la courbe. Voyons tout de même une preuve formelle :

Preuve. Soit $\beta(u)$ ($u \in J$) une reparamétrisation directe de la courbe $\alpha(t)$. Il existe alors une fonction $h : I \rightarrow J$ telle que $h'(t) > 0$ et $\alpha(t) = \beta(h(t))$. On sait que $V_\alpha(t) = V_\beta(u)h'(t)$, par conséquent

$$\begin{aligned} \mathbf{T}_\alpha(t) &= \frac{1}{V_\alpha(t)} \frac{d\alpha}{dt} = \frac{1}{V_\alpha(t)} \frac{d\beta(h(t))}{dt} \\ &= \frac{h'(t)}{V_\alpha(t)} \frac{d\beta(u)}{du} \\ &= \frac{1}{V_\beta(u)} \frac{d\beta(u)}{du} \\ &= \mathbf{T}_\beta(u). \end{aligned}$$

□

Remarque. Si $\beta(u)$ est une reparamétrisation indirecte de $\alpha(t)$, alors on a $\mathbf{T}_\alpha(t) = -\mathbf{T}_\beta(u)$.

La longueur d'une courbe est également une quantité géométrique ; plus généralement, nous avons la proposition suivante.

Proposition 2.2 *Soient α et β deux courbes de classe C^1 . Si β est une reparamétrisation de α , alors $\ell(\beta) = \ell(\alpha)$.*

Preuve. Considérons d'abord le cas où $\beta(u)$ ($a' \leq u \leq b'$) est une reparamétrisation directe de la courbe $\alpha(t)$ ($a \leq t \leq b$). En utilisant la formule de changement de variables dans les intégrales, on a

$$\begin{aligned} \ell(\beta) &= \int_{a'}^{b'} V_\beta(u) du = \int_a^b V_\beta(u) \frac{du}{dt} dt \\ &= \int_a^b V_\alpha(t) dt \\ &= \ell(\alpha). \end{aligned}$$

Dans le cas où $\beta(u)$ est une reparamétrisation indirecte de $\alpha(t)$, alors $\frac{du}{dt} < 0$ et on a

$$\begin{aligned} \ell(\beta) &= \int_{a'}^{b'} V_\beta(u) du = \int_b^a V_\beta(u) \frac{du}{dt} dt \\ &= - \int_b^a V_\alpha(t) dt \\ &= \int_a^b V_\alpha(t) dt \\ &= \ell(\alpha). \end{aligned}$$

□

Considérons par exemple l'arc du cercle unité dans le plan $\mathbb{R}^2$, reliant les points $(1, 0)$ et (x_0, y_0) et contenu dans le demi-plan $y > 0$. La longueur de cet arc est donnée par

$$\ell = \theta = \text{Arcos}(x_0).$$

Si cette courbe est paramétrée comme un graphe, i.e. par $\gamma(x) = (x, \sqrt{1-x^2})$, $(x_0 < x < 1)$, alors la vitesse est $V_\gamma(x) = \frac{1}{\sqrt{1-x^2}}$ et la longueur est donc donnée par

$$\ell = \int_{x_0}^1 \frac{du}{\sqrt{1-u^2}}.$$

La proposition 2.2 nous permet de déduire du résultat précédent l'identité analytique :

$$\int_{x_0}^1 \frac{du}{\sqrt{1-u^2}} = \text{Arcos}(x_0),$$

que nous avons obtenue (presque) sans aucun calcul, mais par un raisonnement purement géométrique.

2.7 Paramétrisation naturelle d'une courbe régulière

Théorème 2.9. Soit $\alpha : I \rightarrow \mathbb{R}^n$ une courbe régulière de classe C^1 et $t_0 \in I$ une valeur du paramètre. Alors il existe une unique reparamétrisation directe $h : I \rightarrow J$, telle que $0 \in J$, $h(t_0) = 0$ et $\beta := \alpha \circ h^{-1} : J \rightarrow \mathbb{R}^n$ est de vitesse 1, i.e. $V_\beta(s) = 1$.

Preuve. Montrons d'abord l'unicité de cette reparamétrisation. On a vu plus haut (p. 23) que

$$V_\alpha(t) = h'(t)V_\beta(s).$$

Comme $V_\beta(s) = 1$ et $h' > 0$, on a donc $h'(t) = V_\alpha(t)$ et comme $h(t_0) = 0$, on doit avoir

$$s = h(t) = \int_{t_0}^t V_\alpha(\tau) d\tau.$$

Ainsi la fonction $h(t)$ coïncide avec l'abscisse curviligne $s(t)$.

Pour montrer l'existence de cette reparamétrisation, on définit à présent h par $h(t) = s(t) = \int_{t_0}^t V_\alpha(\tau) d\tau$ et l'intervalle J par $J = h(I)$. Alors $h(t_0) = 0$ et $h'(t) = V_\alpha(t)$. En utilisant la formule (2.1) de la page 23, on voit que la courbe $\beta := \alpha \circ h^{-1} : J \rightarrow \mathbb{R}^n$ vérifie

$$V_\beta(s) = V_\alpha(t) \frac{1}{h'(t)} = 1.$$

□

Définition 2.10. On dit qu'une courbe régulière γ est *paramétrée naturellement* si $\|\dot{\gamma}(s)\| = 1$ pour tout s , i.e. si sa vitesse vaut 1. Le théorème précédent nous dit que toute courbe régulière de classe C^1 peut être reparamétrée naturellement.

Dès qu'un point initial et un sens de parcours ont été choisis sur la courbe, la paramétrisation naturelle est unique et elle est donnée par l'abscisse curviligne.

Méthode Pour trouver la paramétrisation naturelle d'une courbe α , il faut effectuer les opérations suivantes :

- (1) Identifier ou choisir le point initial u_0 .
- (2) Calculer la vitesse $V_\alpha(u) = \|\dot{\alpha}(u)\|$.
- (3) Intégrer V_α pour obtenir l'abscisse curviligne $s : s(u) = \int_{u_0}^u V_\alpha(\tau) d\tau$.
- (4) Inverser la relation $s = s(u)$ (i.e. exprimer u en fonction de $s : u = u(s)$).
- (5) On obtient alors la paramétrisation naturelle $\beta(s) = \alpha(u(s))$.

Dans la pratique, les points qui peuvent être délicats sont les étapes (3) et (4).

Exemple 2.11. La *chaînette*² est la courbe plane paramétrée par $\alpha(u) = (u, \cosh u)$. Le vecteur vitesse est $\dot{\alpha}(u) = (1, \sinh(u))$, et donc

$$V_\alpha(u) = \|\dot{\alpha}(u)\| = \sqrt{1 + \sinh^2(u)} = \cosh(u).$$

L'abscisse curviligne depuis le point initial $\alpha(0) = (0, 1)$ est donnée par l'intégrale

$$s(u) = \int_0^u V_\alpha(t) dt = \int_0^u \cosh(t) dt = \sinh(u),$$

et on a donc

$$u(s) = \operatorname{argsh}(s) = \log(s + \sqrt{1 + s^2}).$$

Remarquons que $\cosh(u) = \sqrt{1 + \sinh^2(u)} = \sqrt{1 + s^2}$. En substituant cette relation dans la paramétrisation de α , on obtient la paramétrisation naturelle de la chaînette :

$$\beta(s) = \alpha(u(s)) = (u(s), \cosh u(s)) = (\operatorname{argsh}(s), \sqrt{1 + s^2}).$$

On vérifie facilement que $\|\dot{\beta}(s)\| = 1$.

Les courbes pour lesquelles on peut effectivement calculer la paramétrisation naturelle sont plutôt rares ; mais cette notion joue un rôle théorique fondamental. Il faut en particulier se souvenir des relations suivantes qui relient le paramètre naturel s au paramètre donné u .

$$\boxed{ds = V(u) \cdot du \quad \text{et} \quad \frac{d}{ds} = \frac{1}{V(u)} \frac{d}{du}.} \quad (2.3)$$

Remarque. L'abscisse curviligne joue un rôle fondamental en théorie des courbes, car c'est dans la paramétrisation naturelle que les relations fondamentales entre les différentes quantités géométriques liées à une courbe sont le plus clairement mises en évidence. Pour cette raison, les livres traitant de courbes choisissent souvent d'écrire les formules relativement à la seule abscisse curviligne. Nous n'avons pas fait ce choix et avons préféré écrire les formules par rapport à un paramètre général en raison de la difficulté pratique de calculer la paramétrisation naturelle pour la plupart des courbes. Nous invitons toutefois le-la lecteur-ice à récrire elle-lui-même les formules des prochains paragraphes dans le cas spécial d'une courbe paramétrée naturellement ; on constatera ainsi combien les formules et les calculs théoriques se simplifient.

2. La chaînette est ainsi appelée car elle modélise la forme que prend naturellement une chaîne ou un câble suspendu entre deux points fixes sous l'effet de la gravité

2.8 Courbure d'une courbe de $\mathbb{R}^n$

Définition. Le *vecteur de courbure* d'une courbe régulière $\alpha : I \rightarrow \mathbb{R}^n$ de classe C^2 est le champ de vecteurs le long de cette courbe défini par

$$\mathbf{K}_\alpha(t) := \frac{1}{V_\alpha(t)} \frac{d}{dt} \mathbf{T}_\alpha(t),$$

où $\mathbf{T}_\alpha(t) = \frac{1}{V_\alpha(t)} \dot{\alpha}(t)$ est le vecteur tangent à la courbe.

REMARQUE. Le corollaire 2.1 entraîne que le vecteur de courbure est toujours orthogonal au vecteur tangent :

$$\mathbf{K}_\alpha(t) \perp \mathbf{T}_\alpha(t).$$

On définit aussi la *courbure* de la courbe α . C'est par définition la norme du vecteur de courbure :

$$\kappa_\alpha(u) := \|\mathbf{K}_\alpha(u)\|.$$

Il est facile de voir que la courbure d'une droite est nulle, voici un autre exemple simple.

Exemple 2.12. Une paramétrisation d'un cercle de centre p et rayon r est donnée par

$$c(\theta) = p + r \cos(\theta) \mathbf{u}_1 + r \sin(\theta) \mathbf{u}_2,$$

où $p, \mathbf{u}_1, \mathbf{u}_2$ est un repère orthonormé du plan contenant le cercle. On a

$$\dot{c}(\theta) = -r \sin(\theta) \mathbf{u}_1 + r \cos(\theta) \mathbf{u}_2,$$

donc $V_c(\theta) = r$ et $\mathbf{T}(c, \theta) = -\sin(\theta) \mathbf{u}_1 + \cos(\theta) \mathbf{u}_2$.

En dérivant le vecteur tangent, on a $\dot{\mathbf{T}}(c, \theta) = -\cos(\theta) \mathbf{u}_1 - \sin(\theta) \mathbf{u}_2$, donc

$$\mathbf{K}(c, \theta) = \frac{1}{V_c(\theta)} \dot{\mathbf{T}}(c, \theta) = -\frac{1}{r} (\cos(\theta) \mathbf{u}_1 + \sin(\theta) \mathbf{u}_2)$$

et finalement :

$$\kappa(c, \theta) = \|\mathbf{K}(c, \theta)\| = \frac{1}{r}.$$

La courbure d'un cercle est donc l'inverse de son rayon.

Remarquons aussi qu'on a la relation suivante exprimant le centre du cercle en fonction d'un point sur le cercle et de la courbure :

$$c(\theta) + r^2 \mathbf{K}(c, \theta) = p. \quad (2.4)$$

Proposition 2.13. Le vecteur de courbure $\mathbf{K}_\alpha(t)$ et la courbure $\kappa_\alpha(t) = \|\mathbf{K}_\alpha(t)\|$ sont des quantités géométriques.

Nous laissons la preuve en exercice.

Proposition 2.14 (Formule de l'accélération). Le vecteur accélération d'une courbe régulière $\gamma : I \rightarrow \mathbb{R}^n$ vérifie

$$\ddot{\gamma}(u) = (V_\gamma(u))^2 \mathbf{K}_\gamma(u) + \dot{V}_\gamma(u) \mathbf{T}_\gamma(u). \quad (2.5)$$

Preuve. Écrivons le vecteur vitesse sous la forme $\dot{\gamma}(u) = V_\gamma(u)\mathbf{T}_\gamma(u)$ et dérivons ce vecteur :

$$\begin{aligned}\ddot{\gamma}(u) &= \frac{d}{du}(V_\gamma(u)\mathbf{T}_\gamma(u)) = V_\gamma(u)\dot{\mathbf{T}}_\gamma(u) + \dot{V}_\gamma(u)\mathbf{T}_\gamma(u) \\ &= (V_\gamma(u))^2 \mathbf{K}_\gamma(u) + \dot{V}_\gamma(u)\mathbf{T}_\gamma(u).\end{aligned}$$

□

On dit que $\dot{V}_\gamma(u)\mathbf{T}_\gamma(u)$ est l'*accélération tangentielle* et $(V_\gamma(u))^2 \mathbf{K}_\gamma(u)$ est l'*accélération normale* de γ .

En mécanique, cette formule signifie que la force subie par une particule en mouvement est fonction de l'accélération tangentielle $\dot{V}_\gamma$ et du carré de la vitesse multiplié par la courbure.

Corollaire 2.15. *Si α est paramétrée naturellement, i.e. si $V_\alpha \equiv 1$, alors*

$$\ddot{\alpha}(s) = \mathbf{K}(\alpha, s).$$

Preuve. Puisque $V_\alpha \equiv 1$, on a $\dot{V}_\alpha = 0$ et le corollaire se déduit immédiatement de la formule de l'accélération. □

Proposition 2.16. *Une courbe $\alpha : [a, b] \rightarrow \mathbb{R}^n$ est de courbure nulle si et seulement si c'est une droite ou un segment de droite (qui peut être paramétrée arbitrairement).*

Preuve. On peut supposer grâce au théorème 2.9 que α est paramétrée naturellement. Le corollaire précédent entraîne alors que $\ddot{\alpha}(s) = \mathbf{K}(\alpha, s)$ et comme $\kappa(\alpha, s) = \|\mathbf{K}(\alpha, s)\| = 0$, on a donc $\ddot{\alpha}(s) = \mathbf{0}$. Le vecteur $\mathbf{v} := \dot{\alpha}$ est alors constant et on obtient donc en intégrant

$$\alpha(s) = p + s\mathbf{v}$$

où $p = \alpha(0)$. □

Lemme 2.3 *Le vecteur de courbure d'une courbe de classe C^2 en un point est un multiple du vecteur normal principal en ce point :*

$$\mathbf{K}_\gamma(u) = \kappa_\gamma(u)\mathbf{N}_\gamma(u).$$

Preuve. Nous laissons la preuve en exercice. Rappelons que le vecteur normal principal est défini par

$$\mathbf{N}_\gamma(u) = \frac{\ddot{\gamma}(u) - \langle \ddot{\gamma}(u), \mathbf{T}_\gamma(u) \rangle \cdot \mathbf{T}_\gamma(u)}{\|\ddot{\gamma}(u) - \langle \ddot{\gamma}(u), \mathbf{T}_\gamma(u) \rangle \cdot \mathbf{T}_\gamma(u)\|}.$$

□

Proposition 2.4 *La courbure d'une courbe de classe C^2 est la variation angulaire de la direction de cette courbe par rapport au paramètre naturel.*

La signification exacte de cette proposition sera précisée dans la preuve.

Preuve. Soit $\gamma : I \rightarrow \mathbb{R}^3$ une courbe de classe C^2 paramétrée naturellement et notons

$$\varphi(s_0, s) := \angle(\mathbf{T}(s_0), \mathbf{T}(s))$$

l'angle entre $\mathbf{T}(s_0)$ et $\mathbf{T}(s)$ (où $s_0, s \in I$). Comme $\|\mathbf{T}(s_0)\| = \|\mathbf{T}(s)\| = 1$, on a par la trigonométrie élémentaire que

$$\|\mathbf{T}(s_0) - \mathbf{T}(s)\| = 2 \sin\left(\frac{\varphi(s_0, s)}{2}\right).$$

On a donc

$$\begin{aligned} \lim_{s \rightarrow s_0^+} \frac{\varphi(s_0, s)}{s - s_0} &= \lim_{s \rightarrow s_0^+} \left(\frac{\varphi(s_0, s)}{2 \sin\left(\frac{\varphi(s_0, s)}{2}\right)} \cdot \frac{2 \sin\left(\frac{\varphi(s_0, s)}{2}\right)}{s - s_0} \right) \\ &= \lim_{s \rightarrow s_0^+} \left(\frac{\varphi(s_0, s)}{2 \sin\left(\frac{\varphi(s_0, s)}{2}\right)} \right) \cdot \lim_{s \rightarrow s_0^+} \left(\frac{2 \sin\left(\frac{\varphi(s_0, s)}{2}\right)}{s - s_0} \right) \\ &= \lim_{s \rightarrow s_0^+} \left\| \frac{\mathbf{T}(s_0) - \mathbf{T}(s)}{s - s_0} \right\| = \|\dot{\mathbf{T}}(s_0)\|, \end{aligned}$$

car $\lim_{s \rightarrow s_0^+} \left(\frac{\varphi(s_0, s)}{2 \sin\left(\frac{\varphi(s_0, s)}{2}\right)} \right) = 1.$

On a ainsi montré que la courbure est la dérivée à droite de l'angle, on peut noter

$$\kappa(s_0) = \lim_{s \rightarrow s_0^+} \frac{\varphi(s_0, s)}{s - s_0} = \left. \frac{d}{ds} \right|_{s_0^+} \varphi(s_0, s) \quad (2.6)$$

□

2.9 Contact entre deux courbes

Définition. On dit que deux courbes de classe C^k

$$\alpha, \beta : I \rightarrow \mathbb{R}^n$$

ayant le même paramètre $t \in I$ ont un *contact d'ordre k* en $t_0 \in I$ si $\alpha(t_0) = \beta(t_0)$ et si leurs dérivées en t_0 coïncident jusqu'à l'ordre k :

$$\frac{d^m \alpha}{dt^m}(t_0) = \frac{d^m \beta}{dt^m}(t_0),$$

pour $m = 1, 2, \dots, k$.

Ainsi, deux courbes α, β ont un contact d'ordre 0 en t_0 si elles passent par le même point en t_0 . Elles ont un contact d'ordre 1 si elles passent par le même point et elles ont le même vecteur vitesse en ce point :

$$\alpha(t_0) = \beta(t_0), \quad \frac{d\alpha}{dt}(t_0) = \frac{d\beta}{dt}(t_0).$$

Concernant les courbes ayant un contact d'ordre 2, nous avons le résultat suivant :

Théorème 2.17. *Deux courbes α, β de classe C^2 ont un contact d'ordre 2 en t_0 si et seulement si elles passent par le même point et si elles ont le même vecteur vitesse, la même accélération tangentielle et le même vecteur de courbure en ce point :*

$$\alpha(t_0) = \beta(t_0), \quad \frac{d\alpha}{dt}(t_0) = \frac{d\beta}{dt}(t_0), \quad \dot{V}_\alpha(t_0) = \dot{V}_\beta(t_0) \quad \text{et} \quad \mathbf{K}_\alpha(t_0) = \mathbf{K}_\beta(t_0).$$

En particulier, si ces deux courbes sont birégulières alors elles ont le même plan osculateur en t_0 .

Preuve. C'est une conséquence directe de la proposition 2.14.

□

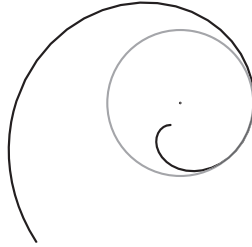
A titre d'exemple important, nous avons le corollaire suivant :

Corollaire 2.18. *Deux cercles de $\mathbb{R}^n$ parcourus à vitesse constante ont un contact d'ordre 2 si et seulement si ces deux cercles coïncident.*

Preuve. Par le théorème précédent, les deux cercles ont le même vecteur vitesse et le même vecteur de courbure en leur point de contact. En particulier les deux cercles se situent dans le même plan, qui est le plan osculateur commun.. On sait en outre par l'exemple 2.12 que la courbure d'un cercle est égale à l'inverse de son rayon. Les deux cercles ont donc même rayon r . Mais on sait aussi par l'équation (2.4) que les deux cercles doivent avoir même centre, donc ils coïncident. $\square$

Théorème 2.19. *Soit $\alpha : I \rightarrow \mathbb{R}^n$ une courbe de classe C^2 qui est birégulière en $t_0 \in I$. Alors il existe un cercle $c : I \rightarrow \mathbb{R}^n$ ayant un contact d'ordre 2 avec α en t_0 . Ce cercle est unique, son rayon est l'inverse de $|\kappa_\alpha(t_0)|$ et son centre est donné par*

$$p = \alpha(t_0) + \frac{1}{\kappa_\alpha(t_0)^2} \mathbf{K}_\alpha(t_0). \quad (2.7)$$



Cercle osculateur.

Définition. Ce cercle s'appelle le *cercle osculateur*³, aussi appelé le *cercle de courbure* à α en t_0 . C'est parmi tous les cercles celui qui approxime le mieux la courbe au voisinage de $\alpha(t_0)$. Il est contenu dans le plan osculateur à la courbe en ce point. Son centre est appelé le *centre de courbure* et son rayon est le *rayon de courbure* de α en t_0 . On le note

$$\rho_\alpha(t_0) = \frac{1}{\kappa_\alpha(t_0)}.$$

Preuve. Supposons pour la preuve que la courbe α est paramétrée naturellement (et notons selon l'usage s le paramètre naturel). On supposera aussi que le point considéré correspond à la valeur $s = 0$ du paramètre.

Notons $\rho = \frac{1}{\kappa_\alpha(0)}$, $\mathbf{T} := \mathbf{T}_\alpha(0) = \dot{\alpha}(0)$ et $\mathbf{N} := \rho \mathbf{K}_\alpha(0)$, puis posons

$$p := \alpha(0) + \rho \mathbf{N} \quad (2.8)$$

et considérons le cercle de centre p et rayon ρ que nous paramétrisons par

$$\gamma(s) = p - \rho \cos\left(\frac{s}{\rho}\right) \mathbf{N} + \rho \sin\left(\frac{s}{\rho}\right) \mathbf{T}.$$

Il s'agit bien d'un cercle, puisque $\mathbf{T}$ et $\mathbf{N}$ sont orthogonaux et de longueur 1. Il est alors clair que $\gamma(0) = \alpha(0)$ et que $\dot{\gamma}(0) = \mathbf{T} = \dot{\alpha}(0)$. On sait d'autre part que la courbure du cercle de rayon ρ est constante et égale à $\frac{1}{\rho} = \kappa_\alpha(0)$. Le théorème 2.17 entraîne donc que la courbe α et le cercle γ ont un contact d'ordre 2 en $s = 0$.

3. Le terme *osculateur* nous vient du latin et signifie *embrasser* : le cercle osculateur embrasse la courbe au point de contact.

L'unicité de ce cercle découle immédiatement du corollaire 2.18 (et du fait que la notion de contact entre courbes est clairement transitive).

□

Définition. On appelle développée de la courbe birégulière $\alpha : I \rightarrow \mathbb{R}^3$ la courbe

$$\beta(u) = \alpha(u) + \rho_\alpha(u)\mathbf{N}_\alpha(u),$$

où ρ_α est le rayon de courbure de α . On dit aussi que β est la développante de α . La développée suit le mouvement du centre du cercle osculateur lorsqu'on parcourt la courbe α .

2.10 Le repère de Frenet d'une courbes dans $\mathbb{R}^3$

Rappelons que le vecteur tangent et le vecteur normal principal d'une courbe birégulière $\gamma : I \rightarrow \mathbb{R}^3$ de classe C^2 sont les champs de vecteurs le long de cette courbe définis par

$$\mathbf{T}_\gamma(u) = \frac{\dot{\gamma}(u)}{V_\gamma(u)} \quad \text{et} \quad \mathbf{N}_\gamma(u) = \frac{\ddot{\gamma}(u) - \langle \ddot{\gamma}(u), \mathbf{T}_\gamma(u) \rangle \cdot \mathbf{T}_\gamma(u)}{\|\ddot{\gamma}(u) - \langle \ddot{\gamma}(u), \mathbf{T}_\gamma(u) \rangle \cdot \mathbf{T}_\gamma(u)\|}.$$

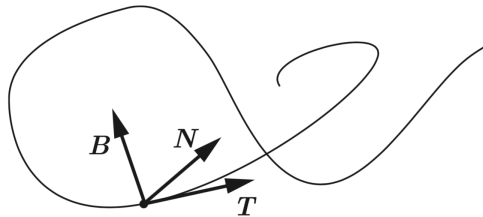
Définitions 1. Le vecteur binormal de γ est le produit vectoriel du vecteur unitaire tangent et du vecteur normal principal :

$$\mathbf{B}_\gamma(u) = \mathbf{T}_\gamma(u) \times \mathbf{N}_\gamma(u).$$

2. Le Repère de Frenet⁴ de γ en u est le repère défini par les trois champs de vecteurs

$$\{\mathbf{T}_\gamma(u), \mathbf{N}_\gamma(u), \mathbf{B}_\gamma(u)\}.$$

Le repère de Frenet est uniquement défini aux points où la courbe est birégulière. C'est un repère mobile (les trois vecteurs sont des champs qui dépendent du paramètre u), il est orthonormé et direct. Il suit la courbe en ce sens que le premier vecteur de ce repère, $\mathbf{T}_\gamma(u)$, est toujours tangent à celle-ci.



Rappelons que le plan passant par $\gamma(u)$ de directions $\mathbf{T}_\gamma(u)$ et $\mathbf{N}_\gamma(u)$ est le *plan osculateur*. Le plan de directions $\mathbf{B}_\gamma(u)$ et $\mathbf{N}_\gamma(u)$ s'appelle le *plan normal* et le plan de directions $\mathbf{T}_\gamma(u)$ et $\mathbf{B}_\gamma(u)$ est le *plan rectifiant*.

Lemme 2.20. Le vecteur binormal à la courbe γ peut aussi s'écrire

$$\mathbf{B}_\gamma(u) = \frac{\dot{\gamma}(u) \times \ddot{\gamma}(u)}{\|\dot{\gamma}(u) \times \ddot{\gamma}(u)\|}.$$

Preuve. On a

$$\dot{\gamma} \times (\ddot{\gamma} - \langle \ddot{\gamma}, \mathbf{T} \rangle \cdot \mathbf{T}) = \dot{\gamma} \times \ddot{\gamma}$$

4. Jean Frédéric Frenet, mathématicien et astronome français 1816-1900.

car $\dot{\gamma} \times \mathbf{T} = 0$. Donc

$$\mathbf{B} = \mathbf{T} \times \mathbf{N} = \frac{\dot{\gamma}}{V_\gamma} \times \frac{(\ddot{\gamma} - \langle \ddot{\gamma}, \mathbf{T} \rangle \cdot \mathbf{T})}{\|\ddot{\gamma} - \langle \ddot{\gamma}, \mathbf{T} \rangle \cdot \mathbf{T}\|} = \frac{\dot{\gamma} \times \ddot{\gamma}}{V_\gamma \|\ddot{\gamma} - \langle \ddot{\gamma}, \mathbf{T} \rangle \cdot \mathbf{T}\|} = \frac{\dot{\gamma}(u) \times \ddot{\gamma}(u)}{\|\dot{\gamma}(u) \times \ddot{\gamma}(u)\|}.$$

□

Définition. (i) On dit qu'une courbe $\gamma : I \rightarrow \mathbb{R}^3$ est *régulière au sens de Frenet*, ou *Frenet régulière* si elle est de classe C^2 , birégulière, et que le vecteur normal principal est un champ $u \mapsto \mathbf{N}_\gamma(u)$ de classe C^1 . ,

(ii) La *torsion* d'une courbe $\gamma : I \rightarrow \mathbb{R}^3$ régulière au sens de Frenet est la fonction $\tau : I \rightarrow \mathbb{R}$ définie par

$$\tau(u) := \frac{\langle \dot{\mathbf{N}}(u), \mathbf{B}(u) \rangle}{V_\gamma(u)}.$$

Il est clair que toute courbe de classe C^3 birégulière est régulière au sens de Frenet, mais la réciproque n'est pas vraie.

Théorème 2.21 (Formules de Serret-Frenet). *Soit $\gamma : I \rightarrow \mathbb{R}^3$ une courbe régulière au sens de Frenet, Alors le repère de Frenet est de classe C^1 et ses dérivées sont données par les formules*

$$\begin{aligned} \frac{1}{V_\gamma(u)} \dot{\mathbf{T}}(u) &= \kappa(u) \mathbf{N}(u), \\ \frac{1}{V_\gamma(u)} \dot{\mathbf{N}}(u) &= -\kappa(u) \mathbf{T}(u) + \tau(u) \mathbf{B}(u), \\ \frac{1}{V_\gamma(u)} \dot{\mathbf{B}}(u) &= -\tau(u) \mathbf{N}(u). \end{aligned}$$

On verra au théorème 2.25 que la courbure et la torsion déterminent complètement la géométrie d'une courbe dans $\mathbb{R}^3$, par conséquent les formules de Serret-Frenet englobent la totalité de la théorie des courbes de $\mathbb{R}^3$.

Preuve. Le vecteur tangent $\mathbf{T}(u)$ est une fonction de classe C^1 du paramètre u car la courbe est supposée de classe C^2 . Le vecteur normal principal $\mathbf{N}(u)$ est une fonction de classe C^1 par hypothèse et le vecteur binormal est une fonction de classe C^1 car $\mathbf{B}(u) = \mathbf{T}(u) \times \mathbf{N}(u)$.

La première équation est une conséquence immédiate des égalités

$$\dot{\mathbf{T}}(u) = V_\gamma(u) \mathbf{K}(u) = V_\gamma(u) \kappa(u) \mathbf{N}(u).$$

Pour prouver la deuxième équation, on remarque d'abord que

$$\dot{\mathbf{N}}(u) = \langle \dot{\mathbf{N}}(u), \mathbf{T}(u) \rangle \mathbf{T}(u) + \langle \dot{\mathbf{N}}(u), \mathbf{N}(u) \rangle \mathbf{N}(u) + \langle \dot{\mathbf{N}}(u), \mathbf{B}(u) \rangle \mathbf{B}(u),$$

car $\{\mathbf{T}_\gamma(u), \mathbf{N}_\gamma(u), \mathbf{B}_\gamma(u)\}$ est un repère orthonormé. D'autre part, on a $\langle \dot{\mathbf{N}}(u), \mathbf{N}(u) \rangle = 0$ car la norme de $\mathbf{N}(u)$ est constante et

$$\langle \dot{\mathbf{N}}(u), \mathbf{T}(u) \rangle \mathbf{T}(u) = -\langle \mathbf{N}(u), \dot{\mathbf{T}}(u) \rangle \mathbf{T}(u) = -V_\gamma(u) \kappa(u) \mathbf{T}(u).$$

On a donc

$$\dot{\mathbf{N}}(u) = -V_\gamma(u) \kappa(u) \mathbf{T}(u) + V_\gamma(u) \tau(u) \mathbf{B}(u),$$

car $\langle \dot{\mathbf{N}}(u), \mathbf{B}(u) \rangle = V_\gamma(u) \tau(u)$ par définition de la torsion.

Pour prouver la troisième équation, on part de

$$\dot{\mathbf{B}}(u) = \langle \dot{\mathbf{B}}(u), \mathbf{T}(u) \rangle \mathbf{T}(u) + \langle \dot{\mathbf{B}}(u), \mathbf{N}(u) \rangle \mathbf{N}(u) + \langle \dot{\mathbf{B}}(u), \mathbf{B}(u) \rangle \mathbf{B}(u),$$

on a $\langle \dot{\mathbf{B}}(u), \mathbf{B}(u) \rangle = 0$ car $\|\mathbf{B}\|$ est constante et

$$\langle \dot{\mathbf{B}}(u), \mathbf{T}(u) \rangle = -\langle \mathbf{B}(u), \dot{\mathbf{T}}(u) \rangle = -V_\gamma(u)\kappa(u)\langle \mathbf{B}(u), \mathbf{N}(u) \rangle = 0.$$

Donc

$$\dot{\mathbf{B}}(u) = \langle \dot{\mathbf{B}}(u), \mathbf{N}(u) \rangle \mathbf{N}(u) = -\langle \mathbf{B}(u), \dot{\mathbf{N}}(u) \rangle \mathbf{N}(u) = -V_\gamma(u)\tau(u)\mathbf{N}(u).$$

Le théorème est démontré. □

Exemple. Rappelons que l'hélice circulaire est la courbe $\gamma(u) = (a \cos(u), a \sin(u), bu)$, on a

$$\dot{\gamma}(u) = (-a \sin(u), a \cos(u), b), \quad \ddot{\gamma}(u) = a(-\cos(u), -\sin(u), 0) \quad \text{et} \quad V_\gamma = \sqrt{a^2 + b^2}.$$

Dans la suite, on suppose $a > 0$ et on notera $c := V_\gamma = \sqrt{a^2 + b^2}$. Le repère de Frenet est donc

$$\mathbf{T} = \frac{1}{c} \begin{pmatrix} -a \sin(u) \\ a \cos(u) \\ b \end{pmatrix} \quad \mathbf{N} = - \begin{pmatrix} \cos(u) \\ \sin(u) \\ 0 \end{pmatrix} \quad \mathbf{B} = \frac{1}{c} \begin{pmatrix} b \sin(u) \\ -b \cos(u) \\ a \end{pmatrix}$$

On trouve la courbure et la torsion en dérivant $\mathbf{N}$:

$$\kappa = -\frac{1}{c} \langle \dot{\mathbf{N}}, \mathbf{T} \rangle = \frac{a}{c^2} \quad \text{et} \quad \tau = \frac{1}{c} \langle \dot{\mathbf{N}}, \mathbf{B} \rangle = \frac{b}{c^2}.$$

Les résultats qui suivent vont mener à une interprétation géométrique de la torsion.

Proposition 2.22. Une courbe $\gamma : I \rightarrow \mathbb{R}^3$ régulière au sens de Frenet est située dans un plan si et seulement si sa torsion est identiquement nulle.

Preuve. Il est clair à partir du Lemme 2.20 que si la courbe γ est située dans un plan $\Pi \subset \mathbb{R}^3$, alors le vecteur binormal est constant (c'est l'un des deux vecteurs unitaires orthogonal à Π). La troisième formule de Serret-Frenet entraîne alors que $\tau_\gamma(u) = 0$.

Réciproquement, supposons $\tau(u) \equiv 0$, alors par la troisième formule de Serret-Frenet, le vecteur binormal est constant. Notons ce vecteur par $\mathbf{B}$ et posons

$$h(u) := \langle \gamma(u) - \gamma(u_0), \mathbf{B} \rangle,$$

et remarquons que

$$\frac{dh}{du} := \langle \dot{\gamma}(u), \mathbf{B} \rangle = V_\gamma(u) \langle \mathbf{T}(u), \mathbf{B} \rangle = 0.$$

Par conséquent h est constante, et comme $h(u_0) = 0$, la fonction h est identiquement nulle, ce qui montre que la courbe γ est contenue dans le plan d'équation

$$\langle \mathbf{x} - \gamma(u_0), \mathbf{B} \rangle = 0,$$

qui n'est autre que le plan orthogonal à $\mathbf{B}$ passant par $\gamma(u_0)$. □

2.10.1 Variation angulaire du plan osculateur

La proposition suivante donne l'interprétation géométrique de la torsion :

Proposition 2.23. *La torsion d'une courbe régulière au sens de Frenet mesure la variation angulaire de son plan osculateur.*

Preuve. Soit $\gamma : I \rightarrow \mathbb{R}^3$ une courbe régulière au sens de Frenet, que nous supposons paramétrée naturellement. Notons $\theta(s_0, s)$ l'angle entre les plans osculateurs à γ en s_0 et en s , et remarquons que

$$\theta(s_0, s) := \angle(\mathbf{B}(s_0), \mathbf{B}(s))$$

car le vecteur normal au plan osculateur en un point de γ est le vecteur binormal en ce point. Le lemme est alors une conséquence de l'égalité suivante :

$$|\tau(s_0)| = \|\dot{\mathbf{B}}(s_0)\| = \lim_{s \rightarrow s_0^+} \frac{\theta(s_0, s)}{s - s_0},$$

qui se prouve de la même manière que la formule (2.6)).

