

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 $M'_1 M'_1' \mu_0$ et $M'_2 M'_2' \mu_0$. Les cordes correspondantes $M'_1 M'_1'$ et $M'_2 M'_2'$ 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 $M'_1 M'_1'$ sera alors égale à celle de l'arc correspondant $M'_2 M'_2'$, 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'_1' de points de C_1 et un couple M'_2, M'_2' de points de C_2 , tels que les cordes $M'_1 M'_1'$ et $M'_2 M'_2'$ 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 appelons *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.

Nous appelons *transversale* d'un nœud N un arc simple (ligne polygonale) qui n'a en commun avec N que ses deux extrémités α et β , supposées distinctes. Deux transversales T et T' de N sont isotopes lorsqu'il existe une application semi-linéaire de E^3 sur lui-même qui transforme T en T' (sans égard à l'orientation), tout en laissant inchangé le nœud N dans son ensemble (N pouvant glisser sur lui-même). Cette relation d'isotopie partage l'ensemble des transversales en *classes de transversales*, deux transversales d'une même classe étant isotopes.

Une transversale T décompose le nœud N en deux nœuds N_1 et N_2 de la manière suivante : les extrémités α et β de T partagent N en deux arcs C_1 et C_2 qui se trouvent orientés, du fait que N est orienté. L'arc C_1 avec la transversale T orientée convenablement forme un nœud N_1 , et de même C_2 avec T orientée inversement forme un nœud N_2 . Ce sont les deux nœuds *générés* par la transversale T dans N .

Theoreme I (Kinoshita-Terasaka). *Chaque nœud possède une transversale qui le décompose en deux cercles.*

Ce théorème a été énoncé par MM. Kinoshita et Terasaka ([3], annexe), qui en ont donné une démonstration par induction.

Le théorème affirme l'existence d'une transversale T telle que, si N_1 et N_2 sont les nœuds générés par T dans N , N_1 pris séparément est un cercle, et de même N_2 . Nous donnons ici une démonstration intuitive assez simple, en obtenant aussi un procédé de construction d'une transversale qui décompose N en deux cercles, lorsque N est donné par une projection dans le plan. Nous appellerons *transversale DC* toute transversale qui décompose N en deux cercles (les initiales *DC* rappellent « décomposition en cercles »).

Figurons la projection γ du nœud N dans le plan AOy (fig. 1), et supposons que la courbe fermée ainsi obtenue ne possède que des points doubles, aucun point double ne coïncidant avec la projection d'un sommet de N ; on sait [1] que ces conditions peuvent toujours être réalisées par un choix convenable du plan de projection AOy . Le dessin de γ laisse voir en chaque point double de cette courbe quelle est la branche de N qui passe en dessous de l'autre; c'est celle qui est figurée par un trait interrompu au voisinage du point double. Ce détail nous précise l'ordre de passage au point double apparent en question (indication de la branche qui passe en dessous de l'autre). Construisons le cylindre ayant γ pour base et les génératrices parallèles à Oz . Sur ce cylindre traçons une courbe (polygonale) $\Gamma = A'B'C'D'E'F'A''A'$, qui est une hélice de ce cylindre, et que nous fermerons par un segment $A'A''$ parallèle à Oz . A étant un point d'abscisse minima sur γ , on voit que la projection de Γ sur AOz n'a aucun point double, donc Γ est un cercle.

La projection de Γ sur AOy coïncide avec γ , mais, comparée à la projection de N , elle en diffère par l'ordre des passages en certains points doubles (B, C, D, E, F sur la fig. 1). Pour reconnaître ces points, on a le critère suivant : en parcourant la projection de Γ à partir de A , dans le sens qui correspond à un point qui monte le long de Γ , chaque fois que l'on arrive pour la première fois à un point double de γ , la branche que l'on suit passe en dessous de l'autre (c'est une conséquence du fait que Γ monte toujours de A' à A'').

Pour transformer Γ en un nœud isotope à N , il suffit de rétablir l'ordre des passages en certains points doubles de sa projection, conformément au schéma du nœud N . Nous réaliserons cette opération par la construction suivante : B' et B'' (fig. 1) étant les points de Γ qui se pro-

jettent au point double B , où l'ordre de passage des branches de Γ doit être inversé, construisons un parallélogramme situé sur le cylindre projetant, dont la base B_1B_2 est un petit segment de Γ contenant B' à son intérieur, ce parallélogramme contenant à son intérieur le segment demi-fermé ($B'B''$); il n'aura donc en commun avec Γ que le segment B_1B_2 et le point B'' . Dans ce qui suit, un tel parallélogramme sera appelé un

