

JOURNAL
DE
MATHÉMATIQUES

PURES ET APPLIQUÉES

FONDÉ EN 1836 ET PUBLIÉ JUSQU'EN 1874

PAR JOSEPH LIOUVILLE

H. LEBESGUE

Sur les fonctions représentables analytiquement

Journal de mathématiques pures et appliquées 6^e série, tome 1 (1905), p. 139-216.

http://www.numdam.org/item?id=JMPA_1905_6_1__139_0



NUMDAM

Article numérisé dans le cadre du programme
Gallica de la Bibliothèque nationale de France
<http://gallica.bnf.fr/>

et catalogué par Mathdoc
dans le cadre du pôle associé BnF/Mathdoc
<http://www.numdam.org/journals/JMPA>

*Sur les fonctions représentables analytiquement ;***PAR M. H. LEBESGUE.****I. — Introduction.**

Bien que, depuis Dirichlet et Riemann, on s'accorde généralement à dire qu'il y a fonction quand il y a correspondance entre un nombre y et des nombres $x_1, x_2, \dots, x_n$, sans se préoccuper du procédé qui sert à établir cette correspondance, beaucoup de mathématiciens semblent ne considérer comme de vraies fonctions que celles qui sont établies par des correspondances analytiques. On peut penser qu'on introduit peut-être ainsi une restriction assez arbitraire ; cependant il est certain que cela ne restreint pas pratiquement le champ des applications, parce que, seules, les fonctions représentables analytiquement sont effectivement employées jusqu'à présent.

Dans certaines théories générales, dans la théorie de l'intégration au sens de Riemann, par exemple, on ne se préoccupe pas de savoir si les fonctions que l'on considère sont ou non représentables analytiquement. Mais cela ne veut pas dire qu'elles ne le sont pas toutes ⁽¹⁾ et, dans tous les cas, quand on applique effectivement ces théories, c'est toujours sur des fonctions représentables analytiquement qu'on opère.

(¹) On verra cependant, à la fin du paragraphe VIII, que l'intégration au sens de Riemann s'applique à des fonctions non représentables analytiquement.

On peut dire plus : quand on emploie les expressions analytiques que M. Baire a considérées dans sa Thèse [*Sur les fonctions de variables réelles* (*Annali di Matematica*, 1900)], on reconnaît facilement que toutes les fonctions qui ont été citées jusqu'ici, qu'elles se soient présentées naturellement ou qu'elles aient été construites de toutes pièces pour fournir des exemples de singularités, sont toutes représentables analytiquement. Ce résultat, surprenant au premier abord, étonnera moins si l'on se rappelle que la fonction $\chi(x)$, si souvent citée comme exemple de singularités, égale à un pour x rationnel, à zéro pour x irrationnel, admet la représentation analytique suivante :

$$\chi(x) = \lim_{m \rightarrow \infty} \left[\lim_{n \rightarrow \infty} (\cos m! \pi x)^{2n} \right].$$

Ainsi il n'est pas évident qu'il existe des fonctions non représentables analytiquement; il y a donc lieu de rechercher s'il existe de telles fonctions et, s'il en existe, il y a lieu de rechercher des propriétés communes à toutes les fonctions représentables analytiquement. Non seulement parce que ces propriétés permettront peut-être de reconnaître si certaines fonctions sont ou non représentables analytiquement, mais surtout parce qu'il faut pouvoir supposer que toutes les fonctions sur lesquelles on raisonne possèdent certaines propriétés particulières, appartenant à toutes les fonctions représentables analytiquement sans appartenir à toutes les fonctions, pour que cela ait un sens de dire qu'on se restreint à la considération des fonctions représentables analytiquement (¹).

(¹) On verra que, pour qu'une propriété appartienne à toutes les fonctions représentables analytiquement, il suffit qu'elle appartienne aux polynômes et qu'elle soit vraie de la somme et du produit de deux fonctions, et de la limite d'une suite de fonctions, dès qu'elle est vraie de chacune d'elles. J'ai donné (*Comptes rendus*, 29 avril 1902) et dans ma Thèse [*Intégrale, longueur, aire* (*Annali di Matematica*, 1902)] une telle propriété, d'où j'ai déduit une autre propriété de même nature (*Comptes rendus*, 28 décembre 1903), sur laquelle je ne reviendrai pas dans ce Mémoire, et que M. Borel avait obtenue de son côté (*Comptes rendus*, 7 décembre 1903).

M. Baire (*Comptes rendus*, 11 décembre 1899) avait énoncé le premier une

La recherche de conditions nécessaires pour qu'une fonction soit représentable analytiquement d'une manière particulière a été l'occasion de nombreux travaux. Je ne citerai que deux résultats dus à Dirichlet et Weierstrass :

Toute fonction continue n'ayant qu'un nombre fini de maxima et minima est représentable par la série de Fourier ;

Toute fonction continue est représentable par une série uniformément convergente de polynomes.

Donc, tandis que l'on aurait pu craindre ne pouvoir exprimer que par des relations analytiques compliquées des conditions nécessaires pour qu'il y ait une représentation analytique, il suffit que certaines conditions très simples relatives à la variation soient remplies pour qu'il en soit ainsi. C'est encore un fait du même genre qu'a mis en évidence M. Baire en faisant connaître à quelles conditions doit satisfaire une fonction pour être la somme d'une série de polynomes (*voir plus loin*, théorème XV). M. Baire a fait connaître de plus, dans sa Thèse, une importante classification des fonctions qui permet d'aborder la recherche de propriétés caractéristiques des fonctions admettant une représentation analytique quelconque (¹).

Dans cette étude j'ai employé des raisonnements simples qui m'avaient déjà permis de retrouver et de compléter certains des résultats de M. Baire (²) ; j'ai surtout précisé une remarque, déjà citée en note, faite dans ma Thèse, page 27 ; cela m'a conduit à étudier la nature de l'ensemble des points pour lesquels une fonction f satisfait à l'inégalité $a \leq f \leq b$. J'ai obtenu des conditions caractéristiques pour que des fonctions soient de chacune des classes que considère M. Baire (théorèmes IV, VII, XVII) ou pour qu'une fonction admette une représentation analytique (théorème VI).

propriété appartenant à toutes les fonctions représentables analytiquement. Mais on ne sait pas si toutes ces propriétés n'appartiennent pas à toutes les fonctions. C'est en modifiant convenablement la première propriété citée que j'ai obtenu une propriété caractéristique des fonctions représentables analytiquement.

(¹) Outre ce qui se trouve dans sa Thèse, M. Baire a publié deux Notes à ce sujet (*Comptes rendus*, 4 et 11 décembre 1899).

(²) Voir *Bulletin des Sciences mathématiques*, novembre 1898 et *Comptes rendus*, 27 mars 1899.

Comme application de cette dernière condition, j'ai distingué dans l'ensemble des fonctions déterminées analytiquement celles qui le sont explicitement ($y =$ expression analytique de $x_1, x_2, \dots, x_n$) et celles qui le sont implicitement (expression analytique de $x_1, x_2, \dots, x_n, y = 0$). On ne distingue pas ordinairement entre ces deux catégories de fonctions, parce que, dans la pratique, le procédé même qui prouve l'existence d'une fonction implicite en fournit un développement; mais l'identité de ces deux familles de fonctions n'est pas évidente. Elle résulte du théorème VI (théorème XVIII).

Enfin (§ VIII), j'ai cité des exemples de fonctions de toutes les classes et de fonctions échappant à tout mode de représentation analytique.

Dans le paragraphe suivant je donnerai quelques définitions indispensables et la classification de M. Baire. Avant cela je veux dire pourquoi l'emploi, fait dans cette classification, des nombres transfinis ne soulève, à mon avis, aucune difficulté.

Si l'on étudie la croissance des fonctions et si, ayant caractérisé la croissance de x^n par n , on constate que e^x croît plus vite que x^n , on pourra éprouver le désir de caractériser cette nouvelle croissance par un nouveau symbole, ω (¹). Nul n'y verra d'inconvénients. Si l'on dit que ω est un nombre transfini et est plus grand que n , on pourra trouver que c'est là un langage bien mal choisi, mauvais pratiquement, mais on ne pourra le déclarer mauvais logiquement. L'emploi que l'on fait, dans la classification de M. Baire, des nombres transfinis est analogue à celui que je viens de rappeler. Les nombres transfinis y sont des signes, des *symboles de classe* permettant de distinguer ces diverses classes.

D'ailleurs, la classification de M. Baire, dont toutes les classes existent effectivement comme on le verra au paragraphe VIII, peut, comme la théorie de la croissance des fonctions (²), ou comme la théorie des ensembles dérivés, fournir une base solide pour construire

(¹) Voir BOREL, *Leçons sur la théorie des fonctions*, Paris, Gauthier-Villars, 1898, Note II.

(²) Et même plus simplement à cause de l'indétermination des échelles de types de croissance.

la théorie des nombres transfinis. Sans développer tous les raisonnements nécessaires, c'est ce point de vue que j'ai adopté; en d'autres termes je n'ai jamais emprunté une proposition à la théorie abstraite des nombres transfinis, j'ai toujours indiqué rapidement comment l'on démontrerait cette propriété pour l'ensemble des symboles de classe.

Si cet emploi déguisé des nombres transfinis était encore gênant pour quelque lecteur, celui-ci devrait remarquer que les nombres transfinis n'interviennent jamais dans les raisonnements ⁽¹⁾, seulement les propriétés ne seraient démontrées que pour les classes de fonctions dont les symboles de classe sont des entiers, puisque le lecteur se refuserait à considérer les autres; j'ajoute que ce refus serait, à mon avis, tout à fait injustifié, logiquement du moins.

II. — Définitions.

Intervalle. Domaine. — Je vais m'occuper de fonctions d'un nombre quelconque, fini, de variables réelles $x_1, x_2, \dots, x_n$. Ces fonctions seront définies pour certains systèmes de valeurs des variables, c'est-à-dire, en adoptant un langage souvent employé, pour certains points de l'espace à n dimensions $x_1, x_2, \dots, x_n$. Je m'occuperai surtout des fonctions définies dans tout un intervalle ou un domaine.

Un *intervalle* est l'ensemble des points satisfaisant aux conditions

$$a_1 \leq x_1 \leq b_1, \quad a_2 \leq x_2 \leq b_2, \quad \dots, \quad a_n \leq x_n \leq b_n.$$

Une transformation de la forme

$$\begin{aligned} X_1 &= X_1(x_1, x_2, \dots, x_n), & X_2 &= X_2(x_1, x_2, \dots, x_n), & \dots, \\ & & X_n &= X_n(x_1, x_2, \dots, x_n), \end{aligned}$$

où les X_i sont des fonctions continues, fait correspondre à tout point

⁽¹⁾ Les nombres transfinis interviennent au contraire dans les raisonnements de M. Baire; cependant, à mon avis, comme je le dirai plus loin, la méthode de M. Baire a certains avantages que ne possède pas celle du texte.

m de l'espace $(x_1, x_2, \dots, x_n)$ un point M de l'espace $(X_1, X_2, \dots, X_n)$; si à tout point M correspond au plus un point m , par définition, cette transformation fera correspondre un *domaine* à un intervalle.

Un intervalle sera fini si tous les a_i et b_i sont finis, les domaines correspondant à ces intervalles sont aussi dits *finis*. Un intervalle est dégénéré si, pour certaines valeurs de i , a_i est égal à b_i ; aux intervalles dégénérés correspondent les domaines dégénérés.

Je ne raisonnerai, en général, que sur les domaines finis non dégénérés; la plupart des propositions obtenues sont cependant vraies pour tous les domaines. Pour s'en assurer, ou pour voir comment on doit les modifier, il suffira de remarquer qu'un domaine infini est la réunion d'une infinité dénombrable de domaines finis, qu'un domaine dégénéré est la partie commune à une infinité de domaines non dégénérés.

Les 2^n points d'un intervalle dont, quel que soit i , la coordonnée x_i est égale à a_i ou b_i sont les sommets de l'intervalle. Les points pour lesquels l'une au moins des coordonnées x_i a l'une de ses valeurs limites a_i ou b_i constituent la *frontière* de l'intervalle. A la frontière d'un intervalle correspond la frontière d'un domaine. Quand un point appartient à un domaine je dirai qu'il est *contenu dans ce domaine*; si, de plus, il ne fait pas partie de la frontière du domaine, je dirai qu'il est *contenu à l'intérieur du domaine* ⁽¹⁾.

Dans la suite, je supposerai toujours qu'un domaine D non dégénéré a été choisi et je ne m'occuperai que des points de ce domaine. C'est ainsi que, lorsque je dirai qu'une fonction f est partout définie, cela voudra dire qu'elle est définie pour tous les points de D , mais f ne sera pas nécessairement partout définie dans l'espace. Lorsque je dirai

(1) J'adopte les définitions du texte pour éviter les difficultés qu'on rencontre dans la démonstration des propriétés des domaines quand on définit ceux-ci par la considération des variétés fermées à $n - 1$ dimensions.

Pour que les définitions du texte soient acceptables, il faut démontrer que la frontière d'un domaine ne dépend pas de la manière dont on l'a déduite d'un intervalle; on y arrivera en prouvant que la définition du texte rentre comme cas particulier dans la définition classique : un point M est point frontière d'un ensemble E si tout intervalle contenant M à son intérieur contient à la fois des points de E et des points n'appartenant pas à E ; la frontière de E est l'ensemble de ses points frontières (JORDAN, *Cours d'Analyse*, t. I).

qu'une expression e représente une fonction f cela voudra dire que, pour les points de D où e a un sens, f est définie et égale à e ; que, pour les points de D où e n'a pas de sens, f n'est pas définie; mais il se pourrait que, en certains points n'appartenant pas à D , e ait un sens sans que f soit définie ou inversement, ou encore que e ait une valeur différente de f ⁽¹⁾. Quant au domaine D ce sera un domaine quelconque fini ou infini, ce pourra être tout l'espace.