□

2.10.2 Courbes de pente constante

Définition. On dit qu'une courbe de classe C^1 dans $\mathbb{R}^3$ est de *pente constante* si elle est régulière et si son vecteur tangent fait un angle constant avec une direction fixe. Une telle courbe s'appelle aussi une *hélice généralisée*.

Théorème 2.24. *Une courbe $\gamma : I \rightarrow \mathbb{R}^3$ régulière au sens de Frenet est de pente constante si et seulement si le rapport*

$$\frac{\tau_\gamma(u)}{\kappa_\gamma(u)}$$

est constant.

Preuve. On peut supposer sans perdre de généralité que γ est paramétrée naturellement. Supposons qu'il existe un vecteur constant non nul $\mathbf{A} \in \mathbb{R}^3$ tel que le produit scalaire $a = \langle \mathbf{T}_\gamma(s), \mathbf{A} \rangle$ est constant. En dérivant cette relation et en utilisant la première équation de Serret-Frenet, on trouve que

$$0 = \langle \dot{\mathbf{T}}_\gamma(s), \mathbf{A} \rangle = \kappa_\gamma(s) \langle \mathbf{N}_\gamma(s), \mathbf{A} \rangle.$$

Nous avons supposé que la courbe est birégulière, donc sa courbure est non nulle et on a donc $\langle \mathbf{N}_\gamma(s), \mathbf{A} \rangle = 0$ pour tout s .

Ceci implique que $b = \langle \mathbf{B}_\gamma(s), \mathbf{A} \rangle$ est également constant, car la troisième équation de Serret-Frenet nous dit que

$$\frac{d}{ds} \langle \mathbf{B}_\gamma(s), \mathbf{A} \rangle = \langle \dot{\mathbf{B}}_\gamma(s), \mathbf{A} \rangle = -\tau_\gamma(s) \langle \mathbf{N}_\gamma(s), \mathbf{A} \rangle = 0.$$

La seconde équation de Serret-Frenet nous dit maintenant que

$$0 = \frac{d}{ds} \langle \mathbf{N}_\gamma(s), \mathbf{A} \rangle = \langle \dot{\mathbf{N}}_\gamma(s), \mathbf{A} \rangle = -\kappa_\gamma(s) \langle \mathbf{T}_\gamma(s), \mathbf{A} \rangle + \tau_\gamma(s) \langle \mathbf{B}_\gamma(s), \mathbf{A} \rangle,$$

et donc

$$\frac{\tau_\gamma(s)}{\kappa_\gamma(s)} = \frac{\langle \mathbf{T}_\gamma(s), \mathbf{A} \rangle}{\langle \mathbf{B}_\gamma(s), \mathbf{A} \rangle} = \frac{a}{b}$$

est constante.

Supposons inversement que $\lambda := \frac{\tau_\gamma(s)}{\kappa_\gamma(s)}$ est constant et considérons le champ de vecteurs

$$\mathbf{A}(s) := \lambda \mathbf{T}_\gamma(s) + \mathbf{B}_\gamma(s).$$

Il est clair que l'angle entre $\mathbf{T}_\gamma$ et $\mathbf{A}$ est constant car $\langle \mathbf{A}, \mathbf{T}_\gamma \rangle = \lambda$. Vérifions que ce vecteur est constant :

$$\frac{d}{ds} \mathbf{A} = \lambda \dot{\mathbf{T}}_\gamma + \dot{\mathbf{B}}_\gamma = \lambda \kappa_\gamma(s) \mathbf{N}_\gamma - \tau_\gamma(s) \mathbf{N}_\gamma = 0.$$

La preuve de la proposition est complète. □

Remarque. La preuve montre que pour une courbe de pente constante, l'angle θ entre le vecteur tangent $\mathbf{T}_\gamma$ et la direction fixe $\mathbf{A}$ est donné par

$$\cos(\theta) = \frac{\langle \mathbf{T}, \mathbf{A} \rangle}{\|\mathbf{T}(u)\| \|\mathbf{A}(u)\|} = \frac{\lambda}{\sqrt{1 + \lambda^2}} = \frac{\tau_\gamma}{\sqrt{\kappa_\gamma^2 + \tau_\gamma^2}}.$$

On remarque aussi que le vecteur $\mathbf{A}$ appartient au plan rectifiant de γ .

2.10.3 Le théorème fondamental de la théorie des courbes de $\mathbb{R}^3$.

Le *théorème fondamental* de la théorie des courbes de $\mathbb{R}^3$ dit que l'on peut prescrire arbitrairement la courbure et la torsion d'une courbe birégulière de $\mathbb{R}^3$. Cette courbe est unique à un déplacement près.

Théorème 2.25. Soient $\kappa, \tau : I \rightarrow \mathbb{R}$ sont deux fonctions continues et si $\kappa(s) > 0$ pour tout $s \in I$, alors il existe une courbe $\gamma : I \rightarrow \mathbb{R}^3$, régulière au sens de Frenet, paramétrée naturellement et telle que

$$\kappa_\gamma(s) = \kappa(s) \quad \text{et} \quad \tau_\gamma(s) = \tau(s)$$

pour tout s . Cette courbe est unique à un déplacement près.

Par exemple toute courbe de $\mathbb{R}^3$ ayant courbure constante $\kappa > 0$ et torsion constante $\tau \neq 0$ est isométrique à une hélice circulaire droite.

Démonstration. Nous prouvons d'abord l'unicité. Supposons que $\gamma_1, \gamma_2 : I \rightarrow \mathbb{R}^3$ sont deux courbes régulières au sens de Frenet, paramétrées naturellement et dont la courbure et la torsion valent respectivement $\kappa(s)$ et $\tau(s)$. Notons $\{\mathbf{T}_1(s), \mathbf{N}_1(s), \mathbf{B}_1(s)\}$ et $\{\mathbf{T}_2(s), \mathbf{N}_2(s), \mathbf{B}_2(s)\}$ leur repère de Frenet respectifs.

Sans perdre de généralité, on peut supposer que l'intervalle I contient 0. Quitte à composer l'une ou l'autre (ou les deux) courbes par un déplacement, on peut supposer que $\gamma_1(0) = \gamma_2(0) = 0$ et qu'en $s = 0$ les deux repères de Frenet coïncident avec la base canonique $\{e_1, e_2, e_3\}$ de $\mathbb{R}^3$. Notons alors $\mathbf{F}_i(s) \in SO(3)$ la matrice orthogonale dont les colonnes sont les composantes des vecteurs $\mathbf{T}_i(s), \mathbf{N}_i(s), \mathbf{B}_i(s)$ pour $i = 1, 2$. Les équations de Serret-Frenet s'écrivent alors

$$\frac{d}{ds} \mathbf{F}_i(s) = \mathbf{F}_i(s) \Omega(s), \quad \text{où} \quad \Omega(s) = \begin{pmatrix} 0 & -\kappa(s) & 0 \\ \kappa(s) & 0 & -\tau(s) \\ 0 & \tau(s) & 0 \end{pmatrix}.$$

Nous avons alors

$$\begin{aligned}
 \frac{d}{ds}(\mathbf{F}_1 \mathbf{F}_2^{-1}) &= \frac{d}{ds}(\mathbf{F}_1 \mathbf{F}_2^\top) = \dot{\mathbf{F}}_1 \mathbf{F}_2^\top + \mathbf{F}_1 \dot{\mathbf{F}}_2^\top \\
 &= (\mathbf{F}_1 \Omega) \mathbf{F}_2^\top + \mathbf{F}_1 (\mathbf{F}_2 \Omega)^\top \\
 &= \mathbf{F}_1 \Omega \mathbf{F}_2^\top + \mathbf{F}_1 \Omega^\top \mathbf{F}_2^\top \\
 &= 0,
 \end{aligned}$$

car la matrice Ω est antisymétrique, i.e. $\Omega^\top = -\Omega$. Par conséquent la matrice $\mathbf{F}_1 \mathbf{F}_2^{-1}$ est constante. Mais on a supposé que $\mathbf{F}_1(0) = \mathbf{F}_2(0) = \mathbf{I}_3$ (la matrice identité). Donc $\mathbf{F}_1(s) \mathbf{F}_2(s)^{-1} = \mathbf{I}_3$ pour tout s c'est-à-dire $\mathbf{F}_1(s) = \mathbf{F}_2(s)$. En particulier $\mathbf{T}_1(s) = \mathbf{T}_2(s)$ pour tout s et donc

$$\gamma_1(s) = \int_0^s \mathbf{T}_1(u) du = \int_0^s \mathbf{T}_2(u) du = \gamma_2(s).$$

Prouvons maintenant l'existence. Pour cela on se donne deux fonctions continues $\kappa, \tau : I \rightarrow \mathbb{R}$ et on considère le problème de Cauchy linéaire

$$\frac{d}{ds} \mathbf{F}(s) = \mathbf{F}(s) \Omega(s), \quad \mathbf{F}(0) = \mathbf{I}_3, \quad (2.9)$$

où $\Omega(s)$ est la matrice définie plus haut. Le théorème de Cauchy–Lipschitz global du cours d'analyse II nous dit qu'il existe une solution globale $\mathbf{F} : I \rightarrow M_3(\mathbb{R})$ de classe C^1 de ce problème. Nous affirmons que $\mathbf{F}(s) \in SO(3)$ pour tout s . En effet, on a

$$\frac{d}{ds} \mathbf{F} \mathbf{F}^\top = \dot{\mathbf{F}} \mathbf{F}^\top + \mathbf{F} \dot{\mathbf{F}}^\top = \mathbf{F} \Omega \mathbf{F}^\top + \mathbf{F} \Omega^\top \mathbf{F}^\top = \mathbf{F} (\Omega + \Omega^\top) \mathbf{F}^\top = 0$$

par antisymétrie de Ω . Or $\mathbf{F}(0) \mathbf{F}^\top(0) = \mathbf{I}_3$ (à cause de la condition initiale dans (2.9)), donc $\mathbf{F}(s) \mathbf{F}^\top(s) = \mathbf{I}_3$ pour tout $s \in I$, ce qui signifie que $\mathbf{F}(s) \in SO(3)$.

Notons respectivement $\mathbf{T}(s), \mathbf{N}(s), \mathbf{B}(s)$ les trois colonnes de la matrice $\mathbf{F}(s)$ et définissons $\gamma : I \rightarrow \mathbb{R}^3$ par

$$\gamma(s) = \int_0^s \mathbf{T}(u) du.$$

Alors γ est clairement une courbe de classe C^2 car $s \rightarrow \mathbf{T}(s)$ est de classe C^1 . De plus cette courbe est paramétrée naturellement puisque $\dot{\gamma}(s) = \mathbf{T}(s)$ est un vecteur unitaire. L'équation différentielle (2.9) est équivalente aux équations de Serret-Frenet. Cela implique que $s \rightarrow \mathbf{N}(s)$ est aussi de classe C^1 et que la courbure et la torsion de γ sont données par les fonctions κ et τ , ce qui complète notre démonstration. □

2.11 Courbes dans un plan orienté

Le repère de Frenet et la courbure d'une courbe dans un plan orienté est défini en tenant compte de l'orientation du plan :

Définition. (a) Le *repère de Frenet orienté* d'une courbe régulière $\gamma : I \rightarrow \mathbb{R}^2$ de classe C^2 dans le plan orienté est le repère mobile d'origine $\gamma(u)$ qui est formé par les deux vecteurs

$$\mathbf{T}_\gamma(u) = \frac{\dot{\gamma}(u)}{V_\gamma(u)}, \quad \mathbf{N}_\gamma^{\text{or}}(u) = \mathbf{J}(\mathbf{T}_\gamma(u)),$$

où $\mathbf{J}$ est l'opérateur de rotation défini au paragraphe (1.7).

(b) La *courbure orientée* de γ en u est définie par

$$\kappa_\gamma^{\text{or}}(u) = \frac{1}{V_\gamma(u)} \langle \dot{\mathbf{T}}_\gamma(u), \mathbf{N}_\gamma^{\text{or}}(u) \rangle.$$

Remarques.

- (i) Le repère $\{\mathbf{T}_\gamma(u), \mathbf{N}_\gamma^{\text{or}}(u)\}$ est un repère orthonormé direct.
- (ii) La courbure non orientée de γ est égale à la valeur absolue de $\kappa_\gamma^{\text{or}}(u)$.
- (iii) La courbure orientée peut aussi s'écrire

$$\kappa_\gamma^{\text{or}}(u) = \frac{\mathbf{T}_\gamma(u) \wedge \dot{\mathbf{T}}_\gamma(u)}{V_\gamma(u)}.$$

(iv) Si la courbe γ est birégulière, on a

$$\kappa_\gamma^{\text{or}}(u) \mathbf{N}_\gamma^{\text{or}}(u) = \kappa_\gamma(u) \mathbf{N}_\gamma(u) = \mathbf{K}_\gamma(u) \quad (= \text{le vecteur de courbure}).$$

Cette égalité vient du fait que si on change l'orientation du plan, alors $\kappa_\gamma^{\text{or}}(u)$ et $\mathbf{N}_\gamma^{\text{or}}(u)$ changent tous les deux de signe.

Dans la suite de ce paragraphe, nous n'utiliserons que le vecteur normal orienté, nous noterons donc $\mathbf{N}_\gamma(u)$ au lieu de $\mathbf{N}_\gamma^{\text{or}}(u)$, nous noterons aussi $k_\gamma(u)$ pour la courbure orientée.

Proposition 2.26. Avec ces notations, les formules de Serret-Frenet pour une courbe plane de classe C^2 s'écrivent

$$\begin{aligned} \frac{1}{V_\gamma} \frac{d}{du} \mathbf{T}_\gamma(u) &= k_\gamma(u) \mathbf{N}_\gamma(u) \\ \frac{1}{V_\gamma} \frac{d}{du} \mathbf{N}_\gamma(u) &= -k_\gamma(u) \mathbf{T}_\gamma(u). \end{aligned}$$

Preuve. On a d'une part

$$\dot{\mathbf{T}} = \langle \dot{\mathbf{T}}, \mathbf{T} \rangle \mathbf{T} + \langle \dot{\mathbf{T}}, \mathbf{N} \rangle \mathbf{N} = \langle \dot{\mathbf{T}}, \mathbf{N} \rangle \mathbf{N} = V k \mathbf{N}.$$

par définition de la courbure orientée k (et en utilisant $\langle \dot{\mathbf{T}}, \mathbf{T} \rangle = 0$). D'autre part

$$\dot{\mathbf{N}} = \langle \dot{\mathbf{N}}, \mathbf{T} \rangle \mathbf{T} + \langle \dot{\mathbf{N}}, \mathbf{N} \rangle \mathbf{N} = \langle \dot{\mathbf{N}}, \mathbf{T} \rangle \mathbf{T} = -\langle \dot{\mathbf{T}}, \mathbf{N} \rangle \mathbf{T} = -V k \mathbf{T}.$$

□

Proposition 2.27. La courbure orientée d'une courbe plane $\gamma : I \rightarrow \mathbb{R}^2$ de classe C^2 est donnée par

$$k_\gamma(u) = \frac{\dot{\gamma}(u) \wedge \ddot{\gamma}(u)}{V_\gamma^3(u)}.$$

Preuve. On a

$$\dot{\gamma} \wedge \ddot{\gamma} = (V\mathbf{T}) \wedge (\dot{V}\mathbf{T} + V^2k\mathbf{N}) = V^3k,$$

car $\mathbf{T} \wedge \mathbf{T} = 0$ et $\mathbf{T} \wedge \mathbf{N} = 1$.

□

La courbure orientée de $\gamma(t) = (x(t), y(t))$ est donc donnée par

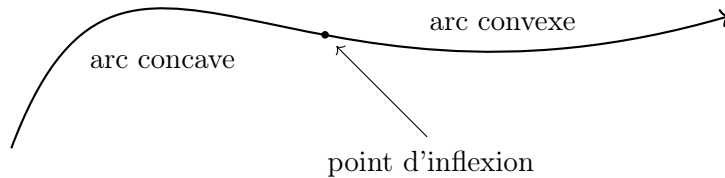
$$k(t) = \frac{\dot{x}\ddot{y} - \ddot{x}\dot{y}}{(\dot{x}^2 + \dot{y}^2)^{3/2}}.$$

En particulier, si γ est le graphe de la fonction $f : [a, b] \rightarrow \mathbb{R}$, i.e. $\gamma(x) = (x, f(x))$, alors on a

$$k(x) = \frac{f''(x)}{(1 + f'(x)^2)^{3/2}}.$$

Définitions. On dit qu'un arc $\gamma(u)$ ($a < u < b$) est *convexe* si la courbure orientée k_γ est positive sur cet arc. L'arc est *concave* si la courbure orientée est négative. Un *point d'inflexion* est un point séparant un arc convexe d'un arc concave (en particulier la courbure est nulle en un point d'inflexion).

On dit qu'un arc est une *spirale* si la courbure est strictement monotone sur cet arc. Un point de la courbe est un *sommet* si c'est un maximum local ou un minimum local de la courbure.



La fonction angulaire

La fonction angulaire mesure l'inclinaison en chaque point d'une courbe par rapport à la direction horizontale.

Définition. Soit $\gamma : [a, b] \rightarrow \mathbb{R}^2$ une courbe régulière de classe C^1 . La *fonction angulaire* de la courbe γ avec point initial $p = \gamma(u_0)$ est la fonction $\varphi : [a, b] \rightarrow \mathbb{R}$ telle que

- (a) $\varphi(u_0)$ est l'angle orienté entre $\dot{\gamma}(u)$ et le vecteur $\mathbf{e}_1 = (1, 0)$.
- (b) φ est continue.
- (c) L'angle orienté entre $\dot{\gamma}(u)$ et $\mathbf{e}_1 = (1, 0)$ est égal à $\varphi(u)$ modulo 2π pour tout $u \in [a, b]$

REMARQUE. Dans le concept de fonction angulaire d'une courbe plane, on n'identifie pas $\varphi(u)$ à $\varphi(u) + 2\pi$. Au contraire, le paramètre angulaire mesure le nombre de tours effectués (entre u_0 et u) par le vecteur tangent. Ce nombre peut être supérieur à 2π .

Le nombre $\varphi(b) - \varphi(a)$ est la *variation angulaire totale* de la courbe. Si $\gamma : [a, b] \rightarrow \mathbb{R}^2$ est une courbe périodique (i.e. une courbe fermée régulière), alors le nombre entier

$$\frac{\varphi(b) - \varphi(a)}{2\pi} \in \mathbb{Z}$$

s'appelle le *nombre de rotations* de γ .

Lemme 2.28. *Le repère de Frenet orienté d'une courbe régulière $\gamma : [a, b] \rightarrow \mathbb{R}^2$ de classe C^1 peut s'écrire*

$$\mathbf{T}_\gamma(u) = (\cos(\varphi(u)), \sin(\varphi(u))) \quad \text{et} \quad \mathbf{N}_\gamma(u) = (-\sin(\varphi(u)), \cos(\varphi(u))),$$

où $\varphi : [a, b] \rightarrow \mathbb{R}$ est la *fonction angulaire* de γ .

Preuve. La formule pour $\mathbf{T}$ est évidente, puisque φ (modulo 2π) mesure l'angle du vecteur tangent $\mathbf{T}$ avec $\mathbf{e}_1$. La formule pour $\mathbf{N}$ se déduit alors de la définition $\mathbf{N} = \mathbf{J}(\mathbf{T})$. □

Théorème 2.29. *La courbure orientée d'une courbe $\gamma : I \rightarrow \mathbb{R}^2$ de classe C^2 vérifie*

$$k_\gamma(u) = \frac{1}{V_\gamma(u)} \frac{d\varphi}{du}.$$

Preuve. Par le lemme précédent, on a $\dot{\mathbf{T}} = (-\sin(\varphi(u)), \cos(\varphi(u)))\dot{\varphi}(u) = \mathbf{N}\dot{\varphi}$, donc $k_\gamma = \frac{1}{V} \langle \dot{\mathbf{T}}, \mathbf{N} \rangle = \frac{1}{V} \dot{\varphi}$. □

Lorsque la courbe est paramétrée naturellement, on a $k_\gamma(s) = \frac{d\varphi}{ds}$. On écrit souvent cette relation sous la forme différentielle :

$$d\varphi = k ds.$$

Le diagramme de courbure

Soit $\gamma : I \rightarrow \mathbb{R}^2$ une courbe régulière de classe C^2 . Choisissons un point initial sur γ et un sens de parcours. Le *diagramme de courbure* de γ est la courbe dans un plan de coordonnées s, k donnée par

$$u \mapsto (s(u), k(u)),$$

où $s(u)$ est l'abscisse curviligne de γ correspondant aux choix du point initial et du sens de parcours, et k est la courbure orientée.

Le diagramme de courbure est toujours un graphe (c'est le graphe de la fonction courbure $k = k(s)$ exprimée à partir de l'abscisse curviligne). Les éléments de la courbe γ que l'on peut facilement mettre en correspondance avec le diagramme de courbure sont :

- sa longueur $\ell(\gamma)$;
- le signe de la courbure ;
- les points d'inflexions de γ (ce sont les points où $k(s)$ change de signe) ;
- les sommets de γ (i.e. les extremums locaux de la courbure orientée).

D'autre part, l'aire $\int_0^\ell k(s) ds = \int_0^\ell d\varphi$ limitée par le diagramme de courbure correspond à la variation angulaire totale de la courbe. Hormis la position de la courbe dans le plan, le diagramme de courbure contient toutes les informations géométriques sur une courbe de $\mathbb{R}^2$.

Théorème 2.30 (Théorème fondamental de la théorie des courbes planes). *Toute fonction continue $k : [0, \ell] \rightarrow \mathbb{R}$ est la courbure orientée d'une courbe plane de classe C^2 paramétrée naturellement. Cette courbe est unique à un déplacement près.*

Ce théorème est la version bidimensionnelle du Théorème 2.25, mais la preuve est plus élémentaire.

Preuve. Montrons d'abord l'unicité. Supposons que $\gamma : [0, \ell] \rightarrow \mathbb{R}^2$ est une courbe de classe C^2 paramétrée naturellement dont la courbure orientée est $k(s)$. Le vecteur tangent est donné par

$$\mathbf{T}(s) = \dot{\gamma}(s) = (\dot{x}(s), \dot{y}(s)) = (\cos(\varphi(s)), \sin(\varphi(s)))$$

où $\varphi : [0, \ell] \rightarrow \mathbb{R}$ est la fonction angulaire. Les trois fonctions $(x(s), y(s), \varphi(s))$ forment alors une solution du système d'équations différentielles

$$\begin{cases} \frac{dx}{ds} = \cos(\varphi) \\ \frac{dy}{ds} = \sin(\varphi) \\ \frac{d\varphi}{ds} = k(s) \end{cases} \quad (2.10)$$

La courbe γ est donc déterminée à partir de la fonction $k(s)$ en résolvant ces équations.

Pour résoudre ce système, on calcule φ par intégration : $\varphi(s) = \varphi_0 + \int_0^s k(\sigma) d\sigma$. Puis on trouve $x(s)$ et $y(s)$ par une nouvelle intégration :

$$x(s) = x_0 + \int_0^s \cos(\varphi(\sigma)) d\sigma \quad , \quad y(s) = y_0 + \int_0^s \sin(\varphi(\sigma)) d\sigma$$

les constantes x_0, y_0 , et φ_0 sont des constantes d'intégration et peuvent être choisies arbitrairement (ce sont les *conditions initiales* du système d'équations différentielles).

En changeant les valeurs de x_0 et y_0 , on modifie la courbe par une translation ; si l'on change φ_0 , alors la courbe γ subit une rotation. L'argument montre à la fois l'existence et l'unicité de la courbe γ à un déplacement près.

□

REMARQUE. La relation $k = k(s)$ entre l'abscisse curviligne et la courbure orientée s'appelle l'*équation intrinsèque* de la courbe. Elle contient la même information que le diagramme de courbure.

Exemple. Considérons la courbe dont le diagramme de courbure est une droite oblique (i.e. l'équation intrinsèque est linéaire : $k(s) = ms + n$ avec $m \neq 0$). Alors la fonction angulaire est donnée par

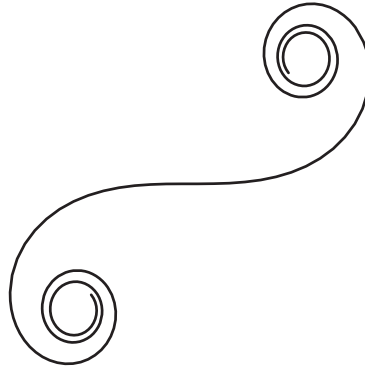
$$\varphi(s) = \int k(s) ds = \frac{m}{2}s^2 + ns + c,$$

et la courbe est donc donnée par

$$x(s) = \int \cos\left(\frac{m}{2}s^2 + ns + c\right) ds \quad , \quad y(s) = \int \sin\left(\frac{m}{2}s^2 + ns + c\right) ds.$$

Ces intégrales s'appellent les *fonctions de Fresnel*. Elle ne peuvent pas être exprimées à partir des fonctions élémentaires.

Cette courbe s'appelle une *chlotoïde* ou *spirale de Cornu*, elle permet par exemple de passer d'une droite à un cercle sans discontinuité de la courbure. Pour cette raison, elle est utilisée dans la conception des tracés ferroviaires ou autoroutiers.



Chloïde.

2.12 Le théorème des quatre sommets

Définition. On dit que $\gamma : [a, b] \rightarrow \mathbb{R}^n$ est une *courbe fermée de classe C^m* si la fonction γ peut s'étendre à un intervalle ouvert $[a - \varepsilon, b + \varepsilon]$ et si on a $\gamma(a) = \gamma(b)$ et

$$\frac{d^k \gamma}{du^k}(a) = \frac{d^k \gamma}{du^k}(b),$$

pour $1 \leq k \leq m$. On dit aussi que γ est une courbe *périodique* de classe C^m car on peut l'étendre en une fonction périodique $\gamma : \mathbb{R} \rightarrow \mathbb{R}^n$ de période $(b - a)$.

En particulier, si $\gamma : [a, b] \rightarrow \mathbb{R}^n$ est une courbe fermée de classe C^1 alors $\gamma(a) = \gamma(b)$ et $\dot{\gamma}(a) = \dot{\gamma}(b)$. Le vecteur vitesse est donc le même en $t = a$ et en $t = b$. Et si la courbe est fermée de classe C^2 , alors la courbure en $t = a$ est égale à la courbure en $t = b$.

Théorème 2.31. *Toute courbe fermée de classe C^3 dans un plan orienté, qui n'est pas un cercle, possède au moins quatre sommets.*

On rappelle qu'un sommet d'une courbe de classe C^2 est un maximum local ou un minimum local de la courbure orientée. A titre d'exemple, une ellipse possède deux minimums et deux maximums de courbure. La preuve utilisera le lemme suivant :

Lemme 2.32. *Soit $\gamma : [0, \ell] \rightarrow \mathbb{R}^2$ une courbe fermée de classe C^3 paramétrée naturellement. Alors*

$$\int_0^\ell x(s) \dot{k}(s) ds = \int_0^\ell y(s) \dot{k}(s) ds = 0.$$

Preuve du Lemme. Examinons la seconde intégrale, on a

$$\int_0^\ell y(s) \dot{k}(s) ds = - \int_0^\ell \dot{y}(s) k(s) ds = \int_0^\ell \ddot{x}(s) ds = 0.$$

En effet la première égalité est une intégration par parties, la seconde égalité vient de la relation $\ddot{x} = -k\dot{y}$ qui se déduit des équations de Serret-Frenet et la dernière égalité est évidente. $\square$

Preuve du Théorème. La preuve dans le cas général est assez élaborée, nous ne la donnerons que dans le cas où la courbe est le bord d'un domaine convexe du plan. Par hypothèse, γ est une courbe fermée de classe C^3 , par conséquent la dérivée de la courbure vérifie $\dot{k}(0) = \dot{k}(\ell)$ et

la fonction $k(s)$ doit donc avoir au moins un maximum local et un minimum local. Nous allons d'abord prouver par l'absurde que $k(s)$ doit avoir au moins un troisième extremum local.

Supposons donc par l'absurde que $k(s)$ a exactement deux extremums locaux, et faisons également les hypothèses suivantes sans perte de généralité :

1. On suppose que $\gamma : [0, \ell] \rightarrow \mathbb{R}^2$ est paramétrée naturellement.
2. Le minimum local de $k(s)$ est en $s = 0$ et le maximum local est en $s_0 \in (0, \ell)$.
3. $\gamma(0) = (0, 0)$ et $\gamma(s_0) = (x_0, 0)$.

Ces hypothèses entraînent que $k(s)$ est strictement croissante sur l'intervalle $(0, s_0)$ et strictement décroissante sur (s_0, ℓ) . Donc $\dot{k}(s) > 0$ sur le premier intervalle et $\dot{k}(s) < 0$ sur le deuxième intervalle.

Puisque γ borde un domaine convexe, les deux arcs $\gamma|_{[0, s_0]}$ et $\gamma|_{[s_0, \ell]}$ sont situés l'un dans le demi-plan $\{y \geq 0\}$ et l'autre dans le demi-plan $\{y \leq 0\}$. Supposons par exemple que $y(s) > 0$ sur l'intervalle $(0, s_0)$ et $y(s) < 0$ sur l'intervalle (s_0, ℓ) , alors nous avons $y(s)\dot{k}(s) > 0$ pour tous $s \notin \{0, s_0\}$. Mais ceci entre en contradiction avec le lemme précédent car ce lemme implique que

$$\int_0^{s_0} y(s)\dot{k}(s)ds = - \int_{s_0}^{\ell} y(s)\dot{k}(s)ds.$$

L'argument est le même (avec le signe opposé) si $y(s) < 0$ sur l'intervalle $(0, s_0)$ et $y(s) > 0$ sur (s_0, ℓ) .

Nous avons montré que $\dot{k}(s)$ doit avoir au moins trois changements de signe. Mais comme on a $\dot{k}(0) = \dot{k}(\ell)$, cette fonction ne peut pas avoir un nombre impair de changements de signe. Il y a donc au moins quatre changements de signe.

□

Note historique. Le théorème des quatre sommets a été démontré par le mathématicien indien Syamadas Mukhopadhyaya en 1909 pour les courbes fermées convexes, puis par le mathématicien allemand A. Kneser dans le cas général en 1912. Il existe une réciproque, démontrée d'abord en 1971 par Herman Gluck dans le cas des courbes convexes, puis en 2005 par Björn Dahlberg dans le cas général.

Chapitre 3

Calcul différentiel et sous-variétés différentiables de $\mathbb{R}^n$

Les sous-variétés différentiables sont des parties de $\mathbb{R}^n$ qui généralisent les courbes et surfaces en toutes dimensions et codimensions. On les suppose assez régulières pour qu'on puisse appliquer les concepts et outils du calcul différentiel.

3.1 Rappels de calcul différentiel

3.1.1 Dérivées directionnelles et dérivées partielles

Soit U un domaine de $\mathbb{R}^m$ et $f : U \rightarrow \mathbb{R}^n$ une application. Pour un point p de U et un vecteur $v \in \mathbb{R}^m$ on définit la *dérivée directionnelle de f en direction de v au point p* par

$$D_v f(p) = \left. \frac{d}{dt} \right|_{t=0} f(p + tv) = \lim_{t \rightarrow 0} \frac{f(p + tv) - f(p)}{t} \in \mathbb{R}^n, \quad (3.1)$$

si cette limite existe.

Si $\{e_1, e_2, \dots, e_m\}$ est la base canonique de $\mathbb{R}^m$ et $x_1, x_2, \dots, x_m$ sont les coordonnées associées (i.e. un vecteur $x \in \mathbb{R}^m$ s'écrit $x = \sum_{i=1}^m x_i e_i$), alors la dérivée directionnelle de f en direction du vecteur e_i s'appelle la *dérivée partielle de f au point p en direction de la $i^{\text{ème}}$ coordonnée* (ou en direction du vecteur e_i), et on note

$$\frac{\partial f}{\partial x_i}(p) = D_{e_i} f(p) = \lim_{t \rightarrow 0} \frac{f(p + te_i) - f(p)}{t} \in \mathbb{R}^n. \quad (3.2)$$

Remarque. Il est important de noter que l'existence des dérivées partielles d'une fonction en un point donné ne garantit pas l'existence des dérivées directionnelles dans toutes les directions. Par exemple la fonction $f : \mathbb{R}^2 \rightarrow \mathbb{R}$ définie par

$$f(x, y) = \begin{cases} \frac{xy}{(x^2 + y^2)^{3/4}} & \text{si } (x, y) \neq (0, 0), \\ 0 & \text{si } (x, y) = (0, 0) \end{cases}$$

est continue et possède les dérivées partielles $\frac{\partial f}{\partial x}(0, 0) = \frac{\partial f}{\partial y}(0, 0) = 0$, mais la dérivée directionnelle en $(0, 0)$ en direction de $v = (v_1, v_2)$ n'existe pas si v_1 et v_2 sont non nuls car

$$\lim_{t \rightarrow 0^+} \frac{f(tv_1, tv_2) - f(0, 0)}{t} = \lim_{t \rightarrow 0^+} \frac{t^2 v_1 v_2}{t((tv_1)^2 + (tv_2)^2)^{3/4}} = \frac{v_1 v_2}{(v_1^2 + v_2^2)^{3/4}} \lim_{t \rightarrow 0^+} \frac{1}{\sqrt{t}} = \pm\infty.$$

3.1.2 Applications de classe C^k sur un ouvert de $\mathbb{R}^m$

On dit que l'application $f : U \rightarrow \mathbb{R}^n$ est de classe C^1 si elle est continue et si toutes les dérivées partielles du premier ordre

$$\frac{\partial f}{\partial x_i} : U \rightarrow \mathbb{R}^n$$

existent en tout point de U et sont continues. La fonction est de classe C^k (k un entier ≥ 2) si les $m + 1$ fonctions $f, \frac{\partial f}{\partial x_1}, \dots, \frac{\partial f}{\partial x_m} : U \rightarrow \mathbb{R}^n$ sont de classe C^{k-1} .

On note $C^0(U, \mathbb{R}^n)$ l'ensemble des applications continues sur U et $C^k(U, \mathbb{R}^n)$ l'ensemble des applications de classe C^k . Une application est de classe C^∞ si elle est de classe C^k pour tout k et on note $C^\infty(U, \mathbb{R}^n) = \bigcap_{k \geq 0} C^k(U, \mathbb{R}^n)$. On dit parfois que f est *lisse* si $f \in C^\infty(U, \mathbb{R}^n)$.

Lorsque $n = 1$, i.e. lorsque f est à valeurs dans $\mathbb{R}$ on note simplement $C^k(U) = C^k(U, \mathbb{R})$, on appelle les éléments de $C^k(U)$ des *fonctions* (ainsi les fonctions sont les applications à valeurs dans $\mathbb{R}$).

La matrice à n lignes et m colonnes contenant les dérivées partielles

$$Df = \begin{pmatrix} \frac{\partial f_1}{\partial x_1} & \cdots & \frac{\partial f_1}{\partial x_m} \\ \vdots & \ddots & \vdots \\ \frac{\partial f_n}{\partial x_1} & \cdots & \frac{\partial f_n}{\partial x_m} \end{pmatrix}$$

s'appelle la *matrice Jacobienne*¹ de f . Noter que c'est une fonction du point $p \in U$. Lorsque $n = m$, le déterminant de cette matrice est alors bien défini, on l'appelle le *Jacobien* de f et on note²

$$J_f(p) = \det \left(\frac{\partial f_i}{\partial x_j}(p) \right).$$

Définitions. 1. Un *difféomorphisme* de classe C^k entre deux ouverts U et V de même dimension est une application bijective $f : U \rightarrow V$ telle que f et f^{-1} sont de classe C^k . Lorsque $k = 0$, on dit que f est un *homéomorphisme*.

2. Une application $f : U \rightarrow \mathbb{R}^n$ est un *difféomorphisme local* de classe C^k en $p \in U$ s'il existe un voisinage ouvert $U' \subset U$ de p tel que $V' = f(U')$ est ouvert et la restriction $f|_{U'}$ est un C^k -difféomorphisme de U' sur V' .