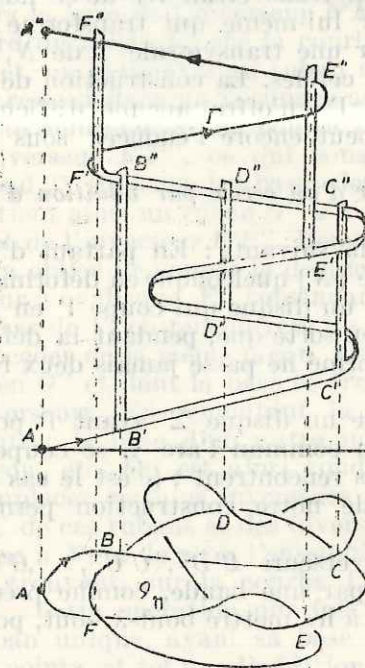


Fig. 1

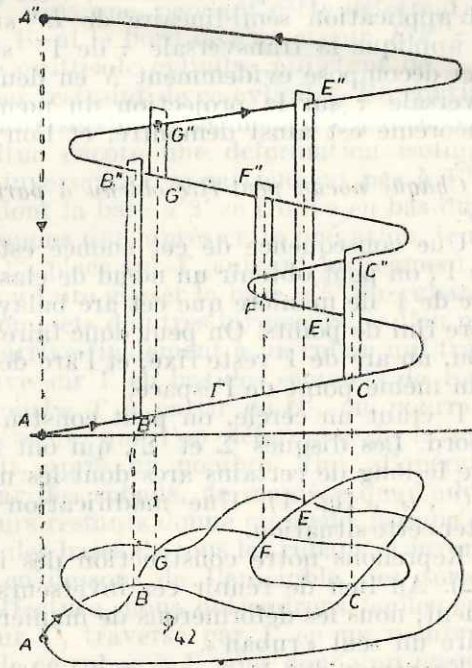


Fig. 2

inverseur. Nous ferons une construction analogue pour chacun des points doubles, C, D, E, F, G (fig. 2), où l'ordre des passages doit être inversé, et nous prendrons ces inverseurs assez minces pour qu'ils soient disjoints deux à deux. On voit alors que l'addition à Γ des inverseurs ainsi construits, convenablement orientés, fournit une courbe Γ' ayant exactement la projection de N , avec les mêmes ordres de passage aux points doubles de γ ; donc Γ' est isotope à N .

Remarquons que les bases de tous ces inverseurs sont « en bas » par rapport à leurs sommets $B'', C'', D'', E'', F'', G''$, car en chacun des points doubles B, C, D, E, F, G de γ on a à remplacer un passage en dessous par un passage en dessus. Joignons maintenant le premier inverseur, placé en B_1B_2 , au dernier, placé en F_1F_2 , par une bande s'appuyant sur Γ , située sur le cylindre projetant, au-dessus de Γ (fig. 1).

Prenons cette bande assez mince pour que, si elle traverse certains inverseurs, elle les traverse intérieurement, au voisinage de leurs sommets. Tous les inverseurs se trouvent alors réunis en une surface connexe Σ . On voit que le bord de cette surface est un cercle, car, en considérant la surface Σ indépendamment de la courbe Γ , on peut commencer par

«retirer» le dernier inverseur $F''F''$ et la bande $F'B''$ à travers l'inverseur $B'B''$, en supprimant ainsi l'intersection de la bande avec cet inverseur. C'est une déformation isotope du bord de la surface Σ . En continuant à retirer ainsi la surface $F''F''D''$ à travers l'inverseur $D'D''$, et ainsi de suite on obtient un élément superficiel dont le bord est manifestement un cercle. Le bord de la surface Σ elle-même est donc un cercle K . Ce cercle K a en commun avec le cercle Γ un arc unique τ , et l'on voit que $K \cup \Gamma - \tau$ est un nœud Γ'' isotope à N . Ainsi, τ est une transversale DC de ce nœud Γ'' . L'application semi-linéaire de E^3 sur lui-même qui transforme Γ'' en N applique la transversale τ de Γ'' sur une transversale T de N , et celle-ci décompose évidemment N en deux cercles. La construction de la transversale τ sur la projection du nœud Γ'' n'offre aucune difficulté. Le théorème est ainsi démontré, et l'on peut encore l'énoncer sous la forme :

Chaque nœud peut être obtenu à partir d'un cercle par addition d'un disque.

Une conséquence de cet énoncé est la suivante : En partant d'un cercle Γ , on peut obtenir un nœud de classe $[N]$ quelconque en déformant un arc de Γ de manière que cet arc balaye un disque qui coupe Γ en un nombre fini de points. On peut donc faire en sorte que, pendant la déformation, un arc de Γ reste fixe, et l'arc déformé ne passe jamais deux fois par un même point de l'espace.

Γ étant un cercle, on peut construire un disque Σ' ayant Γ pour son bord. Les disques Σ et Σ' , qui ont en commun l'arc τ , se coupent encore le long de certains arcs dont les uns rencontrent τ (c'est le cas en B'' , C'' , D'' , fig. 1). Une modification de notre construction permet d'éviter cette situation.

Reprenons notre construction des inverseurs $B'B''$, $C'C''$, ..., $G'G''$ (fig. 2). Au lieu de réunir ces inverseurs par une bande, comme précédemment, nous les déformerons de manière à les mettre bout-à-bout, pour en faire un seul « ruban ».

En remontant la courbe Γ à partir de A' , on rencontre des bases et des sommets d'inverseurs, et, en général, il y aura des inverseurs dont les bases sont situées plus haut sur Γ que les sommets d'autres inverseurs (c'est le cas pour $G'G''$, fig. 2). Nous ferons en sorte que chaque base d'un inverseur se trouve plus bas sur Γ que les sommets de tous les inverseurs.