Expression analytique. Fonctions représentables ou exprimables analytiquement. Fonctions définies ou données analytiquement. — Je dirai qu'une fonction est représentable ou exprimable analytiquement lorsqu'on peut la construire en effectuant, suivant une loi donnée, certaines opérations; cette loi de construction constitue une *expression analytique*. Il faut évidemment préciser les opérations que l'on admet; si l'on veut que les fonctions $x_1, x_2, \dots, x_n$, respectivement égales aux variables, rentrent dans l'ensemble des fonctions représentables analytiquement, ensemble que je désigne avec M. Baire par E ; si l'on veut que la somme et le produit de deux fonctions de E appartiennent à E ; si l'on veut que la somme d'une série convergente de fonctions de E soit une fonction de E , il faut que toute fonction que l'on peut construire en effectuant suivant une loi déterminée un nombre fini ou dénombrable d'additions, de multiplications, de passages à la limite, à partir des variables et de constantes, rentre dans l'ensemble des fonctions exprimables analytiquement.

C'est aux fonctions ainsi définies que je réserverai le nom de *fonctions représentables analytiquement*, mais, pour que les lois de construction considérées puissent conduire à des fonctions non partout définies, je n'exclurai pas celles de ces lois qui conduiraient à

(¹) La convention faite ici est celle qui est adoptée le plus généralement; on peut même dire que c'est celle qui est toujours adoptée pour les fonctions définies en tous les points de D . Pour les fonctions f définies seulement en certains points de D on fait parfois la convention qu'une expression e représente f pourvu que, en tous les points de D où f est définie, e ait un sens et soit égale à f ; on ne se préoccupe pas de la valeur de e aux points de D où f n'est pas définie. Il importe de ne pas confondre cette convention avec celle du texte.

prendre la limite d'une suite

$$u_1, u_2, \dots$$

qui ne serait pas convergente pour tous les points où tous les u_i ont un sens. Bien entendu la limite des u_i n'existe que si la suite est convergente, mais cette suite n'est pas supposée convergente partout où les u_i existent tous.

Les expressions analytiques considérées ordinairement contiennent d'autres signes que ceux qui ont été employés; on y rencontre par exemple les signes $-$, $:$, $\sqrt[n]{}$, $\sin$, $\log$, etc. Il est facile de voir qu'en adjoignant les opérations correspondantes à celles qui ont été employées, on n'élargit pas l'ensemble des fonctions représentables analytiquement. Par exemple, on peut remplacer $u - v$ par $u + v.(-1)$; $\frac{u}{v}$ peut être remplacé par $u \times \frac{1}{v}$ et l'on peut nommer une série de polynomes en v convergente, sauf pour $v = 0$, et représentant $\frac{1}{v}$; par suite la division est remplacée par des additions, des multiplications et un passage à la limite. Une démonstration analogue peut être faite pour chacun des autres signes, car ce sont tous des symboles représentant des fonctions définies dans certains domaines et continues dans les domaines où elles sont définies; de telles fonctions peuvent toujours être représentées par des séries de polynomes ⁽¹⁾.

On emploie aussi quelquefois des symboles d'intégration et de dérivation; l'intégration et la dérivation ne font pas correspondre un nombre à un ensemble fini de nombres donnés, mais une fonction à une fonction donnée; il est donc nécessaire de les étudier à part et il serait nécessaire de même d'étudier toutes les autres opérations fonctionnelles que l'on conviendrait d'employer. Je ne ferai pas cette étude; je me contenterai d'affirmer que, même si l'on donne à l'intégrale le sens que j'ai adopté dans ma Thèse, l'intégration d'une fonction représentable analytiquement donne une fonction de même nature et que la dérivation, même si on l'applique à une fonction non partout

(1) Souvent même elles sont définies par de telles séries, mais il arrive parfois que les séries de définition ne les représentent que dans certains domaines.

dérivable et non partout continu, même si on la remplace par la dérivation supérieure ou inférieure, à droite ou à gauche, de Dini, conduit à une fonction exprimable analytiquement si elle est appliquée à une fonction de même nature. D'ailleurs, dans les expressions analytiques ordinairement employées, l'intégration et la dérivation conduisent à des fonctions continues et alors toute démonstration est inutile (¹).

Une loi de construction constituera une expression analytique si elle n'emploie qu'un nombre fini ou dénombrable d'opérations définies par des symboles de fonctions exprimables analytiquement et un nombre fini ou dénombrable de passages à la limite; ces opérations devant être effectuées à partir de constantes et de fonctions exprimables analytiquement données. Il est évident, d'après la définition même, qu'une expression analytique peut toujours être remplacée par une expression équivalente et n'employant que des additions, des multiplications et des passages à la limite effectués à partir des variables et de constantes.

Une fonction $y_1(x_1, x_2, \dots, x_n)$ sera dite déterminée, donnée ou définie analytiquement, si elle est définie en même temps qu'un nombre fini (²) de fonctions $y_2, \dots, y_p$, comme l'une des solutions d'un système de p équations obtenues en égalant à zéro des fonctions exprimables analytiquement de $x_1, x_2, \dots, x_n; y_1, y_2, \dots, y_p$.

Classification des fonctions. — M. Baire considère les expressions analytiques dans lesquelles les signes $+$ et $\times$ ne sont appliqués qu'à des constantes et aux variables. Ces expressions représentent donc les

(¹) On pourrait croire que toute démonstration est toujours inutile pour l'intégration; il n'en est rien. Le symbole $\int_0^x f(x, y) dx$ donne bien, quand il a un sens pour $x = X$, une fonction continue en x entre 0 et X , mais nous ne savons pas pour quel ensemble de valeurs (x, y) il a un sens, ni comment varie sa valeur quand y varie.

Les affirmations du texte se justifient facilement si l'on emploie les propriétés des ensembles et des fonctions mesurables B qui sont démontrées plus loin.

(²) On pourrait, comme je le dirai plus loin, supposer que les y forment une infinité dénombrable.

fonctions qui se déduisent des polynomes par des passages répétés à la limite ⁽¹⁾.

M. Baire appelle *fonctions de classe zéro* les fonctions partout définies qui sont continues par rapport à l'ensemble des variables. Les fonctions partout définies, discontinues, qui sont limites de fonctions continues, ou, ce qui revient au même, qui sont limites de polynomes, constituent *la classe 1*.

Les fonctions partout définies, qui ne sont pas de classe 1, et qui sont limites de fonctions des classes 0 et 1, forment *la classe 2*.

D'une manière générale, un ensemble $\mathcal{C}$ de classes ayant été défini, la première classe de fonctions venant après cet ensemble de classes sera formée des fonctions partout définies qui n'appartiennent pas à cet ensemble de classes et qui sont limites de fonctions de cet ensemble de classes s'il existe de telles fonctions. Pour distinguer cette classe des précédentes nous créerons un signe α que nous appellerons un *symbole de classe*. La classe définie sera dite *la classe α* . Le symbole α sera dit *plus grand* que les symboles des classes qui font partie de $\mathcal{C}$, je dirai aussi que les classes constituant $\mathcal{C}$ sont *antérieures* ou *inférieures* à la classe α ou encore que les fonctions des classes de $\mathcal{C}$ sont des classes inférieures à α , etc. ⁽²⁾.

Pour noter les premières classes nous emploierons, comme symboles de classes, les nombres entiers successifs; puis, s'il existe effectivement une classe venant après l'ensemble des classes à symboles entiers, nous créerons un symbole représentant cette classe. Le symbole ordinairement adopté est ω . Si, après la classe ω , existe une classe nouvelle, nous la nommerons à l'aide d'un nouveau symbole. La manière d'écrire ces symboles n'est pas donnée, elle importe peu

(1) Les passages à la limite que M. Baire considère, ainsi que ceux qui ont été employés précédemment, ne sont pas nécessairement appliqués à des suites convergeant uniformément. En d'autres termes, dire que u est la limite de la suite $u_1, u_2, \dots$, c'est dire que u est la somme de la série convergente $u_1 + (u_2 - u_1) + (u_3 - u_2) + \dots$; ce n'est pas dire que cette série est uniformément convergente.

(2) Il faudrait démontrer que l'emploi, à la manière ordinaire, des mots *plus grand* et *plus petit*, *inférieur* et *supérieur*, *antérieur* et *postérieur*, etc., ne conduira jamais à une contradiction; cela est à peu près évident.

pour la suite. Il nous sera toutefois commode de convenir que, si α est un symbole de classe, par la classe $\alpha + 1$ nous désignerons celle qui suit immédiatement l'ensemble $\mathcal{C}$ formé de la classe α et des classes inférieures. D'ailleurs, si β est le symbole de la classe que nous venons de désigner par $\alpha + 1$, $\beta - 1$ devra être considéré comme équivalent à β .

Il faut remarquer que, si α est un symbole de classe, $\alpha + 1$ s'applique toujours à une classe déterminée (¹), mais il se peut que $\alpha - 1$ ne s'applique à aucune classe déterminée; c'est le cas de $\alpha = \omega$, par exemple. Lorsque $\alpha - 1$ s'appliquera à une classe déterminée, α sera dit de *première espèce*, sinon il sera dit de *seconde espèce*.

La différence entre les deux espèces de classes est assez grande; nous la rencontrerons constamment dans la suite. La conséquence la plus importante de l'existence de ces deux espèces de classes est qu'il ne suffit pas du raisonnement par récurrence ordinaire pour démontrer une proposition pour toutes les classes. Nous *admettrons* qu'une proposition est démontrée pour toutes les classes, si elle est vérifiée pour la classe 0, et s'il est prouvé qu'elle est vraie de la classe α dès qu'elle est vraie de toutes les classes inférieures à α . Cela ne veut pas dire qu'il suffit de la vérifier pour $\alpha = 0$ et de démontrer que son exactitude pour $\alpha = \beta$ entraîne son exactitude pour $\alpha = \beta + 1$. Il faut, de plus, prouver que, si α est de seconde espèce, la propriété est vraie pour α dès qu'elle est vraie de tous les symboles inférieurs à α . Il arrive très souvent, cependant, quand on donne à un raisonnement la dernière forme indiquée, que l'on néglige de donner explicitement la démonstration relative au cas où α est de seconde espèce parce que ce raisonnement est souvent très simple et qu'on le sous-entend (²).

(¹) L'existence effective de toutes les classes logiquement conçues est admise ici, comme dans toute la suite. La démonstration de cette existence pourrait être donnée immédiatement, j'ai pu la présenter sous une forme plus simple en la rejetant à la fin de ce Mémoire, au paragraphe VIII. On verra, dans un instant, quelle est la limitation logique de l'ensemble des classes.

(²) Il n'est peut-être pas inutile de remarquer que, tandis que le raisonnement par récurrence ordinaire, applicable à l'ensemble des nombres entiers, fournit un procédé régulier permettant de vérifier la proposition pour un entier quelconque donné à l'aide d'un nombre fini de syllogismes, le raisonnement par ré-

Au point de vue de la représentation des fonctions, la différence entre les deux espèces de classes se manifeste de la manière suivante : Soit f une fonction de classe α limite des fonctions $f_1, f_2, \dots$ des classes $\alpha_1, \alpha_2, \dots$ inférieures à α . Si α est de première espèce ceux des α_i qui ne sont pas égaux à $\alpha - 1$ sont en nombre fini ; de sorte que, en supprimant les fonctions correspondantes, f est définie comme limite de fonctions de classe $\alpha - 1$. Si α est de seconde espèce, nous ne pouvons plus supposer que toutes les f_i sont de même classe, mais nous pouvons supposer qu'elles sont de classes croissantes ; car si l'on supprime tous les $f_{p+1}, f_{p+2}, \dots$ qui ne sont pas de classes supérieures à α_p , et cela quel que soit p , il restera certainement des fonctions et elles formeront une suite définissant f comme limite. D'ailleurs α est le premier symbole surpassant tous les α_p conservés ; donc *tout symbole α peut être déterminé comme le plus petit de ceux qui suivent une suite croissante, finie ou dénombrable, de symboles de classe*. Aux suites finies correspondent les symboles de première espèce, aux suites infinies correspondent ceux de seconde espèce.

La définition des symboles entraîne une conséquence importante : *Les symboles inférieurs à α sont en nombre fini ou dénombrable*. Conservons les notations précédentes ; si l'on a $\beta < \alpha$, β est égal ou inférieur à l'un des α_i ; par suite, si la propriété est vraie des α_i , elle est vraie de α . Les passages à la limite appliqués à des suites dénombrables ne permettent donc pas de passer de classes dont les symboles jouissent de la propriété indiquée à des classes dont les symboles ne possèdent plus cette propriété. Puisque toutes les classes sont formées par des passages à la limite à partir de suites de fonctions de classe 0, tous les symboles jouissent bien de la propriété indiquée. Cela ne veut pas dire que l'ensemble des classes soit dénombrable, tout au contraire, la définition générale des classes nous permet de concevoir une classe venant après une infinité dénombrable quelconque de classes. *L'ensemble e des classes, conçu logiquement, jouit donc, entre autres propriétés, des suivantes :*

1° *Il n'est pas dénombrable ;*

currence applicable à l'ensemble des symboles de classe ne fournit rien d'analogue.

2° Avant une classe quelconque, il n'existe qu'un nombre fini ou dénombrable d'autres classes;

3° Après tout ensemble dénombrable de classes, il existe une nouvelle classe.

Les fonctions des différentes classes de e existent réellement, on le verra au paragraphe VIII; pour le moment les différentes classes de e ne sont que logiquement conçues comme pouvant résulter de la définition générale des classes. Il faut bien comprendre comment il se fait que cette définition ne permet pas de concevoir une classe venant après toutes celles de e ; c'est parce que, quelle que soit la suite convergente $f_1, f_2, \dots$ de fonctions de e , sa limite fera partie de e puisqu'elle sera au plus de la classe de e qui, d'après 3°, suit immédiatement les classes de $f_1, f_2, \dots$.

Ainsi, notre définition contient en elle-même un principe logique de limitation, qui restera le même, quel que soit le procédé de classification que l'on emploie, et quel que soit le contenu de la classe o , pourvu qu'une fonction soit définie par une infinité *dénombrable* de fonctions des classes antérieures (¹).