3. Finalement on dit que $f : U \rightarrow \mathbb{R}^n$ est un *difféomorphisme local* si c'est un difféomorphisme local en chaque point de U .

Observer qu'un difféomorphisme local n'est pas forcément une application injective (ni surjective d'ailleurs).

Remarque. Un homéomorphisme de classe C^k n'est pas toujours un difféomorphisme. Par exemple la fonction $f(x) = x^3$ décrit un homéomorphisme C^∞ de $\mathbb{R}$ vers $\mathbb{R}$ mais l'inverse $f^{-1}(y) = \sqrt[3]{y}$ n'est pas dérivable en $y = 0$.

Il y a deux façons de concevoir un difféomorphisme $f : U \rightarrow V$. Dans le premier point de vue, on considère que f déplace les points de U (éventuellement en déformant l'ensemble U). Ainsi, si p est un point de U , on considère que $q = f(p)$ est un autre point, qui appartient à V .

1. Du nom de Carl Gustav Jacob Jacobi, mathématicien allemand (1804–1851).

2. Une autre notation, un peu désuète mais assez explicite, est

$$\frac{\partial(f_1, \dots, f_n)}{\partial(x_1, \dots, x_m)} = \det \left(\frac{\partial f_i}{\partial x_j} \right).$$

Dans le second point de vue, les points ne “bougent” pas, mais on considère que $(y_1, y_2, \dots, y_n) = f(x_1, x_2, \dots, x_n)$ représente un nouveau système de coordonnées sur U . Ceci nous mène à la définition suivante :

Définition. Un *système de coordonnées curviligne* de classe C^k sur l’ouvert $U \subset \mathbb{R}^n$ est la donnée de n fonctions $y_1, y_2, \dots, y_n : U \rightarrow \mathbb{R}$ telles que

$$\phi : (x_1, x_2, \dots, x_n) \mapsto (y_1, y_2, \dots, y_n)$$

décrit un difféomorphisme de classe C^k de U vers un ouvert $V = \phi(U) \subset \mathbb{R}^n$.

3.1.3 Applications Différentiables au sens de Fréchet

Définition. L’application $f : U \subset \mathbb{R}^m \rightarrow \mathbb{R}^n$ est *différentiable au sens de Frechet* en $p \in U$ s’il existe une application linéaire $\ell : \mathbb{R}^m \rightarrow \mathbb{R}^n$ telle que

$$f(x) - f(p) - \ell(x - p) = o(\|x - p\|)$$

Intuitivement, une application f est donc différentiable (au sens de Frechet) en p si $f(x) - f(p)$ est tangente à une application linéaire :

Lemme 3.1. *Si elle existe, l’application linéaire de la définition précédente est unique.*

On appelle alors cette application la *différentielle* de f en p et on note³

$$df_p := \ell.$$

Preuve du lemme. Supposons que ℓ_1 et ℓ_2 soient deux applications linéaires telles que

$$f(x) - f(p) - \ell_1(x - p) = f(x) - f(p) - \ell_2(x - p) = o(\|x - p\|).$$

Soit v un vecteur quelconque de $\mathbb{R}^n$ et $t \in \mathbb{R}$ un réel assez petit pour que $x + tv \in U$, alors on a

$$f(p + tv) - f(p) - \ell_1(tv) = o(\|tv\|) = o(t),$$

et de même

$$f(p + tv) - f(p) - \ell_2(tv) = o(\|tv\|) = o(t).$$

Par conséquent, on a

$$t(\ell_1(v) - \ell_2(v)) = \ell_1(tv) - \ell_2(tv) = o(t),$$

ce qui signifie que

$$\lim_{t \rightarrow 0} \frac{\|\ell_1(tv) - \ell_2(tv)\|}{t} = 0 = \lim_{t \rightarrow 0} \|\ell_1(v) - \ell_2(v)\|,$$

c’est-à-dire $\ell_1(tv) = \ell_2(tv)$. □

Remarque. Il est fréquent de noter h le vecteur $h = x - p$. On pense alors à h comme un “accroissement” de p . On a alors

$$f(p + h) = f(p) + df_p(h) + o(\|h\|).$$

On remarque aussi que $df_p(h)$ peut se calculer par la formule suivante :

$$df_p(h) = \lim_{t \rightarrow 0} \left(\frac{f(p + th) - f(p)}{t} \right) = \frac{d}{dt} \Big|_{t=0} f(p + th),$$

Il s’agit donc de la dérivée directionnelle de f au point p en direction de h .

3. La notation Df_p est également souvent utilisée, mais nous préférons garder cette notation pour la matrice jacobienne.

Exemples

- (i) Si $I \subset \mathbb{R}$ est un intervalle ouvert, alors $f : I \rightarrow \mathbb{R}$ est différentiable en $p \in I$ si et seulement si f est dérivable en p et

$$df_p(h) = f'(p) \cdot h.$$

- (ii) Soit $f : \mathbb{R}^m \rightarrow \mathbb{R}^n$ une application affine, i.e. une application du type $f(x) = Ax + b$ (où A est une $n \times m$ matrice à coefficients réels). Alors f est différentiable en tout point et $df_p = A$ pour tout $p \in \mathbb{R}^m$.
- (iii) Si $\beta : E_1 \times E_2 \rightarrow E_3$ est une application bilinéaire (où E_i sont des espaces normés de dimension finie), alors

$$d\beta_{(p_1, p_2)}(h_1, h_2) = \beta(p_1, h_2) + \beta(h_1, p_2)$$

- (iv) Considérons l'application $\psi : M_n(\mathbb{R}) \rightarrow M_n(\mathbb{R})$ définie par $\psi(A) = A^2$, alors

$$d\psi_A(H) = AH + HA.$$

- (v) L'application $\phi : GL_n(\mathbb{R}) \rightarrow GL_n(\mathbb{R})$ définie par $\phi(A) = A^{-1}$ est différentiable et on a

$$d\phi_A(H) = -A^{-1}HA^{-1}.$$

- (vi) La différentielle de l'application $\det : GL_n(\mathbb{R}) \rightarrow \mathbb{R}$ est donnée par

$$d\det_A(H) = \text{Trace}(\text{Cof}(A)^\top H).$$

où $\text{Cof}(A)$ est la matrice des cofacteurs de A .

Proposition 3.2 (Différentiation en chaîne). *Soient $U \subset \mathbb{R}^m$, $V \subset \mathbb{R}^n$ deux ouverts et $f : U \rightarrow V$, $g : V \rightarrow \mathbb{R}^s$ deux applications telles que f est Fréchet différentiable en $p \in U$ et g est Fréchet différentiable en $q = f(p) \in V$, alors $g \circ f : U \rightarrow \mathbb{R}^s$ est Fréchet différentiable en p et*

$$d(g \circ f)_p = dg_q \circ df_p$$

Cette proposition est l'une des raisons qui rend la notion de différentiabilité au sens de Fréchet efficace et importante.

Preuve. Par hypothèse, on a

$$f(p+h) - f(p) = df_p(h) + o(\|h\|) \quad \text{et} \quad g(q+k) - g(q) = dg_q(k) + o(\|k\|).$$

Donc

$$\begin{aligned} g \circ f(p+h) - g \circ f(p) &= g \circ f(p+h) - g(q) \\ &= g(f(p) + df_p(h) + o(\|h\|)) - g(q) \\ &= g(q + df_p(h) + o(\|h\|)) - g(q) \\ &= dg_q(df_p(h) + o(\|h\|)) + o(df_p(h) + o(\|h\|)) \\ &= dg_q \circ df_p(h) + o(\|h\|), \end{aligned}$$

ce qui démontre que $d(g \circ f)_p = dg_q \circ df_p$. □

Proposition 3.3. *Soit $f : U \rightarrow \mathbb{R}^n$ une application différentiable en chaque point de U , où U est un ouvert convexe de $\mathbb{R}^m$. Supposons que la différentielle de f est bornée sur U , i.e. il existe $C > 0$ tel que $\|df_p\| \leq C$ pour tout $p \in U$ (ici on utilise la norme d'opérateur pour df). Alors f est C -Lipschitzienne, i.e.*

$$\|f(y) - f(x)\| \leq C\|y - x\|.$$

Preuve. Le segment de droite reliant x à y est contenu dans le domaine U puisque celui-ci est supposé convexe. On paramétrise ce segment par $\gamma(t) = x + t(y - x) \in U$. On a alors

$$f(y) - f(x) = f(\gamma(1)) - f(\gamma(0)) = \int_0^1 \frac{d}{dt} f(\gamma(t)) dt.$$

Or la règle de dérivation en chaîne nous dit que

$$\frac{d}{dt} f(\gamma(t)) = df_{\gamma(t)}(\dot{\gamma}(t)) = df_{\gamma(t)}(y - x),$$

donc

$$\begin{aligned} \|f(y) - f(x)\| &= \left\| \int_0^1 df_{\gamma(t)}(\dot{\gamma}(t)) dt \right\| \\ &\leq \int_0^1 \|df_{\gamma(t)}(y - x)\| dt \\ &\leq \int_0^1 \|df_{\gamma(t)}\| \cdot \|y - x\| dt \\ &\leq C \|y - x\|. \end{aligned}$$

□

Théorème 3.4. Si $f : U \subset \mathbb{R}^m \rightarrow \mathbb{R}^n$ est de classe C^1 , alors f est différentiable au sens de Frechet en tout point p de U . De plus la matrice de la différentielle df_p est la matrice Jacobienne de f en p :

$$df_p = \left(\frac{\partial f_i}{\partial x_j}(p) \right)$$

Démonstration. Nous donnons la preuve pour $m = 2$, le cas général est semblable. Écrivons

$$f(p_1 + h_1, p_2 + h_2) - f(p_1, p_2) = f(p_1 + h_1, p_2 + h_2) - f(p_1 + h_1, p_2) + f(p_1 + h_1, p_2) - f(p_1, p_2).$$

On a d'une part

$$f(p_1 + h_1, p_2) - f(p_1, p_2) = \frac{\partial f(p_1, p_2)}{\partial x_1} h_1 + o(h_1).$$

D'autre part, en appliquant le théorème des accroissements fini à la fonction

$$\phi(t) = f(p_1 + h_1, p_2 + th_2),$$

on sait qu'il existe $s \in [0, 1]$ tel que

$$\phi(1) - \phi(0) = \phi'(s) = h_2 \cdot \frac{\partial f(p_1 + h_1, p_2 + sh_2)}{\partial x_2},$$

c'est-à-dire

$$f(p_1 + h_1, p_2) - f(p_1, p_2) = h_2 \cdot \frac{\partial f(p_1 + h_1, p_2 + sh_2)}{\partial x_2}.$$

Par continuité de $\frac{\partial f}{\partial x_2}$, on a

$$\frac{\partial f(p_1 + h_1, p_2 + sh_2)}{\partial x_2} = \frac{\partial f(p_1, p_2)}{\partial x_2} + o(h_1, h_2)$$

En regroupant toute ces identités, on obtient

$$f(p_1 + h_1, p_2 + h_2) - f(p_1, p_2) = \frac{\partial f(p_1, p_2)}{\partial x_1} \cdot h_1 + \frac{\partial f(p_1, p_2)}{\partial x_2} \cdot h_2 + o(h_1, h_2).$$

On a donc montré que f est différentiable en p et que

$$df_p(h) = \frac{\partial f}{\partial x_1} \cdot h_1 + \frac{\partial f}{\partial x_2} \cdot h_2. \quad (3.3)$$

Cette dernière relation signifie que la matrice de df est la matrice Jacobienne de f . $\square$

Corollaire 3.5. *Si $f : U \subset \mathbb{R}^m \rightarrow \mathbb{R}^n$ est de classe C^1 , alors l'application $p \mapsto df_p$ est continue.*

3.1.4 Une autre interprétation de la différentielle

La formule (3.3) suggère une autre façon de voir la différentielle d'une fonction. Remarquons d'abord que si $g : U \subset \mathbb{R}^m \rightarrow \mathbb{R}$ est une fonction différentiable à valeurs scalaires, alors dg_p est une *forme linéaire*, c'est-à-dire un élément du dual de $\mathbb{R}^m$ pour tout point p de U . Si, en particulier, g est elle-même une forme linéaire, alors on a $dg_p = g$ pour tout point p . On a donc la remarque suivante :

Pour tout système de coordonnées linéaires $x_1, \dots, x_m$ sur $\mathbb{R}^m$ on a en tout point⁴ $dx_i|_p = x_i$.

Ainsi pour tout vecteur $v = v_1 e_1 + \dots + v_m e_m$, on a

$$dx_i(v) = v_i = \langle e_i, v \rangle,$$

A condition toutefois que $x_1, \dots, x_m$ soit le système de coordonnées linéaires associé à la base $e_1, \dots, e_m$.

Considérons maintenant une application différentiable $f : U \subset \mathbb{R}^m \rightarrow \mathbb{R}^n$, alors nous avons en tout point p et pour tout $i = 1, \dots, m$:

$$df_p(e_i) = \lim_{t \rightarrow 0} \left(\frac{f(p + te_i) - f(p)}{t} \right) = \frac{\partial f}{\partial x_i}(p).$$

Pour le vecteur $v = v_1 e_1 + \dots + v_m e_m$, on a donc par linéarité de df :

$$df_p(v) = \sum_{i=1}^m df_p(e_i) v_i = \sum_{i=1}^m \frac{\partial f}{\partial x_i}(p) v_i = \sum_{i=1}^m \frac{\partial f}{\partial x_i}(p) dx_i(v)$$

car $v_i = dx_i(v)$. On écrit cette formule sous la forme classique suivante :

$$df = \sum_{i=1}^m \frac{\partial f}{\partial x_i} dx_i \quad (3.4)$$

Remarque 3.6. (1) Les raisonnements précédents montrent que l'image de df_p est le sous-espace vectoriel engendré par les vecteurs $\frac{\partial f}{\partial x_1}(p), \dots, \frac{\partial f}{\partial x_m}(p)$.

Il est important de ne pas oublier que si $n > 1$, alors f est une fonction à valeurs vectorielles, nous pouvons donc encore développer df dans la base canonique de $\mathbb{R}^n$ et écrire la formule précédente sous la forme

$$df = \sum_{i=1}^m \frac{\partial \vec{f}}{\partial x_i} dx_i = \sum_{i=1}^m \sum_{j=1}^n \frac{\partial f_j}{\partial x_i} \vec{u}_j dx_i,$$

où on a noté ici $\{\vec{u}_1, \dots, \vec{u}_n\}$ la base canonique de $\mathbb{R}^n$.

4. Noter que cette formule ne s'applique pas pour des coordonnées curvilignes (non linéaires). Par exemple on ne peut pas écrire $dr = r$ ou $d\theta = \theta$ dans le cas des coordonnées polaires (r, θ) du plan.

La notion de gradient

Considérons à nouveau une fonction différentiable $g : U \subset \mathbb{R}^m \rightarrow \mathbb{R}$ à valeurs scalaires. On a vu que sa différentielle en tout point p est la forme linéaire

$$dg_p = \sum_{i=1}^m \frac{\partial g}{\partial x_i} dx_i.$$

Définition 3.7. On appelle *gradient* de g en p le vecteur dual de la forme linéaire dg_p , où la dualité est induite par le produit scalaire standard de $\mathbb{R}^m$. Le gradient s'écrit

$$\vec{\nabla} g(p) = \sum_{i=1}^m \frac{\partial g}{\partial x_i} \mathbf{e}_i,$$

et il se caractérise par la condition

$$dg_p(\mathbf{v}) = \langle \vec{\nabla} g(p), \mathbf{v} \rangle, \quad \forall \mathbf{v} \in \mathbb{R}^m.$$

Notons encore que la matrice jacobienne de g en p est naturellement une matrice-ligne car dg_p est un élément du dual de $\mathbb{R}^n$. Puisque le gradient de g est un vecteur, il est représenté par une matrice colonne. Ainsi nous avons

$$dg = \left(\frac{\partial g}{\partial x_1} \quad \cdots \quad \frac{\partial g}{\partial x_m} \right) \quad \text{et} \quad \vec{\nabla} g = (dg)^\top = \begin{pmatrix} \frac{\partial g}{\partial x_1} \\ \vdots \\ \frac{\partial g}{\partial x_m} \end{pmatrix}$$

3.1.5 Le théorème d'inversion locale

Théorème 3.8. Soit U un ouvert de $\mathbb{R}^n$ et $f \in C^k(U, \mathbb{R}^n)$ (avec $k \geq 1$). Alors f est un C^k -difféomorphisme local au voisinage de $p \in U$ si et seulement si $J_f(p) \neq 0$.

La preuve a été vue au cours d'analyse 2, nous la donnons ci-dessous par souci de complétude.

Preuve. Quitte à remplacer f par l'application $x \mapsto f(x + p) - f(p)$, on se ramène au cas $p = f(p) = 0$. En composant ensuite f avec l'application linéaire df_0^{-1} , on peut supposer que $df_0 = \text{Id}$. Avec ces hypothèses, on a donc

$$f(x) = x + g(x),$$

où $g \in C^k(U, \mathbb{R}^n)$ vérifie $g(0) = 0$ et $dg_0 = 0$, c'est-à-dire $g(x) = o(\|x\|)$.

Nous devons construire un voisinage de 0 dans $\mathbb{R}^m$ sur lequel f est inversible. Comme $x \mapsto dg_x$ est continu et $dg_0 = 0$, il existe $r > 0$ tel que pour tout x dans la boule fermée $\bar{B}_r$ de centre 0 et de rayon r on a $\|dg\|_x \leq \frac{1}{3}$ (on prend r assez petit pour que $\bar{B}_r \subset U$). Cela implique que g est $\frac{1}{3}$ -Lipschitz sur cette boule, c'est-à-dire

$$x, x' \in \bar{B}_r \quad \Rightarrow \quad \|g(x') - g(x)\| \leq \frac{1}{3} \|x' - x\|.$$

En particulier $\|g(x)\| \leq \frac{r}{3}$ pour tout $x \in \bar{B}_r$. On va montrer que tout point $y \in B_{r/2}$ appartient à l'image de f par la méthode du point fixe. Observons que

$$f(x) = x + g(x) = y \quad \Leftrightarrow \quad x = y - g(x).$$

Fixons donc $y_0 \in B_{r/2}$ et définissons $T : \bar{B}_r \rightarrow \mathbb{R}^n$ par

$$T(x) = y_0 - g(x).$$

Observons d'abord que $T(\bar{B}_r) \subset \bar{B}_r$ car

$$\|x\| \leq r \Rightarrow \|T(x)\| = \|y_0 - g(x)\| \leq \|y_0\| + \|g(x)\| \leq \frac{r}{2} + \frac{r}{3} < r.$$

Ainsi T définit une transformation $T : \bar{B}_r \rightarrow \bar{B}_r$. Montrons qu'elle est strictement contractante :

$$x, x' \in \bar{B}_r \Rightarrow \|T(x') - T(x)\| = \|g(x') - g(x)\| \leq \frac{1}{3}\|x' - x\|.$$

L'application T possède donc un unique point fixe $x_0 \in \bar{B}_r$ tel que $T(x_0) = x_0$, c'est-à-dire $f(x_0) = y_0$. On a montré que $B_{r/2}$ est contenu dans l'image de f .

Considérons maintenant l'ouvert $V = B_r \cap f^{-1}(B_{r/2})$, alors, par construction, $f : V \rightarrow B_{r/2}$ est surjective. Cette application est aussi injective par unicité du point fixe de T .

Notons $h : B_{r/2} \rightarrow V$ l'inverse de $f|_V$ et montrons d'abord que h est continue en 0. Observons que pour $x \in V$ on a $f(x) = x + g(x)$, donc

$$\|x\| = \|f(x) - g(x)\| \leq \|f(x)\| + \|g(x)\| \leq \|f(x)\| + \frac{1}{3}\|x\|,$$

ce qui implique, en posant $x = h(y)$, que

$$\|h(y)\| = \|x\| \leq \frac{3}{2}\|f(x)\| = \frac{3}{2}\|y\|.$$

Nous pouvons montrer maintenant que h est différentiable en 0 et que sa différentielle en 0 est l'identité. Nous avons en effet :

$$\frac{\|h(y) - y\|}{\|y\|} = \frac{\|x - f(x)\|}{\|y\|} = \frac{\|g(x)\|}{\|y\|} \leq \frac{3}{2} \frac{\|g(x)\|}{\|x\|}$$

qui tend vers 0 lorsque $x \rightarrow 0$ (et cette condition est équivalente à $y \rightarrow 0$).

Nous avons démontré que si $f : U \rightarrow \mathbb{R}^n$ est C^k et si df_p est inversible en un point $p \in U$, alors f définit une bijection dans un voisinage de p et l'inverse f^{-1} est différentiable en $f(p)$. Il est clair que si df_p est inversible $p \in U$, alors df est inversible en tout point d'un voisinage de p (car le jacobien est une fonction continue). L'inverse f^{-1} est alors de classe C^k car la différentielle de f^{-1} au point $f(x)$ admet pour matrice jacobienne l'inverse de la matrice jacobienne de df_x . $\square$

Corollaire 3.9. *Une application $f : U \rightarrow V$ de classe C^k , avec $k \geq 1$ entre deux ouverts U et V de $\mathbb{R}^n$ est un difféomorphisme (global) si et seulement si f est bijective et $J_f(p) \neq 0$ pour tout $p \in U$.*

3.1.6 Le théorème du rang constant

Rappels d'algèbre linéaire. Rappelons que le *rang* d'une application linéaire $\ell : \mathbb{R}^m \rightarrow \mathbb{R}^n$ est la dimension de l'image $\text{Im}(\ell) \subset \mathbb{R}^n$. L'application linéaire ℓ est de rang r si et seulement si, après changement de bases sur $\mathbb{R}^m$ et sur $\mathbb{R}^n$, sa matrice prend la forme

$$A = \begin{pmatrix} I_r & 0 \\ 0 & 0 \end{pmatrix}.$$

Dans ces coordonnées, l'application ℓ s'écrit

$$\ell(x_1, \dots, x_m) = (x_1, \dots, x_r, 0, \dots, 0).$$

On montre aussi que le rang est le plus grand entier $r \in \mathbb{N}$ tel que la matrice A admet un $r \times r$ mineur non nul (i.e. une sous-matrice carrée de taille $r \times r$ et de déterminant non nul).

Définitions 3.10. ◦ Le *rang* d'une application $f : U \subset \mathbb{R}^m \rightarrow \mathbb{R}^n$ de classe C^1 est la fonction $\text{rang}_f : U \subset \mathbb{R}^m \rightarrow \mathbb{N}$ définie par $\text{rang}_f(p) = \text{rang}(df_p)$.

◦ On dit que f est de *rang maximal* en p si $\text{rang}_f(p) = \min(m, n)$.

◦ f est une *immersion* si $\text{rang}_f(p) = m$ pour tout $p \in U$ (de façon équivalente df_p est injective pour tout $p \in U$).

◦ f est une *submersion* si $\text{rang}_f(p) = n$ pour tout $p \in U$ (de façon équivalente df_p est surjective pour tout $p \in U$).

Exemple 3.11. a) Soit $F : \mathbb{R} \rightarrow \mathbb{R}^2$ l'application définie par $F(x) = (x^2, x^3)$. On a $DF_x = \begin{pmatrix} 2x \\ 3x^2 \end{pmatrix}$, et donc le rang de F vaut 0 en $(0, 0)$ et 1 en tout autre point.

b) Soit $G : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ l'application définie par $G(x_1, x_2) = (x_1^2 + x_2, x_1^3)$. Alors $DG = \begin{pmatrix} 2x_1 & 1 \\ 3x_1^2 & 0 \end{pmatrix}$.

Ainsi, le rang de G vaut 1 si $x_1 = 0$ et 2 sinon.

Lemme 3.12. Soit $f : U \subset \mathbb{R}^m \rightarrow \mathbb{R}^n$ une application de classe C^1 , alors la fonction $U \rightarrow \mathbb{N}$ définie par $p \mapsto \text{rang}_f(p)$ est *semi-continue inférieurement*. En particulier si f est de rang maximal en un point p , alors f est de rang maximal dans un voisinage de ce point.

Rappelons qu'une fonction $\rho : U \rightarrow \mathbb{R}$ est *semi-continue inférieurement* si pour tout $\alpha \in \mathbb{R}$ l'ensemble $\{x \in U \mid \rho(x) > \alpha\}$ est ouvert.

Preuve. L'application f vérifie $\text{rang}_f(p) \geq r$ si et seulement si la matrice jacobienne de df_p admet un $r \times r$ mineur non nul (i.e. si cette matrice contient une sous-matrice carrée de taille $r \times r$ dont le déterminant est non nul). Par continuité de la matrice jacobienne, ce même mineur est non nul dans un voisinage de p . □

Le théorème du rang constant affirme qu'une application de classe C^k dont le rang est constant est localement C^k -équivalente à une application linéaire :

Théorème 3.13 (Théorème du rang constant). Soit $f : U \subset \mathbb{R}^m \rightarrow \mathbb{R}^n$ une application de classe C^k et de rang constant $= r$. Pour tout point $p \in U$ il existe des voisinages V de p et W de $q = f(p)$ ainsi que des C^k -difféomorphismes

$$\Phi : V \rightarrow V' \subset \mathbb{R}^m \quad \text{et} \quad \Psi : W \rightarrow W' \subset \mathbb{R}^n$$

tels que

(i) $V \subset U$ et $f(V) \subset W$,

(ii) $\Phi(p) = 0 \in \mathbb{R}^m$ et $\Psi(q) = 0 \in \mathbb{R}^n$,

(iii) l'application $F = \Psi \circ f \circ \Phi^{-1} : V' \rightarrow W'$ s'écrit

$$F(x_1, \dots, x_n) = (x_1, \dots, x_r, 0, \dots, 0).$$

Notons en particulier que F est une application linéaire, alors que f ne l'est pas en général. La situation peut se représenter sur le diagramme suivant :

$$\begin{array}{ccc} V & \xrightarrow{f} & W \\ \Phi \downarrow & & \downarrow \Psi \\ V' & \xrightarrow{F} & W' \end{array}$$

Preuve. On peut supposer, quitte à faire des translations, que $p = 0 \in \mathbb{R}^m$ et $q = f(p) = 0 \in \mathbb{R}^n$. Quitte à faire des changements de bases sur $\mathbb{R}^m$ et $\mathbb{R}^n$, on peut aussi supposer que df_0 prenne la forme normale, c'est-à-dire que la matrice jacobienne de f en 0 est la $n \times m$ matrice

$$DF_0 = \left(\frac{\partial f_i}{\partial x_j}(0) \right) = \begin{pmatrix} I_r & 0 \\ 0 & 0 \end{pmatrix}$$

Définissons l'application suivante $\Phi : U \rightarrow \mathbb{R}^m$

$$\Phi(x_1, \dots, x_m) = (f_1(x), \dots, f_r(x), x_{r+1}, \dots, x_m),$$

et observons que la matrice jacobienne de Φ en 0 est la $m \times m$ matrice identité, car

$$D\Phi_0 = \left(\frac{\partial \Phi_i}{\partial x_j}(0) \right) = \begin{pmatrix} I_r & 0 \\ 0 & I_{m-r} \end{pmatrix}.$$

Par le théorème d'inversion locale, on sait que Φ définit un difféomorphisme de classe C^k d'un voisinage V de $0 \in \mathbb{R}^m$ sur un autre voisinage V' de $0 \in \mathbb{R}^m$.

On considère maintenant l'application $f \circ \Phi^{-1} : V' \rightarrow \mathbb{R}^n$, cette application s'écrit

$$f \circ \Phi^{-1}(x_1, \dots, x_m) = (x_1, \dots, x_r, f^{r+1} \circ \Phi^{-1}(x), \dots, f^n \circ \Phi^{-1}(x)),$$

et sa matrice jacobienne en x est une $m \times n$ matrice de la forme

$$\begin{pmatrix} I_r & 0 \\ * & \Delta(x) \end{pmatrix}, \quad \text{avec} \quad \Delta(x) = \left(\frac{\partial}{\partial x_j} f_i \circ \Phi^{-1} \right)_{(r+1) \leq i, j \leq n}.$$

Or nous savons que cette matrice est de rang r pour tout $x \in V'$ car $\text{rang } f \circ \Phi^{-1} = \text{rang } f = r$. Par conséquent $\Delta(x)$ est la matrice nulle pour tout $x \in V'$. Ainsi $f \circ \Phi^{-1}$ ne dépend que des variables $x_1, \dots, x_r$ et on peut donc écrire

$$f \circ \Phi^{-1}(x) = (x_1, \dots, x_r, h(x_1, \dots, x_r))$$

On définit maintenant une application $\Psi : f(U) \rightarrow \mathbb{R}^n$ par

$$\Psi(y) = (y_1, \dots, y_r, y_{r+1} - h_{r+1}(y_1, \dots, y_r), y_n - h_n(y_1, \dots, y_r)).$$

La matrice jacobienne de Ψ en 0 est une $n \times n$ matrice du type

$$\left(\frac{\partial \Psi_i}{\partial y_j} \right) = \begin{pmatrix} I_r & 0 \\ * & I_{n-r} \end{pmatrix}.$$

Cette matrice est inversible, donc Ψ définit un C^k -difféomorphisme d'un voisinage W de $0 \in \mathbb{R}^n$ vers un voisinage W' de 0. On vérifie finalement que $F := \Psi \circ f \circ \Phi^{-1} : V' \rightarrow W'$ est donné par

$$\Psi \circ f \circ \Phi^{-1}(x_1, \dots, x_n) = (x_1, \dots, x_r, 0, \dots, 0).$$

□

Une conséquence importante du théorème du rang constant est le

Théorème 3.14 (Théorème des fonctions implicites). *Soit $f : U \subset \mathbb{R}^m \rightarrow \mathbb{R}^n$ une application C^∞ où $m = n + k$. Soit p un point de U et supposons que la matrice Jacobienne partielle (de taille $n \times n$)*

$$\left(\frac{\partial f_i}{\partial x_j} \right)_{1 \leq i \leq n, (k+1) \leq j \leq m}$$

est inversible en p . Alors il existe un voisinage de p de type $V \times W \subset U$ avec $V \subset \mathbb{R}^k$ et $W \subset \mathbb{R}^n$ ainsi qu'une application $\phi : V \rightarrow W$ de classe C^∞ telle que pour tout $x \in V \times W$ on a

$$f(x) = q \quad \Leftrightarrow \quad (x_{k+1}, \dots, x_m) = \phi(x_1, \dots, x_k)$$

où $q = f(p)$.

Nous laissons la preuve en exercice. Remarquons que le rang de f en p est égale à n , donc par le Lemme 3.12, le rang est constant, égal à n dans un voisinage de p . On peut donc supposer que f est une submersion.

3.2 Sous-Variétés de $\mathbb{R}^n$

Une sous-variété de dimension m dans l'espace $\mathbb{R}^n$ est un sous-ensemble qui peut localement, c'est-à-dire au voisinage de chaque point, être approximé par un sous-espace affine de dimension $m \leq n$. On suppose que des conditions de régularité (i.e. de différentiabilité) sont vérifiées. Voici la définition précise.

Définition. Un sous-ensemble $M \subset \mathbb{R}^n$ est une *sous-variété de dimension m de classe C^k* si pour tout point p de M il existe un voisinage $U \subset \mathbb{R}^n$ de p et un C^k -difféomorphisme $\phi : U \rightarrow V$, où V est un ouvert de $\mathbb{R}^n$, tel que

$$\phi(U \cap M) = V \cap E$$

où $E \subset \mathbb{R}^n$ est un sous-espace affine de dimension m .

On dit aussi que $(n - m)$ est la *codimension* de M . Une sous-variété de dimension 2 s'appelle une *surface* et une sous-variété de dimension 1 est une *courbe*. Une sous-variété de codimension 1, donc de dimension $(n - 1)$, s'appelle une *hypersurface*.

Remarque. Sans perdre de généralité, On peut remplacer dans cette définition les mots *sous-espace affine* par *sous-espace vectoriel*. On peut même supposer que $E \subset \mathbb{R}^m$ est le sous-espace vectoriel engendré par les m premiers vecteurs de la base canonique :

$$E = \{y \in \mathbb{R}^m \mid y_{m+1} = \dots = y_n = 0\}.$$

Dans la pratique, vérifier que M est une sous-variété de dimension m revient à montrer que pour tout point $p \in M$, il existe un système de coordonnées curvilignes $y_1, \dots, y_n$ dans un voisinage $U \subset \mathbb{R}^n$ de p tel que

$$\phi(U \cap M) = V \cap \{y \in V \mid y_{m+1} = \dots = y_n = 0\},$$

où $\phi(x_1, \dots, x_n) = (y_1, \dots, y_n)$ est le difféomorphisme qui représente le changement de coordonnées.

On peut alors considérer que les m premières coordonnées $y_1, \dots, y_m$ sont des “coordonnées curvilignes locales” sur la sous variété M ; elle servent à paramétrer une région de la variété au voisinage du point p .

Observons aussi que les coordonnées restantes $y_{m+1}, \dots, y_n$ sont nulles sur la sous-variété, elle représentent donc un système de $n - m$ équations (en général non linéaires) qui définissent localement la variété.

Premiers exemples.

- (i) Une sous-variété de dimension 0 de $\mathbb{R}^n$ est un sous-ensemble discret (tous ses points sont isolés).
- (ii) Une sous-variété de dimension n de $\mathbb{R}^n$ est un sous-ensemble ouvert de $\mathbb{R}^n$.
- (iii) L'ensemble vide est une sous-variété de dimension m pour tout entier m .
- (iv) Si $f : U \rightarrow \mathbb{R}$ est une fonction de classe C^k définie sur un ouvert $U \subset \mathbb{R}^m$, alors son graphe

$$M = \{(x, t) \in U \times \mathbb{R} \mid t = f(x)\}$$

est une sous-variété de $\mathbb{R}^{m+1}$.

Les trois premiers exemples sont banals. S'agissant du quatrième exemple, pour montrer que le graphe M de la fonction $f \in C^k(U)$ est une sous-variété, on considère l'application $\phi : U \times \mathbb{R} \rightarrow U \times \mathbb{R}$ définie par

$$\phi(x_1, \dots, x_m, x_{m+1}) = (x_1, \dots, x_m, x_{m+1} - f(x_1, \dots, x_m)).$$

Cette application est clairement de classe C^k et c'est un difféomorphisme dont l'inverse est explicitement donné par

$$\phi^{-1}(y_1, \dots, y_m, y_{m+1}) = (y_1, \dots, y_m, y_{m+1} + f(y_1, \dots, y_m)).$$

Il est clair que $\phi(M) = U \times \{0\} \subset \mathbb{R}^{m+1}$, qui est un ouvert d'un sous-espace vectoriel de dimension m .

Théorème 3.15. Soit $f : U \rightarrow \mathbb{R}^n$ une application de classe C^k et de rang constant r , où U est un ouvert de $\mathbb{R}^m$. Alors on a les conclusions suivantes :

- A)** Pour chaque point $q \in \mathbb{R}^n$, la préimage $f^{-1}(q) \subset U$ est une sous-variété différentiable de codimension r (i.e. de dimension $m - r$).
- B)** Chaque point $p \in U$ admet un voisinage $V_p \subset U$ tel que l'image directe $f(V_p) \subset \mathbb{R}^n$ est une sous-variété de dimension r .

En particulier :

- a) Si $f : U \subset \mathbb{R}^m \rightarrow \mathbb{R}^n$ est une submersion de classe C^k , alors $f^{-1}(q)$ est une sous-variété de codimension n .
- b) Si $f : U \subset \mathbb{R}^m \rightarrow \mathbb{R}^n$ est une immersion de classe C^k , alors chaque point $p \in U$ admet un voisinage $V_p \subset U$ tel que l'image directe $f(V_p) \subset \mathbb{R}^n$ est une sous-variété de dimension m .

Remarque. En général, l'image $f(U) \subset \mathbb{R}^n$ d'une immersion $f : U \rightarrow \mathbb{R}^n$ n'est pas une sous-variété, même si f est injectif. Il est par contre facile de voir que $f(U)$ est une sous-variété si f est une immersion et f définit un homéomorphisme de U vers M .