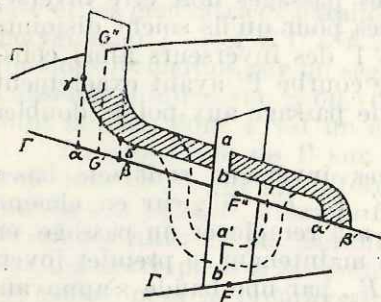


Fig. 3

En remontant la courbe Γ à partir de A' , on rencontre des bases et des sommets d'inverseurs, et, en général, il y aura des inverseurs dont les bases sont situées plus haut sur Γ que les sommets d'autres inverseurs (c'est le cas pour $G'G''$, fig. 2). Nous ferons en sorte que chaque base d'un inverseur se trouve plus bas sur Γ que les sommets de tous les inverseurs.

Pour y arriver, nous déformerons ces inverseurs de manière que la courbe Γ , qui forme un nœud isotope à N , ne se traverse pas elle-même, en restant donc isotope à N . Dufait de ces déformations, les inverseurs déformés prendront l'aspect de rubans à structure simpliciale. En remontant la courbe Γ à partir de A' , arrêtons-nous à la première base d'inverseur qui est précédée par des sommets d'autres inverseurs ($G'G''$, fig. 2). « Faire glisser » la base $\alpha\beta$ de l'inverseur $G'G''$ vers le bas

nous appellerons désormais un ruban, la partie $\alpha\beta\gamma\delta$ de l'inverseur initial $G'G''$ ayant été supprimée. Il est clair que cette déformation de l'inverseur $G'G''$ est une déformation isotope du nœud Γ' tant que le bord du ruban ne touche le bord d'aucun autre inverseur, et ne touche Γ que le long de sa base $\alpha'\beta'$. On pourra donc faire glisser $\alpha'\beta'$ en bas du sommet F'' , le ruban traversant l'inverseur $F''F''$ suivant un segment ab intérieur à ce dernier. En maintenant fixe la base $\alpha'\beta'$, nous déformerons maintenant le ruban de manière que le segment d'intersection ab glisse en $a'b'$ (fig. 3) vers la base de l'inverseur $F''F''$, sans que, pendant cette déformation, le bord du ruban touche la courbe Γ , ni le bord de l'inverseur $F''F''$; à cet effet, on permettra au ruban de quitter le cylindre projetant de Γ , mais en restant dans un voisinage assez restreint de ce cylindre. En continuant cette déformation, on pourra « faire sortir » le segment $a'b'$ par la base de l'inverseur $F''F''$, ce qui constitue encore une déformation isotope du nœud Γ' (puisque les bases des inverseurs n'appartiennent pas à Γ'). On obtient ainsi un ruban $G''\alpha'\beta'$, dont la base $\alpha'\beta'$ se trouve en bas du sommet de l'inverseur $F''F''$. Remarquons que, après cette opération, le ruban sera obligé de couper le disque Σ' , placé sur Γ , suivant un segment intérieur à ce disque. En continuant à faire glisser la base $\alpha'\beta'$ du ruban vers le bas de Γ , on rencontrera les sommets d'autres inverseurs, et l'on pourra procéder de la même façon. On arrive finalement à un ruban qui traverse Γ en G'' et dont la base se trouve sur Γ en bas des sommets de tous les inverseurs. En remontant la courbe Γ à partir de G' , on pourra rencontrer la base d'un autre inverseur, que l'on déformera de la même façon, etc. On est ainsi conduit, après un nombre fini d'opérations, à remplacer certains inverseurs par des rubans, de manière que l'addition à Γ de ces rubans et des inverseurs restants donne un nœud isotope à Γ' , donc à N , et de plus, l'ensemble des bases de tous les rubans et inverseurs se trouvent, sur la courbe Γ , en dessous de l'ensemble des sommets.

Cette opération une fois effectuée, nous obtiendrons facilement un ruban unique, ayant sa base sur Γ , traversé par Γ en un nombre fini de points, et tel que l'addition de ce ruban à Γ nous donne un nœud isotope à N .

Partons de A' et remontons la courbe Γ . En passant sur la base du premier inverseur rencontré (B' , fig. 2), arrivons à la base du second (C'). En faisant glisser la base de $C'C''$ le long de Γ vers B' , puis sur le bord de l'inverseur $B'B''$ (fig. 4), on pourra amener cette base au sommet de l'inverseur $B'B''$. On obtient ainsi un ruban $B'B''C''$, de base B_1B_2 , qui traverse Γ en B'' et C'' . On pourra ensuite passer au troisième inverseur (ou ruban) $D'D''$, faire glisser sa base vers B' , puis le long du bord du ruban $B'B''C''$, jusqu'à ce que la base de $D'D''$ soit amenée au sommet de ce ruban, et ainsi de suite, en épuisant tous les inverseurs et rubans. Le ruban final R sera donc sans autosections¹⁾, et traversera le disque Σ' placé sur Γ suivant des arcs de deux espèces :

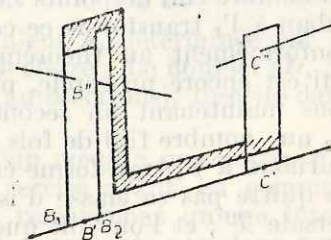


Fig. 4

1. Arcs pq appartenant à l'intérieur de Σ' , tandis que p et q appartiennent au bord de R . C'est ce que nous appellerons une traversée intérieure.