(¹) La classification donnée par M. Baire fournit une occasion d'utiliser les symboles imaginés par M. Cantor et que l'on appelle *nombres transfinis* de la première classe de nombres transfinis ou de la deuxième classe numérique. Pour utiliser de même les nombres transfinis de la troisième classe numérique il faudrait, d'après ce qui est dit dans le texte, employer un procédé de classification dans lequel une fonction serait définie par une infinité non dénombrable de fonctions des classes antérieures. On peut citer de tels procédés.

J'appelle avec M. Borel (*Leçons sur la théorie des fonctions*, p. 122) *suite transfinie* une suite de nombres ayant pour indices les symboles de e . On imagine facilement ce qu'on entendra par la limite d'une telle suite. Ceci posé, partons des fonctions étudiées par M. Baire comme classe initiale, et convenons que chaque classe nouvelle sera formée des fonctions limites des suites transfinies de fonctions des classes antérieures. Une telle classification fournirait peut-être l'occasion d'utiliser les nombres transfinis de la deuxième classe de nombres transfinis.

Il est à remarquer que la première des classes ainsi définie, celle qui vient après l'ensemble des fonctions exprimables analytiquement, existe certainement. On verrait, en effet, facilement que les fonctions définies au paragraphe VIII et qui échappent à toute représentation analytique font partie de cette classe. Mais il faut remarquer aussi que cette classe existerait seule, ce qui enlèverait tout

Jusqu'ici, nous ne nous sommes occupés que des fonctions partout définies; il n'y a évidemment aucune difficulté à appliquer la classification de M. Baire à toutes les fonctions. Cependant, pour simplifier, je ne m'occuperai dans la suite, sauf avis contraire, que de la classification des fonctions partout définies. Certains des théorèmes que l'on obtiendra, relatifs aux fonctions d'une classe déterminée, devraient être légèrement modifiés pour s'appliquer aux fonctions non partout définies (¹).

M. Baire désigne par E l'ensemble des fonctions auxquelles s'applique sa classification; cet ensemble de fonctions ne diffère pas de celui que j'ai appelé plus haut E , c'est-à-dire de l'ensemble des fonctions représentables analytiquement. Cela se démontre facilement pour les fonctions partout définies. Raisonnons sur ces fonctions; les fonctions représentables analytiquement sont formées à l'aide d'additions, de multiplications et de passages à la limite à partir de fonctions de M. Baire. Il suffit donc de démontrer que la limite d'une suite convergente de fonctions de l'ensemble E de M. Baire est une fonction du même ensemble, ce qui est évident par définition même de E , et de démontrer que la somme et le produit de deux fonctions de E sont aussi une fonction de E . Or, soient f et g deux fonctions de l'ensemble E de M. Baire; on peut trouver un symbole de classe α tel que f et g soient au plus de classe α . Alors f et g sont respectivement limites des fonctions f_p et g_p de classes inférieures à α ; puisque $f + g$ et fg sont respectivement limites de $f_p + g_p$ et $f_p g_p$, il sera démontré que le produit et la somme de deux fonctions de classe β au plus sont des fonctions de classe β au plus pour $\beta = \alpha$, si cela est démontré pour $\beta < \alpha$.

intérêt à la classification, si, comme on l'a prétendu quelquefois, l'ensemble e avait la puissance du continu.

(¹) Pour la simplicité des énoncés il serait d'ailleurs bon de faire rentrer dans la classe 0, non seulement les fonctions continues partout définies, mais encore les fonctions définies seulement pour les points d'un ensemble fermé et continues sur cet ensemble. La classe 0 contiendrait donc, si le domaine D est fini, les séries de polynômes uniformément convergentes dans l'ensemble des points de convergence.

Nous pouvons appliquer le raisonnement par récurrence généralisé, puisque la propriété considérée est vraie pour $\beta = 0$. Donc, le produit et la somme de deux fonctions rentrant dans la classification de M. Baire sont une fonction de même nature. Il y a bien identité, pour les fonctions partout définies, entre les fonctions représentables analytiquement et celles auxquelles s'appliquent la classification de M. Baire.

Cela est moins évident pour les fonctions non partout définies, parce qu'il se peut que $\lim (f_p + g_p)$ ou $\lim f_p g_p$ ait un sens sans que $\lim f_p$ et $\lim g_p$ en aient un. La propriété est cependant vraie, elle sera démontrée incidemment dans la suite (*voir* p. 170).

III. — Opérations effectuées sur des fonctions de classe α .

Pardes raisonnements analogues à ceux du paragraphe précédent, on démontrerait que $f + \varphi$, $f - \varphi$, $f\varphi$, $\sqrt[m]{f}$, $\cos f$, ..., sont de classe α au plus, si f et φ sont de classe α au plus. Tous ces résultats se déduisent d'ailleurs du théorème suivant :

I. Si la fonction $\varphi(t_1, t_2, \dots, t_p)$ est de classe 0, et si $f_1, f_2, \dots, f_p$ sont de classe α au plus, $\Phi = \varphi(f_1, f_2, \dots, f_p)$ est de classe α au plus ⁽¹⁾.

Conformément à nos conventions $f_1, f_2, \dots, f_p$ sont partout définies dans un certain domaine D; quant à φ nous admettons qu'elle est définie dans tout un certain domaine Δ de l'espace $(t_1, t_2, \dots, t_p)$ contenant l'ensemble des points $(f_1, f_2, \dots, f_p)$. Il est alors évidemment possible de prolonger φ en dehors de Δ tout en respectant la continuité, de sorte que l'on pourra raisonner comme si φ était partout définie.

Ceci posé, supposons $f_1, f_2, \dots, f_p$ de classe β et supposons-les respectivement limites des fonctions $f'_1, f'_2, \dots, f'_p$ des classes inférieures

⁽¹⁾ On pourrait remplacer dans cet énoncé *classe α au plus* par *classe inférieure à α* .

En considérant le cas où φ serait de classe β , on serait conduit à la définition de la somme des symboles de classe.

à β . En vertu de la continuité de φ , Φ est évidemment la limite, pour r infini, de $\psi_r = \varphi(f_1^r, f_2^r, \dots, f_p^r)$. Si nous admettons que le théorème est vrai pour $\alpha < \beta$, ψ_r est de classe inférieure à β , donc Φ est de classe β au plus, le théorème est vrai pour $\alpha = \beta$. D'ailleurs, il est évident pour $\alpha = 0$, donc il est vrai pour toutes les classes.

Pour les applications il peut être nécessaire de faire sur φ d'autres hypothèses; par exemple, il peut être nécessaire d'étudier le cas de $\varphi = \frac{t_1}{t_2}$, pour démontrer que le quotient de deux fonctions f_1 et f_2 de classe α au plus est de classe α au plus quand on suppose que la fonction diviseur ne s'annule jamais. C'est à ce cas particulier que je vais me borner.

Je suppose donc que $\varphi(t_1, t_2)$ est définie et continue partout, sauf pour $t_2 = 0$. J'appelle φ_p la fonction continue de (t_1, t_2) égale à φ pour $|t_2| \geq \frac{1}{p}$ et variant linéairement par rapport à t_2 , quand t_2 varie de $-\frac{1}{p}$ à $+\frac{1}{p}$. Alors, f_2 étant différente de zéro par hypothèse, pour que $\Phi = \varphi(f_1, f_2)$ ait un sens, Φ est la limite, pour r infini, des fonctions partout définies, $\psi_r = \varphi_r(f_1^r, f_2^r)$; f_1^r et f_2^r ayant les mêmes significations que précédemment. A l'aide de ces nouvelles fonctions ψ_r on démontrera le théorème comme plus haut ⁽¹⁾.

J'utiliserai aussi le théorème I dans le cas de la fonction composée $\varphi(f_1, f_2)$, f_1 et f_2 ne s'annulant pas en même temps, et φ étant partout définie et continue sauf à l'origine; on légitimera facilement le théorème I dans ce cas et, d'une manière générale, toutes les fois que l'ensemble des points où φ n'est pas définie est fermé.

Voici maintenant une conséquence du théorème I, peu importante par elle-même, mais qui nous sera fort utile. Je considère la fonction $\varphi(t)$ égale à t pour $a \leq t \leq b$, égale à a pour $t \leq a$, égale à b pour $t \geq b$; c'est une fonction continue; donc $\varphi(f)$ est au plus de la classe de f . $\varphi(f)$ est égale à f pour $a \leq f \leq b$, à a pour $f \leq a$, à b pour $f \geq b$; je dirai que $\varphi(f)$ est la fonction f limitée à a et b . Si $a = -\infty$, je

(1) Il est à remarquer que ce raisonnement ne serait plus suffisant s'il s'agissait de fonctions f_1, f_2 non partout définies, car la limite de ψ_r pourrait avoir un sens, sans que Φ en ait un.

dirai que $\varphi(f)$ est la fonction f limitée supérieurement à b ; si $b = +\infty$, je dirai que $\varphi(f)$ est la fonction limitée inférieurement à a .

II. En limitant une fonction à a et b , on n'augmente pas sa classe.

Dans cette opération, l'ensemble des points où f n'a pas changé est l'ensemble des points pour lesquels on a $a \leq f \leq b$, ensemble que je désignerai par $E(a \leq f \leq b)$ ⁽¹⁾. Cette simple remarque fait prévoir l'importance des ensembles $E(a \leq f \leq b)$, c'est à leur étude qu'est consacré le paragraphe suivant. Voici une conséquence de II.

Si f est de classe α , la fonction f limitée inférieurement à b , f_b , la fonction f limitée à a et b , f_a^b , la fonction f limitée supérieurement à a , f_a , sont des fonctions de classe α au plus, et l'une d'elles est effectivement de classe α . Si l'on a $a < 0 < b$, on peut dire la même chose des trois fonctions $\frac{1}{f_b}$, f_a^b , $\frac{1}{f_a}$; mais, de plus, ces trois fonctions sont bornées. La classe de f étant exactement la plus grande des classes de ces trois fonctions, dans l'étude des fonctions de classe α au plus, on pourra, si on le juge utile, se borner à la considération des fonctions bornées ⁽²⁾.

Plus généralement on pourra diviser comme on le voudra l'intervalle de variation de f , c'est-à-dire l'intervalle à une dimension qui a pour extrémités les limites inférieure et supérieure de f , et remplacer l'étude de f par l'étude de fonctions ayant pour intervalles de variation les intervalles partiels arbitrairement choisis.

Soit f une fonction de classe α , ayant l et L pour limites inférieure et supérieure; f est la limite d'une suite de fonctions $f_1, f_2, \dots$, de classe inférieure à α ; en limitant ces fonctions à l et L , je n'augmente pas leur classe, je ne change pas leurs limites; donc, je pourrai toujours

⁽¹⁾ J'emploierai aussi les notations $E(a < f < b)$, $E(f \neq 0)$, ..., sans qu'il soit besoin d'insister sur leurs significations. Pour éviter des confusions possibles, les parenthèses des symboles précédents seront parfois remplacées par des crochets ou des accolades.

⁽²⁾ Je rappelle qu'il est question de fonctions partout définies.

supposer que les f_i sont toutes comprises entre ℓ et L . Cette remarque nous conduit facilement au théorème III.

III. Une suite, uniformément convergente, de fonctions de classe α au plus, a pour limite une fonction de classe α au plus.

Soient $f_1, f_2, \dots$ la suite considérée, f sa limite. On peut, quel que soit ε_p , trouver n_p tel que l'on ait, pour tous les points,

$$|f - f_{n_p}| < \varepsilon_p.$$

J'écris

$$f = f_{n_1} + (f_{n_2} - f_{n_1}) + (f_{n_3} - f_{n_2}) + \dots;$$

c'est une série uniformément convergente de fonctions de classe α au plus. Son terme général $(f_{n_p} - f_{n_{p-1}})$ est, en valeur absolue, inférieur à $\varepsilon_p + \varepsilon_{p-1}$; on peut le considérer comme la limite d'une suite convergente de fonctions $\varphi_p^1, \varphi_p^2, \dots$ des classes inférieures à α , et même supposer que $|\varphi_p^i|$ ne surpasse pas $\varepsilon_p + \varepsilon_{p-1}$. Posons

$$\psi_q = \varphi_1^q + \varphi_2^q + \dots + \varphi_q^q,$$

ψ_q est de classe inférieure à α ; supposons que q augmente indéfiniment, la somme des p premiers termes de ψ_p tend vers f_{n_p} , la somme des autres est moindre, en valeur absolue, que

$$\varepsilon_p + \varepsilon_q + 2(\varepsilon_{p+1} + \varepsilon_{p+2} + \dots + \varepsilon_{q-1});$$

donc, si la série $\varepsilon_1 + \varepsilon_2 + \dots$ a été choisie convergente, ψ_q tend vers f ; le théorème est démontré.

IV. — Classification des ensembles mesurables B.

Nous dirons qu'un ensemble de points est F, de classe α , s'il peut être considéré comme l'ensemble $E(\alpha \leq f \leq b)$, relatif à une fonction f de classe α , et si cela est impossible à l'aide d'une fonction de classe inférieure à α . Nous dirons qu'un ensemble de points est O,

de classe α , s'il peut être considéré comme l'ensemble $E(a < f < b)$ relatif à une fonction f , de classe α , et si cela est impossible à l'aide d'une fonction de classe inférieure à α .

Voici la raison de ces dénominations : pour $\alpha = 0$, les ensembles F sont les ensembles *fermés*, c'est-à-dire ceux qui contiennent leurs dérivés ou encore, si l'on veut, ceux qui contiennent leurs frontières; les ensembles O , de classe 0 , sont les ensembles *ouverts*, c'est-à-dire ceux qui sont les complémentaires des ensembles fermés, ou encore, si l'on veut, ceux qui ne contiennent aucun point de leurs frontières.

Soient $\varphi_1(t)$, $\varphi_2(t)$, $\varphi_3(t)$ des fonctions continues qui sont nulles pour les valeurs suivantes, et pour celles-là seulement :

$$\begin{aligned} \varphi_1(t) &= 0 && \text{pour} && a \leq t \leq b, \\ \varphi_2(t) &= 0 && \text{pour} && t \leq a \quad \text{et} \quad t \geq b, \\ \varphi_3(t) &= 0 && \text{pour} && t = 0, \end{aligned}$$

et supposons de plus que $\varphi_3(t)$ n'est jamais négative. Alors, d'après I, $\varphi_1(f)$, $\varphi_2(f)$, $\varphi_3(f)$ sont de classe α au plus, si f est de classe α .

