

En résolvant par rapport à a_2 et a_3 , on trouve

$$\frac{a_3}{a_2} = \frac{U_2(B_1A_2 - B_3) + U_3B_2}{U_2C_3 - U_3C_2}.$$

Il suffit maintenant que $U_3 \equiv 0$ pour que

$$\frac{a_3}{a_2} = \frac{B_1A_2 - B_3}{C_3}, \text{ d'où } \omega = \arctg \frac{\rho\tau^2 - \rho''}{\rho\tau_1 + 2\rho'}.$$

Un procédé analogue peut être appliqué dans le cas d'un contact d'ordre quelconque, et met en évidence des invariants de la forme $(1/2\pi)[\omega]_c$ d'ordre plus grand. Le comportement de ces invariants nous est encore inconnu.

LITTÉRATURE

1. K. Reidemesiter *Knotentheorie*. Berlin, 1932.
2. H. Seifert-Threlfall, *Lehrbuch der Topologie*. Berlin, 1934.
3. H. Seifert: *Algebraische Approximation von Mannigfaltigkeit*. Math. Zeitschrift, 41 (1936), 1-17.
4. G. Polya, G. Szegő, *Aufgaben u. Lehrsätze aus der Analysis*. T. II, '66, 230.
5. G. Călugăreanu, *L'intégrale de Gauss et l'analyse des noeuds tridimensionnels*. Revue de math. pures et appl., Bucarest, IV (1959) 5-20.
6. R. Lillienthal *Vorlesungen über Differentialgeometrie*. T. I, 255-257.

UN THÉORÈME ÉLÉMENTAIRE SUR LES NŒUDS ¹⁾

(C.R. Acad. Sci., 252 (1961), 2172-2173).

Nous appelons *noeud* toute courbe fermée simple, orientée, de l'espace euclidien tridimensionnel, ayant en chaque point une tangente qui varie continûment (en direction et sens) en fonction de l'arc. Deux noeuds C et C_1 étant représentés par $x(t), y(t), z(t)$ et $x_1(t), y_1(t), z_1(t)$ pour $0 \leq t \leq T$, une déformation continue de C en C_1 s'écrit $X = X(t, \lambda)$, $Y = Y(t, \lambda)$, $Z = Z(t, \lambda)$, $0 \leq \lambda \leq 1$, les fonctions X, Y, Z étant continues en (t, λ) et telles que $x(t) = X(t, 0), \dots, x_1(t) = X(t, 1), \dots$. Appelons *déformation isotope* une déformation continue telle que pour chaque $\lambda = 0, 1$ la courbe $C_\lambda(X, Y, Z)$ soit un noeud, et que la tangente à C_λ varie continûment pendant la déformation, ce qui revient à admettre que $\partial X / \partial t, \partial Y / \partial t, \partial Z / \partial t$ sont continues en (t, λ) dans le rectangle $t \in [0, T], \lambda \in [0, 1]$. Deux noeuds seront dits *isotopes* s'il existe une déformation isotope de l'un dans l'autre. L'isotopie étant une relation d'équivalence, elle divise l'ensemble des noeuds en *classes d'isotopie*. L'hypothèse de l'existence et continuité des dérivées du premier ordre, impliquée par nos définitions du noeud et de la déformation isotope, nous permet d'éviter la fusion des classes d'isotopie en une seule classe (équivalence de deux noeuds quelconques) ²⁾.

Le théorème suivant, dont la démonstration est immédiate, fournit un critère nécessaire pour que deux noeuds appartenant à des classes d'isotopie différentes :

I. C_1 et C_2 étant deux noeuds de classes différentes et $M_2 = h(M_1)$ une homéomorphie dérivable, à dérivée continue, entre C_1 et C_2 , il existe un couple M'_1, M'_2 de points de C_1 tel que les cordes correspondantes $M'_1 M'_2$ et $M_2 M'_2$ soient parallèles [$M'_2 = h(M'_1), M'_2 = h(M'_1)$]. Ceci a lieu quelle que soit la position relative des noeuds C_1 et C_2 , et quelle que soit l'homéomorphie h , à dérivée continue. Lorsque les noeuds C_1 et C_2 sont de même classe, on trouve facilement des exemples pour lesquels l'énoncé est en défaut (par exemple deux cercles égaux dont les plans sont parallèles, les points se correspondant par égalité des arcs, avec une différence constante, non nulle, des angles au centre). Pour la démonstration, joignons chaque couple $M_1, M_2 = h(M_1)$ par le segment rectiligne $M_1 M_2$. Nous obtenons une bande de surface réglée s'appuyant sur C_1 et C_2 . Sur chaque segment $M_1 M_2$ prenons le point μ qui partage ce segment dans le rapport $\lambda \in (0, 1)$. Les points μ forment une courbe fermée $\Gamma(\lambda)$, et cette courbe devra présenter des points multiples pour une valeur $\lambda = \lambda_0$, sans quoi C_1 et C_2 appartiendraient à une même classe d'isotopie, contrairement à l'hypothèse, car $\Gamma(\lambda)$ donnerait alors une déformation isotope de C_1 en-

¹⁾ Présentée par M. Arnaud Denjoy dans la Séance du 20 mars 1961.