¹⁾ Nous employons le terme « autosection » pour « intersection avec lui-même ».

2. Arcs pq tels que p est point intérieur sur Σ' et appartient au bord de R , tandis que q appartient au bord de Σ' et est un point intérieur du ruban R . Une telle traversée sera appelée *traversée marginale*.

La construction que nous venons d'indiquer démontre encore une fois le théorème I, puisque le ruban R est un disque dont le bord a en commun avec Γ le segment B_1B_2 . Ce segment est ici la traversée qui décompose le nœud en deux cercles. Mais on voit cette fois que les disques Σ' et R n'ont en commun, dans un voisinage assez restreint du segment B_1B_2 , aucun point qui n'appartienne à ce segment.

Une application semi-linéaire de E^3 sur lui-même, qui transforme Γ' en N , transforme Γ en un cercle Γ_1 , et le ruban R en un ruban R_1 ayant sa base sur Γ_1 ; ainsi, N peut être obtenu par addition d'un ruban R_1 à un cercle Γ_1 .

Théorème II. On peut passer d'un nœud N à un nœud de classe $[N']$ quelconque par addition à N d'un seul disque qui traverse N un nombre fini de fois.

Le nœud N étant donné, construisons $\Gamma' \sim N$ comme précédemment, à l'aide du cercle Γ et d'un ruban R . Sur le ruban R plaçons un autre ruban R' , ayant sa base au sommet de R , et assez mince pour qu'il s'étale le long de R tout en restant à son intérieur (sauf sa base), et de manière qu'il traverse Γ aux mêmes points que R . L'addition de R' à Γ' ramène ce nœud dans la classe $[0]$. Or, Γ' étant isotope à N , il existe une application semi-linéaire de E^3 qui transforme Γ' en N . Cette même application transforme R' en un ruban R_2 appliqué à N , tel que l'addition de R_2 à N fournit un nœud N_1 de classe nulle. Conformément au théorème I, on peut passer du cercle N_1 au nœud N' par l'addition d'un ruban R_3 , dont la base peut être placée sur un arc quelconque de N_1 . Plaçons cette base au sommet de R_2 . On obtient ainsi un seul ruban dont l'addition à N fournit un nœud de classe $[N']$, et le théorème est démontré.

Ceci nous permet une généralisation du théorème de Kinoshita-Terasaka :

Théorème III. $[N_1]$ et $[N_2]$ étant deux classes d'isotopie choisies à volonté, le nœud N possède une transversale qui le divise en deux nœuds, l'un appartenant à $[N_1]$ et l'autre à $[N_2]$.

Soit T une transversale qui divise N en deux cercles Γ_1 et Γ_2 . Construisons un ruban ayant sa base sur T et qui traverse Γ_1 en un nombre fini de points sans toucher Γ_2 , de manière que l'addition de ce ruban à Γ_1 transforme ce cercle en un nœud $\Gamma'_1 \in [N_1]$, ce qui est possible conformément au théorème II. Ceci transforme T en T' et Γ_2 en Γ'_2 qui est encore un cercle, puisque le ruban ne traverse pas Γ_2 . Construisons maintenant un second ruban ayant sa base sur T' , qui traverse Γ_2 un nombre fini de fois sans toucher Γ_1 , de manière que l'addition de ce ruban à Γ'_2 transforme ce cercle en un nœud $\Gamma'_2 \in [N_2]$. De ce fait, Γ'_1 ne quitte pas sa classe d'isotopie, T' est remplacée par une autre transversale T'' , et l'on voit que T'' répond aux conditions de l'énoncé.

2. Génération des nœuds par un doublet

Il résulte des théorèmes précédents que, $[N]$ étant une classe d'isotopie, on obtient un représentant de cette classe en ajoutant à un disque un ruban sans autosections (donc un second disque) qui traverse le premier disque un nombre fini de fois, marginalement et intérieurement. Or, on peut toujours éliminer les traversées intérieures en augmentant le

nombre des traversées marginales, ce qui a l'avantage de nous ramener à un seul genre de traversées. Pour vérifier la possibilité de cette transformation, il suffit de regarder la figure 5 qui montre que, en repliant convenablement la surface du ruban, puis en y pratiquant deux entailles (hachurées sur la figure) on arrive à remplacer une traversée intérieure par deux traversées marginales qui se succèdent sur la longueur du ruban dans un ordre que l'on peut choisir à volonté.

