

THÈSES DE L'ENTRE-DEUX-GUERRES

TADÉ WAŻEWSKI

**Sur les courbes de Jordan ne renfermant aucune courbe
simple fermée de Jordan**

Thèses de l'entre-deux-guerres, 1923

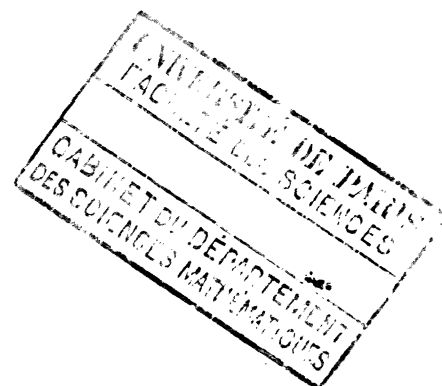
[<http://www.numdam.org/item?id=THESE_1923__37__1_0>](http://www.numdam.org/item?id=THESE_1923__37__1_0)

L'accès aux archives de la série « Thèses de l'entre-deux-guerres » implique l'accord avec les conditions générales d'utilisation (<http://www.numdam.org/conditions>). Toute utilisation commerciale ou impression systématique est constitutive d'une infraction pénale. Toute copie ou impression de ce fichier doit contenir la présente mention de copyright.

NUMDAM

*Thèse numérisée dans le cadre du programme
Numérisation de documents anciens mathématiques*
<http://www.numdam.org/>

Numéro d'ordre 135.



THÈSES

PRÉSENTÉES

A LA FACULTÉ DES SCIENCES
DE L'UNIVERSITÉ DE PARIS

POUR OBTENIR

LE TITRE DE DOCTEUR DE L'UNIVERSITÉ

PAR

TADÉ WAŻEWSKI

1^{re} THÈSE. — SUR LES COURBES DE JORDAN NE RENFER-
MANT AUCUNE COURBE SIMPLE FERMÉE DE
JORDAN.

2^{re} THÈSE. — PROPOSITIONS DONNÉES PAR LA FACULTÉ.

SONTENUES LE DEVANT LA COMMISSION
D'EXAMEN

MM. BOREL	. . .	<i>Président.</i>
MONTEL	}	. . . <i>Examineurs.</i>
DENJOY		

CRACOVIE 1923.
IMPRIMERIE DE L'UNIVERSITÉ.

FACULTÉ DES SCIENCES DE L'UNIVERSITÉ DE PARIS

MM.

Doyen MOLLIARD, *Professeur*. Physiologie végétale.

Doyen honoraire P. APPELL.

Professeurs honoraires. { P. PUISEUX.
VÉLAIN.
BOUSSINESQ.
PRUVOT.

	EMILE PICARD . . .	Analyse supérieure et algèbre supérieure.
	KOENIGS	Mécanique physique et expérimentale.
	GOUBSAT	Calcul différentiel et calcul intégral.
	HALLER	Chimie organique.
	JOANNIS	Chimie (Enseignement P. C. N.)
	JANET	Electrotechnique générale.
	WALLEKANT	Minéralogie.
	ANDOYER	Astronomie.
	PAINLEVÉ	Mécanique analytique et mécanique céleste.
	HAUG	Géologie.
	H. LE CHATELIER . .	Chimie générale.
	GABRIEL BERTRAND .	Chimie biologique.
	Mme P. CURIE . . .	Physique générale et radioactivité.
	CAULLERY	Zoologie (Evolution des êtres organisés).
	C. CHABRIE	Chimie appliquée.
	G. URBAIN	Chimie minérale.
	EMILE BOREL	Calcul des probabilités et Physique mathém.
	MARCHIS	Aviation.
	JEAN FERRIN	Chimie physique.
	ABRAHAM	Physique.
	CARTAN	Mécanique rationnelle.
	Cl. GUICHARD	Géométrie supérieure.
	LAPICQUE	Physiologie.
<i>Professeurs</i>	GENTIL	Géographie physique.
	VESSIOT	Théorie des groupes et calcul des variations.
	COTTON	Physique générale.
	DRACH	Application de l'analyse à la géométrie.
	C. FABRY	Physique.
	CHARLES PÉREZ . . .	Zoologie.
	LÉON BERTRAND . . .	Géologie appliquée et géologie régionale.
	DANGEARD	Botanique.
	LESPIEAU	Théories chimiques.
	LEDUC	Physique théorique et physique céleste.
	MONTÉL	Mathématiques générales.
	MAURAIN	Physique du globe.
	RABAUD	Biologie expérimentale.
	WINTREBERT	Anatomie et physiologie comparées.
	HÉROUARD	Zoologie.
	RÉMY PERRIER	Zoologie (Enseignement P. C. N.)
	SAGNAC	Physique théorique et physique céleste.
	PORTIER	Physiologie.
	BLAISE	Chimie organique.
	PÉCHARD	Chimie (Enseignement P. C. N.)
	AUGER	Chimie analytique.
	M. GUICHARD	Chimie minérale.
	GUILLET	Physique.

Secrétaire DANIEL TOMBECK.

À mon Maître

M. Jan Sleszyński

Professeur de l'Université de Cracovie

en hommage et reconnaissance.

PPN 089 244 287



Avant-Propos.

C'est avec une profonde gratitude que je présente mes remerciements à M. M. Denjoy, Fréchet et Lebesgue pour leurs précieux conseils et indications concernant ce travail et à mes Maîtres M. M. Hoborski et Zaremba pour avoir rendu possible son impression et pour leur encouragement.

A tous mes amis ayant contribué à mon progrès à différents points de vue et tout particulièrement à M^{lle} Jeanne Prohet et à M. Pierre Aubertin qui s'est gracieusement chargé de la revision des épreuves j'offre mes sentiments les meilleurs.

Sur les courbes de Jordan ne renfermant aucune courbe simple fermée de Jordan.

par

Tadé Ważewski

Introduction.

Nous nous proposons de faire comprendre pourquoi les notions de correspondance homéomorphe, de type topologique, et de courbe de Jordan se sont imposées aux recherches mathématiques.

Nous invoquerons pour cette raison quelques faits historiques sans prétendre cependant donner le développement historique complet de cette question.

Nous commençons par donner la définition de la correspondance homéomorphe, de quelques notions qui s'y rattachent et de quelques théorèmes qui la concernent.

§ 1.

Def. 1. On dit que la fonction

$$p = \varphi(q)$$

constitue une correspondance homéomorphe entre les ensembles P et Q , si

1°) l'équation précédente fait correspondre à tout point p de P un seul point q de Q et à tout point q de Q un seul point p de P ; c. à d., si cette correspondance est biunivoque et réciproque

2°) si

$$\lim p_\nu = p$$

entraîne pour les points correspondants de Q

$$\lim q_\nu = q$$

et inversement, c. à d. si la correspondance est bicontinue.

Déf. 2. Deux ensembles P et Q entre lesquels se laisse établir une correspondance homéomorphe s'appellent *homéomorphes*.

Déf. 3. On appelle *invariant topologique* toute propriété P qui se conserve, si l'ensemble qui en jouit subit une transformation homéomorphe. La classe de tous les ensembles possédant une propriété étant invariant topologique s'appelle *type topologique*.

$$\text{Th. 4. Si } p = \varphi(q), q = \psi(r)$$

sont deux correspondances homéomorphes, alors leur produit c. à d. la correspondance

$$p = \varphi[\psi(r)]$$

est homéomorphe.

Th. 5. La correspondance biunivoque et *continue* entre deux ensembles P et Q situés dans deux espaces cartésiens dont un est fermé et borné est *bicontinue*.

(On entend par espace cartésien C_n la classe de toutes les suites $x_1, x_2, \dots, x_n$ composées de nombres réels).

§ 2.

Correspondance homéomorphe dans la question du nombre de dimensions.

Les difficultés contre lesquelles on s'est heurté en cherchant la définition précise d'une multiplicité à n dimensions s'offrant aussi dans le cas des espaces cartésiens C_n , il nous suffit de nous borner aux ensembles faisant partie de ces espaces.

Riemann est un des premiers qui ont cherché à préciser la notion du nombre de dimensions.

Afin de pouvoir se rendre compte du rapport du concept intuitif du nombre de dimensions à la notion découlant de la définition riemannienne nous donnons quelques théorèmes sur ce sujet puisés dans l'intuition.

Th. a) Le segment de droite à une dimension.

a') L'équation

$$y = f(x)$$

où f est une fonction continue définie dans l'intervalle $(0, 1)$ représente une ligne.

a'') La ligne a une dimension.

- b) L'intérieur du cercle et du carré (frontière comprise) ont chacun deux dimensions.
- c) L'intérieur de la sphère (frontière comprise) a trois dimensions.
- d) Aucun sous ensemble d'une multiplicité à n dimensions n'a $n + k$ dimensions ($k > 0$).

Dans les énoncés précédents deux notions ne sont pas précisées: celle du nombre de dimensions et celle de la ligne. Pour ne pas s'éloigner de l'intuition, il convient de les définir d'une façon conforme aux propositions *a* et *d*.

Voici la définition riemanienne que nous citons *in extenso* afin d'éviter toute altération de ses idées¹⁾.

„Etant donné, dit-il, un concept dont les modes de détermination [les points²⁾] forment une variété continue, si l'on passe, suivant une manière déterminée, d'un mode de détermination à un autre les modes de détermination parcourus forment une variété étendue dans un seul sens dont le caractère essentiel est que, dans cette variété, on ne peut, en partant d'un point s'avancer d'une manière continue que dans deux directions en avant et en arrière. Imaginons maintenant que cette variété se transporte à son tour sur une autre variété complètement distincte et cela encore d'une manière déterminée, c'est à dire tellement que chacun de ses points se transporte en un point déterminé de l'autre variété. l'ensemble des modes de détermination ainsi obtenus formera une variété de deux dimensions. On obtient semblablement une variété de trois dimensions, si l'on conçoit qu'une variété de deux dimensions se transforme d'une façon déterminée sur une autre complètement distincte, et il est aisé de voir comment on peut poursuivre cette construction“...

La définition suivante de „l'arc simple spécial“ donne une ligne au sens de Riemann, sans les comprendre cependant toutes.

Déf. On appelle arc simple spécial chaque classe de tous les points p de l'espace C_n satisfaisant aux conditions suivantes.

$$\alpha) \quad p = \varphi(t) \quad (0 \leq t \leq 1)$$

où $\varphi(t)$ est une fonction continue de t ;

$\beta)$ $\varphi(t)$ n'a pas de points multiples c. à d.

¹⁾ Über die Hypothesen die der Geometrie zu Grunde liegen (Trad. française dans les Oeuvres mathématiques de Riemann.

²⁾ La parenthèse n'est pas du texte

$$\varphi(t_1) \neq \varphi(t_2), \text{ si } t_1 \neq t_2;$$

γ) Il n'existe pas trois arcs simples L_1, L_2, L_3 , tels que

$$L_1 + L_2 + L_3 \subset L$$

et dont les équations sont

$$l_1 = \varphi_1(t)$$

$$l_2 = \varphi_2(t)$$

$$l_3 = \varphi_3(t) \quad (0 \leq t \leq 1)$$

et, tels que

$$L_1 L_2 + L_2 L_3 + L_3 L_1 = \varphi_1(0) = \varphi_2(0) = \varphi_3(0).$$

En abrégé on pourrait dire: l'arc simple spécial est l'image homéomorphe du segment $(0,1)$ dans l'espace C_1 qui satisfait en outre à la condition γ .

On n'est pas a priori certain si γ résulte de α et β , la nécessité de la démonstration va paraître dans la suite.

L'exemple suivant jettera un peu de lumière sur la définition riemannienne d'une multiplicité à deux dimensions.

On a donné l'exemple d'une fonction

$$z = f(x, y)$$

qui, dans tout point du carré aux sommets opposés

$$(0, 0), (1, 1),$$

est continue par rapport à x seul et par rapport à y seul et possède néanmoins les points de discontinuité dans tout cercle intérieur au carré¹⁾.

La classe de tous les points x, y, z satisfaisant à l'équation précédente est une multiplicité à deux dimensions au sens de Riemann, car elle est décrite par la famille d'arcs simples spéciaux que l'on obtient en regardant la variable x comme paramètre. On pourrait démontrer que cette multiplicité n'est pas homéomorphe

¹⁾ Soit $a_1, a_2, \dots$ une suite dont tous les points forment un ensemble dense sur le carré. La fonction

$$\sum_{\nu=1}^{\infty} \frac{1}{4^\nu} \varphi(x - x_\nu, y - y_\nu),$$

$$\text{où } \varphi(x, y) = \frac{xy}{x^2 + y^2} \quad (x^2 + y^2 \neq 0)$$

$$\varphi(0, 0) = 0$$

$$(x_\nu, y_\nu) = a_\nu$$

possède la propriété en question.

avec aucun ensemble plan ainsi que tout son sous-ensemble correspondant à un carré partiel

La correspondance homéomorphe ne joue donc pas le rôle qu' on lui pourrait attribuer à première vue dans la définition de Riemann.

Cet exemple suggère l'idée de chercher à remplacer la définition riemanienne par une autre qui ne caractériserait pas comme multiplicités à deux dimensions les ensembles si différents des surfaces ordinaires.

Il fait entrevoir les difficultés qui se présentent, si l'on cherche à démontrer le théorème *d* (page 3). Ces difficultés ont été mises en évidence par les résultats de M. M. Cantor et Peano.

La marque caractéristique de la multiplicité à deux dimensions au sens de Riemann consiste dans une certaine correspondance entre ses points et les points de l'espace cartésien C_2 .

Cantor s'est posé le problème de savoir si la propriété, pour cette correspondance, d'être biunivoque, suffirait à caractériser le nombre de dimensions. Il a montré que la réponse est négative¹⁾, en construisant une correspondance biunivoque entre les points d'un segment et ceux d'un carré c. à d. entre des ensembles dont les nombres de dimensions sont différents. Il a remarqué que cette correspondance n'était pas continue et a énoncé l'opinion que le postulat de continuité rendrait un exemple pareil impossible.

Des recherches du même genre ont conduit M. Peano à construire une courbe remplissant le carré. Cette courbe a des points multiples — de là l'existence d'une courbe pareille mais sans points multiples (d'un arc simple) apparaît a priori comme possible. Si l'on parvenait à en construire une, il en résulterait que la propriété γ introduit une restriction essentielle dans la définition de l'arc simple spécial.

Le problème posé par Cantor se généralise de la façon suivante:

Est-ce qu'est vrai le théorème: (*T*) L'image homéomorphe d'un hypersolide sphérique à n dimensions c. à d. de la classe de tous les points $x_1, x_2, \dots, x_n$ satisfaisant à la relation

$$\sum (x_n - a_n)^2 \leq \varepsilon, \quad (\varepsilon > 0)$$

ne peut pas renfermer un hypersolide sphérique à $n+k$ dimensions ($k > 0$).

¹⁾ Journal de Crelle t. 84, aussi Mat. Ann. 1897 (Zur Begründung der transfiniten Mengenlehre I).

Tout d'abord ce théorème a été démontré, dans le cas particulier, par M. Lüroth; la démonstration générale a été donnée par M. Brouwer et par M. Lebesgue ¹⁾.

Le rapprochement des théorèmes (d) et (T) suggère l'idée de considérer comme ensembles à m dimensions, dans un espace cartésien C_n , les ensembles qui ont leurs homéomorphes dans l'espace C_m et n'en ont pas dans les espaces d'ordre inférieur.

On peut constater que les théorèmes $a - d$ sont vrais pour une telle définition (ligne γ étant considérée comme synonyme d'arc simple). Si l'on remarque qu'elle est plus simple que celle de Riemann, nous sommes conduits à lui donner la préférence, d'autant plus qu'en tant que nous le savons bien la propriété d n'a pas été démontrée pour la définition de Riemann.

Remarque. La circonférence suivant cette définition a deux dimensions et pourtant elle est considérée comme une ligne. De là vient l'idée de donner à part de la définition précédente aussi une définition du nombre de dimensions dans un autre sens, définition „in micro“. A tout point m de la circonférence correspond un sous-ensemble (un petit arc) qui est d'une dimension et qui contient tous les points de la circonférence renfermés dans un certain voisinage de m . On peut exprimer cela en disant que la circonférence a une dimension en chacun de ses points. (Un ensemble à m dimensions in micro peut être décomposé, s'il est fermé et borné, en un nombre fini d'ensembles à n dimension in marco — théorème de Borel Lebesgue.)

§ 3.

Courbes de Jordan.

Déf 1. On appelle *courbe de Jordan* la classe de tous les points $x_1, \dots, x_n$ satisfaisant aux équations:

$$x_v = \varphi_v(t) \quad (v/1, \dots, n)$$

où $\varphi_v(t)$ sont des fonctions continues et définies dans un intervalle fermé (t_0, t_1) .

Déf 2. On appelle *orbe de Jordan* ou *orbe* l'image homéomorphe de la circonférence. C'est évidemment une courbe de Jordan.

¹⁾ Math. Ann. t. 70. Voir aussi Fund. Math. t. II.

Th 3. La propriété d'être une courbe de Jordan est un invariant topologique, c. à d. toute image homéomorphe d'une courbe de Jordan est une courbe de Jordan.

L'exemple de la courbe de M. Peano qui est une courbe de Jordan est un de premiers qui a apporté de l'intérêt à cette classe d'ensembles.

On a recherché les conditions auxquelles doivent satisfaire les ensembles pouvant être remplis par une courbe de Jordan.

Deux critères très précis et très maniables ont été donnés par M. Mazurkiewicz et par M. Sierpiński¹⁾.

La classe des courbes de Jordan doit aussi son intérêt au théorème de Jordan suivant lequel toute orbée J divise le plan en deux domaines E_1 et E_2 , tels que deux points du même domaine E_i se laissent joindre par une courbe de Jordan contenue dans E_i et toute courbe reliant un point de E_1 avec un point de E_2 empiète sur J .

Le rôle joué par ce théorème dans les théorèmes de Cauchy et Green est bien connu.

L'étude de ces courbes a été l'objet des travaux de M. Schönflies (et M. Riesz)²⁾. Il s'est proposé inversement de chercher, si tout ensemble de points divisant le plan est une courbe de Jordan.

L'exemple de l'ensemble E suivant fait voir que la réponse est négative. E est composé de la ligne

$$y = \sin \frac{1}{x} \quad \left(0 < x \leq \frac{1}{\pi} \right)$$

et des segments

$$\left(\frac{1}{\pi}, 0 \right); \left(\frac{1}{\pi}, 2 \right)$$

$$(0, -1); (0, 2)$$

$$(0, 2); \left(\frac{1}{\pi}, 2 \right)$$

(Il existe même des ensembles découpant le plan et ne renfermant aucun arc simple fermé. On en obtient un en combinant

¹⁾ Fundamenta Mathematicae, Tome I, Varsovie 1920, p. 166 et 44.

²⁾ Math. Annalen, tomes 58 et 59.

plusieurs continus ne reufermant aucun sous-continu representable paramétriquement ¹⁾.

Les travaux de M. Schönflies sont suggéré à M. Brouwer la construction d'un domaine dont la frontière est fermée et continue; elle divise le plan en trois domaines dont elle est frontiere commune ²⁾. M. Denjoy a donné un exemple pareil dans lequel le nombre trois est remplacé par l'infinité dénombrable.

La recherche des conditions auxquelles doit satisfaire un domaine fini plan pour que sa frontière soit une orbe a abouti au théorème suivant ³⁾:

La condition nécessaire et suffisante est que la frontière de ce domaine découpe le plan en deux parties et que tout point f de cette frontière soit *accessible* de l'intérieur et de l'extérieur, c. a. d. qu'il existe des arcs simples situés à l'extérieur et des autres à l'intérieur à l'exception d'une de ces extrémités qui coïncide avec f . M. Denjoy a donné un autre critère consistant à ce qu'à tout $\varepsilon > 0$ il existe $\delta > 0$ tel que deux points du domaine se laissent joindre par un sous-arc simple du domaine de diamètre inférieur à ε pourvu que la distance de ces points soit inférieure à δ .

L'arc simple a été l'objet de plusieurs travaux. On a recherché des conditions pour qu'un ensemble soit arc simple, des conditions dans lesquelles n'intervienne pas la représentation paramétrique. M. M. Lennes, Zoretti, Janiszewski, Sierpiński, Knaster et Kuratowski y ont consacré des notes et des mémoires ⁴⁾.

Tous les resultats cités montrent que les courbes de Jordan constituent un groupe d'ensembles jouissant de propriétés tout à fait particulières.

¹⁾ Les continus de la sorte ont été construits par M. Janiszewski et par M. Brouwer. Une construction développée d'un tel continu se trouve dans le mémoire de M. Knaster Fund. Math. t. III.

²⁾ Mat. Ann. t. 68.

³⁾ Les articles de M. Schönflies et de M. Riesz précités (Math. Ann. t. 58 et 59).

⁴⁾ Lennes Ann. Journ. Math. 1911.

Zoretti Annales de l'Ecole Norm. 1909.

Janiszewski Journal de l'Ecole polytechnique 1912.

Sierpiński Fund. Math. 1.

Mazurkiewicz ibid. 1 et 2.

Knaster et Kuratowski ibid 2.

Kuratowski ibid 1.

M. Mazurkiewicz¹⁾ en train à rechercher si sur chaque courbe de Jordan il existe des points qui ne la découpent pas (comme tous les points pour la circonférence et les extrémités pour l'arc simple) a été obligé de s'occuper d'une famille particulière de courbes de Jordan qui ne renferment comme sous-ensemble (l'ensemble étant regardé comme un de ses sous-ensembles) aucune orbe de Jordan. (Nous les appelons *dendrites*).

Il a posé à cette occasion le problème: est-ce que toutes les dendrites sont homéomorphes d'un ensemble plan? ou, dans la terminologie du chapitre précédent: est-ce que le nombre de dimensions de toute dendrite est deux au plus?

Nous démontrerons que

I) Toute dendrite est une image homéomorphe d'un ensemble plan et l'homéomorphie en question peut être effectivement indiquée.

II. La classe des points de branchement de toute dendrite est effectivement énumérable.

III. On peut construire effectivement une dendrite D^* , telle que toute autre dendrite est effectivement homéomorphe avec une partie de D^* .

Une partie de la démonstration consiste à établir un certain ordre sur la dendrite. Pour le faire de façon à éviter les nombres transfinis et l'axiome de M. Zermelo, nous nous sommes servis d'un certain espace abstrait. Cependant cette voie n'est pas la plus simple. M. Denjoy au cours de l'expertise de ce travail a indiqué une méthode beaucoup plus simple. Elle permet de démontrer, en plus, le théorème II et donne un critère permettant de reconnaître l'ordre de multiplicités des points de branchement. Elle se trouve développée à la fin de ce travail.

Idée directrice de la démonstration.

Traçons sur le plan un segment a, b et considérons un ensemble dénombrable de points (a_i) dense sur ce segment et ne contenant aucune de ses extrémités. Choisissons un des côtés de ce segment comme positif. On peut, comme il est facile de le voir, tracer à partir de chaque point a_i deux demi-droites différentes et situées du côté positif de a, b .

¹⁾ Fund. Math. t. II. p. 130.

$$Ra_1, R'a_1$$

de façon que les angles $\angle (Ra_1, R'a_1)$ correspondant aux a_1 différents soient deux à deux disjoints. Dans chacun de ces angles choisissons une infinité dénombrable de demi-droites issues du point a_1 correspondant. Sur chacune de ces demi-droites nous découpons un petit segment d'origine a_1 de façon que la longueur de ces segments tende vers 0, et que ces segments forment avec le segment a_1, b_1 un ensemble fermé et bien enchaîné. Désignons ces nouveaux segments par a_1, b_1' .

Nous avons obtenu une sorte d'arbre dont a, b est la tige et dont les feuilles a_1, b_1' ne se touchent pas mutuellement en dehors de la tige.

Sur chacune des feuilles a_1, b_1' nous construisons un feuillage de second ordre d'une façon analogue, tout en ayant soin que les feuilles nouvelles ne se touchent pas mutuellement en dehors des feuilles anciennes et que chacune touche une seule feuille ancienne. En continuant ce procédé et en prenant des précautions convenables concernant la longueur des feuilles, on obtiens après une infinité d'opérations un ensemble E . En le fermant on obtient une dendrite D^* .

Soit d'autre part une dendrite quelconque D . Elle renferme des arcs simples saturés c. à d. ne se laissant pas prolonger sans quitter D . Nous en choisissons un, soit a, b . Dans ce qui reste nous choisissons un ensemble d'arcs simples a_1, b_1' ayant seulement a_1 sur a, b , ne se laissant pas prolonger au delà de b_1' et ne se touchant pas mutuellement en dehors de a, b . Nous traçons autant de „feuilles“ a_1, b_1' qu'il est possible, en prenant quelques précautions sur lesquelles nous n'insistons pas pour le moment. Nous recommençons cette opération à partir des a_2, b_2' et continuons jusqu'à épuiser toute la dendrite.

L'analogie de la structure des feuillages dans les cas de D^* et D suggère l'idée que D est homéomorphe avec une partie de D^* et fait entrevoir comment on peut établir cette homéomorphie.

Les difficultés consistent (1°) à décider parmi tous les choix possibles (2°) à ne pas laisser hors des raisonnements les points s'introduisant à la limite, (3°) à épuiser toute la dendrite D par une suite infinie ordinaire d'opérations (abstraction faite de des points qui s'introduisent pendant qu'on ferme le résultat).

A°. Choix effectif.

§ 1. *Faire le choix effectif* (ch. e) dans une classe A veut dire établir une loi à laquelle satisfait un seul élément de la classe A .

Au lieu de dire que nous savons faire le choix effectif dans A , nous dirons souvent qu'il existe dans A un élément parfaitement déterminé.

§ 2. Pour démontrer que toute classe peut être bien ordonnée, M. Zermelo a été obligé de se servir du principe suivant („axiome du choix“):

(A) étant une classe d'ensembles sans points communs deux à deux, il existe un ensemble B ayant avec tout ensemble A un seul point commun et ne contenant pas d'autres éléments.

Ce principe n'est pas reconnu par tous les mathématiciens. Ses adversaires sont d'avis que l'existence de B n'est établie que si l'on sait faire le choix effectif dans tout ensemble A .

Toutes les fois que nous aurons besoin de nous servir de l'existence de l'ensemble B , nous ferons les choix effectifs dans les A , ce qui sera possible en raison du caractère spécial des classes (A) que nous rencontrerons. Nous ne développerons cependant la loi du choix dans les cas où elle est banale ou bien, si elle découle d'une façon simple des considérations antérieures.

A.

I. Classe fréchétienne normale $E(\rho)$.

§ 1. Nous appelons espace fréchétien *normal* toute classe $E(\rho)$ satisfaisant aux conditions:

1₀) A tout couple (a, b) d'éléments („points“) de $E(\rho)$ correspond un nombre $\rho(a, b)$ („distance“), tel'que

$$\alpha) \rho(a, b) = 0 \text{ équivaut à } a = b$$

$$\beta) \rho(a, b) = \rho(b, a)$$

$$\gamma) \rho(a, b) + \rho(b, c) \geq \rho(a, c)$$

2₀) $E(\rho)$ est compact c. à d. toute infinité de points admet au moins un point d'accumulation¹⁾.

¹⁾ On appelle sphère de rayon r ($r > 0$) et au centre a la classe de tous les points x de $E(\rho)$ pour lesquels $\rho(a, x) < r$. Toute sphère de cette sorte $Sp(a, r)$ s'appelle voisinage de a . On dit que b est un point d'accumulation d'un ensemble B , si tout voisinage de b contient un sous-ensemble infini de B . Nous désignons dans la suite les points par minuscules, les ensembles par majuscules.

§ 0) $E(\varphi)$ renferme une partie dénombrable (e_ν) effectivement énumérable et partout dense sur $E(\varphi)$ c. à. d. empiétant sur tout voisinage de chaque point.

Voici quelques conséquences de cette définition.

§ 2. La classe de tous les points de l'espace cartésien à n dimensions, tels que

$$\varphi[(x_1, x_2, \dots, x_n), (0, 0, \dots, 0)] \leq \alpha$$

où α est un nombre positif et

$$\varphi[(x_1, \dots, x_n), (y_1, \dots, y_n)] = \sqrt{\sum_{v=1}^n (x_v - y_v)^2}$$

forme un espace $E(\varphi)$.

§ 3. $\varphi(a, b) \geq 0$.

§ 4. Déf. Nous désignons par $\bar{A}$ la classe composée de tous les points de A et de leur points d'accumulation.

§ 5. Si $A \subset B$, alors $\bar{A} \subset \bar{B}$.

§ 6. Déf. „ A est fermé“ signifie: $A = \bar{A}$.

§ 7. La partie commune d'une famille quelconque d'ensembles fermés est fermée.

§ 8. $E(\varphi)$ est fermée.

§ 9. $\bar{\bar{A}} = \bar{A}$.

§ 10. $A + B = \bar{A} + \bar{B}$.

§ 11. Si $A = \bar{A}$, $B = \bar{B}$, alors $A \cdot B = \overline{A \cdot B}$.

§ 12. Une fonction $f(x)$ réelle est dite continue au point a , si pour tout $\varepsilon > 0$ existe un voisinage de a , tel que pour les points de ce voisinage pour lesquels $f(x)$ est définie, on a

$$f(x) - f(a) \leq \varepsilon.$$

§ 13. $\varphi(a, x) = \varphi(r, a)$ est une fonction continue en tout point de $E(\varphi)$.

§ 14. Si $f(x)$ est continue dans un ensemble fermé $\{x\}$, alors la classe de tous les x pour lesquels $f(x) \geq a$ (ou $f(x) \leq \alpha$) est fermée. Autrement: la classe

$$L(f(x) \geq \alpha)$$

est fermée.

§ 15. Les bornes supérieure et inférieure d'une fonction continue sont les mêmes pour A et $\bar{A}$.

§ 16. On appelle diamètre de A

$$\tau(A)$$

la borne supérieure de $\rho(x, y)$, où x et y sont deux points variables dans A .

§ 17. Le diamètre de toute sphère de rayon r est au plus $2r$.

§ 18. On appelle sphère rationnelle toute sphère au centre e_ν (comp § 1, 3_0) et dont le rayon est rationnel. e_ν s'appelle point rationnel.

§ 19. Toutes les sphères rationnelles constituent une classe effectivement énumérable.

§ 20. Tout point a de $E(\varepsilon)$ est contenu dans une sphère rationnelle de rayon $\varepsilon > 0$ (ch. e) où ε est arbitraire.

Cela résulte des § 1, 3_0 et § 19.

§ 21. Si
$$\left. \begin{array}{l} F_\nu = \overline{F}_\nu \neq 0 \\ F_{\nu+1} \subset F_\nu \end{array} \right\} (\nu/1, 2, \dots, \infty)$$

alors

$$\bigcap_{\nu/1}^\infty F_\nu = \bigcap_{\nu/1}^\infty \overline{F}_\nu \neq 0$$

et on sait faire le choix effectif d'un élément de cet ensemble.

Dém. La classe de toutes les sphères rationnelles étant effectivement énumérable, on peut facilement (§ 20) trouver une suite de sphères rationnelles (ch. e)

$$\begin{aligned} & Sx_1, Sx_2, \dots, \\ \text{telle que} \quad & Sx_{\nu+1} \subset Sx_\nu \quad (\nu/1, \dots, \infty) \\ & \lim \text{ rayon } Sx_\nu = 0 \\ & Sx_\nu \times F_\nu \neq 0 \quad (\nu/1, \dots, \infty) \end{aligned}$$

On démontre facilement que le produit de toutes ces sphères n'est pas vide et qu'il ne comprend qu'un seul point et que ce point appartient à F_ν ($\nu/1, \dots, \infty$).

§ 22. Dans chaque ensemble F fermé et non vide on peut choisir effectivement un point.

Dém. Cela résulte du § 21., si l'on y pose $F_\nu = F$.

§ 23. A tout ensemble infini A correspond (par un choix effectif) un point d'accumulation.

La démonstration est analogue à celle du § 21.

§ 24. A toute suite infinie de points correspond une suite choisie (ch. e) convergente. Cela résulte du § 23.

§ 25. La classe des points où une fonction continue et définie dans un ensemble fermé atteint son maximum (ou minimum) est non vide et fermée. Une telle fonction est donc bornée.

Dém. Considérons une suite décroissante $\{z_\nu\}$ tendant vers le minimum de la fonction qui à priori est fini ou égale ∞ . On a

$$E(f(x) \leq z_\nu) \neq \emptyset$$

$$E(f(x) \leq z_{\nu+1}) \subset E(f(x) \leq z_\nu)$$

et, d'après le § 14, les ensembles E sont fermés.

$$\bigcap_{\nu=1}^{\infty} E(f(x) \leq z_\nu)$$

est donc fermé et non vide (§ 21). d'où résulte le théorème.

§ 26. L'espace $E(\cdot)$ est borné, c. à. d. son diamètre est fini (§ 13, § 25).

§ 27. Dans le cas des espaces normaux l'image univoque et continue d'un ensemble fermé est fermé.

Autrement:

Si $y = \varphi(x)$ est une fonction continue subordonnant à tout point x d'un ensemble fermé X appartenant à un espace normal $E(\rho_1)$ un point y d'un autre espace normal, alors $(\varphi(x))$ est fermé.

Nous n'insistons pas sur la définition de continuité dans ce cas, c'est une simple généralisation de la définition du § 12.

On ramène comme il suit la proposition précédente au théorème du § 25.

Supposons que

$$y_1 \in \overline{(\varphi(x))}^1 \quad (1)$$

et considérons la fonction

$$f(x) = \rho(y_1, \varphi(x)). \quad (2)$$

$f(x)$ est une fonction continue non négative. D'après (1) son minimum est 0.

Il est atteint en un point x_1 de X (ch. e § 25) et on a

$$0 = f(x_1) = \rho(y_1, \varphi(x_1)),$$

d'où

$$y_1 = \varphi(x_1)$$

c. à. d.

$$y_1 \in \{\varphi(x)\}$$

c. q. f. d.

§ 28. On appelle

$$\rho(a, A)$$

¹⁾ $a \in A$ signifie: a est un élément de la classe A .

le minimum de la fonction $\rho(a, x)$ définie pour tout x appartenant à A .