Des exemples pour le cas A :

- 1) Si $F : \mathbb{R}^m \rightarrow \mathbb{R}$ est une fonction C^∞ telle que $dF_p \neq 0$ quel que soit p , alors $F^{-1}(q)$ est une hypersurface de $\mathbb{R}^m$.
- 2) La sphère $\mathbb{S}^{n-1}$ est une hypersurface (car la sphère est définie par l'équation $\sum_i x_i^2 = 1$).
- 3) Si $F, G : \mathbb{R}^m \rightarrow \mathbb{R}$ sont tels que dF_p et dG_p sont linéairement indépendants quel que soit p , alors $F^{-1}(q_1) \cap G^{-1}(q_2)$ est une sous-variété de codimension 2.
- 4) Le groupe orthogonal $O(n)$ est une sous-variété de $M_n(\mathbb{R})$ (car $O(n)$ est défini par l'équation $X^\top X = I_n$).

Preuve du théorème. Le théorème du rang constant nous dit que l'application f est localement équivalente à une application linéaire de rang r . Donc la préimage d'un point par f et son image directe sont localement équivalents à des sous-espaces vectoriels, or c'est précisément cela la définition d'une sous-variété.

Commençons par démontrer l'affirmation (A). Fixons $q \in \mathbb{R}^n$, si $q \notin f(U)$, alors $f^{-1}(q) = \emptyset$ et il n'y a rien à montrer. On suppose donc que $q \in f(U)$ et on choisit un point $p \in M = f^{-1}(q)$. Par le théorème du rang constant on sait qu'il existe des ouverts $V, V' \subset \mathbb{R}^m$, et $W, W' \subset \mathbb{R}^n$ tels que $V \subset U$, $f(V) \subset W$, $p \in V$, $q \in W$, ainsi que des difféomorphismes $\Psi : W \rightarrow W'$ et $\Phi : V \rightarrow V'$ tels que $\Phi(p) = 0$, $\Psi(q) = 0$ et $F = \Psi \circ f \circ \Phi^{-1} : V' \rightarrow W'$ s'écrit $F(x_1, \dots, x_n) = (x_1, \dots, x_r, 0, \dots, 0)$. On a alors

$$\Phi(M \cap V) = V' \cap \{x \in \mathbb{R}^m \mid x_1 = x_2 = \dots = x_r = 0\} \subset \mathbb{R}^m.$$

Ceci démontre que $M = f^{-1}(q)$ est différentiablement équivalent à un ouvert d'un sous-espace vectoriel de dimension $m - r$ au voisinage de chacun de ses points. Par définition M est donc une sous-variété de dimension $m - r$ de $\mathbb{R}^m$.

Montrons maintenant l'affirmation (B). Fixons un point quelconque $p \in U$ et considérons des voisinages V de p et W de $q = f(p)$ ainsi que des difféomorphismes $\Psi : W \rightarrow W'$ et $\Phi : V \rightarrow V'$ comme plus haut. Notons $M = f(V)$, alors

$$\Psi(M \cap W) = W' \cap \{x \in \mathbb{R}^m \mid x_{r+1} = \dots = x_{n-1} = x_n = 0\} \subset \mathbb{R}^n,$$

ce qui prouve que M est une sous-variété de dimension $n - (n - r) = r$ de $\mathbb{R}^n$. □

3.3 L'espace tangent à une sous-variété

Un vecteur $\mathbf{v} \in \mathbb{R}^n$ est tangent en un point p à une sous-variété $M \subset \mathbb{R}^n$ si c'est le vecteur vitesse d'une courbe différentiable contenue dans la variété et passant par p . Plus précisément ;

Définition. Soit $M \subset \mathbb{R}^n$ une sous-variété de classe C^1 de dimension m et soit p un point de M . On dit qu'un vecteur $\mathbf{v} \in \mathbb{R}^n$ est un *vecteur tangent* à M s'il existe une courbe $\gamma : (-\varepsilon, \varepsilon) \rightarrow M$ de classe C^1 telle que

$$\gamma(0) = p \quad \text{et} \quad \dot{\gamma}(0) = \mathbf{v}.$$

Dans ce cas on dit que la courbe γ *représente* le vecteur tangent $\mathbf{v}$. On note $T_p M$ l'ensemble des vecteurs tangents à M en p .

Exemple. Si $U \subset \mathbb{R}^n$ est un ouvert, alors $T_p U$ est canoniquement isomorphe à $\mathbb{R}^n$, car tout vecteur $\mathbf{v} \in \mathbb{R}^n$ est le vecteur vitesse de la courbe $\gamma_{\mathbf{v}}(t) = p + t\mathbf{v}$ (cette courbe est contenu dans l'ouvert U pour $|t| < \varepsilon$ assez petit).

Proposition 3.16. *En chaque point p d'une sous-variété différentiable $M \subset \mathbb{R}^n$ de dimension m , l'espace tangent $T_p M$ est un sous-espace vectoriel de dimension m de $\mathbb{R}^n$.*

Preuve. On sait par définition de la notion de variété qu'il existe un voisinage U de p et un difféomorphisme $\phi : U \rightarrow V$ tel que $\phi(p) = 0$ et

$$\phi(U \cap M) = V \cap E$$

où $E \subset \mathbb{R}^n$ est un sous-espace vectoriel de dimension m .

Soit $\mathbf{v} \in T_p M$ un vecteur tangent à M en p . Par définition il existe une courbe $\alpha : (-\varepsilon, \varepsilon) \rightarrow M \cap U$ de classe C^1 telle que $\alpha(0) = p$ et $\mathbf{v} = \dot{\alpha}(0)$. Définissons la courbe $\beta : (-\varepsilon, \varepsilon) \rightarrow V \cap E$ par $\beta(t) = \phi \circ \alpha(t)$, alors on a

$$\mathbf{v} = \left. \frac{d}{dt} \right|_{t=0} \alpha(t) = \left. \frac{d}{dt} \right|_{t=0} \phi^{-1}(\beta(t)) = d\phi_0^{-1}(\dot{\beta}(0)) \in d\phi_0^{-1}(E),$$

ce qui montre que $T_p M$ est inclus dans l'espace vectoriel $d\phi_0^{-1}(E)$.

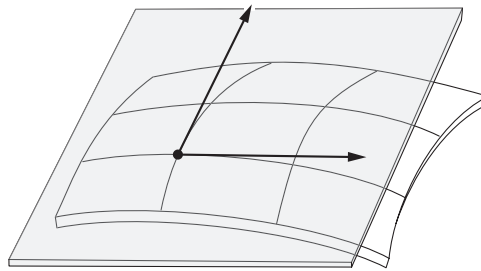
Pour montrer l'inclusion inverse, on considère un vecteur quelconque $\mathbf{w} \in E$, alors pour $\varepsilon > 0$ assez petit, la courbe $\beta : (-\varepsilon, \varepsilon) \rightarrow E$ définie par $\beta(t) = t\mathbf{w}$ prend ses valeurs dans $V \cap E$. Notons $\alpha = \phi^{-1} \circ \beta$, alors $\alpha : (-\varepsilon, \varepsilon) \rightarrow M \cap U$ est une courbe C^1 telle que $\alpha(0) = p$, par conséquent $\dot{\alpha}(0) \in T_p M$. Mais on a

$$d\phi_0^{-1}(\mathbf{w}) = d\phi_0^{-1}(\dot{\beta}(0)) = \left. \frac{d}{dt} \right|_{t=0} \phi^{-1}\beta(t) = \dot{\alpha}(0) \in T_p M,$$

par conséquent $d\phi_0^{-1}(E) \subset T_p M$. On a montré que

$$T_p M = d\phi_0^{-1}(E),$$

qui est bien un sous-espace vectoriel de dimension m car $d\phi_0^{-1}$ est un isomorphisme de $\mathbb{R}^n$ dans lui-même. □



Plan tangent à une surface.

Proposition 3.17. *Si $f : U \subset \mathbb{R}^n \rightarrow \mathbb{R}^k$ une application différentiable de rang constant r , alors pour tout point p de la sous-variété $M = f^{-1}(q)$, on a $T_p M = \text{Ker}(df_p)$.*

Preuve. Sous les hypothèses de la proposition, M est une sous-variété de $\mathbb{R}^n$ de dimension $m = n - r$. En particulier l'espace tangent $T_p M$ est un sous-espace vectoriel de dimension m .

Montrons que $T_p M \subset \text{Ker}(df_p)$. Soit $\mathbf{v} \in T_p M$ un vecteur tangent M en p et $\alpha : (-\varepsilon, \varepsilon) \rightarrow M$ une courbe qui représente $\mathbf{v}$, alors

$$df_p(\mathbf{v}) = df_p(\dot{\alpha}(0)) = \left. \frac{d}{dt} \right|_{t=0} f(\alpha(t)) = 0,$$

car $f(\alpha(t)) = q$ pour tout t puisque $\alpha(t) \in M$. Ceci montre que $T_p M \subset \text{Ker}(df_p)$. Mais d'autre part, $df_p : \mathbb{R}^n \rightarrow \mathbb{R}^k$ est une application linéaire de rang r , donc son noyau $\text{Ker}(df_p)$ est un sous-espace vectoriel de dimension $n - r = m$. On conclut que $T_p M = \text{Ker}(df_p)$. $\square$

Exemple. Considérons le cas d'une hypersurface $M = f^{-1}(0)$ où $f : U \subset \mathbb{R}^n \rightarrow \mathbb{R}$ est une submersion. Alors pour tout point $p \in M = f^{-1}(0)$ on a

$$T_p M = \{\mathbf{v} \in \mathbb{R}^n \mid df_p(\mathbf{v}) = 0\} = \{\mathbf{v} \in \mathbb{R}^n \mid \langle \vec{\nabla} f(p), \mathbf{v} \rangle = 0\},$$

l'espace tangent en p à M est donc le sous-espace vectoriel orthogonal au gradient $\vec{\nabla} f(p)$.

Proposition 3.18. Si $f : U \subset \mathbb{R}^k \rightarrow \mathbb{R}^n$ est une application différentiable de rang constant r telle que $M = f(U) \subset \mathbb{R}^n$ est une sous-variété (de dimension r), alors pour tout $q = f(p) \in M$, on a $T_q M = \text{Im}(df_p)$.

Preuve. Observons que la courbe $t \mapsto f(p + t\mathbf{e}_i)$ est tracée sur la sous-variété M et passe par le point q en $t = 0$, par conséquent le vecteur

$$\mathbf{b}_i := \left. \frac{d}{dt} \right|_{t=0} f(p + t\mathbf{e}_i) = \frac{\partial f}{\partial x_i}(p) = df_p(\mathbf{e}_i)$$

est un élément de $T_q M$ pour tout $i = 1, \dots, k$. Les vecteurs $\{\mathbf{b}_1, \dots, \mathbf{b}_k\}$ engendrent le sous-espace vectoriel $\text{Im}(df_p)$, par conséquent $\text{Im}(df_p) \subset T_q M$. Or ces deux sous-espaces vectoriels sont de dimension r , ils sont par conséquent égaux. $\square$

Remarque. Les vecteurs $\{\mathbf{b}_1, \dots, \mathbf{b}_k\}$ de la preuve précédente forment les colonnes de la matrice Jacobienne de f en p . Ils sont linéairement indépendants si et seulement si le rang de f en p est égale à k .

Exemple. Considérons le graphe de la fonction différentiable $\varphi : U \rightarrow \mathbb{R}$ où U est un ouvert de $\mathbb{R}^2$, notons S cette surface et p un point de S . Il y a deux façons de comprendre le plan tangent $T_p S$.

- (i) Point de vue implicite : La surface S est l'ensemble des zéros de la fonction $f : U \times \mathbb{R} \rightarrow \mathbb{R}$ définie par $f(x, y, z) = z - \varphi(x, y)$. Le gradient de f est $\vec{\nabla} f = (-\varphi_x, -\varphi_y, 1)$ et le plan tangent en $p = (x, y, \varphi(x, y))$ admet l'équation

$$T_p S = \left(\vec{\nabla} f(p) \right)^\perp = \{\mathbf{v} = (v_1, v_2, v_3) \mid v_3 = \varphi_x v_1 + \varphi_y v_2\}$$

où on a noté pour simplifier $\varphi_x = \frac{\partial \varphi}{\partial x}$ et $\varphi_y = \frac{\partial \varphi}{\partial y}$.

- (ii) Point de vue paramétrique : La surface S est l'image de U par l'application $h : U \rightarrow \mathbb{R}^3$ définie par $h(x, y) = (x, y, \varphi(x, y))$. Alors $T_p S$ est le sous-espace vectoriel engendré par

$$\mathbf{b}_1 = \frac{\partial h}{\partial x} = (1, 0, \varphi_x) \quad \text{et} \quad \mathbf{b}_2 = \frac{\partial h}{\partial y} = (0, 1, \varphi_y).$$

On vérifie facilement que ces deux descriptions donnent le même sous-espace vectoriel.

Remarque. L'espace tangent $T_p M$ d'une sous-variété $M \subset \mathbb{R}^n$ est un sous-espace vectoriel de $\mathbb{R}^n$. Ce sous-espace n'est pas *géométriquement* tangent à la sous-variété M (il peut même être disjoint de M). Pour cette raison on introduit la notion suivante :

Définition 3.19. L'espace affine tangent à une sous-variété différentiable $M \subset \mathbb{R}^n$ en un point $p \in M$ est le sous-espace affine

$$A_p M = \{q \in \mathbb{R}^n \mid (q - p) \in T_p M\} = p + T_p M.$$

Observons que le point p appartient à l'intersection $M \cap A_p M$ et le sous-espace affine est géométriquement tangent à M en ce point.

Exemple. Si $f : U \subset \mathbb{R}^n \rightarrow \mathbb{R}$ est une fonction de classe C^1 telle que $df_p \neq 0$ en tout $p \in U$, alors l'espace affine tangent à l'hypersurface $M = f^{-1}(0)$ en p est l'hyperplan

$$A_p M := \left\{ x \in \mathbb{R}^n \mid \sum_{i=1}^n \frac{\partial f}{\partial x_i}(p)(x_i - p_i) = 0 \right\}.$$

A titre d'exemple concret, le plan tangent en $p = (x_0, y_0, z_0)$ à l'ellipsoïde $\frac{x^2}{a^2} + \frac{y^2}{b^2} + \frac{z^2}{c^2} = 1$ est le plan d'équation

$$\frac{x_0(x - x_0)}{a^2} + \frac{y_0(y - y_0)}{b^2} + \frac{z_0(z - z_0)}{c^2} = 0.$$

3.4 Applications différentiables entre deux sous-variétés

La notion d'application différentiable entre des ouverts $U \subset \mathbb{R}^m$ et $V \subset \mathbb{R}^n$ se généralise au cas des applications entre deux sous-variétés.

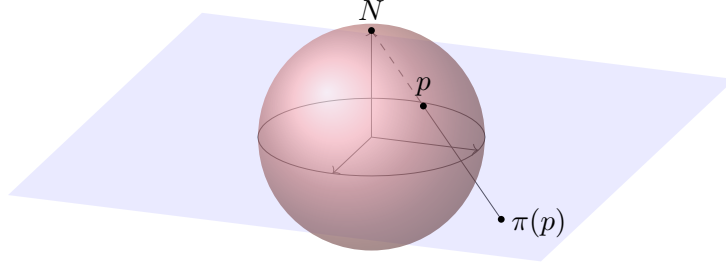
Définitions. Soient $M \subset \mathbb{R}^d$ et $N \subset \mathbb{R}^\ell$ deux sous-variétés différentiables de classe C^k et $f : M \rightarrow N$ une application entre ces deux sous-variétés. On dit que f est *différentiable de classe C^k* au voisinage du point $p \in M$ s'il existe un voisinage ouvert $U \subset \mathbb{R}^d$ de p et une application $F : U \rightarrow \mathbb{R}^\ell$ de classe C^k telle que F et f coïncident sur l'intersection $U \cap M$, c'est-à-dire qu'on a $F|_{M \cap U} = f|_{M \cap U}$. L'application F s'appelle alors une *extension locale de f au voisinage de p* . On dit $f : M \rightarrow N$ différentiables de classe C^k et on note $f \in C^k(M, N)$ si f est différentiable au voisinage de tout point de M .

On distingue certaines applications différentiables particulières :

- (a) On dit que $f \in C^k(M, N)$ est un *difféomorphisme* de classe C^k si f est bijective et $f^{-1} \in C^k(N, M)$.
- (b) On dit que $\psi : \Omega \rightarrow M$ est une *paramétrisation locale* de classe C^k si Ω est un ouvert de $\mathbb{R}^m$ et ψ est un difféomorphisme de classe C^k de l'ouvert $\Omega \subset \mathbb{R}^m$ vers son image $W = \psi(\Omega) \subset M$. Lorsque ψ est bijective, on dit que c'est une *paramétrisation globale* de M .
- (c) On dit que $\varphi : W \rightarrow U$ est une *carte locale* de classe C^k pour la variété M si W est un ouvert relatif⁵ de M (appelé le *domaine* de la carte), U est un ouvert de $\mathbb{R}^m$ et φ est un difféomorphisme de classe C^k . Dans ce cas on doit avoir $m = \dim(M)$. Une carte locale est donc l'inverse d'une paramétrisation locale.

Exemple. La projection stéréographique est l'application $\pi : S^n \setminus \{N\} \rightarrow \mathbb{R}^n$ définie sur le complémentaire du "pôle nord" N de la sphère unité $S^{n-1} \subset \mathbb{R}^{n+1}$ (c'est-à-dire le point $(0, \dots, 0, 1)$) qui envoie le point $P \in S^n \setminus \{N\}$ sur l'intersection P' de la droite par N et P avec $\mathbb{R}^n$ (vu comme l'hyperplan de $\mathbb{R}^{n+1}$ défini par $\{x_{n+1} = 0\}$). La projection stéréographique est donc une carte de la sphère dont le domaine est le complémentaire du pôle nord.

5. Un sous-ensemble $W \subset M$ est un *ouvert relatif* de M s'il existe un ouvert $V \subset \mathbb{R}^n$ tel que $W = V \cap M$.



La différentielle d'une application de classe C^1 est définie a priori pour les applications entre des ouverts de $\mathbb{R}^n$. La proposition suivante permet de généraliser cette notion importante au cas des applications entre deux sous-variétés.

Proposition 3.20. *Soit $f : M \rightarrow N$ une application entre deux variétés différentiables et $p \in M$. Supposons que f soit de classe C^1 au voisinage de p . Alors pour toute extension locale $F : U \rightarrow \mathbb{R}^\ell$ de f au voisinage de p , on a $dF_p(T_p M) \subset T_q N$, où $q = f(p)$. De plus la restriction $dF_p|_{T_p M} : T_p M \rightarrow T_q N$ ne dépend pas de l'extension locale de f choisie.*

On notera $df_p : T_p M \rightarrow T_q N$ l'application ainsi définie et on dit que c'est la *différentielle de l'application $f : M \rightarrow N$ en p* .

Preuve. Par définition de la notion de vecteur tangent, il existe une courbe $\gamma : (-\varepsilon, \varepsilon) \rightarrow M$ de classe C^1 telle que $\gamma(0) = p$ et $\dot{\gamma}(0) = v$. Si on suppose $\varepsilon > 0$ assez petit, alors $\gamma(-\varepsilon, \varepsilon) \subset U \cap M$ et la règle de dérivation en chaîne appliquée à $F \circ \gamma$ nous dit que

$$w = dF_p(v) = dF_p(\dot{\gamma}(0)) = \left. \frac{d}{dt} \right|_{t=0} F(\gamma(t)) = \left. \frac{d}{dt} \right|_{t=0} f(\gamma(t)),$$

car par définition on a $F|_{M \cap U} = f|_{M \cap U}$. Cela montre d'une part que l'image w ne dépend que de f et non de l'extension F choisie et d'autre part que w appartient à $T_q N$ puisque la courbe $\alpha = f \circ \gamma : (-\varepsilon, \varepsilon) \rightarrow N$ vérifie $\alpha(0) = q$ et $\dot{\alpha}(0) = w$. □

Remarquons que cette preuve nous donne une interprétation très naturelle de la différentielle $df_p : T_p M \rightarrow T_q N$: si le vecteur tangent $v \in T_p M$ est représenté par la courbe γ tracée sur M , alors $df_p(v) \in T_q N$ est le vecteur tangent représenté par la courbe $f \circ \gamma$.

Proposition 3.21 (Règle de dérivation en chaîne). *Si $f : M_1 \rightarrow M_2$ et $g : M_2 \rightarrow M_3$ sont des applications différentiables de classe C^k entre des sous-variétés différentiables, alors $g \circ f : M_1 \rightarrow M_3$ est différentiable de classe C^k et pour tout point $p \in M_1$ on a*

$$d(g \circ f)_p = dg_{f(p)} \circ df_p,$$

qui est une application linéaire de $T_p M_1$ vers $T_{g(f(p))} M_3$.

Preuve. Il suffit d'appliquer la règle de dérivation en chaîne classique à des extensions F et G des applications f et g à des voisinages de p , respectivement $f(p)$. □

3.5 Le fibré tangent à une sous-variété

Définition. On appelle *espace tangent total* ou *fibré tangent* à la sous-variété $M \subset \mathbb{R}^n$ le sous-ensemble de $\mathbb{R}^n \times \mathbb{R}^n$ défini par

$$TM = \{(p, \mathbf{v}) \in \mathbb{R}^n \times \mathbb{R}^n \mid p \in M \text{ et } \mathbf{v} \in T_p M\}.$$

Proposition 3.22. (a) Si $M \subset \mathbb{R}^n$ est une sous-variété de dimension m et de classe C^k avec $k \geq 2$, alors $TM \subset \mathbb{R}^{2n}$ est une sous-variété de dimension $2m$ et de classe C^{k-1} .

(b) Si $f : M \rightarrow N$ est une application différentiable de classe C^k entre deux sous-variétés de classe C^k , alors l'application

$$Tf : TM \rightarrow TN \quad \text{définie par } Tf(p, v) = (f(p), df_p(v))$$

est une application différentiable de classe C^{k-1} .

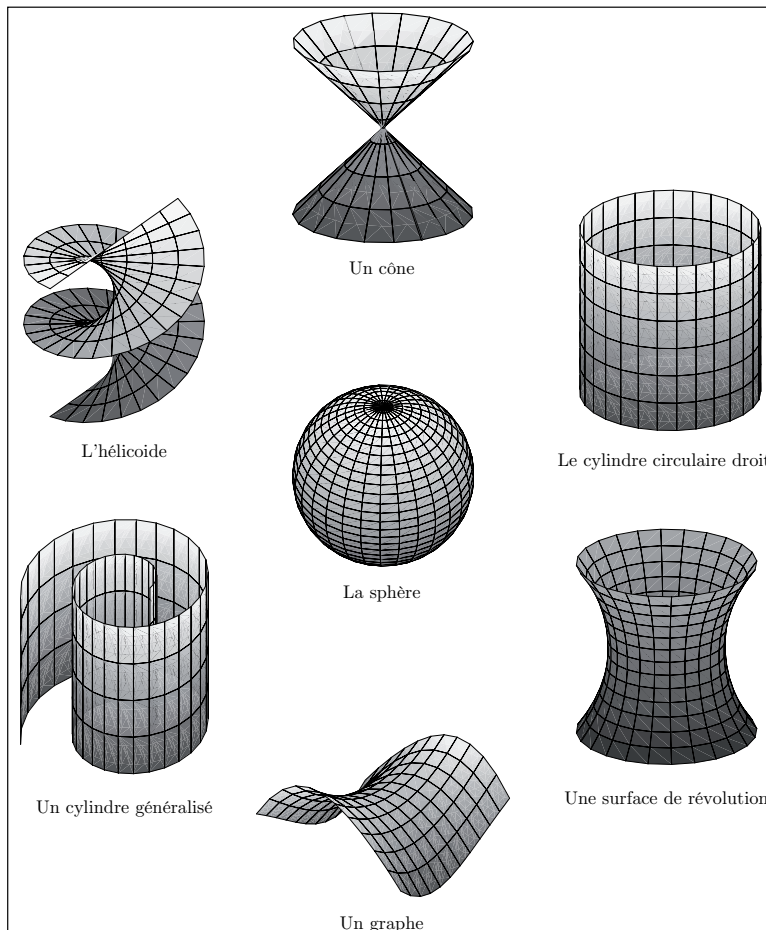
(c) Les sous-variétés différentiables de classe C^∞ forment une catégorie dont les morphismes sont les applications de classe C^∞ . La correspondance $(M \rightarrow TM, f \rightarrow Tf)$ définit un foncteur covariant de cette catégorie dans elle-même.

Nous laissons la preuve en exercice.

Chapitre 4

Géométrie des sous-variétés

Dans ce chapitre nous étudions la géométrie des sous-variétés du point de vue des distances, aires et volumes, angles etc. La notion de courbure des surfaces sera étudiée au chapitre 5



4.1 Distances extrinsèque et intrinsèque sur une sous-variété

Définitions. (i) La *distance extrinsèque* entre deux points p et q sur une sous-variété $M \subset \mathbb{R}^n$ est la distance euclidienne $\|q - p\|$ entre ces deux points.

(ii) La *distance intrinsèque* entre deux points p et q sur une sous-variété connexe $M \subset \mathbb{R}^n$ de classe C^1 est l'infimum des longueurs des courbes de classe C^1 par morceaux tracées sur la sous-variété et qui relient ces deux points. On note cette distance par

$$d_M(p, q) = \inf\{\ell(\gamma) \mid \gamma : [a, b] \rightarrow M, \gamma \text{ est } C^1 \text{ par morceaux, } \gamma(a) = p, \gamma(b) = q\}. \quad (4.1)$$

(iii) Deux sous-variété connexes $M_1 \subset \mathbb{R}^n$ et $M_2 \subset \mathbb{R}^d$ de classe C^1 dans sont dites *intrinsèquement isométriques* s'il existe une application bijective $f : M_1 \rightarrow M_2$ telle que $d_{M_2}(f(p), f(q)) = d_{M_1}(p, q)$ pour tous $p, q \in M_1$. Dans ce cas on dit que l'application f est une *isométrie* entre les deux sous-variétés.

Le lecteur vérifiera sans difficulté que la distance ainsi définie satisfait aux propriétés habituelles d'une distance, et donc (M, d_M) est un espace métrique. Remarquons aussi que

$$d_M(p, q) \geq \|q - p\|$$

pour toute paire de points p et q de M .

Exemple. La distance intrinsèque entre deux points p et q d'une sphère S de centre c et de rayon R dans $\mathbb{R}^n$ est égale à $d_S(p, q) = R\alpha$, où $\alpha = \angle_c(p, q)$ est l'angle entre les vecteurs $(p - c)$ et $(q - c)$. La distance extrinsèque est égale à $\|q - p\| = 2R \sin(\alpha/2)$. L'inégalité précédente est donc dans ce cas l'inégalité

$$\alpha \geq 2 \sin(\alpha/2),$$

qui est vérifiée pour tout $\alpha \geq 0$.

Lemme 4.1. Soit $f : M_1 \rightarrow M_2$ un difféomorphisme entre deux sous-variétés connexes de classe C^1 . Supposons que pour tout point $p \in M_1$, et tout vecteur tangent $\mathbf{v} \in T_p M_1$, on a $\|df_p(\mathbf{v})\| = \|\mathbf{v}\|$. Alors f est une isométrie entre M_1 et M_2 pour les distances intrinsèques d_{M_1} et d_{M_2} , c'est-à-dire

$$d_{M_2}(f(p), f(q)) = d_{M_1}(p, q)$$

pour tous $p, q \in M_1$.

Remarquons que l'hypothèse de cette proposition signifie que df_p est une isométrie linéaire entre les espaces tangents $T_p M_1$ et $T_{f(p)} M_2$. En particulier on a $\langle df_p(\mathbf{v}_1), df_p(\mathbf{v}_2) \rangle = \langle \mathbf{v}_1, \mathbf{v}_2 \rangle$ pour tous $\mathbf{v}_1, \mathbf{v}_2 \in T_p M_1$,

Preuve. Soit $\gamma : [a, b] \rightarrow M_1$ un chemin de classe C^1 par morceaux reliant p à q , alors $\tilde{\gamma} = f \circ \gamma : [a, b] \rightarrow M_2$ un chemin de classe C^1 par morceaux dans la variété M_2 qui relie $f(p)$ à $f(q)$. Nous avons alors $\dot{\tilde{\gamma}}(t) = df_{\gamma(t)}(\dot{\gamma}(t))$, et donc, par hypothèse

$$\|\dot{\tilde{\gamma}}(t)\| = \|df_{\gamma(t)}(\dot{\gamma}(t))\| = \|\dot{\gamma}(t)\|$$

pour tout $t \in [a, b]$. Par conséquent

$$\ell(\tilde{\gamma}) = \int_a^b \|\dot{\tilde{\gamma}}(t)\| dt = \int_a^b \|\dot{\gamma}(t)\| dt = \ell(\gamma).$$

Ceci implique que

$$d_{M_2}(f(p), f(q)) \leq \ell(\tilde{\gamma}) = \ell(\gamma).$$

En prenant l'infimum des chemins γ qui relient p à q on conclut que $d_{M_2}(f(p), f(q)) \leq d_{M_1}(p, q)$. Finalement, en remplaçant f par le difféomorphisme f^{-1} et en répétant le même argument, on obtient également l'inégalité $d_{M_1}(p, q) \leq d_{M_2}(f(p), f(q))$, ce qui prouve que f est une isométrie. $\square$

La réciproque du lemme précédent est également vraie, nous l'énonçons sous forme d'un théorème que nous admettons sans démonstration :

Théorème 4.2. *Une application $f : M_1 \rightarrow M_2$ entre deux sous-variétés connexes de classe C^2 est une isométrie entre M_1 et M_2 pour les distances intrinsèques d_{M_1} et d_{M_2} si et seulement si f est un difféomorphisme de classe C^1 tel que $\|df_p(\mathbf{v})\| = \|\mathbf{v}\|$ pour tout $p \in M_1$ et tout $\mathbf{v} \in T_p M_1$.*

Remarque. La partie difficile de ce théorème est de prouver qu'une application f qui préserve les distances (i.e. telle que $d_{M_2}(f(p), f(q)) = d_{M_1}(p, q)$ pour tous $p, q \in M_1$) est différentiable. Ce résultat a été démontré par S. B. Myers et N. E. Steenrod en 1939.

Exemple. Deux sous-variétés M_1 et M_2 de $\mathbb{R}^n$ sont dites *congruentes* s'il existe une isométrie globale $f : \mathbb{R}^n \rightarrow \mathbb{R}^n$ telle que $f(M_1) = M_2$. Lorsque c'est le cas, il est clair que la restriction $f|_{M_1}$ est une isométrie de M_1 vers M_2 pour la distance intrinsèque, et aussi pour la distance extrinsèque (ce qui n'est en général pas le cas pour les isométries intrinsèques).

4.2 Le tenseur métrique associé à une paramétrisation locale

On a vu que la distance intrinsèque entre deux points d'une sous-variété connexe est l'infimum des longueurs des courbes reliant ces deux points. Il sera donc important de pouvoir calculer la longueur d'une courbe lorsqu'elle est décrite dans une paramétrisation (locale) de la variété.

Rappelons qu'une *paramétrisation locale* d'une sous-variété $M \subset \mathbb{R}^n$ de classe C^k est la donnée :

- (i) d'un domaine $\Omega \subset \mathbb{R}^m$, où $m = \dim(M)$,
- (ii) et d'une application $\psi : \Omega \rightarrow M$, de classe C^k , qui est un difféomorphisme sur son image $\psi(\Omega) \subset M$. En particulier, ψ est une immersion de Ω dans $\mathbb{R}^n$.

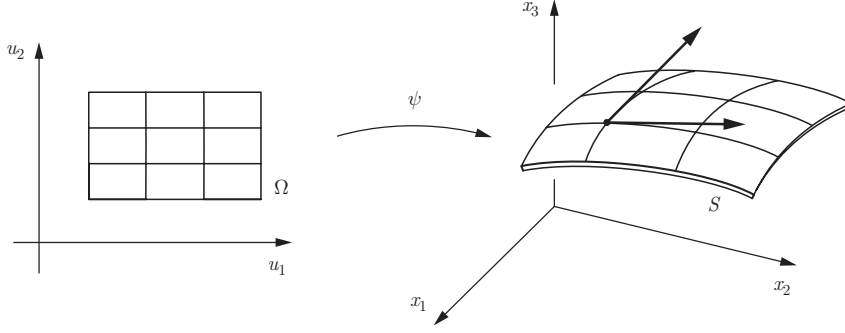
Par la condition (ii), on sait que ψ est en particulier une immersion. Si on note $(u_1, \dots, u_m)$ les coordonnées sur Ω , alors pour tout $u \in \Omega$, les vecteurs

$$\mathbf{b}_1(u) = \frac{\overrightarrow{\partial \psi}}{\partial u_1}(u), \dots, \mathbf{b}_m(u) = \frac{\overrightarrow{\partial \psi}}{\partial u_m}(u)$$

sont donc linéairement indépendants et forment ainsi une base de l'espace tangent $T_{\psi(u)}M$.

Définition. On dit que $\{\mathbf{b}_1(u), \dots, \mathbf{b}_m(u)\}$ est la *base adaptée* à la paramétrisation ψ de l'espace tangent $T_p S$, où $p = \psi(u)$.

Les paramètres $u_1, \dots, u_m$ s'appellent les *coordonnées curvilignes locales* associées à la paramétrisation locale ψ de M . Les courbes sur M obtenues en fixant toutes les coordonnées u_i sauf une s'appellent les *lignes de coordonnées* sur la sous-variété paramétrée M . Ensemble, elles forment le *réseau de coordonnées* associé à la paramétrisation ψ .



Définition. Les vecteurs $\{\mathbf{b}_1, \dots, \mathbf{b}_m\}$ ne forment en général pas une base orthonormée de l'espace tangent $T_p M$. La matrice de Gram $\mathbf{G}$ de cette base s'appelle le *tenseur métrique* de la sous-variété M dans la paramétrisation ψ .

Notons que, puisque les vecteurs $\mathbf{b}_i$ dépendent de $u \in \Omega$, le tenseur métrique $\mathbf{G} = (g_{ij}) : \Omega \rightarrow M_m(\mathbb{R})$ est une fonction à valeurs matricielle définie sur le domaine Ω . Ses coefficients sont

$$g_{ij}(u) = \langle \mathbf{b}_i(u), \mathbf{b}_j(u) \rangle = \sum_{k=1}^m \frac{\partial \psi_k}{\partial u_i} \frac{\partial \psi_k}{\partial u_j}. \quad (4.2)$$

Le tenseur métrique est donc une fonction $\mathbf{G} \in C^{k-1}(\Omega, M_m(\mathbb{R}))$ de classe C^{k-1} définie sur Ω à valeurs dans l'espace vectoriel des $m \times m$ matrices symétriques. Le tenseur métrique est associé à la paramétrisation locale $\psi : \Omega \rightarrow M$ et non pas uniquement à la sous-variété $M \subset \mathbb{R}^n$.

Le tenseur métrique s'appelle aussi la *première forme fondamentale* de associée à ψ . On verra plus loin qu'il y a aussi une deuxième et une troisième formes fondamentales.

Exemple 1 (graphe d'une fonction). Comme premier exemple, considérons le graphe S de la fonction $f : \Omega \rightarrow \mathbb{R}$, où Ω est un ouvert de $\mathbb{R}^2$. Une paramétrisation $\psi : \Omega \rightarrow S$ est donné par

$$\psi(x, y) = (x, y, f(x, y)).$$

La base du plan tangent en un point de la surface adaptée à cette paramétrisation est

$$\mathbf{b}_1 = \frac{\partial \psi}{\partial x} = \begin{pmatrix} 1 \\ 0 \\ f_x \end{pmatrix}, \quad \mathbf{b}_2 = \frac{\partial \psi}{\partial y} = \begin{pmatrix} 0 \\ 1 \\ f_y \end{pmatrix},$$

où on a noté $f_x = \frac{\partial f}{\partial x}$ et $f_y = \frac{\partial f}{\partial y}$. Les coefficients du tenseur métrique sont donc

$$g_{11} = \langle \mathbf{b}_1, \mathbf{b}_1 \rangle = 1 + f_x^2, \quad g_{12} = \langle \mathbf{b}_1, \mathbf{b}_2 \rangle = f_x f_y, \quad g_{22} = \langle \mathbf{b}_2, \mathbf{b}_2 \rangle = 1 + f_y^2,$$

c'est-à-dire

$$\mathbf{G}(x, y) = \begin{pmatrix} 1 + f_x^2 & f_x f_y \\ f_x f_y & 1 + f_y^2 \end{pmatrix}.$$

Exemple 2 (Surface de révolution)

Considérons une courbe régulière $\alpha(v) = (r(v), z(v))$ de classe C^1 , où $v \in I \subset \mathbb{R}$ dans un plan muni de coordonnées (r, z) et supposons que $r(v) > 0$ pour tout $v \in I$.