De l'identité évidente

$$E(a \leq f \leq b) \equiv E[\varphi_1(f) = 0],$$

on déduit que les ensembles F de classe α sont ceux qui peuvent être considérés comme ensemble $E(f = 0)$, où f est de classe α et ne peut être de classe inférieure à α . Des identités

$$\begin{aligned} E(a < f < b) &\equiv E[\varphi_2(f) \neq 0], \\ E(f \neq 0) &\equiv E[0 < \varphi_3(f) < +\infty], \end{aligned}$$

on déduit que les ensembles O de classe α sont ceux qui peuvent être considérés comme ensembles $E(f \neq 0)$, où f est de classe α et ne peut être de classe inférieure à α ⁽¹⁾.

De ces nouvelles définitions il résulte que le complémentaire d'un

⁽¹⁾ On peut même supposer que le module de f est limité comme on le veut, d'après le théorème II.

ensemble F ⁽¹⁾ est un ensemble O de même classe, et inversement. Cela me permettra de ne considérer dans la suite que les ensembles F de classe α que j'appellerai alors simplement *ensembles de classe α* . Dans ce paragraphe, où je vais étudier la formation des ensembles de classe quelconque, je donnerai les propriétés sous les deux formes qu'elles prennent suivant qu'on les énonce en parlant d'ensembles F ou d'ensembles O .

Je désignerai par $\mathcal{E}$ l'ensemble de tous les ensembles F et O précédemment définis, c'est-à-dire l'ensemble de tous les ensembles $E(f=o)$, $E(f \neq o)$, où f est exprimable analytiquement. Je vais étudier la nature des ensembles obtenus en effectuant sur des ensembles de $\mathcal{E}$ les deux opérations suivantes :

L'opération I donne l'ensemble formé des points appartenant à l'un au moins des ensembles d'une famille donnée d'ensembles.

L'ensemble ainsi obtenu s'appelle la somme des ensembles donnés; si les ensembles donnés $E_1, E_2, \dots$ sont en nombre fini ou dénombrable, je représenterai l'ensemble somme par la notation

$$E_1 + E_2 + \dots$$

L'opération II donne l'ensemble formé des points communs à tous les ensembles d'une famille donnée d'ensembles. L'ensemble ainsi obtenu s'appelle la partie commune aux ensembles donnés; si ceux-ci $E_1, E_2, \dots$ sont en nombre fini ou dénombrable, je représenterai la partie commune par la notation $(E_1, E_2, \dots)$ ou $[E_1, E_2, \dots]$ ou $\{E_1, E_2, \dots\}$.

Si C_e désigne le complémentaire de e , on a

$$C(E_1, E_2, \dots) = CE_1 + CE_2 + \dots,$$

l'opération II peut donc être remplacée par l'opération I et *une nouvelle opération III permettant le passage d'un ensemble à son complémentaire*. Cette remarque sera souvent utilisée.

Si $f_1, f_2, \dots, f_n$ sont des fonctions de classe α au plus,

$$\varphi_1 = f_1^2 + f_2^2 + \dots + f_n^2, \quad \varphi_2 = f_1 f_2 \dots f_n$$

(1) C'est-à-dire l'ensemble des points n'en faisant pas partie.

sont de classe α au plus. Des identités

$$\begin{aligned} E(\varphi_1 = 0) &\equiv [E(f_1 = 0), E(f_2 = 0), \dots, E(f_n = 0)], \\ E(\varphi_2 = 0) &= E(f_1 = 0) + E(f_2 = 0) + \dots + E(f_n = 0) \end{aligned}$$

et de la remarque précédente on déduit que :

Les opérations I et II appliquées à un nombre fini d'ensembles F (ou O) de classe α au plus donnent des ensembles F (ou O) de classe α au plus.

Occupons-nous du cas où les opérations I et II sont appliquées à une infinité dénombrable d'ensembles donnés. Soient $E_1, E_2, \dots$ des ensembles F de classe α au plus, on peut trouver f_i de classe α au plus et telle que E_i soit identique à $E(f_i = 0)$, on peut même supposer que la série

$$\varphi_1 = f_1^2 + f_2^2 + \dots + f_n^2 + \dots$$

est uniformément convergente puisqu'on peut assujettir $|f_i|$ à être inférieur à $\frac{1}{i}$. Alors φ_1 est encore de classe α au plus, donc :

La partie commune à une infinité dénombrable d'ensembles F de classe α au plus est F de classe α au plus, ou encore la somme d'une infinité dénombrable d'ensembles O de classe α au plus est O de classe α au plus.

La somme d'une infinité dénombrable d'ensembles F de classe α n'est pas nécessairement un ensemble F de classe α ni même un ensemble F de classe $\alpha + 1$. Par exemple, l'ensemble des valeurs rationnelles de t est la somme d'une infinité dénombrable d'ensembles F de classe 0 formés chacun d'un seul point; ce n'est cependant ni un ensemble F de classe 0, ni un ensemble F de classe 1 ⁽¹⁾. Pour pré-

(1) En effet, une fonction $f(t)$ qui s'annulerait pour les valeurs rationnelles de t et celles-là seulement ne pourrait être continue que pour des valeurs rationnelles de t . Or, si $f(t)$ était de classe 1, d'après un théorème de M. Baire (théorème XV), $f(t)$ aurait, dans tout intervalle, des points de continuité; par suite,

ciser la nature d'une telle somme, je vais démontrer qu'un ensemble F (ou O) de classe α est O (ou F) de classe $\alpha + 1$ au plus.

Soit $E(f=0)$ un ensemble F de classe α , f étant de classe α . Soit $\varphi(t)$ une fonction définie par $\varphi(t) = 0$ pour $t \neq 0$ et $\varphi(0) = 1$. On a $\varphi(t) = \lim_{n \rightarrow \infty} 2^{-nt}$, d'où

$$\varphi(f) = \lim_{n \rightarrow \infty} 2^{-nf},$$

la fonction du second membre étant de classe α au plus, $\varphi(f)$ est de classe $\alpha + 1$ au plus. De là il résulte que les ensembles

$$E(f=0) \equiv E[\varphi(f)=1] \equiv E[\varphi(f) \neq 0],$$

$$E(f \neq 0) \equiv E[\varphi(f) \neq 1] \equiv E[\varphi(f) = 0]$$

sont à la fois F et O de classe $\alpha + 1$ au plus, cela démontre la proposition. Faire la somme d'une infinité dénombrable d'ensembles F de classe α au plus c'est donc faire la somme d'une infinité dénombrable d'ensembles O de classe $\alpha + 1$ au plus, par suite *la somme d'une infinité dénombrable d'ensembles F de classe α au plus est un ensemble O de classe $\alpha + 1$ au plus; donc un ensemble F de classe $\alpha + 2$ au plus.* Ou encore *la partie commune à une infinité dénombrable d'ensembles O de classe α au plus est un ensemble F de classe $\alpha + 1$ au plus, donc un ensemble O de classe $\alpha + 2$ au plus.*

L'ensemble $\mathcal{E}$, qui est formé par la réunion de tous les ensembles F et O , peut donc être aussi considéré comme l'ensemble des ensembles F , ou celui des ensembles O . De plus, si l'on applique l'une ou l'autre des opérations I, II à partir d'un nombre fini ou dénombrable d'ensembles de $\mathcal{E}$ dont les classes sont $\alpha_1, \alpha_2, \dots$, comme on peut tou-

quel que soit ε , l'ensemble des points où $|f(t)|$ surpasserait ε serait partout non dense; l'ensemble des points où $f(t)$ différerait de zéro serait donc de première catégorie (voir p. 184), ce qui est impossible, puisque c'est le complémentaire de l'ensemble des valeurs rationnelles de t , lequel est de première catégorie.

Sous une autre forme, cette remarque a été faite par M. Volterra dans le *Giornale de Battaglini*, 1881 (*Alcune osservazione sulle funzioni punteggiate discontinue*).

jours trouver un symbole α supérieur à tous ces α_i , ce sera appliquer cette opération à des ensembles de classe α au plus, donc cela conduira à un ensemble de $\mathcal{C}$. Ainsi l'application répétée des opérations I et II (auxquelles on peut, si l'on veut, joindre l'opération III) ne permet pas de sortir de l'ensemble $\mathcal{C}$. Ces opérations vont nous permettre de former tous les ensembles de $\mathcal{C}$, mais il est utile pour cela de donner une nouvelle définition.

J'appellerai *ensemble de rang α* tout ensemble qui peut être considéré comme la partie commune à un nombre fini ou à une infinité dénombrable d'ensembles F des classes inférieures à α , cela étant supposé impossible lorsque l'on remplace α par un symbole plus petit. Si α est de première espèce, les ensembles de rang α sont compris parmi les ensembles de classe $\alpha - 1$; examinons le cas où α est de seconde espèce. Soient $f_1, f_2, \dots$ des fonctions des classes inférieures à α et telles que la série

$$\varphi_1 = f_1^2 + f_2^2 + \dots$$

soit uniformément convergente; nous ne pouvons pas affirmer que φ_1 est de classe inférieure à α , mais φ_1 est de classe α au plus. L'ensemble $E(\varphi_1 = 0) \equiv [E(f_1 = 0), E(f_2 = 0), \dots]$ est l'ensemble le plus général de rang α au plus, donc un tel ensemble est toujours F de classe α au plus. On peut aller un peu plus loin; reprenons la fonction $\varphi(t) = \lim_{n \rightarrow \infty} 2^{-nt^t}$ et posons

$$\varphi_2 = \varphi(f_1) \times \varphi(f_2) \times \dots$$

Ce produit infini est convergent et de classe α au plus, car si f_i est de classe $\alpha_i < \alpha$, $\varphi(f_i)$ est de classe $\alpha_{i+1} < \alpha$. L'ensemble $E(\varphi_2 = 1)$, qui est l'ensemble de rang α au plus précédemment considéré, peut aussi être noté $E(\varphi_2 \neq 0)$; donc, si α est de seconde espèce, les ensembles de rang α sont à la fois F et O de classe α ; cela est vrai quel que soit α ⁽¹⁾.

Remarquons encore que les fonctions $1 - \varphi(f_1) \varphi(f_2) \dots \varphi(f_n)$, qui

(1) Puisque, si α est de première espèce, un ensemble de rang α est F de classe $\alpha - 1$.

sont de classes inférieures à α et qui ne décroissent jamais quand n croît constamment, tendent vers la fonction $1 - \varphi_2$ qui ne prend que les valeurs 0 et 1 et qui s'annule pour les points de l'ensemble considéré. Cette propriété, qui suppose α de seconde espèce, sera utile plus tard.

D'après la définition même des ensembles de rang α , que α soit de première ou de seconde espèce, la partie commune à un nombre fini ou à une infinité dénombrable d'ensembles de rang α au plus est un ensemble de rang α au plus. La définition des ensembles de rang α peut d'ailleurs être un peu modifiée; soient $E_1, E_2, \dots$ des ensembles F de classes inférieures à α dont la partie commune est un ensemble c de rang α . Si $e_i = (E_1, E_2, \dots, E_i)$, e_i est, on le sait, F de classe inférieure à α , mais $c = (e_1, e_2, \dots)$; donc on peut, dans la définition des ensembles de rang α , supposer que l'on applique l'opération II à une suite dénombrable d'ensembles de classes inférieures à α , telle que chaque ensemble contienne tous ceux qui le suivent. En d'autres termes, on remplace l'opération II par l'opération II' :

L'opération II' fournit la partie commune à une suite dénombrable d'ensembles donnés $e_1, e_2, \dots$, telle que chaque ensemble contienne tous ceux qui le suivent.

De la définition ainsi modifiée résulte immédiatement que la somme d'un nombre fini d'ensembles de rang α est un ensemble de rang α . D'ailleurs un ensemble de rang α étant toujours O de classe α au plus, la somme d'une infinité d'ensembles de rang α au plus est un ensemble O de classe α au plus (¹).

(¹) Il reste à démontrer que, pour α de seconde espèce, il existe bien des ensembles de rang α ; pour la suite il est même utile de démontrer qu'il existe, pour α de seconde espèce, des ensembles de rang α qui ne sont pas somme d'une infinité dénombrable d'ensembles de classes inférieures à α , c'est-à-dire qu'il existe des ensembles de rang α dont les complémentaires ne sont pas de rang α . Des propriétés qui vont suivre il résultera que, si tout ensemble de rang α avait pour complémentaire un ensemble de rang α , auquel cas les opérations I, II, III appliquées à des ensembles de rang α au plus donneraient des ensembles de rang α au plus, il n'existerait pas d'ensembles de classe supérieure à α , donc pas de fonctions de classe supérieure à α . Or on verra au paragraphe VIII qu'il

Voici maintenant comment on passe des ensembles des classes inférieures à α à ceux de classe α .

Soit f une fonction de classe α limite des fonctions $f_1, f_2, \dots$ de classes inférieures à α . Si l'on pose

$$E_{n,h} = [E(a + h \leq f_n \leq b - h), E(a + h \leq f_{n+1} \leq b - h), \dots],$$

on a

$$\begin{aligned} E(a < f < b) = & E_{1, \frac{1}{4}} + E_{2, \frac{1}{4}} + E_{3, \frac{1}{4}} + \dots \\ & + E_{1, \frac{1}{2}} + E_{2, \frac{1}{2}} + E_{3, \frac{1}{2}} + \dots \\ & + E_{1, \frac{3}{4}} + E_{2, \frac{3}{4}} + E_{3, \frac{3}{4}} + \dots \end{aligned}$$

On obtient donc tout ensemble O de classe α en effectuant, dans l'ordre A, B , les deux opérations :

A. On prend la partie commune à une infinité dénombrable d'ensembles F de classes inférieures à α ; cela donne les ensembles de rang α au plus.

B. On fait la somme d'une infinité dénombrable d'ensembles de rang α au plus; cela donne, nous le savons, des ensembles O de classe α au plus.