§ 29. Suivant le § 15

$$\rho(a, A) = \rho(a, \overline{A}).$$

§ 30. D'après le § 25, il existe un a_1 (ch. e), tel que

$$\rho(a, A) = \rho(a, a_1)$$

$$a_1 \in \overline{A}$$

§ 31. Si $F = \overline{F}$, alors

$$\rho(a, F) = 0$$

équivalent à

$$a \in F.$$

C'est une conséquence du § 25.

§ 32. *Déf.* Pour $0 \neq A \neq B \neq 0$, nous posons

$$\rho(A, B) = \min (\rho(x, y))$$

$$(x \in A, y \in B)$$

§ 33. On a pour $0 \neq A \neq B \neq 0$

$$\rho(A, B) = \min_{(x \in A, y \in B)} \rho(x, y) = \min_{(y \in B)} \rho(A, y) = \min_{(x \in A)} \rho(x, B)$$

$$(x \in A, y \in B) \quad (y \in B) \quad (x \in A)$$

Dém. Evidemment pour $x \in A, y \in B$ on a

$$\rho(A, y) \leq \rho(x, y), \quad (1)$$

donc

$$\min_{(y \in B)} \rho(A, y) \leq \min_{(x \in A, y \in B)} \rho(x, y)$$

$$(y \in B) \quad (x \in A, y \in B)$$

D'autre part à tout y ($y \in B$) et à tout $\varepsilon > 0$, il existe x' ($x' \in A$), tel que

$$\rho(A, y) \geq \rho(x', y) - \varepsilon$$

d'où ε étant arbitraire on a

$$\min_{(y \in B)} \rho(A, y) \geq \min_{(x \in A, y \in B)} \rho(x, y) \quad (2)$$

$$(y \in B) \quad (x \in A, y \in B)$$

De (1) et (2) il résulte,

$$\min_{(x \in A, y \in B)} \rho(x, y) = \min_{(y \in B)} \rho(A, y).$$

D'une façon analogue on démontre

$$\min_{(x \in A, y \in B)} \rho(x, y) = \min_{(x \in A)} \rho(x, B)$$

$$\S 34. \quad \rho(A, B) = \rho(\overline{A}, B) = \rho(\overline{A}, \overline{B}) = \rho(B, A)$$

C'est une conséquence des § 33 et § 29.

$$\S 35. \quad \rho(A, b) + \rho(b, C) \geq \rho(A, C)$$

Il suffit de démontrer (§§ 34, 29) que

$$\rho(\overline{A}, b) + \rho(b, \overline{C}) \geq \rho(\overline{A}, \overline{C})$$

D'après § 30, il existe un a_1 et un b_1 ($a_1 \in \overline{A}$, $b_1 \in \overline{B}$), tels que

$$\rho(\overline{A}, b) = \rho(a_1, b)$$

$$\rho(b, \overline{C}) = \rho(b, c_1)$$

donc

$$\rho(\overline{A}, b) + \rho(b, \overline{C}) \geq \rho(a_1, b_1) \geq \rho(\overline{A}, \overline{B}) \text{ c. q. f. d.}$$

§ 36. $\rho(x, A)$ est une fonction continue dans tout point x_0 de E (ρ)

En effet

$$\rho(x, A) \leq \rho(x, x_0) + \rho(x_0, A)$$

$$\rho(x_0, A) \leq \rho(x_0, x) + \rho(x, A),$$

donc

$$-\rho(x, x_0) \leq \rho(x_0, A) - \rho(x, A) \leq \rho(x, x_0)$$

c. à. d.

$$\rho(x, A) - \rho(x_0, A) \leq \rho(x_0, x) \text{ c. q. f. d.}$$

§ 37. Pour deux ensembles fermés non vides

$$\rho(F_1, F_2) = 0 \quad (1)$$

équivalent à

$$F_1 \times F_2 \neq \emptyset \quad (2)$$

car la fonction continue $\rho(x, F_2)$ atteint son minimum dans un point a , de F_1 . On a selon (1)

$$\rho(a_1, F_2) = 0,$$

donc

$$a_1 \in F_2,$$

d'où résulte (2). L'inverse est évidente.

§ 38. Si F_1 et F_2 sont fermés et $F_1 - F_2 \neq \emptyset$, alors $F_1 - F_2$ est une somme d'une suite d'ensembles fermés (ch. e).

Dém. On a

$$F_1 - F_2 = \sum_{n=1}^{\infty} \left(E(\rho(x, F_2) \geq \frac{1}{n}) \right) \cdot F_1$$

et les ensembles $\left(E(\rho(x, F_2) \geq \frac{1}{n}) \right) \cdot F_1$ sont fermés en raison de la continuité de $\rho(x, F_2)$ [§ 14, § 36].

§ 39. Si F_1 et F_2 sont fermés et $F_1 - F_2 \neq 0$, alors on sait faire le choix effectif dans $F_1 - F_2$.

Cela résulte des §§ 38 et 22.

§ 40. Toute classe finie d'ensembles fermés est effectivement énumérable.

Cela résulte du lemme:

Dans toute classe finie d'ensembles fermés on sait faire le choix effectif de l'un d'ensembles. Prenons le cas de deux ensembles. Leur somme S et leur produit P sont des ensembles fermés et $S - P$ n'est pas vide, car on a deux ensembles différents. Nous y savons choisir un point (ch. e. § 39). Ce point n'appartient qu'à un seul ensemble qui se trouve ainsi parfaitement déterminé. On généralise par l'induction mathématique.

§ 41. Théorème de M. M. Borel et Lebesgue (cas spécial).

Si $\varepsilon > 0$, il existe une classe finie et effectivement énumérable de sphères rationnelles dont les rayons sont inférieurs à ε et dont la somme renferme $E(q)$.

Dém. Les sphères rationnelles de rayon $< \varepsilon$ forment une suite infinie (ch. e. § 19):

$$S_1, S_2, \dots$$

Tout point de $E(q)$ appartient au moins à une d'elles (§ 20). Je dis qu'il existe un N , tel que

$$E(q) \subset S_1 + S_2 + \dots + S_N.$$

Dans le cas contraire, on aurait pour tout n naturel

$$P_n = C(S_1) \times C(S_2) \times \dots \times C(S_n) \neq 0,$$

où $C(S_v)$ désigne le complémentaire de S_v par rapport à $E(q)$.

Suivant la définition de S_v

$$S_v = E[\varrho(e_{\alpha_v}, x) < r_{\alpha_v}],$$

où e_{α_v} est son centre, r_{α_v} son rayon.

On a par conséquent

$$C(S_v) = E[\varrho(e_{\alpha_v}, x) > r_{\alpha_v}]$$

et cet ensemble est fermé (§ 14) c. à d.

$$P_n = \overline{P_n}.$$

Comme

$$P_n \neq 0 \quad \text{et} \quad P_{n+1} \subset P_n \quad (n/1, \dots, \infty)$$

donc (§ 21)

$$\prod_{\nu=1}^{\infty} P_{\nu} = \prod_{\nu=1}^{\infty} C(S_{\nu}) \neq 0,$$

contrairement à ce que tout point de $E(\varrho)$ appartient à une au moins des S_{ν} .

§ 42. Déf. On dit qu'un ensemble A contenu dans B est dense sur l'ensemble B, si tout voisinage de tout point de B contient des éléments de A.

§ 43. Dans tout ensemble fermé F il existe une suite (finie ou non) de points différents

$$a_1, a_2, \dots$$

dense sur F (ch. e)

Dém. Couvrons l'espace $E(\varrho)$ avec un nombre fini sphères de rayon $< 1/n$ (ch. e. § 41)

$$S_1, S_2, \dots, S_m$$

et choisissons dans tout $\overline{S_{\nu}} \times F$ qui est non vide un point (ch. e. § 22). Rangeons ces points en une suite (ch. e. § 40)

$$(1)_n \quad s_1, s_2, \dots, s_k.$$

Soit $b_1, b_2, \dots$ la suite infinie dans la quelle on a écrit tout d'abord les points de la suite $(1)_1$, ensuite ceux de la suite $(1)_2$ etc. Si l'on élimine de cette suite les éléments qui y interviennent pour la 2-e, 3-e, etc. fois, on obtient la suite aux propriétés annoncées.

§ 44. Définitions du continu.

I. Déf. de Cantor.

On appelle continu un ensemble fermé, bien enchaîné et non vide (tout point est un continu).

Un ensemble E est dit bien enchaîné, si à toute paire de ses points a et a' et à tout $\varepsilon > 0$, il existe une suite finie de points

$$a_0, a_1, \dots, a_n,$$

telle que

$$\begin{aligned} a_0 &= a, \quad a_n = a', \quad a_{\nu} \in E & (\nu/1, \dots, n) \\ \varrho(a_{\nu}, a_{\nu+1}) &\leq \varepsilon & (\nu/0, \dots, n-1). \end{aligned}$$

II. Définition de Jordan-Schönfliess. On appelle continu un ensemble fermé non vide qui n'est pas somme de deux ensembles fermés non vides et sans points communs.

On pourrait facilement démontrer l'équivalence de ces définitions pour l'espace normal.

II. Espace $E(q^*)$.

§ 45. Déf. Considérons deux ensembles A et B appartenant à un espace $E(q)$.

Si $A = \overline{A} \neq 0$, $\overline{B} = B \neq 0$, nous appelons $q^*(A, B)$ le plus grand des nombres

$$\frac{\overline{\text{borne}}}{x \in A} q(x, B) \\ \frac{\overline{\text{borne}}}{y \in B} q(A, y)^1$$

et nous appelons $E(q^*)$ la famille de tous les ensembles fermés et non vides de $E(q)$.

§ 46. On a pour des ensembles fermés et non vides de $E(q)$:

α) $q^*(F_1, F_2) = 0$ équivaut à $F_1 = F_2$,

β) $q^*(F_1, F_2) = q^*(F_2, F_1)$,

γ) $q^*(F_1, F_2) + q^*(F_2, F_3) \geq q^*(F_1, F_3)$.

Dém. α) Supposons que $q^*(F_1, F_2) = 0$. De là

$$\frac{\overline{\text{borne}}}{x \in F_1} q(x, F_2) = 0$$

par conséquent (§ 31) $x \in F_1$ entraîne $x \in F_2$ c. à d.

$$F_1 \subset F_2.$$

Analoguement on a $F_2 \subset F_1$, donc $F_1 = F_2$. L'inverse est évidente

β) Résulte de l'égalité $q(a, A) = q(A, a)$.

Avant d'aborder la démonstration de γ nous prouverons deux lemmes:

Lemme 1. Si $F_1 = \overline{F_1} \neq 0$, $F_2 = \overline{F_2} \neq 0$, alors

$$q(a, F_2) \leq q^*(F_1, F_2)$$

(pour $a \in F_1$) (Déf. § 45).

Lemme 2.

$$q(a, F_2) + q^*(F_2, F_3) \geq q(a, F_3).$$

Dém. Il existe un point f_2 (ch. e.), tel que

$$f_2 \in F_2,$$

$$(1) \quad q(a, F_2) = q(a, f_2).$$

¹ Cette définition paraîtra plus naturelle, si l'on remarque que, dans le cas où $q^*(A, B)$ est infiniment petit, à tout point de A correspond au moins un point de B infiniment voisin et inversement. Elle est due à M. Hausdorff. *Grundsätze Mengenlehre* p. 293.

On a suivant le lemme précédent

$$(2) \quad \varphi^*(F_2, F_3) \geq \varphi(f_2, F_3).$$

En ajoutant (1) et (2), on obtient (§ 35)

$$\varphi(a, F_2) + \varphi^*(F_2, F_3) \geq \varphi(a, f_2) + \varphi(f_2, F_3) \geq \varphi(a, F_3) \text{ c. q. f. d.}$$

Pour démontrer γ remarquons que pour tout a , tel que $a \in F_1$, on a suivant les lemmes précédents

$$\varphi^*(F_1, F_2) + \varphi^*(F_2, F_3) \geq \varphi(a, F_3),$$

done

$$\varphi^*(F_1, F_2) + \varphi^*(F_2, F_3) \geq \overline{\text{borne}}_{(a \in F_1)} \varphi(a, F_3)$$

et d'une façon analogue

$$\varphi^*(F_1, F_2) + \varphi^*(F_2, F_3) \geq \overline{\text{borne}}_{(c \in F_1)} \varphi(F_1, c),$$

done suivant la définition de φ^*

$$\varphi^*(F_1, F_2) + \varphi^*(F_2, F_3) \geq \varphi^*(F_1, F_3) \text{ c. q. f. d.}$$

§ 47. Si $F_{\nu+1} \subset F_\nu$, $F_\nu = \overline{F_\nu} \neq 0$ ($\nu/1 \dots \infty$),
alors

$$\lim_{\nu/\infty} \varphi^*(F_\nu, \prod_{\mu/1}^\infty F_\mu) = 0.$$

Dém. Comme $0 \neq \prod_{\nu/1}^\infty F_\nu = \overline{\prod_{\nu/1}^\infty F_\nu}$, donc suivant la définition de φ^*

$$(1) \quad \varphi^*(F_\nu, \prod_{\mu/1}^\infty F_\mu) = \overline{\text{borne}}_{\mu/1} (f_\nu, \prod_{\mu/1}^\infty F_\mu).$$

Soit $\varepsilon > 0$. L'ensemble¹⁾

$$(2) \quad (x) = E[\varphi(x, \prod_{\mu/1}^\infty F_\mu) \geq \varepsilon]$$

est fermé d'après §§ 36 et 14.

Suivant (2) et § 31 on a

$$(x) \times \prod_{\mu/1}^\infty F_\mu = 0.$$

L'ensemble $G_\nu = F_\nu \times (x)$ est fermé et on a

$$(3) \quad G_{\nu+1} \subset G_\nu \quad (\nu/1, \dots \infty)$$

$$\prod_{\nu/1}^\infty G_\nu = (x) \times \prod_{\nu/1}^\infty F_\nu = 0,$$

done (§ 21) il existe un ν' , tel que

$$G_{\nu'} = 0.$$

¹⁾ Le lecteur est prié de ne pas confondre les parenthèses () et ().

Soit N la plus petite valeur d'un tel ν' .

On a d'après (3)

$$G_\nu = 0 \quad (\nu \geq N)$$

c. à. d.

$$F_\nu \times (x) = 0 \quad (\nu \geq N),$$

donc d'après (2) pour tout point x de F_ν ($\nu \geq N$), on a

$$\varrho(x, \overline{\sum_{\mu/1}^\infty F_\mu}) \leq \varepsilon$$

et par conséquent

$$\overline{\text{borne}} \varrho(x, \overline{\sum_{\mu/1}^\infty F_\mu}) \leq \varepsilon.$$

De là d'après (1)

$$\varrho^*(F_\nu, \overline{\sum_{\nu/1}^\infty F_\mu}) \leq \varepsilon. \quad (\nu \geq N),$$

ce qui prouve le théorème.

§ 48. Toute suite croissante d'ensembles fermés est convergente. Autrement:

Si $F_\nu = \overline{F_\nu} \neq 0$,

$$(1) \quad F_\nu \subset F_{\nu+1} \quad (\nu/1 \dots \infty),$$

alors

$$\lim \varrho^*(F_\nu, \overline{\sum_{\mu/1}^\infty F_\mu}) = 0.$$

Dém. Soit $\varepsilon > 0$. Posons

$$(2) \quad G_\nu = E[\varrho(x, F_\nu) \geq \varepsilon] \times \overline{\sum_{\mu/1}^\infty F_\mu}.$$

On a (§§ 14 et 7)

$$G_\nu = \overline{G_\nu}$$

et d'après (1) et (2)

$$G_{\nu+1} \subset G_\nu.$$

Je dis qu'il existe au moins un ν , tel que $G_\nu = 0$. Supposons en effet que

$$G_\nu \neq 0 \quad (\nu/1, \dots, \infty).$$

On aurait selon le § 21

$$\prod G_\nu \neq 0$$

et pour tout point s de ce produit on aurait

$$(4) \quad \varrho(s, \overline{\sum_{\nu/1}^\infty F_\nu}) \geq \varepsilon, \quad (\nu/1, \dots, \infty),$$

donc

$$\varrho(s, \overline{\sum_{\nu/1}^{\infty} F_{\nu}}) \geq \varepsilon$$

et par suite

$$\varrho(s, \overline{\sum_{\nu/1}^{\infty} F_{\nu}}) > 0$$

contrairement à (4).

Il existe donc un N , tel que $G_N = 0$.

On a selon (3)

$$G_{\nu} = 0 \quad (\nu \geq N)$$

c. à. d. d'après (2) tout point x de $\overline{\sum_{\mu/1}^{\infty} F_{\mu}}$ satisfait à

$$\varrho(x, F_{\nu}) < \varepsilon, \quad (\nu \geq N)$$

donc

$$\overline{\text{borne}}_{(x \in \overline{\sum_{\mu/1}^{\infty} F_{\mu}})} \varrho(x, F_{\nu}) \leq \varepsilon.$$

D'autre part selon (1)

$$\overline{\text{borne}}_{(f_{\nu} \in F_{\nu})} \varrho(\overline{\sum_{\mu/1}^{\infty} F_{\mu}}, f_{\nu}) = 0,$$

$$\varrho^*(\overline{\sum_{\mu/1}^{\infty} F_{\mu}}, F_{\nu}) \leq \varepsilon \quad (\nu \geq N) \quad \text{c. q. f. d.}$$

§ 49. Déf. On dit que l'ensemble A (fermé ou non) est bien distribué au degré n sur une suite de sphères (fermées ou non)¹

$$S_1, S_2, \dots, S_m$$

si

$$\begin{aligned} 1^{\circ}) \quad & A \times S_{\mu} \neq 0 & (\mu/1, \dots, m) \\ 2^{\circ}) \quad & A \subset S_1 + \dots + S_m \\ 3^{\circ}) \quad & \tau(S_{\mu}) \leq 1/n & (\mu/1, \dots, m). \end{aligned}$$

§ 50. Si chaque ensemble A appartenant à une classe d'ensemble (A) est bien distribué au degré n sur une même suite $S_1, \dots, S_m$, alors il en est de même de la somme de ces ensembles.

¹ Nous dirons aussi qu'une suite est bien distribuée sur un ensemble.

§ 51. A tout n naturel, il appartient une suite finie de sphères rationnelles fermées (ch. e.)

$$(1) \quad \overline{S_1}, \dots, \overline{S_m}$$

telle que tout ensemble non vide est bien distribué sur une suite particulière de (1).

Dém. D'après le théorème de MM. Borel et Lebesgue, il existe une suite finie de sphères rationnelles (ch. e.) $S_1, \dots, S_m$ dont la somme couvre $E(\varrho)$ et dont les rayons sont inférieurs à $1/2n$, donc les diamètres inférieurs à $1/n$. Il est clair que la suite $S_1, \dots, S_m$ a les propriétés proposées.

§ 52. Si (A) est une classe infinie d'ensembles et que n est naturel, alors il existe une suite finie de sphères rationnelles (ch. e.) qui est bien distribuée au degré n sur une infinité de A .

Pour le voir il suffit de remarquer que la classe de suites choisies dans la suite $S_1, \dots, S_m$ intervenant dans le théorème précédent est finie.

§ 53. Si $\mathcal{A}$ est une classe infinie d'ensembles A appartenant à $E(\varrho)$, alors il existe 1°) une suite infinie de classes d'ensembles

$$\mathcal{A}_1, \mathcal{A}_2, \dots$$

telle que $\mathcal{A}_\nu$ est une classe infinie et

$$\mathcal{A}_{\nu+1} \subset \mathcal{A}_\nu \subset \mathcal{A},$$

2°) une suite infinie de suites finies de sphères rationnelles (ch. e.).

$$\Sigma_1, \Sigma_2, \dots,$$

telle que tout A faisant partie de $\mathcal{A}_n$ est bien distribué sur Σ_n au degré n .

C'est une conséquence immédiate du § 51 et du § 52.

§ 54. Si A est bien distribué sur $S_1, \dots, S_m$ au degré n alors

$$\varrho^*(\overline{A}, \overline{S_1 + \dots + S_m}) \leq 1/n.$$

Pour le voir il suffit de rapprocher la définition de ϱ^* (§ 45) et celle de l'ensemble bien distribué sur $S_1, \dots, S_m$ au degré n .

§ 55. $E(\varrho^*)$ renferme une partie dénombrable effectivement énumérable dense sur $E(\varrho^*)$. Cela résulte immédiatement des § 51 et § 54, si l'on remarque que la classe de toutes les suites finies choisies dans la suite de toutes les sphères rationnelles se laisse ranger en une suite infinie (ch. e.).

§ 56. $E(q^*)$ forme une classe fréchétienne compacte (comp. § 1).

Dém. Soit $\mathcal{A}$ une classe infinie d'ensembles A éléments de $E(q^*)$, donc étant des sous-ensembles non vides et fermés de $E(q)$.

Reprenons les notations du § 53 et désignons par (Σ_ν) la somme des sphères de la suite Σ_ν , par $(\mathcal{A}_\nu)$ la somme des A appartenant à $\mathcal{A}_\nu$.

Comme $\mathcal{A}_{\nu+1} \subset \mathcal{A}_\nu$, donc

$$(\mathcal{A}_{\nu+1}) \subset (\mathcal{A}_\nu)$$

$$\text{et} \quad (\overline{\mathcal{A}_{\nu+1}}) \subset (\overline{\mathcal{A}_\nu}).$$

D'autre part $(\mathcal{A}_\nu)$ étant bien distribué sur Σ_ν au degré ν (selon § 50), on a suivant § 54

$$(1) \quad q^*[(\overline{\mathcal{A}_\nu}), (\overline{\Sigma_\nu})] \leq 1/\nu.$$

D'après § 47

$$(2) \quad \lim_{\nu/1} q^*[(\overline{\mathcal{A}_\nu}), \overline{\Pi}(\overline{\mathcal{A}_\nu})] = 0.$$

On a

$$q^*[(\overline{\Sigma_\nu}), \overline{\Pi}(\overline{\mathcal{A}_\nu})] \leq q^*[(\overline{\Sigma_\nu}), (\overline{\mathcal{A}_\nu})] + q^*[(\overline{\mathcal{A}_\nu}), \overline{\Pi}(\overline{\mathcal{A}_\nu})],$$

d' où selon (1), et (2)

$$(3) \quad \lim_{\nu/1} q^*[(\overline{\Sigma_\nu}), \overline{\Pi}(\overline{\mathcal{A}_\nu})] = 0.$$

Pour démontrer le théorème, il suffit de prouver que pour tout $\varepsilon > 0$ il existe une infinité des A , tels que $A \in \mathcal{A}$ et pour lesquels

$$q^*(A, \overline{\Pi}(\overline{\mathcal{A}_\nu})) \leq \varepsilon.$$

Cela résulte de (3) et de cette circonstance que pour tout élément A de la sous-classe infinie $\mathcal{A}_\nu$ de $\mathcal{A}$ on a

$$q^*(A, (\overline{\Sigma_\nu})) \leq 1/\nu,$$

car A est bien distribué sur Σ_ν au degré ν .

§ 57. $E(q^*)$ forme un espace normal.

[§1, § 46, § 56, § 55].

Pour distinguer les notions relatives aux classes $E(q)$ et $E(q^*)$ nous nous exprimerons: Ensemble fermé q^* , ensemble fermé q , fonction continue q^* etc.

§ 58. Si $q^*(F_1, F_2) < \varepsilon$, alors tout voisinage q de rayon ε de tout point f_1 de F_1 empiète sur F_2 . Cette proposition subsiste, si

l'on échange le rôle de F_1 et de F_2 . La démonstration découle de la définition de ϱ^* .

§ 59. La classe de tous les ensembles fermés ϱ appartenant à un ensemble fermé ϱ est fermée ϱ^* .

C'est une conséquence immédiate de la proposition précédente.

§ 60. La classe de tous les ensembles fermés ϱ contenant un point a est fermée ϱ^* (selon § 58).

§ 61. La classe de tous les sous-ensembles A fermés ϱ d'un ensemble F fermé ϱ , tels que tout A contient le même ensemble (fermé ou non) B est fermée ϱ^* .

Il suffit de remarquer qu'elle est la partie commune des classes fermées ϱ^* intervenant dans les § 59 et § 60.

§ 62. Si

$$(1) \quad \varrho^*(F_1 + F_2, F) < \frac{\varrho(F_1, F_2)}{2}$$

alors F n'est pas continu.

Dém. Considérons les ensembles A et B de tous les points f de F pour lesquels on a respectivement

$$(2) \quad \begin{cases} \varrho(f, F_1) \leq \varrho^*(F_1 + F_2, F) \\ \varrho(f, F_2) \leq \varrho^*(F_1 + F_2, F). \end{cases}$$

Evidemment

$$F_1 \subset A, F_2 \subset B.$$

A et B sont donc non vides. Ils sont fermés [(2), § 14]. Ils sont disjoints, car dans le cas contraire en ajoutant les inégalités précédentes on obtiendrait (§ 35):

$$\varrho(F_1, F_2) \leq \varrho(F_1, f) + \varrho(f, F_2) \leq 2\varrho^*(F_1 + F_2, F)$$

contrairement à (1).

Suivant la définition de A et de B on a

$$A + B \subset F$$

On a aussi

$$F \subset A + B.$$

En effet dans le cas contraire il existerait un point f de F pour lequel les inégalités (2) seraient inexactes. On aurait

$$\varrho(f, F_1) > \varrho^*(F_1 + F_2, F)$$

$$\varrho(f, F_2) > \varrho^*(F_1 + F_2, F).$$

donc

$$\varrho(f, F_1 + F_2) > \varrho^*(F_1 + F_2, F)$$

et comme $f \in F$ on en déduit d'après § 46 lem 1:

$$\varrho^*(F, F_1 + F_2) > \varrho^*(F_1 + F_2, F)$$

ce qui est impossible.

On a donc $F = A + B$. C'est une décomposition de F en deux ensembles fermés disjoints et non vides donc F n'est pas continu.

§ 63. La classe de tous les continus ϱ est fermée ϱ^* .

C'est une conséquence du § 62.

§ 64. On appelle un ensemble fermé F saturé par rapport à une classe $\mathcal{F}$ d'ensembles fermés, s'il fait partie de cette classe et s'il n'est pas un vrai sous-ensemble d'aucun ensemble de la classe $\mathcal{F}$.

§ 65. Dans toute classe $\mathcal{F}$ non vide et fermée ϱ^* il existe un ensemble saturé (ch. e.).

Dem. Choisissons dans $\mathcal{F}$ un ensemble F_1 (ch. e. d'après § 22). La classe $\mathcal{F}_1$ de tous les ensembles fermés contenant F_1 est fermée ϱ^* (§ 61), la classe $\mathcal{F}_1 \times \mathcal{F}$ l'est donc aussi et la fonction

$$\varrho^*(F_1, X)$$

définie pour

$$X \in \mathcal{F}_1 \times \mathcal{F}$$

atteint son maximum pour un élément F_2 de $\mathcal{F}_1 \times \mathcal{F}$ (ch. e. § 25).

A partir de F_2 nous construisons l'ensemble F_3 d'une façon analogue. Nous obtenons ainsi une suite infinie (F_ν) , telle que

$$F_\nu \subset F_{\nu+1}, \quad F_\nu \in \mathcal{F}.$$

D'après § 48

$$(1) \quad \lim_{\mu \uparrow \infty} \varrho^*(F_\nu, \overline{\sum_{\mu=1}^{\infty} F_\mu}) = 0$$

et comme la classe $\mathcal{F}$ est fermée ϱ^*

$$\overline{\sum_{\nu=1}^{\infty} F_\nu} \in \mathcal{F}.$$

$\overline{\sum_{\nu=1}^{\infty} F_\nu}$ est saturé. Supposons par l'impossible qu'il existe un ensemble fermé $F (F \in \mathcal{F})$, tel que

$$(2) \quad \begin{aligned} \overline{\sum_{\mu=1}^{\infty} F_\mu} &\subset F \in \mathcal{F} \\ \overline{\sum_{\mu=1}^{\infty} F_\mu} &\neq F. \end{aligned}$$

On a

$$(3) \quad \varphi^*(\overline{\Sigma F_\mu}, F) > 0 \\ F_\nu \subset F \quad (\nu/1, \dots, \infty)$$

Soit N le plus petit indice pour lequel $\nu \geq N$ entraîne

$$(4) \quad \varphi^*(F_\nu, \overline{\Sigma F_\mu}) < \varphi^*(\overline{\Sigma F_\mu}, F)$$

N existe en raison de (1) et (3).

On a selon (4)

$$(5) \quad \varphi^*(F_N, F_{N+1}) \leq \varphi^*(F_N, \overline{\Sigma F_\mu}) + \varphi^*(\overline{\Sigma F_\mu}, F_{N+1}) < \frac{2}{3} \varphi^*(\overline{\Sigma F_\mu}, F).$$

D'autre part

$$\varphi^*(F, \overline{\Sigma F_\mu}) \leq \varphi^*(F, F_N) + \varphi^*(F_N, \overline{\Sigma F_\mu}),$$

donc suivant (4)

$$(1 - \frac{1}{3}) \varphi^*(F, \overline{\Sigma F_\mu}) < \varphi^*(F, F_N),$$

d'où d'après (5)

$$\varphi^*(F_N, F_{N+1}) < \varphi^*(F_N, F),$$

et c'est contraire à la définition de F_{N+1} pour lequel la fonction $\varphi^*(F_N, X)$ définie pour $F_N \subset X \in \mathcal{F}$ atteint son maximum.

B.

Classe faible. Courbe de Jordan. Homéomorphie.

Bornons nous à la classe $E(\varrho)$.

§ 1. On appelle une classe d'ensembles faible, si

1°) pour tout ensemble F de cette classe

$$\tau(F) > 0,$$

2°) à tout $\varepsilon > 0$ correspond un nombre fini seulement d'ensembles pour lesquels

$$\tau(F) \geq \varepsilon.$$

$\tau(F)$ est le diamètre de F .

§ 2. Toute classe faible d'ensembles fermés est effectivement énumérable.

En effet elle se laisse représenter comme la somme

$$(F)_1 + (F)_2 + \dots$$

où $(F)_1$ est la classe de tous les F en question pour lesquels

$\tau(F) \geq 1$. et $(F)_\nu$ ($\nu \geq 2$) est la classe de tous les F pour lesquels

$$1/\nu \leq \tau(F) < 1/\nu - 1.$$

Toute classe $(F)_\nu$ est vide ou finie, donc dans le dernier cas effectivement énumérable (A § 40). Le théorème en résulte immédiatement.

§ 3. Si (K) est une classe faible de continus empiétant sur un continu C , alors

$$C + (K)$$

est un continu

Dém. α) $C + (K)$ est fermé.

Supposons le contraire. Il existe donc au moins un point c , tel que

$$(1) \quad c \in \overline{C + (K)} - [C + (K)].$$

On a donc

$$(2) \quad \varrho(c, C) > 0$$

$$(3) \quad \varrho(c, K) > 0 \quad (\text{pour tout } K).$$

Imaginons la classe (K) rangée en une suite (c'est possible selon § 2)

$$K_1, K_2, \dots$$

Il existe au plus un nombre fini de continus K_ν , tels que

$$\tau(K_\nu) \geq \frac{\varrho(c, C)}{2}.$$

Soit K_ν celui d'entre eux qui a l'indice le plus grand. D'après (2) et (3)

$$(4) \quad \varrho(c, C + K_1 + \dots + K_N) > 0.$$

Entourons c d'une sphère S de rayon $r > 0$ inférieur à $\frac{\varrho(c, C)}{2}$ et $\frac{\varrho(c, C + K_1 + \dots + K_N)}{2}$.

Comme

$$C \times K_{N+\mu} \neq 0 \quad (\mu \geq 1)$$

et

$$\tau(K_{N+\mu}) < \frac{\varrho(c, C)}{2},$$

il est clair que

$$\sum_{\mu=1}^{\infty} K_{N+\mu} \times S = 0,$$

done selon (3) et (4)

$$\varrho(c, C + (K)) \geq r > 0$$

et par suite

$$\varrho(c, \overline{C + (K)}) > 0$$

contrairement à (1).

$\beta)$ $C + (K)$ est évidemment bien enchaîné, ce qui avec α prouve le théorème.

Orbe de Jordan.

§ 4. Nous appelons orbe de Jordan ou simplement orbe l'image homéomorphe de la circonférence.

On a le théorème:

Toute somme de deux arcs simples a, b et a, b_1 non identiques ayant les extrémités communes contient une orbe.

Dém. En partant d'un point de a, b par exemple qui n'appartient pas à a, b_1 et en s'avancant d'abord vers a et après vers b , on rencontre l'arc a, b_1 aux points a' et b' . La somme des arcs a', b' et a', b'_1 découpés par a', b' sur a, b et a, b_1 donne une orbe de Jordan.

§ 5. Si une orbe est contenue dans une somme de deux continus dont le produit se réduit à un seul point, alors elle est contenue dans un seul d'eux.

La démonstration se fait par une simple application de la définition du continu.

Courbe de Jordan.

§ 6. La définition de la courbe de Jordan a été donnée dans l'introduction.

Pour qu'un ensemble E soit une courbe de Jordan chacune de couples de conditions $\alpha, \beta; \alpha, \beta'; \alpha, \beta''$ est nécessaire et suffisante:¹

$\alpha)$ E est un continu.

$\beta)$ Pour tout $\varepsilon > 0$ il existe $\delta > 0$, tel que chaque paire de points de E dont la distance est inférieure à δ se laisse joindre par un sous-continu de E au diamètre $\leq \varepsilon$.

¹ St. Mazurkiewicz: Sur les lignes de Jordan théorème V et IX. *Fund. Math.* T. I. 1920.

W. Sierpiński: Sur une condition etc. *ibid.*

β') Cette condition ne diffère de la précédente que par le remplacement du mot sous-continu par sous-arc simple.

β'') A tout $\varepsilon > 0$ il existe une décomposition de E en un nombre fini de continus de diamètre $\leq \varepsilon$.

Remarque. En partant des propriétés des espaces $E(\varrho)$ et $E(\varrho^*)$, on pourrait facilement prouver que tous les choix auxquels donnent lieu les théorèmes précédents peuvent être faits effectivement. Leurs démonstrations se laissent aussi affranchir de l'axiome de M. Zermelo.