²⁾ Voir à ce sujet une remarque de H. Seifert- W. Threlfall, *Lehrbuch der Topologie*, note 39, p. 329-

C_2 . Or, μ_0 étant un point multiple de $\Gamma(\lambda_0)$, par μ_0 doivent passer au moins deux segments $M'_1 M'_2$ et $M''_1 M''_2$. Le point μ_0 partage ces segments dans le même rapport λ_0 , et l'on en déduit la similitude des triangles $M'_1 M'_2 \mu_0$ et $M''_1 M''_2 \mu_0$. Les cordes correspondantes $M'_1 M'_2$ et $M''_1 M''_2$ sont donc parallèles.

Si C_1 et C_2 ont même longueur L (ce qu'on peut toujours réaliser sans changement des classes d'isotopie, par une homothétie appliquée à l'un des nœuds) on pourra choisir h de manière que les points de même abscisse curviligne se correspondent sur C_1 et C_2 . Pour chaque couple M'_1, M'_2 , la longueur de l'arc $M'_1 M'_2$ sera alors égale à celle de l'arc correspondant $M''_1 M''_2$, et l'on a l'énoncé :

II. C_1 et C_2 étant deux nœuds de même longueur, de classes différentes, il existe un couple M'_1, M'_2 de points de C_1 et un couple M''_1, M''_2 de points de C_2 , tels que les cordes $M'_1 M'_2$ et $M''_1 M''_2$ soient parallèles et sous-tendent des arcs égaux.

On doit remarquer que ce second énoncé, où l'homéomorphie h a été particularisée, pourra rester vrai pour des couples de nœuds de même classe. On se demande si le critère nécessaire offert par I est aussi suffisant pour que les nœuds soient de classes différentes, problème qui reste ouvert.

CONSIDÉRATIONS DIRECTES SUR LA GÉNÉRATION DES NŒUDS

(Revue Roum. Math. Pures et Appl., X, 4 (1965), 389—403)

1. Définitions et notations

Nous entendons par *nœud* [1] un contour polygonal fermé N , simple, orienté, à nombre fini de côtés, situé dans l'espace euclidien tridimensionnel E^3 . Appelons, avec M. W. Graeb [2], *application semi-linéaire* de E^3 une homéomorphie de E^3 sur lui-même qui, pour une certaine subdivision de E^3 en 3-simplexes, se réduit à une application linéaire affine dans chacun de ces 3-simplexes, l'orientation de E^3 étant conservée. Appelons *arc simple* toute image semi-linéaire d'un segment (1-simplexe) de E^3 , et *élément superficiel* e^2 toute image semi-linéaire d'un triangle (2-simplexe) de E^3 . L'intérieur de e^2 sera alors l'image de l'intérieur de ce 2-simplexe, et l'image de sa frontière sera appelée le *bord* de e^2 . Afin de simplifier le langage, nous employerons de préférence le terme *disque* pour élément superficiel.

N étant un nœud, soit e un disque dont le bord a en commun avec N un arc simple unique γ , tandis que l'intérieur de e est traversé par N en un nombre fini de points. Une *déformation* de N sur e est l'opération qui consiste à remplacer l'arc γ de N par l'arc γ' complémentaire de γ sur le bord de e , γ' étant orienté convenablement.

A cette opération nous pouvons associer la déformation continue de γ en γ' sur le disque e , et il y a lieu de remarquer que cette déformation peut toujours être faite de manière que la courbe déformée ne passe jamais deux fois par un même point de l'espace. Cette opération sera encore appelée *addition* du disque e au nœud N . La déformation est *isotope* si l'intérieur de e ne contient aucun point de N (le bord de e ayant en commun avec N l'arc simple unique γ). La déformation sera appelée une *autotraversée* du nœud N lorsque l'intérieur de e est traversé $\sigma > 0$ fois par N . Deux nœuds N et N' sont isotopes si l'on peut passer de l'un à l'autre par un nombre fini de déformations isotopes, et nous écrirons alors $N \sim N'$. Cette relation d'équivalence partage l'ensemble des nœuds en *types* ou *classes d'isotopie*, et nous désignerons par $[N]$ la classe dont N est un représentant. Nous désignerons par $[0]$ la classe du cercle, que nous appellerons *classe nulle*, et nous appellerons *cercle* chaque nœud de la classe $[0]$. Le bord d'un disque est donc un cercle; on sait que la réciproque est vraie. On sait encore ([2], p. 37) que cette définition de l'isotopie équivaut à la suivante: on a $N \sim N'$, lorsqu'il existe une application semi-linéaire de E^3 sur lui-même qui transforme N en N' , et alors seulement. Nous employerons encore l'extension suivante de la définition donnée plus haut: nous appellerons *addition* de deux surfaces orientables S et S' l'opération qui consiste à mettre en contact les bords de ces surfaces le long de plusieurs arcs disjoints, de manière que S et S' induisent sur chacun de ces arcs des orientations opposées. Nous désignerons par $S + S'$ la surface qui en résulte.