Ainsi les deux opérations A, B donnent tous les ensembles O de classe α et ne donnent que des ensembles O de classe α au plus. D'ailleurs, l'opération A n'est utile que si α est de seconde espèce.

Pour avoir les ensembles F de classe α , il faut effectuer, après A et B , l'opération C :

C. On prend les complémentaires des ensembles O de classe α au plus; ce qui donne les ensembles F de classe α au plus.

On peut, en se servant du passage d'un ensemble à son complémen-

existe des fonctions de toute classe, donc il existe des ensembles de rang α jouissant de la propriété indiquée.

On pourrait d'ailleurs nommer un ensemble de rang α en employant des procédés analogues à ceux du paragraphe VIII.

taire, remplacer ces opérations par d'autres de bien des manières possibles, j'en cite un seul exemple.

Pour passer des ensembles O des classes inférieures à α aux ensembles O de la classe α au plus, il suffit d'effectuer les trois opérations A_1 , B_1 , C_1 .

A_1 . On fait la somme d'une infinité dénombrable d'ensembles O de classes inférieures à α .

B_1 . On prend la partie commune à toutes les suites dénombrables d'ensembles formées avec les ensembles fournis par A_1 .

C_1 . On prend les complémentaires des ensembles que donne B_1 .

A_1 n'est nécessaire que si α est de seconde espèce; A_1 et B_1 fournissent les ensembles F de classe α au plus.

Avant d'aller plus loin supposons, pour un instant, qu'on se soit servi dans la classification des ensembles de toutes les fonctions, partout définies ou non. Sans raisonnements nouveaux nous pouvons affirmer que les opérations A , B , C , ou A_1 , B_1 , C_1 , appliquées aux ensembles F , ou O , des classes inférieures à α , nous auraient encore fait connaître ceux de classes α ; seulement, nous ne pouvons pas affirmer sans raisonnement qu'elles ne nous auraient donné que ces ensembles. Mais cela est presque évident; les ensembles de classe α étant les mêmes dans les deux classifications, de ce qui précède il résulte que l'ensemble des ensembles de classe α au plus relatif à la seconde classification est contenu dans celui relatif à la première; or le passage de la première à la seconde classification ne peut qu'étendre le contenu de l'ensemble des ensembles de classe α au plus, quel que soit α ; donc les deux classifications donnent les mêmes résultats.

Les opérations C et C_1 ne sont que l'opération III appliquée à des ensembles particuliers, les opérations B et A_1 ne diffèrent pas de I, B_1 et A ne diffèrent pas de II. Puisqu'on peut se passer des opérations C et C_1 , on voit que l'on peut former, par l'application répétée de I et II, tous les ensembles de $\mathcal{C}$ à partir des ensembles fermés et ouverts, et même on peut toujours remplacer II par II'.

Remarquons que tout ensemble O de classe α étant O de classe inférieure à $\alpha + 1$ est la somme d'une infinité dénombrable d'ensembles F de classe α , donc les ensembles ouverts sont sommes d'une infinité dénombrable d'ensembles fermés, ce qui d'ailleurs se prouve immé-

diatement de bien des manières. Soit maintenant E un ensemble fermé, j'appelle E_p l'ensemble somme de ceux des intervalles

$$\frac{a_1}{p} \leq x_1 \leq \frac{a_1+1}{p}, \quad \frac{a_2}{p} \leq x_2 \leq \frac{a_2+1}{p}, \quad \dots, \quad \frac{a_n}{p} \leq x_n \leq \frac{a_n+1}{p},$$

où les a_i sont entiers, qui contiennent des points de E. On a évidemment

$$E = (E_1, E_{\frac{1}{2}}, E_{\frac{1}{3}}, \dots);$$

l'opération indiquée dans le second membre est l'opération II'.

Ainsi, *ε est formé des ensembles que l'on obtient par l'application répétée des opérations I et II (ou I et II') à partir d'intervalles*. Les ensembles ainsi construits sont ceux que j'ai déjà eu l'occasion de considérer ailleurs et que j'ai appelés les ensembles mesurables B parce que ce sont ceux qu'on peut mesurer par les procédés donnés par M. Borel ⁽¹⁾.

J'indiquerai dans le paragraphe suivant quelle importance il y aurait à savoir résoudre cette question : un ensemble étant donné, reconnaître s'il est ou non mesurable B et quelle est sa classe? Je n'aborderai pas cette question, je me contenterai de montrer comment parfois une telle recherche peut être faite ⁽²⁾.

(1) Voir ma Thèse et mes *Leçons sur l'intégration*. J'abandonne ici une restriction que j'avais adoptée aux endroits indiqués, savoir que les opérations I et II ne sont appliquées qu'un nombre fini de fois. Cette restriction n'est jamais intervenue dans mes raisonnements, je l'avais indiquée parce qu'elle conduit à une famille d'ensembles plus voisine, à ce qu'il m'a semblé, de celle que considère M. Borel dans ses *Leçons sur la Théorie des fonctions*.

De ce qui sera démontré dans la suite, il résulte que les ensembles mesurables B sont ceux qui peuvent être définis par des égalités ou inégalités analytiques; pour cette raison ils mériteraient d'être nommés *ensembles analytiques*.

(2) Il est évident qu'il serait illusoire de chercher à répondre à la question posée sous la forme trop générale que je lui ai donnée; mais, examinant successivement les différents procédés que l'on emploie ordinairement pour nommer des ensembles, il faudrait donner pour chacun d'eux un moyen de faire l'étude en question (comparer avec la page 183). Au point de vue auquel je me suis

Un premier procédé souvent commode résulte de l'emploi de la définition même des ensembles mesurables B. Si l'on démontre qu'un ensemble peut être construit à partir d'intervalles, à l'aide des opérations I et II, on sait qu'il est mesurable B et l'on a une limite supérieure de sa classe en tant qu'ensemble F ou O. J'ai donné un exemple de cette méthode dans mes *Leçons sur l'intégration et la recherche des fonctions primitives* aux pages 109 et 110.

Un autre procédé, donnant les mêmes renseignements et qui résulte de ce qui précède, est de démontrer que l'ensemble donné est un ensemble F ou O attaché à une fonction f de classe déterminée (¹).

En utilisant des résultats ultérieurs on pourrait même dire qu'il suffit que f soit représentable analytiquement ou soit définie analytiquement. De l'application de ce second procédé résulte donc immédiatement que tout ensemble qui est défini par un nombre fini ou dénombrable d'égalités ou d'inégalités entre fonctions exprimables analytiquement à l'aide des coordonnées, est mesurable B.

V. — Étude, sur certains ensembles de points, des fonctions de classe donnée.

J'appelle *fonction mesurable B* toute fonction f telle que, quels que soient a et b , l'ensemble $E(a \leq f \leq b)$ soit mesurable B. D'après ce qui précède, on pourrait tout aussi bien dire que f est mesurable B si, quels que soient a et b , $E(a < f < b)$ est mesurable B ou encore si, quel que soit a , $E(a < f)$ est mesurable B.

placé, c'est cette étude qu'il importerait de faire pour que les résultats du texte aient une valeur pratique.

Les méthodes que je viens d'indiquer, pour reconnaître qu'un ensemble donné est mesurable B, semblent s'appliquer à tout ensemble donné par l'un des procédés classiques; c'est pourquoi j'ai dit que tous les ensembles que l'on considère ordinairement sont mesurables B, que toutes les fonctions considérées ordinairement sont représentables analytiquement.

(¹) Il faut remarquer qu'il n'y a pas cercle vicieux à employer ce second procédé, même si le but final est de reconnaître si une fonction φ , auquel l'ensemble donné est attaché comme F ou O, est représentable analytiquement. Car on peut prendre f différente de φ .

On peut aussi remarquer qu'il suffit que tous les $E(r_1 \leq f \leq r_2)$, où les r_1 et r_2 sont rationnels, soient mesurables B pour que f le soit. En effet, on peut prendre des nombres rationnels r_1^p, r_2^p , tendant vers a et b quand p augmente indéfiniment et tels que l'on ait

$$r_1^p \leq a < b \leq r_2^p,$$

a et b étant des nombres donnés. Alors $E(a \leq f \leq b)$ est la partie commune à tous les ensembles $E(r_1^p \leq f \leq r_2^p)$, lesquels sont mesurables B, donc $E(a \leq f \leq b)$ est mesurable B.

Cette propriété nous montre que, si l'on savait reconnaître si un ensemble est, ou non, mesurable B, on reconnaîtrait, par une infinité dénombrable d'opérations, si une fonction est, ou non, mesurable B.

Reprenons les $E(r_1 \leq f \leq r_2)$; s'ils sont tous mesurables B, on peut tous les considérer comme des ensembles de classes déterminées (¹). Mais ces classes sont en nombre fini ou dénombrable; par suite, on peut citer un symbole α tel que tous les $E(r_1 \leq f \leq r_2)$ soient de classe α au plus. Mais $E(a \leq f \leq b)$ étant la partie commune à certains

$$E(r_1 \leq f \leq r_2)$$

est aussi de classe α . Par suite, si f est mesurable B, on peut citer un symbole α tel que tous les $E(a \leq f \leq b)$ soient de classe α au plus. Ce nombre α va servir à caractériser la nature de f , lorsque f est partout définie.

IV. *Pour qu'une fonction f partout définie soit de classe α , il faut et il suffit que, quels que soient a et b , l'ensemble $E(a \leq f \leq b)$ soit de classe α au plus, et qu'il soit effectivement de classe α pour certaines valeurs de a et b (²).*

La condition est évidemment nécessaire, c'est la définition même des ensembles de classe α .

(¹) Sauf indications contraires, je dirai désormais *ensemble de classe α* au lieu de *ensemble F de classe α* .

(²) Cette proposition n'est pas vraie si les fonctions sont partout définies; pour ces fonctions, il y a lieu de tenir compte de la nature de l'ensemble des points en lesquels la fonction est définie.

de multiplications et de passages à la limite, la démonstration résulte des trois propriétés a , b , c qui suivent.

a . La somme de deux fonctions mesurables B est une fonction mesurable B.

Soient f_1 et f_2 deux fonctions mesurables B; quels que soient a et r l'ensemble $[E(f_1 > r), E(f_2 > a - r)]$ est mesurable B. Or la somme de ceux de ces ensembles qui correspondent aux valeurs rationnelles de r est l'ensemble $E(f_1 + f_2 > a)$; $f_1 + f_2$ est donc mesurable B.

b . Le produit de deux fonctions mesurables B est une fonction mesurable B.

La démonstration est analogue à celle de a .

c . La limite d'une suite de fonctions mesurables B est une fonction mesurable B.

Soit f la limite de la suite $f_1, f_2, \dots$; on ne suppose pas que les f_i soient partout définies ni soient convergentes partout dans l'ensemble où elles sont toutes définies. En posant

$$E_n = [E(a < f_n < b), E(a < f_{n+1} < b), \dots],$$

on a

$$E(a < f < b) = E_1 + E_2 + \dots;$$

donc, f est mesurable B.

La condition énoncée est donc nécessaire; j'en déduis que l'ensemble des points sur lequel existe une fonction f exprimable analytiquement est mesurable B. En effet, l'ensemble considéré est la somme de tous les $E(-n \leq f \leq n)$, où n est entier, et tous ces ensembles sont mesurables B (¹). L'ensemble des points en lesquels f n'a pas de sens est donc mesurable B.

La condition énoncée est suffisante. Cela résulte de IV et V pour les fonctions partout définies. Si f n'est pas partout définie mais est mesurable B, la fonction f_p , égale à f quand f existe et à p quand f

(¹) En particulier, les ensembles de points de convergence des suites de fonctions continues, qui ont été parfois employés pour citer des exemples de certaines particularités, rentrent tous dans la famille C des ensembles mesurables B.

La question posée incidemment par M. Borel, dans la Note 1 de la page 67 de ses *Leçons sur la Théorie des fonctions*, est donc résolue affirmativement.

n'existe pas, est mesurable B, donc représentable analytiquement; et il en est de même de f limite des f_p .

Le théorème est donc démontré; mais, en se reportant à la démonstration, on voit de plus que *toute fonction représentable analytiquement fait partie de E, c'est-à-dire rentre dans la classification de M. Baire.*

Voici de nouvelles définitions (¹). Je dirai qu'une fonction f est, à ε près, représentable analytiquement lorsqu'il existe une fonction φ représentable analytiquement et telle que $|f - \varphi|$ ne surpasse jamais ε . Si φ peut être choisie de classe α et pas de classe inférieure, je dirai que f est, à ε près, de classe α ; si, la fonction φ étant toujours supposée définie dans tout le domaine dont on s'occupe, on a

$$|f - \varphi| \leq \varepsilon$$

pour les points d'un ensemble E, je dirai que f est, à ε près, de classe α sur E. Ces définitions posées, j'énonce une proposition à laquelle on est conduit en remarquant que la fonction F, employée dans la démonstration du théorème IV, est encore de classe α au plus si les a_i ne sont plus des constantes, mais sont des fonctions de classe α au plus (²).

VII. *Pour qu'une fonction f soit de classe α au plus, il faut et il suffit que, quel que soit le nombre positif ε , on puisse considérer le domaine où f est définie comme la somme d'un nombre fini d'ensembles de classe α au plus sur chacun desquels f est, à ε près, de classe α au plus.*

Démontrons que la condition est suffisante. Soient $\varphi_1, \varphi_2, \dots, \varphi_n$ des fonctions de classe α au plus. Supposons que, sur $E(\varphi_i = 0)$, f ne

(¹) Bien qu'il n'y ait pas de difficultés sérieuses à considérer toutes les fonctions, sauf avis contraire, je ne m'occuperai que des fonctions partout définies, conformément à ce que j'ai dit dans le paragraphe II.

(²) Cette remarque ne suffit pas pour démontrer le théorème VII parce que la démonstration du théorème IV suppose que trois des fonctions φ_i ne peuvent s'annuler en même temps et la condition analogue n'est plus remplie dans le cas du théorème VII.