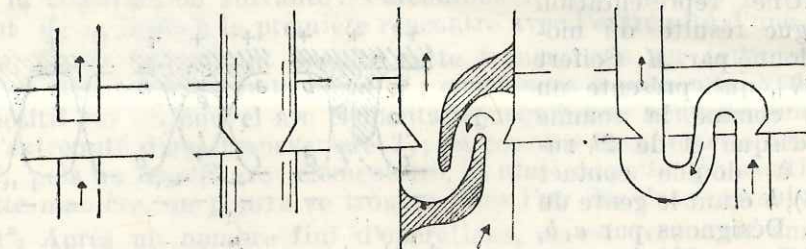


Fig. 5

Supposons donc que, l'élimination des traversées intérieures étant effectuée, on est ramené à un disque au bord duquel est collé un ruban sans autosections, qui traverse le disque marginalement un nombre fini de fois. Nous appelons *doublet* un pareil assemblage, et l'on voit que la distinction entre disque et ruban n'a rien d'essentiel (après suppression des traversées intérieures), ces deux éléments jouant un rôle symétrique dans le doublet. On peut donner au doublet une forme canonique en prenant pour disque D un rectangle allongé $abcd$ (fig. 6), auquel est attaché le ruban R , le long d'un côté bc du rectangle. Toute classe de nœuds correspond à un doublet (D, R) au moins. On voit que le nombre des segments d'intersection disque-ruban peut être augmenté à volonté, sans changer la classe d'isotopie du nœud N formé par le bord du doublet. Mais ce nombre ne peut être diminué au delà d'une valeur σ , déterminée par la classe $[N]$. σ est un invariant d'isotopie de cette classe. Le doublet (D, R) forme une surface orientable à singularités (autosections), inscrite dans le nœud N .

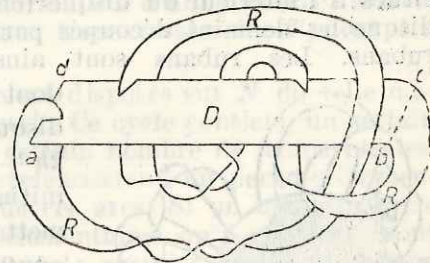


Fig. 6

On peut passer de cette représentation d'un nœud à une autre, qui consiste à réduire le nœud à la somme de deux cercles ayant en commun un nombre impair d'arcs disjoints. En effet, remarquons qu'une intersection marginale des disques D et R peut être remplacée par deux nouveaux rubans à double contact. Ceci est visible sur la fig. 7.

Les dessins I concernent le cas d'une intersection marginale suivant le segment ef , l'une des deux parties présentant sa face positive marquée du signe $+$, et l'autre sa face négative. En coupant ces deux parties le long des arcs aeb et efd , on pourra remplacer les parties $acbf$ et $cfde$ par deux rubans tordus $aecfa$ et $bfdeb$. On voit que par cette modification la surface $D + R$ reste orientable, et son bord n'aura subi aucun chan-

gement. Les dessins II concernent le cas où les deux parties en intersection présentent leurs faces positives. En appliquant une torsion à l'une des deux parties, on se ramène au cas précédent et, après application du procédé que nous venons d'indiquer, on refera en sens inverse la torsion appliquée. On obtient finalement deux rubans croisés sans torsion, qui remplacent l'intersection marginale. En appliquant cette modification, la surface obtenue sera la somme de deux disques ayant en commun $2\sigma + 1$ arcs.

Une représentation analogue résulte du modèle donné par M. Seifert [5], [7], qui présente un nœud comme la somme d'un disque et de $2h$ rubans à double contact (fig. 8), h étant le genre du nœud. Désignons par a_i, b_i ($i = 1, 2, \dots, 2h$) les segments de contact des rubans avec le bord du disque, dans leur ordre de succession sur ce bord. Joignons a_1 à b_1 , a_2 à b_3 , a_4 à b_5 , $\dots$, a_{2h-2} à b_{2h-1} , a_{2h} à b_{2h} , par des arcs simples

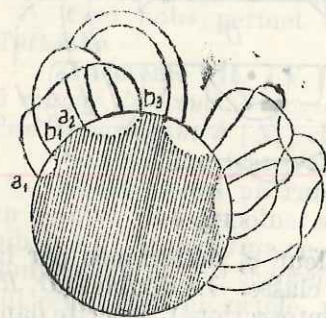


Fig. 8

situés à l'intérieur du disque (en pointillé sur la fig. 8). Retranchons du disque les domaines découpés par ces arcs et ajoutons ces domaines aux rubans. Les rubans sont ainsi incorporés à un disque unique dont le bord a en commun avec le disque amputé (hachuré sur la fig. 8) les $2h + 1$ arcs simples tracés sur le disque initial. Or, h étant le nombre minimum permettant une telle représentation, il résulte que $2h + 1 \leq 2\sigma + 1$, donc $h \leq \sigma$. On peut passer inversement du modèle de M. Seifert à un doublet à intersections marginales (fig. 9), ce qui permet de calculer une limite supérieure de σ en fonction des torsions des rubans sur le modèle de M. Seifert.