§ 7. Si deux point a, b d'une courbe de Jordan J se laissent joindre par un sous-continu K ne renfermant pas un point c de J , alors ils se laissent joindre par un arc simple (ch. e.) ayant la même propriété.

Dém. On a

$$\varrho(c, K) > 0.$$

Posons

$$0 < \varepsilon < \frac{\varrho(c, K)}{2}.$$

Il existe (§ 6) un $\delta > 0$, tel que toute couple de points x, y appartenant à J se laisse joindre par un arc simple au diamètre $\leq \varepsilon$ (ch. e.)

Le continu K étant bien enchaîné, il existe une suite de points (on en pourrait indiquer une qui est parfaitement déterminée)

$$a_0, a_1, \dots, a_n,$$

telle que

$$(2) \quad \begin{array}{ll} a_0 = a, & a_n = b \\ a_\nu \in K & (\nu/0, \dots, n) \\ \varrho(a_\nu, a_{\nu+1}) \leq \delta & (\nu/0, \dots, n-1). \end{array}$$

En joignant a_ν et $a_{\nu+1}$ par un arc simple $a_\nu, a_{\nu+1}$ contenu dans J au diamètre $\leq \varepsilon$ (ch. e.), nous avons d'après (1) et (2)

$$(3) \quad (|a_0, a_1| + \dots + |a_{n-1}|) \times \varepsilon = 0.$$

On pourrait facilement démontrer que la somme figurant dans (3) contient un arc simple (ch. e.) joignant a et b et contenu dans J .

§ 8. Tout continu qui est la somme d'un nombre fini de courbes de Jordan est une courbe de Jordan (§ 6, β'').

Homéomorphie.

§ 9. Si $\sim$ est une correspondance réciproque biunivoque et continue entre deux ensembles F et F' appartenant à deux espaces normaux $E(\rho)$ et $E(\rho')$ et si F est fermé, alors

1°) F' est fermé,

2°) la correspondance $\sim$ est biunivoque et bicontinue c. a. d. homéomorphe.

Nous n'insistons pas sur la démonstration.

§ 10. Si

1°) F est un continu et (F_ν) une classe faible de continus que nous supposons rangés en une suite,

2°) $F \times F_\nu \neq 0$ pour tout F_ν de (F_ν) ,

3°) (Φ_ν) est une classe faible de continus que nous supposons rangés en une suite,

4°) (Φ_ν) et (F_ν) ont le même nombre d'éléments,

5°) $(F_\nu - F) \times (F_\mu - F) = 0$; $(\Phi_\nu - \Phi) \times (\Phi_\mu - \Phi) = 0$
($\mu \neq \nu$),

6°) H et H_ν sont des homéomorphies entre les ensembles suivants:

$$\begin{array}{ccc} F & H & \Phi \\ F_\nu & H_\nu & \Phi_\nu \\ (F_\nu \cdot F) & H & (\Phi_\nu \cdot \Phi). \end{array}$$

telles que H_ν est identique à H sur $F_\nu \cdot F$, alors la correspondance $\sim$ définie comme identique à H , s'il s'agit des ensembles F et Φ et comme identique à H_ν , s'il s'agit des ensembles F_ν et Φ_ν , constitue une homéomorphie entre

$$(1) \quad F + (F_\nu) \quad \text{et} \quad \Phi + (\Phi_\nu).$$

Dém. La correspondance $\sim$ est biunivoque d'après 5° et la troisième condition de 6°.

Les ensembles (1) étant fermés (§ 3), il suffit de démontrer que la correspondance est continue (§ 9). C'est presque évident, si (F_ν) contient un nombre fini d'éléments

Soit dans le cas de (F_ν) infini f un point de $F + (F_\nu)$, φ son correspondant dans $\Phi + (\Phi_\nu)$ et $\varepsilon > 0$. Il existe un N tel que

$$\begin{array}{l} \varphi \varepsilon \Phi + \sum_{\nu=1}^N \Phi_\nu \\ (2) \quad \tau(\Phi_{N+\mu}) \leq \varepsilon/2 \quad (\mu \geq 1) \end{array}$$

$\sim$ établit suivant ce qui vient d'être dit une homéomorphie entre

$$F + \sum_{v=1}^{\infty} F_v \quad \text{et} \quad \Phi + \sum_{v=1}^{\infty} \Phi_v.$$

Traçons autour du point φ une sphère de rayon $\varepsilon/2$.

Il existe un $\delta > 0$, tel que l'image de

$$(3) \quad (F + \sum_{v=1}^{\infty} F_v), \text{ Sph. } (f, \delta)$$

est contenu dans

$$(4) \quad (\Phi + \sum_{v=1}^{\infty} \Phi_v), \text{ Sph. } (\varphi, \varepsilon/2).$$

L'image de tout F_{v+1} qui empiète sur (3) empiète sur (4), donc d'après (2) cet image est contenue dans la sphère Sph. (φ, ε) , et cela prouve la continuité.

§ 11. Si a, b et a_1, b_1 sont deux arcs simples aux extrémités indiquées dans leurs symboles, si (j) est une classe de points effectivement énumérable contenue dans

$$: a, b : = a, b | - a - b$$

et que (j_1) est une classe de points effectivement énumérable dense sur a_1, b_1 et comprise dans $: a_1, b_1 :$, alors il existe une homéomorphie $\sim$ entre $: a, b$ et $: a_1, b_1$ telle que

$$(1) \quad \begin{aligned} a &\sim a_1 \\ b &\sim b_1 \end{aligned}$$

et suivant laquelle tout j correspond à un certain j_1 (ch. e.).

Dém. Il est possible de trouver un sous ensemble D effectivement énumérable et dense sur $: a, b$ (selon A § 43).

$(j)^* = D + (j) - a - b$ forme un ensemble effectivement énumérable dense sur $: a, b$ et contenu dans $: a, b :$. Il suffit de démontrer l'existence d'une homéomorphie (ch. e.) dans laquelle (1) a lieu et qui fait correspondre à tout point de $(j)^*$ un point de (j_1) .

Suivant un procédé classique on range $(j)^*$ et (j_1) en deux suites infinies

$$(2) \quad \begin{aligned} j^1, j^2, \dots \\ j_1^1, j_1^2, \dots \end{aligned}$$

On fait correspondre à a le point a_1 et à b le point b_1 , à j^1 , le point j_1^1 . Au point j^2 on associe le premier point j_1^2 , tel que les quatre points a, b, j^1, j^2 soient ordonnés de la même façon sur $: a, b$ que

les points $a_1, b_1, j_1^1, j_1^{21}$, sur a_1, b_1 . On continue de cette façon jusqu'à épuiser tous les j^v . On démontre facilement qu'en faisant correspondre au point limite de toute suite particulière de (2)

$$j^{j_1}, j^{j_2}, \dots$$

le point limite de la suite

$$j_1^{\alpha\beta_1}, j_1^{\alpha\beta_2}, j_1^{\alpha\beta_3}, \dots$$

On obtient une homéomorphie satisfaisant aux conditions du théorème.

C.

Propriétés générales de la dendrite.

J'appelle dendrite toute courbe de Jordan contenue dans un espace normal $E(z)$ (p. exemple dans un morceau fini d'un espace cartésien à n dimensions) et ne contenant aucune orbe de Jordan.

Nous désignons par D une dendrite quelconque supposée la même au courant des considérations de ce chapitre. Nous traiterons seulement les sous-ensembles de D .

§ 1. Si a et b sont deux points de D , alors il existe un seul arc simple a, b joignant ces points et les ayant comme extrémités¹⁾.

Dém. C'est évident, si $a = b$. On a alors $a, b = a$. Si, pour $a \neq b$, il existait deux arcs a, b et a, b_1 , de cette sorte, alors $a, b + a, b_1$ contiendrait une orbe (B § 4), ce qui est contraire à la définition de la dendrite.

D'autre part il existe au moins un arc a, b d'après B § 6 β'.

§ 2. Déf. Nous désignerons par a, b , dans la suite, sauf mention expresse, les arcs positifs c. à d. pour lesquels $\varphi(a, b) > 0$.

Nous posons par définition

$$\begin{aligned} :a, b = b, a: &= a, b - a \\ :a, b: &= a, b - a - b. \end{aligned}$$

§ 3. Si $a \neq b$,
 $a + b \subset K$,
 K est un sous-continu de D ,

alors on a

$$a, b \subset K^2).$$

¹⁾ St. Mazurkiewicz: Un théorème sur les lignes de Jordan Fund. Math. T. II p. 123, lemme 2.

²⁾ St. Mazurkiewicz, l. c.

Dém. Dans le cas contraire, il existerait un point c pour lequel

$$\begin{aligned} c \varepsilon :a, b: \\ c \times K = 0. \end{aligned}$$

Il existerait donc (B § 7) un sous-arc simple joignant a, b et n'empiétant pas sur c partant différent de a, b contrairement au § 1.

§ 4. Si $b \times K = 0$ (où K est un continu), alors il existe un seul arc a, b pour lequel

$$(1) \quad a, b \times K = a.$$

Nous le désignerons par K, b et nous posons

$$:K, b := :a, b, \quad K, b := a, b: \quad :K, b := :a, b:$$

Nous appelons K, b feuille complète poussant sur K ,

$:K, b$ feuille restreinte poussant sur K ,

K tige de ces feuilles.

$K \times K, b = K, b + :K, b = a$ jonction de ces feuilles:

$$a = j(K, b).$$

Dém. Il existe au moins un arc de la sorte. En effet soit k un point quelconque de K et soit a le premier point dans lequel on rencontre K en partant de b et en s'avancant sur b, k . Évidemment a, b satisfait à (1).

Supposons qu'il existe deux arcs de cette sorte a, b et a_1, b .

D'après § 1 on a $a \neq a_1$

et de (1) résulte

$$(2) \quad (a_1, b + a, b) \times K = a + a_1$$

donc suivant le § 3

$$(3) \quad a, a_1 \subset a, b + a_1, b$$

$$(4) \quad a, a_1 \subset K.$$

En multipliant (3) par (4) on obtient

$$a, a_1 \subset a + a_1,$$

ce qui est impossible.

Remarque. On a évidemment

$$\begin{aligned} :K, b \times K &= 0 \\ j(K, b) &\varepsilon K \\ K, b &= j(K, b), b. \end{aligned}$$

§ 5. Tout sous-continu de toute dendrite est une dendrite¹⁾.

Dém. Considérons un sous-continu K quelconque de 0 et $\varepsilon < 0$.

D'après le théorème de M. Mazurkiewicz (B § 6 β'), il existe $\delta > 0$, tel que toute couple k_1, k_2 de points de K dont la distance est inférieure à δ se laisse joindre par un sous-arc simple de D au diamètre $< \varepsilon$. Or ce sous-arc est contenu dans K , car ces extrémités le sont (§ 3). Suivant le théorème cité K est une courbe de Jordan. K comme un sous-ensemble d'une dendrite ne renferme aucune orbe, il est donc dendrite.

§ 6. La classe de tous les sous-arcs simples d'une dendrite est fermée φ^* (Le point est considéré comme un arc simple).

Nous démontrerons d'abord le

Lemme. L'arc a, b est une fonction continue φ^* de la somme $a + b$.

Il faut démontrer qu'à tout $\varepsilon > 0$, il existe $\delta > 0$, tel que

$$(1) \quad \varphi^*(a + b, a_0 + b_0) \leq \delta$$

entraîne

$$\varphi^*(a, b, a_0, b_0) \leq \varepsilon.$$

Cas 1°. $a_0 = b_0$.

Le diamètre de $a + b$ comme une fonction continue φ^* tend vers 0, si $a + b$ tend vers $a_0 + b_0$ et par suite (B § 6 β') le diamètre de a, b tend vers 0. Il est évident d'après la définition de φ^* (A § 45) que a, b tend vers a_0, b_0 .

Cas 2°. Supposons $\varphi(a_0, b_0) > 0$ et posons

$$(2) \quad 0 < \varepsilon_1 < \frac{\varphi(a_0, b_0)}{3}$$

$$0 < \varepsilon_1 < \varepsilon.$$

Il existe un

$$\delta < 0$$

$$\delta < \frac{\varphi(a_0, b_0)}{3},$$

tel que

$$\varphi(x, y) \leq \delta$$

entraîne

$$\tau(x, y) \leq \varepsilon_1.$$

Supposons (1) vrai pour un tel δ . En permutant au besoin a et b

¹⁾ St Mazurkiewicz, l. c.

nous avons selon la définition de ϱ^* (A § 45)

$$(3) \quad \begin{aligned} \varrho(a, a_0) &\leq \delta \\ \varrho(b, b_0) &\leq \delta \\ \tau(a, a_0) &\leq \varepsilon_1 \\ \tau(b, b_0) &\leq \varepsilon_1 \end{aligned}$$

De (2) et (3) on obtient

$$a, a_0 \times b, b_0 = 0.$$

Désignons par c le point dans lequel l'arc a, b quitte pour la dernière fois l'arc a, a_0 et par d le point dans lequel il rencontre pour la première fois l'arc b, b_0 .

On a

$$(4) \quad \begin{aligned} a_0, b_0 &= a_0, c + c, d + d, b_0 \\ a, b &= a, c + c, d + d, b. \end{aligned}$$

Ces décompositions et (3) mettent en évidence que

$$\varrho^*(a, b, a_0, b_0) \leq \varepsilon_1 < \varepsilon \quad \text{c. q. f. d.}$$

La classe de toutes les sommes de points $a + b$ (le cas $a = b$ étant admis) est évidemment fermée ϱ^* . La classe de tous les a, b comme l'image continue de cette classe (Lemme précédent, A § 27) est donc aussi fermée ϱ^* .

§ 7. Corrolaire. Si $\varrho(a_0, b_0) > 0$, alors pour $a + b$ assez approché ϱ^* de $a_0 + b_0$ on a

$$a_0, b_0 \times a, b \neq 0.$$

C'est une conséquence de la décomposition (4), § précédent.

§ 8. Est faible toute classe (a, b) composée d'arcs positifs disjoints deux à deux (comp. B § 1).

Ce théorème reste évidemment vrai, si l'on remplace la classe (a, b) par la classe $(:a, b)$ ou la classe $(:a, b:)$.

Dém. Dans le cas contraire il existerait évidemment une infinité d'arcs a_1, b_1 , tels que

$$\varrho(a_1, b_1) \geq \varepsilon_0 > 0.$$

La classe de sommes $(a_1 + b_1)$ étant infinie, une somme $a^* + b^*$ existe dont tout voisinage contient une infinité des sommes $a_1 + b_1$ (A §§ 1 et 56). ce qui est impossible d'après le corrolaire précédent.

§ 9. A tout a, b correspond un arc saturé (ch. e.) qui le contient¹⁾.

Dém. La classe des arcs simples contenant a, b est fermée ϱ^* comme le produit de la classe fermée ϱ^* de tous les ensembles fermés ϱ qui contiennent a, b (A § 59) et de la classe de tous les sous-arcs simples de D qui est fermée ϱ^* (§ 6).

Cette classe admet donc un arc saturé contenant a, b (ch. c. A § 65).

§ 10. Déf. On dit que l'arc $[a, b]$ est saturé au point b , s'il est impossible de le prolonger au delà de b .

§ 11. Tout arc $[a, c]$ qui n'est pas saturé au point c peut être prolongé au delà de c de façon à être saturé à la nouvelle extrémité. C'est une conséquence du § 9.

§ 12. La condition nécessaire et suffisante que l'arc $[a, b]$ (positif) soit saturé au point b est qu'il n'existe aucun arc $[c, d]$ contenant b entre ses extrémités.

Dém. La condition est évidemment suffisante. Elle est nécessaire. Supposons en effet qu'un arc $[c, d]$ de la sorte existe. On aurait, $[a, b]$ étant saturé dans b

$$\begin{aligned} :b, c] \times [a, b] &\neq 0 \\ :b, d] \times [a, b] &\neq 0. \end{aligned}$$

Soient e et f un point du premier et un point du second produit. On aurait selon le § 3

$$(1) \quad \begin{aligned} :b, e] &\subset :b, c] \\ :b, f] &\subset :b, d] \\ :b, e] + :b, f] &\subset [a, b]. \end{aligned}$$

La dernière relation donne (e et f étant $\neq b$ d'après leur définition)

$$(2) \quad :b, e] \times :b, f] \neq 0$$

et d'autre part en multipliant les relations (1) l'une par l'autre on obtient

$$:b, e] \times :b, f] \subset :b, c] \times :b, d] = 0,$$

ce qui est contraire à (2).

¹⁾ St. Mazurkiewicz, l. c. p. 129. La démonstration de M. Mazurkiewicz est basée sur l'axiome de M. Zermelo et fait l'usage des nombres transfinis. Notre procédé qui évite les nombres transfinis est une généralisation d'un procédé adopté par M. A. Denjoy dans le cas d'une famille d'intervalles fermée ϱ^* . (Cnf. An. Scient. Ec. norm 1916).

§ 13. Si un arc a, b (positif) est saturé au point b , alors tout arc contenant b y est saturé aussi (§ 12). Nous appelons pour cette raison tout point tel que b point de saturation de la dendrite D .

§ 14. Si un continu K ne se réduit pas à un seul point, alors $j(K, b)$ ne coïncide avec aucun point de saturation de K (comp. § 4).

En effet dans le cas contraire on pourrait prolonger K, b au delà de sa jonction en ajoutant un arc $j(K, b), c$ où

$$\begin{aligned} c &\in K \\ c &\neq j(K, b). \end{aligned}$$

§ 15. On dit qu'un arc a, b est en T^1 par rapport à a et c, d , si

$$a, b \times : c, d = a.$$

On vérifie facilement que dans ce cas

$$a, b \times [c, d] = a.$$

On dit que: a, b est en T par rapport à c, d , si $[a, b]$ est en T par rapport à a et c, d .

§ 16. a, b est en T par rapport à a et c, d , s'il est en T par rapport à a et à un arc partiel de c, d . Cette proposition résulte du lemme:

Lemme. Si K est un continu ne se réduisant pas à un seul point et que

$$\begin{aligned} K \times : c_1, d_1 &= b \\ c_1 d_1 &\subset [c, d], \end{aligned}$$

alors

$$: c, d : \times K = b.$$

L'hypothèse que le lemme est faux conduit à l'existence d'une orbe contenue dans D , ce qui est impossible.

§ 17. On dit qu'un point b d'une dendrite D (ou d'un continu K est un point de branchement au sens large, s'il existe au moins

¹⁾ La barre verticale de la lettre T est en T par rapport à son extrémité supérieure et par rapport à la barre horizontale.

trois continus K_1, K_2, K_3 ne se réduisant pas à un point, tels que

$$\begin{aligned} K_i \times K_j &= b \quad (i \neq j) \\ K_i &\subset D \quad (i/1, 2, 3). \end{aligned}$$

§ 18. Si un point b , tel que

$$(1) \quad b \varepsilon : c, d :$$

est un point de branchement au sens large, il existe un arc b, a étant en T par rapport à b et c, d (ch. e.).

Dém. A lieu au moins une de relations suivantes dans lesquelles K_1, K_2, K_3 ont la même signification que dans § 17

$$(2) \quad \begin{aligned} c, d \times K_1 &= b \\ c, d \times K_2 &= b \\ c, d \times K_3 &= b \end{aligned}$$

Dans le cas contraire deux K empièteraient sur un des arcs b, c, d p. ex. on aurait

$$\begin{aligned} :b, c \times K_1 &\neq 0 \\ :b, c \times K_2 &\neq 0. \end{aligned}$$

Il existerait donc deux points c_1 et c_2 , tels que

$$\begin{aligned} c_1 &\varepsilon :b, c \times K_1 \\ c_2 &\varepsilon :b, c \times K_2 \end{aligned}$$

et on aurait selon § 3

$$\begin{aligned} :b, c_1 &\subset :b, c \times K_1 \\ :b, c_2 &\subset :b, c \times K_2, \end{aligned}$$

d'où

$$\begin{aligned} :b, c_1 \times :b, c_2 &\subset :b, c \\ :b, c_1 \times :b, c_2 &\subset K_1 \times K_2 \end{aligned}$$

La première de ces relations donne

$$:b, c_1 \times :b, c_2 \neq 0$$

done d'après la seconde relation

$$K_1 \times K_2 - b \neq 0$$

contrairement à la définition de K_1, K_2, K_3 (§ 17).

Nous pouvons supposer que la relation (2) a lieu.

Soit k_1 un point de $K - b$ (ch. e. A § 39).

On a selon § 3

$$b, k_1 \subset K_1$$

et en multipliant les deux membres de cette relation par c, d et en tenant compte de (1) et (2) on obtient

$$:c, d: \times b, k_1 = b \quad \text{c. q. f. d.}$$

§ 19. Si b est un point de branchement au sens large, alors il existe un arc c, d , tel que

$$b \varepsilon :c, d: \quad (\text{ch. e.}).$$

La démonstration se déduit facilement de la définition § 17.

§ 20. Aucun point de saturation n'est point de branchement au sens large.

C'est une conséquence immédiate de la proposition précédente et du § 12.

D.

Constituant¹⁾.

§ 1. On appelle constituant par rapport à un continu K (pouvant se réduire à un seul point) d'une dendrite D correspondant à un point b de D n'appartenant pas à K la somme de tous les sous-continus de D renfermant b et n'empiétant par sur K . Nous le désignerons par

$$:S(K, b)$$

et nous posons

$$S(K, b) = \overline{:S(K, b)}.$$

Pour simplifier nous appelons $S(Kb)$ constituant complet, $:S(K, b)$ constituant restreint.

Voici quelques conséquences de cette définition :

§ 2. $K \times :S(K, b) = 0$.

$$b \varepsilon :S(K, b).$$

$S(K, b)$ est un continu.

§ 3. La condition nécessaire et suffisante pour qu'un continu K_1 , tel que

$$K_1 \times K = 0$$

¹⁾ Nous nous occupons dans ce chapitre et dans les suivants seulement des sous-ensembles d'une seule dendrite D .

soit contenu dans : $S(K, b)$ est que

$$: S(K, b) \times K_1 \neq 0.$$

§ 4. $: S(K, b) = : S(K_1, b_1)$

équivalent à la paire de conditions

1°) existe un point c , tel que

$$c \varepsilon : S(K, b) \times : S(K_1, b_1)$$

2°) tout continu contenant c et n'empiétant pas sur K n'empiète pas sur K_1 et inversement.

§ 5. Si $c \varepsilon : S(K, b)$,
alors

$$: S(K, b) = : S(K, c).$$

§ 6. La condition nécessaire et suffisante que

$$: S(K, b) = : S(K, c)$$

est que

$$: S(K, b) \times : S(K, c) \neq 0.$$

§ 7. Si $c + d \subset : S(K, b)$,
alors

$$[c, d] \subset : S(K, b).$$

§ 8. $: a, b] \subset : S(a, b)$
 $: K, b] \subset : S(K, b)$
 $[a, b \subset S(a, b)$
 $[K, b \subset S(K, b)$ (comp. § 4 C).

§ 9. Tout point c , tel que

$$c \varepsilon D - K$$

appartient à un seul : $S(K, b)$

et on a

$$D = K + (: S(K, b))$$

où $()$ est la somme de tous les constituants par rapport à K .

§ 10. α) A tout point c d'un : $S(K, b)$ correspond un voisinage de c n'empiétant pas sur $D - : S(K, b)$.

Autrement nous exprimerons ce fait :

$\beta) : S(K, b)$ est un ensemble ouvert ¹⁾ par rapport à la dendrite D .

Dém. Soit $c \in : S(K, b)$,
selon le § 2

$$\varrho(c, K) > 0.$$

D'après le théorème de M. Mazurkiewicz, il existe un $\delta > 0$, tel que tout point de la sphère $\text{Sph}(c, \delta)$ appartenant à D se laisse joindre au point c par un continu au diamètre inférieur à $\varrho(c, K)$. Suivant la définition du constituant on a donc

$$\text{Sph}(c, \delta) \times D \subset : S(K, b)$$

d'où résulte facilement le théorème.

§ 11. $K + (: S(K, b))'$ où la somme $()'$ est étendue à quelques constituants est un continu.

Dém. Cette somme est égale à la différence

$$D - (: S(K, b))''$$

où $()''$ est la somme de tous les constituants restants, car les constituants restreints sont disjoints deux à deux (§ 6).

Cette différence est fermée d'après la note ¹⁾.

Il reste à démontrer que $K + (: S(K, b))'$ est bien enchaîné. On a suivant ce qui précède :

$$K + (: S(K, b))' = K + (\overline{(: S(K, b))})' = K + (S(K, b))'.$$

On a (§ 8 et C § 4)

$$(1) \quad 0 \neq K, K, b' \subset K, S(K, b).$$

$(S(K, b))'$ forme donc une classe de continus empiétant sur K . $K + ()'$ est donc bien enchaîné c. q. f. d.

§ 12. L'ensemble

$$(1) \quad D - : S(K, b)$$

est un continu et $: S(K, b)$ est son unique constituant restreint. On a

$$(2) \quad : S(K, b) = : S([D - : S(K, b)], b).$$

Dém. (1) est un continu selon le § 11.

¹⁾ On a sans peine: Le complémentaire par rapport à D d'un ensemble ouvert par rapport à lui est fermé. De là (A § 7): Toute somme d'ensembles ouverts est ouverte (relativement à D).

Cela étant l'égalité (2) résulte immédiatement de la définition du constituant restreint.

Ce théorème nous permettra souvent de se borner au cas où il n'existe qu'un seul constituant.

$$\begin{aligned} \S 13. \text{ Déf. jonction } [:S(K, b)] &= \overline{ :S(K, b) } - :S(K, b) = \\ &= S(K, b) - :S(K, b). \end{aligned}$$

Nous la désignons : $j[:S(K, b)]$.

§ 14. Évidemment, si $:S(K, b) = :S(K_1, b_1)$, alors

$$j(:S(K, b)) = j(:S(K_1, b_1)).$$

§ 15. $j(:S(K, b)) = K \times S(K, b) = 1$ point.

Dém. D'après le § 14 et le § 12 il suffit de se borner au cas où $:S(K, b)$ est l'unique constituant par rapport à K .

On a

$$\begin{aligned} D &= K + :S(K, b) \\ K \times :S(K, b) &= 0 \end{aligned}$$

d'où facilement

$$(1) \quad \overline{ :S(K, b) } - :S(K, b) = K \times \overline{ :S(K, b) }$$

c. à. d. (§ 13)

$$j(:S(K, b)) = K \times S(K, b).$$

Selon (1) et la relation (1) du § 11

$$j(:S(K, b)) \neq 0.$$

Il reste à démontrer que la jonction en question ne peut pas contenir deux points différents. Supposons en effet qu'il existe c et d , tels que

$$c + d \subset \overline{ :S(K, b) } - :S(K, b) = K \times \overline{ :S(K, b) }.$$

On aurait de là (C § 3)

$$(2) \quad c, d \subset K$$

et d'autre part, il existerait deux points c_1 et d_1 variables dans $:S(K, b)$ et tendant respectivement vers c et d . On aurait (selon § 7)

$$(3) \quad c_1, d_1 \subset :S(K, b)$$

et pour c_1 et d_1 assez voisins de c et d on aurait

$$c, d \times c_1, d_1 \neq 0 \quad (\text{C § 7}),$$

donc selon (2) et (3)

$$K \times :S(K, b) \neq 0$$

ce qui est impossible.

§ 16. De la relation (1) du § 11 on a selon le théorème précédent

$$j(:S(K, b)) = j(K, b) \quad (\text{comp. C § 4})$$

et plus particulièrement

$$j(:S(a, b)) = a.$$

§ 17.

$$(1) \quad :S(K, b) \times :S(K, c) = 0$$

équivalent à

$$(2) \quad :K, b \times :K, c = 0.$$

Dém. Cette condition est nécessaire, car, si (2) n'avait pas lieu, on aurait (§ 5):

$$0 \neq :K, b \times :K, c \subset :S(K, b) \times :S(K, c).$$

(2) est aussi condition suffisante de (1). Supposons en effet que (1) est faux et (2) est vrai. On a donc (§ 6 et § 8)

$$:S(K, c) = :S(K, b)$$

$$(3) \quad :K, b + :K, c \subset :S(K, b)$$

et d'après (2)

$$b \neq c$$

Suivant § 7

$$(4) \quad \subset :S(K, b).$$

De (3) on a selon § 16:

$$j(K, b) = j(K, c) = j[:S(K, b)]$$

donc d'après (2)

$$K, b + K, c$$

est un arc simple aux extrémités b et c et renfermant $j[:S(K, b)]$.

On a par conséquent selon (4)

$$j[:S(K, b)] \subset b, c \subset :S(K, b),$$

ce qui est impossible suivant la définition de la jonction.

§ 18. La classe de tous les $:S(K, b)$

$$(1) \quad (:S(K, b))$$

est a) faible b) effectivement énumérable;

c) La classe $(j[:S(K, b)])$ est effectivement énumérable.

Dém. Évidemment

$$\tau[:S(K, b)] > 0.$$

Supposons que la classe (1) n'est pas faible. Il existe donc une sous-classe infinie de (1), telle que pour tout son élément

$$\tau[:S(K, b)] > \varepsilon_0 > 0.$$

Comme : $S(K, b)$ est une différence de deux ensemble fermés :

$$:S(K, b) = S(K, b) - \text{jonc. } [:S(K, b)],$$

donc en se basant sur A § 39 on peut choisir dans tout : $S(K, b)$ en question deux points c et d (ch. e.), tels que

$$\begin{aligned} \varrho(c, d) &\geq \varepsilon_0/2 \\ c + d &\subset :S(K, b). \end{aligned}$$

On a (§ 7 et § 6)

$$\begin{aligned} c, d &\subset :S(K, b) \\ c, d \times c, d &= 0 \end{aligned}$$

(pour deux arcs différents) et c'est impossible d'après C § 7.

a) est ainsi démontré. b) et c) en résultent immédiatement (B § 2).

§ 19. Si deux b différents d'une classe $(b)'$ sont contenus dans deux : $S(K, b)$ différents, alors

$$K + (:K, b)' = K + (:K, b)'$$

et les deux membres de cette égalité sont des continus.

Dém. L'égalité est vrai, car (C § 4, § 16, § 15)

$$K, b = :K, b + j(K, b)$$

$$(1) \quad j(K, b) = j[:S(K, b)] = K \times S(K, b) \neq 0.$$

Suivant

$$:K, b \subset :S(K, b)$$

la classe $(:K, b)'$ est faible. D'après (1) tout : K, b empiète sur K , donc $K + (:K, b)'$ est un continu (B § 3).

§ 20. Tout continu K_1 joignant un point d'un constituant (restreint ou complet) par rapport à K , avec un point du complémentaire de K relativement à D passe par la jonction de : $S(K, b)$.

Dém. Il suffit évidemment (§ 12) de traiter le cas. où il existe un seul constituant. On a suivant la prémisse du théorème

dans le cas d'un constituant restreint et complet :

$$(1) \quad K \times K_1 \neq 0$$

$$(2) \quad K_1 \times S(K, b) \neq 0.$$

On a

$$(3) \quad D = K + S(K, b).$$

Supposons que $K_1 \times j[:S(K, b)] = 0$ c. à d.

$$(4) \quad K_1 \times K \times S(K, b) = 0.$$

On a (3)

$$K_1 = K_1 \cdot K + K_1 \cdot S(K, b)$$

et c'est suivant (1), (2), (4) une décomposition du continu K_1 en deux ensembles fermés, non vides et disjoints, ce qui est impossible.

Corr. 1. Si C est un continu et

$$C \times S(K, b) \neq 0$$

$$C \times j[:S(K, b)] = 0,$$

alors

$$C \subset :S(K, b).$$

Corr. 2. Si C est un continu et

$$C \times \text{complémentaire de } :S(K, b) \neq 0$$

$$C \times j[:S(K, b)] = 0,$$

alors

$$C \subset \text{complémentaire de } S(K, b).$$

§ 21. $:S(a, b) = :S(K, b)$
équivalent à

$$:S(a, b) \times K = 0$$

$$a \in K.$$

Dém. C'est une conséquence de l'équivalence du § 4.

§ 22. Comme $j[:S(a, b)] = a$ (d'après § 16), on déduit du théorème précédent

$$:S(K, b) = :S(j[:S(K, b)], b).$$

§ 23. (1) $:S(K_1, b) = :S(K_2, b)$

équivalent à chacun de systèmes de relations

$$(2) \quad :S(K_2, b) \times K_1 = 0$$

$$S(K_2, b) \times K_1 \neq 0$$

$$(3) \quad \begin{aligned} & :S(K_2, b) \times K_1 = 0 \\ & j[:S(K_2, b)] \varepsilon K_1. \end{aligned}$$

Dém. (2) et (3) sont équivalents. car en soustrayant la première des relations (2) de la seconde, on obtient la deuxième relation (3).

(1) entraîne évidemment (3). Il reste à prouver l'inverse.

Suivant § 22 on peut transformer (3) en

$$\begin{aligned} & :S(j[:S(K_2, b)], b) \times K_1 = 0 \\ & j[:S(K_2, b)] \varepsilon K_1 \end{aligned}$$

et cela équivaut d'après le § 21 à

$$:S(j[:S(K_2, b)], b) = :S(K_1, b)$$

c. à. d. selon § 22.

$$:S(K_2, b) = :S(K_1, b) \quad \text{c. q. f. d.}$$

§ 24.

$$(1) \quad :S(K, c) = :S(a, c)$$

équivaut à chacun de systèmes de relations :

$$(2) \quad \begin{aligned} & :a, c \times K = 0 \\ & a \varepsilon K \end{aligned}$$

$$(3) \quad :K, c = :a, c.$$

D'après C § 4 (2) équivaut à (3). Suivant (1) on a

$$j[:S(K, c)] = j[:S(a, c)],$$

done (§ 16)

$$j(K, c) = a$$

et par suite (C § 4)

$$:K, c = :a, c.$$

L'inverse se démontre en reculant.

§ 25. Si K est un continu ne se réduisant pas à un seul point, alors $j[:S(K, b)]$ ne coïncide avec aucun point de saturation contenu dans K .