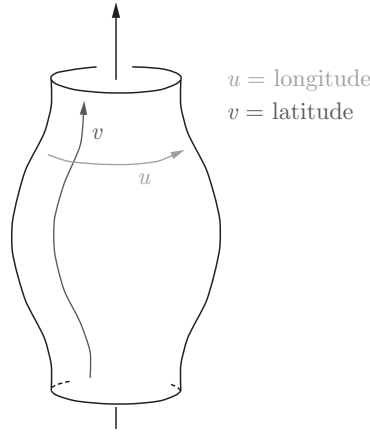
On appelle *surface de révolution de profil α autour de l'axe Oz* la surface $S \subset \mathbb{R}^3$ paramétrée par $\psi : \Omega \rightarrow S$, ($\Omega = [0, 2\pi] \times I$), où

$$\psi(u, v) = (x(u, v), y(u, v), z(u, v))$$

est donné par

$$\begin{cases} x(u, v) = r(v) \cos(u) \\ y(u, v) = r(v) \sin(u) \\ z(u, v) = z(v). \end{cases}$$

La coordonnée u s'appelle *longitude* et la coordonnée v s'appelle *latitude*. Les courbes $u = \text{const.}$ sont les *méridiens* et les courbes $v = \text{const.}$ sont les *parallèles* de S .



La base du plan tangent associée à cette paramétrisation est

$$\mathbf{b}_1 = \frac{\partial \psi}{\partial u} = \begin{pmatrix} -r(v) \sin(u) \\ r(v) \cos(u) \\ 0 \end{pmatrix}; \quad \mathbf{b}_2 = \frac{\partial \psi}{\partial v} = \begin{pmatrix} r'(v) \cos(u) \\ r'(v) \sin(u) \\ z'(v) \end{pmatrix}$$

Les coefficients du tenseur métrique sont

$$g_{11} = \left\| \frac{\partial \psi}{\partial u} \right\|^2 = r(v)^2, \quad g_{12} = \left\langle \frac{\partial \psi}{\partial u}, \frac{\partial \psi}{\partial v} \right\rangle = 0, \quad g_{22} = \left\| \frac{\partial \psi}{\partial v} \right\|^2 = r'(v)^2 + z'(v)^2.$$

On a donc

$$\mathbf{G}(u, v) = \begin{pmatrix} r(v)^2 & 0 \\ 0 & (r'(v)^2 + z'(v)^2) \end{pmatrix} = \begin{pmatrix} r(v)^2 & 0 \\ 0 & \|\dot{\alpha}(v)\|^2 \end{pmatrix}.$$

Remarquons que le réseau des coordonnées longitude-latitude est partout orthogonal puisque $g_{12} \equiv 0$.

Application à la sphère : La sphère S_a de rayon a centrée à l'origine est la surface d'équation $x^2 + y^2 + z^2 = a^2$. C'est la surface de révolution dont le profil est le demi-cercle

$$\gamma(v) = (r(v), z(v)) = (a \cos(v), a \sin(v)) \quad \left(-\frac{\pi}{2} \leq v \leq \frac{\pi}{2}\right).$$

La paramétrisation de la sphère est donc donné par

$$\begin{cases} x = a \cos(u) \cos(v) \\ y = a \sin(u) \cos(v) \\ z = a \sin(v) \end{cases}$$

où (u, v) parcourt le domaine Ω défini par les inégalités $0 \leq u \leq 2\pi$, $-\frac{\pi}{2} \leq v \leq \frac{\pi}{2}$ (les paramètres utilisés dans cet exemple, s'appellent les *coordonnées géographiques* sur la sphère). Les formules précédentes nous donnent le tenseur métrique suivant :

$$g_{11} = a^2 \cos^2(v), \quad g_{12} = 0, \quad g_{22} = a^2 \quad (4.3)$$

c'est-à-dire

$$\mathbf{G}(u, v) = \begin{pmatrix} a^2 \cos^2(v) & 0 \\ 0 & a^2 \end{pmatrix}.$$

Exemple 3 (Surface réglée). Une surface est dite *réglée* si elle est une réunion de droites ou de segments de droites, ces droites sont appelées les *génératrices*. Le plan, le cylindre et le cône sont les exemples les plus simples de surfaces réglées.

Pour paramétrer une surface réglée, on se donne une courbe $\alpha : I \rightarrow \mathbb{R}^3$, qu'on appellera une *directrice* de la surface réglée, et qu'on suppose transverse aux génératrices, ainsi qu'un champ de vecteurs $\mathbf{w}(u)$ le long de α , ce champ indique la direction des génératrices. La surface est alors paramétrée par

$$\psi(u, v) = \alpha(u) + v \mathbf{w}(u)$$

où $(u, v) \in \Omega := I \times \mathbb{R}$.

On a alors

$$\mathbf{b}_1 = \frac{\partial \psi}{\partial u} = \dot{\alpha}(u) + v \dot{\mathbf{w}}(u), \quad \mathbf{b}_2 = \frac{\partial \psi}{\partial v} = \mathbf{w}(u).$$

et le tenseur métrique est donné par

$$\begin{cases} g_{11} = \|\dot{\alpha}\|^2 + 2 \langle \dot{\alpha}, \dot{\mathbf{w}} \rangle v + v^2 \|\dot{\mathbf{w}}\|^2 \\ g_{12} = \langle \dot{\alpha}, \mathbf{w} \rangle + v \langle \dot{\mathbf{w}}, \mathbf{w} \rangle \\ g_{22} = \|\mathbf{w}\|^2. \end{cases}$$

Voyons deux cas particuliers de surface réglée où ce tenseur métrique prend une forme simple. Supposons d'abord que α est une courbe plane paramétrée naturellement et que le vecteur $\mathbf{w}$ est constant, unitaire et orthogonal au plan contenant α . On dit alors que S est un *cylindre généralisé*.

On a $\|\dot{\alpha}\| = \|\mathbf{w}\| = 1$, $\dot{\mathbf{w}} = 0$ et $\langle \dot{\alpha}, \mathbf{w} \rangle = 0$. Par conséquent $g_{11} = g_{22} = 1$ et $g_{12} = 0$ et le tenseur métrique est alors $\mathbf{G} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$.

L'autre cas particulier est la *surface des tangentes* à une courbe birégulière $\alpha : I \rightarrow \mathbb{R}^3$ que l'on suppose paramétrée naturellement. C'est la surface réglée obtenue en prenant le champ de vecteur $\mathbf{u} = \mathbf{T}(u) = \dot{\alpha}(u)$, on a donc

$$\begin{cases} g_{11} = \|\dot{\alpha}\|^2 + 2 \langle \dot{\alpha}, \dot{\mathbf{T}} \rangle v + v^2 \|\dot{\mathbf{T}}\|^2 = 1 + (\kappa(u)v)^2 \\ g_{12} = \langle \dot{\alpha}, \mathbf{T} \rangle + v \langle \dot{\mathbf{T}}, \mathbf{T} \rangle = 1 \\ g_{22} = \|\mathbf{T}\|^2 = 1. \end{cases}$$

où $\kappa(u)$ est la courbure de α . Le tenseur métrique est alors

$$\mathbf{G}(s, v) = \begin{pmatrix} 1 + (\kappa(u)v)^2 & 1 \\ 1 & 1 \end{pmatrix}.$$

Remarque. Il est utile de remarquer que les vecteurs $\mathbf{b}_j(u) = \frac{\partial \psi}{\partial u_j}(u)$ forment les colonnes de la matrice jacobienne $D\psi(u) = \left(\frac{\partial \psi_i}{\partial u_j}(u) \right)$. Par conséquent on a

$$\mathbf{G}(u) = D\psi(u)^\top \cdot D\psi(u). \quad (4.4)$$

Cette identité peut nous donner une façon rapide de calculer un tenseur métrique.

Exemple. Reprenons l'exemple du tenseur métrique d'un graphe M , mais cette fois dans le cas d'une fonction de m variables $f \in C^k(\Omega, \mathbb{R})$, où Ω est un ouvert de $\mathbb{R}^m$. La variété M est donc paramétrée par l'application $\psi : \Omega \rightarrow \mathbb{R}^{n+1}$ définie par $\psi(x_1, \dots, x_m) = (x_1, \dots, x_m, f(x_1, \dots, x_m))$. La matrice jacobienne de ψ en un point de Ω est

$$D\psi = \begin{pmatrix} 1 & 0 & \cdots & 0 \\ 0 & 1 & \cdots & \vdots \\ \vdots & \vdots & \cdots & 1 \\ f_{x_1} & f_{x_2} & \cdots & f_{x_m} \end{pmatrix},$$

où $f_{x_i} = \frac{\partial f}{\partial x_i}$. Le tenseur métrique est donc

$$\mathbf{G} = D\psi^\top \cdot D\psi = \begin{pmatrix} 1 + f_{x_1}^2 & f_{x_1}f_{x_2} & \cdots & f_{x_1}f_{x_m} \\ f_{x_2}f_{x_1} & 1 + f_{x_2}^2 & \cdots & f_{x_2}f_{x_m} \\ \vdots & \vdots & \cdots & \vdots \\ f_{x_m}f_{x_1} & f_{x_m}f_{x_2} & \cdots & 1 + f_{x_m}^2 \end{pmatrix},$$

que l'on peut aussi écrire $g_{ij} = \delta_{ij} + f_{x_i}f_{x_j}$.

4.3 Signification géométrique du tenseur métrique

Le rôle du tenseur métrique est de nous permettre de calculer la norme d'un vecteur ainsi que le produit scalaire entre deux vecteurs tangents en un point d'une sous-variété différentiable $M \subset \mathbb{R}^n$ lorsque ces vecteurs sont exprimés dans la base adaptée à une paramétrisation locale $\psi : \Omega \rightarrow M$.

En effet, si $\xi, \eta \in T_p M$ sont deux vecteurs tangents en un point $p = \psi(u) \in M$, alors on peut écrire $\xi = \sum_{i=1}^m \xi_i \mathbf{b}_i$ et $\eta = \sum_{j=1}^m \eta_j \mathbf{b}_j$, où $\{\mathbf{b}_1, \dots, \mathbf{b}_m\} \subset T_p M$ est la base de $T_p M$ adaptée à la paramétrisation ψ . On a donc

$$\langle \xi, \eta \rangle = \sum_{i,j=1}^m \xi_i \eta_j \langle \mathbf{b}_i, \mathbf{b}_j \rangle = \sum_{i,j=1}^m g_{ij}(u) \xi_i \eta_j,$$

où les g_{ij} sont les coefficients du tenseur métrique. On notera souvent ce produit scalaire sous la forme

$$\mathbf{g}_u(\xi, \eta) = \sum_{i,j=1}^m g_{ij}(u) \xi_i \eta_j.$$

En particulier la norme du vecteur $\xi \in T_p M$ et l'angle θ entre les vecteurs $\xi, \eta \in T_p M$ (supposés non nuls) sont donnés par

$$\|\xi\| = \sqrt{\mathbf{g}_u(\xi, \xi)} = \sqrt{\sum g_{ij} \xi_i \xi_j},$$

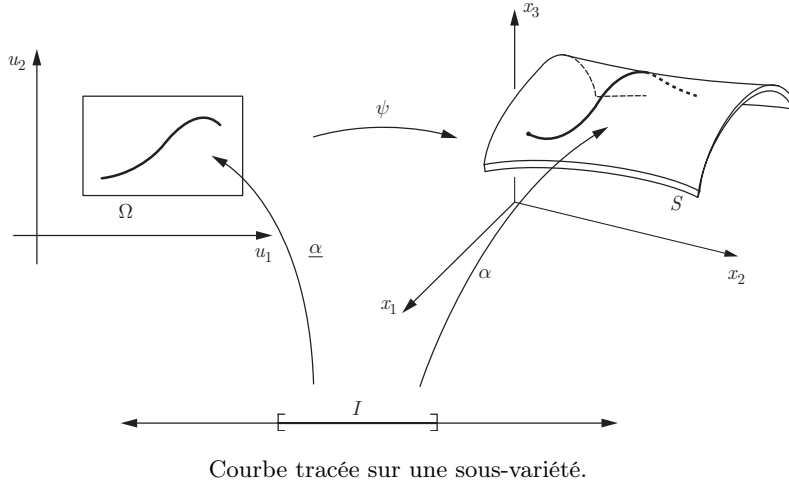
et

$$\cos(\theta) = \frac{\mathbf{g}_u(\xi, \eta)}{\sqrt{\mathbf{g}_u(\xi, \xi)} \sqrt{\mathbf{g}_u(\eta, \eta)}} = \frac{\sum g_{ij} \xi_i \eta_j}{\sqrt{\sum g_{ij} \xi_i \xi_j} \cdot \sqrt{\sum g_{ij} \eta_i \eta_j}}.$$

Ces formules nous permettent de calculer la longueur d'une courbe $\alpha : [a, b] \rightarrow M$ de classe C^1 tracée dans l'image d'une paramétrisation $\psi : \Omega \rightarrow M$. La représentation de la courbe α dans la carte Ω s'écrit

$$\underline{\alpha}(t) = \psi^{-1}(\alpha(t)) = (u_1(t), \dots, u_m(t)) \in \Omega, \quad t \in [a, b]$$

(i.e. nous avons une courbe auxiliaire $\underline{\alpha}(t) \in \Omega$ dans le domaine de la paramétrisation Ω telle que $\psi \circ \underline{\alpha}(t) = \alpha(t)$ pour tout $t \in I$).



Le vecteur vitesse de α est donné par

$$\dot{\alpha}(t) = \frac{d}{dt} \psi(\underline{\alpha}(t)) = \sum_{i=1}^m \frac{\partial \psi}{\partial u_i} \frac{du_i}{dt} = \sum_{i=1}^m \dot{u}_i(t) \mathbf{b}_i(u),$$

et la longueur de α est finalement donnée par

$$\ell(\alpha) = \int_a^b \|\dot{\alpha}(t)\| dt = \int_a^b \sqrt{\sum g_{ij}(u(t)) \dot{u}_i(t) \dot{u}_j(t)} dt.$$

Rappelons que l'abscisse curviligne le long de la courbe α est la longueur de l'arc $\alpha|_{[a,s]}$. En particulier on a $\frac{ds}{dt} = \|\dot{\alpha}(t)\|$. On peut donc écrire

$$\left(\frac{ds}{dt}\right)^2 = \sum_{i,j=1}^m g_{ij} \frac{du_i}{dt} \frac{du_j}{dt}.$$

En multipliant formellement cette égalité par dt^2 , on obtient

$$ds^2 = \sum_{i,j=1}^m g_{ij} du_i du_j.$$

Cette identité exprime le *carré de l'élément de longueur infinitésimal* d'une courbe α sur la sous-variété dans les coordonnées u_i associée à une carte. On voit que cette expression (qu'on appelle simplement "*le ds^2* ") contient la même information que le tenseur métrique.

L'étude du ds^2 nous donne une troisième façon de calculer le tenseur métrique qui est très efficace dans certains cas. Considérons une courbe $\alpha(t) = (x_1(t), \dots, x_n(t))$ sur la variété $M = \psi(\Omega)$. Sa représentation dans la carte est $\psi^{-1}(\alpha(t)) = (u_1(t), \dots, u_m(t))$, et on a le long de cette courbe

$$ds^2 = \sum_{i=1}^n dx_i^2 = \sum_{i,j=1}^m g_{ij} du_i du_j.$$

Il suffit alors de calculer $dx_i = \sum_{j=1}^m \frac{\partial x_i}{\partial u_j} du_j$ pour trouver les coefficients g_{ij} du tenseur métrique.

Exemple a. Les coordonnées polaires dans le plan sont données par les formules $x = r \cos(\theta)$, $y = r \sin(\theta)$. On a donc

$$ds^2 = dx^2 + dy^2 = (\cos(\theta)dr - r \sin(\theta)d\theta)^2 + (\sin(\theta)dr + r \cos(\theta)d\theta)^2 = dr^2 + r^2 d\theta^2.$$

Le tenseur métrique associé est donc donné par $\mathbf{G}(r, \theta) = \begin{pmatrix} 1 & 0 \\ 0 & r^2 \end{pmatrix}$.

Exemple b. La surface de révolution autour de l'axe Oz dont le profil est la courbe $\gamma(v) = (r(v), z(v))$ admet la paramétrisation

$$(x, y, z) = (r(v) \cos(u), r(v) \sin(u), z(v)).$$

On a donc

$$\begin{aligned} ds^2 &= dx^2 + dy^2 + dz^2 \\ &= (r'(v) \cos(u)dv - r(v) \sin(u)du)^2 + (r'(v) \sin(u)dv + r(v) \cos(u)du)^2 + (z'(v))^2 dv^2 \\ &= r^2(v) du^2 + (r'(v)^2 + z'(v)^2) dv^2. \end{aligned}$$

Le tenseur métrique associé est donc donné par

$$\mathbf{G}(u, v) = \begin{pmatrix} r^2(v) & 0 \\ 0 & (r'(v)^2 + z'(v)^2) \end{pmatrix} = \begin{pmatrix} r^2(v) & 0 \\ 0 & \|\gamma'(v)\|^2 \end{pmatrix}.$$

4.4 Sur les isométries entre sous-variétés paramétrées

Le résultat suivant nous dit comment se comparent les tenseurs métriques de deux sous-variétés paramétrées qui sont isométriques :

Théorème 4.3. Soient $\psi_1 : \Omega_1 \rightarrow M_1$ et $\psi_2 : \Omega_2 \rightarrow M_2$ deux sous-variétés paramétrées de classe C^1 . Alors il existe une isométrie intrinsèque $f : M_1 \rightarrow M_2$ si et seulement s'il existe un difféomorphisme $h : \Omega_1 \rightarrow \Omega_2$ tel que pour tout $u \in \Omega_1$ on a

$$\mathbf{G}_1(u) = Dh(u)^\top \mathbf{G}_2(h(u)) Dh(u), \quad (4.5)$$

où $\mathbf{G}_1$ est le tenseur métrique de ψ_1 , $\mathbf{G}_2$ est le tenseur métrique de ψ_2 et Dh est la matrice jacobienne du difféomorphisme h .

La situation est représentée par le diagramme commutatif suivant :

$$\begin{array}{ccc} M_1 & \xrightarrow{f} & M_2 \\ \psi_1 \uparrow & & \uparrow \psi_2 \\ \Omega_1 & \xrightarrow{h} & \Omega_2 \end{array}$$

Remarquons aussi que si on note $h(u_1, \dots, u_m) = (v_1, \dots, v_m)$, alors la formule (4.5) peut s'écrire

$$g_{ij}(u) = \sum_{\mu, \nu=1}^m \tilde{g}_{\mu\nu}(v) \frac{\partial v_\mu}{\partial u_i} \frac{\partial v_\nu}{\partial u_j}, \quad (4.6)$$

où $\mathbf{G}_1 = (g_{ij})$ et $\mathbf{G}_2 = (\tilde{g}_{ij})$ et $m = \dim(M_1) = \dim(M_2)$. Cette formule peut aussi s'écrire sous forme différentielle :

$$ds^2 = \sum_{i,j=1}^m g_{ij}(u) du_i du_j = \sum_{\mu, \nu=1}^m \tilde{g}_{\mu\nu}(v) dv_\mu dv_\nu. \quad (4.7)$$

Preuve. Supposons qu'il existe une isométrie intrinsèque $f : M_1 \rightarrow M_2$, et remarquons que $f \circ \psi_1 : \Omega_1 \rightarrow M_2$ est une paramétrisation de M_2 (en général différente de ψ_2). Par le théorème 4.2, on sait que f est une isométrie si et seulement si

$$\langle df_p(\xi), df_p(\eta) \rangle = \langle \xi, \eta \rangle$$

pour tout $p \in M_1$ et $\xi, \eta \in T_p M_1$. Ceci implique que pour tout $u \in \Omega_1$, les coefficients du tenseur métrique de ψ_1 vérifient :

$$\begin{aligned} g_{ij}(u) &= \langle d\psi_{1_u}(\mathbf{e}_i), d\psi_{1_u}(\mathbf{e}_j) \rangle \\ &= \langle df_{\psi_1(u)}(d\psi_{1_u}(\mathbf{e}_i)), df_{\psi_1(u)}(d\psi_{1_u}(\mathbf{e}_j)) \rangle \\ &= \langle d(f \circ \psi_1)_u(\mathbf{e}_i), d(f \circ \psi_1)_u(\mathbf{e}_j) \rangle \end{aligned}$$

où $\{\mathbf{e}_1, \dots, \mathbf{e}_m\}$ est la base canonique de $\mathbb{R}^m$. Cette condition s'écrit matriciellement

$$\mathbf{G}_1(u) = D\psi_1(u)^\top \cdot D\psi_1(u) = D(f \circ \psi_1(u))^\top \cdot D(f \circ \psi_1(u)).$$

On considère maintenant l'application $h : \Omega_1 \rightarrow \Omega_2$ définie par $h = \psi_2^{-1} \circ f \circ \psi_1$. Cette application est un difféomorphisme de Ω_1 vers Ω_2 car c'est la composition de trois difféomorphismes, de plus on a $\psi_2 \circ h = f \circ \psi_1$, par conséquent

$$\begin{aligned} \mathbf{G}_1(u) &= D(\psi_2 \circ h(u))^\top \cdot D(\psi_2 \circ h(u)) \\ &= Dh(u)^\top D\psi_2(h(u))^\top D\psi_2(h(u)) Dh(u) \\ &= Dh(u)^\top \mathbf{G}_2(h(u)) Dh(u). \end{aligned}$$

Inversément, supposons qu'il existe un difféomorphisme $h : \Omega_1 \rightarrow \Omega_2$ tel que la condition (4.5) est vérifiée, alors on définit $f : M_1 \rightarrow M_2$ par $f = \psi_2 \circ h \circ \psi_1^{-1} : M_1 \rightarrow M_2$. Le calcul précédent prouve que f est une isométrie de M_1 vers M_2 . □

Le théorème 4.3 contient les cas particuliers suivants :

Corollaire 4.4. *Si les deux sous-variétés paramétrées de classe C^1 $\psi_1 : \Omega \rightarrow M_1$ et $\psi_2 : \Omega \rightarrow M_2$, avec même domaine de paramétrisation, $\Omega \subset \mathbb{R}^m$ ont le même tenseur métrique, alors elles sont isométriques*

Preuve. Ce corollaire correspond au cas où $h : \Omega \rightarrow \Omega$ est l'identité. □

Corollaire 4.5. Si $\psi_1 : \Omega_1 \rightarrow M$ et $\psi_2 : \Omega_2 \rightarrow M$ sont deux paramétrisations de la même sous-variété M , alors il existe un difféomorphisme $h : \Omega_1 \rightarrow \Omega_2$ tel que (4.5) soit satisfaite pour tout $u \in \Omega_1$

Preuve. Correspond au cas où f est l'identité. □

4.5 Intégration sur une sous-variété

Définition. Si $M \subset \mathbb{R}^n$ est une sous-variété de dimension m et $\psi : \Omega \rightarrow M$ est une paramétrisation régulière globale de classe C^1 , alors l'intégrale

$$\text{Vol}_m(M) = \int_{\Omega} \sqrt{\det(\mathbf{G}(u))} du_1 \cdots du_m. \quad (4.8)$$

s'appelle le *volume m -dimensionnel de la variété M* .

Nous ne pouvons pas prouver ici cette formule, que nous prenons donc comme une définition. Toutefois elle peut se justifier heuristiquement de la façon suivante : Considérons une sous-variété paramétrée $\psi : \Omega \rightarrow M$ de classe C^1 et de dimension m . Pour estimer son volume, on peut subdiviser le domaine $\Omega \subset \mathbb{R}^m$ en sous-domaines Ω_i :

$$\Omega = \cup_i \Omega_i, \quad \text{tels que si } i \neq j, \text{ alors } \Omega_i \cap \Omega_j \text{ est de mesure nulle,}$$

en sorte que

$$\text{Vol}_m(M) = \text{Vol}_m(\psi(\Omega)) = \sum_i \text{Vol}_m(\psi(\Omega_i)).$$

Si les sous-domaines Ω_i sont suffisamment petits, on peut approximer la restriction de ψ à Ω_i par sa différentielle, ainsi

$$\text{Vol}_m(\psi(\Omega_i)) \cong \text{Vol}_m(d\psi_{u_i}(\Omega_i)) = \sqrt{\det(d\psi_{u_i}^\top d\psi_{u_i})} \text{Vol}_m(\Omega_i) = \sqrt{\det \mathbf{G}(u)} \text{Vol}_m(\Omega_i),$$

où $u_i \in \Omega_i$ est arbitraire. On a donc

$$\text{Vol}_m(M) \cong \sum_i \sqrt{\det \mathbf{G}(u)} \text{Vol}_m(\Omega_i).$$

En raffinant la subdivision de Ω , et en supposant que $\max\{\text{diam}(\Omega_i)\} \rightarrow 0$, cette somme converge vers l'intégrale (4.8).

Remarque. Lorsque M est une surface, i.e. $m = 2$, on note $\text{Vol}_2(M) = \text{Aire}(M)$ et on dit que c'est l'aire de M . Lorsque $m = 1$, $\text{Vol}_1(M)$ n'est rien d'autre que la longueur de la courbe M .

Le résultat suivant nous dit que le volume est une notion géométrique, c'est-à-dire indépendante de la paramétrisation choisie.

Proposition 4.6. Soient $\psi_1 : \Omega_1 \rightarrow M$ et $\psi_2 : \Omega_2 \rightarrow M$ sont deux paramétrisations régulières globales de classe C^1 d'une même sous-variété $M \subset \mathbb{R}^n$ de dimension m . Alors on a

$$\int_{\Omega_1} \sqrt{\det(\mathbf{G}_1(u))} du = \int_{\Omega_2} \sqrt{\det(\mathbf{G}_2(v))} dv.$$

Ce résultat se déduit du corollaire 4.5 en appliquant la formule de changement de variables dans les intégrales multiples.

Exemple 1. Si S est le graphe de la fonction $f : \Omega \rightarrow \mathbb{R}$, alors

$$\text{Aire}(S) = \iint_{\Omega} \sqrt{1 + f_x^2 + f_y^2} \, dx dy.$$

Exemple 2. La paramétrisation par longitude et latitude de la sphère $S_a \subset \mathbb{R}^3$ de rayon a admet $\Omega = [0, 2\pi] \times [-\frac{\pi}{2}, \frac{\pi}{2}]$ comme domaine de paramétrisation, et on a

$$G(u, v) = \begin{pmatrix} a^2 \cos(v)^2 & 0 \\ 0 & a^2 \end{pmatrix},$$

donc $dA = a^2 \cos(v) du dv$ et l'aire de cette sphère est

$$\text{Aire}(S_a) = \iint_{S_a} dA = \iint_{\Omega} \sqrt{\det G(u, v)} du dv = \int_{v=-\frac{\pi}{2}}^{\frac{\pi}{2}} \int_{u=0}^{2\pi} a^2 \cos(v) du dv = 4\pi a^2.$$

Généralisation. Plus généralement, si $\psi : \Omega \rightarrow M \subset \mathbb{R}^n$ est une paramétrisation globale de classe C^1 d'une sous-variété M et $\rho : M \rightarrow \mathbb{R}$ est une fonction continue non négative, alors l'intégrale de la fonction ρ sur M est définie par

$$\int_M \rho(x) dV := \int_{\Omega} (\rho \circ \psi)(u) \sqrt{\det G(u)} du.$$

Nous pouvons aussi considérer le cas des fonctions à valeurs vectorielles. Par exemple, le *centre de gravité* de la variété M (pour une distribution de masse homogène) est le point $C \in \mathbb{R}^n$ défini par

$$C = \frac{1}{\text{Vol}(M)} \int_S \mathbf{x} dV = \frac{1}{\text{Vol}(M)} \int_{\Omega} \psi(u) \sqrt{\det \mathbf{G}(u)} du. \quad (4.9)$$

4.6 Domaines riemanniens

Définition. On appelle *métrique riemannienne* de classe C^k sur un domaine $\Omega \subset \mathbb{R}^m$ la donnée en chaque point $u \in \Omega$, d'un produit scalaire

$$\mathbf{g}_u : \mathbb{R}^m \times \mathbb{R}^m \rightarrow \mathbb{R},$$

qui varie de façon différentiable par rapport à u . Cette condition signifie que la fonctions $g_{ij} : \Omega \rightarrow \mathbb{R}$ définies par $g_{ij}(u) = \mathbf{g}_u(\mathbf{e}_i, \mathbf{e}_j)$ sont de classe C^k , où $\{\mathbf{e}_1, \dots, \mathbf{e}_m\}$ est la base canonique de $\mathbb{R}^m$. On note parfois

$$\mathbf{g}_u = \sum_{i,j=1}^m g_{ij}(u) du_i du_j,$$

et on dit que cette expression est le *tenseur métrique* sur le domaine Ω . Un *domaine riemannien* est une couple $(\Omega, \mathbf{g})$, où Ω est un domaine de $\mathbb{R}^m$ et $\mathbf{g}$ est une métrique riemannienne définie sur ce domaine.

Cette structure permet de faire de la géométrie (dite *géométrie intrinsèque* ou *géométrie riemannienne*) dans le domaine Ω indépendamment d'une éventuelle réalisation de ce domaine comme plongement dans un espace euclidien. En particulier

- La *norme riemannienne* d'un vecteur ξ au point u est

$$\|\xi\|_u = \sqrt{\mathbf{g}_u(\xi, \xi)} = \sqrt{\sum_{i,j=1}^m g_{ij}(u) \xi_i \xi_j}.$$

- L'*angle* en $u \in \Omega$ entre les vecteurs non nuls ξ et η est

$$\cos(\angle_u(\xi, \eta)) = \frac{\mathbf{g}_u(\xi, \eta)}{\|\xi\|_{\mathbf{g}} \|\eta\|_{\mathbf{g}}}$$

- La *longueur riemannienne* de la courbe $\gamma : [a, b] \rightarrow \Omega$ de classe C^1 par morceaux est définie par

$$\ell_{\mathbf{g}}(\gamma) = \int_a^b \|\dot{\gamma}(t)\|_{\gamma(t)} dt = \int_a^b \sqrt{\mathbf{g}_{\gamma(t)}(\dot{\gamma}(t), \dot{\gamma}(t))} dt.$$

- La *distance intrinsèque* entre deux points est définie comme l'infimum des longueurs des courbes de classe C^1 par morceaux contenues dans le domaine Ω qui rejoignent ces deux points.
- Le *volume* du domaine Riemannien $(\Omega, \mathbf{g})$ est l'intégrale

$$\text{Vol}(\Omega, \mathbf{g}) = \int_{\Omega} \sqrt{\det(g_{ij}(u))} du_1 \cdots du_m.$$

Voyons quelques exemples :

- 1) Si $\psi : \Omega \rightarrow \mathbb{R}^n$ est une immersion, alors le tenseur métrique défini par

$$\mathbf{g}_u(\xi, \eta) = \langle d\psi_u(\xi), d\psi_u(\eta) \rangle$$

définit une structure Riemannienne sur Ω pour laquelle on a $g_{ij} = \left\langle \frac{\partial \psi}{\partial u_i}, \frac{\partial \psi}{\partial u_j} \right\rangle$.

- 2) La *métrique hyperbolique* de Poincaré sur le demi-espace $\mathbb{H}^m = \{x \in \mathbb{R}^m \mid x_m > 0\}$ est la métrique riemannienne définie par

$$\mathbf{h}_x(\xi, \eta) = \frac{\langle \xi, \eta \rangle_{\mathbb{R}^m}}{x_m^2}.$$

Pour cette métrique on a $h_{ij}(x) = \frac{1}{x_m^2} \delta_{ij}$.

- 3) La métrique hyperbolique de Poincaré dans la boule $\mathbb{B}^m = \{x \in \mathbb{R}^m \mid \|x\| < 1\}$ est la métrique riemannienne définie par

$$\mathbf{g}_x(\xi, \eta) = \frac{4\langle \xi, \eta \rangle_{\mathbb{R}^m}}{(1 - \|x\|^2)^2}.$$

Pour cette métrique on a $h_{ij}(x) = \frac{\delta_{ij}}{(1 - \|x\|^2)^2}$.

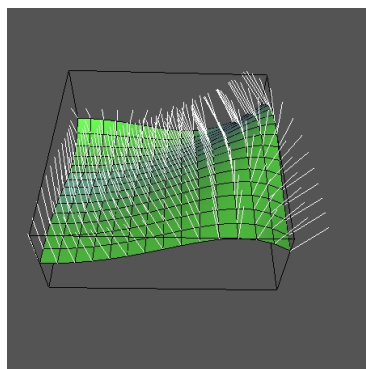
On peut démontrer que les domaines $\mathbb{H}^m$ et $\mathbb{B}^m$ sont isométriques pour leur métriques hyperboliques respectives.

Chapitre 5

Les surfaces et leur courbure

5.1 Co-orientation d'une surface et application de Gauss

Définition. On appelle *co-orientation* d'une surface régulière $S \subset \mathbb{R}^3$ de classe C^1 la donnée d'un champ de vecteurs continu $\nu : S \rightarrow \mathbb{R}^3$ tel que $\|\nu(p)\| = 1$ et $\nu(p) \perp T_p S$ pour tout point $p \in S$. La surface S est *co-orientable* si elle admet une co-orientation.



Remarques. (i) Si elle existe, une co-orientation ν de la surface S est unique au signe près. De plus le champ ν est de classe C^{k-1} si la surface S est de classe C^k .

(ii) Toute surface de classe C^1 est localement co-orientable, i.e. elle admet une co-orientation au voisinage de chacun de ses points.

(iii) Un exemple de surface qui n'est pas co-orientable globalement est le ruban de Möbius. On peut d'ailleurs prouver que toute surface qui n'est pas co-orientable contient un ouvert qui est homéomorphe au ruban de Möbius.

(iv) Le choix d'une co-orientation d'une surface régulière S permet de définir une orientation du plan tangent $T_p S$ pour tout $p \in S$ qui dépend continûment du point. On dit alors que la surface est *orientée* (les deux termes sont donc essentiellement synonymes).

Une co-orientation est concrètement obtenue de la façon suivante : Si la surface S est définie par l'équation $f(x, y, z) = 0$, alors une co-orientation est donnée par le champ

$$\nu(x) = \frac{\nabla f(x)}{\|\nabla f(x)\|},$$

en supposant que le gradient de f ne s'annule pas sur S . Si la surface est paramétrée de façon régulière par l'application injective $\psi : \Omega \rightarrow S \subset \mathbb{R}^3$, alors une co-orientation est donnée par

$$\nu = \frac{\frac{\partial \psi}{\partial u} \times \frac{\partial \psi}{\partial v}}{\left\| \frac{\partial \psi}{\partial u} \times \frac{\partial \psi}{\partial v} \right\|}.$$

Mentionnons pour finir, que l'application $\boldsymbol{\nu}$ est souvent vue non comme un champ de vecteurs mais comme une application de la surface S vers la sphère unité. Dans ce cas, l'application

$$\boldsymbol{\nu} : S \rightarrow \mathbb{S}^2$$

s'appelle l'*application de Gauss*.