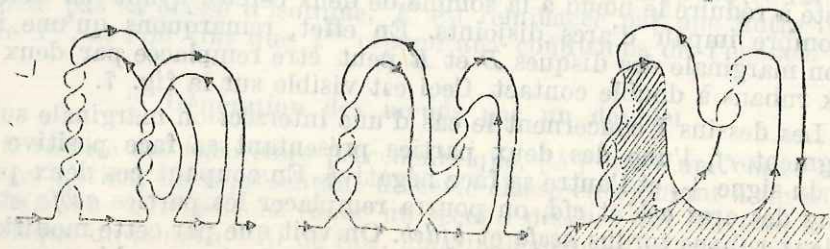


Fig. 9

3. Transversales et cycles sur un nœud

Considérons k transversales deux à deux disjointes du nœud N : $T_1, T_2, \dots, T_k$, et désignons par A_i, B_i les extrémités de T_i , $i = 1, 2, \dots, k$. Les $2k$ points A_i, B_i partagent N en $2k$ arcs, que nous appellerons arcs élémentaires, chacun de ces arcs ayant pour extrémités deux des points A_i, B_i et laissant à son extérieur tous les autres points A_i, B_i . Le nœud N étant orienté, prenons un point A sur l'un des arcs élémentaires et effectuons la construction suivante: Parcourons N dans le sens positif, en partant de A , jusqu'à la première rencontre avec l'extrémité d'une transversale T_i ; en parcourant ensuite cette transversale sur toute sa longueur on sera ramené sur le nœud N ; continuons à parcourir N dans le sens positif sur un nouvel arc élémentaire jusqu'à une nouvelle rencontre avec l'extrémité d'une transversale T_j ; parcourons ensuite cette transversale T_j , puis un nouvel arc élémentaire, et ainsi de suite. En continuant de cette manière, on pourra se trouver dans l'un des deux cas suivants:

1° Après un nombre fini d'opérations, on est reconduit sur l'arc élémentaire qui contient le point A , après avoir parcouru une fois un certain nombre d'arcs élémentaires et de transversales T_i . On obtient alors une courbe fermée simple que nous appellerons le cycle du point A . Il est clair que cette courbe est déterminée d'une manière univoque par un point quelconque de l'un des arcs élémentaires qu'elle contient.

2° Par l'application du procédé on se trouve conduit à parcourir une seconde fois l'une des transversales T_i , avant d'arriver au point initial A . Si l'on s'interdit de parcourir plus d'une fois une même transversale, on n'arrive pas, dans ce cas, à la formation d'un cycle (ceci arrive, par ex., dans le cas d'un cercle C avec deux transversales T_1, T_2 dont les extrémités sont entrelacées sur C).

Supposons que les points A_i, B_i soient disposés sur N de telle manière que la formation du cycle de A réussisse. Ce cycle contient un certain nombre d'arcs élémentaires de N et un certain nombre de transversales. Si des arcs élémentaires extérieurs à ce cycle existent, on peut recommencer la construction en partant de l'un de ces arcs. Si un nouveau cycle est ainsi obtenu, on voit que les arcs élémentaires qu'il contient sont différents des arcs élémentaires situés sur le premier cycle; le second cycle peut contenir un certain nombre des transversales appartenant au premier. Ainsi, deux cycles n'ont en commun qu'un certain nombre de transversales T_i , s'ils ne sont pas disjoints. Si la formation d'un troisième cycle est possible, on voit que ce cycle ne peut contenir aucune des transversales communes aux deux premiers, les arcs élémentaires qui aboutissent aux extrémités d'une telle transversale appartenant aux deux premiers cycles. Donc, chaque transversale T_i appartient à deux cycles au plus.

Si, en appliquant le procédé, on arrive à décomposer N en cycles, chaque arc élémentaire appartenant donc à un cycle et à un seul, chaque transversale T_i devra appartenir à deux cycles exactement, puisque chaque extrémité d'une transversale est extrémité commune de deux arcs élémentaires de N ; de plus, chaque cycle étant orienté en concordance avec N , on voit que, sur les deux cycles qui contiennent une transversale T_i , celle-ci se trouve orientée de deux manières opposées.

Théorème IV. Pour que le système de transversales $\{T_i\}$, $i = 1, 2, \dots, k$, décompose N en cycles, il est nécessaire et suffisant que, en plaçant

les points A_i, B_i dans leur ordre de succession sur N , à partir d'une origine arbitraire, les points A_i, B_i occupent dans cette suite des rangs de parité différente, et ceci pour chaque $i = 1, 2, \dots, k$.

Notre démonstration est une adaptation du raisonnement employé par J. v. Nagy [4] pour démontrer un théorème de Gauss relatif aux courbes planes à points doubles.

La condition de parité indiquée dans l'énoncé est nécessaire. Supposons donc que N se trouve décomposé en cycles par le système $\{T_i\}$. Considérons l'une de ces transversales T_i . Ses extrémités A_i, B_i décomposent N en deux arcs C_1 et C_2 . Parmi les transversales $T_j, j \neq i$, nous en aurons n dont les deux extrémités se trouvent sur C_1 , et m transversales dont une seule extrémité est située sur C_1 , l'autre étant donc sur C_2 ; nous dirons alors que cette transversale est « à cheval » sur C_1 et C_2 . On voit alors que l'arc C_1 contient $2n + m + 1$ arcs élémentaires. Or, m est un nombre pair. En effet, T_j étant une transversale dont l'extrémité A_j est sur C_1 et l'autre B_j est sur C_2 , il existe un cycle qui contient T_j avec son orientation $A_j B_j$. Il existe donc un arc élémentaire $\alpha_1 \subset C_1$ et un autre $\alpha_2 \subset C_2$, ces deux arcs ayant chacun une extrémité commune avec T_j . En parcourant ce cycle à partir d'un point de α_1 , on devra bien revenir sur α_1 ; ce cycle contient donc une seconde transversale qui est à cheval sur C_1 et C_2 , soit T_l , avec $l \neq j$. S'il existe une troisième transversale à cheval sur C_1 et C_2 , appartenant au même cycle, on voit de même qu'il existe nécessairement une quatrième, distincte des autres, et ainsi de suite. Ainsi, les transversales d'un cycle qui sont à cheval sur C_1 et C_2 sont en nombre pair. Chaque arc élémentaire appartenant à C_1 étant situé sur un cycle et un seul, il en résulte que m est un nombre pair, donc $2n + m + 1$ est impair, ce qui montre que la condition de l'énoncé est nécessaire.