En effet dans le cas contraire en choisissant un point dans $K - j[:S(K, b)]$ et un point dans $:S(K, b)$

$$\begin{aligned} & k \varepsilon K - j[:S(K, b)] \\ & s \varepsilon :S(K, b) \end{aligned}$$

on pourrait démontrer que

$$(1) \quad |k.j[:S(K,b)]| + |j[:S(K,b),s]| = |k,s|$$

$$j[:S(K,b)] \varepsilon :k,s:.$$

(Pour le voir il suffit de rappeler que

$$j[] = K.S \quad \text{et} \quad K \times :S = 0).$$

La relation (1) est impossible pour tout point saturation (C § 12).

E.

Point principal d'un constituant restreint.

§ 1. Déf. On appelle point principal d'un constituant restreint $:S(K,b)$ dont la jonction est a tout point c , tel que

$$1^0) \quad c \varepsilon :S(K,b)$$

$$2^0) \quad \tau(|a,c|) \geq \frac{\tau(:S(K,b))}{3}$$

$$3^0) \quad \text{l'arc } |a,c| \text{ est saturé au point } c.$$

On appelle la paire de points (a,c) composée de la jonction du constituant et d'un de ces points principaux paire principale de $:S(K,b)$.

§ 2. A tout constituant restreint correspond au moins un point principal (ch. e.). Nous appellerons désormais principal ce point choisi.

Dém. Il suffit de traiter le cas où $K = 1$ point (D § 22). Il existe dans $S(a,b)$ qui est fermé deux points e et d , tels que

$$(1) \quad e + d \subset S(a,b)$$

$$(2) \quad \varrho(e,d) = \tau[:S(a,b)].$$

Au moins pour une des paires (a,e) ; (a,d) on a selon (1)

$$\varrho(a,e) \geq \frac{1}{2} \tau[:S(a,b)]$$

$$\varrho(a,d) \geq \frac{1}{2} \tau[:S(a,b)]$$

Supposons par exemple que la dernière égalité a lieu. On a $a \neq d$, donc (selons (2))

$$d \varepsilon S(a,b) - a = :S(a,b).$$

Il existe un arc $|a,c|$ contenant $|a,d|$ et saturé au point c (ch. e. C § 11).

Il est évident que $|a, c|$ satisfait aux conditions 2° et 3° du § précédent. Il reste à démontrer

$$(3) \quad c \varepsilon :S(a, b).$$

On a

$$|a, c| = |a, d| + |d, c|.$$

$|d, c|$ est donc un continu empiétant sur $:S(a, b)$ et ne contenant pas sa jonction a , par conséquent (D § 20 corr. 1)

$$|d, c| \subset :S(a, b)$$

d'où résulte (3).

§ 3. Si a, b est la paire principale de $:S(K, c)$, alors

$$:S(a, b) = :S(K, c).$$

C'est une conséquence de D §§ 22 et 5.

§ 4. Entre les constituants restreints et leurs paires principales, il existe une correspondance biunivoque.

En effet on a pour le point principal de $:S(a, b)$

$$:S(a, b) = :S(a, c).$$

Inversement à tout $:S(K, b)$ correspond une seule jonction et un seul point principal (§ 2).

§ 5. Si a, b est la paire principale de $:S(K, c)$, alors $:a, b|$ est une feuille restreinte poussant sur K c. à. d.

$$:a, b| = :K, b|.$$

Pour cette raison j'appelle $:a, b|$ feuille principale restreinte poussant sur K .

Un simple rapprochement des § 3 et D § 24 le fait voir.

§ 6. Deux feuilles principales restreintes poussant sur un continu K sont disjointes. En effet (selon § 3, § 4 et D §§ 6, 8) elles sont contenues dans deux constituants restreints différents, donc disjointes.

F.

Arc et constituant en T . — Arbre.

§ 1. Déf. On dit qu'un constituant $:S(a, b)$ est en T par rapport à $|c, d|$, si

$$S(a, b) \times :c, d : = a = j[:S(a, b)].$$

On voit facilement que dans ce cas

$$S(a, b) \times |c, d| = j[:S(a, b)].$$

§ 2. La condition nécessaire et suffisante que $|a, b|$ soit en T par rapport à a et $|c, d|$ est que $:S(a, b)$ soit en T par rapport à $|c, d|$ c. à. d. que

$$(1) \quad :c, d: \times S(a, b) = a.$$

Dém. La condition est suffisante. En effet de (1) et D § 8 on a

$$:c, d: \times |a, b| = a.$$

Elle est nécessaire. En effet $|a, b|$ étant en T par rapport à a et $|c, d|$ on a

$$(2) \quad \begin{aligned} |a, b| \times |c, d| &= a \\ a \varepsilon :c, d: \end{aligned}$$

De là d'après D § 24

$$:S(|c, d|, b) = :S(a, b),$$

donc en multipliant cette égalité par $|c, d|$ on obtient

$$:S(a, b) \times |c, d| = 0,$$

d'où selon (2)

$$S(a, b) \times :c, d: = a \quad \text{c. q. f. d.}$$

§ 3. Si $:S(a, b)$ est en T par rapport à $|K, c|$, alors

$$S(a, b) \times K = 0.$$

Dém. Il suffit de prouver (§ 20 D corr. 2) que

$$(1) \quad K \times j(:S(a, b)) = 0$$

$$(2) \quad K \times |\text{complémentaire de } :S(a, b)| \neq 0.$$

On a (§ 1, C § 4)

$$(3) \quad j[:S(a, b)] \subset :K, b:$$

$$(4) \quad :S(a, b) \times |K, b| = 0$$

$$(5) \quad |K, b| \times K \neq 0$$

$$(6) \quad :K, b: \times K = 0.$$

De (3) et (6) résulte (1), et de (4) et (5) résulte (2).

§ 4. Si $|a, b|$ (ou $:S(a, b)$) est en T par rapport à $|K, d|$, alors

$$\varphi(K, S(a, b)) > 0.$$

Dém. Il suffit de remarquer que $S(a, b)$ est fermé et que $S(a, b) \times K = 0$ (§ 3).

§ 5. Si $|a, b|$ [ou, ce qui est équivalent : $S(a, b)$] est en T par rapport à a et $|c, d|$, alors

$$S(a, b) \subset :S(c, d).$$

Dém. D'après D § 10, corr. 1, il suffit de prouver que

$$(1) \quad S(a, b) \times j[:S(c, d)] = 0$$

$$(2) \quad S(a, b) \times :S(c, d) \neq 0.$$

Selon le § 3 on a

$$c \times S(a, b) = 0$$

c. à. d. (1) est vrai.

Comme d'autre part

$$a \varepsilon :c, d \subset :S(c, d),$$

donc (2) est aussi vrai.

§ 6. Si $|a, b|$ (ou : $S(a, b)$) est en T par rapport à a et $|c, d|$, alors

$$:S(a, b) = :S(|c, d|, b).$$

Dém. D'après D § 24 il suffit de prouver que

$$:a, b| \times |c, d| = 0, \\ a \varepsilon |c, d|.$$

Or c'est vrai en raison de la définition de la relation T .

§ 7. Tout constituant d'un arc saturé $|c, d|$ est en T par rapport à sa jonction et $|c, d|$.

Dém. Suivant D § 2

$$:S(|c, d|, b) \times |c, d| = 0$$

et selon D § 15

$$a = j[:S(|c, d|, b)] \varepsilon |c, d|.$$

Or a ne peut coïncider ni avec c ni avec d , car ce sont des points de saturation. (D § 25).

§ 8. Déf. On appelle une classe d'arcs $(:a, b|)$ feuillage restreint poussant sur K („la tige“), si

1°) Tout $:a, b|$ est une feuille restreinte poussant sur K , c. à. d.

$$:a, b| = :K, b|. \quad (\text{comp. C § 4}).$$

2°) Les arcs $:a, b|$ sont disjoints deux à deux.

$K + (:a, b|)$ s'appelle arbre,
 $(|a, b|)$ feuillage complet.

§ 9. Tout feuillage poussant sur un continu K forme une classe faible et par suite effectivement énumérable.

C'est une conséquence de C § 8.

§ 10. Tout arbre est un continu et

$$K + (:a, b|) = K + (|a, b|)$$

(selon D § 19).

§ 11. Pour deux feuilles différentes $:a, b|$, $:a', b'|$ d'un feuillage $(:a, b|)$ poussant sur K on a

- (1) $:S(a, b) = :S(K, b)$
- (2) $:S(a', b') = :S(K, b')$
- (3) $:S(a, b) \times :S(a', b') = 0$.

Dém. (1) et (2) sont des conséquences de D § 24. En multipliant (1) par (2) on obtient (3) selon D § 17.

§ 12. Si sur chaque feuille complète $|a_1, b_1|$ d'un feuillage $(|a_1, b_1|)$ poussant sur K on construit un feuillage $(:a_2, b_2|)$ dont toute feuille est en T par rapport à a_2 et $|a_1, b_1|$, alors la somme de tous ces nouveaux feuillages

$$((:a_2, b_2|))$$

constitue un feuillage poussant sur l'arbre

$$K + (:a_1, b_1|)$$

considéré comme une nouvelle tige.

Nous appelons $((:a_2, b_2|))$ feuillage en T poussant sur le feuillage

$$(|a_1, b_1|) \quad \text{ou} \quad (:a_1, b_1|).$$

Dém. α) Deux feuilles $:a_2, b_2|$ différentes sont sans points communs.

En effet c'est vrai pour deux feuilles $:a_2, b_2|$ poussant sur un même $|a_1, b_1|$.

Dans le cas où $:a_2, b_2|$ pousse sur $:a_1, b_1|$ et $:a'_2, b'_2|$ pousse sur $:a'_1, b'_1|$, tels que

$$:a_1, b_1| \neq :a'_1, b'_1|,$$

on a (§ 3, § 5)

$$(1) \quad |a_2, b_2| \subset :S(a_1, b_1)$$

$$(2) \quad |a'_2, b'_2| \subset :S(a'_1, b'_1).$$

En multipliant (1) par (2), on a selon le § 11

$$|a_2, b_2| \times |a'_2, b'_2| = 0.$$

α est donc démontré.

Remarquons qu'en multipliant les deux termes de (1) par $:S(a'_1, b'_1)$, on obtient d'après le § 11

$$|a_2, b_2| \times :S(a'_1, b'_1) = 0,$$

d'où

$$|a_2, b_2| \times :a'_1, b'_1| = 0,$$

ce qui exprime que toute feuille nouvelle n'empiète que sur une seule feuille de $(:a_1, b_1|)$ voire sur celle sur laquelle elle pousse.

D'autre part en multipliant (1) par K on a pour toute feuille $|a_2, b_2|$ (selon le § 11)

$$(3) \quad |a_2, b_2| \times K = 0.$$

Ces deux remarques entraînent

$$[K + (:a_1, b_1|)] \times |a_2, b_2| = a_2$$

et cette égalité exprime que $|a_2, b_2|$ est une feuille restreinte poussant sur $K + (:a_1, b_1|)$. Deux feuilles différentes $|a_2, b_2|$ $|a'_2, b'_2|$ étant disjointes (d'après α) $((:a_2, b_2|))$ constitue un feuillage poussant sur $K + (:a_1, b_1|)$ (§ 8).

§ 13. Aucune feuille $|a_2, b_2|$ d'un feuillage $((|a_2, b_2|))$ qui est en T par rapport à un feuillage $(|a_1, b_1|)$ sur K n'empiète sur K .

Plus généralement

$$K \times S(a_2, b_2) = 0.$$

C'est une conséquence du § 3.

§ 14. Si $((:a_2, b_2|))$ est un feuillage en T par rapport à un feuillage $(|a_1, b_1|)$ poussant sur un continu K , alors

$$((S(a_2, b_2))) \times K = 0.$$

C'est le théorème précédent autrement exprimé.

§ 15. $K_1 = K + (:a_1, b_1|)$ étant un arbre et $((:a_2, b_2|))$ un feuillage en T par rapport au feuillage $(|a_1, b_1|)$ et, si pour un

$S(a_2, b_2)$ variable

$$\lim \varrho(K, S(a_2, b_2)) = 0,$$

on aura aussi

$$\lim \tau[S(a_2, b_2)] = 0.$$

Dém. La classe $((S(a_2, b_2)))$ est faible comme identique à la classe $((S(K_1, b_2)))$ (d'après § 12, § 11 et D § 18).

Il existe donc au plus un nombre fini de $S(a_2, b_2)$ pour lesquels

$$\tau(S(a_2, b_2)) > \varepsilon.$$

Soient $S_1, S_2, \dots, S_n$ ces constituants.

D'après le § 14

$$(S_1 + S_2 + \dots + S_n) \times K = 0,$$

done, comme $S_1 + \dots + S_n$ est fermé,

$$\varrho(K, S_1 + \dots + S_n) > 0.$$

Posons

$$\delta = \frac{\varrho(S_1 + \dots + S_n, K)}{2}.$$

Il est évident d'après ce qui précède que

$$\varrho(K, S(a_2, b_2)) \leq \delta$$

entraîne

$$\tau(S(a_2, b_2)) \leq \varepsilon \quad \text{c. q. f. d.}$$

§ 16. Du théorème précédent résultent les corrolaires suivants:
 $K_1 = K + (|a_1, b_1|)$ étant un arbre et $((:a_2, b_2|))$ un feuillage en T par rapport au feuillage $(|a_1, b_1|)$, alors pour un $:S(a_2, b_2)$ variable, on a

$$1^0) \quad \lim [j(:S(a_2, b_2)), K] = 0$$

équivalent à

$$\lim \varrho[S(a_2, b_2), K] = 0.$$

2°) Si $k \in K$, alors

$$\lim \varrho[j(:S(a_2, b_2), k) = 0$$

équivalent à

$$\lim \varrho[S(a_2, b_2), k] = 0.$$

3°) Si $k \in K$ et que

$$\lim \varrho[j(:S(a_2, b_2), k) = 0,$$

alors tout point de $S(a_2, b_2)$ tend vers k .

Ce théorème s'exprime d'une façon stricte comme il suit :

A tout point k de K et à tout $\varepsilon > 0$, correspond un $\delta > 0$, tel que, si

$$\varrho[j(:S(a_2, b_2)), k] \leq \delta,$$

alors pour tout point c' de $S(a_2, b_2)$ on a

$$\varrho(c', k) \leq \varepsilon.$$

§ 17. La somme

$$K + (:a_1, b_1|)$$

où $()$ est étendu sur toutes les feuilles principales restreintes poussant sur K (E § 5) est un arbre. Nous l'appelons **arbre principal** à la tige K et nous appelons $(:a_1, b_1|)$ feuillage principal restreint poussant sur K .

Dém. D'après E § 4 $(:a_1, b_1|)$ est une classe de feuilles restreintes poussant sur K . Cette classe forme un feuillage restreint (§ 8), car suivant (E § 6 et E § 4) les feuilles différentes sont disjointes.

§ 18. Si $(:a_1, b_1|)$ est le feuillage principal poussant sur K_1 , alors

$$D = K_1 + (:S(a_1, b_1))$$

où D désigne la dendrite en question. (E § 5, F § 17, D § 5, D § 9).

§ 19. Si K_2 est l'arbre principal à la tige K_1 , alors

α) aucun constituant complet par rapport à K_2 n'empiète sur K_1 ,

β) aucune feuille complète poussant sur K_2 n'empiète sur K_1 .

Dém. Supposons qu'il existe un $:S(K_2, b)$, tel que

$$(1) \quad S(K_2, b) \times K_1 \neq 0.$$

On a

$$(2) \quad K_1 \times :S(K_2, b) \subset K_2 \times :S(K_2, b).$$

De (1) et (2) on obtient d'après D § 23

$$:S(K_1, b) = :S(K_2, b|),$$

d'où en multipliant par K_2 on a

$$:S(K_1, b) \times K_2 = 0,$$

ce qui est impossible, car K_2 contient le point principal de $:S(K_1, b)$.

§ 20. Si K_2 est l'arbre principal à la tige K_1 ne se réduisant pas à un seul point, alors

α) tout constituant restreint par rapport à K_2 est en T par rapport à une seule feuille principale poussant sur K_1 ,

β) il en est de même de toute feuille restreinte poussant sur K_2 .

Dém. β résulte de α d'après le § 2.

On a

$$K_2 = K_1 + (|a_2, b_2|)$$

où $()$ est la somme étendue sur toutes les feuilles principales poussant sur K_1 .

Soit : $S(K_2, b_3)$ un constituant quelconque par rapport à K_2 .

On a d'après § 19

$$S(K_2, b_3) \times K_1 = 0,$$

donc à fortiori

$$S(K_2, b_3) \times (a_2) = 0.$$

D'autre part comme les b_2 sont des points de saturation, on a (D § 25)

$$S(K_2, b_3) \times (b_2) = 0.$$

Les quatre égalités précédentes donnent

$$K_2 \times S(K_2, b_3) = (:a_2, b_2:) \times S(K_2, b_3)$$

c. à d.

$$j[:S(K_2, b_3)] = (:a_2, b_2:) \times S(K_2, b_3).$$

Les feuilles restreintes différentes $:a_2, b_2|$ étant disjointes deux à deux, il en existe une seule soit $:a'_2, b'_2|$, telle que

$$j[:S(K_2, b_3)] = :a'_2, b'_2: \times S(K_2, b_3)$$

et cela exprime (§ 1) que $:S(K_2, b_3)$ est en T par rapport à $|a'_2, b'_2|$. Le théorème est donc démontré.

§ 21. Si K_2 est l'arbre principal à la tige K_1 et, si $(|a_1, b_1|)$ est son feuillage, alors la classe de toutes les feuilles principales $(:a_2, b_2|)$ poussant sur K_2 constitue un feuillage en T par rapport au feuillage $(|a_1, b_1|)$ et par conséquent les théorèmes des §§ 8—16 sont vrais pour lui.

Dém. C'est une conséquence immédiate des § 17, § 20 et § 12.

G.

Suite $\mathcal{Q}$.

§ 1. Déf. Nous appelons suite $\mathcal{Q}$ déduite d'une dendrite D toute suite finie ou non

$$(a_1, b_1), (a_2, b_2), \dots$$

jouissant des propriétés suivantes :

$\alpha)$ (a_n, b_n) est une classe de paires de points (paire = suite composée de deux éléments).

Nous appelons n rang de (a_n, b_n) ¹⁾.

$\beta)$ (a_1, b_1) contient une seule paire et $|a_1, b_1|$ est un arc saturé.

Tout $b_n (n \geq 2)$ est le point principal de $:S(a_n, b_n)$.

$\gamma)$ $|a_n, b_n|$ est en T au moins par rapport à un des $|a_{n-1}, b_{n-1}|$.

On dit dans ce cas que la paire (a_{n-1}, b_{n-1}) précède immédiatement (a_n, b_n) .

$\delta)$ Pour deux paires différentes du même rang on a

$$|a_n, b_n| \times |a'_n, b'_n| = 0.$$

§ 2.

$$(a_n) \subset \sum_{\nu=1}^{n-1} (|a_\nu, b_\nu|)$$

C'est une conséquence du § 1.

Remarque. Cette relation est fausse pour $n = 1$. Nous supprimerons dans la suite, à quelques exceptions près, les restrictions relatives au cas de $n = 1$. — aucune difficulté n'en résultera. Nous supprimerons aussi les démonstrations concernant ce cas comme étant particulièrement simples.

§ 3. La classe de tous les $|a_n, b_n|$ de rang n est faible et effectivement énumérable, car $(|a_n, b_n|)$ est une classe d'arcs positifs et disjoints (§ 1, C § 8).

De là :

§ 4. La classe (a_n) est effectivement énumérable.

¹⁾ Nous désignerons dans la suite par (a_n, b_n) , (a'_n, b'_n) , $(a_n^{(\mu)}, b_n^{(\mu)})$ seulement les paires de (a_n, b_n) .

Si $\varphi(a_n, b_n)$ est un ensemble de points, nous désignerons par

$$(\varphi(a_n, b_n))$$

suivant le besoin la classe ou la somme de ces ensembles correspondant à la classe (a_n, b_n) .

§ 5. $(:a_n, b_n|)$ constitue un feuillage poussant sur $\sum_{\nu/1}^{n-1}(|a_\nu, b_\nu|)$ ($n \geq 2$) et $(:a_\nu, b_\nu|)$ forme un feuillage en T par rapport au feuillage $(|a_{\nu-1}, b_{\nu-1}|)$ ($n \geq 3$)

Dém. On établit ce théorème de proche en proche en se basant sur le § 1 et F § 12.

§ 6. $:S(a_n, b_n)$ est un constituant par rapport à $\sum_{\nu/1}^{n-1}(|a_\nu, b_\nu|)$ c. à. d.

$$:S(a_n, b_n) = :S(\sum_{\nu/1}^{n-1}(|a_\nu, b_\nu|), b_n).$$

Dém. D'après le § 5 $:a_n, b_n|$ est une feuille poussant sur $\sum_{\nu/1}^{n-1}$, donc l'égalité précédente a lieu en raison de F § 11.

§ 7. $\sum_{\nu/1}^n(|a_\nu, b_\nu|) \times (S(a_{n+2}, b_{n+2})) = 0$.

C'est une application de F § 13.

§ 8. $\sum_{\nu/1}^n(|a_\nu, b_\nu|) \times (S(a_{n+k}, b_{n+k})) = 0 \quad (k \geq 2),$

car d'après le § 7

$$\sum_{\nu/1}^{n+k-2}(|a_\nu, b_\nu|) \times (S(a_{n+k}, b_{n+k})) = 0.$$

§ 9. Pour $(a_n, b_n) \neq (a'_n, b'_n)$ on a

$$:S(a_n, b_n) \times :S(a'_n, b'_n) = 0.$$

En effet les feuilles $:a_n, b_n|$ et $:a'_n, b'_n|$ étant disjointes les deux $:S$ le sont aussi (F § 11).

§ 10. D'après le § 6

$$m \neq n$$

entraîne

$$(a_m, b_m) \neq (a_n, b_n).$$

§ 11. Si,
alors on a

$$(a_m, b_m) \neq (a_n, b_n),$$

$$\varphi(a_m, b_m) \neq \varphi(a_n, b_n)$$

pour chacun des cas suivants

$$\begin{array}{ll} \alpha) & \varphi(a_k, b_k) = |a_k, b_k| \\ \beta) & \quad \quad \quad = :a_k, b_k| \\ \gamma) & \quad \quad \quad = :a_k, b_k: \\ \delta) & \quad \quad \quad = |a_k, b_k| \end{array}$$

$$\begin{aligned}\varepsilon) \quad \varphi(a_k, b_k) &= :S(a_k, b_k) \\ \zeta) \quad \quad \quad &= S(a_k, b_k) \\ \eta) \quad \quad \quad &= b_k.\end{aligned}$$

Dém. Les cas α - δ sont évidents. Les cas ε , ζ , η résultent sans peine des §§ 6 et 9.

Corrolaire. La correspondance dans laquelle à la paire (a_ν, b_ν) est associé $\varphi(a_\nu, b_\nu)$ est biunivoque.

§ 12.

$$(1) \quad a_m \times :S(a_{m+k}, b_{m+k}) = 0 \quad (k \geq 1).$$

Dém. Pour $k \geq 2$ c'est vrai selon le § 8.

Suivant les §§ 2 et 5 on a

$$\begin{aligned}a_m \varepsilon \sum_{\nu/1}^{m-1} (|a_\nu, b_\nu|) \\ \sum_{\nu/1}^{m-1} (|a_\nu, b_\nu|) \times :S(a_{m+1}, b_{m+1}) = 0,\end{aligned}$$

donc (1) est aussi vrai pour $k = 1$.

§ 13. Pour a_m et (a'_m, b'_m) quelconques on a

$$a_m \times :S(a'_m, b'_m) = 0.$$

Dém. Il suffit de rappeler que $a_m \varepsilon \sum_{\nu/1}^{m-1} |a_\nu, b_\nu|$ (§ 2) et que $:S(a'_m, b'_m)$ est un constituant par rapport à $\sum_{\nu/1}^{m-1} |a_\nu, b_\nu|$ (§ 6).

§ 14. Si $(a_m, b_m) \neq (a'_m, b'_m)$,
alors

$$S(a_m, b_m) \times :S(a'_m, b'_m) = 0.$$

Dém. Il suffit de prouver que

$$\begin{aligned}(1) \quad a_m \times :S(a'_m, b'_m) &= 0 \\ (2) \quad :S(a_m, b_m) \times :S(a'_m, b'_m) &= 0.\end{aligned}$$

Or (1) est vrai selon le § 13 et (2) selon le § 9.

§ 15.

$$(\alpha) \quad :a_n, b_n| \times :a_{n+k}, b_{n+k}| = 0. \quad (k \geq 1).$$

Dém. On a (D § 8)

$$\begin{aligned}(1) \quad :a_{n+k}, b_{n+k}| &\subset :S(a_{n+k}, b_{n+k}) \\ (2) \quad :a_n, b_n| &\subset \sum_{\nu/1}^{n+k-1} (|a_\nu, b_\nu|).\end{aligned}$$

$:S(a_{n+k}, b_{n+k})$ étant un constituant par rapport à $\sum_{\nu/1}^{n+k-1}(|a_\nu, b_\nu|)$ (§ 6),

on a $\sum_{\nu/1}^{n+k-1}(|a_\nu, b_\nu|) \times :S(a_{n+k}, b_{n+k}) = 0$,

donc en multipliant (1) par (2) on obtient l'égalité (α).

§ 16. Tout point de

$$\sum_{\nu/1}^{\infty}(|a_\nu, b_\nu|) - \sum_{\nu/1}^{\infty}(a_\nu)$$

appartient à un seul $|a_n, b_n|$.

$\sum_{\nu/1}^{\infty}$ désigne la somme correspondant à tous les éléments de la suite $\mathcal{Q}$ même, si elle est finie.

Dém. C'est une conséquence du § 15 et du § 1 d.

§ 17. Si $(a_n, b_n) \neq (a'_n, b'_n)$, alors

$$(I) \quad S(a_n, b_n) \times S(a'_n, b'_n) \subset a_n \times a'_n.$$

Dém. On a d'après le § 14.

$$\begin{aligned} S(a_n, b_n) \times S(a'_n, b'_n) &= [a_n + :S(a_n, b_n) \times S(a'_n, b'_n) = \\ &= a_n \cdot S(a'_n, b'_n) \subset a_n \end{aligned}$$

et analogiquement

$$S(a_n, b_n) \times S(a'_n, b'_n) \subset a'_n$$

d'où résulte (1).

§ 18. La classe de tous les $|a_n, b_n|$ de tous les rangs est faible et effectivement énumérable, car il en est ainsi de la classe de tous les $:a_n, b_n|$ (§ 1 d, § 15, C § 8).

§ 19. La classe de tous les $:S(a_n, b_n)$ de tous les rangs est faible et effectivement énumérable.

Dém. b_n est point principal de $:S(a_n, b_n)$ (§ 1), on a donc (E § 1)

$$\tau(:S(a_n, b_n)) \leq 3 \tau(|a_n, b_n|)$$

et ainsi on est ramené au théorème précédent.

§ 20. Déf. Nous appelons rang d'une fonction $\chi(a_n, b_n)$ le rang de la paire (a_n, b_n) c. à. d. n . (comp. § 1).

§ 21. Chacune des fonctions $\chi(a_n, b_n)$ suivantes

$$\begin{aligned} (a_n, b_n), \quad :a_n, b_n|, \quad :a_n \cdot b_n:, \quad |a_n, b_n|:, \quad |a_n, b_n|, \\ :S(a_n, b_n), \quad S(a_n, b_n), \quad b_n, a_n \end{aligned}$$

a un seul rang et notamment n .

Dém. Pour la dernière fonction c'est vrai en raison du § 12. Pour les autres en raison de la correspondance biunivoque entre (a_n, b_n) et $\chi(a_n, b_n)$ (§ 11).

§ 22. La relation :

(1) (a'_n, b'_n) précède immédiatement (a'_{n+1}, b'_{n+1})
(comp. § 1) équivaut à chacune des relations suivantes

$$(2) \quad a'_{n+1} \varepsilon : a'_n, b'_n :$$

$$(3) \quad a'_{n+1} \varepsilon [a'_n, b'_n].$$

Dém. (1) entraîne (2) et (3) suivant le § 1. Pour prouver l'inverse nous établirons trois lemmes:

$$\text{Lemme 1.} \quad \sum_{\nu/1}^n (a_\nu) \times a'_{n+1} = 0. \quad (\S 21)$$

$$\text{Lemme 2.} \quad \sum_{\nu/1}^\infty (b_n) \times a'_{n+1} = 0. \quad (\S 15 \text{ et } \S 1)$$

Lemme 3. $[a'_{n+1}, b'_{n+1}]$ empiète sur un seul $[a_n, b_n]$ et il est en T par rapport à ce $[a_n, b_n]$ et a'_{n+1} .

Dém. On a (§ 5, § 8)

$$\begin{aligned} \sum_{\nu/1}^n |(a_\nu, b_\nu)| \times |a'_{n+1}, b'_{n+1}| &= a'_{n+1} \\ \sum_{\nu/1}^{n-1} |(a_\nu, b_\nu)| \times |a'_{n+1}, b'_{n+1}| &= 0, \end{aligned}$$

$$\text{donc} \quad |(a_n, b_n)| \times |a'_{n+1}, b'_{n+1}| = a'_{n+1}$$

et selon les lemmes précédents

$$(:a'_n, b'_n:) \times [a'_{n+1}, b'_{n+1}] = a'_{n+1}.$$

Comme deux $:a'_n, b'_n:$ sont disjoint (§ 1), donc il existe un seul $:a_n, b_n:$, tel que

$$:a_n, b_n: \times [a'_{n+1}, b'_{n+1}] = a'_{n+1} \quad \text{c. q. f. d.}$$

Ce lemme montre que (3) entraîne (1) et que (2) entraîne (1).

§ 23. On dit que $\psi(a_n, b_n)$ précède immédiatement $\chi(a_m, b_m)$, alors et seulement alors, si (a_n, b_n) précède immédiatement (a_m, b_m) .

§ 24. Si ψ et χ sont chacune une des fonctions $(a_n, b_n); :a_n, b_n|; :a_n, b_n:: |a_n, b_n:: |a_n, b_n|; :S(a_n, b_n); S(a_n, b_n); b_n, a_n;$ alors à $\chi(a_m, b_m)$ ($m \geq 2$) correspond un seul $\psi(a_n, b_n)$ qui le précède immédiatement.

Dém. C'est vrai, si

$$\chi(a_k, b_k) = \psi(a_k, b_k) = (a_k, b_k) \quad (\S 22, \text{lemme 3}).$$

La généralisation est immédiate (§ 22, § 11 corr.).

Remarque: Le théorème devient faux, si l'on remplace le mot „précède“ par le mot „suit“.

§ 25. Si $\psi(a_n, b_n)$ précède immédiatement $\chi(a_m, b_m)$, alors $m - n = 1$.

C'est vrai pour $\psi = \chi = (a_k, b_k)$ (§ 1).

La généralisation est vraie, car ψ (et χ) a un seul rang.

§ 26. Gardons les notations du § 24.

On dit que

$$\psi(a_n, b_n) < \varphi(a_m, b_m),$$

s'il existe une suite

$$(a_n, b_n), (a_{\alpha_{n+1}}, b_{\alpha_{n+1}}), \dots, (a_{\alpha_{m-1}}, b_{\alpha_{m-1}}), (a_m, b_m),$$

telle que de deux éléments voisins le premier précède immédiatement (au sens du § 1) le suivant

§ 27. La relation $<$ est transitive.

§ 28. Si $\psi(a_n, b_n) < \varphi(a_m, b_m)$, alors

$$m > n$$

(notations du § 24). Cela résulte du § 25.

§ 29. Si $n < m$, alors dans la classe $(\psi(a_n, b_n))$ (notations du § 24) il existe un seul $\psi(a_n, b_n)$ qui précède $\varphi(a_m, b_m)$.

C'est une généralisation du § 24.

§ 30. Si $\chi(a_k, b_k)$ et $\psi(a_k, b_k)$ sont deux fonctions ainsi désignées dans le § 24, mais

$$\chi \neq a_k, \quad \psi \neq a_k,$$

alors

$$(a_m, b_m) < (a_n, b_n)$$

équivalent à

$$\chi(a_m, b_m) < \psi(a_n, b_n),$$

car entre (a_k, b_k) et $\chi(a_k, b_k)$ (et $\psi(a_k, b_k)$) il existe la correspondance biunivoque mentionnée dans le § 11.

§ 31. (Notations du § 24). Si

$$(1) \quad \varphi(a_m, b_m) < a'_n$$

et que $(a'_n, b'_n) \varepsilon ((a_n, b_n))$, alors

$$(2) \quad \varphi(a_m, b_m) < \psi(a'_n, b'_n).$$

Dém. On a (§ 28)

$$m < n.$$

Il suffit (Déf. § 23, § 27) de démontrer le théorème dans le cas

$$n - m = 1.$$

Dans ce cas il existe selon (1) un b''_n , tel que

$$(a_m, b_m) \text{ précède imméd. } (a'_n, b''_n),$$

donc (§ 22)

$$a'_n \varepsilon |a_m, b_m|$$

et par conséquent (§ 22) (a_m, b_m) précède immédiatement (a'_n, b'_n) , d'où résulte (2) selon la définition du § 23.

§ 32. $:S(a_n, b_n)$ contient tous les $S(a_m, b_m)$ qui le suivent.

Dém. C'est vrai pour le cas, où le rang de $S(a_m, b_m)$ est supérieur de l'unité à celui de $:S(a_n, b_n)$ (§ 1 γ , F § 5, F § 2). On généralise facilement, car les relations $<$ et $\subset$ sont transitives.

$$\S 33. \quad \sum_{\nu/n+1}^{\infty} (S(a_\nu, b_\nu)) \subset (:S(a_n, b_n)),$$

car tout $S(a_{n+\nu}, b_{n+\nu})$ a un seul précédent de rang n .

§ 34. Si (a_{n-1}, b_{n-1}) précède (a_n, b_n) , alors

$$\begin{aligned} |a_{n-1}, b_{n-1}| \times S(a_n, b_n) &= a_n \\ |a_{n-1}, b_{n-1}| \times |a_n, b_n| &= a_n \\ :a_{n-1}, b_{n-1}: \times S(a_n, b_n) &= a_n \\ :a_{n-1}, b_{n-1}: \times |a_n, b_n| &= a_n. \end{aligned}$$

Dém. Suivant § 1 γ et F § 2 $|a_n, b_n|$ et $:S(a_n, b_n)$ sont un arc et un constituant en T par rapport à a_n et $|a_{n-1}, b_{n-1}|$. De là résultent facilement les quatre relations précédentes.