5.2 Courbure d'une courbe tracée sur une surface

5.2.1 Géodésiques

Définition 5.1. Une courbe $\gamma : I \rightarrow S$ de classe C^2 tracée sur une surface régulière $S \subset \mathbb{R}^3$ est une *géodésique* de cette surface si son accélération est normale à la surface :

$$\ddot{\gamma}(t) \perp T_{\gamma(t)}S, \quad \text{pour tout } t \in I.$$

Si la surface est co-orientée par le champ $\boldsymbol{\nu}$, alors $\gamma : I \rightarrow S$ est géodésique si et seulement si elle vérifie l'équation différentielle suivante sur l'intervalle I :

$$\ddot{\gamma}(t) \times \boldsymbol{\nu}(\gamma(t)) = 0.$$

Lemme 5.2. *Toute géodésique sur une surface régulière est parcourue à vitesse constante.*

Nous laissons la preuve en exercice.

Exemples. 1) Les géodésiques d'un plan sont les droites de ce plan paramétrées affinement (i.e. parcourues à vitesse constante).

2) Les géodésiques d'une sphère sont les grands cercles de cette sphère, paramétrés à vitesse constante.

5.2.2 Repère de Darboux, courbures normale et géodésique

Définition 5.3. Soit $\gamma : I \rightarrow S$ une courbe régulière de classe C^2 tracée sur une surface régulière co-orientée $S \subset \mathbb{R}^3$.

- (i) On appelle *repère de Darboux*¹ le long de γ relatif à la surface S le repère mobile orthonormé $\{\boldsymbol{\nu}(t), \mathbf{T}_\gamma(t), \boldsymbol{\mu}(t)\}$ où $\mathbf{T}_\gamma(t) = \frac{1}{V_\gamma(t)}\dot{\gamma}(t)$ est le vecteur tangent unitaire à γ , $\boldsymbol{\nu}(t)$ est la co-orientation de S évaluée au point $\gamma(t) \in S$ et $\boldsymbol{\mu}(t) = \boldsymbol{\nu}(t) \times \mathbf{T}_\gamma(t)$.
- (ii) La *courbure normale* et la *courbure géodésique* de γ sont les fonctions du paramètre t définies respectivement par

$$k_n(t) = \langle \mathbf{K}_\gamma(t), \boldsymbol{\nu}(t) \rangle \quad \text{et} \quad k_g(t) = \langle \mathbf{K}_\gamma(t), \boldsymbol{\mu}(t) \rangle.$$

où $\mathbf{K}_\gamma(t)$ est le vecteur de courbure de γ . Ces courbures représentent les composantes normale et tangentielle de la courbure de γ .

- (iii) La *torsion géodésique* de γ par

$$\tau_g(t) = \frac{1}{V_\gamma(t)} \langle \dot{\boldsymbol{\nu}}(t), \boldsymbol{\mu}(t) \rangle.$$

Remarquons qu'en tout point $p = \gamma(t)$ de la courbe, les vecteurs $\{\mathbf{T}_\gamma(t), \boldsymbol{\mu}(t)\}$ forment une base orthonormée du plan tangent $T_p S$. La courbe γ est géodésique si et seulement si cette courbe est paramétrée à vitesse constante et sa courbure géodésique est nulle. Il est par ailleurs clair que

$$\mathbf{K}_\gamma(t) = k_n(t)\boldsymbol{\nu}(t) + k_g(t)\boldsymbol{\mu}(t) \quad \text{et} \quad k_n(t)^2 + k_g(t)^2 = \kappa(t)^2 = \|\mathbf{K}_\gamma(t)\|^2.$$

1. Attention, il n'y a pas de lien entre cette notion et le vecteur de Darboux défini au chapitre 2.

Proposition 5.4 (Équations de Darboux). *Le repère de Darboux vérifie les équations différentielles suivantes :*

$$\begin{cases} \frac{1}{V_\gamma} \dot{\mathbf{T}} &= k_g \boldsymbol{\mu} + k_n \boldsymbol{\nu}, \\ \frac{1}{V_\gamma} \dot{\boldsymbol{\nu}} &= -k_n \mathbf{T} + \tau_g \boldsymbol{\mu}, \\ \frac{1}{V_\gamma} \dot{\boldsymbol{\mu}} &= -k_g \mathbf{T} - \tau_g \boldsymbol{\nu}. \end{cases}$$

Nous laissons la preuve de cette proposition en exercice.

La courbure géodésique apparaît naturellement lorsqu'on dérive la fonctionnelle de longueur d'une courbe. De façon plus précise, considérons une courbe $\gamma : [a, b] \rightarrow S$ de classe C^2 sur une surface régulière S , que l'on suppose paramétrée naturellement : $\|\dot{\gamma}(u)\| \equiv 1$ (on notera ici u le paramètre de γ). Une *déformation* de γ sur S est la donnée d'une application $\psi : [a, b] \times (-\varepsilon, \varepsilon) \rightarrow S \subset \mathbb{R}^3$, de classe C^2 telle $\gamma(u) = \psi(u, 0)$ pour tout $u \in [a, b]$. On note alors $\gamma_v(u) = \psi(u, v)$, que l'on considère comme une famille à un paramètre de courbes tracées sur S et qui déforment la courbe initiale $\gamma = \gamma_0$. Nous avons alors le résultat suivant :

Théorème 5.5 (Formule de variation première pour la longueur). *Dans les conditions ci-dessus, la dérivée en $v = 0$ de la fonctionnelle longueur $v \rightarrow \ell(\gamma_v)$ en 0 est donnée par*

$$\left. \frac{\partial}{\partial v} \right|_{v=0} \ell(\gamma_v) = \left\langle \frac{\partial \psi}{\partial v}, \dot{\gamma}(u) \right\rangle \Big|_{u=a}^b - \int_a^b k_g(u) \left\langle \frac{\partial \psi}{\partial v}(u, 0), \boldsymbol{\mu}(u) \right\rangle du. \quad (5.1)$$

Preuve. On a

$$\left. \frac{\partial}{\partial v} \right|_{v=0} \ell(\gamma_v) = \int_a^b \left. \frac{\partial}{\partial v} \right|_{v=0} \left(\sqrt{\left\langle \frac{\partial \psi}{\partial u}, \frac{\partial \psi}{\partial u} \right\rangle} \right) du.$$

Pour simplifier cette intégrale, on observe que

$$\frac{\partial}{\partial v} \left(\sqrt{\left\langle \frac{\partial \psi}{\partial u}, \frac{\partial \psi}{\partial u} \right\rangle} \right) = \frac{\left\langle \frac{\partial^2 \psi}{\partial v \partial u}, \frac{\partial \psi}{\partial u} \right\rangle}{\sqrt{\left\langle \frac{\partial \psi}{\partial u}, \frac{\partial \psi}{\partial u} \right\rangle}} = \frac{1}{\left\| \frac{\partial \psi}{\partial u} \right\|} \left\langle \frac{\partial^2 \psi}{\partial u \partial v}, \frac{\partial \psi}{\partial u} \right\rangle.$$

Nous avons supposé que pour $v = 0$, on $\left\| \frac{\partial \psi}{\partial u} \right\| = \|\dot{\gamma}(u)\| = 1$, par conséquent

$$\left. \frac{\partial}{\partial v} \right|_{v=0} \ell(\gamma_v) = \int_a^b \left\langle \frac{\partial^2 \psi}{\partial u \partial v}, \frac{\partial \psi}{\partial u} \right\rangle du = \left\langle \frac{\partial \psi}{\partial v}, \frac{\partial \psi}{\partial u} \right\rangle \Big|_{u=a}^b - \int_a^b \left\langle \frac{\partial \psi}{\partial v}, \frac{\partial^2 \psi}{\partial u^2} \right\rangle du.$$

(on a intégré par parties). En $v = 0$, nous avons $\frac{\partial \psi}{\partial u} = \dot{\gamma}(u)$ et

$$\left\langle \frac{\partial \psi}{\partial v}, \frac{\partial^2 \psi}{\partial u^2} \right\rangle = \left\langle \frac{\partial \psi}{\partial v}, \ddot{\gamma}(u) \right\rangle = \left\langle \frac{\partial \psi}{\partial v}, \mathbf{K}_\gamma(u) \right\rangle = \left\langle \frac{\partial \psi}{\partial v}, k_g(u) \boldsymbol{\mu}(u) \right\rangle$$

car $\frac{\partial \psi}{\partial v}$ est un champ de vecteurs tangent à la surface. On a donc finalement

$$\left. \frac{\partial}{\partial v} \right|_{v=0} \ell(\gamma_v) = \left\langle \frac{\partial \psi}{\partial v}(u, 0), \dot{\gamma}(u) \right\rangle \Big|_{u=a}^b - \int_a^b k_g(u) \left\langle \frac{\partial \psi}{\partial v}, \boldsymbol{\mu}(u) \right\rangle du.$$

□

Une conséquence importante de ce théorème dit qu'une courbe sur une surface S qui minimise la distance intrinsèque entre ses extrémités est une géodésique si elle est parcourue à vitesse constante :

Corollaire 5.6. Soit $\gamma : [a, b]$ une courbe de classe C^2 paramétrée à vitesse constante sur la surface S . Si la longueur de γ est égale à la distance entre les points $p = \gamma(a)$ et $q = \gamma(b)$, alors γ est une géodésique de S .

Preuve. Soit $\psi : [a, b] \times (-\varepsilon, \varepsilon) \rightarrow S \subset \mathbb{R}^3$ une déformation quelconque de classe C^2 de γ dont les extrémités sont fixées, i.e. $\psi(a, v) = p$ et $\psi(b, v) = q$ pour tout $v \in (-\varepsilon, \varepsilon)$, alors $\frac{\partial \psi}{\partial v}$ s'annule lorsque $u = a$ et $u = b$. La formule de variation première (5.1) s'écrit donc

$$\left. \frac{\partial}{\partial v} \right|_{v=0} \ell(\gamma_v) = - \int_a^b k_g(t) \langle \xi(u), \mu(u) \rangle du,$$

où on a noté pour simplifier $\xi(u) = \frac{\partial \psi}{\partial v}(u, 0)$. Mais par hypothèse $\ell(\gamma_0) = d(p, q)$ est la longueur minimale parmi toutes les courbes sur S qui relient p à q , par conséquent $\left. \frac{\partial}{\partial v} \right|_{v=0} \ell(\gamma_v) = 0$ et on a donc

$$\int_a^b k_g(t) \langle \xi(u), \mu(u) \rangle du = 0,$$

Dans cette égalité, ξ est un champ de vecteurs quelconque le long de γ qui s'annule aux extrémités de la courbe (car nous avons considéré une déformation ψ à extrémités fixes quelconque de γ). Cette condition implique que la courbure géodésique k_g de γ est identiquement nulle et donc γ est géodésique. □

Une géodésique ne minimise pas toujours la distance entre ses extrémités, toutefois c'est le cas localement ; et cette propriété caractérise les géodésiques :

Théorème 5.7. Une courbe de classe C^2 sur une surface régulière S est une géodésique de cette surface si et seulement si

- (i) La vitesse V de γ est constante.
- (ii) La courbe γ réalise localement les distances minimales entre les points qu'elle parcourt. De façon plus précise, pour tout $t_0 \in I$ il existe $\varepsilon > 0$ tel que si $t_1, t_2 \in [t_0 - \varepsilon, t_0 + \varepsilon]$, alors la distance $d_S(\gamma(t_1), \gamma(t_2))$ est égale à la longueur de l'arc $\gamma|_{[t_1, t_2]}$.

Il suit du corollaire précédent que si la courbe γ vérifie les conditions (i) et (ii), alors c'est une géodésique. La preuve de l'affirmation sort du cadre de ce cours.

Remarque. La notion de géodésique est *a priori* une notion *cinématique* puisqu'elle fait intervenir l'accélération de la courbe. On définit parfois une géodésique comme une courbe qui réalise localement la distance entre les points de cette courbe. Cela revient à garder la condition (ii) du corollaire et à oublier la condition (i) ; avec cette définition alternative la notion de géodésique devient une notion *géométrique*, équivalente à la condition que la courbure géodésique s'annule.

5.2.3 Le théorème de Meusnier

Le théorème de Meusnier² dit que la courbure normale d'une courbe tracée sur une surface en un point p ne dépend que de la direction de cette courbe en ce point :

Théorème 5.8 (Meusnier, 1785). *Soit $\gamma : I \rightarrow S$ une courbe régulière de classe C^2 tracée sur une surface co-orientée S . Alors sa courbure normale en $t \in I$ ne dépend que de la direction de $\dot{\gamma}(t)$. Plus précisément, nous avons la formule suivante :*

$$k_n(t) = -\frac{\langle d\boldsymbol{\nu}(\dot{\gamma}(t)), \dot{\gamma}(t) \rangle}{\|\dot{\gamma}(t)\|^2}. \quad (5.2)$$

En particulier la courbure normale ne dépend pas de l'accélération de la courbe.

Remarquons que la formule précédente peut aussi s'écrire

$$k_n(t) = -\langle d\boldsymbol{\nu}(\mathbf{T}_\gamma(t)), \mathbf{T}_\gamma(t) \rangle, \quad (5.3)$$

où $\mathbf{T}_\gamma(t) = \frac{\dot{\gamma}(t)}{\|\dot{\gamma}(t)\|}$ est le vecteur tangent à γ en t .

Preuve. On a clairement $\langle \boldsymbol{\nu}(\gamma(t)), \dot{\gamma}(t) \rangle = 0$ pour tout $t \in I$. En dérivant cette relation et en appliquant la formule de l'accélération on obtient

$$\begin{aligned} 0 &= \frac{d}{dt} \langle \boldsymbol{\nu}(\gamma(t)), \dot{\gamma}(t) \rangle = \left\langle \frac{d}{dt} \boldsymbol{\nu}(\gamma(t)), \dot{\gamma}(t) \right\rangle + \langle \boldsymbol{\nu}(\gamma(t)), \ddot{\gamma}(t) \rangle \\ &= \left\langle \frac{d}{dt} \boldsymbol{\nu}(\gamma(t)), \dot{\gamma}(t) \right\rangle + \dot{V}_\gamma(t) \langle \boldsymbol{\nu}(\gamma(t)), \mathbf{T}_\gamma(t) \rangle + V_\gamma(t)^2 \langle \boldsymbol{\nu}(\gamma(t)), \mathbf{K}_\gamma(t) \rangle, \end{aligned}$$

où $V_\gamma(t) = \|\dot{\gamma}(t)\|$. Nous avons $\langle \boldsymbol{\nu}(\gamma(t)), \mathbf{T}_\gamma(t) \rangle = 0$ et $k_n(t) = \langle \boldsymbol{\nu}(\gamma(t)), \mathbf{K}_\gamma(t) \rangle$ est la courbure normale de γ . Notons aussi que

$$\frac{d}{dt} \boldsymbol{\nu}(\gamma(t)) = d\boldsymbol{\nu}(\dot{\gamma}(t)).$$

Par conséquent le calcul précédent montre que

$$k_n(t) = -\frac{\langle d\boldsymbol{\nu}(\dot{\gamma}(t)), \dot{\gamma}(t) \rangle}{V_\gamma(t)^2}.$$

□

Nous avons immédiatement le corollaire suivant :

Corollaire 5.9. *Si $\gamma_1 : (-\varepsilon, \varepsilon) \rightarrow S$ et $\gamma_2 : (-\varepsilon, \varepsilon) \rightarrow S$ sont deux courbes régulières de classe C^2 sur une surface régulière co-orientée $S \subset \mathbb{R}^3$ de classe C^2 telles que $\gamma_1(0) = \gamma_2(0) = p$ et $\dot{\gamma}_1(0) = \lambda \dot{\gamma}_2(0)$ avec $\lambda \neq 0$, alors γ_1 et γ_2 ont même courbure normale en $t = 0$.*

5.3 L'application de Weingarten et la deuxième forme fondamentale

Dans cette section, nous présentons différentes notions de courbure liées à une surface. Nous commençons par la définition suivante, qui est motivée par la preuve du théorème de Meusnier :

2. Jean-Baptiste Meusnier (1754–1793).

Définition 5.10. On appelle *application de Weingarten* en un point p d'une surface $S \subset \mathbb{R}^3$ co-orientée, régulière de classe C^2 , la différentielle en ce point de l'application de Gauss. On note cette application

$$L_p = d\boldsymbol{\nu}_p.$$

Remarque 5.11. (1) L'application de Weingarten s'appelle souvent *the shape operator* dans les livres en anglais.

(2) Certain livres définissent l'application de Weingarten avec le signe opposé (i.e. $L_p = -d\boldsymbol{\nu}_p$).

(3) L'application de Weingarten est a priori une application linéaire entre le plan tangent en p à la surface S et le plan tangent à $\boldsymbol{\nu}(p)$ à la sphère unité $\mathbb{S}^2$. Cependant le vecteur $\boldsymbol{\nu}(p)$ est à la fois vecteur normal de $T_p S$ et vecteur normal à $T_{\boldsymbol{\nu}(p)} \mathbb{S}^2$, donc ces deux plans tangents coïncident et on peut donc considérer que l'application de Weingarten au point $p \in S$ est un endomorphisme du plan tangent $T_p S$:

$$L_p = d\boldsymbol{\nu}_p : T_p S \rightarrow T_p S.$$

Définition 5.12. La *seconde forme fondamentale* en un point p d'une surface régulière de classe C^2 co-orientée $S \subset \mathbb{R}^3$ est l'application bilinéaire $\mathbf{h}_p : T_p S \times T_p S \rightarrow \mathbb{R}$ définie sur le plan tangent $T_p S$ par

$$\mathbf{h}_p(\xi, \eta) = -\mathbf{g}_p(L_p(\xi), \eta)$$

où $\mathbf{g}$ est le tenseur métrique associé à ψ .

La formule (5.2) nous dit que la courbure normale d'une courbe de C^2 sur la surface S peut s'écrire

$$k_n(t) = -\frac{\mathbf{g}_p(L_p(\dot{\gamma}(t)), \dot{\gamma}(t))}{\mathbf{g}_p(\dot{\gamma}(t), \dot{\gamma}(t))} = \frac{\mathbf{h}_p(\dot{\gamma}(t), \dot{\gamma}(t))}{\mathbf{g}_p(\dot{\gamma}(t), \dot{\gamma}(t))} \quad (5.4)$$

Proposition 5.13. Soit $\psi : \Omega \rightarrow S$ une surface paramétrée régulière de classe C^2 et notons $\mathbf{b}_i = \frac{\partial \psi}{\partial u_i}$, alors les coefficients de la seconde forme fondamentale dans la base $\{\mathbf{b}_1(u), \mathbf{b}_2(u)\}$ de $T_p S$ adaptée à la paramétrisation ψ en un point $p = \psi(u)$ sont donnés par

$$h_{ij}(u) = \mathbf{h}_p(\mathbf{b}_i, \mathbf{b}_j) = \left\langle \boldsymbol{\nu}(p), \frac{\partial^2 \psi}{\partial u_i \partial u_j}(u) \right\rangle. \quad (5.5)$$

Preuve. Pour simplifier la suite, on notera $\boldsymbol{\nu}(u)$ pour $\boldsymbol{\nu}(\psi(u))$. Pour tout $u \in \Omega$, nous avons $\langle \boldsymbol{\nu}(u), \mathbf{b}_j(u) \rangle = 0$, par conséquent

$$\left\langle \frac{\partial \boldsymbol{\nu}}{\partial u_i}, \mathbf{b}_j \right\rangle + \left\langle \boldsymbol{\nu}, \frac{\partial \mathbf{b}_j}{\partial u_i} \right\rangle = 0.$$

Or, par définition de l'application de Weingarten, on a

$$L(\mathbf{b}_i) = d\boldsymbol{\nu}(\mathbf{b}_i) = d\boldsymbol{\nu} \left(\frac{\partial \psi}{\partial u_i} \right) = \frac{\partial \boldsymbol{\nu} \circ \psi}{\partial u_i}.$$

et

$$\frac{\partial \mathbf{b}_j}{\partial u_i} = \frac{\partial}{\partial u_i} \frac{\partial \psi}{\partial u_j} = \frac{\partial^2 \psi}{\partial u_i \partial u_j}.$$

Par conséquent :

$$\mathbf{h}_p(\mathbf{b}_i, \mathbf{b}_j) = -\langle L(\mathbf{b}_i), \mathbf{b}_j \rangle = -\left\langle \frac{\partial \boldsymbol{\nu} \circ \psi}{\partial u_i}, \mathbf{b}_j \right\rangle = \left\langle \boldsymbol{\nu}, \frac{\partial \mathbf{b}_j}{\partial u_i} \right\rangle = \left\langle \boldsymbol{\nu}, \frac{\partial^2 \psi}{\partial u_i \partial u_j} \right\rangle.$$

□

Remarque. Il est commode de noter

$$\mathbf{b}_{ij} = \frac{\partial \mathbf{b}_j}{\partial u_i} = \frac{\partial^2 \psi}{\partial u_i \partial u_j}.$$

La proposition précédente nous dit qu'en tout point de la surface paramétrée par ψ , on a

$$h_{ij} = \langle \boldsymbol{\nu}, \mathbf{b}_{ij} \rangle.$$

Corollaire 5.14. Si $\psi : \Omega \rightarrow S$ une surface paramétrée régulière de classe C^2 , alors la seconde forme fondamentale h est une forme bilinéaire symétrique en tout point de S :

$$\mathbf{h}_p(\xi, \eta) = \mathbf{h}_p(\eta, \xi),$$

pour tous $\xi, \eta \in T_p S$. De façon équivalente, l'application de Weingarten L est auto-adjointe, i.e. on a

$$\mathbf{g}_p(L_p(\xi), \eta) = \mathbf{g}_p(\xi, L_p(\eta)).$$

Preuve. Il suffit de vérifier que $h(\mathbf{b}_1, \mathbf{b}_2) = h(\mathbf{b}_2, \mathbf{b}_1)$, ce qui se déduit immédiatement de la proposition précédente car $\frac{\partial^2 \psi}{\partial u_2 \partial u_1} = \frac{\partial^2 \psi}{\partial u_1 \partial u_2}$.

□

Proposition 5.15. Notons $\mathbf{G}$ la matrice du tenseur métrique dans la base $\{\mathbf{b}_1, \mathbf{b}_2\}$ en un point donné de la surface paramétrée $\psi : \Omega \rightarrow S$. Notons de même $\mathbf{H}$ la matrice de la seconde forme fondamentale et $\mathbf{L}$ la matrice de l'application de Weingarten. Alors on a

$$\mathbf{H} = -\mathbf{G}\mathbf{L} = -\mathbf{L}^\top \mathbf{G}$$

Cette formule est utile en pratique car il est souvent plus facile de calculer la deuxième forme fondamentale que l'application de Weingarten. On peut donc calculer d'abord $\mathbf{G}$ et $\mathbf{H}$, puis

$$\mathbf{L} = -\mathbf{G}^{-1}\mathbf{H}.$$

Preuve. Pour deux vecteurs tangents quelconques $\xi = \xi_1 \mathbf{b}_1 + \xi_2 \mathbf{b}_2$ et $\eta = \eta_1 \mathbf{b}_1 + \eta_2 \mathbf{b}_2$, nous avons $\mathbf{h}(\xi, \eta) = -\mathbf{g}(L(\xi), \eta) = -\mathbf{g}(\xi, L(\eta))$. Cette relation s'écrit matriciellement

$$\xi^\top \mathbf{H} \eta = -\xi^\top \mathbf{G}(\mathbf{L} \eta) = -\xi^\top (\mathbf{G}\mathbf{L}) \eta.$$

Comme ξ et η sont quelconques, on doit avoir $\mathbf{H} = -\mathbf{G}\mathbf{L}$.

On a également $\xi^\top \mathbf{H} \eta = -(\mathbf{L}\xi)^\top \mathbf{G} \eta = -\xi^\top (\mathbf{L}^\top \mathbf{G}) \eta$, qui entraîne $\mathbf{H} = -\mathbf{L}^\top \mathbf{G}$.

□

5.4 Les différentes courbures d'une surface

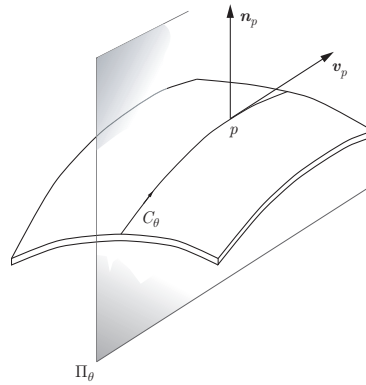
5.4.1 La courbure normale

Le théorème de Meusnier, ou plus précisément la formule (5.2), nous suggère la définition suivante :

Définition 5.16. Soit p un point d'une surface régulière de classe C^2 co-orientée. La *courbure normale en direction du vecteur tangent* non nul $\mathbf{v} \in T_p S \setminus \{\mathbf{0}\}$ est définie par

$$k_n(\mathbf{v}) = \frac{\mathbf{h}_p(\mathbf{v}, \mathbf{v})}{\mathbf{g}_p(\mathbf{v}, \mathbf{v})} = -\frac{\mathbf{g}_p(L_p(\mathbf{v}), \mathbf{v})}{\mathbf{g}_p(\mathbf{v}, \mathbf{v})}$$

La preuve du théorème de Meusnier montre que la courbure normale d'une courbe régulière γ tracée sur S est précisément égale à $k_n(\dot{\gamma}(t))$. En particulier, la courbure normale $k_n(\mathbf{v})$ en p est la courbure de l'intersection de la surface S avec le plan $\Pi_{\mathbf{v}}$ passant par p et de directions ν et $\mathbf{v}$. Une telle courbe s'appelle une *section normale* de la surface S .



Section normale d'une surface

Rappelons que l'application de Weingarten L_p est un point p d'une surface régulière S de classe C^2 est un endomorphisme autoadjoint du plan tangent $T_p S$. Par le théorème spectral, on sait donc que les valeurs propres de L_p sont réelles et qu'il existe une base orthonormée de $T_p S$ formée de vecteurs propres de L_p .

5.4.2 Courbures principales, moyenne et de Gauss

Définitions.

- 1.) Les valeurs propres de $-L_p$ s'appellent les *courbures principales* de S au point p . On les note $k_1(p)$ et $k_2(p)$; on supposera que $k_1 \leq k_2$.
- 2.) Le déterminant de L_p s'appelle la *courbure de Gauss* de S au point p . On note

$$K(p) = \det(L_p) = k_1(p)k_2(p).$$

- 3.) Le point $p \in S$ est dit
 - *elliptique* si $K(p) > 0$, c'est-à-dire si les courbures principales en p ont le même signe.

- *hyperbolique* si $K(p) < 0$, c'est-à-dire si les courbures principales en p ont des signes opposés.
 - *parabolique* si l'une (et une seule) des courbures principales en p est nulle.
 - *planaire* si les deux courbures principales en p sont nulles : $k_1(p) = k_2(p) = 0$.
 - *ombilique* si les deux courbures principales en p sont égales : $k_1(p) = k_2(p)$ (de façon équivalente, le point p est ombilique si L_p est scalaire).
- 4.) La *courbure moyenne* de S en p est la moyenne des courbures principales, on la note³

$$H(p) = \frac{1}{2}(k_1(p) + k_2(p)) = -\frac{1}{2}\text{Trace}(L_p).$$

- 5.) Les *directions principales* de S en un point non ombilique sont les directions des vecteurs propres de L_p .
- 6.) Une courbe de classe C^1 tracée sur la surface S est une *ligne de courbure* de S si elle est tangente en chaque point à une direction principale.

Remarques.

1. Les directions principales en un point non ombilique p sont orthogonales car ce sont des vecteurs propres de l'opérateur autoadjoint L_p .
2. Il suit immédiatement de la proposition 5.15 que la courbure de Gauss est donnée par

$$K(p) = \frac{\det(\mathbf{H}(p))}{\det(\mathbf{G}(p))}, \quad (5.6)$$

où $\mathbf{G}$ et $\mathbf{H}$ sont les matrices de la première et la seconde forme fondamentale de S .

Le résultat suivant, dû à Euler, nous dit que la courbure normale d'une surface dans une direction non nulle s'exprime en fonction des courbures principales et de l'angle que fait la direction considérée avec les directions principales :

Proposition 5.17 (Euler). *Soit p un point non ombilique d'une surface régulière de classe C^2 . On note $\mathbf{v}_1$ et $\mathbf{v}_2$ les vecteurs unités de $T_p S$ dans les direction principales. Alors la courbure normale du vecteur $\mathbf{v}_\theta = \cos(\theta)\mathbf{v}_1 + \sin(\theta)\mathbf{v}_2 \in T_p S$ est donnée par*

$$k_n(\mathbf{v}_\theta) = k_1 \cos(\theta)^2 + k_2 \sin(\theta)^2,$$

où k_1, k_2 sont les courbures principales de S en p .

Nous laissons la preuve de cette proposition en exercice.

Corollaire 5.18. *Les courbures principales en un point p d'une surface régulière S de classe C^2 sont les valeurs minimale et maximale de la courbure normale de S en ce point.*

Preuve. Notons k_1 et k_2 les courbures principales de p en S , et supposons que $k_1 \geq k_2$. La formule précédente peut aussi s'écrire

$$k_n(\mathbf{v}_\theta) = k_1 + (k_2 - k_1) \sin(\theta)^2.$$

Nous avons $0 \leq \sin(\theta)^2 \leq 1$, donc $k_1 \leq k_n(\mathbf{v}_\theta) \leq k_2$ et on a $k_n(\mathbf{v}_\theta) = k_1$ lorsque $\sin(\theta) = 0$ et $k_n(\mathbf{v}_\theta) = k_2$ lorsque $\sin(\theta) = \pm 1$. □

3. Attention aux notations, ne pas confondre la seconde forme fondamentale et la courbure moyenne, ça devrait être clair dans chaque cas selon le contexte.

5.4.3 Interprétation locale des courbures principales.

Pour étudier la géométrie locale d'une surface S de classe C^2 au voisinage d'un point régulier p , il est commode d'introduire un système de coordonnées cartésien $Oxyz$ dont l'origine 0 coïncide avec le point p et le plan tangent à S en 0 est le plan Oxy . Dans ce cas, on dit que le système de coordonnées cartésien est *adapté à la surface S en point p* . La surface est alors localement représentée comme le graphe $z = \varphi(x, y)$ d'une fonction $\varphi : \Omega \rightarrow \mathbb{R}$ de classe C^2 où $\Omega \subset \mathbb{R}^2$ est un voisinage de $(0, 0)$ dans le plan. De plus on a

$$\varphi(0, 0) = 0, \quad \frac{\partial \varphi}{\partial x}(0, 0) = \frac{\partial \varphi}{\partial y}(0, 0) = 0.$$

Une paramétrisation locale de la surface est alors donnée par l'application $\psi : \Omega \rightarrow \mathbb{R}^3$ définie par $\psi(x, y) = (x, y, \varphi(x, y))$ et la base adaptée en 0 est

$$\mathbf{b}_1(0, 0) = \frac{\partial \psi}{\partial x}(0, 0) = (1, 0, 0) = \mathbf{e}_1, \quad \mathbf{b}_2(0, 0) = \frac{\partial \psi}{\partial y}(0, 0) = (0, 1, 0) = \mathbf{e}_2.$$

Pour la suite nous choisirons la co-orientation définie par le vecteur normal $\boldsymbol{\nu} = \mathbf{b}_1 \times \mathbf{b}_2 = \mathbf{e}_3$.

Le tenseur métrique à l'origine prend la valeur $\mathbf{G} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$.

Les coefficients de la seconde forme fondamentale en $(0, 0)$ sont les produits scalaires

$$h_{11} = \langle \boldsymbol{\nu}, \psi_{xx} \rangle = \langle \mathbf{e}_3, \psi_{xx} \rangle = \varphi_{xx},$$

et de même $h_{12} = \varphi_{xy}$ et $h_{22} = \varphi_{yy}$, on a donc

$$\mathbf{H} = \begin{pmatrix} \varphi_{xx} & \varphi_{xy} \\ \varphi_{xy} & \varphi_{yy} \end{pmatrix}$$

(c'est la matrice hessienne de φ en $(0, 0)$). La matrice de l'application de Weingarten est alors donnée par

$$\mathbf{L} = -\mathbf{G}^{-1}\mathbf{H} = -\begin{pmatrix} \varphi_{xx} & \varphi_{xy} \\ \varphi_{xy} & \varphi_{yy} \end{pmatrix}$$

On a finalement

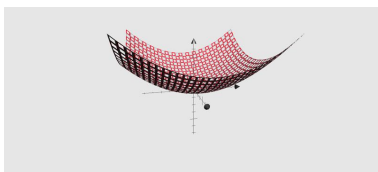
$$K = \det(\mathbf{L}) = \varphi_{xx}\varphi_{yy} - \varphi_{xy}^2 \quad \text{et} \quad H = -\frac{1}{2}\text{Trace}(\mathbf{L}) = \frac{1}{2}(\varphi_{xx} + \varphi_{yy}).$$

Noter que tous ces calculs sont valable en 0 , et a priori uniquement en 0 .

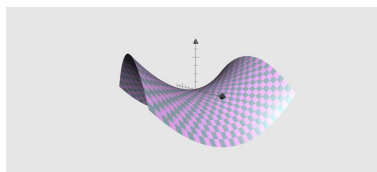
Quitte à effectuer une rotation de notre système de coordonnées autour de l'axe Oz , on peut supposer que les directions principales de S en 0 sont les directions des vecteurs $\mathbf{e}_1$ et $\mathbf{e}_2$. On a donc le développement de Taylor

$$\varphi(x, y) = \frac{1}{2}(ax^2 + by^2) + o(x^2 + y^2),$$

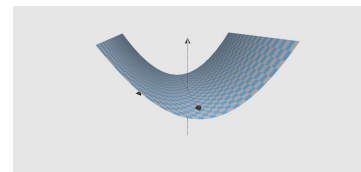
avec $a = \varphi_{xx}(0, 0)$ et $b = \varphi_{yy}(0, 0)$ (et on a $\varphi_{xy}(0, 0) = 0$). Ainsi les courbures principales de S en 0 sont $k_1 = a$ et $k_2 = b$, la courbure de Gauss est $K = ab$ et la courbure moyenne est $H = \frac{1}{2}(a + b)$.



Point elliptique



Point hyperbolique



Point parabolique

La représentation locale de la surface comme un graphe dans un système de coordonnées nous permet donc de facilement interpréter la géométrie des différents types de point : le point 0 est elliptique si a et b ont le même signe (non nul), il est hyperbolique si a et b ont des signes opposés (non nuls), il est parabolique si a ou b est nul mais pas les deux et il est plat si $a = b = 0$. Le point est ombilique si $a = b$.

Une remarque sur l'orientation des surfaces.

Si on change le signe de la co-orientation ν de la surface S , alors l'application de Weingarten L change de signe. Par conséquent le signe des courbures principales et de la courbure moyenne est sensible au choix de la co-orientation. Un calcul montre que dans le cas de la sphère, ces courbures sont positives pour le choix de la normale intérieure à la sphère et elles sont négatives pour le choix de la normale extérieure. Cela s'explique géométriquement par le fait que l'accélération d'une courbe tracée sur une sphère pointe vers l'intérieur de cette sphère. D'une manière générale, si la surface S est le bord d'un domaine de $\mathbb{R}^3$, il est préférable de choisir le champ normal ν pointant du côté intérieur, avec cette convention le bord d'un domaine convexe de $\mathbb{R}^3$ est de courbure moyenne positive. Observons en revanche que la courbure de Gauss $K = \det(L)$ ne dépend pas du choix de la co-orientation.