On peut cependant démontrer ce dernier théorème par une méthode analogue à celle employée pour le théorème IV, comme on le verra facilement.

diffère de la fonction a_i de classe α au plus, que de ε au plus. L'ensemble $E(\varepsilon)$,

$$E(\varepsilon) = \sum_{i=1}^{i=n} [E(\alpha - \varepsilon \leq a_i \leq \beta + \varepsilon), E(\varphi_i = 0)],$$

est de classe α au plus. $E(\alpha \leq f \leq \beta)$, qui est la partie commune à tous les $E\left(\frac{1}{n}\right)$, où n est entier, est donc de classe α au plus et le théorème est démontré.

On peut remarquer que, si, au lieu d'un nombre *fini* d'ensembles $E(\varphi_i = 0)$, nous en avons eu une infinité dénombrable, l'ensemble $E(\varepsilon)$, somme d'une infinité dénombrable d'ensembles de classe α , aurait été de classe $\alpha + 2$. De cela on déduit que f aurait été de classe $\alpha + 2$ au plus. Ceci nous fait voir comment nous pouvons espérer établir une relation entre les fonctions de classe α et les fonctions des classes inférieures à α . En reprenant par une autre méthode le cas où les $E(\varphi_i = 0)$ forment une infinité dénombrable, nous allons démontrer que :

Si, quel que soit ε , on peut considérer l'ensemble dans lequel une fonction f est définie comme la somme d'une infinité dénombrable d'ensembles des classes inférieures à α , sur chacun desquels f est, à ε près, de classe inférieure à α , alors f est de classe α au plus.

Soient $\varphi_1, \varphi_2, \dots$ des fonctions de classe inférieure à α . Supposons que, sur $E_i = E(\varphi_i) = 0$, f ne diffère que de ε au plus de la fonction a_i de classe inférieure à α . Soient, d'autre part, $\Psi_n(t)$ des fonctions continues nulles pour $t = 0$ et telles que, pour $t \neq 0$, les fonctions $\Psi_n(t)$ tendent vers 1 quand n augmente indéfiniment; par exemple

$$\Psi_n(t) = 1 - e^{-nt}.$$

Formons les fonctions

$$F_1 = a_1,$$

$$F_2 = a_1 + (a_2 - a_1)\Psi_1(\varphi_1),$$

$$F_3 = a_1 + (a_2 - a_1)\Psi_2(\varphi_1)$$

$$+ [a_3 - [a_1 + (a_2 - a_1)\Psi_2(\varphi_1)]]\Psi_2(\varphi_1, \varphi_2),$$

et, d'une manière générale,

$$F_{p+1} = a_1 + (a_2 - s_1) \Psi_p(\varphi_1) \\ + (a_3 - s_2) \Psi_p(\varphi_1, \varphi_2) + \dots + (a_{p+1} - s_p) \Psi_p(\varphi_1, \varphi_2, \dots, \varphi_p),$$

où s_i désigne la somme des i premiers termes de F_{p+1} . Ce sont des fonctions de classe inférieure à α ; elles tendent, quand leur indice augmente indéfiniment, vers la fonction F égale à a_1 sur E_1 , égale à a_2 sur $E_2 - (E_1, E_2)$, égale à a_3 sur $E_3 - (E_1 + E_2, E_3)$, et, d'une manière générale, égale à a_p sur $E_p - (E_1 + E_2 + \dots + E_{p-1}, E_p)$. F est donc de classe α au plus et, puisque f ne diffère de F que de ε au plus, f est au plus de classe α .

Recherchons si la réciproque est vraie. Soit f une fonction de classe α limite de fonctions $f_1, f_2, \dots$ de classe inférieure à α . Le domaine où f est définie est la somme des ensembles E_n ,

$$E_n = \{E[(f_n - f_{n+1})^2 \leq \varepsilon^2], E[(f_n - f_{n+2})^2 \leq \varepsilon^2], E[(f_n - f_{n+3})^2 \leq \varepsilon^2], \dots\},$$

puisque la suite des f_n est convergente dans ce domaine. E_n , étant la partie commune à une infinité dénombrable d'ensembles de classe inférieure à α , est au plus de rang α . Le domaine considéré est donc la somme d'une infinité dénombrable d'ensembles de rang α au plus sur chacun desquels f est, à ε près, de classe inférieure à α , puisque, dans E_n , on a

$$|f - f_n| \leq \varepsilon.$$

Mais, si α est de première espèce, E_n est au plus de classe $\alpha - 1$, donc :

VIII. *Pour qu'une fonction f soit de classe $\alpha + 1$, il faut et il suffit que, quel que soit ε , le domaine où f est définie puisse être considéré comme la somme d'une infinité dénombrable d'ensembles de classe α au plus sur chacun desquels f est, à ε près, de classe α au plus.*

En rapprochant cette proposition du théorème IV, on voit que :

IX. *Pour qu'une fonction f soit de classe $\alpha + 1$, il faut et il suffit que, quel que soit ε , le domaine où f est définie puisse être considéré comme la somme d'une infinité dénombrable d'ensembles de*

classe α au plus sur chacun desquels f est d'oscillation au plus égale à ε (').

Voici, maintenant, un énoncé qui convient à tous les cas :

X. *Pour qu'une fonction f soit de classe $\alpha > 0$ il faut et il suffit que, quel que soit ε , le domaine où f est définie puisse être considéré comme la somme d'une infinité dénombrable d'ensembles de rang α au plus sur chacun desquels f est, à ε près, de classe inférieure à α ; ou encore sur chacun desquels f est d'oscillation au plus égale à ε .*

Il suffit de démontrer le premier de ces deux énoncés. Il est déjà démontré si α est de première espèce; on sait de plus que la condition énoncée est nécessaire quel que soit α . Il reste à démontrer que, dans le cas où α est de seconde espèce, cette condition est suffisante.

Soient $\varphi_1, \varphi_2, \dots$ des fonctions de classe α au plus, ne prenant que les valeurs 0 et 1 et définissant des ensembles $E(\varphi_1 = 0)$, $E(\varphi_2 = 0), \dots$, de rang α au plus. On sait que φ_i est la limite d'une suite convergente de fonctions ψ_i^p de classe inférieure à α , ne prenant que les valeurs 0 et 1 et ne décroissant jamais avec p (p. 162).

Supposons que, sur $E(\varphi_i = 0)$, f ne diffère que de ε au plus de la fonction a_i de classe inférieure à α . Formons la fonction F_{p+1} ,

$$F_{p+1} = a_1 + (a_2 - s_1)\psi_1^p + (a_3 - s_2)\psi_1^p\psi_2^p + \dots + (a_{p+1} - s_p)\psi_1^p\psi_2^p\dots\psi_p^p,$$

où s_i désigne la somme des i premiers termes de F_p . On voit, comme précédemment, que les fonctions F_p , de classe inférieure à α , tendent vers une limite F , qui est par suite au plus de classe α , et qui ne diffère de f que de ε au plus.

(') En remplaçant dans VIII et IX classe $\alpha + 1$ par classe α et classe α au plus par classe inférieure à α , on obtient deux énoncés équivalents qui sont évidemment exacts quand α est de première espèce, puisque alors ces énoncés ne diffèrent pas des théorèmes VIII et IX, mais qui sont tous deux inexacts quand α est de seconde espèce.

En les appliquant, en effet, à une fonction φ de classe α ne prenant que les valeurs 0 et 1 et définissant un ensemble $E(\varphi = 0)$ de rang α , on en déduirait que tout ensemble de rang α est la somme d'une infinité dénombrable d'ensembles de classe inférieure à α . Or, on sait que cela est inexact quand α est de seconde espèce. Voir la note de la page 162.

Le théorème est donc démontré (¹).

Aux énoncés qui précèdent on peut évidemment ajouter le suivant, qui ne nous apprend rien de nouveau :

XI. Pour qu'une fonction f soit exprimable analytiquement il faut et il suffit que, quel que soit ε , l'ensemble dans lequel f est définie puisse être considéré comme la somme d'une infinité dénombrable d'ensembles mesurables B sur chacun desquels f est, à ε près, représentable analytiquement.

VI. — Étude en certains points des fonctions de classe donnée.

On vient de voir comment on peut déduire la classe d'une fonction de la nature de cette fonction sur certains ensembles; nous allons maintenant nous placer à un point de vue différent et rechercher comment on peut déduire cette classe de la nature de la fonction en certains points.

Je dirai qu'une fonction est, à ε près, de classe α en un point M s'il existe un intervalle contenant M à son intérieur et dans lequel la fonction est, à ε près, de classe α au plus, cela étant impossible si l'on remplace α par $\beta < \alpha$. Au point M , on peut toujours attacher un nombre positif ou nul, fini ou infini, η_α , tel que la fonction considérée soit, en M , de classe α au plus, à ε près, quand on a $\varepsilon > \eta_\alpha$ et tel que cela ne soit plus vrai pour $\varepsilon < \eta_\alpha$.

Pour $\alpha = 0$, $2\eta_0$ est l'oscillation en M ; de sorte que, pour $\eta_0 = 0$, la

(¹) Il est à remarquer que le seul cas où les théorèmes VIII et IX ne sont pas équivalents au théorème X est celui où l'on sait former le moins facilement les ensembles de classe α à partir des ensembles de classe inférieure.

Il faut aussi remarquer que le théorème V présentait, à un certain point de vue, un avantage que n'ont pas les théorèmes VIII, IX, X : celui de fournir un procédé opératoire régulier pour connaître si une fonction est, ou non, de classe α . On verra plus loin ce qu'il faut penser, à mon avis, de ce genre d'avantage (p. 183).

La démonstration des théorèmes qui précèdent aurait pu être simplifiée si l'on avait employé les résultats du paragraphe IV relativement à la formation des ensembles de classe α , à partir de ceux des classes inférieures.

fonction est continue en M . Je dirai qu'une fonction est de classe α en M si, en ce point, $\eta_\alpha = 0$. Cela revient à dire que, quel que soit $\varepsilon > 0$, on peut trouver une fonction φ de classe α au plus qui, dans un certain intervalle contenant M à son intérieur, diffère de la fonction donnée de ε au plus, cela étant impossible pour $\beta < \alpha$ ⁽¹⁾. Dire qu'une fonction est de classe 0 en un point, c'est dire qu'elle est continue en ce point.

Si les conditions précédentes sont réalisées quand on ne s'occupe que des valeurs de la fonction f donnée pour les points d'un ensemble E , nous dirons, suivant les cas, que f est, à ε près, de classe α en M sur E , ou que f est de classe α en M sur E . Cela ne suppose pas que M appartienne à E , mais la définition ne présente quelque intérêt que si M est l'un des points limites de E .

Si, à ε donné, on peut faire correspondre α tel que f soit, à ε près, de classe α en M , nous dirons que f est, à ε près, représentable analytiquement en M ; si, quel que soit $\varepsilon > 0$, f est, à ε près, représentable analytiquement en M , nous dirons que f est représentable analytiquement en M ⁽²⁾.

Considérons la fonction $f(x)$ définie dans $(-1, +1)$ comme étant nulle pour x irrationnel ou nul, comme étant égale à 1 pour $|x| = \frac{1}{2^p}$, où p est entier positif ou nul, comme étant égale à $\frac{1}{p}$ pour x rationnel et tel que $\frac{1}{2^{p+1}} < |x| < \frac{1}{2^p}$; on verrait facilement que $f(x)$ est de classe 2.

$f(x)$ est de classe 2 pour $x \neq 0$; elle est de classe 1 à l'origine, car la fonction $\varphi(x)$, égale à $f(x)$ pour $f(x) = 1$, et à zéro aux autres points, est de classe 1 et, quel que soit ε , on peut trouver $(-a, +a)$

⁽¹⁾ Dans cette définition, φ varie avec ε , on verra facilement qu'on peut prendre φ indépendamment de ε , dès que M est donné, mais l'intervalle à considérer varie en général avec ε .

⁽²⁾ On aurait pu rattacher cette définition à celle d'un certain nombre η analogue à η_α . En appelant oscillation analytique en M le nombre 2η et oscillation de classe α en M le nombre $2\eta_\alpha$, on aurait pu donner, aux énoncés qui suivent, des formes simples qui, peut-être, auraient été préférables à celles que j'ai adoptées.

dans lequel on a

$$0 \leq f(x) - \varphi(x) \leq \varepsilon.$$

A l'origine $f(x)$ est de classe 0 sur l'ensemble des points $\pm \frac{1}{2^p}$, de classe 1 sur l'ensemble des points à abscisses rationnelles, de classe 0 sur l'ensemble des points à abscisses irrationnelles.

Je vais maintenant démontrer un théorème qui généralise la propriété bien connue relative à la continuité uniforme. J'aurai besoin, pour cela, d'une proposition très simple et très importante de M. Borel : *Si l'on a une famille de domaines Δ tels que tout point d'un certain domaine D, y compris les points frontières de D, soit intérieur à l'un au moins des Δ , il existe une famille formée d'un nombre fini de domaines choisis parmi les domaines Δ et qui jouissent de la même propriété (tout point de D est intérieur à l'un d'eux) (1).*

(1) M. Borel a énoncé ce théorème pour le cas où l'ensemble des Δ est dénombrable; il a donné pour ce cas deux démonstrations simples qui s'étendent aux domaines à un nombre quelconque de dimensions. [Voir la Thèse de M. BOREL, *Sur quelques points de la Théorie des fonctions* (Annales de l'École normale, 1895) et ses *Leçons sur la Théorie des fonctions*; voir aussi *Contribution à l'analyse arithmétique du continu* (Journal de Liouville, 1903).]

Pour les applications du texte il est nécessaire que le théorème soit démontré pour le cas où l'ensemble des Δ n'est pas nécessairement dénombrable; on arrive à cette démonstration en modifiant légèrement celle que M. Borel avait donnée dans sa Thèse. Voir mes *Leçons sur l'intégration et la recherche des fonctions primitives*, p. 105 et 117. La démonstration ainsi modifiée ne donne pas un procédé régulier permettant, par une infinité dénombrable d'opérations, d'effectuer parmi les Δ le choix dont elle prouve la possibilité; mais le procédé opératoire indiqué par M. Borel pour le cas des domaines à une dimension reste toujours applicable à ce cas, quelle que soit la puissance de l'ensemble des Δ , pourvu qu'à chaque point de D on sache associer un Δ le contenant. D'ailleurs, d'après la démonstration même (*loc. cit.*, p. 117), le cas d'un domaine à plusieurs dimensions se ramène au cas de la droite, grâce à l'emploi d'une courbe remplissant tout le domaine D. On peut donc, par une infinité dénombrable d'opérations effectuées suivant un procédé régulier, faire le choix qui nous occupe toutes les fois qu'à chaque point de D on sait faire correspondre un Δ le contenant, et cette correspondance sera établie, dans les applications qui suivent, par la définition même des Δ .