§ 35. Tout $S(a_m, b_m)$ et tout $S(a_n, b_n)$ qui ne le précède pas sont disjoints ($n > m$).

Autrement, si

$$\begin{aligned} (1) \quad & S(a_m, b_m) \text{ non } < S(a_{m+k}, b_{m+k}), \quad (k \geq 1) \text{ on a} \\ (2) \quad & S(a_m, b_m) \times S(a_{m+k}, b_{m+k}) = 0. \end{aligned}$$

Dém. Il existe (§ 29) un seul $:S(a'_m, b'_m)$, tel que

$$(3) \quad :S(a'_m, b'_m) < S(a_{m+k}, b_{m+k})$$

et de là (§ 32)

$$(4) \quad S(a_{m+k}, b_{m+k}) \subset :S(a'_m, b'_m).$$

Comme selon (1) et (3) $(a_m, b_m) \neq (a'_m, b'_m)$, donc (§ 14)

$$S(a_m, b_m) \times :S(a'_m, b'_m) = 0$$

et par conséquent en multipliant les deux membres de (4) par $S(a_m, b_m)$, on obtient (2).

§ 36. Parmi les $[a_m, b_m]$ de rang différent de n il existe un seul $[a'_k, b'_k]$ (et notamment son précédent) qui renferme a_n . On a de plus

$$a_n \varepsilon :a'_k, b'_k:.$$

Dém. a_n n'appartient ni à aucun $[a_m, b_m]$ de rang supérieur à n (§ 12), ni à aucun $[a_m, b_m]$, tel que $n - m \geq 2$ (§ 8), ni à aucun $[a_{n-1}, b_{n-1}]$ qui ne le précède (§ 35). a_n appartient donc seulement à $[a_m, b_m]$ qui le précède immédiatement et le théorème est vrai d'après le § 22.

§ 37.

$$(1) \quad \overline{\sum_{\nu/1}^{\infty} (a_{\nu}, b_{\nu})} \subset \sum_{\nu/1}^n ([a_{\nu}, b_{\nu}] + (:S(a_{n+1}, b_{n+1}))).$$

Dém. D'après le § 2

$$\sum_{\nu/1}^n ([a_{\nu}, b_{\nu}]) + (a_{n+1}) = \sum_{\nu/1}^n ([a_{\nu}, b_{\nu}]).$$

D'autre part

$$(:a_{n+1}, b_{n+1}) \subset (:S(a_{n+1}, b_{n+1})),$$

et selon le § 33

$$\sum_{\nu/n+2}^{\infty} ([a_{\nu}, b_{\nu}]) \subset (:S(a_{n+1}, b_{n+1})).$$

En ajoutant ces trois relations et en remarquant que le second membre du résultant est fermé (D § 11), on obtient (1).

§ 38. Pour tout point x appartenant à

$$\overline{\sum_{\nu/1}^{\infty} ([a_{\nu}, b_{\nu}])} - \sum_{\nu/1}^{\infty} ([a_{\nu}, b_{\nu}])$$

et pour tout $n > 1$ il existe un seul $:S(a_n, b_n)$, tel que

$$x \varepsilon :S(a_n, b_n).$$

Dém. Selon le § 37 il existe au moins un $:S(a_n, b_n)$ de la sorte. Il existe au plus un $:S$ en question, car deux $:S$ de même rang sont disjoints (§ 14).

§ 39. Si

$$(1) \quad x \in \overline{\sum_{\nu/1}^{\infty}(|a_{\nu}, b_{\nu}|)} - \sum_{\nu/1}^{\infty}(|a_{\nu}, b_{\nu}|),$$

alors il n'existe aucun arc $|c, d|$, tel que

$$x \in c, d \subset \overline{\sum_{\nu/1}^{\infty}(|a_{\nu}, b_{\nu}|)}$$

Dém. Dans le cas contraire on aurait

$$(2) \quad \begin{aligned} |a, c| &= |a, x| + |x, c| \\ |a, x| \times |x, c| &= x \end{aligned}$$

$\varrho(a, x)$ et $\varrho(x, c)$ étant positifs, il existe un δ , tel que

$$(3) \quad \begin{aligned} 0 &< \delta < \varrho(a, x) \\ 0 &< \delta < \varrho(b, x) \end{aligned}$$

D'autre part d'après le § 19 il existe un n , tel que pour tout $:S(a_n, b_n)$ de rang n

$$(4) \quad \tau(:S(a_n, b_n)) \leq \delta.$$

x appartient à un seul: $S(a_n, b_n)$ (§ 38), par exemple

$$(5) \quad x \in :S(a'_n, b'_n).$$

$|a, x|$ et $|x, c|$ empiètent donc sur $:S(a'_n, b'_n)$.

Ils empiètent aussi sur son complémentaire (d'après (3) et (4)). Ils passent donc (D § 20) par $j(:(a', b'))$ et cela est impossible, car cette jonction est différente de x (d'après (5)) et x est selon (2) le seul point commun à $|a, x|$ et $|x, c|$.

Corrolaire 1. $\overline{\sum_{\nu/1}^{\infty}(|a_{\nu}, b_{\nu}|)} - \sum_{\nu/1}^{\infty}(|a_{\nu}, b_{\nu}|)$ ne renferme aucun point de branchement (C § 19).

Corrolaire 2. Tout point de cet ensemble est un point de saturation (C § 12).

§ 40 Si

$$(1) \quad |c, d| \subset \overline{\sum_{\nu/1}^{\infty}(|a_{\nu}, b_{\nu}|)}$$

et que $|c, d|$ est en T par rapport à c et $|a'_n, b'_n|$, alors

$$c \in (a_{n+1}).$$

Dém. $:c, d|$ constitue une feuille en T par rapport à une feuille $|a'_n, b'_n|$ poussant sur $\sum_{\nu/1}^{n-1}(|a_{\nu}, b_{\nu}|)$ et appartenant au feuillage

$(:a_n, b_n|)$ poussant sur la même tige (§ 5), donc (F § 12) $:c, d|$ est une feuille restreinte poussant sur $\sum_{\nu/1}^n (|a_\nu, b_\nu|)$ c. à d.

$$(2) \quad :c, d| \times \sum_{\nu/1}^n (|a_\nu, b_\nu|) = 0$$

$$(3) \quad c \in \sum_{\nu/1}^n (|a_\nu, b_\nu|).$$

Suivant (1), (2) § 33 et § 12, il existe un seul constituant $:S(a'_{n+1}, b'_{n+1})$, tel que

$$d \in :S(a'_{n+1}, b'_{n+1}).$$

De (2) et (3) on conclut facilement que

$$\begin{aligned} :c, d| &\subset :S(a'_{n+1}, b'_{n+1}) \\ c &= a'_{n+1} \quad \text{c. q. f. d.} \end{aligned}$$

Suite régulière.

§ 41. Une suite $\chi(a_1, b_1), \chi(a_2, b_2), \dots$ (notations du § 24) est appelée régulière, alors et seulement alors, si

1° elle est infinie

2° $\chi(a_\nu, b_\nu)$ précède immédiatement $\chi(a_{\nu+1}, b_{\nu+1})$ ($\nu/1, \dots, \infty$).

§ 42 Toute suite infinie régulière „converge“ vers un seul

point x de $\overline{\sum_{\nu/1}^{\infty} |a_\nu, b_\nu|} = \sum_{\nu/1}^{\infty} |a_\nu, b_\nu|$ c. à d. on a

$$x \in \prod_{\nu/1}^{\infty} [:S(a_\nu, b_\nu)] = \prod_{\nu/1}^{\infty} S(a_\nu, b_\nu).$$

Dém. On a (§ 41, § 32)

$$:S(a_2, b_2) < :S(a_3, b_3) < \dots$$

$$(1) \quad S(a_{\nu+1}, b_{\nu+1}) \subset :S(a_\nu, b_\nu) \subset S(a_\nu, b_\nu),$$

done

$$(2) \quad \prod_{\nu/2}^{\infty} S(a_\nu, b_\nu) \subset \prod_{\nu/1}^{\infty} :S(a_\nu, b_\nu) \subset \prod_{\nu/1}^{\infty} S(a_\nu, b_\nu)$$

d'où résulte

$$\prod_{\nu/1}^{\infty} S(a_\nu, b_\nu) = \prod_{\nu/1}^{\infty} :S(a_\nu, b_\nu),$$

car suivant (1) et A § 21

$$\prod_{\nu/1}^{\infty} S(a_\nu, b_\nu) \neq 0$$

et ce produit se réduit à un seul point, car le diamètre de $S(a_\nu, b_\nu)$ tend vers 0 (§ 19).

§ 43. A tout point x de $\overline{\sum_{\nu/1}^{\infty}(|a_\nu, b_\nu|)} = \sum_{\nu/1}^{\infty}(|a_\nu, b_\nu|)$, il appartient une seule suite régulière convergeant vers x .

Dém. D'après le § 38 x est contenu dans un seul constituant restreint de rang n soit dans : $S(a_n, b_n)$

$$x \in S(a_n, b_n) \quad (n | 2, 3, \dots)$$

Cette égalité entraîne en raison du § 35

$$: S(a_n, b_n) < : S(a_{n+1}, b_{n+1}).$$

Il est facile à voir que

$$: S(a_1, b_1), : S(a_2, b_2), \dots$$

est l'unique suite correspondant à la question.

§ 44. En rapprochant les §§ 42 et 43, on obtient:

Entre les points de $\overline{\sum_{\nu/1}^{\infty}(|a_\nu, b_\nu|)} = \sum_{\nu/1}^{\infty}(|a_\nu, b_\nu|)$ et les suites régulières tendant vers ces points, il existe une correspondance biunivoque.

§ 45.

$$(1) \quad : S(a_n, b_n) \times \overline{\sum_{\nu/1}^{\infty}(|a_\nu, b_\nu|)}$$

est composé

1°) de : a_n, b_n ,

2°) de tous les $|a_n, b_n|$, qui le suivent,

3°) des points vers lesquels convergent les suites régulières

($: S(a_\nu, b_\nu)$) ayant $S(a_n, b_n)$ comme élément.

Dém. 1°) $: a_n, b_n | \subset : S(a_n, b_n)$

2°) Si $: a_n, b_n | < |a_m, b_m|$, alors

$$|a_m, b_m| \subset : S(a_n, b_n)$$

(selon le § 32).

3°) Pour tout point x vers lequel converge une suite régulière en question on a

$$x = \overline{\sum_{\nu/1}^{\infty} : S(a_\nu, b_\nu)} \subset : S(a_n, b_n).$$

Il reste à démontrer qu'il n'existe pas d'autres points contenus dans (1).

En effet n'empiètent pas sur : $S(a_n, b_n)$

α) les $|a_\nu, b_\nu|$ d'ordre inférieur (§ 6),

β) les $|a_\nu, b_\nu|$ d'ordre n mais pour lesquels $(a_\nu, b_\nu) \neq (a_n, b_n)$ (§ 14),

γ) les points vers lesquels convergent les suites régulières ne contenant pas : $S(a_n, b_n)$, car dans le produit

$$\prod_{\nu=1}^{\infty} S(a'_\nu, b'_\nu)$$

figure : $S(a'_n, b'_n)$ qui est disjoint de : $S(a_n, b_n)$ (§ 14).

Il est facile de se rendre compte que nous avons passé en revue tout les points de $\overline{\sum_{\nu=1}^{\infty} (|a_\nu, b_\nu|)}$.

§ 46. A toute dendrite D correspond une suite Ω (ch. e)

$$((a_1, b_1)), ((a_2, b_2)), \dots$$

telle que

$$(1) \quad D = \overline{\sum_{\nu=1}^{\infty} (|a_\nu, b_\nu|)}.$$

Dém. Désignons par $K_1 = |a_1, b_1|$ un arc saturé de D (ch. e. C § 9) et posons

$$K_{\nu+1} = K_\nu + (| : a_{\nu+1}, b_{\nu+1} |),$$

où $(| : a_{\nu+1}, b_{\nu+1} |)$ est le feuillage principal restreint poussant sur K_ν (F § 17).

Les propriétés α et β de la définition du § 1 pour la suite

$$((a_1, b_1)), ((a_2, b_2)), \dots$$

sont évidentes.

Selon F § 10 les K_ν sont des continus. Suivant la définition du feuillage (F § 8) on a pour deux feuilles différentes

$$: a_\nu, b_\nu | \times : a'_\nu, b'_\nu | = 0$$

et cela prouve la propriété γ de la définition du § 1.

Les feuilles $: a_2, b_2 |$ sont toutes en T par rapport à $|a_1, b_1|$, car leurs jonctions ne peuvent pas coïncider avec les points de saturation a_1 et b_1 (D § 25).

Toute feuille $: a_\nu, b_\nu |$ ($\nu \geq 2$) est en T au moins par rapport à un des $|a_{\nu-1}, b_{\nu-1}|$ (F § 20) et la propriété δ de la définition du § 1 se trouve ainsi démontrée.

Il reste à démontrer (1). Si pour un certain ν la somme $(|a_\nu, b_\nu|)$ est nulle, alors (1) est évident.

Dans le cas contraire on a selon F § 18

$$(2) \quad D = K_\nu + (S(a_{\nu+1}, b_{\nu+1})).$$

La classe $((S(a_\nu, b_\nu)))$

est faible (§ 19), donc le maximum τ_ν des diamètres des $S(a_\nu, b_\nu)$ tend vers 0 avec $\frac{1}{\nu}$

$$(3) \quad \lim \tau_\nu = 0.$$

On a selon (2) et selon la définition de ϱ^*

$$(4) \quad \varrho^*(D, K_\nu) \leq \tau_{\nu+1}$$

D'autre part à cause de l'identité

$$\varrho^*(K_\nu, \overline{\sum_{\mu/1}^\infty (|a_\mu, b_\mu|)}) = \varrho^*(\sum_{\mu/1}^\nu (|a_\mu, b_\mu|), \overline{\sum_{\mu/1}^\infty (|a_\mu, b_\mu|)})$$

on a selon A § 48

$$(5) \quad \lim \varrho^*(K_\nu, \overline{\sum_{\mu/1}^\infty (|a_\mu, b_\mu|)}) = 0.$$

Comme

$$\varrho^*(D, \overline{\Sigma}) \leq \varrho^*(D, K_\nu) + \varrho^*(K_\nu, \overline{\Sigma}),$$

donc d'après (3), (4), (5)

$$\varrho^*(D, \overline{\Sigma}) = 0$$

et par conséquent

$$D = \overline{\sum_{\nu/1}^\infty (|a_\nu, b_\nu|)} \quad \text{c. q. f. d.}$$

§ 47. La classe de points de branchement de toute dendrite est effectivement énumérable.

Dém. La classe $((a_n))$ étant effectivement énumérable (§ 18) et toute dendrite étant développable en une suite Ω (§ 46), il suffit de prouver que tout point de branchement fait partie de la classe $((a_n))$.

Supposons qu'il existe un point de branchement c disjoint de la classe $((a_n))$. Comme (§ 39 corr 1, C § 20) aucun point de branchement n'appartient à aucune des classes

$$\overline{\sum_{\nu/1}^\infty (|a_\nu, b_\nu|)} - \sum_{\nu/1}^\infty (|a_\nu, b_\nu|), \quad ((b_n)),$$

done

$$c \in \sum_{\nu/1}^\infty (|a_\nu, b_\nu|).$$

Il existe par conséquent (§ 15) un seul $: a_n, b_n :$ contenant c

$$c \varepsilon : a_n, b_n :$$

et selon C § 18, il existe un arc $[c, d]$ étant en T par à c et $[a_n, b_n]$, donc (§ 40) c appartient à $((a_n))$, ce qui est impossible.

H.

Suites Ω semblables. Homéomorphie de deux dendrites développables en deux suites Ω semblables. Suite Ω spéciale.

§ 1. Deux suites Ω et Ω'

$$((a_1, b_1)), ((a_2, b_2)), \dots$$

$$((\alpha_1, \beta_1)), ((\alpha_2, \beta_2)), \dots$$

correspondant à deux dendrites D_a et D_α sont „semblables“, s'il existe une correspondance biunivoque $\sim$ (ch. e.) entre les paires $(a_\nu, b_\nu), (\alpha_\mu, \beta_\mu)$ satisfaisant aux conditions:

$\alpha)$ Si $(a_\mu, b_\mu) \sim (\alpha_\nu, \beta_\nu)$, alors

$$u = v$$

c. à. d. la transformation $\sim$ ne change pas le rang de (a_μ, b_μ) .

$\beta)$ A tout a_μ correspond ($\sim$) un seul α_μ et inversement. (On dit que $\chi(a_\mu, b_\mu) \sim \psi(\alpha_\mu, \beta_\mu)$, si

$$(a_\mu, b_\mu) \sim (\alpha_\mu, \beta_\mu).$$

$\gamma)$ Entre $[a_\mu, b_\mu]$ et $[\alpha_\mu, \beta_\mu]$ correspondants ($\sim$) il existe une homéomorphie (ch. e) $h(a_\mu, b_\mu)$ dans laquelle

$$1^\circ) \quad a_\mu \quad h(a_\mu, b_\mu) \quad \alpha_\mu$$

$$2^\circ) \quad \text{Si } a_{\mu+1} \varepsilon [a_\mu, b_\mu],$$

alors

$$a_{\mu+1} \quad h(a_\mu, b_\mu) \quad \alpha_{\mu+1}.$$

§ 2. La relation $<$ est invariante par rapport à la transformation $\sim$.

Dém. On peut se borner au cas, où

$$\chi(a_k, b_k) = \psi(a_k, b_k) = (a_k, b_k)$$

(notations du § 24, G).

Supposons

$$\begin{aligned} (a_n, b_n) &< (a_m, b_m) \\ (a_n, b_n) &\sim (\alpha'_n, \beta'_n) \\ (1) \quad (a_n, b_n) &\sim (\alpha_{m'}, \beta_{m'}), \end{aligned}$$

$\sim$ laissant le rang invariant, on a $n = n'$, $m = m'$. D'après la définition de $<$ (G § 26) le théorème se ramène au cas, où

$$(a_n, b_n) \text{ précède immédiatement } (a_m, b_m)$$

c. à. d. (§ 22, G)

$$\begin{aligned} m &= n + 1, \\ (2) \quad a_{n+1} &\varepsilon |a_n, b_n|. \\ \text{De (1) on a} \\ (3) \quad a_{n+1} &\sim \alpha_{n+1}. \end{aligned}$$

La correspondance $h(a_n, b_n)$ est homéomorphe entre $|a_n, b_n|$ et $|\alpha_n, \beta_n|$ [(1), § 1]. et elle fait correspondre au point a_{n+1} , le point α_{n+1} . [(2), (3), § 1 γ].

On a donc

$$\alpha_{n+1} \varepsilon |\alpha_n, \beta_n|$$

et cela signifie (G § 23) que (α_n, β_n) précède immédiatement $(\alpha_{n+1}, \beta_{n+1})$ c. q. f. d.

§ 3. D'après § 2, § 1 et G § 41 la correspondance $\sim$ constitue une correspondance biunivoque entre les suites régulières

$$\chi(a_1, b_1), \dots$$

$$\psi(a_1, b_1), \dots$$

(notations du § 24 G)

§ 4. Définition

$$x H \xi$$

signifie:

il existe une paire (a_n, b_n) , telle que

$$x h(a_n, b_n) \xi.$$

§ 5. La correspondance H est biunivoque entre

$$\sum_{\nu/1}^k (|a_\nu, b_\nu|) \text{ et } \sum_{\nu/1}^k (|\alpha_\nu, \beta_\nu|).$$

Dém. Si $x \varepsilon \sum_{\nu/1}^k (|a_\nu, b_\nu|) - \sum_{\nu/2}^k (a_\nu)$, alors x est contenu dans un seul $|a_n, b_n|$ (G § 16). La correspondance H se réduit donc à la homéomorphie $h(a_n, b_n)$ qui fait correspondre à x un seul point ξ

appartenant à un $|\alpha_n, \beta_n|$ de même rang (§ 1, G § 21).

Si $x = a'_{m+1}$ ($1 < m \leq k-1$),

alors x appartient aux $|a_{m+1}, b_{m+1}|$ pour lesquels $a_{m+1} = a'_{m+1}$ et à $|a_m, b_m|$ qui précède a'_{m+1} (G § 35, G § 17).

A a'_{m+1} correspond ($\sim$) un seul α'_{m+1} (§ 1 β), donc d'après le § 1 γ , la correspondance $h(a_m, b_m)$ et les correspondances $h(a_{m+1}, b_{m+1})$ subordonnent à a'_{m+1} un seul point α'_{m+1} .

D'une manière pareille on démontre qu'à tout ξ de $\sum_{\nu/1}^k |\alpha_\nu, \beta_\nu|$ correspond un seul point x de $\sum_{\nu/1}^k (|a_\nu, b_\nu|)$.

§ 6. H est une correspondance biunivoque entre $\sum_{\nu/1}^\infty (|a_\nu, b_\nu|)$ et $\sum_{\nu/1}^\infty (|\alpha_\nu, \beta_\nu|)$ (selon § 5).

§ 7. $a_n \sim \alpha_n$ équivaut à $a_n H \alpha_n$. (§ 4, § 1 γ).

§ 8. Si $(a_n, b_n) \sim (\alpha_n, \beta_n)$,

alors H constitue une homéomorphie entre $|a_n, b_n|$ et $|\alpha_n, \beta_n|$ dans laquelle a_n et α_n ; b_n et β_n sont associés.

De là résulte le

§ 9.

$$(a_m, b_m) \sim \alpha_m, \beta_m)$$

équivaut à

$$(a_m, b_m) H (\alpha_m, \beta_m).$$

§ 10. H constitue une homéomorphie entre $\sum_{\nu/1}^k (|a_\nu, b_\nu|)$ et $\sum_{\nu/1}^k (|\alpha_\nu, \beta_\nu|)$.

Dém. C'est vrai dans le cas de $k=1$ (§ 1, § 4). Supposons le théorème vrai pour $k=n$ et considérons le cas $k=n+1$.

A chaque $|a_{n+1}, b_{n+1}|$ H fait correspondre son image homéomorphe $|\alpha_{n+1}, \beta_{n+1}|$ (§ 1, § 4, § 5).

Les classes d'arcs simples $(|a_{n+1}, b_{n+1}|)$, $(|\alpha_{n+1}, \beta_{n+1}|)$ sont faibles et empiètent sur $\sum_{\nu/1}^n (|a_\nu, b_\nu|)$ et $\sum_{\nu/1}^n (|\alpha_\nu, \beta_\nu|)$ (selon § 18 G, § 2 G).

H fait une correspondance biunivoque entre $\sum_{\nu/1}^{n+1} (|a_\nu, b_\nu|)$ et $\sum_{\nu/1}^{n+1} (|\alpha_\nu, \beta_\nu|)$.

On en déduit suivant B § 10 que H constitue une homéomorphie entre ces ensembles.

§ 11. Nous étendons la définition de la correspondance H à l'ensemble

$$\overline{\sum_{\nu/1}^{\infty} (|a_{\nu}, b_{\nu}|)} - \sum_{\nu/1}^{\infty} (|a_{\nu}, b_{\nu}|)$$

en subordonnant à tout x de cet ensemble le point ξ vers lequel tend la suite régulière correspondant $(\sim)$ à la suite régulière convergeant vers x .

§ 12. L'extension précédente n'altère pas la définition de H pour $\sum_{\nu/1}^{\infty} (|a_{\nu}, b_{\nu}|)$; elle fait correspondre d'une façon biunivoque aux points de

$$(1) \quad \overline{\sum_{\nu/1}^{\infty} (|a_{\nu}, b_{\nu}|)} - \sum_{\nu/1}^{\infty} (|a_{\nu}, b_{\nu}|)$$

les points de

$$\overline{\sum_{\nu/1}^{\infty} (|\alpha_{\nu}, \beta_{\nu}|)} - \sum_{\nu/1}^{\infty} (|\alpha_{\nu}, \beta_{\nu}|).$$

Dém. Il suffit de remarquer que $\sim$ établit une correspondance biunivoque entre les suites régulières (§ 3) et que d'autre part entre les points de (1) et les suites régulières, il existe une correspondance biunivoque (G § 44).

§ 13. Des §§ 6 et 12 on a:

La correspondance H généralisée est biunivoque entre

$$\overline{\sum_{\nu/1}^{\infty} (|a_{\nu}, b_{\nu}|)} \text{ et } \overline{\sum_{\nu/1}^{\infty} (|\alpha_{\nu}, \beta_{\nu}|)}.$$

§ 14. Si $(a_n, b_n) H (\alpha_n, \beta_n)$,

alors

$$(1) \quad \left[: S(a_n, b_n) \times \overline{\sum_{\nu/1}^{\infty} (|a_{\nu}, b_{\nu}|)} \right] H \left[: S(\alpha_n, \beta_n) \times \overline{\sum_{\nu/1}^{\infty} (|\alpha_{\nu}, \beta_{\nu}|)} \right].$$

Dém. En effet d'après G § 45 le premier membre de cette relation est composé de $: a_n, b_n[$, des $|a_{\nu}, b_{\nu}|$ suivants $: a_n, b_n[$ et des points vers lesquels convergent les suites régulières

$$: S(a_1, b_1), \dots$$

renfermant $: S(a_n, b_n)$.

Or les propriétés de ces objets ne varient pas après la transformation H (§ 9, § 2, § 3), donc (G § 45) l'image du premier membre de (1) est contenue dans le second

L'inverse se prouve de la même façon et (1) est ainsi établi.

§ 15. H établit une homéomorphie entre

$$\overline{\Sigma}_a = \overline{\sum_{\nu=1}^{\infty} (|a_{\nu}, b_{\nu}|)} \text{ et } \overline{\Sigma}_\alpha = \overline{\sum_{\nu=1}^{\infty} (|\alpha_{\nu}, \beta_{\nu}|)}.$$

Dém. H étant une correspondance biunivoque, il suffit de démontrer la continuité de H et de la correspondance inverse H^{-1} . Pour raison de symétrie il suffit de faire la démonstration pour H .

Soit $x_0 \in \overline{\Sigma}_a$

et $x_0 H \xi_0$.

Deux cas se présentent

$$(1) \quad \text{I) } x_0 \in \sum_{\nu=1}^{\infty} (|a_{\nu}, b_{\nu}|) = \Sigma_a$$

$$(2) \quad \text{II) } x_0 \in \overline{\Sigma}_a - \Sigma_a$$

Cas I. D'après (1) il existe un n (ch. e.), tel que

$$x_0 \in \sum_{\nu=1}^n (|a_{\nu}, b_{\nu}|) = \Sigma_a^n.$$

H constitue une homéomorphie entre Σ_a^n et Σ_α^n (§ 10), on a donc

$$\xi_0 \in \Sigma_\alpha^n.$$

Comme (G § 37)

$$\overline{\Sigma}_a = \Sigma_a^{n+1} + (:S(a_{n+1}, b_{n+1})) \times \overline{\Sigma}_a$$

$$\overline{\Sigma}_\alpha = \Sigma_\alpha^{n+1} + (:S(\alpha_{n+1}, \beta_{n+1})) \times \overline{\Sigma}_\alpha,$$

il suffit de démontrer que:

1°) Si un point x variable dans Σ_a^{n+1} tend vers x_0 , alors son image tend vers ξ_0 .

2°) Si un point x variable dans $:S(a_{n+1}, b_{n+1}) \times \overline{\Sigma}_a$ tend vers x_0 , alors son image tend vers ξ_0 .

La démonstration de 1°) résulte de ce que H constitue une homéomorphie entre Σ_a^{n+1} et Σ_α^{n+1} (§ 10).

Pour établir 2°) rappelons que (G § 2)

$$(3) \quad \begin{aligned} (a_{n+2}) &\subset \Sigma_a^{n+1} \\ (\alpha_{n+2}) &\subset \Sigma_\alpha^{n+1} \end{aligned}$$

et que $(|a_{n+1}, b_{n+1}|)$ est un feuillage en T par rapport au feuillage $(|a_n, b_n|)$ (G § 5).

Si un point x variable dans $(S(a_{n+2}, b_{n+2})) \times \bar{\Sigma}_a$ tend vers $x_0 \in \Sigma_a$, alors (F § 16) il entraîne avec lui vers x_0 les jonctions a_{n+2} des $S(a_{n+2}, b_{n+2})$ dans lesquels il est contenu. D'après (3) et la homéomorphie entre Σ_a^{n+1} et Σ_x^{n+1} l'image de ce a_{n+2} variable tend vers ξ_0 . Cette image est un α_{n+2} (§ 8). Il entraîne avec lui vers ξ_0 tous les points des $S(\alpha_{n+2}, \beta_{n+2})$ ayant ces points α_{n+2} comme jonctions [F § 16].

Mais ces $S(\alpha_{n+2}, \beta_{n+2})$ renferment l'image ξ du point variable x (§ 14), donc cet image tend vers ξ_0 c. q. f. d.

Cas II. Dans ce cas nous avons [(2), § 12]

$$\xi_0 \in \bar{\Sigma}_x - \Sigma_a$$

et il existe d'après G § 43 une suite régulière

$$(a_1, b_1), \dots$$

telle que

$$(4) \quad x_0 \in \bigcap_{v/1}^{\infty} S(a_v, b_v)$$

et on a (§ 11, § 12)

$$\xi_0 = \bigcap_{v/1}^{\infty} S(\alpha_v, \beta_v).$$

où

$$(a_v, b_v) \in H(\alpha_v, \beta_v).$$

Le diamètre de $S(\alpha_v, \beta_v)$ tend vers 0 avec $\frac{1}{v}$ (G § 19), donc pour un certain n

$$\tau(S(\alpha_n, \beta_n)) < \varepsilon.$$

D'après (4) et D § 10 un certain voisinage δ du point x_0 ne comprend aucun point de $D_a - S(a_n, b_n)$.

Il suffit de rappeler que (§ 14)

$$S(a_n, b_n) \times \bar{\Sigma}_a \subset H(S(\alpha_n, \beta_n) \times \bar{\Sigma}_x,$$

pour conclure à la continuité.

En rappelant que H est parfaitement déterminé on a le théorème suivant:

§ 16. Si deux suites Ω et Ω'

$$((a_1, b_1)), \dots$$

$$((\alpha_1, \beta_1)), \dots$$

correspondant à deux dendrites D_a et D_α sont semblables, alors il existe une homéomorphie (ch. e) entre

$$\overline{\sum_{v/1}^{\infty}(|a_v, b_v|)} \text{ et } \overline{\sum_{v/1}^{\infty}(|\alpha_v, \beta_v|)}.$$

§ 17. Nous appelons „suite Ω , spéciale“ toute suite Ω

$$((a_1, b_1)), ((a_2, b_2)), \dots$$

se distinguant par les propriétés:

1°) Tout $|a_v, b_v|$ renferme un ensemble dénombrable des a_{v+1} dense sur $|a_v, b_v|$.

2°) A tout a'_μ ($\mu \geq 2$) il existe une infinité dénombrable des (a_μ, b_μ) pour lesquels

$$a_\mu = a'_\mu.$$

§ 18. A toute suite Ω correspond dans toute suite Ω , spéciale une suite Ω^* particulière semblable à Ω (ch. e).

Par la suite particulière, nous comprenons toute suite Ω obtenue de la suite Ω , par élimination de quelques paires.

Dém. Soit $((a_1, b_1)), \dots$

la suite Ω donnée et

$$(|\alpha_1, \beta_1|), \dots$$

une suite spéciale.

Les a_2 placés sur $|a_1, b_1|$ forment un ensemble effectivement énumérable et sont situés sur a_1, b_1 : (G § 1, G § 19).

Les α_2 placés sur $|\alpha_1, \beta_1|$ forment un ensemble dense sur $|\alpha_1, \beta_1|$ et situé sur α_1, β_1 : Il existe donc une homéomorphie (ch. e) $h(a_1, b_1)$ entre $|a_1, b_1|$ et $|\alpha_1, \beta_1|$ dans laquelle se correspondent a_1 et α_1 , b_1 et β_1 et, telle que tout point a_2 correspond à un certain point α_2 (B § 11). Désignons par $(\alpha_2)'$ l'ensemble de ces derniers α_2 . Désignons par $(\alpha_2, \beta_2)^*$ la classe de tous les (α_2, β_2) dont le premier élément appartient à $(\alpha_2)'$. A chaque point α_2 de $(\alpha_2)'$ correspond une infinité effectivement énumérable des (α_2, β_2) ayant cet α_2 comme premier élément (§ 17 et G § 19).

Désignons par a'_2 un point quelconque de la classe (α_2) et par α'_2 son correspondant suivant l'homéomorphie $h(a_1, b_1)$. Faisons correspondre à toute paire pour laquelle $a_2 = a'_1$ un (α_2, β_2) pour lequel $\alpha_2 = \alpha'_2$ de façon à faire correspondre aux (a_2, b_2) différents les (α_2, β_2) différents.

Quelques (α_2, β_2) peuvent rester en dehors de cette corres-

pondance. Cette correspondance ($\sim$) peut être établie (ch. e.), car il y a un ensemble effectivement énumérable des (a_2, b_2) en question et une infinité effectivement énumérable des (α_2, β_2) (§ 17 et G § 19).

Désignons par $(\alpha_2, \beta_2)'$ la classes des (α_2, β_2) ainsi choisis dans $(\alpha_2, \beta_2)^*$. ($\sim$) établit entre $((a_2, b_2))$ et $(\alpha_2, \beta_2)^*$ une correspondance biunivoque

Nous allons construire la fonction $h(a_2, b_2)$.

Soit

$$(a'_2, b'_2) \sim (\alpha'_2, \beta'_2)$$

La classe des a_3 situés sur $[a'_2, b'_2]$ est vide ou effectivement énumérable; elle est contenue dans $[a'_2, b'_2]$. La classe des α_3 situés sur $[\alpha'_2, \beta'_2]$ est infinie, effectivement énumérable et dense sur $[\alpha'_2, \beta'_2]$.