La condition est suffisante. Choisissons un arc élémentaire de N , que nous teindrons en rouge. En parcourant N dans le sens positif, nous teindrons alternativement en blanc et en rouge les arcs élémentaires ainsi rencontrés. Le nombre total de ces arcs étant $2k$, on finira par teindre en rouge k de ces arcs, les k autres étant teints en blanc. La condition de parité de l'énoncé étant satisfaite, appliquons le procédé de construction à partir d'un point A placé sur un arc rouge. En parcourant cet arc, on arrivera à une extrémité de la transversale T_i , soit A_i , puis à son extrémité B_i . Ces points étant séparés, sur N , par un nombre impair d'arcs élémentaires, on voit que l'on sera amené à parcourir, à partir de B_i , un nouvel arc rouge, et ainsi de suite. Le cas 2° est exclu, on ne peut être conduit à parcourir une seconde fois une même transversale, puisque les deux arcs qui aboutissent à une extrémité d'une transversale sont coloriés différemment. Le cycle se ferme donc en A , et le théorème est démontré.

Remarquons encore que, si l'on change le sens de parcours sur N , la décomposition de N en cycles sera la même, seule l'orientation de ces cycles devant être inversée.

En somme, si la condition de l'énoncé est satisfaite, N se trouve décomposé en cycles, chaque arc élémentaire appartenant à un seul cycle, et chaque transversale appartenant à deux cycles, avec des orientations opposées.

Il est clair, d'après ce théorème, que la forme et la position relative des transversales ne comptent pour rien, s'il est question de la possibilité de décomposition en cycles seulement.

Au contraire, si l'on veut que ces cycles soient tous des cercles, la position relative des transversales T_i et leur position par rapport au nœud N , jouent un rôle essentiel. Le théorème de Kinoshita-Terasaka [3] affirme l'existence d'une transversale T qui décompose N en deux cycles, chacun de ces cycles étant un cercle. En plaçant un disque sur chacun de ces cercles, ces deux disques donneront un certain nombre de lignes d'intersection, à moins que N soit un cercle; c'est ce que nous avons appelé un doublet. Mais en employant plusieurs transversales de N , on peut construire ainsi une surface orientable sans singularités ayant pour bord le nœud N .

Supposons que les transversales $\{T_i\}, i = 1, 2, \dots, k$, décomposent N en cycles, mais ces cycles ne sont pas des cercles. On peut alors remplacer les transversales $\{T_i\}$ par d'autres transversales $\{T'_i\}$ de mêmes extrémités, qui décomposent N en cycles circulaires. En effet, soit Γ_1 l'un des cycles déterminés par les $\{T_i\}$, et T_1 l'une des transversales appartenant à Γ_1 . On peut déformer Γ_1 le long d'un ruban r ayant sa base sur T_1 et qui ne traverse que les arcs de N situés sur Γ_1 , de manière que le cycle ainsi déformé soit un cercle Γ'_1 . La transversale T_1 se trouve remplacée par T'_1 , qui résulte de T_1 par addition du ruban r . Remarquons que cette déformation de T_1 en T'_1 ne change pas la classe d'isotopie des cycles différents de Γ_1 , puisque le ruban placé sur T_1 ne coupe que les arcs de N situés sur Γ_1 , qui n'appartiennent pas à d'autres cycles. On peut donc reprendre l'opération pour les autres cycles, en obtenant finalement un système $\{T'_i\}$ de transversales qui décomposent N en cycles circulaires.

Sur chacun de ces cycles plaçons un disque. Peut-on choisir ces disques de manière qu'ils forment une surface orientable sans singularités, inscrite au nœud N ? La possibilité d'une telle construction résulte du procédé indiqué par M. Seifert [5], pour $k = d$ (où d représente le nombre des points doubles d'une projection plane de N), qui a conduit cet auteur à la définition du genre h d'un nœud. Mais en général, on a $k \leq d$.