5.4.4 Courbure des surfaces de révolution

Considérons le cas d'une surface de révolution S autour de l'axe Oz dont le profil est la courbe $\alpha(v) = (r(v), z(v))$ ($v \in I$). On suppose que α est de classe C^2 , paramétrée naturellement, et que $r(v) > 0$ pour tout $v \in I$. La paramétrisation standard de cette surface est donné par $\psi : \Omega = [0, 2\pi] \times I \rightarrow S \subset \mathbb{R}^3$:

$$\psi(u, v) = (r(v) \cos(u), r(v) \sin(u), z(v)).$$

Le repère adapté est

$$\mathbf{b}_1 = \frac{\partial \psi}{\partial u} = \begin{pmatrix} -r(v) \sin(u) \\ r(v) \cos(u) \\ 0 \end{pmatrix}, \quad \mathbf{b}_2 = \frac{\partial \psi}{\partial v} = \begin{pmatrix} \dot{r}(v) \cos(u) \\ \dot{r}(v) \sin(u) \\ \dot{z}(v) \end{pmatrix}, \quad \nu = \frac{\mathbf{b}_1 \times \mathbf{b}_2}{\|\mathbf{b}_1 \times \mathbf{b}_2\|} = \begin{pmatrix} \dot{z}(v) \cos(u) \\ \dot{z}(v) \sin(u) \\ -\dot{r}(v) \end{pmatrix},$$

où on a noté $\dot{\cdot}$ pour la dérivée par rapport à v . Rappelons que $\dot{r}(v)^2 + \dot{z}(v)^2 = 1$ par hypothèse, le tenseur métrique est donc

$$\mathbf{G} = \begin{pmatrix} r^2(v) & 0 \\ 0 & 1 \end{pmatrix}$$

(qu'on peut aussi écrire $ds^2 = r^2(v)du^2 + dv^2$). Les dérivées secondes de ψ sont

$$\mathbf{b}_{11}(u, v) = \frac{\partial \mathbf{b}_1}{\partial u} = \frac{\partial^2 \psi}{\partial u^2} = \begin{pmatrix} -r(v) \cos(u) \\ -r(v) \sin(u) \\ 0 \end{pmatrix},$$

$$\mathbf{b}_{12}(u, v) = \frac{\partial \mathbf{b}_1}{\partial v} = \frac{\partial \mathbf{b}_2}{\partial u} = \frac{\partial^2 \psi}{\partial u \partial v} = \begin{pmatrix} -\dot{r}(v) \sin(u) \\ \dot{r}(v) \cos(u) \\ 0 \end{pmatrix},$$

et

$$\mathbf{b}_{22}(u, v) = \frac{\partial \mathbf{b}_2}{\partial v} = \frac{\partial^2 \psi}{\partial v^2} = \begin{pmatrix} \ddot{r}(v) \cos(u) \\ \ddot{r}(v) \sin(u) \\ \ddot{z}(v) \end{pmatrix}.$$

La matrice de la seconde forme fondamentale est alors donnée par $h_{ij} = \langle \boldsymbol{\nu}, \mathbf{b}_{ij} \rangle$,

$$\mathbf{H} = \begin{pmatrix} -r(v)\dot{z}(v) & 0 \\ 0 & (\ddot{r}(v)\dot{z}(v) - \ddot{z}(v)\dot{r}(v)) \end{pmatrix},$$

ce qui nous donne la courbure de Gauss :

$$K = \frac{\det(\mathbf{H})}{\det(\mathbf{G})} = -\frac{\dot{z}(v)(\ddot{r}(v)\dot{z}(v) - \ddot{z}(v)\dot{r}(v))}{r(v)}.$$

En utilisant la relation $\dot{r}(v)^2 + \dot{z}(v)^2 = 1$, on peut simplifier cette expression. On a

$$0 = \frac{d}{dv}(\dot{r}^2 + \dot{z}^2) = 2\dot{r}\ddot{r} + 2\dot{z}\ddot{z},$$

donc

$$\dot{z}^2 = 1 - \dot{r}^2 \quad \text{et} \quad \dot{z}\ddot{z} = -\dot{r}\ddot{r},$$

d'où l'on déduit que

$$K = -\frac{1}{r}(\ddot{r}\dot{z}^2 - \dot{z}\ddot{z}\dot{r}) = -\frac{1}{r}(\ddot{r}(1 - \dot{r}^2) + \dot{r}^2\ddot{r}) = -\frac{\ddot{r}}{r}.$$

On a donc montré que la courbure de Gauss de notre surface de révolution est

$$K(v) = -\frac{1}{r(v)} \frac{d^2 r(v)}{dv^2}. \quad (5.7)$$

Remarque. La matrice de l'application de Weingarten dans la base $\{\mathbf{b}_1, \mathbf{b}_2\}$ en un point quelconque de la surface de révolution S est donnée par

$$\mathbf{L} = -\mathbf{G}^{-1}\mathbf{H} = \begin{pmatrix} \frac{\dot{z}(v)}{r(v)} & 0 \\ 0 & -(\ddot{r}(v)\dot{z}(v) - \ddot{z}(v)\dot{r}(v)) \end{pmatrix}.$$

En particulier, les vecteurs $\mathbf{b}_1, \mathbf{b}_2$ sont les vecteurs propres de $\mathbf{L}$, cela prouve que les directions principales sur la surface sont les directions des parallèles et des méridiens.

Surfaces de révolution à courbure de Gauss constante.

La formule (5.7), nous permet de déterminer toutes les surfaces de révolution dont la courbure de Gauss K est constante. Il s'agit en effet de résoudre les équations différentielles

$$\ddot{r} + Kr = 0, \quad \dot{z} = \sqrt{1 - \dot{r}^2}.$$

Exemple 1. Supposons $K = 1$, alors une solution simple est donnée par $r(v) = \cos(v)$ et $z(v) = \sin(v)$. Cette solution correspond à la sphère unité standard.

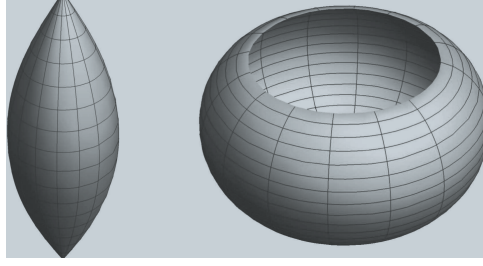
Les autres solutions sont données par

$$r(v) = a$$

Il y a d'autres solutions, qui donnent d'autres surfaces de révolutions à courbure de Gauss constante positives :

Exemple 2. Si $K = -1$, une solution est donnée par

$$r(v) = e^{-v}, \quad z(v) = \int_0^v \sqrt{1 - e^{-2s}} ds.$$



Cette dernière intégrale ne peut pas s'exprimer par des fonctions élémentaires, toutefois la courbe de profil $\alpha(v) = (r(v), z(v))$ peut-être décrite (et donc dessinée) par les propriétés suivantes :

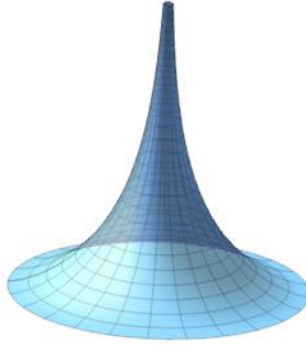
$$(r(0), z(0)) = (1, 0) \quad \text{et} \quad \alpha(v) + \mathbf{T}_\alpha(v) \text{ est un point de l'axe vertical } Oz.$$

En effet, le vecteur tangent $\mathbf{T}_\alpha$ est égal à $\dot{\alpha}$ puisque α est paramétrée naturellement, on a donc

$$\begin{aligned} \alpha(v) + \mathbf{T}_\alpha(v) &= \alpha(v) + \dot{\alpha}(v) \\ &= (r(v), z(v)) + (\dot{r}(v), \dot{z}(v)) \\ &= (r(v) + \dot{r}(v), z(v) + \dot{z}(v)) \\ &= (0, z(v) + \dot{z}(v)), \end{aligned}$$

$$\text{car } (v) + \dot{r}(v) = e^{-v} + (-e^{-v}) = 0.$$

Une telle courbe s'appelle une *tractrice* d'axe Oz et la surface de révolution d'une tractrice autour de son axe est la *pseudo-sphère de Mindina-Beltrami*⁴



4. La pseudo-sphère apparaît dans les travaux de Ferdinand Minding en 1839, puis de Eugenio Beltrami en 1868.

5.5 Quelques théorèmes classiques de la théorie des surfaces

Dans ce paragraphe nous énonçons, sans tous les démontrer, quelques théorèmes importants de la théorie des surfaces dans $\mathbb{R}^3$. Rappelons pour commencer qu'un point d'une surface est *ombilique* si les deux courbures principales en ce point sont égales.

Théorème 5.19. *Si tous les points d'une surface régulière $S \subset \mathbb{R}^3$ de classe C^3 sont ombiliques, alors S est contenue dans un plan ou dans une sphère.*

Preuve. Supposons que tout les points de S soient ombiliques. Alors il existe une fonction $\lambda : S \rightarrow \mathbb{R}$ telle que $L_p(\xi) = \lambda(p)\xi$ pour tout $p \in S$ et tout $\xi \in T_p S$. Soit $\psi : \Omega \rightarrow S$ une paramétrisation locale de S , nous avons alors avec les notations de la proposition 5.13 :

$$\frac{\partial \boldsymbol{\nu}}{\partial u_i} = L(\mathbf{b}_i(u)) = \lambda(\psi(u))\mathbf{b}_i(u) = \lambda(\psi(u))\frac{\partial \psi}{\partial u_i}.$$

Écrivons $\lambda(u) = \lambda(\psi(u))$ pour simplifier, alors

$$\frac{\partial^2 \boldsymbol{\nu}}{\partial u_1 \partial u_2} = \frac{\partial \lambda}{\partial u_1} \frac{\partial \psi}{\partial u_2} + \lambda \frac{\partial^2 \psi}{\partial u_1 \partial u_2} = \frac{\partial \lambda}{\partial u_1} \mathbf{b}_2 + \lambda \frac{\partial^2 \psi}{\partial u_1 \partial u_2}.$$

La même équation est vérifiée en échangeant les indices 1 et 2, on a donc pour tout $u \in \Omega$:

$$\frac{\partial \lambda}{\partial u_1} \mathbf{b}_2 - \frac{\partial \lambda}{\partial u_2} \mathbf{b}_1 = \frac{\partial^2 \boldsymbol{\nu}}{\partial u_1 \partial u_2} - \frac{\partial^2 \boldsymbol{\nu}}{\partial u_2 \partial u_1} + \lambda \frac{\partial^2 \psi}{\partial u_1 \partial u_2} - \lambda \frac{\partial^2 \psi}{\partial u_2 \partial u_1} = 0.$$

Puisque les vecteurs $\mathbf{b}_1$ et $\mathbf{b}_2$ sont en tout points linéairement indépendants, on doit avoir

$$\frac{\partial \lambda}{\partial u_1} = \frac{\partial \lambda}{\partial u_2} = 0,$$

et donc λ est constant. Si $\lambda = 0$, alors $\boldsymbol{\nu}$ est constant et S est contenu dans un plan orthogonal à ce vecteur. Si $\lambda \neq 0$, alors $c = \psi - \frac{1}{\lambda}\boldsymbol{\nu}$ est constant et la surface S est donc contenue dans la sphère de centre c et de rayon $1/|\lambda|$. □

Le résultat le plus important sur la courbure des surfaces est probablement le célèbre *théorème egregium* démontré par K. F. Gauss en 1827. Il dit que la courbure de Gauss est une notion intrinsèque de la géométrie des surfaces (deux surfaces intrinsèquement isométriques ont même courbure de Gauss).

Théorème 5.20 (Théorème Egregium de Carl Friedrich Gauss (1827)). *Si $f : S_1 \rightarrow S_2$ est une isométrie entre deux surfaces de $\mathbb{R}^3$ de classe C^3 pour la distance intrinsèque, alors on a $K_1 = K_2 \circ f$ où K_i est la courbure de Gauss de S_i .*

Nous démontrerons ce théorème plus loin (voir théorème C.3).

Exemple. On sait qu'un cône ou un cylindre sont des surfaces localement isométriques au plan, donc ces surfaces sont de courbure nulle. Le théorème egregium nous dit aussi qu'il n'existe pas d'isométrie entre un ouvert d'un ouvert d'une sphère et un ouvert du plan.

Dans le cas des surfaces à courbure de Gauss constante, F. Minding⁵ a démontré le résultat suivant, qui est une réciproque partielle du théorème egregium :

Théorème 5.21 (Théorème de Ernst Ferdinand Minding, 1839.). *Deux surfaces S_1 et S_2 de $\mathbb{R}^3$ de classe C^3 qui ont même courbure de Gauss constante sont localement isométriques pour la distance intrinsèque.*

5. Ferdinand Minding (1806-1885)

En particulier, ce théorème dit que toute surface de courbure nulle est localement isométrique au plan et toute surface dont la courbure de Gauss est constante positive est localement isométrique à une sphère.

Le résultats précédent concernait la géométrie locale des surfaces, i.e. la géométrie au voisinage d'un point quelconque de la surface. Les théorèmes suivants sont de nature globale.

Théorème 5.22 (Formule de Gauss-Bonnet). *Si S_1 et S_2 sont deux surfaces compactes sans bord de $\mathbb{R}^3$ qui sont homéomorphes, alors elles ont la même courbure totale :*

$$\iint_{S_1} K_1 dA_1 = \iint_{S_2} K_2 dA_2,$$

de plus cette courbure totale appartient à $4\pi\mathbb{Z}$.

Plus précisément, la courbure totale d'une telle surface est égale à sa caractéristique d'Euler, multipliée par 2π . La caractéristique d'Euler est un entier qui ne dépend que de la topologie de la surface.

Théorème 5.23 (H. Liebmann, 1900.). *Soit S une surface compacte (sans bord) de classe C^4 dans $\mathbb{R}^3$ dont la courbure de Gauss est partout positive. Supposons que ou bien la courbure de Gauss est constante ou bien la courbure moyenne est constante. Alors S est une sphère.*

Théorème 5.24 (Théorème de Jacques Hadamard sur les surfaces à courbure positive). *Soit $S \subset \mathbb{R}^3$ une surface régulière compacte de classe C^3 . Alors les conditions suivantes sont équivalentes :*

- (i) *La courbure de Gauss de S est strictement positive en tout point de S .*
- (ii) *S est le bord d'un domaine bornée strictement convexe $D \subset \mathbb{R}^3$.*
- (iii) *L'application de Gauss $\nu : S \rightarrow \mathbb{S}^2$ est un difféomorphisme de classe C^1 .*

Théorème 5.25. *Toute surface compacte sans bord de $\mathbb{R}^3$ admet au moins un point où la courbure de Gauss est strictement positive.*

Par comparaison, nous avons le résultat suivant sur les surfaces complètes à courbure constante négative :

Théorème 5.26 (Théorème de David Hilbert (1901)). *Il n'existe pas de surface régulière de classe C^3 dans $\mathbb{R}^3$ qui soit complète, sans bord, et dont la courbure de Gauss est constante négative.*

Rappelons qu'une surface est dite *complète* si toute suite de Cauchy dans cette surface converge.

Théorème 5.27 (Théorème de N. V. Efimov (1964)). *Il n'existe pas de surface régulière de classe C^2 dans $\mathbb{R}^3$ qui soit complète, sans bord, et dont la courbure de Gauss vérifie $\sup(K) < 0$.*

Annexe A

Notions de topologie et espaces vectoriels normés

A.1 Rappels de topologie

La topologie étudie et formalise les notions de *voisinage*, de *convergence* et de *continuité*.

Définition A.1. Soit X un ensemble. On appelle *topologie* sur X une famille de sous-ensembles $\mathcal{O} \subset \mathcal{P}(X)$ telle que

- i.) $\emptyset, X \in \mathcal{O}$,
- ii.) si $U, V \in \mathcal{O}$, alors $U \cap V \in \mathcal{O}$,
- iii.) si $\{U_\alpha\}_{\alpha \in A} \subset \mathcal{O}$ est une famille quelconque d'éléments de $\mathcal{O}$, alors $\bigcup_{\alpha \in A} U_\alpha \in \mathcal{O}$.

On dit que $U \subset X$ est *ouvert* si $U \in \mathcal{O}$ et que $F \subset X$ est *fermé* si $F^c = X \setminus F \in \mathcal{O}$. L'ensemble $A \subset X$ est un *voisinage* du point $p \in X$ s'il existe un ouvert $U \in \mathcal{O}$ tel que $p \in U \subset A$. Un *espace topologique* est un couple $(X, \mathcal{O})$ où X est un ensemble et $\mathcal{O}$ est une topologie sur X .

Il est clair que l'intersection d'une famille quelconque de fermés d'un espace topologique $(X, \mathcal{O})$ est un fermé. Pour tout $A \in \mathcal{O}$, on note $\bar{A}$ l'intersection de tous les fermés qui contiennent A :

$$\bar{A} = \bigcap_{F \supset A, F \text{ fermé}} F.$$

L'ensemble $\bar{A}$ est donc le plus petit ensemble fermé qui contient A . On l'appelle l'*adhérence* ou la *fermeture* de A . On le note aussi $\text{Cl}(A)$ ("Cl" pour *closure* = fermeture en anglais).

On définit aussi l'*intérieur* de A . C'est le plus grand ouvert qui est contenu dans A , on le note A° ou $\text{Int}(A)$, il est défini par

$$\text{Int}(A) = A^\circ = \bigcup_{U \subset A, U \text{ ouvert}} U.$$

Il est clair que $\text{Int}(A) \subset \text{Cl}(A)$; , la différence s'appelle la *frontière* de A et se note

$$\text{Fr}(A) = \text{Cl}(A) \setminus \text{Int}(A).$$

Une application $f : (X, \mathcal{O}_X) \rightarrow (Y, \mathcal{O}_Y)$ entre deux espaces topologiques est *continue* si l'image inverse d'un ouvert de Y est un ouvert de X , i.e. $f^{-1}(\mathcal{O}_Y) \subset \mathcal{O}_X$. L'application est *ouverte* si l'image directe d'un ouvert de X est un ouvert de Y , i.e. $f(\mathcal{O}_X) \subset \mathcal{O}_Y$.

L'application f est un *homéomorphisme* si elle est bijective, continue et ouverte (et donc $f^{-1} : Y \rightarrow X$ est aussi continue).

L'espace topologique $(X, \mathcal{O})$ est *séparé* (on dit aussi qu'il est *de Hausdorff*) si toute paire de points distincts admet des voisinages disjoints, i.e. si pour tout $p, q \in X$, $p \neq q$, il existe $U, V \in \mathcal{O}$ tels que $U \cap V = \emptyset$ et $p \in U$, $q \in V$. L'espace topologique $(X, \mathcal{O})$ est *connexe* si tout sous-ensemble qui est à la fois ouvert et fermé est égal à X ou $\emptyset$. La réunion de tous les sous ensembles connexes contenant un point $x \in X$ s'appelle la *composante connexe* de x . L'ensemble X est réunion disjointe des ses composantes connexes, et chaque composante connexe est un sous-ensemble connexe et maximal (i.e. qui n'est contenu dans aucun sous-ensemble connexe plus grand). L'espace X est *localement connexe* si tout point admet un voisinage connexe. Lorsque X est localement connexe, les composantes connexes de X sont les sous-ensembles qui sont ouverts, fermés et connexes.

L'espace topologique $(X, \mathcal{O})$ admet une base dénombrable d'ouverts (on dit aussi qu'il vérifie le *second axiome de dénombrabilité*) s'il existe une suite dénombrable d'ouverts $\{U_i\}_{i \in \mathbb{N}}$ telle que tout ouvert est réunion d'éléments de cette suite.

Exemples d'espaces topologiques.

- 1.) Pour tout ensemble X , l'ensemble $\mathcal{O} = \mathcal{P}(X)$ de toutes les parties de X est une topologie séparée appelée la *topologie discrète*.
- 2.) $\mathcal{O} = \{\emptyset, X\}$ est une topologie appelée la *topologie grossière*. Elle est non séparée dès que X contient au moins deux points.
- 3.) La collection des sous-ensembles de X qui sont vides ou de complémentaire fini est une topologie sur X (en général non séparée). On l'appelle la *topologie cofinie*.
- 4.) Si (X, d) est un espace métrique, alors il existe une topologie séparée dont les ouverts sont les parties $U \subset X$ qui sont réunion de boules ouvertes, i.e. d'ensembles du type

$$B(p, \varepsilon) = \{q \in X \mid d(p, q) < \varepsilon\}.$$

Un espace topologique $(X, \mathcal{O})$ est dit *métrisable* s'il existe une distance d sur X induisant la topologie $\mathcal{O}$.

- 5.) Si $(X, \mathcal{O}_X)$ est un espace topologique et $Y \subset X$, alors

$$\mathcal{O}_Y := \{V = U \cap Y \mid U \in \mathcal{O}_X\}$$

est une topologie sur Y . On l'appelle la *topologie relative* ou la *topologie induite* sur Y par $\mathcal{O}_X$.

Le théorème d'invariance du domaine

Un théorème fondamental sur la topologie de $\mathbb{R}^n$ est le suivant :

Théorème A.2 (Théorème d'invariance du domaine de Brouwer (1912)). *Si U est un ouvert de $\mathbb{R}^n$ et $f : U \rightarrow \mathbb{R}^n$ est une application continue et injective, alors f est une application ouverte (et c'est donc un homéomorphisme sur son image).*

Ce théorème a été démontré par le mathématicien néerlandais Luitzen E.J. Brouwer en 1912. Il existe plusieurs preuves, dont certaines utilisent des techniques de topologie algébrique. Nous admettons ce résultat sans démonstration.

Corollaire A.3 (Invariance de la dimension). *Soient U un ouvert de $\mathbb{R}^n$ non vide, et V un ouvert de $\mathbb{R}^m$. Si U et V sont homéomorphes, alors $m = n$. Plus généralement, deux variétés non vides qui sont homéomorphes ont même dimension.*

Rappelons que Cantor avait démontré qu'il existe une bijection entre $\mathbb{R}^n$ et $\mathbb{R}^m$ pour toute paire d'entiers $n, m \geq 1$, le corollaire ci-dessus nous dit qu'une telle bijection ne peut pas être un homéomorphisme si $n \neq m$, ce qui est conforme à notre intuition de la notion de dimension.

Preuve. Supposons que $n > m$ et que $g : U \rightarrow V$ est un homéomorphisme. Considérons l'application $f : U \rightarrow \mathbb{R}^n$ définie par

$$f(x_1, x_2, \dots, x_n) = (g_1(x), g_2(x), \dots, g_m(x), \underbrace{0, \dots, 0}_{n-m}).$$

Alors f est continue et injective, donc $f(U) \subset \mathbb{R}^n$ est ouvert par le théorème précédent. Mais c'est impossible car $f(U) \subset \{y \in \mathbb{R}^n \mid y_n = 0\}$ qui ne contient aucun sous-ensemble ouvert non vide. Donc il est impossible que $n > m$. De même $m \not> n$.

□

A.2 Rappels sur la notion de norme

Soit E un espace vectoriel sur le corps de réels. Rappelons qu'une *norme* sur E est une fonction $\|\cdot\| : E \rightarrow \mathbb{R}$ vérifiant les trois propriétés suivantes pour tous $x, y \in E$ et $\lambda \in \mathbb{R}$:

- i.) $\|x\| \geq 0$ et $\|x\| = 0$ si et seulement si $x = 0$,
- ii.) $\|\lambda x\| = |\lambda| \|x\|$,
- iii.) $\|x + y\| \leq \|x\| + \|y\|$.

A toute norme sur E on définit une distance d sur E définie par $d(x, y) = \|y - x\|$; en particulier une norme définit une topologie sur E et on peut alors parler d'ouverts, de fermés, d'ensembles compacts, de convergence, de continuité etc.

Deux normes $\|\cdot\|_1$ et $\|\cdot\|_2$ sur E sont dites *équivalentes* (ou *topologiquement équivalentes*) si elles définissent la même topologie.

Lemme A.4. *a) Les normes $\|\cdot\|_1$ et $\|\cdot\|_2$ sur l'espace vectoriel E sont équivalentes si et seulement s'il existe une constante $c > 0$ telle que pour tout $x \in E$ on a*

$$\frac{1}{c} \|x\|_2 \leq \|x\|_1 \leq c \|x\|_2.$$

b) Deux normes sur un espace vectoriel de dimension finies sont toujours équivalentes.

La preuve est un simple exercice du cours d'analyse 2.

Exemples de normes.

- 1.) Si $\langle , \rangle$ est un produit scalaire, alors $\|x\| = \sqrt{\langle x, x \rangle}$ est une norme (lorsque c'est le cas, on dit que la norme $\| \cdot \|$ *dérive d'un produit scalaire*).
- 2.) Pour $1 \leq p < \infty$, on définit la norme $\| \cdot \|_p$ sur $\mathbb{R}^n$ par

$$\|x\|_p = \left(\sum_{i=1}^n |x_i|^p \right)^{1/p}.$$

- 3.) Toujours sur $\mathbb{R}^n$, on définit la norme $\| \cdot \|_\infty$ par

$$\|x\|_\infty = \max |x_i|.$$

- 4.) Sur l'espace $C^0([0, 1])$ des fonctions continues sur l'intervalle $[0, 1]$, on définit aussi des normes $\| \cdot \|_p$:

$$\|f\|_p = \left(\int_0^1 |f(t)|^p dt \right)^{1/p} \quad \text{et} \quad \|f\|_\infty = \sup_{t \in [0, 1]} |f(t)|.$$

- 5.) Si $(V_1, \| \cdot \|_1)$ et $(V_2, \| \cdot \|_2)$ sont deux espaces normés de dimensions finies, on définit la *norme d'opérateurs* sur l'espace vectoriel $\mathcal{L}(V_1, V_2)$ des homomorphismes linéaires de V_1 dans V_2 par

$$\|\Phi\|_{\text{Op}} = \sup\{\|\Phi(x)\|_2 \mid \|x\|_1 \leq 1\}.$$

- 6.) La *norme de Hilbert-Schmidt* sur l'espace $M_n(\mathbb{R})$ des matrices carrées de taille n à coefficients réels est définie par

$$\|A\|_{\text{HS}} = \text{Trace}(A^\top A) = \sum_{i,j} \sqrt{A_{ij}^2}.$$

Proposition A.5. *Toute application linéaire entre deux espaces vectoriels réels normés de dimension finie est continue.*

Preuve. Soit $\Phi : V_1 \rightarrow V_2$ une application linéaire entre deux espaces normés de dimensions finies $(V_1, \| \cdot \|_1)$ et $(V_2, \| \cdot \|_2)$, et soit $\{e_1, \dots, e_n\}$ une base de V_1 et notons $C = \max_{1 \leq i \leq n} \|\Phi(e_i)\|_2$. Si $x = \sum_{i=1}^n x_i e_i$ et $y = \sum_{i=1}^n y_i e_i$, alors

$$\begin{aligned} \|\Phi(y) - \Phi(x)\|_2 &= \|\Phi(y - x)\|_2 = \left\| \Phi \left(\sum_{i=1}^n (y_i - x_i) e_i \right) \right\|_2 \\ &= \left\| \sum_{i=1}^n (y_i - x_i) \Phi(e_i) \right\|_2 \\ &\leq \sum_{i=1}^n |y_i - x_i| \|\Phi(e_i)\|_2 \\ &\leq C \sum_{i=1}^n |y_i - x_i|. \end{aligned}$$

Par conséquent, si $y \rightarrow x$ alors $\Phi(y) \rightarrow \Phi(x)$.

□

Annexe B

Sur les notations classiques de la géométrie différentielle des surfaces

Parmi les textes historiquement importants traitant de la géométrie différentielle des surfaces, on doit citer *Recherches sur la courbure des surfaces* par Leonhard Euler en 1760, *Application de l'analyse à la géométrie, à l'usage de l'École impériale polytechnique* par Gaspard Monge en 1807 et les *Leçons sur la théorie générale des surfaces et les applications géométriques du calcul infinitésimal* en 4 volumes par Gaston Darboux publiés entre 1887 et 1896. Ces développements historiques ont conduit à un système de notations assez différent de celui que nous avons exposés dans ces notes de cours, mais qui reste fréquemment utilisé car il est efficace dans les calculs.

On se donne d'abord un système d'axes orthonormés $Oxyz$ dans l'espace euclidien $\mathbb{R}^3$, en sorte qu'un point p peut être représenté par son *vecteur position* (ou *rayon vecteur*), qui est noté $\mathbf{r} = \overrightarrow{Op} = (x, y, z)$. On obtient une courbe $\mathbf{r}(t) = (x(t), y(t), z(t))$ lorsque le point dépend d'un paramètre t . L'abscisse curviligne le long de cette courbe est donnée par l'intégrale

$$s = \int \|\dot{\mathbf{r}}\| dt, \quad \text{où} \quad \|\dot{\mathbf{r}}\| = \sqrt{\left(\frac{dx}{dt}\right)^2 + \left(\frac{dy}{dt}\right)^2 + \left(\frac{dz}{dt}\right)^2}.$$

La différentielle $ds = \|\dot{\mathbf{r}}\| dt$ s'appelle l'*élément linéaire*, et il est commode d'écrire

$$ds^2 = d\mathbf{r} \cdot d\mathbf{r} = dx^2 + dy^2 + dz^2.$$

Lorsque le point dépend de deux paramètres u, v , on obtient une surface

$$\mathbf{r} = \mathbf{r}(u, v) = (x(u, v), y(u, v), z(u, v)).$$

On demandera à cette surface d'être régulière, ce qu'on exprimera par la condition

$$\mathbf{r}_u \times \mathbf{r}_v \text{ est non nul, où } \mathbf{r}_u = \frac{\partial \mathbf{r}}{\partial u} = \left(\frac{\partial x}{\partial u}, \frac{\partial y}{\partial u}, \frac{\partial z}{\partial u} \right), \quad \mathbf{r}_v = \frac{\partial \mathbf{r}}{\partial v} = \left(\frac{\partial x}{\partial v}, \frac{\partial y}{\partial v}, \frac{\partial z}{\partial v} \right).$$

Cette condition nous dit que l'application $(u, v) \mapsto \mathbf{r}(u, v)$ est une immersion d'un ouvert de $\mathbb{R}^2$ dans $\mathbb{R}^3$ et les vecteurs $\mathbf{r}_u, \mathbf{r}_v$ forment la base adaptée du plan tangent à la surface au point $\mathbf{r}(u, v)$. Une co-orientation de la surface est donnée en tout point $p = \mathbf{r}(u, v)$ par

$$\boldsymbol{\nu} = \frac{\mathbf{r}_u \times \mathbf{r}_v}{\|\mathbf{r}_u \times \mathbf{r}_v\|}.$$

L'élément linéaire peut donc se réécrire en fonction des différentielles du et dv :

$$\begin{aligned} ds^2 &= d\mathbf{r} \cdot d\mathbf{r} = \left(\frac{\partial \mathbf{r}}{\partial u} du + \frac{\partial \mathbf{r}}{\partial v} dv \right) \cdot \left(\frac{\partial \mathbf{r}}{\partial u} du + \frac{\partial \mathbf{r}}{\partial v} dv \right) \\ &= E(u, v) du^2 + 2F(u, v) du dv + G(u, v) dv^2, \end{aligned}$$

avec

$$\begin{aligned} E &= \frac{\partial \mathbf{r}}{\partial u} \cdot \frac{\partial \mathbf{r}}{\partial u} = \left(\frac{\partial x}{\partial u} \right)^2 + \left(\frac{\partial y}{\partial u} \right)^2 + \left(\frac{\partial z}{\partial u} \right)^2, \\ F &= \frac{\partial \mathbf{r}}{\partial u} \cdot \frac{\partial \mathbf{r}}{\partial v} = \frac{\partial x}{\partial u} \frac{\partial x}{\partial v} + \frac{\partial y}{\partial u} \frac{\partial y}{\partial v} + \frac{\partial z}{\partial u} \frac{\partial z}{\partial v}, \\ G &= \frac{\partial \mathbf{r}}{\partial v} \cdot \frac{\partial \mathbf{r}}{\partial v} = \left(\frac{\partial x}{\partial v} \right)^2 + \left(\frac{\partial y}{\partial v} \right)^2 + \left(\frac{\partial z}{\partial v} \right)^2. \end{aligned}$$

L'élément d'aire s'écrit alors

$$dA = \|\mathbf{r}_u \times \mathbf{r}_v\| \, dudv = \sqrt{EG - F^2} \, dudv.$$

Les paramètres u et v sont vus comme des coordonnées curvilignes sur la surface et les différentielles du , dv sont des coordonnées linéaires sur l'espace tangent à la surface au point $p = \mathbf{r}(u, v)$. Le ds^2 s'appelle aussi la *première forme fondamentale* et se note $I = d\mathbf{r} \cdot d\mathbf{r}$. En comparant avec les notations du §4.2, on remarque que

$$g_{11} = E, \quad g_{12} = F, \quad g_{22} = G.$$

Comme premier exemple, considérons l'hélicoïde $\mathbf{r}(u, v) = (v \cos(u), v \sin(u), u)$. On a

$$d\mathbf{r} = (-v \sin(u), v \cos(u), 1) du + (\cos(u), \sin(u), 0) dv,$$

donc l'élément linéaire est donné par

$$ds^2 = d\mathbf{r} \cdot d\mathbf{r} = (1 + v^2) du^2 + dv^2.$$

Le tenseur métrique est donc $E = g_{11} = (1 + v^2)$, $F = g_{12} = 0$ et $G = g_{22} = 1$. On a aussi

$$\mathbf{r}_u \times \mathbf{r}_v = (-\sin(u), \cos(u), -v),$$

et

$$dA = \|\mathbf{r}_u \times \mathbf{r}_v\| \, dudv = \sqrt{EG - F^2} \, dudv = \sqrt{1 + v^2} \, dudv.$$

Comme second exemple, on considère maintenant la surface de révolution autour de l'axe Oz paramétrée par

$$\mathbf{r}(u, v) = (\rho(v) \cos(u), \rho(v) \sin(u), z(v)).$$

Alors

$$d\mathbf{r} = (-\rho(v) \sin(u), \rho(v) \cos(u), 0) du + (\rho'(v) \cos(u), \rho'(v) \sin(u), z'(v)) dv,$$

l'élément linéaire est donné par

$$ds^2 = d\mathbf{r} \cdot d\mathbf{r} = \rho(v)^2 du^2 + (\rho'(v)^2 + z'(v)^2) dv^2.$$

On peut écrire le tenseur métrique matriciellement :

$$\begin{pmatrix} E & F \\ F & G \end{pmatrix} = \begin{pmatrix} \rho(v)^2 & 0 \\ 0 & (\rho'(v)^2 + z'(v)^2) \end{pmatrix}.$$

On a aussi

$$\mathbf{r}_u \times \mathbf{r}_v = (\rho(v) z'(v) \cos(u), \rho(v) z'(v) \sin(u), -\rho(v) \rho'(v)),$$

et

$$dA = \|\mathbf{r}_u \times \mathbf{r}_v\| \, dudv = \sqrt{EG - F^2} \, dudv = \rho(v) \sqrt{\rho'(v)^2 + z'(v)^2} \, dudv.$$

La deuxième forme fondamentale s'obtient en dérivant une seconde fois le vecteurs position $\mathbf{r}(u, v)$. On pose

$$\begin{aligned} L &= \mathbf{r}_{uu} \cdot \boldsymbol{\nu} = -\mathbf{r}_u \cdot \boldsymbol{\nu}_u, \\ M &= \mathbf{r}_{uv} \cdot \boldsymbol{\nu} = -\mathbf{r}_u \cdot \boldsymbol{\nu}_v, \\ N &= \mathbf{r}_{vv} \cdot \boldsymbol{\nu} = -\mathbf{r}_v \cdot \boldsymbol{\nu}_v. \end{aligned}$$

Alors la *seconde forme fondamentale* est donnée par

$$\text{II} = Ldu^2 + 2Mdudv + Ndv^2.$$

Les notations (e, f, g) sont parfois utilisées pour les coefficients (L, M, N) . Bien évidemment nous avons $L = h_{11}$, $M = h_{12}$ et $N = h_{22}$. Avec ces notations, la courbure normale d'une courbe $\gamma(t) = \mathbf{r}(u(t), v(t))$ est donnée par

$$k_n(t) = \frac{\text{II}(\dot{\gamma}, \dot{\gamma})}{\text{I}(\dot{\gamma}, \dot{\gamma})} = \frac{L\dot{u}^2 + 2M\dot{u}\dot{v} + N\dot{v}^2}{E\dot{u}^2 + 2F\dot{u}\dot{v} + G\dot{v}^2}.$$

La courbure de Gauss est

$$K = \frac{\det(\text{II})}{\det(\text{I})} = \frac{LN - M^2}{EG - F^2}.$$

Revenons aux exemples : Pour l'hélicoïde, nous avons

$$\mathbf{r}_u = (-v \sin(u), v \cos(u), 1), \quad \mathbf{r}_v = (\cos(u), \sin(u), 0), \quad \boldsymbol{\nu} = \frac{\mathbf{r}_u \times \mathbf{r}_v}{\|\mathbf{r}_u \times \mathbf{r}_v\|} = \frac{(-\sin(u), \cos(u), -v)}{\sqrt{1+v^2}},$$

et les coefficients de la seconde forme fondamentale sont donc

$$\begin{aligned} L &= \mathbf{r}_{uu} \cdot \boldsymbol{\nu} = 0, \\ M &= \mathbf{r}_{uv} \cdot \boldsymbol{\nu} = \frac{1}{\sqrt{1+v^2}}, \\ N &= \mathbf{r}_{vv} \cdot \boldsymbol{\nu} = 0. \end{aligned}$$

La courbure de Gauss de l'hélicoïde est alors

$$K = \frac{LN - M^2}{EG - F^2} = -\frac{1}{(1+v^2)^2}$$

Pour la surface de révolution, on a

$$\mathbf{r}_u = (-\rho(v) \sin(u), \rho(v) \cos(u), 0), \quad \mathbf{r}_v = (\rho'(v) \cos(u), \rho'(v) \sin(u), z'(v)),$$

et

$$\boldsymbol{\nu} = \frac{\mathbf{r}_u \times \mathbf{r}_v}{\|\mathbf{r}_u \times \mathbf{r}_v\|} = \frac{(z'(v) \cos(u), z'(v) \sin(u), -\rho'(v))}{\sqrt{\rho'(v)^2 + z'(v)^2}}.$$

Les coefficients de la seconde forme fondamentale sont alors

$$\begin{aligned} L &= \mathbf{r}_{uu} \cdot \boldsymbol{\nu} = -\frac{\rho(v)z'(v)}{\sqrt{\rho'(v)^2 + z'(v)^2}}, \\ M &= \mathbf{r}_{uv} \cdot \boldsymbol{\nu} = 0, \\ N &= \mathbf{r}_{vv} \cdot \boldsymbol{\nu} = \frac{\rho''(v)z'(v) - \rho'(v)z''(v)}{\sqrt{\rho'(v)^2 + z'(v)^2}}. \end{aligned}$$

La courbure de Gauss est alors

$$K = \frac{LN - M^2}{EG - F^2} = -\frac{z'(v)(\rho''(v)z'(v) - \rho'(v)z''(v))}{\rho(v)}$$

Annexe C

Symboles de Christoffel et preuve du Théorème Egregium

C.1 Les symboles de Christoffel

La seconde forme fondamentale d'une surface paramétrée contrôle les composantes normales des dérivées des vecteurs de bases $\{\mathbf{b}_1, \mathbf{b}_2\}$ adaptés. Les composantes tangentielles de ces dérivées s'expriment à partir des *symboles de Christoffel*¹, que nous introduisons ci-dessous. Ces quantités interviennent dans l'équation intrinsèque des géodésiques et jouent un rôle central dans la preuve du théorème egregium.