On peut donc (B § 11) établir une homéomorphie (ch. e.) entre $[a'_2, b'_2]$ et $[\alpha'_2, \beta'_2]$ suivant laquelle au point a'_2 correspond α'_2 et à tout point a_3 situé sur $[a'_2, b'_2]$ correspond un point α_3 situé sur $[\alpha'_2, \beta'_2]$.

A l'aide de l'induction mathématique on continue ainsi et on obtient une suite qui est facile à reconnaître comme une suite Ω choisie dans la suite spéciale Ω_* . La correspondance $\sim$ forme une ressemblance (suivant la loi même de sa formation) entre la suite Ω donnée et la suite choisie dans Ω_* .

§ 19. S'il existe une dendrite plane D^* développable en une suite spéciale Ω_* , alors toute dendrite est homéomorphe avec une partie de D^* (ch. e.) et par conséquent toute dendrite constitue une multiplicité à deux dimensions (§ 1, § 16).

Les chapîtres suivants seront consacrés à la démonstration de l'existence de D^* .

J.

Arbre généralisé.

Dans ce chapitre nous faisons abstraction des notations et des notions des chapîtres précédents.

Nous allons traiter les ensembles contenus dans un espace normal $E(\varphi)$.

Nous ne nous bornerons pas comme dans les chapîtres qui précèdent aux sous-ensembles d'une certaine dendrite.

§ 1. Déf. On appelle feuillage complet par rapport à un continu G (tige du feuillage) toute classe (F) de continus F ne se réduisant pas à un seul point (feuilles complètes) telle que

1°) $F \times G = 1$ point = jonction de $F = j(F)$

2°) En posant par définition

$$:F = F - j(F)$$

on a pour deux feuilles différentes

$$:F \times F' = 0.$$

3°) la classe (F) est faible.

On appelle arbre la somme

$$A = G + (F).$$

On appelle $:F$ feuille restreinte,

$(:F)$ feuillage restreint

et on pose

$$\text{jonction de } :F = j(F) = j(:F).$$

Tout arbre est un continu (B § 3).

§ 3. Evidemment

$$G + (:F) = G + (F)$$

§ 4 On a

$$A - :F^* = G + (F)',$$

où $(F)'$ est la somme des feuilles différentes de la feuille F^* .

§ 5. La somme de la tige et de quelques feuilles est un arbre, donc un continu.

§ 6. $A - :F^*$ est un continu.

§ 7. Toute classe composée d'une seule feuille F^* d'un arbre A constitue un feuillage par rapport à la classe $A - :F^*$ c. à d. à la classe composée de la tige G de l'arbre A et des feuilles de cet arbre différentes de F^* . La jonction de F^* par rapport à la nouvelle tige $A - :F^*$ reste invariée (§ 1, § 5).

Ce théorème nous permettra bien des fois de ramener les démonstrations au cas de l'arbre ayant une seule feuille.

§ 8. Le produit de tout sous-continu K d'un arbre A et d'une quelconque de ces feuilles complètes F est un continu ou bien forme un ensemble vide. Si K empiète sur la tige de A , alors ce produit y empiète aussi.

La démonstration ne représente pas de difficulté.

§ 9. Tout continu C contenu dans un arbre A et empiétant sur une feuille (complète ou restreinte) et son complémentaire passe par sa jonction.

La démonstration se fait facilement au moyen du § 7.

§ 10. Si un arbre A contient une orbe de Jordan sans qu'elle soit renfermée dans sa tige, alors il existe une seule feuille complète qui la contient (§ 9).

Remarque. Si ni la tige ni aucune feuille complète ne contient aucune orbe, alors l'arbre n'en contient non plus.

§ 11. Est courbe de Jordan tout arbre dont la tige et les feuilles complètes sont des courbes de Jordan.

Dém. Le théorème est vrai pour le cas d'un nombre fini de feuilles (B § 8). Le critère de M. Sierpiński (B § 6 β'') nous permettra de démontrer le cas général.

Soit $\varepsilon > 0$.

Le feuillage étant une classe faible de courbes de Jordan, il existe au plus un nombre fini de feuilles dont le diamètre dépasse $\frac{\varepsilon}{2}$. La somme de la tige et de ces feuilles est une courbe de Jordan, elle se laisse donc décomposer en un nombre fini des continus

$$C_1, \dots, C_n$$

au diamètre $< \frac{\varepsilon}{2}$. Toute feuille restante empiète au moins sur un d'eux. En joignant à C_v , ces feuilles qui empiètent sur C_v , on obtient un continu C_v^* (B § 3). Evidemment

$$A = C_1^* + \dots + C_n^*$$

est une décomposition satisfaisant aux conditions du critère.

§ 12. Si la tige et les feuilles d'un arbre sont dendrites, alors l'arbre l'est aussi. (§ 10, § 11).

§ 13. Si un arc simple $[a, b]$ contenu dans la tige G d'un arbre A est saturé au point a par rapport à la tige G et que la jonction d'aucune feuille ne coïncide pas avec a , alors $[a, b]$ est saturé au point a par rapport à l'arbre A .

Dém. Supposons que $[a, b]$ se laisse prolonger au delà de a sans quitter A et soit $[a, c]$ l'arc simple constituant ce prolongement. $[a, c]$ ne peut pas pénétrer dans deux feuilles restreintes différentes, car il serait obligé de pénétrer dans une d'elle et de la quitter, donc il passerait deux fois par sa jonction (§ 9). Soit j la jonction de cette feuille dans laquelle il pénètre. $j \neq a$. $[a, c]$ passe une fois

par j . L'arc $[a, j]$ est évidemment contenu dans G , donc $[a, b]$ se laisse prolonger dans G , ce qui est impossible.

§ 14. Feuillage enveloppant étroitement un autre feuillage.

Si (F) est un feuillage poussant sur G (c. à d. feuillage par rapport à G) et, si tout point de la somme

$$(:F)$$

est intérieur à un ensemble I (I est le même pour tous les $:F$), alors il existe une classe de continus $(I(F))$ (ch. e.), telle que

1°) Tout point de $:F$ est un point intérieur par rapport à $I(F)$,

$$2^0) \tau(I(F)) \leq 2 \tau(F),$$

3°) $(I(F))$ forme un feuillage comptet poussant sur G ,

$$4^0) j(I(F)) = j(F),$$

$$5^0) :I(F) \subset \bar{I}.$$

Nous appelons le nouveau feuillage feuillage enveloppant étroitement le feuillage (F) , et le nouvel arbre $I(A)$ l'arbre enveloppant étroitement l'arbre précédent A .

Dém. On a pour tout f , tel que

$f \in :F$ la relation:

$$\varrho(f, A - :F) > 0 \quad (\S 1, \S 6, \S 7).$$

Entourons tout point f de cette sorte par la sphère $Sph(f, r(f))$, où

$$(1) \quad \begin{aligned} 0 < r(f) &< \frac{1}{3} \varrho(f, A - :F) \\ 0 < r(f) &< \frac{1}{2} \tau(F) \end{aligned}$$

$r(f)$ peut être effectivement choisi.

Posons

$$I(F) = \overline{(Sph(f, r(f)))}$$

où $()$ est étendu sur tous ces f de $:F$. $I(F)$ satisfait suivant sa définition aux propriétés

$$1^0, 2^0 \text{ et } 5^0,$$

si l'on suppose en plus que

$$Sph(f, r(f)) \subset I$$

et cette hypothèse est évidemment réalisable.

Pour prouver 3° il suffit d'établir (§ 1) que

α) $I(F)$ est un continu,

β) $I(F) \times G = 1$ point $= j$,

γ) $I(F) \times I(F') = 0$ pour

$$F \neq F',$$

δ) $(I(F))$ forme une classe faible.

Or α résulte immédiatement de la définition de $I(F)$. δ est vrai, car (F) est une classe faible et $\tau(I(F)) \leq 2\tau(F)$.

La propriété β se trouvera démontrée, si l'on prouve que

$$(2) \quad I(F) \times G = F \times G = \text{jonc } (F) = j.$$

Selon 1° on a

$$j(F) = F \times G \subset I(F) \times G.$$

Pour démontrer l'inclusion inverse, il suffit de prouver que pour tout point g , tel que

$$(3) \quad g \in G - j,$$

on a

$$q(g, I(F)) > 0.$$

Selon (3) on a

$$(4) \quad q(g, F) > 0.$$

Soit s un point quelconque de la somme déjà introduite

$$(Sph(f, r(f))) \quad (f \in F)$$

et supposons que par exemple

$$s \in Sph(f, r(f)).$$

Selon (1) on a

$$q(f, s) \leq r(f) \leq \frac{1}{2} q[f, A - F],$$

donc selon (3) à fortiori

$$q(f, s) \leq \frac{1}{2} q(f, g).$$

Comme

$$q(f, g) \leq q(f, s) + q(s, g),$$

donc

$$\frac{2}{3} q(f, g) \leq q(s, g)$$

et à fortiori

$$\frac{2}{3} q(F, g) \leq q(s, g).$$

Mais s est un point arbitraire de $(Sph(f, r(f)))$, donc

$$\frac{2}{3} q(F, g) \leq q((Sph), g) = q(\overline{(Sph)}, g)$$

c'est à dire selon la définition de $I(F)$ et selon (4)

$$q(I(F), g) > 0, \quad \text{c. q. f. d.}$$

Observons en passant que (2) exprime la vérité de la propriété 4°.

Avant d'aborder la démonstration de γ nous prouverons le lemme:

Si l'intérieur d'une sphère Σ de rayon r empiète sur $I(F)$ et sur $I(F')$, alors l'intérieur de la sphère concentrique de rayon $3r$ empiète sur F et F' .

Soit c le centre de la sphère Σ . On a

$$I(F) = \overline{(Sph(f, r(f)))}$$

$$I(F') = \overline{(Sph(f', r(f')))}.$$

Il existe donc suivant l'hypothèse au moins une $Sph(f, r(f))$ et au moins une $Sph(f', r(f'))$ et d'autre part au moins une couple de points i et i' , tels que

$$i \in Sph(f, r(f)) \times \Sigma$$

$$i' \in Sph(f', r(f')) \times \Sigma.$$

On peut facilement se rendre compte que tous ces objets peuvent être choisis effectivement.

On a

$$(5) \quad \begin{aligned} \varrho(c, i) &< r \\ \varrho(c, i') &< r \end{aligned}$$

et suivant la définition de $I(F)$

$$(6) \quad \varrho(f, i) \leq r(f) \leq \frac{1}{3} \varrho(f, A - : F) \leq \frac{1}{3} \varrho(f, F') \leq \frac{1}{3} \varrho(f, f')$$

et d'une façon analogue

$$(7) \quad \varrho(f', i') \leq \frac{1}{3} \varrho(f', f).$$

Evidemment

$$\varrho(f, f') \leq \varrho(f, i) + \varrho(i, c) + \varrho(c, i') + \varrho(i', f'),$$

donc (5), (6), (7)

$$\varrho(f, f') < \frac{1}{3} \varrho(f, f') + r + r + \frac{1}{3} \varrho(f, f')$$

c. à. d.

$$(8) \quad \frac{1}{3} \varrho(f, f') < 2r.$$

On a

$$\varrho(c, f) \leq \varrho(f, i) + \varrho(i, c),$$

donc selon (6) et (5)

d'où d'après (8) $\varphi(c, f) < \frac{1}{3} \varphi(f, f') + r,$

et pareillement $\varphi(c, F) < 3r$

$\varphi(c, F') < 3r,$ c. q. f. d.

Il est à présent facile de démontrer la propriété γ c. à d. que

$$:I(F) \times :I(F') = 0 \quad (\text{pour } F \neq F').$$

Pour cela il suffit de prouver que

$$(9) \quad I(F) \times I(F') \subset j \times j'$$

parce que

$$:I(F) = I(F) - j,$$

$$:I(F') = I(F') - j'.$$

Soit c un point, tel que

$$c \in I(F) \times I(F')$$

et soit $\varepsilon > 0$.

Entourons c de la sphère de rayon $\frac{1}{3}\varepsilon$. D'après le lemme précédent on a:

$$\varphi(c, F) < \varepsilon$$

$$\varphi(c, F') < \varepsilon$$

ε étant arbitraire et F et F' fermés, on a:

$$c \in F \times F',$$

donc

$$c \in [j + :F] \times [j' + :F'] = j \times j'$$

suivant les propriétés du feuillage (F), et cela prouve (9), car selon la propriété 4^o (déjà démontrée) on a

$$j(F) = j(I(F)) = j,$$

$$j(F') = j(I(F')) = j'.$$

§ 15. Si un arc $|a, b|$ contenu dans une feuille F d'un arbre A

$$(1) \quad |a, b| \subset F$$

est saturé au point b par rapport à F et est distinct de $j(F)$

$$(2) \quad b \neq j(F),$$

alors $|a, b|$ est saturé par rapport à A au point b .

Dém. Suivant le § 7 on peut ramener la démonstration au cas où il existe une seule feuille. D'après (1) et (2) on a

$$b \varepsilon : F.$$

Il existe évidemment une sphère $Sph(b, \delta)$ n'empiétant pas sur la tige G de l'arbre.

Si $|a, b|$ n'était pas saturé au point b par rapport à A , il se laisserait prolonger au delà de b de façon à ne pas quitter A .

Or un petit morceau de ce prolongement serait situé dans $Sph(b, \delta)$ donc dans $:F$, ce qui est impossible, car $|a, b|$ est saturé au point b par rapport à $:F$.

§ 16. En se basant sur le § 7 et sur le § 14, on obtient

$$A - :F \subset I(A) - :I(F),$$

où $I(A)$ est l'arbre enveloppant l'arbre A .

§ 16. On appelle une feuille F feuille simple, si elle est un arc simple, dont une des extrémités coïncide avec la jonction de F . On appelle l'autre extrémité de cet arc „extrémité libre“.

On appelle feuillage simple et arbre simple tout feuillage et tout arbre aux feuilles exclusivement simples.

§ 18. Tout arc formant une feuille simple est saturé à son extrémité libre. (§ 1).

§ 19 Si A est un arbre simple à la tige G et au feuillage (F) , et si l'on construit sur chaque feuille F considérée comme tige un feuillage simple $(F_1)_F$ dont toute feuille est en T par rapport à F et qui satisfait à la relation

$$F_1 \subset I(F),$$

alors

1°) la classe $((F_1)_F)$ composée des feuilles F_1 de tous ces nouveaux feuillages constitue un feuillage simple poussant sur A considéré comme tige nouvelle.

Nous l'appelons feuillage intermédiaire entre les feuillages (F) et $(I(F))$.

2°) Aucune jonction du nouveau feuillage n'empiète sur la tige ancienne G .

3°) Toute feuille de l'arbre ancien est saturé à son extrémité libre par rapport à l'arbre nouveau, [l'arbre intermédiaire entre les arbres A et $I(A)$.]

La démonstration de 1°) présente une grande analogie avec la démonstration de F § 12. Nous n'insistons donc sur cette démonstration et nous indiquons seulement que l'analogie consiste dans la correspondance des objets suivants,

$$F \text{ et } |a_1, b_1|,$$

$$:F_1 \text{ et } :a_1, b_1|$$

$$:I(F) \text{ et } :S(a_1, b_1).$$

La démonstration de la seconde partie se déduit facilement de F § 13.

Enfin la partie 3° résulte des parties précédentes, si l'on tient compte des §§ 18 et 15.

§ 20. Tout arbre A_1 intermédiaire entre l'arbre A et l'arbre $I(A)$ qui l'enveloppe est contenu dans $I(A)$.

§ 21. On appelle une paire de suites d'arbres

$$A_1, A_2, \dots$$

$$I(A_1) I(A_2), \dots$$

régulière, si

$\alpha)$ A_1 est un arbre simple dont la tige G est un arc simple $|a_0, b_0|$ par rapport auquel toutes les feuilles sont en T .

$\beta)$ A_ν est un arbre simple intermédiaire entre $A_{\nu-1}$ et $I(A_{\nu-1})$ ($\nu \geq 2$).

$\gamma)$ Le maximum des diamètres des feuilles de l'arbre A_ν tend vers 0 avec $\frac{1}{\nu}$.

$$\delta) \quad I(A_{\nu-1}) \subset I(A_\nu).$$

Nous désignons le feuillage de A_ν par $(:a_\nu, b_\nu|)$

§ 22. En gardant les notations du paragraphe précédent, on a

I) $\overline{\sum_{\nu=1}^{\infty} A_\nu}$ est une dendrite (nous la désignons par D).

II) La tige $|a_0, b_0|$ de A_1 est un arc simple saturé par rapport à D dans ses deux extrémités.

III) Toute feuille de tout arbre A_μ est saturée par rapport à D dans son extrémité libre.

IV) Tout b_ν est un point principal de $:S(a_\nu, b_\nu)$ ($\nu \geq 1$).

Dém. Suivant la définition de la suite régulière, celle de

l'arbre intermédiaire et celle de l'arbre enveloppant on a

- (1) $A_n \subset A_{n+k}$
 (2) $I(A_{n+k}) \subset I(A_n)$,
 (3) $A_{n+k} \subset I(A_{n+k})$,
 ($n \geq 1, k \geq 0$).

Suivant (1)

$$\sum_{v/1}^{n+k} A_v = A_{n+k},$$

donc [(2), (3)]

$$\sum_{v/1}^{n+k} A_v \subset I(A_n),$$

d'où en passant à la limite et en observant que le second membre est fermé (§ 5), on a

$$(4) \quad \overline{\sum_{v/1}^{\infty} A_v} \subset I(A_n).$$

Lemme 1. A_{n-1} est la tige de chacun des arbres A_n et $I(A_n)$. (§ 19, 1°; § 14, 3°).

Lemme 2. A_n est une dendrite.

En effet A_1 est une dendrite, car A_1 est un arbre simple

α) dont la tige est un arc simple (§ 21, α), donc A_1 est une dendrite (§ 12).

Si A_v est une dendrite, A_{v+1} l'est aussi, car sa tige A_v est une dendrite (lemme 1) et ses feuilles sont des arcs simples (§ 21, β ; § 12).

Démonstration de I. Nous allons démontrer que $\overline{\sum_{v/1}^{\infty} A_v} = \bar{\Sigma}$ est un continu satisfaisant à la condition de M. Sierpiński (B § 6 β'') et ne contenant aucune orbe et, par cela même, nous démontrerons que $\bar{\Sigma}$ est une dendrite.

$\bar{\Sigma}$ est un continu (A § 48, A § 63). Soit $\varepsilon > 0$. Supposons qu'une orbe J fait partie de $\bar{\Sigma}$. Soit ε_1 , tel que

$$0 < \varepsilon_1 < \frac{\varepsilon}{3},$$

(5)

$$\varepsilon_1 < \tau(J).$$

D'après le § 21, il existe un nombre naturel n tel que

pour les feuilles F_{n+1} de l'arbre A_{n+1} on a

$$\tau(F_{n+1}) < \frac{\varepsilon_1}{2}.$$

$I(F_{n+1})$ enveloppe étroitement F_{n+1} , donc (§ 14, 2°)

$$(6) \quad \tau(I(F_{n+1})) < \varepsilon_1$$

Selon (4)

$$(7) \quad J \subset A_n + (I(F_{n+1})),$$

$$(8) \quad \overline{\sum_{\nu/1}^{\infty} A_{\nu}} \subset A_n + (I(F_{n+1})) = I(A_{n+1}).$$

La tige A_n de $I(A_{n+1})$ est une dendrite, donc elle ne contient aucune orbe et aucune feuille $I(F_{n+1})$ ne contient pas J selon (5). (7) est donc impossible (§ 10, remarque) c. à. d. $\bar{\Sigma}$ ne contient aucune orbe.

(8) donne

$$(9) \quad \overline{\sum_{\nu/1}^{\infty} A_{\nu}} = A_n + (I(F_{n+1}) \times \overline{\sum_{\nu/1}^{\infty} A_{\nu}}).$$

$\overline{\sum_{\nu/1}^{\infty} A_{\nu}}$ est un sous-continu de l'arbre $I(A_{n+1})$, donc d'après le § 8

$$I(F_{n+1}) \times \overline{\sum_{\nu/1}^{\infty} A_{\nu}}$$

est un continu empiétant sur A_n . [Pour appliquer le § 8 il suffit de remarquer que $\bar{\Sigma}$ empiète sur la tige A_n de $I(A_{n+1})$]. Suivant le critère mentionné A_n étant une courbe de Jordan (lemme 2) se laisse représenter comme une somme finie

$$C_1 + \dots + C_n$$

de continus au diamètre $< \varepsilon_1$.

En joignant à C_{μ} ($\mu/1, \dots, n$) tous les continus $I(F_{n+1}) \times \overline{\sum_{\nu/1}^{\infty} A_{\nu}}$ qui empiètent sur lui et en fermant l'ensemble ainsi obtenu nous avons d'après (9) une décomposition de $\overline{\sum_{\nu/1}^{\infty} A_{\nu}}$ en n continus de diamètre au plus égal à $3 \varepsilon_1 = \varepsilon$, c. q. f. d.

Démonstration de II. La tige de A_1 c. à. d. l'arc $[a_0, b_0]$ est saturé par rapport à lui même aux points a_0 et b_0 . Toute feuille F_1 de A_1 est en T par rapport à $[a_0, b_0]$ (§ 21 α), sa jonction

n'empiète donc pas sur $a_0 + b_0$. En raison de

$$j(F_1) = j(I(F_1)) \quad (\S 14, 4^0)$$

les jonctions du feuillage $(I(F_1))$ n'empiètent pas sur $a_0 + b_0$. $|a_0, b_0|$ est donc (selon le § 13) saturé par rapport à l'arbre

$$|a_0, b_0| + (I(F_1)) = I(A_1),$$

donc il est saturé à fortiori par rapport à $\overline{\sum_{\nu/1}^\infty A_\nu}$, car (8)

$$|a_0, b_0| \subset \overline{\sum_{\nu/1}^\infty A_\nu} \subset I(A_1).$$

La démonstration de III. découle d'une façon analogue des §§ 18 et 13.

Démonstration de IV. Pour prouver que b_ν est un point principal de $:S(a_\nu, b_\nu)$, il suffit de prouver selon E § 1 que

$$b_\nu \varepsilon :S(a_\nu, b_\nu)$$

$$(10) \quad \tau(a_\nu, b_\nu) \geq \frac{\tau(:S(a_\nu, b_\nu))}{3}$$

et que $|a_\nu, b_\nu|$ est saturé au point b_ν par rapport à $\overline{\sum_{\nu/1}^\infty A_\nu}$ ($= D$).

Or la première relation est évidente, la dernière relation a lieu en raison de III.

Reste à démontrer (10).

On a pour $\nu > 1$ (le cas de $\nu = 1$ se traite d'une façon analogue) selon (8):

$$(11) \quad D \subset A_{\nu-1} + (I(|a_\nu, b_\nu|)) = I(A_\nu).$$

D'autre part (§ 14, 2°, 3°, 4°)

$$(12) \quad \begin{aligned} a_\nu &= j(I(a_\nu, b_\nu)) \\ \tau(I(|a_\nu, b_\nu|)) &\leq 2 \tau(|a_\nu, b_\nu|). \end{aligned}$$

Pour démontrer (10) il suffit de faire voir que

$$(13) \quad :S(a_\nu, b_\nu) \subset :I(|a_\nu, b_\nu|)$$

et pour cela suivant la définition de $:S(a_\nu, b_\nu)$, il suffit de prouver que pour tout continu K , tel que

$$(14) \quad K \varepsilon :S(a_\nu, b_\nu),$$

$$(15) \quad b_v \in K,$$

$$(16) \quad K \times a_v = 0,$$

on a

$$(17) \quad K \subset I(|a_v, b_v|).$$

Ee (11) et (14) on a

$$K \subset I(A_v)$$

et d'après (15), (12) et (16) K empiète sur la feuille $I(a_v, b_v)$ sans empiéter sur sa jonction, donc, selon le § 9, (17) est vrai. (13) est donc démontré.

Remarquons pour un but ultérieur que selon (12) on a

$$(18) \quad j(I(a_v, b_v)) = j[:S(a_v, b_v)].$$

§ 23, Des résultats du théorème précédent et des relations (13) et (18) on conclut facilement que

$$((a_v, b_v)), (a_1, b_1), \dots$$

est une suite Ω et qu'on a

$$D = \overline{\sum_{v/1}^{\infty} (|a_v, b_v|)} = \overline{\sum_{v/1}^{\infty} A_v}.$$

(Notations du § précédent).

K.

Construction d'une dendrite plane développable en une suite Ω spéciale.

§ 1. Si $|a, b|$ est un segment quelconque du plan et I un ensemble plan, tel que

$\alpha)$ tout point de $:a, b:$ est un point intérieur de I , alors

A) il existe une classe

$$Cl(I, |a, b|, n)$$

de segments rectilignes

$$(|a_1, b_1|)$$

1°) formant un feuillage simple possant sur $|a, b|$,

2°) dont les points de jonction constituent une classe (a_1) effectivement énumérable, dense sur $|a, b|$ et située sur $:a, b:$,

3°) tout point de jonction est commun à une infinité (effectivement énumérable) de feuilles,

4°) chaque feuille a le diamètre inférieur à $\frac{1}{n}$,

5°) chaque feuille est intérieure à I par rapport au plan

Dém. Supposons que $|a, b|$ est situé sur l'axe Ox d'un système de coordonnées rectangulaires et supposons que a est à gauche de b . Cela peut se faire par un changement des coordonnées (ch. e).

Choisissons sur $|a, b|$ une classe effectivement énumérable et dense sur $|a, b|$ (ch. e. A § 43).

Soit une fonction $f(\alpha_v)$, telle que

$$\sum_{v/1}^{\infty} f(\alpha_v) = \frac{\pi}{2}$$

$$f(\alpha_v) > 0 \quad (\text{ch. e.})$$

et considérons deux classes de demidroites

$$(R_\mu) \text{ et } (R'_\mu)$$

dont les équations sont

$$\left. \begin{aligned} x &= \alpha_\mu + \varrho \cos \left[-\pi + \sum_{\alpha_v < \alpha_\mu} f(\alpha_v) \right] \\ y &= \varrho \sin \left[-\pi + \sum_{\alpha_v < \alpha_\mu} f(\alpha_v) \right] \end{aligned} \right\} R_\mu$$

$$\varrho \geq 0$$

$$\left. \begin{aligned} x &= \alpha_\mu + \varrho \cos \left[-\pi + \sum_{\alpha_\mu \leq \alpha_v} f(\alpha_v) \right] \\ y &= \varrho \sin \left[-\pi + \sum_{\alpha_\mu \leq \alpha_v} f(\alpha_v) \right] \end{aligned} \right\} R'_\mu$$

$$\varrho \geq 0.$$

Il est facile de prouver que l'angle $\angle (R_\mu, R'_\mu)$ situé au dessus de l'axe Ox n'empiète sur aucun $\angle (R_\mu, R'_\mu)$ ($\mu \neq \mu_1$).

Du point α_μ traçons une infinité de demi-droites effectivement énumérable (ch. e.) et comprises entre R_μ, R'_μ dans l'angle $\angle (R_\mu, R'_\mu)$ et sur chacun de ces rayons découpons un segment $|\alpha_\mu, \beta_{\mu v}|$ de longueur inférieure à $\frac{1}{n} \cdot \frac{1}{2^{\mu+v}}$ et dont tout point est intérieur à I (ch. e.).

C'est possible, car tout α_μ est intérieur à I .

Il est clair qu'en posant

$$Cl(I, |a, b|, n) = (|\alpha_\mu, \beta_\mu|),$$

on obtient une classe satisfaisant aux conditions 1° — 5°.

§ 2. Soit A_1 un arbre simple sur le plan dont toutes les feuilles sont rectilignes et soit $I(A_1)$ un arbre enveloppant étroitement l'arbre A_1 .

Nous disons que la relation

$$(A_1, I(A_1)) \xrightarrow{n} (A_2, I(A_2))$$

a lieu, si

1°) L'arbre A_2 s'obtient de l'arbre A_1 en construisant sur toute feuille $|a_1, b_1|$ de A_1 le feuillage

$$Cl(|a_1, b_1|, I(|a_1, b_1|), 1/n).$$

2°) Toutes les feuilles de A_1 sont rectilignes.

3°) L'arbre A_2 est un arbre simple intermédiaire entre les arbres A_1 et $I(A_1)$ (conf. J § 19).

4°) Le maximum de diamètre des feuilles de A_2 est $< 1/n$.

5°) On a

$$I(A_2) \subset I(A_1).$$

Les conditions 1° — 5° ne sont pas indépendantes. Nous les avons juxtaposées pour la commodité.

§ 3. Si A_1 est un arbre simple aux feuilles rectilignes et $I(A_1)$ l'arbre qui l'enveloppe étroitement, alors il existe A_2 et $I(A_2)$ (ch. e.) tels que

$$(A_1, I(A_1)) \xrightarrow{n} (A_2, I(A_2)).$$

D é m. a). Pour chaque feuille $|a_1, b_1|$ de A_1 il existe $Cl(|a_1, b_1|, I(a_1, b_1), \frac{1}{n})$.

En effet tout point de $|a_1, b_1|$ est intérieur à $I(|a_1, b_1|)$ par rapport au plan (car c'est assuré par J § 14, 1°), donc en raison du § 1 $Cl(\quad)$ existe.

On: a suivant la définition de Cl (§ 1, 2° — 5°)

b) Toute jonction du feuillage $Cl(|a_1, b_1|, I(a_1, b_1), 1/n)$ est contenue dans

$$:a_1, b_1:$$

c) Chaque „feuille“ $:a_1, b_1:$ est intérieure à $I(a_1, b_1)$. $\bar{b}$ et c) entraînent selon § 19 que

d) La réunion $(|a_2, b_2|)$ de toutes les $Cl(|a_1, b_1|, I(|a_1, b_1|), 1/n)$ est un feuillage intermédiaire entre les feuillages $(|a_1, b_1|)$ et $I(|a_1, b_1|)$ et par conséquent l'arbre A_2 est intermédiaire entre les arbre A_1 et $I(A_1)$

e) D'après la définition de $Cl(|a_1, b_1|, I(|a_1, b_1|), 1/n)$ toute feuille $|a_2, b_2|$ a le diamètre inférieur à $1/n$.

f) D'après c et J § 14 il existe un feuillage enveloppant le feuillage $Cl()$ dont toutes les feuilles sont contenues dans $I(a_1, b_1)$. En réunissant ces feuillages, on obtient un feuillage $(I(|a_2, b_2|))$ et il est qu'en posant

$$I(A_2) = A_1 + (I(|a_2, b_2|))$$

nous obtenons un arbre enveloppant A_2 et pour lequel

$$I(A_1) \subset I(A_2).$$

La proposition est donc démontrée.

§ 4. Voici le théorème auquel nous avons ramené la démonstration du fait que toute dendrite constitue une variété à deux dimensions.

Il existe sur le plan une dendrite développable en une suite Ω spéciale (ch. e.).

Dém. Considérons sur le plan un segment $|a_0, b_0|$ situé à l'intérieur d'un carré I et considérons l'arbre A_1 au feuillage $Cl(|a_0, b_0|, I, 1)$ (ch. e. § 1).

Désignons par $I(A_1)$ un arbre enveloppant étroitement l'arbre A_1 (ch. e. T § 14). A partir de A_1 et $I(A_1)$ nous pouvons construire par l'itération

$$(A_v, I(A_v)) \vartriangleright (A_{v+1}, I(A_{v+1}))$$

une paire de suites

$$A_1, A_2, \dots$$

$$I(A_1), I(A_2), \dots$$

(ch. e. § 3).

C'est une paire régulière de suites d'arbres (J § 21). En effet A_1 est suivant sa définition un arbre simple dont la tige $|a_0, b_0|$ est un arc simple et toutes les feuilles de A_1 sont en T par rapport à $|a_0, b_0|$ (§ 1, 1°, 2°). La propriété α de la définition J § 21 est donc vérifiée. Il en est de même des propriétés β , γ , δ de J § 21 qui sont assurées par les propriétés 3°, 4°, 5° de l'opération $\vartriangleright$ (§ 2).

La somme $\sum_{v=1}^{\infty} A_v = D^*$ constitue donc une dendrite (J § 38).

Suivant la définition de l'opération $\vee$ et d'après J § 23 la suite

$$((a_0, b_0)), ((a_1, b_1)), \dots$$

est une suite spéciale (conf. H § 17) et

$$D^* = |a_0, b_0| + \overline{\sum_{v=1}^{\infty} (|a_v, b_v|)}$$

c. à. d. la dendrite plane D^* est développable en une suite $\mathcal{Q}$ spéciale c. q. f. d.

Représentation réduite de la dendrite.

§ 1. Nous ne nous occuperons dans ce chapitre que des fonctions continues subordonnant à tout nombre d'un segment un point d'une dendrite D située à l'espace à n dimension et ne se réduisant pas à un seul point.

Nous désignerons par miniscules grecques — nombres réels, par majuscules grecques — ensembles de nombres, par minuscules latines points de D , par majuscules latines — ensembles de points de D , par $g(\mathcal{Q})$ l'ensemble de valeurs de la fonction g correspondant aux nombres de l'ensemble par $\mathcal{Q}$, par $|p_1, p_2|$ — l'arc simple faisant partie de D et ayant p_1 et p_2 comme extrémités par $[\tau_1, \tau_2]$ — segment aux extrémités τ_1 et τ_2 , c. à. d. la classe de tous les τ pour lesquels $\tau_1 \leq \tau \leq \tau_2$.

Nous appelons intervalle l'ensemble de points intérieurs d'un segment quelconque.

La dendrite D sera supposée la même au cours de ce chapitre.

§ 2. M. Denjoy a énoncé le théorème suivant:

Il existe une fonction continue $f(\tau)$ définie sur un segment $[\alpha, \beta]$, représentant la dendrite D , cofinale sur ce segment (c. à. d. telle que $f(\alpha) = f(\beta)$) et offrant une représentation réduite de cette dendrite c. à. d. satisfaisant aux conditions:

1°) Il n'existe pas quatre points

$$\tau_1 < \tau_2 < \tau_3 < \tau_4$$

tels que

$$f(\tau_1) = f(\tau_2) \neq f(\tau_3) = f(\tau_4).$$

Nous exprimerons cela en disant que ces quatre points ne s'alternent pas relativement à $f(\tau)$.