De là résulte que, *si une propriété est vraie pour la somme d'un nombre fini de domaines dès qu'elle est vraie de chacun d'eux, elle est vraie pour un domaine D dès qu'autour de chaque point de D existe un domaine dans lequel elle est vraie.*

Pour appliquer cette remarque considérons une fonction f de classe α au plus ($\alpha \geq 1$), à ε près, dans chacun des domaines $D_1, D_2, \dots, D_n$. Cela veut dire qu'il existe une fonction f_i de classe α au plus et différant de f de ε au plus dans D_i . La fonction φ , égale à f_1 dans D_1 , égale à f_2 dans la partie de D_2 extérieure à D_1 , égale à f_3 dans la partie de D_3 extérieure à $D_1 + D_2$, etc., est de classe α au plus d'après le théorème VII, car un domaine est un ensemble de classe 0 et la partie d'un domaine extérieure à la somme d'un nombre fini de domaines est un ensemble de classe 0 ou 1. La fonction f qui diffère de φ de ε au plus est donc, d'après VII, de classe α au plus, à ε près, sur

$$D_1 + D_2 + \dots + D_n \quad (').$$

Notre remarque nous permet de conclure que, *si une fonction est, à ε près, de classe α au plus en tout point d'un domaine D, elle est, à ε près, de classe α au plus dans D.*

Cet énoncé n'est justifié par ce qui précède que pour $\alpha \neq 0$, il est cependant exact pour $\alpha = 0$. Voici comment on peut l'établir. Les conditions de l'énoncé étant remplies, chaque point de D est intérieur à un intervalle Δ dans lequel f est continue à ε près. Prenons un nombre fini de Δ de manière que tout point de D soit *intérieur* à l'un d'eux.

A chaque Δ est attachée une fonction f_Δ continue, différant de f de ε au plus. Soit P un point, il appartient à un ou plusieurs Δ , d'où un ou plusieurs nombres $f_\Delta(P)$. Soit $l(P)$ le plus petit, $L(P)$ le plus grand. On peut toujours construire une fonction continue φ telle que $\varphi(P)$ soit compris entre $l(P)$ et $L(P)$; cela est particulièrement évident pour $n = 1$, mais se voit facilement dans tous les cas.

(') Cette propriété n'est plus vraie pour $\alpha = 0$.

On a évidemment $|f - \varphi| \leq \epsilon$ dans D , la propriété est démontrée (¹).

XII. — *Pour qu'une fonction soit de classe α au plus dans un domaine D il faut et il suffit qu'elle soit de classe α au plus en tout point de D (²).*

La condition évidemment nécessaire est suffisante. Si elle est remplie, en effet, quel que soit ϵ , la fonction est, d'après la proposition précédente, à ϵ près, de classe α au plus et, d'après le théorème III, cela suffit pour que la fonction soit de la classe α au plus.

Ce théorème indique une relation entre la nature d'une fonction dans un domaine et la nature de cette fonction en chaque point du domaine; en ce sens il est à rapprocher du théorème sur la continuité uniforme. Je veux faire une remarque à ce sujet.

L'énoncé XII est pris souvent, pour $\alpha = 0$, comme définition des fonctions continues.

(¹) On peut, en particulier, affirmer que l'on est dans les conditions de l'énoncé si, en tout point de D , f a une oscillation inférieure à 2ϵ .

Dans ce cas on aurait pu choisir les Δ de façon que, dans chacun d'eux, f soit constante à ϵ près. Alors deux points d'un même Δ ou de deux Δ ayant des points frontières en commun auraient donné des valeurs de f différant entre elles de 2ϵ au plus. Or, on peut citer un nombre λ tel que deux points distants de λ au plus soient dans un même δ ou dans deux δ consécutifs, donc : *Si η , est un nombre supérieur au maximum de l'oscillation de f en un point quelconque de D , il existe un nombre positif λ tel que, si deux points P et Q sont à une distance au plus égale à λ l'un de l'autre, on a*

$$|f(P) - f(Q)| \leq \eta.$$

Ce théorème, qui comprend comme cas particulier le théorème sur la convergence uniforme, a été donné par M. Baire à la page 15 de sa Thèse. On peut en donner une démonstration beaucoup plus simple que la précédente (voir *Leçons sur l'intégration et la recherche des fonctions primitives*, p. 31).

(²) D'après nos conventions les points frontières de D sont des points de D . Pour $\alpha > 0$ on peut dire qu'il suffit que la fonction soit de classe α au plus en tout point intérieur à D et de classe α au plus sur la frontière de D en tout point de cette frontière. Cette dernière condition est remplie d'elle-même s'il s'agit d'une fonction à une seule variable, car la frontière de D se compose alors de deux points.

Si les fonctions continues sont définies par la continuité uniforme, il exprime au contraire une propriété très importante des fonctions continues.

Mais, ainsi entendu, le théorème XII n'est pas démontré par les considérations du texte ⁽¹⁾. Les différences qui se manifestent ainsi entre le cas $\alpha = 0$ et le cas $\alpha > 0$ proviennent de ce que les fonctions continues ne sont pas, comme les autres, définies comme limites de fonctions des classes antérieures. Les fonctions continues sont supposées antérieurement définies et antérieurement étudiées; d'ailleurs, sauf peut-être en ce qui concerne la représentation d'une fonction continue par une série de polynômes, la continuité uniforme des fonctions continues n'est guère intervenue dans nos considérations.

Des raisonnements analogues à ceux qui ont donné le théorème XII montrent que :

XIII. Pour qu'une fonction soit représentable analytiquement dans un domaine D, il faut et il suffit qu'elle soit représentable analytiquement en tout point de D.

Voici maintenant une conséquence importante du théorème XIII.

Soit f une fonction qui n'est ni de la classe α , ni d'une classe antérieure. Dès que ϵ est assez petit, il existe des points en lesquels f n'est pas de classe α au plus, à ϵ près; étudions l'ensemble de ces points.

1° Il est évidemment fermé;

2° Pour $\alpha > 0$, il ne contient aucun point isolé. Soit, en effet, une fonction f qui est, à ϵ près, de classe α au plus en tous les points distants d'un point M de $\frac{1}{n_0}$ au plus. Il faut démontrer que f est aussi de classe α au plus, à ϵ près, au point M.

Cela peut se démontrer facilement, sans faire appel aux théorèmes établis dans le paragraphe V, mais, en utilisant ces théorèmes, on peut raisonner un peu plus simplement.

Soit E_n l'ensemble des points dont la distance δ à M vérifie la double

(1) Voir la note 1, page 178, où il est démontré par des raisonnements assez analogues à ceux du texte.

inégalité $\frac{1}{n} \geq \delta \geq \frac{1}{n+1}$. E_n peut toujours être considéré comme la somme de deux domaines; en chaque point de E_n , ($n \geq 0$), f est, à ε près, de classe α au plus, donc f est, à ε près, de classe α au plus sur E_n . Par suite, E_n est la somme d'une infinité dénombrable d'ensembles e_n de rang α au plus sur chacun desquels f est, à ε près, de classe inférieure à α . L'ensemble formé des e_n , pour $n \geq n_0$, et du point M est un ensemble dénombrable d'ensembles de rang α au plus sur chacun desquels f est, à ε près, de classe inférieure à α , donc f est, à ε près, de classe α dans le domaine $M + E(n_0) + E(n_0 + 1) + \dots$

La proposition est établie. Ainsi :

L'ensemble des points en lesquels une fonction f n'est pas, à ε près, de classe $\alpha > 0$ au plus, est un ensemble parfait.

Soit E l'ensemble parfait qui vient d'être défini; je dis que f n'est sur E , à ε près, de classe α au plus en aucun point M de E ⁽¹⁾. Soit, en effet, $\mathfrak{M}$ l'ensemble des points de E distants de M de $\frac{1}{n_0}$ au plus et je suppose, ce qui serait possible si f était, à ε près, de classe α au plus en M sur E , que n_0 est assez grand pour que, sur $\mathfrak{M}$, f soit, à ε près, de classe α au plus. Reprenons le raisonnement précédent en remplaçant M par $\mathfrak{M}$ et E_n par $E_n - (E_n, \mathfrak{M})$; rien n'est changé, sauf que

$$E_n - (E_n, \mathfrak{M})$$

n'est plus en général la somme de deux domaines, mais la somme d'une infinité dénombrable de domaines. Cela suffit pour le raisonnement et nous voyons que, dans les conditions supposées, f serait, à ε près, de classe α au plus en M , ce qui est contraire à la définition de E .

Nous pouvons donc dire en particulier que :

XIV. *Si une fonction n'est ni de la classe $\alpha > 0$, ni d'une classe*

⁽¹⁾ Il n'y a aucune difficulté à établir cette proposition sans se servir des théorèmes du paragraphe V.

antérieure, il existe un ensemble parfait E en tout point duquel elle n'est, sur E, ni de la classe α , ni d'une classe antérieure.

En d'autres termes, si l'on sait que sur tout ensemble parfait il existe des points en lesquels une fonction f est de classe $\alpha > 0$ au plus sur l'ensemble parfait considéré, on peut affirmer que f est de classe α au plus.

Le cas de $\alpha = 0$ a été exclu. Lorsque, sur tout ensemble parfait, il existe des points en lesquels une fonction f est continue sur l'ensemble considéré, c'est-à-dire lorsque f est *ponctuellement discontinue sur tout ensemble parfait*, nos raisonnements nous permettent seulement d'affirmer que f est au plus de classe 1.

D'ailleurs f peut être effectivement de classe 1, comme le montre l'exemple d'une fonction partout nulle, sauf à l'origine, fonction qui est ponctuellement discontinue sur tout ensemble parfait.

Ainsi, pour $\alpha = 1$, il existe des fonctions de classe α telles que, sur tout ensemble parfait, se trouvent des points en lesquels la fonction est de classe inférieure à α sur l'ensemble parfait; d'après le théorème XIV, cela est impossible pour $\alpha > 1$. Nous allons voir, au contraire, que toute fonction de classe 1 jouit de la propriété énoncée.

Soit f la limite de la suite convergente des fonctions continues $f_1, f_2, \dots$, et soit E un ensemble parfait. Appelons E_n l'ensemble défini par l'égalité

$$E_n = \{E, E[f_n - f_{n+1}]^2 \leq \epsilon^2, E(f_n - f_{n+2})^2 \leq \epsilon^2, \dots\};$$

E_n est un ensemble fermé; E est la somme des E_n .

Ou bien E_1 est identique à E où il existe un point A de E ne faisant pas partie de E_1 et, puisque E_1 est fermé, on peut déterminer un intervalle I_1 contenant A et ne contenant aucun point de E_1 . Ou bien, dans I_1 , E et E_2 sont identiques, ou bien on peut déterminer dans I_1 un intervalle I_2 contenant à son intérieur des points de E et ne contenant aucun point de E_2 . Ou bien, dans I_2 , E et E_3 sont identiques, ou bien on peut déterminer dans I_2 un intervalle contenant des points de E et ne contenant pas de points de E_3 , et ainsi de suite.

On arrivera ainsi à un I_n dans lequel E et E_n sont identiques; sans quoi, à l'intérieur de tous ces I_n , qui contiennent tous des points de E

et qui sont contenus dans ceux qui les précèdent, existerait au moins un point de E , puisque E est parfait; or, d'après la définition de ce point, il n'appartiendrait à aucun E_i , ce qui est impossible.

Soit M un point de E contenu à l'intérieur de cet I_n ; autour de ce point on peut choisir un intervalle δ intérieur à I_n dans lequel l'oscillation de f_n est inférieure à ε , et puisque, dans I_n , f et f_n diffèrent de ε au plus sur E , l'oscillation de f sur E est, dans δ , inférieure à 3ε .

Ceci posé, choisissons les nombres $\varepsilon_1, \varepsilon_2, \dots$ tendant vers zéro. On peut trouver un intervalle δ_1 contenant à son intérieur des points de E et dans lequel l'oscillation f , sur E , soit inférieure à ε_1 ; dans δ_1 , on peut trouver un intervalle δ_2 contenant à son intérieur des points de E et dans lequel l'oscillation de f , sur E , soit inférieure à ε_2 , et ainsi de suite.

À l'intérieur de $\delta_1, \delta_2, \dots$ existe au moins un point de l'ensemble parfait E ; en ce point la fonction f , étant d'oscillation inférieure à ε_i sur E , est continue sur E .

De tout cela nous concluons que :

XV. Pour qu'une fonction soit de classe 1 au plus il faut et il suffit qu'elle soit ponctuellement discontinue sur tout ensemble parfait.

Ce théorème a été démontré par M. Baire dans sa Thèse citée (¹). Si l'on ne démontrait les propositions qui précèdent que dans la mesure où elles servent à l'établissement du théorème XV et si, comme j'en ai indiqué la possibilité, on modifiait les raisonnements de manière à ne pas utiliser les théorèmes du paragraphe V, on aurait une démonstration particulièrement simple du théorème de M. Baire. La méthode qui a servi à celui-ci pour établir son théorème est beau-

(¹) Le raisonnement primitif de M. Baire ne s'appliquait qu'aux fonctions d'une variable; l'exactitude du théorème pour le cas général restait douteuse, comme le remarquait M. Baire à la page 88 de sa Thèse.

J'ai indiqué (*Comptes rendus*, 27 mars 1899) que l'emploi d'un théorème démontré plus loin (théorème XIX) prouvait son exactitude dans tous les cas. Depuis (*Bulletin de la Société mathématique*, 1900), M. Baire a donné un raisonnement applicable au cas général.

coup plus compliquée, mais elle a l'avantage de fournir un procédé régulier permettant, par une infinité dénombrable d'opérations, de reconnaître si une fonction est ou non de classe 1.