Rappelons que le repère mobile adapté à une paramétrisation $\psi : \Omega \rightarrow S$ se définit de la façon suivante :

$$\mathbf{b}_1 = \frac{\partial \psi}{\partial u_1}, \quad \mathbf{b}_2 = \frac{\partial \psi}{\partial u_2}, \quad \boldsymbol{\nu} = \frac{\mathbf{b}_1 \times \mathbf{b}_2}{\|\mathbf{b}_1 \times \mathbf{b}_2\|}.$$

Nous noterons $\mathbf{b}_{ij}$ les dérivées de $\mathbf{b}_1$ et $\mathbf{b}_2$:

$$\mathbf{b}_{ij} = \frac{\partial \mathbf{b}_i}{\partial u_j} = \frac{\partial^2 \psi}{\partial u_i \partial u_j}$$

où i, j prennent les valeurs 1 ou 2. Nous pouvons développer les vecteurs $\mathbf{b}_{ij}$ dans la base $\{\mathbf{b}_1, \mathbf{b}_2, \boldsymbol{\nu}\}$:

$\mathbf{b}_{ij} = \Gamma_{ij}^1 \mathbf{b}_1 + \Gamma_{ij}^2 \mathbf{b}_2 + h_{ij} \boldsymbol{\nu}.$

(C.1)

Observer que les h_{ij} sont les coefficients de la deuxième forme fondamentale.

Définition C.1. (i) Les coefficients Γ_{ij}^k s'appellent les *symboles de Christoffel de deuxième espèce* de la surface paramétrée.

(ii) Les *symboles de Christoffel de première espèce* sont les produits scalaires de $\mathbf{b}_{ij}$ avec $\mathbf{b}_k$.
On les notes

$\Gamma_{ijk} = \langle \mathbf{b}_{ij}, \mathbf{b}_k \rangle.$

(C.2)

1. Elwin Bruno Christoffel, mathématicien allemand 1829–1900.

Remarques.

1. Les symboles de Christoffel sont des fonctions des paramètres $(u_1, u_2) \in \Omega$.
2. Les symboles de Christoffel de première et deuxième espèces s'expriment linéairement les uns en fonctions des autres. On a en effet

$$\Gamma_{ijk} = \langle \mathbf{b}_{ij}, \mathbf{b}_k \rangle = \langle \Gamma_{ij}^1 \mathbf{b}_1 + \Gamma_{ij}^2 \mathbf{b}_2 + h_{ij} \boldsymbol{\nu}, \mathbf{b}_k \rangle = \Gamma_{ij}^1 g_{1k} + \Gamma_{ij}^2 g_{2k},$$

de façon spécifique :

$$\Gamma_{ij1} = g_{11}\Gamma_{ij}^1 + g_{12}\Gamma_{ij}^2 \quad \text{et} \quad \Gamma_{ij2} = g_{21}\Gamma_{ij}^1 + g_{22}\Gamma_{ij}^2. \quad (\text{C.3})$$

On peut inverser cette relation :

$$\Gamma_{ij}^1 = \frac{g_{22}\Gamma_{ij1} - g_{12}\Gamma_{ij2}}{g_{11}g_{22} - g_{12}^2} \quad \text{et} \quad \Gamma_{ij}^2 = \frac{g_{11}\Gamma_{ij2} - g_{12}\Gamma_{ij1}}{g_{11}g_{22} - g_{12}^2}. \quad (\text{C.4})$$

3. Les symboles de Christoffel sont symétriques en leur deux premiers indices :

$$\Gamma_{ijk} = \Gamma_{jik} \quad \text{et} \quad \Gamma_{ij}^k = \Gamma_{ji}^k.$$

Cela découle de l'égalité $\mathbf{b}_{21} = \mathbf{b}_{12}$, qui provient de la symétrie des dérivées partielles d'ordre 2 pour une fonction de classe C^2 :

$$\frac{\partial^2 \psi}{\partial u_2 \partial u_1} = \frac{\partial^2 \psi}{\partial u_1 \partial u_2}.$$

Le lemme suivant jouera un rôle fondamental dans la suite. Il nous dit que les symboles de Christoffel ne dépendent que de la géométrie intrinsèque de la surface.

Lemme C.2 (Levi-Civita). *Les symboles de Christoffel d'une surface paramétrée $\psi : \Omega \rightarrow S \subset \mathbb{R}^3$ de classe C^2 ne dépendent que des coefficients g_{ij} du tenseur métrique et de leur dérivées du premier ordre.*

Preuve. En dérivant le coefficient g_{jk} du tenseur métrique, on voit que

$$\frac{\partial g_{jk}}{\partial u_i} = \frac{\partial}{\partial u_i} \langle \mathbf{b}_j, \mathbf{b}_k \rangle = \langle \mathbf{b}_{ij}, \mathbf{b}_k \rangle + \langle \mathbf{b}_j, \mathbf{b}_{ik} \rangle.$$

C'est-à-dire

$$\frac{\partial g_{jk}}{\partial u_i} = \Gamma_{ijk} + \Gamma_{ikj}.$$

De même on a

$$\frac{\partial g_{ik}}{\partial u_j} = \Gamma_{jik} + \Gamma_{jki} \quad \text{et} \quad \frac{\partial g_{ij}}{\partial u_k} = \Gamma_{kij} + \Gamma_{kji}.$$

On a donc, en tenant compte de la symétrie $\Gamma_{ij,k} = \Gamma_{ji,k}$,

$$\frac{\partial g_{jk}}{\partial u_i} + \frac{\partial g_{ik}}{\partial u_j} - \frac{\partial g_{ij}}{\partial u_k} = (\Gamma_{ijk} + \Gamma_{ikj}) + (\Gamma_{jik} + \Gamma_{jki}) - (\Gamma_{kij} + \Gamma_{kji}) = 2\Gamma_{ijk}.$$

Les symboles de Christoffel de première espèce sont donc donnés par la somme suivante de dérivées des coefficients du tenseur métrique :

$$\Gamma_{ijk} = \frac{1}{2} \left(\frac{\partial g_{jk}}{\partial u_i} + \frac{\partial g_{ik}}{\partial u_j} - \frac{\partial g_{ij}}{\partial u_k} \right) \quad (\text{C.5})$$

En appliquant (C.4), on voit que les symboles de Christoffel de deuxième espèce ne dépendent également que des coefficients du tenseur métrique et de leur dérivées.

□

C.2 Accélération des courbes tracées sur une surface

Les courbes sur les surfaces possèdent la propriété remarquable suivante.

Théorème C.1 *Soit $\gamma : I \rightarrow S$ une courbe de classe C^2 tracée sur la surface S supposée également de classe C^2 . Alors son accélération normale en un point est donnée par*

$$\langle \boldsymbol{\nu}, \ddot{\gamma} \rangle = \mathbf{h}(\dot{\gamma}, \dot{\gamma}), \quad (\text{C.6})$$

où h est la seconde forme fondamentale. En particulier, l'accélération normale en un point ne dépend que du vecteur vitesse en ce point.

Ce théorème nous dit que si deux courbes sur S passent par un même point p , et ont le même vecteur vitesse en ce point, alors elles ont aussi la même accélération normale.

Preuve. Nous présentons deux preuves de ce résultat. La première preuve est très courte : on sait que $\langle \boldsymbol{\nu}(\gamma(t)), \dot{\gamma}(t) \rangle = 0$ pour tout $t \in S$, on a donc

$$\langle \boldsymbol{\nu}(\gamma(t)), \ddot{\gamma}(t) \rangle = - \left\langle \frac{d}{dt} \boldsymbol{\nu}(\gamma(t)), \dot{\gamma}(t) \right\rangle = - \langle d\boldsymbol{\nu}(\dot{\gamma}(t)), \dot{\gamma}(t) \rangle = \mathbf{h}(\dot{\gamma}, \dot{\gamma}).$$

La seconde preuve donne plus de détails sur le vecteur accélération : On peut représenter la courbe γ dans la paramétrisation de la surface par $\gamma(t) = \psi(u_1(t), u_2(t))$, et donc

$$\dot{\gamma}(t) = \dot{u}_1 \mathbf{b}_1 + \dot{u}_2 \mathbf{b}_2.$$

Ainsi,

$$\begin{aligned} \ddot{\gamma}(t) &= \ddot{u}_1 \mathbf{b}_1 + \ddot{u}_2 \mathbf{b}_2 + \dot{u}_1 \dot{\mathbf{b}}_1 + \dot{u}_2 \dot{\mathbf{b}}_2 \\ &= \ddot{u}_1 \mathbf{b}_1 + \ddot{u}_2 \mathbf{b}_2 + \dot{u}_1 (\dot{u}_1 \mathbf{b}_{11} + \dot{u}_2 \mathbf{b}_{12}) + \dot{u}_2 (\dot{u}_1 \mathbf{b}_{21} + \dot{u}_2 \mathbf{b}_{22}) \\ &= \ddot{u}_1 \mathbf{b}_1 + \ddot{u}_2 \mathbf{b}_2 + (\dot{u}_1)^2 \mathbf{b}_{11} + 2\dot{u}_1 \dot{u}_2 \mathbf{b}_{12} + (\dot{u}_2)^2 \mathbf{b}_{22}. \end{aligned}$$

Les vecteurs $\mathbf{b}_1$ et $\mathbf{b}_2$ sont orthogonaux à $\boldsymbol{\nu}$, et comme $h_{ij} = \langle \mathbf{b}_{ij}, \boldsymbol{\nu} \rangle$ on a

$$\langle \ddot{\gamma}(t), \boldsymbol{\nu} \rangle = h_{11}(\dot{u}_1)^2 + 2h_{12}\dot{u}_1\dot{u}_2 + h_{22}(\dot{u}_2)^2 = \mathbf{H}(\dot{\gamma}, \dot{\gamma}).$$

□

A l'aide des symboles de Christoffel, nous pouvons développer plus complètement le calcul de l'accélération. Nous avons vu lors de la démonstration précédente que

$$\ddot{\gamma}(t) = \ddot{u}_1 \mathbf{b}_1 + \ddot{u}_2 \mathbf{b}_2 + (\dot{u}_1)^2 \mathbf{b}_{11} + 2\dot{u}_1 \dot{u}_2 \mathbf{b}_{12} + (\dot{u}_2)^2 \mathbf{b}_{22}.$$

En développant les vecteurs $\mathbf{b}_{ij}$ dans la base $\mathbf{b}_1, \mathbf{b}_2, \boldsymbol{\nu}$ via l'équation (C.2), nous trouvons

$$\begin{aligned} \ddot{\gamma}(t) &= (\ddot{u}_1 + \Gamma_{11}^1 \dot{u}_1^2 + 2\Gamma_{12}^1 \dot{u}_1 \dot{u}_2 + \Gamma_{22}^1 \dot{u}_2^2) \mathbf{b}_1 \\ &\quad + (\ddot{u}_2 + \Gamma_{11}^2 \dot{u}_1^2 + 2\Gamma_{12}^2 \dot{u}_1 \dot{u}_2 + \Gamma_{22}^2 \dot{u}_2^2) \mathbf{b}_2 \\ &\quad + \mathbf{h}(\dot{\gamma}, \dot{\gamma}) \boldsymbol{\nu}. \end{aligned}$$

L'équation des géodésiques peut en particulier se récrire sous la forme suivante :

$$\begin{cases} \ddot{u}_1 + \Gamma_{11}^1 \dot{u}_1^2 + 2\Gamma_{12}^1 \dot{u}_1 \dot{u}_2 + \Gamma_{22}^1 \dot{u}_2^2 = 0 \\ \ddot{u}_2 + \Gamma_{11}^2 \dot{u}_1^2 + 2\Gamma_{12}^2 \dot{u}_1 \dot{u}_2 + \Gamma_{22}^2 \dot{u}_2^2 = 0. \end{cases} \quad (\text{C.7})$$

Ces équations, avec le lemme (C.2), montrent en particulier que la notion de géodésique ne dépend que de la géométrie intrinsèque de la surface.

C.3 Preuve du Théorème Egregium

Nous reformulons le théorème Egregium de la façon suivante :

Théorème C.3. *La courbure de Gauss d'une surface paramétrée $\psi : \Omega \rightarrow S \subset \mathbb{R}^3$ de classe C^3 ne dépend que des coefficients g_{ij} du tenseur métrique et de leur dérivées jusqu'à l'ordre 2.*

Le démonstration utilise les deux lemmes suivants qui sont de nature calculatoire. Leur preuve ne présentent pas de difficulté particulière autre qu'une attention soutenue aux indices. L'effort se justifie par l'importance du théorème egregium.

Lemme C.4. *On a*

$$\langle \mathbf{b}_{ij}, \mathbf{b}_{km} \rangle = \Gamma_{km}^1 \Gamma_{ij1} + \Gamma_{km}^2 \Gamma_{ij2} + h_{km} h_{ij}.$$

Preuve. Rappelons que $\mathbf{b}_{km} = \Gamma_{km}^1 \mathbf{b}_1 + \Gamma_{km}^2 \mathbf{b}_2 + h_{km} \boldsymbol{\nu}$, donc

$$\begin{aligned} \langle \mathbf{b}_{ij}, \mathbf{b}_{km} \rangle &= \langle \mathbf{b}_{ij}, \Gamma_{km}^1 \mathbf{b}_1 + \Gamma_{km}^2 \mathbf{b}_2 + h_{km} \boldsymbol{\nu} \rangle \\ &= \Gamma_{km}^1 \langle \mathbf{b}_{ij}, \mathbf{b}_1 \rangle + \Gamma_{km}^2 \langle \mathbf{b}_{ij}, \mathbf{b}_2 \rangle + h_{km} \langle \mathbf{b}_{ij}, \boldsymbol{\nu} \rangle \\ &= \Gamma_{km}^1 \Gamma_{ij1} + \Gamma_{km}^2 \Gamma_{ij2} + h_{km} h_{ij} \end{aligned}$$

□

Lemme C.5. *On a*

$$\langle \mathbf{b}_{11}, \mathbf{b}_{22} \rangle - \|\mathbf{b}_{12}\|^2 = \frac{\partial}{\partial u_1} \Gamma_{221} - \frac{\partial}{\partial u_2} \Gamma_{121}.$$

Preuve. Calculons

$$\frac{\partial}{\partial u_1} \Gamma_{221} = \frac{\partial}{\partial u_1} \langle \mathbf{b}_{22}, \mathbf{b}_1 \rangle = \left\langle \frac{\partial \mathbf{b}_{22}}{\partial u_1}, \mathbf{b}_1 \right\rangle + \left\langle \mathbf{b}_{22}, \frac{\partial \mathbf{b}_1}{\partial u_1} \right\rangle = \left\langle \frac{\partial \mathbf{b}_{22}}{\partial u_1}, \mathbf{b}_1 \right\rangle + \langle \mathbf{b}_{22}, \mathbf{b}_{11} \rangle.$$

On a donc

$$\left\langle \frac{\partial \mathbf{b}_{22}}{\partial u_1}, \mathbf{b}_1 \right\rangle = \frac{\partial}{\partial u_1} \Gamma_{221} - \langle \mathbf{b}_{22}, \mathbf{b}_{11} \rangle.$$

De même

$$\left\langle \frac{\partial \mathbf{b}_{12}}{\partial u_2}, \mathbf{b}_1 \right\rangle = \frac{\partial}{\partial u_2} \Gamma_{121} - \langle \mathbf{b}_{12}, \mathbf{b}_{12} \rangle.$$

La différence de ces deux identités prouve le lemme car ²

$$\frac{\partial \mathbf{b}_{22}}{\partial u_1} - \frac{\partial \mathbf{b}_{12}}{\partial u_2} = \frac{\partial^3 \psi}{\partial u_1 \partial u_2 \partial u_2} - \frac{\partial^3 \psi}{\partial u_2 \partial u_2 \partial u_1} = 0.$$

□

Démonstration du Théorème Egregium. En appliquant les deux lemmes précédents, on voit que

$$\begin{aligned} \frac{\partial}{\partial u_1} \Gamma_{221} - \frac{\partial}{\partial u_2} \Gamma_{121} &= \langle \mathbf{b}_{11}, \mathbf{b}_{22} \rangle - \|\mathbf{b}_{12}\|^2 \\ &= (\Gamma_{11}^1 \Gamma_{221} + \Gamma_{11}^2 \Gamma_{222} + h_{11} h_{22}) - (\Gamma_{12}^1 \Gamma_{121} + \Gamma_{12}^2 \Gamma_{122} + h_{12} h_{12}) \\ &= (h_{11} h_{22} - h_{12}^2) + (\Gamma_{11}^1 \Gamma_{221} + \Gamma_{11}^2 \Gamma_{222} - \Gamma_{12}^1 \Gamma_{121} - \Gamma_{12}^2 \Gamma_{122}). \end{aligned}$$

2. C'est à cet endroit qu'on doit supposer que ψ est de classe C^3 .

On a donc

$$h_{11}h_{22} - h_{12}^2 = \frac{\partial}{\partial u_1}\Gamma_{221} - \frac{\partial}{\partial u_2}\Gamma_{121} - \Gamma_{11}^1\Gamma_{221} - \Gamma_{11}^2\Gamma_{222} + \Gamma_{12}^1\Gamma_{121} + \Gamma_{12}^2\Gamma_{122}.$$

Par conséquent la courbure de Gauss $K = \frac{\det(\mathbf{H})}{\det(\mathbf{G})}$ peut s'écrire

$$K = \frac{1}{g_{11}g_{22} - g_{12}^2} \left(\frac{\partial}{\partial u_1}\Gamma_{221} - \frac{\partial}{\partial u_2}\Gamma_{121} - \Gamma_{11}^1\Gamma_{221} - \Gamma_{11}^2\Gamma_{222} + \Gamma_{12}^1\Gamma_{121} + \Gamma_{12}^2\Gamma_{122} \right). \quad (\text{C.8})$$

Il suit alors du lemme (C.2) que la courbure de Gauss est fonction des coefficients g_{ij} et de leur dérivées premières et secondes. □

Exemple. Supposons que le tenseur métrique est donné par $ds^2 = du_1^2 + a^2 du_2^2$, où a est une fonction positive de (u_1, u_2) . On a donc $g_{11} = 1$, $g_{12} = g_{21} = 0$ et $g_{22} = a^2$. Les symboles de Christoffel de première espèce se calculent à partir de (C.5). On trouve que

$$\Gamma_{221} = -a \frac{\partial a}{\partial u_1}, \quad \Gamma_{122} = \Gamma_{212} = a \frac{\partial a}{\partial u_1}, \quad \Gamma_{222} = a \frac{\partial a}{\partial u_2},$$

et tous les autres Γ_{ijk} sont nuls. Pour les coefficients de seconde espèce, nous avons $\Gamma_{ij}^1 = \Gamma_{ij1}$ et $\Gamma_{ij}^2 = \frac{1}{a^2}\Gamma_{ij2}$. Donc la formule (C.8) se réduit à

$$K = \frac{1}{a^2} \left(\frac{\partial}{\partial u_1}\Gamma_{221} + \Gamma_{12}^2\Gamma_{122} \right) = \frac{1}{a^2} \left(\frac{\partial}{\partial u_1}\Gamma_{221} + \frac{1}{a^2}(\Gamma_{122})^2 \right) = -\frac{1}{a} \frac{\partial^2 a}{\partial u_1^2}. \quad (\text{C.9})$$

C.4 Les équations de Codazzi-Mainardi

Le théorème C.3 est une conséquence de l'identité $\frac{\partial \mathbf{b}_{11}}{\partial u_2} = \frac{\partial \mathbf{b}_{12}}{\partial u_1}$ provenant de la symétrie des dérivées partielles. On a plus généralement

$$\frac{\partial \mathbf{b}_{ii}}{\partial u_k} = \frac{\partial \mathbf{b}_{ik}}{\partial u_i}, \quad (\text{C.10})$$

et cette identité nous permet de trouver de nouvelles relations. On vérifie par un calcul direct que

$$\left\langle \boldsymbol{\nu}, \frac{\partial \mathbf{b}_{ik}}{\partial u_i} \right\rangle = \Gamma_{ik}^1 h_{1i} + \Gamma_{ik}^2 h_{2i} + \frac{\partial h_{ik}}{\partial u_i}, \quad (\text{C.11})$$

et

$$\left\langle \boldsymbol{\nu}, \frac{\partial \mathbf{b}_{ii}}{\partial u_k} \right\rangle = \Gamma_{ii}^1 h_{1k} + \Gamma_{ii}^2 h_{2k} + \frac{\partial h_{ii}}{\partial u_k}. \quad (\text{C.12})$$

En utilisant (C.10), on obtient alors

$$\frac{\partial h_{ik}}{\partial u_i} - \frac{\partial h_{ii}}{\partial u_k} = \Gamma_{ii}^1 h_{1k} + \Gamma_{ii}^2 h_{2k} - \Gamma_{ik}^1 h_{1i} - \Gamma_{ik}^2 h_{2i}. \quad (\text{C.13})$$

C'est l'équation de Codazzi-Mainardi.

Un théorème démontré par P. Bonnet en 1867 nous dit que *si l'on se donne deux matrices (g_{ij}) et (h_{ij}) de taille 2×2 qui dépendent de deux paramètres u, v et qui vérifient les relations données dans les équations de Gauß et de Codazzi-Mainardi, alors il existe un morceau de surface dans $\mathbb{R}^3$ pour laquelle (g_{ij}) est la première forme fondamentale et (h_{ij}) la deuxième. Cette surface est unique à un déplacement près.*³

Le théorème de Bonnet entraîne en particulier que toutes les relations qui existent entre la première et la deuxième forme fondamentales (et leur dérivées) sont des conséquences des équations de Gauß et de Codazzi-Mainardi.

3. C'est le *théorème fondamental* de la théorie des surfaces.

Annexe D

Formulaire

• **Produit scalaire** Un espace vectoriel euclidien est un espace vectoriel réel de dimension finie muni d'un produit scalaire (i.e. une forme bilinéaire symétrique définie positive) qu'on note $\langle \cdot, \cdot \rangle$. Dans une base orthonormée il est donné par $\langle \mathbf{a}, \mathbf{b} \rangle = \sum_{i=1}^n a_i b_i$. A partir de la norme le produit scalaire s'exprime

$$\langle \mathbf{a}, \mathbf{b} \rangle = \frac{1}{4} \left(\|\mathbf{a} + \mathbf{b}\|^2 - \|\mathbf{a} - \mathbf{b}\|^2 \right).$$

On a aussi

$\langle \mathbf{a}, \mathbf{a} \rangle = \ \mathbf{a}\ ^2$	$ \langle \mathbf{a}, \mathbf{b} \rangle \leq \ \mathbf{a}\ \ \mathbf{b}\ $
$\langle \mathbf{a}, \mathbf{b} \rangle = \ \mathbf{a}\ \ \mathbf{b}\ \cos(\vartheta(\mathbf{a}, \mathbf{b}))$	$\text{proj}_{\mathbf{a}}(\mathbf{b}) = \frac{\langle \mathbf{a}, \mathbf{b} \rangle}{\ \mathbf{a}\ ^2} \mathbf{a}$
$\langle \mathbf{a}, \mathbf{b} \rangle = \frac{1}{2} (\ \mathbf{a} + \mathbf{b}\ ^2 - \ \mathbf{a}\ ^2 - \ \mathbf{b}\ ^2)$	$\langle \mathbf{a}, \mathbf{b} \rangle = \frac{1}{2} (\ \mathbf{a}\ ^2 + \ \mathbf{b}\ ^2 - \ \mathbf{a} - \mathbf{b}\ ^2)$

• **Produits vectoriel et mixte dans $\mathbb{R}^3$.**

1.) Dans une base orthonormée d'orientation positive on a

$$\mathbf{a} \times \mathbf{b} = \begin{vmatrix} a_2 & b_2 \\ a_3 & b_3 \end{vmatrix} \mathbf{e}_1 - \begin{vmatrix} a_1 & b_1 \\ a_3 & b_3 \end{vmatrix} \mathbf{e}_2 + \begin{vmatrix} a_1 & b_1 \\ a_2 & b_2 \end{vmatrix} \mathbf{e}_3$$

2.) $(\mathbf{a} \times \mathbf{b}) \times \mathbf{c} = \langle \mathbf{a}, \mathbf{c} \rangle \mathbf{b} - \langle \mathbf{b}, \mathbf{c} \rangle \mathbf{a}$

3.) $\mathbf{a} \times (\mathbf{b} \times \mathbf{c}) = \langle \mathbf{a}, \mathbf{c} \rangle \mathbf{b} - \langle \mathbf{a}, \mathbf{b} \rangle \mathbf{c}$

4.) $\langle \mathbf{a} \times \mathbf{b}, \mathbf{c} \times \mathbf{d} \rangle = \langle \mathbf{a}, \mathbf{c} \rangle \langle \mathbf{b}, \mathbf{d} \rangle - \langle \mathbf{a}, \mathbf{d} \rangle \langle \mathbf{b}, \mathbf{c} \rangle$

5.) $\langle \mathbf{a} \times \mathbf{b}, \mathbf{c} \times \mathbf{d} \rangle = \langle (\mathbf{a} \times \mathbf{b}) \times \mathbf{c}, \mathbf{d} \rangle$.

6.) Le *produit mixte* de trois vecteurs $\mathbf{a}, \mathbf{b}, \mathbf{c} \in \mathbb{V}^3$ est défini par $[\mathbf{a}, \mathbf{b}, \mathbf{c}] = \langle \mathbf{a} \times \mathbf{b}, \mathbf{c} \rangle = \langle \mathbf{a}, \mathbf{b} \times \mathbf{c} \rangle$ il est trilinéaire et est donné par le déterminant 3×3 formé par la matrice dont les colonnes sont les coefficients des 3 vecteurs.

7.) on a $(\mathbf{a} \times \mathbf{b}) \times (\mathbf{c} \times \mathbf{d}) = [\mathbf{a}, \mathbf{b}, \mathbf{d}] \mathbf{c} - [\mathbf{a}, \mathbf{b}, \mathbf{c}] \mathbf{d}$.

• **Produits extérieur dans le plan.** Dans le plan orienté $\mathbb{R}^2$, le produit extérieur de deux vecteurs est

$$\mathbf{a} \wedge \mathbf{b} = \langle \mathbf{J}(\mathbf{a}), \mathbf{b} \rangle$$

où $\mathbf{J}$ est l'opérateur de rotation d'angle $\pi/2$ dans le sens positif.

• **Courbes.** Le vecteur vitesse d'une courbe γ de $\mathbb{R}^n$ se note $\dot{\gamma}$. La vitesse est $V = V_\gamma(u) = \|\dot{\gamma}(u)\|$ et l'abscisse curviligne depuis le point initial $\gamma(u_0)$ est

$$s(u) = \int_{u_0}^u V_\gamma(\tau) d\tau.$$

La formule de l'accélération est

$$\ddot{\gamma}(u) = \dot{V}\mathbf{T} + V^2\mathbf{K}$$

où $\mathbf{T} = \frac{1}{V}\dot{\gamma}$ et $\mathbf{K} = \frac{1}{V}\dot{\mathbf{T}}$ est le vecteur de courbure. La courbure de γ est la fonction scalaire $\kappa(u) = \|\mathbf{K}(u)\|$.

• **Repère de Frenet.** Si $\gamma(u) \in \mathbb{R}^3$ est C^3 et birégulière, le repère mobile de Frenet est le repère repère orthonormé direct d'origine $\gamma(u)$ et de base

$$\mathbf{T} = \frac{1}{V}\dot{\gamma}, \quad \mathbf{N} = \frac{\dot{\mathbf{T}}}{\|\dot{\mathbf{T}}\|} = \frac{1}{\kappa}\mathbf{K}, \quad \mathbf{B} = \frac{\dot{\gamma} \times \ddot{\gamma}}{\|\dot{\gamma} \times \ddot{\gamma}\|}.$$

La torsion est $\tau = \frac{1}{V}\langle \mathbf{B}, \dot{\mathbf{N}} \rangle$ et on a les équations de Serret-Frenet :

$$\frac{1}{V}\dot{\mathbf{T}} = \kappa\mathbf{N}, \quad \frac{1}{V}\dot{\mathbf{N}} = -\kappa\mathbf{T} + \tau\mathbf{B}, \quad \frac{1}{V}\dot{\mathbf{B}} = -\tau\mathbf{N}$$

On a aussi

$$\kappa = \frac{\|\dot{\gamma} \times \ddot{\gamma}\|}{V^3}, \quad \tau = \frac{[\dot{\gamma}, \ddot{\gamma}, \ddot{\gamma}]}{\|\dot{\gamma} \times \ddot{\gamma}\|^2} = \frac{[\dot{\gamma}, \ddot{\gamma}, \ddot{\gamma}]}{\kappa^2 V^6}$$

Le vecteur de Darboux est le champ de vecteurs le long de γ défini par

$$\mathbf{D} = \tau\mathbf{T} + \kappa\mathbf{B}$$

• **Surfaces paramétrée.** Si $\psi : \Omega \rightarrow \mathbb{R}^3$ est une surface paramétrée, le repère mobile adapté est

$$\mathbf{b}_1 = \frac{\partial \vec{\psi}}{\partial u_1}(u_1, u_2), \quad \mathbf{b}_2 = \frac{\partial \vec{\psi}}{\partial u_2}(u_1, u_2), \quad \boldsymbol{\nu} = \frac{\mathbf{b}_1 \times \mathbf{b}_2}{\|\mathbf{b}_1 \times \mathbf{b}_2\|}$$

$\mathbf{b}_1, \mathbf{b}_2$ engendrent le plan tangent à la surface au point $p = \psi(u_1, u_2)$ et $\boldsymbol{\nu}$ est le vecteur normal. Si $f(x, y, z) = 0$ est une équation pour la surface alors on a aussi

$$\boldsymbol{\nu} = \pm \frac{\vec{\nabla} f}{\|\vec{\nabla} f\|}.$$

Le tenseur métrique $G = (g_{ij})$ est la matrice de Gram de $\{\mathbf{b}_1, \mathbf{b}_2\}$, i.e. $g_{ij} = \langle \mathbf{b}_i, \mathbf{b}_j \rangle$.

L'élément d'aire infinitésimale est

$$dA = \sqrt{g_{11}g_{22} - g_{12}^2} \cdot du_1 du_2 = \|\mathbf{b}_1 \times \mathbf{b}_2\| du_1 du_2$$

et l'élément de longueur infinitésimale est

$$ds = \sqrt{g_{11} du_1^2 + 2g_{12} du_1 du_2 + g_{22} du_2^2}$$

• **Repère de Darboux, courbures normales et géodésiques.** Si γ est tracée sur la surface S , on note $\boldsymbol{\mu} = \boldsymbol{\nu} \times \mathbf{T}$. Le repère de Darboux est $\{\boldsymbol{\nu}, \mathbf{T}_\gamma, \boldsymbol{\mu}\}$. La courbure normale, la courbure géodésique et la torsion géodésique de γ sont définis par

$$k_n(u) = \langle \mathbf{K}_\gamma(u), \boldsymbol{\nu}(u) \rangle, \quad k_g(u) = \langle \mathbf{K}_\gamma(u), \boldsymbol{\mu}(u) \rangle \quad \text{et} \quad \tau_g(u) = \frac{1}{V_\gamma(u)} \langle \dot{\boldsymbol{\nu}}(u), \boldsymbol{\mu}(u) \rangle.$$

Les équations de Darboux sont :

$$\frac{1}{V} \dot{\mathbf{T}} = k_g \boldsymbol{\mu} + k_n \boldsymbol{\nu}, \quad \frac{1}{V} \dot{\boldsymbol{\nu}} = -k_n \mathbf{T} + \tau_g \boldsymbol{\mu}, \quad \frac{1}{V} \dot{\boldsymbol{\mu}} = -k_g \mathbf{T} - \tau_g \boldsymbol{\nu}.$$

• **Application de Weingarten et deuxième forme fondamentale.** L'application de Weingarten L_p en un point d'une surface S est l'endomorphisme de $T_p S$ défini par $L_p = d\boldsymbol{\nu}_p$.

La deuxième forme fondamentale est la forme bilinéaire sur $T_p S$ définie par $h_p(\xi, \eta) = -\langle L_p(\xi), \eta \rangle$. Les coefficients de h_p dans la base adaptée $\{\mathbf{b}_1, \mathbf{b}_2\}$ se calculent par :

$$h_{ij} = h(\mathbf{b}_i, \mathbf{b}_j) = \langle \boldsymbol{\nu}, \mathbf{b}_{ij} \rangle,$$

où

$$\mathbf{b}_{ij} = \frac{\partial \mathbf{b}_j}{\partial u_i} = \frac{\partial^2 \psi}{\partial u_i \partial u_j}$$

Les coefficients de la matrice de l'application de Weingarten dans la même base sont définis par

$$L(\mathbf{b}_i) = \frac{\partial \boldsymbol{\nu}}{\partial u_i} = \ell_{1i} \mathbf{b}_1 + \ell_{2i} \mathbf{b}_2.$$

Pour calculer cette matrice il est commode d'utiliser la relation $\mathbf{H} = -\mathbf{GL}$, qui implique

$$\mathbf{L} = -\mathbf{G}^{-1} \mathbf{H}.$$

Les courbures principales, de Gauss et moyenne de S en p sont les valeurs propres, le déterminant et la demi trace de $-L_p$. Le point p est ombilique si les deux courbures principales coïncident en ce point.