2°) La fonction $f(\tau)$ n'est constante sur aucun sous-intervalle de $[\alpha, \beta]$.

Si l'on supprime la dernière condition on obtiendra la définition de la représentation sémiréduite.

Le but de ce chapitre sera de démontrer ce théorème, d'en tirer quelques conséquences et d'indiquer comment, à partir de cette représentation, on peut effectuer l'application homéomorphe de D sur le plan ¹⁾.

1. Propriétés des fonctions représentant une dendrite.

§ 1. Si $g(\tau)$ définie sur $|\alpha, \beta|$ représente D et que

$$|\gamma, \delta| \subset |\alpha, \beta|,$$

alors tout arc simple joignant deux points de $g(|\gamma, \delta|)$ y est contenu tout entier; $g(|\gamma, \delta|)$ est aussi une dendrite.

§ 2. Si $g(\tau)$ représente D sur $|\alpha, \beta|$, est cofinale sur $|\alpha, \beta|$ et admet a comme extrémité (c. à. d. $g(\alpha) = g(\beta) = a$), si $|a, c|$ est un sous-arc simple de D , Λ la classe de tous les τ pour lesquels

$$g(\tau) \in |a, c|$$

alors Λ est fermé et g est cofinale sur chacun de ses intervalles contigus (fonction cofinale sur l'intervalle = cofinale sur le segment aux mêmes extrémités).

Dém. Soit $|\delta_1, \delta_2|$ un segment contigu de Λ . D'après le § précédent

$$|g(\delta_1), g(\delta_2)| \subset g(|\delta_1, \delta_2|).$$

D'autre part

$$|a, c| \times g(|\delta_1, \delta_2|)$$

contient au plus deux points: $g(\delta_1)$ et $g(\delta_2)$. Comme en outre

$$|g(\delta_1), g(\delta_2)| \subset |a, c|$$

done d'arc simple $|g(\delta_1), g(\delta_2)|$ contient au plus deux points c. à. d. $g(\delta_1) = g(\delta_2)$ c. q. f. d.

§ 3. Si Δ_1 et Δ_2 sont deux intervalles contigus de Λ (cf. § 2) et que $g(\tau)$ admet sur Δ_1 une autre extrémité que sur Δ_1 , alors $g(\overline{\Delta_1}) \times g(\overline{\Delta_2}) = 0$.

¹⁾ Les raisonnements qui suivent sont, dans leurs principes, dus à M. Denjoy.

Dé m. Soit

$$\Delta_1 = |\alpha_1, \beta_1|, \Delta_2 = |\alpha_2, \beta_2|$$

$$a_1 = g(\alpha_1) = g(\beta_1) \neq g(\alpha_2) = g(\beta_2) = a_2$$

Supposons par l'impossible que

$$g(\overline{\Delta_1}) \times g(\overline{\Delta_2}) \neq 0$$

Soit p^* un point, tel que

$$p^* \in g(\overline{\Delta_1}) \times g(\overline{\Delta_2}).$$

Evidemment

$$a_1 \neq p^* \neq a_2.$$

On a (cf. § 1)

$$|a_1, p^*| \subset g(\overline{\Delta_1})$$

$$|a_2, p^*| \subset g(\overline{\Delta_2})$$

d'où en multipliant ces relations par $|a, c|$ et en rappelant que Δ_1 et Δ_2 sont des intervalles contigus de l'ensemble de points représentatif τ de $|a, c|$ on obtient:

$$|a, c| \times |a_1, p^*| = a_1,$$

$$|a, c| \times |a_2, p^*| = a_2.$$

D'autre part

$$|a_1, a_2| \subset |a, c|,$$

$$a_1 \neq a_2.$$

(car a_1 et a_2 sont sur $|a, c|$).

Des quatre dernières relations il résulte [si l'on remarque que $p^* \times |a, c| = 0$] que $|a, c| + |a, p^*| + |a_2, p^*|$ contient une orbe de Jordan, ce qui est impossible.

§ 4. Toute dendrite est représentable par une fonction cofinale.

En effet, si $g(\tau)$ représente D sur $|\alpha, \beta|$, ($\alpha < \beta$) il suffit de compléter la définition de y pour l'intervalle $|\beta, \beta + (\beta - \alpha)|$ par la formule

$$g(\tau) = g(2\beta - \tau).$$

§ 5. Nous dirons qu'une fonction $g(\tau)$ est régulièrement cofinale sur un segment (ou intervalle) $|\alpha, \beta|$, si elle y est cofinale et qu'elle ne prend pour aucun τ entre α et β la valeur $g(\alpha)$. Nous avons le théorème:

L'extrémité d'une fonction régulièrement cofinale sur $|\alpha, \beta|$

et représentant une dendrite D ne découpe pas cette dendrite, c.à.d. toute couple de points de D se laisse joindre par un sous-arc simple de D et ne passant pas par e .

§ 6. Si g cofinale et définie sur $[\alpha, \beta]$ ($\alpha < \beta$) prend une valeur p^* pour un seul τ de $[\alpha, \beta]$, soit τ^* , alors p^* ne découpe pas la dendrite $g([\alpha, \beta])$.

Dém. On a $\alpha < \tau^* < \beta$. Complétons la définition de g pour les τ de l'intervalle $\beta, \beta + \tau^* - \alpha$ par la formule

$$g(\tau) = g(\tau - \beta + \alpha) \quad (\beta < \tau \leq \beta + \tau^* - \alpha).$$

$g(\tau)$ est continue, régulièrement cofinale, a p^* comme extrémité et représente D — tout cela sur le segment $[\tau^*, \beta + \tau^* - \alpha]$. Nous sommes ainsi ramenés au cas du § précédent.

Rappelons que point ne découplant pas D et point de saturation de D — c'est la même chose.

II. Alternance.

§ 1. Déf. Un ensemble Ω s'alterne relativement à une fonction $g(\tau)$ veut dire qu'il contient 4 nombres qui s'alternent relativement à g .

§ 2. Si g est définie dans $[\alpha, \beta]$ et que Π est un ensemble fermé aux extrémités α et β , si en plus $\Delta_1, \Delta_2, \dots$ sont les intervalles contigus de Π tels que

$$g(\bar{\Delta}_v) \times g([\alpha, \beta] - \bar{\Delta}_v) = 0,$$

alors du moment que Π ne s'alterne pas relativement à g , tout ensemble de 4 points qui s'alternent relativement à g est situé sur un seul Δ_v .

§ 3. On dit que $g(\tau)$ fait une représentation progressive d'un arc $[a, c]$ sur un ensemble fermé Π , progressive dans le sens $a \rightarrow c$, si

$$1^0) g(\Pi) = [a, c]$$

2⁰) lorsque croît sur Π , alors $g(\tau)$ ne recule jamais dans le sens $c \rightarrow a$.

§ 4. On dit que $g(\tau)$ fait une représentation aller et retour de $[a, c]$ sur Π , s'il existe Π_1 et Π_2 fermés telsque

$$\alpha) \Pi_1 \times \Pi_2 = 1 \text{ point} = \gamma$$

$$\beta) \text{ aucun point de } \Pi_1 \text{ ne pas à droite d'aucun point de } \Pi_2$$

$$\gamma) g(\gamma) = c$$

δ) g fait une représentation progressive dans le sens $a \rightarrow c$ de $|a, c|$ sur Π_1 et une représentation progressive dans le sens inverse de $|a, c|$ sur Π_2 .

§ 5. Si g fait une représentation aller et retour de $|a, c|$ sur un ensemble fermé Π , alors Π ne s'alterne pas relativement à Π .

III. Fonctions $g(\tau)$ et $h(\tau)$.

§ 1. Soit $g(\tau)$ une fonction cofinale sur $|\alpha, \beta|$ et représentant la dendrite D . Nous supposons cette fonction $g(\tau)$ toujours la même dans la suite. Nous désignons par a l'extrémité de g sur $|\alpha, \beta|$

$$g(\alpha) = g(\beta) = a$$

et nous définissons le point c comme il suit:

La fonction

$$\varrho[g(\alpha), g(\tau)]$$

où ϱ désigne la distance de $g(\alpha)$ à $g(\tau)$ est une fonction continue de τ . Elle atteint donc son maximum pour de certains τ . Soit τ' le plus petit d'eux (il existe évidemment). Nous posons

$$c = g(\tau').$$

Evidemment

$$\varrho(a, c) \geq \frac{1}{2} \text{ diamètre de } g(|\alpha, \beta|).$$

Nous désignerons par Λ la classe de tous les τ pour lesquels

$$g(\tau) \varepsilon |a, c|$$

et nous désignerons par $\Delta_1, \Delta_2, \dots$ les intervalles contigus de Λ .

§ 2. $g(\tau)$ est cofinale sur tout Δ_ν et, si l'extrémité de g sur Δ_ν est différente de son extrémité sur Δ_μ . ($\mu \neq \nu$), alors

$$g(\overline{\Delta}_\mu) \times g(\overline{\Delta}_\nu) = 0$$

(Cf. § 2, I et § 3 I).

§ 3. Définissons la fonction $h(\tau)$ par les conditions:

$$h(\tau) = g(\tau) \text{ sur } \Lambda$$

et sur tout segment contigu $\overline{\Delta}_\nu = |\alpha_\nu, \beta_\nu|$

$$h(\tau) = g(\alpha_\nu) = g(\beta_\nu).$$

$$\S 4. \quad h(|a, \beta|) = |a, c|.$$

$$\S 5. \quad h(\tau') = h(\tau'')$$

signifie qu'au moins un de cas se présente

$$1^0) \tau' \in \Delta, \tau'' \in \Delta, g(\tau') = g(\tau'') = h(\tau') = h(\tau''),$$

2^o) τ' est contenu dans un segment contigu $\overline{\Delta}_\nu$ sur lequel la fonction g a $g(\tau'')$ comme extrémité, τ'' est contenu dans Δ ,

3^o) la condition précédente où l'on a fait changer de place τ' et τ'' ,

4^o) ni τ' ni τ'' n'appartiennent à Δ . τ' appartient à un intervalle contigu Δ_ν aux extrémités α_ν, β_ν et τ'' à Δ_μ aux extrémités α_μ, β_μ et

$$h(\tau') = h(\tau'') = g(\alpha_\nu) = g(\beta_\nu) = g(\alpha_\mu) = g(\beta_\mu).$$

$\S 6.$ p étant un point de $|a, c|$ les classes de nombres

$$E_\tau(h(\tau) \in |a, p|),$$

$$E_\tau(h(\tau) = p),$$

$$E_\tau(h(\tau) \in |a, p|),$$

[c. à. d. les classes de tous les τ satisfaisant aux conditions marquées entre parenthèses] sont mesurables.

$$|a, p| \text{ désigne } |a, p| - p.$$

En effet les deux classes dernières sont fermées, la première est la différence de la seconde et de la dernière.

$\S 7.$ L'ensemble de points p de $|a, c|$ pour lesquels

$$\text{mesure } E_\tau(h(\tau) = p) \neq 0$$

est dénombrable ou finie, car les classes en question correspondant à deux p différents sont disjointes.

On peut ordonner ces points suivant les valeurs décroissantes de leur mesure, en cas d'ambiguïté nous affectons d'indice inférieur celui dont l'ensemble représentatif E_τ a l'extrémité droite plus à droite que les autres.

Nous désignons les points ainsi numérotés par

$$p_1, p_2, \dots$$

$$\S 8. \quad \text{mes } E_\tau(h(\tau) \in |a', c'|) > 0$$

où $[a', c']$ est un sous-arc de $[a, c]$ qui ne se réduit pas à un seul point.

Cela résulte de la continuité de la fonction $h(\tau)$.

§ 9 Si

$$(1) \quad g(\tau') = g(\tau''),$$

alors

$$(2) \quad h(\tau') = h(\tau'').$$

Dém, Il suffit de discuter les cas

$\alpha)$ $\tau' \in \Delta, \tau'' \in \Delta$

$\beta)$ $\tau' \in \Delta, \tau'' \in \Delta_v$

$\gamma)$ $\tau' \in \Delta_\mu, \tau'' \in \Delta_v$.

Dans le premier cas (2) a lieu suivant la définition de h .

Le cas β ne peut pas se présenter, car $g(\Delta) \times g(\Delta) = 0$ et par suite (1) est impossible.

En posant dans le troisième cas

$$\overline{\Delta}_\mu = [\alpha_\mu, \beta_\mu], \quad \overline{\Delta}_v = [\alpha_v, \beta_v],$$

nous avons selon le § 3. I:

$$g(\alpha_\mu) = g(\beta_\mu) = g(\alpha_v) = g(\beta_v)$$

donc suivant la définition de h et selon γ (2) est vérifié.

§ 10. Soit p un point de l'arc $[a, c]$.

Nous désignerons par

$$\underline{\omega}(p) \text{ et } \overline{\omega}(p)$$

respectivement la plus petite et la plus grande valeur de τ pour laquelle

$$h(\tau) = p,$$

h étant continue, $\underline{\omega}$ et $\overline{\omega}$ existent. On a

$$h(\underline{\omega}(p)) = h(\overline{\omega}(p)) = p,$$

§ 11. $\underline{\omega}(p)$ et $\overline{\omega}(p)$ font partie de Δ .

Démontrons p. ex. que c'est le cas de $\underline{\omega}(p)$. La fonction h est constante sur tout segment $\overline{\Delta}_v$ contigu de $\overline{\Delta}$, donc le minimum de τ satisfaisant à la relation

$$h(\tau) = p$$

ne peut pas être atteint dans $\Sigma \Delta_v$, il l'est donc sur Δ .

§ 12. $\underline{\omega}(p)$ et $\overline{\omega}(p)$ donnent les extrema des τ pour lesquels

$$g(\tau) = p \quad (p \in |a, c|).$$

En effet $\underline{\omega}(p)$ faisant partie de Λ on a

$$g(\underline{\omega}(p)) = h(\underline{\omega}(p));$$

d'autre comme pour

$$\tau < \underline{\omega}(p)$$

$$h(\tau) \neq h(\underline{\omega}(p)) = p,$$

donc (§ 9)

$$g(\tau) \neq h(\underline{\omega}(p)).$$

Le cas de $\overline{\omega}$ se traite pareillement.

IV. La fonction $\varphi(\tau)$.

§ 1. Nous posons:

$$\begin{aligned} \varphi(\tau^*) &= \text{mesure } E_\tau(h(\tau) \in |a, h(\tau^*)|) + \\ &+ \text{mes}([E_\tau(h(\tau) = h(\tau^*))] \times |a, \tau^*|) + \alpha \end{aligned}$$

De cette définition résultent immédiatement des propriétés différentes:

$$\text{§ 2.} \quad \varphi(\alpha) = \varphi(\underline{\omega}[h(\alpha)]) = \varphi(\underline{\omega}(\alpha)) = \alpha$$

$$\varphi[\overline{\omega}(c)] = \beta.$$

$$\text{§ 3.} \quad \alpha \leq \varphi(\tau) \leq \beta$$

pour tous les τ pour lesquels $\varphi(\tau)$ existe.

$$\text{§ 4.} \quad \varphi(\underline{\omega}(p)) = \text{mes } E_\tau(h(\tau) \in |a, p|) + \alpha,$$

$$\varphi(\overline{\omega}(p)) = \text{mes } E_\tau(h(\tau) \in |a, p|) + \alpha$$

pour tout p de $|a, c|$.

$$\text{§ 5.} \quad \varphi(\overline{\omega}(p)) - \varphi(\underline{\omega}(p)) = \text{mes } E_\tau(h(\tau) = p) \geq 0.$$

§ 6. Nous désignons par $\Theta(p)$ le segment $[\varphi(\underline{\omega}(p)), \varphi(\overline{\omega}(p))]$,

§ 7. Trois faits suivants sont équivalents

α) existe τ tel que

$$h(\tau) = p,$$

$$\varphi(\tau) = \psi.$$

$$\beta) \varphi(\underline{\omega}(p)) \leq \psi \leq \varphi(\overline{\omega}(p)).$$

$$\gamma) \psi \in \overline{\Theta}(p).$$

§ 8. Rappelons que l'on a désigné par

$$p_1, p_2, \dots$$

la suite de tous les points p pour lesquels

$$\text{mes } E_\tau(h(\tau) = p) \neq 0. \quad (p \in |a, c|).$$

On a donc pour p satisfaisant à

$$p \in |a, c|,$$

la relation

$$p \neq p_v,$$

$$\varphi(\underline{\omega}(p)) = \varphi(\overline{\omega}(p)) \quad (\text{cf. § 5}).$$

§ 9. Soit Δ_v un contigu de Δ ; l'équation $\psi = \varphi(\tau)$ ($\tau \in \overline{\Delta}_v$) détermine une translation de $\overline{\Delta}_v$ à la position finale $\varphi(\overline{\Delta}_v)$, car $h(\tau)$ est constante sur $\overline{\Delta}_v$ (cf. les définitions de h et φ).

§ 10. Si

$$h(\tau') = p',$$

$$h(\tau'') = p'',$$

$$p' < p'' \quad (\text{c. à. d. } p' \text{ et } p'')$$

sans coïncider se succèdent sur $|a, c|$ dans le même ordre que a et c), alors

$$\varphi(\tau') < \varphi(\tau'').$$

Corrol. 1. Si

$$\varphi(\tau') = \varphi(\tau''),$$

alors

$$h(\tau') = h(\tau'').$$

Corrol. 2 Si

$$\varphi(\tau') \leq \varphi(\tau''),$$

alors

$$h(\tau') \leq h(\tau'').$$

Corrol. 3. Si

$$p' < p'',$$

alors le segment

$$|\varphi(\underline{\omega}(p')), \varphi(\overline{\omega}(p'))| = \overline{\Theta}(p')$$

est situé à gauche du segment $\overline{\Theta}(p'')$ et ces segments sont disjoints.

§ 11.

$$\lim_{p \rightarrow p^*} \varphi(\overline{\omega}(p)) = \varphi(\underline{\omega}(p^*)).$$

$$p \rightarrow p^*$$

$$p < p^*.$$

Dém. On a pour toute suite de points

$$l_1, l_2, l_3, \dots$$

telle que $l_v \in]a, c[$, $l_v < l_{v+1}$, $\lim l_v = p^*$, les relations

$$\varphi(\overline{\omega}(l_v)) - \alpha = \text{mes } E_\tau(h(\tau) \varepsilon | a, l_v|), \quad E_\tau(h(\tau) \varepsilon | a, l_v| \subset E_\tau(h(\tau) \varepsilon | a, l_{v+1}|),$$

donc suivant un théorème sur la suite des ensembles mesurales dont chacun est cotenu dans le suivant

$$\lim \varphi(\overline{\omega}(l_v)) - \alpha = \text{mes } \bigcup_{v/1}^\infty E_\tau(h(\tau) \varepsilon | a, l_v|]$$

c. à. d.

$$\lim \varphi(\overline{\omega}(l_v)) - \alpha = \text{mes } E_\tau(h(\tau) \varepsilon | a, p^*|)$$

d'où selon le § 4

$$\lim \varphi(\overline{\omega}(l_v)) = \varphi(\overline{\omega}(p^*)), \quad \text{c. q. f. d.}$$

§ 12. Si

$$p^* < p,$$

$$\lim p = p^*,$$

alors

$$\lim \varphi(\overline{\omega}(p)) = \varphi(\overline{\omega}(p^*)),$$

c. à. d. la fonction $\varphi(\overline{\omega}(p))$ est continue „à droite“ de tout point de $]a, c[$.

Dém. Soit

$$p^* < l_{v+1} < l_v \quad (v/1, \dots)$$

$$\lim l_v = p^*.$$

On a

$$\varphi(\overline{\omega}(l_v)) - \alpha = \text{mes } E(h(\tau) \varepsilon | a, l_v|) = \text{mes } E_v.$$

Comme

$$E_{v+1} \subset E_v = \overline{E_v}, \text{ donc}$$

$$\lim \text{mes } E_v = \text{mes } \bigcap_{v/1}^\infty E_v = \text{mes } E_\tau(h(\tau) \varepsilon | a, p^*|)$$

donc (§ 4)

$$\lim \varphi(\overline{\omega}(l_v)) = \varphi(\overline{\omega}(p^*)) \quad \text{c. q. f. d.}$$

§ 13. Si

$$\alpha \leq \psi \leq \beta,$$

alors il existe un seul point p de $]a, c[$, tel que

$$(1) \quad \psi \in \overline{\Theta}(p).$$

Dém. Deux segments $\Theta(p)$ étant disjoints (§ 10, corr. 3), il

existe au plus un segment satisfaisant à (1). Il reste à démontrer qu'il en existe au moins un. Si

$$\alpha \leq \psi \leq \overline{\omega}(\alpha),$$

alors (§ 2)

$$\psi \in \overline{\Theta}(\alpha).$$

Supposons que

$$(2) \quad \varphi(\overline{\omega}(\alpha)) < \psi \leq \beta = \varphi(\overline{\omega}(c)).$$

Soit p^* la „limite supérieure“ sur $|\alpha, c|$ des points p pour lesquels

$$\varphi(\overline{\omega}(p)) < \psi.$$

En rappelant que $\varphi(\overline{\omega}(p))$ est une fonction croissante de p qui avance sur $|\alpha, c|$ dans le sens $\alpha \rightarrow c$ (cf. § 9 corr. 3) on obtient selon le § 11

$$(3) \quad \varphi(\underline{\omega}(p^*)) \leq \psi.$$

Je dis que

$$(4) \quad \psi \leq \varphi(\overline{\omega}(p^*)).$$

En effet dans le cas contraire on aurait

$$(5) \quad \varphi(\overline{\omega}(p^*)) < \psi$$

donc selon (2) $p^* < c$.

La fonction $\varphi(\omega(p))$ étant continue à droite en tout point de $|\alpha, c|$ (§ 12), l'inégalité (5) serait vérifiée pour de certains points de l'arc $:p^*, c|$ contrairement à la définition de p^* . (4) est donc vrai. (3) et (4) entraînent (1).

§ 14. On a

$$|\alpha, \beta| = (\overline{\Theta}(p))$$

où () signifie la somme de tous les $\overline{\Theta}(p)$ que l'on obtient si p parcourt $|\alpha, c|$.

On le voit en rapprochant les §§ 3 et 13. De là:

$$\S 15. \quad \varphi(|\alpha, \beta|) = |\alpha, \beta|.$$

§ 16. Soit $p \in |\alpha, c|$. La fonction $\varphi(\tau)$ est non décroissante sur l'ensemble E_τ ($h(\tau) = p$). Elle est croissante et continue sur tous segment $\overline{\Delta}_v$ contigu de Δ . Cela résulte de la définition de $\varphi(\tau)$.

§ 17. Si un segment $\overline{\Delta}_v$ contigu de Δ empiète sur E_τ ($h(\tau) = p$), alors $\overline{\Delta}_v$ est contenu dans cet ensemble. En effet $h(\tau)$ est constante sur $\overline{\Delta}_v$.

§ 18. $\varphi(\Sigma\Delta_\nu) \times \varphi(\Delta) = 0$.

Dém. Il suffit de démontrer que pour tout Δ_ν on a

$$\varphi(\Delta_\nu) \times \varphi(\Delta) = 0$$

et pour cela il est suffisant de prouver que $\tau' \in \Delta$ entraîne

$$(1) \quad \varphi(\tau') \times \varphi(\Delta_\nu) = 0.$$

On a

$$h(\Delta_\nu) = p'',$$

$$h(\tau') = p',$$

donc d'après le § 7

$$\varphi(\Delta_\nu) \subset \overline{\Theta}(p''),$$

$$\varphi(\tau') \in \overline{\Theta}(p').$$

Si $p' \neq p''$, alors en multipliant les relations précédentes l'une par l'autre on obtient (1) selon le § 10 cor. 3. Dans le cas $p' = p''$ (1) est vrai aussi, car $\varphi(\tau)$ est une fonction non décroissante sur $E_\tau(h(\tau) = p)$ et croissante sur Δ_ν (cf. § 16).

Corrol. $\varphi(\Delta)$ est complémentaire de $\varphi(\Sigma\Delta_\nu)$ relativement à $[\alpha, \beta]$ (cf. § 15).

§ 19. Si $\mu \neq \nu$, on a

$$\varphi(\Delta_\mu) \times \varphi(\Delta_\nu) = 0.$$

On le prouve d'une façon analogue à celle du § précédent.

§ 20. Il existe un seul τ' satisfaisant aux relations

$$(1) \quad \begin{aligned} \psi &\in \varphi(\Delta_\nu) \\ \psi &= \varphi(\tau') \end{aligned}$$

où τ' est considéré comme inconnu. De plus pour ce τ' on a

$$(2) \quad \tau' \in \Delta_\nu.$$

Dém. Nous allons démontrer que toute solution de (1) satisfait à (2).

En effet $\tau' \in \Delta_\mu$, où $\mu \neq \nu$, est exclu d'après le § 19 et $\tau' \in \Delta$ est exclu selon le § 18, donc (2) a lieu pour tous les τ' satisfaisant à (1).

Il existe au plus un τ' qui satisfait à (1), car s'il existait un autre soit τ'' , on aurait $\tau'' \in \Delta_\nu$ et, comme $\varphi(\tau)$ est croissante sur Δ_ν (§ 16), on aurait $\varphi(\tau'') \neq \varphi(\tau') = \psi$, τ'' ne satisfairait donc pas à (1).

Il existe, enfin, au moins un τ' de cette sorte suivant l'égalité

$$[\alpha, \beta] = \varphi([\alpha, \beta]) \quad (\S 15).$$

§ 21, Si $\psi \in \varphi(\Lambda)$
 alors $\psi = \varphi(\tau)$,
 $\tau \in \Lambda$.

Cela découle du § précédent.

§ 22. $\psi = \varphi(\tau)$ établit une homéomorphie entre $\Sigma \Delta_v$ et $\varphi(\Sigma \Delta_v)$.

Dém. La correspondance en question est biunivoque (§ 20). Elle est bicontinue (§ 16).

§ 23. 1° $\varphi(\Lambda)$ est fermé

2° La totalité des intervalles contigus de $\varphi(\Lambda)$ et formée par tous les $\varphi(\Delta_v)$.

Dém. $\varphi(\Sigma \Delta_v)$ est un ensemble ouvert comme étant homéomorphe de $\Sigma \Delta_v$; $\varphi(\Lambda)$ est donc fermé comme complémentaire de $\varphi(\Sigma \Delta_v)$ par rapport à $|\alpha, \beta|$ (§ 18 corr.)

2°) Selon le § 18 $\varphi(\Delta_v) \times \varphi(\Lambda) = 0$.

D'autre part comme les extrémités α_v et β_v de Δ_v font partie de Λ et que $\varphi(\tau)$ effectue une translation sur $\overline{\Delta_v}$ (§ .), donc $\varphi(\alpha_v)$ et $\varphi(\beta_v)$ font partie de $\varphi(\Lambda)$ et sont extrémités de $\varphi(\Delta_v)$. Tout $\varphi(\Delta_v)$ est donc un intervalle contigu de $\varphi(\Lambda)$. Il n'y a pas d'autres intervalles contigus, car $\varphi(\Lambda)$ est complémentaire de $\varphi(\Sigma \Delta_v)$ relativement à $|\alpha, \beta|$ (§ 18, corr.).

§ 24 La classe

$$II = |\alpha, \beta| - \Sigma \Theta_v,$$

où Θ_v désigne l'intervalle aux extrémités $\varphi(\underline{\omega}(p_v)), \varphi(\overline{\omega}(p_v))$, constitue un ensemble fermé contenant α et β , admettant les intervalles Θ_v comme intervalles contigus et n'ayant pas d'autres intervalles contigus.

On a

$$(1) \quad II = (\varphi(\underline{\omega}(p))) + (\varphi(\overline{\omega}(p))),$$

où () désigne la classe de tous les $\varphi(\underline{\omega}(p))$ et $\varphi(\overline{\omega}(p))$ que l'on obtient, quand p parcourt l'arc $|\alpha, \beta|$.

Dém. II est fermé comme complémentaire de l'ensemble ouvert $\Sigma \Theta_v$.

Les Θ_v étant deux à deux disjoints, forment bien les intervalles contigus de II .

Enfin l'égalité (1) est vraie, car, selon le § 14, on a

$$|\alpha, \beta| = (\overline{\Theta}(p))$$

et que d'autre part tous les $\overline{\theta}(p)$ à l'exception des $\overline{\theta}(p_v)$ se réduisent à un seul point (Cf. § 8).

α et β font partie de Π d'après (1), car $\varphi(\underline{\omega}(a)) = \alpha$, $\varphi(\underline{\omega}(c)) = \beta$ (§ 2).

§ 25. $\Pi = |\alpha, \beta| - \Sigma \theta_v \subset \varphi(\Lambda)$.

En effet, les $\underline{\omega}(p)$ et $\overline{\omega}(p)$ font partie de Λ (III, § 11), donc les $\varphi(\underline{\omega}(p))$ et $\varphi(\overline{\omega}(p))$ font partie de $\varphi(\Lambda)$ et, d'après la formule (1) du § précédent, il en est de même de Π .

V. Fonction $k(\psi)$.

§ 1. Nous définissons la fonction

$$p = k(\psi)$$

par la condition: p est subordonné à ψ veut dire qu'il existe un τ , tel que

$$p = h(\tau)$$

$$\psi = \varphi(\tau).$$

§ 2. La fonction $k(\psi)$ ainsi définie est uniforme, définie sur le segment $|\alpha, \beta|$ et là seulement, elle fait une représentation continue et progressive de l'arc $|a, c|$ de façon que

$$k(\alpha) = a, k(\beta) = c.$$

Dém. α) $k(\psi)$ est une fonction uniforme. En effet, dans le cas contraire il existerait τ_1 et τ_2 tels que

$$p_1 = h(\tau_1) \neq h(\tau_2) = p_2,$$

$$\psi = \varphi(\tau_1) = \varphi(\tau_2),$$

ce qui est impossible selon IV, § 10, corr. 1.

β) $k(\psi)$ est défini sur $|\alpha, \beta|$ et là seulement et représente l'arc $|a, c|$.

Pour le voir, il suffit de rappeler que $\varphi(|\alpha, \beta|) = |\alpha, \beta|$ (IV § 10) et d'observer que $h(\tau)$ représente l'arc $|a, c|$ sur $|\alpha, \beta|$ (III, § 4).

γ) le ψ fait une représentation progressive de $|a, c|$ sur $|\alpha, \beta|$ (Cf. II § 3) et on a $k(\alpha) = a$, $k(\beta) = c$.

Les dernières égalités sont vraies, car

$$a = \varphi(\alpha). \quad \beta = \varphi(\overline{\omega}(c)),$$

$$a = h(\alpha) \quad c = h(\overline{\omega}(c)).$$

(Cf. IV § 2, § III § 10).

δ) $k(\psi)$ est une fonction continue sur $[\alpha, \beta]$.

Dém. Soit

$$(1) \quad p^* = k(\psi^*)$$

et soit $\varepsilon > 0$.

Nous allons, d'abord, traiter le cas:

$$a < p^* < c.$$

Selon (1) il existe un τ^* tel que

$$p^* = h(\tau^*),$$

$$\psi^* = \varphi(\tau^*).$$

Considérons un sous-arc simple $[p', p'']$ de $[a, c]$, tel que

$$(2) \quad \begin{aligned} p' &< p^* < p'', \\ \text{diamètre de } [p', p''] &< \varepsilon. \end{aligned}$$

Il existe τ' et τ'' , tels que

$$p' = h(\tau'),$$

$$p'' = h(\tau'').$$

Posons

$$\psi' = \varphi(\tau'),$$

$$\psi'' = \varphi(\tau'').$$

De (2) on a selon le § 10, IV

$$\varphi(\tau') < \varphi(\tau^*) < \varphi(\tau'').$$

Posons δ égal à la plus petite des différences

$$\varphi(\tau'') - \varphi(\tau'), \quad \varphi(\tau^*) - \varphi(\tau').$$

Nous allons démontrer que

$$(3) \quad |\psi - \psi^*| \leq \delta$$

entraîne

$$(4) \quad \text{distance}(k(\psi), k(\psi^*)) \leq \varepsilon.$$

D'après (3) et la définition de δ

$$\varphi(\tau') \leq \psi \leq \varphi(\tau'').$$

Il existe un τ pour lequel

$$(5) \quad \psi = \varphi(\tau)$$

et on a

$$\varphi(\tau') \leq \varphi(\tau) \leq \varphi(\tau'')$$

d'où selon IV, § 10, corr. 2

$$h(\tau') \leq h(\tau) \leq h(\tau''),$$

c. à d.

$$p' \leq h(\tau) \leq p''$$

$h(\tau)$ est donc situé sur l'arc $[p', p'']$ donc selon (2)

$$\text{distance } (h(\tau), p^*) \leq \varepsilon$$

est comme, d'après (5),

$$h(\tau) = k(\psi)$$

par conséquent

$$\text{dist } (k(\psi), k(\psi^*)) \leq \varepsilon.$$

(3) entraîne donc (4) c. q. f. d.

§ 3. Pour ψ situé sur $\overline{\Theta}(p^*)$ (c. à d. sur $[\varphi(\underline{\omega}(p^*)), \varphi(\overline{\omega}(p^*))]$) on a

$$k(\psi) = p^* = \text{const.}$$

Dém. Comme

$$\varphi(\underline{\omega}(p^*)) \leq \psi \leq \varphi(\overline{\omega}(p^*)),$$

donc, selon IV, § 7 il existe un τ , tel que

$$h(\tau) = p^*,$$

$$\varphi(\tau) = \psi$$

et de là (§ 1)

$$k(\psi) = p^*, \quad \text{c. q. f. d.}$$

§ 4. Si $p' \neq p''$, $p' \in]a, c[$, $p'' \in]a, c[$, alors

$$k(\overline{\Theta}(p')) \times k(\overline{\Theta}(p'')) = 0.$$

En particulier pour $\mu \neq \nu$

$$k(\overline{\Theta}_\nu) \times k(\overline{\Theta}_\mu) = 0 \quad (\text{Cf. IV, § 24}).$$

C'est une conséquence du § précéd.