4. Surfaces fermées orientables contenant un nœud

Soit S une surface orientable, sans singularités, ayant le nœud N pour bord. Grâce à l'orientabilité de S , on peut donner à cette surface une petite déformation dans le sens des normales positives, de manière à obtenir une seconde surface S' , ayant encore N pour bord, et ne rencontrant S qu'aux points de ce bord. S et S' forment alors une surface fermée orientable Σ , sans singularités, donc une sphère à anses en nombre fini. N divise la surface Σ en deux domaines S et S' , homéomorphes l'un à l'autre. Il existe une infinité de surfaces telles que Σ , de genres aussi grands que l'on veut, puisque l'on peut ajouter simultanément à S et S' un même nombre arbitraire de nouvelles anses. Désignons par Σ^* la surface Σ de genre minimum sur laquelle N se trouve tracé de manière à diviser Σ^* en deux domaines homéomorphes (si cette dernière condition était abandonnée, le genre de Σ^* pourrait être diminué; à ce sujet, voir une construction indiquée dans [6], pp. 135-137).

Dans le cas du modèle dû à M. Seifert [5, 7], chaque ruban présente, entre ses deux extrémités attachées au disque, des torsions paires quelconques (le bord d'un ruban présente, en projection plane, un nombre pair de points doubles), mais — ce qui complique davantage l'étude de

ce modèle — les rubans peuvent former des nœuds, et ces nœuds peuvent être enlacés entre eux (fig. 10a). La surface Σ^* correspondante sera donc une sphère à $2h$ anses, ces anses étant nouées et enlacées entre elles. Il est avantageux de remplacer Σ^* par une surface fermée orientable en position normale [6, p. 130], c'est-à-dire une sphère à anses non nouées et non enlacées. Cette simplification s'obtient facilement, au prix d'une augmentation du nombre des anses.

Pour y arriver, il suffit « d'aplatir » la surface Σ^* dans la direction perpendiculaire au plan de projection (fig. 10b), en permettant aux parties des anses qui se touchent de se joindre en un canal ramifié (à quatre branches). La figure 11 précise le détail de cette opération. Après que deux

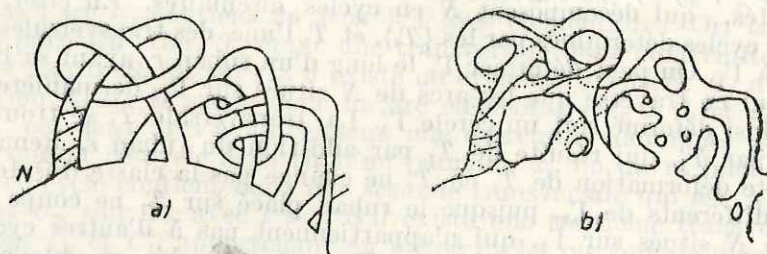


Fig. 10

branches de l'anse se soient touchées, on déplacera les arcs de N qui parcourent ces branches (fig. 11a, b) de manière que l'on puisse joindre ces deux branches en un canal à quatre branches sans que les arcs de N situés sur l'une des branches touchent les arcs de N situés sur l'autre. De plus, on peut faire en sorte que les parties des deux branches d'anses qui viennent en contact appartiennent à un même domaine S (ou S').

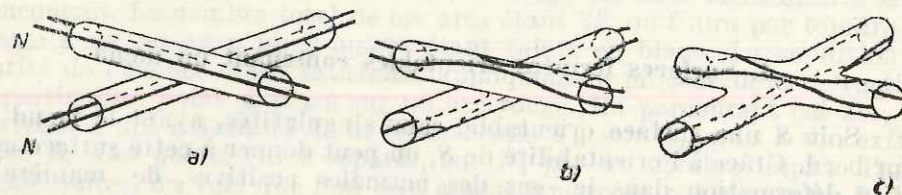


Fig. 11

La surface Σ^{**} ainsi obtenue sera une sphère à anses en position normale, sur laquelle le nœud N détermine deux domaines disjoints, l'un d'eux étant homéomorphe au disque à rubans initial. Nous pouvons dire encore que le modèle de M. Seifert peut être placé sur une sphère à anses en position normale. Soit Σ la sphère à anses en position normale de genre minimum p , sur laquelle on puisse tracer un nœud isotope à N , qui divise Σ en deux domaines. Le nombre p est un invariant d'isotopie de la classe $[N]$ que nous appellerons *genre normal* de cette classe. La construction des nœuds de genre normal p donné est un problème bidimensionnel, puisque l'on peut considérer une sphère à anses en position normale comme une variété assez simple, dont la structure est résumée par les propriétés du polygone canonique correspondant. C'est un point de vue que nous nous proposons d'analyser dans un autre travail.

1. K. Reidemeister, *Knotentheorie*. Springer, Berlin, 1932.
2. W. Graeb, *Die semilinearen Abbildungen*. Sitzungsab. d. Heidelberger Akad. d. Wiss., 4-te Abh., 1950.
3. S. Kinoshita, H. Terasaka, *On unions of knots*. Osaka math. Journ. 9, (1957), 131—153.
4. J. v. Sz. Nagy, *Über ein topologisches problem von Gauss*. Math. Zeitschr. 26, (1927), 579 — 592.
5. H. Seifert, *Über das Geschlecht von Knoten*. Math. Annalen, 110, (1934), 571—592.
6. H. Terasaka, *On the non-triviality of some kinds of knots*. Osaka math. Journ., 12, (1960), 113—144.
7. H. Seifert, *La théorie des nœuds*. L'enseignement math., 35, (1936), 201—212.