§ 5. $k(\overline{\Theta}(p)) \times k([\alpha, \beta] - \overline{\Theta}(p)) = 0.$

Dém. Supposons que ce n'est pas vrai. Il existe donc ψ' et ψ'' , tels que

$$(1) \quad \psi' \in \overline{\Theta}(p),$$

$$(2) \quad \psi'' \in [\alpha, \beta] - \overline{\Theta}(p),$$

$$(3) \quad k(\psi') = k(\psi'').$$

Or ψ'' comme élément de $[\alpha, \beta]$ est contenu dans un seul $\overline{\Theta}(p'')$.

(Cf. IV, § 13)

$$(4) \quad \psi'' \in \overline{\Theta}(p'').$$

D'après (2) $\overline{\Theta}(p'')$ n'est pas identique avec $\overline{\Theta}(p)$ donc $p \neq p''$ et par conséquent (IV, § 10, corr. 3)

$$(5) \quad k(\overline{\Theta}(p)) \times k(\overline{\Theta}(p'')) = 0$$

D'autre part d'après (1), (3), (4)

$$k(\psi') = k(\psi'') \subset k(\overline{\Theta}(p)) \times k(\overline{\Theta}(p''))$$

contrairement à (5).

§ 6. La fonction k est cofinale sur tout Θ_v et fait une représentation continue et progressive de l'arc $|a, c|$ sur l'ensemble fermé

$$II = |a, \beta| - \Sigma \Theta_v \quad \text{c. à. d.}$$

$$(1) \quad k(II) = |a, c|.$$

On a

$$\alpha \in II, \beta \in II, k(\alpha) = a, k(\beta) = c.$$

Dém. En tenant compte des résultats de IV, § 24 et du § 2 de ce chapitre on voit qu'il suffit de démontrer (1).

On a (IV § 24)

$$II = (\varphi(\underline{\omega}(p))) + (\varphi(\overline{\omega}(p)))$$

où () sont les classes de φ que l'on obtient, si p parcourt l'arc $[a, c]$

On a donc

$$k(II) = (k[\varphi(\underline{\omega}(p))]) + (k[\varphi(\overline{\omega}(p))])$$

et selon le § 3

$$k(II) = (p) + (p) = (p) = |a, c|. \quad \text{c. q. f. d.}$$

VI. Fonction $l(\psi)$.

§ 1. Nous définissons la fonction

$$p = l(\psi)$$

par la condition: p est subordonné à ψ veut dire qu'il existe un τ , tel que

$$p = g(\tau),$$

$$\psi = \varphi(\tau).$$

(Cf III § 1, IV § 1).

§ 2. $l(\psi)$ est une fonction uniforme.

Dém. Supposons, par l'im possible, qu'il existe τ' et τ'' tels que

$$(1) \quad \tau' \neq \tau'',$$

$$p' = g(\tau') \neq g(\tau'') = p,$$

$$(2) \quad \psi' = \varphi(\tau') = \varphi(\tau'').$$

Trois cas sont possibles (Cf. IV § 18 corr.).

1°) τ' et τ'' appartiennent à Λ

2°) τ' et τ'' appartiennent à $\Sigma\Delta_v$

3°) un d'eux p. ex. τ' appartient à Λ l'autre à $\Sigma\Delta_v$.

Dans le cas 1°) on a

$$g(\tau') = h(\tau'),$$

$$g(\tau') = h(\tau'')$$

donc

$$p_1 = h(\tau') \neq h(\tau'') = p_2$$

d'où selon (2) $p_1 = k(\psi'_1)$, $p_2 = k(\psi'_2)$, ce qui est impossible, car $k(\psi)$ est une fonction uniforme (V, § 2).

Dans les cas 2°) et 3°), l'inégalité (1) donne selon IV § 18 et IV § 23

$$\varphi(\tau') \neq \varphi(\tau'')$$

contrairement à (2).

§ 3. La fonction $l(\psi)$ est définie sur $|\alpha, \beta|$ et là seulement et elle représente la dendrite D , c. à. d. $l(|\alpha, \beta|) = D$.

Pour le voir, il suffit, en tenant compte de la définition de $l(\psi)$, de remarquer que selon IV, § 15

$$\varphi(|\alpha, \beta|) = |\alpha, \beta|.$$

§ 4. L'ensemble de points représentatifs ψ de $|a, c|$ c. à. d. tels que

$$l(\psi) \varepsilon |a, c|$$

est identique à $\varphi(\Lambda)$ et $k(\varphi(\Lambda)) = |a, c|$. Sur $\varphi(\Lambda)$

on a:

$$k(\psi) = l(\psi).$$

Dém. La première partie du théorème est vraie, car $\Lambda = E_\tau [g(\tau) \varepsilon |a, c|]$, la seconde découle du fait que sur Λ on a

$$g(\tau) = h(\tau).$$

Corrol. 1.

$$l(\psi) \times |a, c| = l(\psi) \times k(\varphi(\Lambda)) = k(\psi \times \varphi(\Lambda))$$

Corrol 2.

$k(\Pi) = (|a, c|)$, (V § 6), $\Pi \subset \varphi(\Lambda)$ (IV, § 25), donc $l(\Pi) = |a, c|$.

§ 5. $l(\psi)$ fait une représentation progressive de $|a, c|$ sur $\varphi(\Lambda)$.

On a

$$l(\alpha) = a, l(\beta) = c.$$

Dém. En effet sur $\varphi(\Lambda)$ k et l sont identiques (§ précéd.)

et k fait une représentation progressive de $|a, c|$ sur $\varphi(\Delta)$ (V § 2; VI, § 4).

§ 6 $l(\psi)$ est une fonction continue.

Dém. $l(\psi)$ est continue sur $\varphi(\Delta)$ comme identique sur cet ensemble à $k(\psi)$.

Elle est continue sur tout segment $\varphi(\overline{\Delta})$, c. à. d. sur tout segment contigu de $\varphi(\Delta)$ (Df. IV, § 16, corr.), car $k(\psi)$ s'obtient de $g(\tau)$ par le changement de variables $\psi = \varphi(\tau)$ qui appliqué à $\overline{\Delta}$, revient à sa translation à la position nouvelle $\varphi(\overline{\Delta}_v)$. Pour conclure de là à la continuité, il suffit de remarquer que le diamètre de $k(\varphi(\overline{\Delta}_v))$ tend vers 0 avec $\frac{1}{v}$ (si Δ_v sont en nombre infini). Or il en est ainsi du diamètre de $g(\overline{\Delta}_v)$ qui est égal au précédent (translation).

De là:

$$\S 7. \quad l(\varphi(\Delta_v)) = g(\Delta_v)$$

$$l(\varphi(\overline{\Delta}_v)) = g(\overline{\Delta}_v).$$

§ 8.

$$(1) \quad \text{Si } k(\psi_1) \neq k(\psi_2),$$

alors

$$l(\psi_1) \neq l(\psi_2).$$

Dém. Il existe τ_1 et τ_2 tels que

$$(2) \quad \psi_1 = \varphi(\tau_1),$$

$$\psi_2 = \varphi(\tau_2),$$

$$k(\psi_1) = h(\tau_1),$$

$$k(\psi_2) = h(\tau_2)$$

donc selon (1)

$$h(\tau_1) \neq h(\tau_2)$$

d'où suivant III, § 9

$$g(\tau_1) \neq g(\tau_2)$$

ce qui donne avec (2)

$$l(\psi_1) \neq l(\psi_2). \quad \text{c. q. f. d.}$$

Corrol. $\mathcal{W}_1$ et $\mathcal{W}_2$ étant deux sous-ensembles de $|a, \beta|$, l'égalité

$$k(\mathcal{W}_1) \times k(\mathcal{W}_2) = 0$$

entraîne

$$l(\mathcal{W}_1) \times l(\mathcal{W}_2) = 0.$$

§ 9. $p' \neq p''; p' \varepsilon |a, c|; p'' \varepsilon |a, c|$ entraîne

$$l(\overline{\Theta}(p')) \times l(\overline{\Theta}(p'')) = 0$$

(Corrolaire précédent et V, § 4).

§ 10. Pour $p \in |a, c|$ on a

$$k(\overline{\Theta}(p) \times k(|a, \beta| - \overline{\Theta}(p))) = 0.$$

(Corrol du § 8 et V, § 5).

§ 11. $l(\psi)$ est cotinale sur $\overline{\Theta}_v$ et à p_v comme extrémité.

Dém. $k(\psi)$ est constante sur $\overline{\Theta}_v$ et égale à p_v (V, § 3). Les extrémités de $\overline{\Theta}_v$ appartiennent à $\varphi(\Lambda)$ (IV, §§ 24 et 25) donc pour ces extrémités l prend la même valeur que k (Cf. § 4) c. à. d. p_v .

§ 12.

$$(1) \quad |a, c| \times l(\overline{\Theta}(p)) = p$$

et en particulier

$$|a, c| \times l(\overline{\Theta}_v) = p_v.$$

Dém. On a selon le § 4

$$l(\overline{\Theta}(p)) \times |a, c| = l[\overline{\Theta}(p) \times \varphi(\Lambda)] =$$

$$(2) \quad = k[\overline{\Theta}(p) \times \varphi(\Lambda)] \subset k(\overline{\Theta}(p)) = p.$$

(cf. V § 3).

D'autre part, comme $\varphi(\underline{\omega}(p))$ appartient à Λ (IV, § 24) on a selon le § 4 et V § 3.

$$l(\varphi(\underline{\omega}(p))) = k(\varphi(\underline{\omega}(p))) = p$$

donc

$$(3) \quad p \in l(\overline{\Theta}(p))$$

(2) et (3) entraînent (1).

§ 13. Définissons la fonction $g^*(\tau)$ comme il suit:

$$\text{Sur} \quad \left| a, \frac{\alpha + \beta}{2} \right|$$

$$g^*(\tau) = l(\alpha + 2(\tau - \alpha))$$

$$\text{et sur} \quad \left| \frac{\alpha + \beta}{2}, \beta \right|$$

$$g^*(\tau) = k(\alpha + 2(\beta - \tau)).$$

Remarquons que, si τ croît de α à $\frac{\alpha + \beta}{2}$, $\alpha + 2(\tau - \alpha)$ croît de α à β et, si τ croît de $\frac{\alpha + \beta}{2}$ à β , $\alpha + 2(\beta - \tau)$ décroît de β à α .

La fonction $g^*(\tau)$ est uniforme et continue sur $|a, \beta|$, car l et k le sont et $l(\beta) = k(\beta) = c$ (§ 5, § 6 et V, § 2).

Sur $\left| \frac{\alpha + \beta}{2}, \beta \right|$ g^* fait une représentation progressive de $|a, c|$ dans le sens $c \rightarrow a$ (cf. V § 2).

Donnons aux symboles II et Θ_v une autre signification ce qui ne causera pas de confusion. Désignons par Θ_v ce que devient Θ_v d'auaravant par le changement de variables.

$$\tau' = \alpha + 2(\tau - \alpha)$$

et désignons par II le segment $\left| \frac{\alpha + \beta}{2}, \beta \right|$ augmenté de ce que devient le II ancien par la transformation précédente. $g^*(\tau)$ étant ainsi définie nous pouvons énoncer le théorème:

§ 14. A toute fonction $g(\tau)$ définie et continue sur le segment $[\alpha, \beta]$, cofinale sur ce segment $[\alpha, \beta]$, ayant a comme une extrémité et représentant une dendrite D , il appartient un arc $|a, c|$, une fonction $g^*(\tau)$ et un ensemble II , parfaitement déterminés jouissant des propriétés:

1°) diamètre $|a, c| \geq \frac{1}{2} \text{diam } D = \frac{1}{2} \text{diam } g([\alpha, \beta])$. (cf. III § 1).

2°) $g^*(\tau)$ est continue et définie sur le même segment $[\alpha, \beta]$, représente sur ce segment la même dendrite D , est cofinale sur $[\alpha, \beta]$ et admet le même point a comme extrémité. (§ 13; VI, § 3; V, § 6; VI, § 5).

3°) II est un sous-ensemble fermé de $[\alpha, \beta]$, contenant α et β et ayant de certains intervalles $\Theta_1, \Theta_2, \dots$ comme intervalles contigus (§ 13; IV, § 24).

4°) II ne s'alterne pas relativement à $g^*(\tau)$ (car $g^*(\tau)$ fait une représentation aller et retour de $|a, c|$ sur II (cf. § 13; VI, § 4, corr II; VI, § 5; V § 2).

$$5_0) \quad g^*(\overline{\Theta}_v) \times g^*([\alpha, \beta] - \overline{\Theta}_v) = 0.$$

$$g^*(\overline{\Theta}_v) \times g^*(\overline{\Theta}_\mu) = 0.$$

(§ 13; VI, §§ 9 et 10).

6°) $g^*(\tau)$ est cofinale sur Θ_v et admet un point p_v comme extrémité. (§§ 13 et 11).

$$7_0) \quad g^*(\overline{\Theta}_v) \times g^*(II) = g^*(\overline{\Theta}_v) \times |a, c| = p_v$$

(cf. 4°) ce § et le § 12).

Nous désignons par $P(g(\tau), [\alpha, \beta])$ l'opération donnant, à partir de la fonction $g(\tau)$, l'arc $|a, c|$, la fonction $g^*(\tau)$ et l'ensemble II .

VII. Construction d'une représentation cofinale et sémiré- duite d'une dendrite.

§ 1. A partir d'une fonction $g(\tau)$ définie sur $|\alpha, \beta|$, cofinale et admettant un point a comme extrémité et représentant une dendrite D , construisons une suite de fonctions

$$g_1(\tau), g_2(\tau), \dots,$$

une suite d'ensembles

$$II_1, II_2, \dots,$$

une suite double d'arcs simples

$$(S_1) \quad |a_1^1, c_1^1|, |a_1^2, c_1^2|, \dots,$$

$$(S_2) \quad |a_2^1, c_2^1|, |a_2^2, c_2^2|, \dots,$$

$$\begin{array}{c} \cdot \cdot \cdot \cdot \cdot \cdot \cdot \\ \cdot \cdot \cdot \cdot \cdot \cdot \cdot \end{array},$$

une suite double d'intervalles

$$(\Omega_1) \quad \Theta_1^1, \Theta_1^2, \dots,$$

$$(\Omega_2) \quad \Theta_2^1, \Theta_2^2, \dots,$$

$$\begin{array}{c} \cdot \cdot \cdot \cdot \cdot \cdot \cdot \\ \cdot \cdot \cdot \cdot \cdot \cdot \cdot \end{array},$$

de la façon suivante:

$g_1(\tau)$, II_1 , $|a_1^1, c_1^1|$ sont déduits de g par l'opération $P(g(\tau), |\alpha, \beta|)$ (cf. VI, § 14).

La suite S_1 se réduit, par définition, à son premier terme. Les intervalles Θ_1^v sont les intervalles contigus de II_1 ; a_1^v est l'extrémité de g_1 sur Θ_1^v .

Supposons que l'on a définie g_v , II_v , (S_v) , et (Ω_v)

Nous posons

$$g_{v+1}(\tau) = g_v(\tau) \text{ sur } II_v.$$

Nous désignons par

$$g_{v+1}^\mu, II_{v+1}^\mu, |a_{v+1}^\mu, c_{v+1}^\mu|$$

$$(\Omega_{v+1}^\mu) \quad \Theta_{v+1}^{\mu,1}, \Theta_{v+1}^{\mu,2}, \dots$$

la fonction, l'ensemble, l'arc et la suite d'intervalles contigus de ces ensembles obtenus par l'opération

$$P(g_v(\tau), \overline{\Theta}_v^\mu).$$

Nous remplaçons sur tout Θ_ν^μ la fonction g_ν par $g_{\nu+1}^\mu$ et nous désignons la fonction ainsi obtenue par $g_{\nu+1}$.

Nous posons

$$\Pi_{\nu+1} = \Pi_\nu + \Sigma \Pi_{\nu+1}^\mu,$$

où la somme contient autant de termes que la suite Ω_ν .

Le μ -ème terme de la suite $S_{\nu+1}$ a déjà été défini.

En faisant varier μ dans $\Omega_{\nu+1}^\mu$ nous obtenons une suite double que nous pouvons ordonner en une suite simple suivant la méthode des diagonales et nous prenons cette suite comme la suite $\Omega_{\nu+1}$. $a_{\nu+1}^\mu$ est extrémité de la fonction g_ν sur $\Theta_{\nu+1}^\mu$.

[Il se peut que le pas de ν à $\nu + 1$ peut rater faute des Θ_ν^μ , nous écartons ce cas, en remarquant seulement que g_ν donne la représentation sémiréduite cherchée; la démonstration de ce cas particulier résulte facilement de celle du cas général].

En s'appuyant sur les propriétés de l'opération P (VI, § 14) et en tenant compte des définitions précédentes, on démontre de proche en proche que

$$\alpha) \quad g_\nu(\tau) = g_{\nu+q}(\tau) \text{ sur } \Pi_\nu, \quad (q = 1, 2, \dots)$$

$$\beta) \quad g_\nu(\overline{\Theta}_\nu^\mu) = g_{\nu+q}(\overline{\Theta}_\nu^\mu),$$

$$g_\nu(|\alpha, \beta| - \overline{\Theta}_\nu^\mu) = g_{\nu+q}(|\alpha, \beta| - \overline{\Theta}_\nu^\mu)$$

$$(q = 1, 2, \dots)$$

$$\gamma) \quad \Pi_\nu \subset \Pi_{\nu+1}.$$

$$1^\circ) \text{ diam } |\alpha_\nu^\mu, c_\nu^\mu| \geq \frac{1}{2} \text{ diam } g_\nu(\overline{\Theta}_\nu^\mu) = \frac{1}{2} \text{ diam } g_{\nu+q}(\overline{\Theta}_\nu^\mu).$$

2°) $g_\nu(\tau)$ est définie et continue sur $|\alpha, \beta|$, représente D , est cofinale sur $|\alpha, \beta|$ et admet a comme extrémité

3°) Π_ν est un sous-ensemble fermé de $|\alpha, \beta|$, contenant α et β et ayant les intervalles contigus $\Theta_\nu^1, \Theta_\nu^2, \dots$

4°) Π_ν ne s'alterne pas relativement à $g_\nu(\tau)$ et $|\alpha_{\nu+1}^\mu, c_{\nu+1}^\mu| \subset g_\nu(\overline{\Theta}_\nu^\mu); |\alpha_\nu^\mu, c_\nu^\mu| \subset g_\nu(\Pi_\nu)$.

$$5^\circ) \quad g_\nu(\overline{\Theta}_\nu^\mu) \times g_\nu(|\alpha, \beta| - \overline{\Theta}_\nu^\mu) = 0$$

6°) g_ν est cofinale sur Θ_ν^μ et admet le point a_ν^μ comme extrémité.

$$7^\circ) \quad g_\nu(\overline{\Theta}_\nu^\mu) \times g_\nu(\Pi_\nu) = a_\nu^\mu$$

α, β et γ résultent des définition de ce §.

1°) — 7°) résultent des propriétés portant les mêmes numéros et insérées au § 14, VI.

On obtient aussi sans peine

$$8^0) \quad |a_\nu^\mu, c_\nu^\mu| \times |a_{\nu+q}^\lambda, c_{\nu+q}^\lambda| \subset a_{\nu+q}^\lambda \\ (\lambda, \nu, q = 1, 2, \dots).$$

En effet selon ce qui précède

$$|a_\nu^\mu, c_\nu^\mu| \times |a_{\nu+q}^\lambda, c_{\nu+q}^\lambda| \subset g_\nu(\Pi_\nu) \times g_{\nu+q}(\bar{\Theta}_{\nu+q}^\lambda) \\ \subset g_{\nu+q}(\Pi_{\nu+q}) \times g_{\nu+q}(\bar{\Theta}_{\nu+q}^\lambda) = a_{\nu+q}^\lambda.$$

§ 2. En s'appuyant sur les remarques précédentes nous allons démontrer que $g_\nu(\tau)$ tend sur $[\alpha, \beta]$ vers une fonction $f(\tau)$ continue, cotinale et ayant α comme extrémité, représentant la dendrite D et que cette représentation est sémiréduite.

Dém. Comme la dendrite D ne peut pas renfermer une infinité d'arcs simples disjoints aux diamètres dépassant un nombre fixe positif, donc d'après 8⁰)

$$\lim_{\nu/\infty} \text{diam } |a_\nu^\mu, c_\nu^\mu| = 0,$$

donc selon 1⁰) et β

$$\lim_{\nu/\infty} \text{diam } g_{\nu+q}(\bar{\Theta}_\nu^\mu) = 0 \quad (q \geq 0)$$

et cette relation ne dépend pas de la façon suivant laquelle q et μ dépendent de ν .

Soit $\varepsilon > 0$. Il existe ν_0 , tel que

$$(1) \quad \text{diamètre } g_{\nu_0+q}(\bar{\Theta}_{\nu_0}^\mu) \leq \varepsilon.$$

On a sur Π_{ν_0}

$$(2) \quad \text{distance } (g_{\nu_0+q}(\tau), g_{\nu_0}(\tau)) = 0.$$

Pour τ contenu dans un contigu de Π_{ν_0} on a

$$g_{\nu_0+q}(\tau) \subset g_{\nu_0+q}(\bar{\Theta}_{\nu_0}^\mu) = g_{\nu_0}(\bar{\Theta}_{\nu_0}^\mu)$$

$$g_{\nu_0+q}(\tau) \subset g_{\nu_0}(\bar{\Theta}_{\nu_0}^\mu)$$

donc selon (1) distance

$$(g_{\nu_0+q}(\tau), g_{\nu_0}(\tau)) \leq \varepsilon.$$

En tenant compte de (2), on voit que cette égalité est valable pour tout τ du segment $[\alpha, \beta]$, elle exprime que la suite g_ν est uniformément convergente et par suite la fonction

$$f(\tau) = \lim g_\nu(\tau)$$

est continue.

Comme $g_\nu(\alpha) = g_\nu(\beta) = a$, donc f est cofinale et admet a comme extrémité.

f représente la dendrite D . C'est une conséquence de ce que g_ν représente D et tend uniformément vers f .

Reste à prouver que le segment $[\alpha, \beta]$ ne s'alterne pas relativement à f .

Nous allons, d'abord, démontrer que

$$(3) \quad f(\overline{\Theta}_\nu^\mu) = g_\nu(\overline{\Theta}_\nu^\mu)$$

$$(4) \quad f([\alpha, \beta] - \overline{\Theta}_\nu^\mu) = g_\nu([\alpha, \beta] - \overline{\Theta}_\nu^\mu).$$

On a selon β

$$(5) \quad g_{\nu+q}(\overline{\Theta}_\nu^\mu) = g_\nu(\overline{\Theta}_\nu^\mu).$$

Or (5) entraîne (3) car $\overline{\Theta}_\nu^\mu$ est fermé et la convergence est uniforme.

L'ensemble $[\alpha, \beta] - \overline{\Theta}_\nu^\mu$ est égal à

$$\Sigma'(\overline{\Theta}_\nu^\mu) + \Phi$$

Σ' étant certaine somme et Φ une partie de Π_ν . D'après (3)

$$f(\Sigma'(\overline{\Theta}_\nu^\mu)) = g_\nu(\Sigma'(\overline{\Theta}_\nu^\mu))$$

et comme sur Π_ν $g_\nu = g_{\nu+q}$ donc

$$f(\Phi) = g_\nu(\Phi).$$

En ajoutant ces égalités on obtient (4). (3) et (4) donnent selon 5°

$$(6) \quad f(\overline{\Theta}_\nu^\mu) \times f([\alpha, \beta] - \overline{\Theta}_\nu^\mu) = 0$$

Cela étant supposons qu'il existe

$$\tau_1^* < \tau_2 < \tau_3 < \tau_4$$

$$p' = f(\tau_1) = f(\tau_3) \neq f(\tau_2) = f(\tau_4) = p''.$$

Choisissons ν_0 si grand que

$$\text{diam } g_{\nu_0}(\overline{\Theta}_{\nu_0}^\mu) = \text{diam } f(\overline{\Theta}_{\nu_0}^\mu) < \text{distance } (p', p'').$$

C'est possible. Comme sur Π_ν on a $f = g_\nu$ donc Π_{ν_0} ne s'alterne pas relativement à f . Selon (6) et II, § 2 $\tau_1, \tau_2, \tau_3, \tau_4$ sont contenus dans un même segment contigu $\overline{\Theta}_{\nu_0}^\mu$ et c'est impossible suivant la dernière inégalité.

§ 3 A toute dendrite D qui ne se réduit pas à un seul point, appartient une représentation réduite et cofinale.

Nous allons indiquer la démonstration sans la développer complètement.

D'après I, § 4 il existe une représentation cofinale de D donc, suivant le § précédent, il en existe une qui est cofinale et sémiréduite, définie sur un segment $[\alpha, \beta]$. Cette représentation $f(\tau)$ peut admettre des intervalles dans lesquels elle est constante. f étant continue, à tout intervalle de cette sorte appartient l'intervalle le plus grand jouissant de la même propriété. Ces segments sont évidemment sans points commun deux à deux, désignons les par

$$\overline{\Delta}_1, \overline{\Delta}_2, \dots$$

Δ_ν étant intervalle aux mêmes extrémités que $\overline{\Delta}_\nu$. f est cofinale sur tout Δ_ν . $\Phi = [\alpha, \beta] - \Sigma \Delta_\nu$ est un ensemble fermé admettant au plus deux points isolés: α et β . La fonction f n'est pas constante sur aucun morceau de Φ contenu entre deux points de Φ n'étant pas extrémité d'un même Δ_ν (dans le cas contraire les Δ_ν ne seraient pas saturés). La fonction f représente D sur Φ .

D'après les propriétés précédentes de l'ensemble Φ , il existe une fonction continue non décroissant

$$\psi = \varphi(\tau)$$

$$\varphi(\alpha) = \alpha, \varphi(\beta) = \beta$$

et admettant les intervalles Δ_ν et ceux-là seulement comme intervalles saturés sur lesquels elle est constante.

En éliminant τ entre les relations

$$p = f(\tau)$$

$$\psi = \varphi(\tau)$$

on obtient une fonction

$$p = f_1(\psi)$$

qui correspond à la question. Nous laissons de côté le développement rigoureux de la démonstration qui vient d'être ébauchée.

§ 4. Soit $g(\tau)$ une représentation réduite et cofinale de D définie sur un segment $[\tau', \tau'']$. Sur D il existe au moins un point p^* ayant un seul point représentatif dans $[\tau', \tau'']$.

Dém. Tout d'abord nous démontrerons quelques lemmes.

Lemme I. Si $\tau_1 < \tau_2$ et que $g(\tau_1) = g(\tau_2)$, il existe entre τ_1 et τ_2 deux nombres différents pour lesquels g prend la même valeur, car dans le cas contraire g représenterait sur $[\tau_1, \tau_2]$ une orbe.

Lemme II. A tout $\varepsilon > 0$ il existe, sous les hypothèses du lemme précédent, deux nombres $\tau_I < \tau_{II}$ parfaitement déterminés, contenus entre τ_1 et τ_2 tels que $0 < \tau_{II} - \tau_I < \varepsilon$ et tels que $g(\tau)$ est régulièrement cofinale sur $[\tau_I, \tau_{II}]$.

En effet d'après le lemme précédent on voit facilement qu'il existe un entier n auquel correspond un intervalle Δ contenu entre τ_1 et τ_2 dont la longueur est au moins égale à $\frac{1}{n}$ et au plus égale à ε , sur lequel la fonction g est cofinale et admet un point e comme extrémité. Cet intervalle peut être déterminé sans ambiguïté: Les extrémités d'un intervalle contigu convenable de l'ensemble de points représentatifs de e qui se trouvent sur $\overline{\Delta}$ répondent à la question.

La démonstration du théorème est à présent facile. Suivant le lemme précédent il existe une suite infinie d'intervalles emboîtés dont la longueur tend vers 0 avec l'indice croissant infiniment et sur chacun desquels g est régulièrement cofinale. Soient $\tau'_\nu < \tau''_\nu$ les extrémités du ν -ème intervalle et soit τ^* le point commun à tous ces intervalles. On a

$$\tau'_\nu < \tau^* < \tau''_\nu$$

$$\lim \tau'_\nu = \lim \tau''_\nu = \tau^*$$

τ^* est bien l'unique point représentatif de $g(\tau^*)$, car, selon les relations précédentes et la remarque que $g(\tau^*) \neq g(\tau'_\nu) = g(\tau''_\nu)$, la représentation g ne serait pas réduite, s'il en existait un autre.

Corrol. I. Sur toute dendrite il existe au moins un point de saturation (I § 6).

Corrol. II. En effectuant sur la fonction $g(\tau)$ [cofinale et offrant une représentation réduite de D] une opération convenable du type de celle du § 6, I, nous obtenons une nouvelle représentation réduite de D régulièrement cofinale et admettant un point de saturation comme extrémité. Selon le lemme II il existe sur D un point de saturation différent du précédent, par conséquent:

Toute dendrite admet au moins deux points de saturation.

§ 5. Si p est un point de saturation de D et que $g(\tau)$ est une représentation cofinale et réduite de D sur un segment $[\tau', \tau'']$, alors, soit p a un seul point représentatif soit $g(\tau)$ est régulièrement cofinale et

$$g(\tau') = g(\tau'') = p.$$

Dém. Supposons d'abord que

$$g(\tau') = g(\tau'') = p.$$

Il s'agit de démontrer que g est régulièrement cofinale c. à d. qu'il n'existe pas τ''' ($\tau' < \tau''' < \tau''$) pour lequel

$$g(\tau''') = p.$$

Supposons le contraire. Les dendrites $g([\tau', \tau''])$, $g([\tau''', \tau''])$ n'ont pas de points communs excepté le point p , car dans le cas contraire la représentation g ne serait pas réduite, le point p découpe donc D au moins en deux parties, ce qui est impossible parce qu'il est un point de saturation.

Le cas où p admet un seul point représentatif se ramène au précédent, si l'on transforme $g(\tau)$ par le procédé du § 6, I.

§ 6. Le résultat du § précédent montre la différence du rôle joué par les extrémités du segment $[\tau', \tau'']$ et ses points intérieurs et suggère l'idée de remplacer, afin d'écartier cette différence, la représentation réduite sur le segment par la représentation „réduite“ sur la circonférence.

En déformant $[\tau', \tau'']$ d'une façon continue sur un plan et en confondant ses extrémités nous pouvons le faire coïncider avec une circonférence.

La fonction $g(\tau)$ passe par cette transformation de τ en une fonction continue définie sur la circonférence; nous dirons que cette nouvelle fonction fait une représentation réduite de la dendrite sur la circonférence.

Les résultats des §§ précédents s'expriment: Dans toute représentation réduite de D sur une circonférence C à tout point de saturation de D correspond un seul point représentatif sur C et tout point de D ayant un seul point représentatif sur C est un point de saturation.

En généralisant convenablement les raisonnements concernant les points de saturation on pourrait démontrer que tout point de D ayant n (ou une infinité dénombrable) de points représentatifs sur C divise D en n (ou une infinité dénombrable) parties et inversement.

En s'adressant à la construction de la représentation réduite de D sur un segment on pourrait démontrer que la classe de points p de D qui partagent D en plus de 2 parties est effectivement énumérable.

On pourrait à partir des arcs $[a_\mu'', c_\mu'']$ (VII § 1) d'une façon analogue à cette qui nous a servi dans les chapitres précédents, démontrer l'existence, sur le plan, d'un ensemble homéomorphe avec D .

§ 18 Imaginons une dendrite D quelconque tracée sur une feuille de papier. Découpons la feuille suivant tous les sous-arcs simples de D . Le trou ainsi pratiqué à une orbe de Jordan comme frontière. En subordonnant à chaque point cette orbe (qui est homéomorphe avec circonférence) le point de D qui l'a engendré nous obtenons une représentation réduite de D sur cette orbe.

Nous laissons de côté la précision et la démonstration de cette remarque

L'élégance de la méthode de M Denjoy consiste à construire une telle représentation sans l'application homéomorphe préalable de la dendrite sur le plan.

Notes.

§ 1. Comme il est aisé de le voir, le point de branchement d'une dendrite et le point qui la découpe en plus de deux parties c'est la même chose.

Cette remarque conduit au théorème

La classe de points découpant une courbe de Jordan quelconque [faisant partie d'un espace normal $E(q)$] en plus de deux parties est au plus dénombrable.

La démonstration est basée sur le principe de choix et sur le critère B § 6 β .

§ 2. Il existe sur le plan deux courbes de Jordan A et B se correspondant suivant une homéomorphie $\sim$ et telles que pour aucune couple de points associés a et b ($a \sim b$), la correspondance $\sim$ ne se laisse compléter sur aucun voisinage de a , de façon à ne pas changer son caractère de homéomorphie.

Soit en effet A une dendrite développable en une suite Ω spéciale.

$$((a_1, b_1)) ((a_2, b_2)), \dots$$

Nous allons déformer A de façon à obtenir B répondant à la question.

Soit $a'_2 \varepsilon (a_2)$. Les $[a_2, b_2]$ pour lesquelles $a_2 = a'_2$ forment un feuillage composé d'une infinité de feuilles et poussant sur a'_2 .

Faisons changer de place, dans ce feuillage, deux feuilles qui

ne sont pas voisines c. à. d. qui sont séparées par deux autres feuilles. En procédant pareillement avec les feuilles d'ordre supérieur on obtient à la limite l'ensemble B en question.

§ 3. Le théorème D § 18 se généralise de la façon suivante.

La classe des constituants d'une courbe de Jordan plane A qui correspondent à une courbe partielle A_1 est faible.

Ce théorème est étroitement lié avec le théorème:

Si A est une courbe de Jordan contenue dans un plan P ; K_ν un continu tel que

$$\begin{aligned} K_\nu &\subset P & (\nu/1, \dots, \infty) \\ K_\nu \times A &= 0 \\ \varphi^*(K_\nu, K) &= 0 \\ K &\subset A, \end{aligned}$$

alors K est une courbe de Jordan. De là résulte le corollaire:

La frontière de tout domaine plan complémentaire d'une courbe de Jordan plane est une courbe de Jordan.

Errata.

Page	ligne	au lieu	doit être
89	14		$((a_1, b_1)), ((a_2, b_2))$
90	22	$\alpha_\mu \leq \alpha_\mu$	$\alpha_\nu \leq \alpha_\mu$
90	23	$- II$	$- \pi$
91	36	§ 19	I § 19
92	36	I § 38	I § 36