Tout procédé opératoire relatif aux fonctions les plus générales suppose que l'on sache effectuer certaines opérations relatives à ces fonctions. Comme il n'y a aucune question, si simple qu'elle soit, que l'on puisse résoudre pour la fonction la plus générale, donnée d'une manière quelconque, tout procédé opératoire est illusoire quand on cherche à l'appliquer au cas général. Le procédé de M. Baire n'échappe pas à cette critique, car il suppose que l'on sache trouver les points de discontinuité d'une fonction sur un ensemble parfait, ce que l'on ne sait pas faire dans le cas général. Mais, comme le plus souvent on sait effectuer cette recherche, le procédé de M. Baire est pratiquement utile toutes les fois qu'il ne demande qu'un nombre fini d'opérations. Quand il exige un nombre infini d'opérations, il peut encore être utile, non plus à proprement parler comme procédé opératoire, mais comme guide du raisonnement.

Ce n'est pas là, à mon avis, l'unique avantage du procédé de M. Baire (le théorème XV lui-même permet le plus souvent de reconnaître facilement si une fonction donnée est ou non de classe 1). Mais, et l'on peut dire quelque chose d'analogue pour chaque procédé opératoire, tandis que le théorème XV montre seulement qu'il y aurait contradiction à supposer à la fois qu'une fonction est ponctuellement discontinue pour tout ensemble parfait et qu'elle n'est ni de classe 0, ni de classe 1, le procédé de M. Baire fournit une définition précise d'un ensemble sur lequel la fonction est totalement discontinue si elle n'est ni de classe 0, ni de classe 1 (¹). J'ajoute que, si l'on a pu démontrer, par les procédés de M. Baire, qu'une fonction f est de classe 1, on sait construire une suite de fonctions continues tendant vers f .

Essayons de généraliser l'énoncé de M. Baire. Le théorème XIV

(¹) La méthode qui nous a donné le théorème de M. Baire fournit bien aussi une définition d'un tel ensemble, mais cette définition suppose connues des propriétés des fonctions de classe 1, tandis que celle que fournit le procédé de M. Baire ne suppose connues que des propriétés relatives à la classe 0.

nous apprend, nous l'avons déjà remarqué, qu'il faut renoncer à trouver sur tout ensemble parfait des points en lesquels une fonction donnée de classe $\alpha > 1$ est de classe inférieure à α . Par exemple, la fonction $\chi(x)$ (p. 140) qui est de classe 2 dans tout intervalle n'est, en aucun point, ni de classe 0, ni de classe 1. Nous serons conduit à la généralisation cherchée en utilisant des notions importantes introduites par M. Baire et qui ont conduit celui-ci à une condition nécessaire pour qu'une fonction soit représentable analytiquement, la suivante : la fonction doit être ponctuellement discontinue sur tout ensemble parfait, quand on néglige les ensembles de première catégorie par rapport à l'ensemble parfait considéré.

Pour comprendre cet énoncé, quelques explications sont nécessaires. Considérons un ensemble $\mathcal{C}$ formé à l'aide de points d'un ensemble parfait E ; M. Baire (Thèse, p. 65 et 67) dit qu'il est *de première catégorie par rapport à E* s'il est la somme d'une infinité dénombrable d'ensembles partout non denses sur E , soient $E_1, E_2, \dots$ (1). Si $\mathcal{C}$ est de première catégorie, il ne contient pas tous les points de E ; en effet, appelons D_0 un intervalle contenant à son intérieur tout E et D_i un intervalle intérieur à D_{i-1} , contenant à son intérieur des points de E , mais aucun point de E_i . Ces D_i peuvent être choisis de bien des manières. Leur partie commune contient au moins un point de E qui, n'appartenant à aucun E_i , n'appartient pas à $\mathcal{C}$.

Il existe donc des ensembles qui ne sont pas de première catégorie sur E ; M. Baire les appelle les *ensembles de deuxième catégorie sur E* ; E lui-même est un tel ensemble. Remarquons encore que, si E est la somme d'une infinité dénombrable d'ensembles e , l'un au moins des e est de deuxième catégorie sur E , car la somme d'une infinité dénombrable d'ensembles de première catégorie est évidemment un ensemble de première catégorie. En particulier, le complémentaire par rapport à E d'un ensemble de première catégorie sur E est un ensemble de deuxième catégorie sur E .

Une autre définition est nécessaire pour faire comprendre l'énoncé

(1) C'est-à-dire que, quel que soit i , dans tout intervalle contenant à son intérieur des points de E , on peut trouver un autre intervalle jouissant de la même propriété et ne contenant pas de points de E_i .

de M. Baire. Nous dirons qu'une fonction f est continue sur E en un point P de E , lorsque l'on néglige les ensembles d'une certaine famille d'ensembles, s'il existe un ensemble e de cette famille tel que f soit continue en P sur l'ensemble $E - (e, E)$. Cette définition ne présente quelque intérêt que si P est un point limite de $E - (e, E)$. Enfin nous dirons que f est ponctuellement discontinue sur E , quand on néglige les ensembles d'une famille F , si E est l'ensemble dérivé de l'ensemble des points de E en lesquels f est continue sur E quand on néglige les ensembles de la famille F . Sans chercher pour le moment à légitimer l'énoncé de M. Baire (¹), nous allons déduire quelques conséquences des définitions qui précèdent.

Je vais d'abord montrer que si $\mathcal{C}$ est de deuxième catégorie sur E , il existe un domaine D tel que, si Δ est un domaine quelconque contenu dans D et contenant des points de E , $(\Delta, \mathcal{C})$ est de seconde catégorie sur E ; ce que j'exprime en disant que, dans D , $\mathcal{C}$ est partout de seconde catégorie sur E . Considérons les intervalles $I(a_p \leq x_p \leq b_p)$, où les a_p et b_p sont rationnels. Enlevons de $\mathcal{C}$ tous les $(\mathcal{C}, I)$ qui sont de première catégorie; l'ensemble enlevé est de première catégorie; l'ensemble restant $\mathcal{C}_1$ est de même catégorie que $\mathcal{C}$. Si $\mathcal{C}$ est de seconde catégorie, il en est de même de $\mathcal{C}_1$, et alors il existe un domaine D tel que, quel que soit le domaine Δ intérieur à D , et contenant des points de E , $\mathcal{C}_1$ a des points dans Δ . Donc aucun des I intérieurs à D ou Δ et con-

(¹) M. Baire a démontré le théorème dont il s'agit ici, aux pages 81 à 87 de sa Thèse, pour les fonctions d'une variable et de classe 2. Depuis, M. Baire a énoncé le théorème général dans les *Comptes rendus* du 11 décembre 1899.

Je me suis un peu écarté dans le texte du langage et des définitions adoptés par M. Baire, aux pages 72 et 73, 81 et 82 de sa Thèse. M. Baire ne définit pas la continuité lorsque l'on néglige les ensembles d'une famille F , mais il dit ce que l'on doit entendre dans ces conditions par l'*oscillation de la fonction*. Si la fonction est continue, au sens du texte, l'oscillation, au sens de M. Baire, est nulle. La réciproque n'est pas toujours vraie. Elle l'est cependant si la famille F est celle des ensembles dénombrables, ou des ensembles de classe ou de rang α , ou des ensembles de première catégorie, ou des ensembles parfaits ou fermés, ou des ensembles de mesure nulle, etc. Cela m'a permis d'adopter la définition du texte suffisante pour mon objet. Cette définition n'est d'ailleurs que provisoire; elle sera complétée plus loin (p. 189).

tenant des points de E n'a été enlevé, $\mathcal{C}$, et $\mathcal{C}$ sont de seconde catégorie dans Δ , c'est-à-dire partout de seconde catégorie dans D .

Étant donné l'ensemble $\mathcal{C}$ sur E , enlevons de $\mathcal{C}$ tous les ensembles $(\mathcal{C}, I)$ qui sont partout de seconde catégorie sur E dans I , les I étant les mêmes intervalles que précédemment. Il reste un ensemble $\mathcal{C}'$ de première catégorie sur $\mathcal{C}$. Donc tout ensemble est la somme d'un ensemble de première catégorie et d'une infinité dénombrable d'ensembles qui sont partout de seconde catégorie dans certains domaines les contenant.

Les ensembles partout de seconde catégorie et les ensembles de première catégorie sont, par suite, les éléments constitutants de tout ensemble. Comme ensembles partout de seconde catégorie, nous connaissons déjà les complémentaires d'ensembles de première catégorie; cherchons ce que sont les autres. Si un ensemble est partout de seconde catégorie et a pour complémentaire un ensemble de seconde catégorie, il existe un intervalle dans lequel l'ensemble considéré et son complémentaire sont tous deux partout de seconde catégorie sur l'ensemble parfait E considéré.

M. Baire a appelé mon attention, il y a trois ou quatre ans, sur l'intérêt qu'il y aurait à nommer un ensemble A partout de seconde catégorie ainsi que son complémentaire. En effet, d'après l'énoncé de M. Baire cité précédemment, la fonction nulle aux points de A , égale à 1 aux autres points, étant totalement discontinue sur E quand on néglige les ensembles de première catégorie, ne serait pas représentable analytiquement. Concluons de là, en nous servant du théorème VI, que A ne serait pas mesurable B. Nous allons démontrer cette propriété sans nous servir de l'énoncé de M. Baire et nous en déduirons cet énoncé.

Convenons pour un instant de dire qu'un ensemble est Z s'il n'existe aucun intervalle dans lequel cet ensemble et son complémentaire soient tous deux partout de seconde catégorie. Si un ensemble est Z , son complémentaire, par rapport à l'ensemble parfait E à l'aide des points duquel on forme les ensembles que nous considérons, est aussi Z . La somme S d'un nombre fini ou d'une infinité dénombrable d'ensembles M , qui sont Z , est aussi Z ; car, si S est partout de seconde catégorie dans l'intervalle I , l'un des M est de seconde catégorie dans

I; par suite, il existe dans I un intervalle I_1 dans lequel l'un des M est partout de seconde catégorie, le complémentaire de M n'est pas de seconde catégorie dans I_1 , celui de S n'est donc pas partout de seconde catégorie dans I. Cette propriété, appliquée aux complémentaires des ensembles M, nous montre que la partie commune à un nombre fini ou à une infinité dénombrable d'ensembles M, qui sont Z, est aussi Z⁽¹⁾.

En résumé, les ensembles formés par les deux opérations I et II précédemment étudiées, à partir d'ensembles qui sont Z, sont aussi des ensembles qui sont Z. Tout ensemble contenu dans l'ensemble parfait E et qui est mesurable B est la partie commune à E et à un ensemble mesurable B; par suite, c'est un ensemble obtenu par les deux opérations citées, à partir des ensembles (E, I) où I est un intervalle quelconque. Un ensemble (E, I) est un ensemble parfait contenu dans E; un tel ensemble ne peut être partout dense sur E dans un certain intervalle que s'il est identique à E dans cet intervalle; par suite, un ensemble (E, I) est toujours Z. Donc tous les ensembles mesurables B, contenus dans E, sont Z. C'est ce que l'on peut énoncer de la manière suivante :

Lorsqu'un ensemble $\mathcal{E}$, mesurable B et formé à l'aide de points d'un ensemble parfait E, est de seconde catégorie sur E, il existe un intervalle contenant des points de E et dans lequel le complémentaire de $\mathcal{E}$ par rapport à E est de première catégorie.

Ceci posé, soient E un ensemble parfait, f une fonction de classe α . f est, à ε près, constante sur chacun des ensembles M d'une infinité dénombrable d'ensembles de classe α au plus, d'après le théorème IV. L'un au moins des ensembles (E, M) est de seconde catégorie sur E; donc il existe un intervalle I contenant à son intérieur des points de E et dans lequel cet ensemble (E, M) est partout de seconde catégorie. Remarquons que sur (E, M) f est constante à ε près dans I et que l'ensemble négligé $[I, E - (E, M)]$ est de première catégorie sur E.

Prenons $\varepsilon = \frac{1}{i}$ et faisons-lui correspondre, comme il vient d'être

(¹) Je ne sais pas s'il existe des ensembles qui ne sont pas Z.

dit, un intervalle I_1 et un ensemble M_1 ; prenons ensuite $\varepsilon = \frac{1}{2}$ et faisons-lui correspondre I_2 intérieur à I_1 et un ensemble M_2 ; à $\varepsilon = \frac{1}{3}$ correspondra I_3 intérieur à I_2 et M_3 , etc. A l'intérieur de tous ces intervalles existe au moins un point P de E ; en P , f est, à ε près, continue quand on néglige l'ensemble de première catégorie N

$$N = [I_1, E - (E, M_1)] + [I_2, E - (E, M_2)] + \dots;$$

donc :

XVI. Toute fonction représentable analytiquement est ponctuellement discontinue sur tout ensemble parfait quand on néglige les ensembles de première catégorie par rapport à cet ensemble parfait.

C'est l'énoncé de M. Baire.

La démonstration précédente montre que l'ensemble négligé est mesurable B , elle fournit même une limite supérieure de sa classe, mais cela est de peu d'intérêt. Soient, en effet, $N_1, N_2, \dots$ les ensembles partout non denses sur E dont la somme est l'ensemble N négligé, la fonction f est encore continue en P si l'on néglige l'ensemble

$$(N_1 + N'_1) + (N_2 + N'_2) + \dots,$$

lequel, étant la somme d'une infinité dénombrable d'ensembles fermés, partout non denses sur E , est de classe 2 au plus et est de première catégorie par rapport à E .

Maintenant une question très importante se pose : la condition nécessaire fournie par le théorème XVI est-elle suffisante ? Et, si elle ne l'est pas, existe-t-il des fonctions qui n'y satisfont pas ? Je n'essaierai pas de répondre à ces difficiles questions ; il me suffit d'avoir montré par les considérations précédentes, très incomplètes, que des raisonnements simples, de la nature de ceux qui m'ont servi dans les précédents paragraphes, permettent d'obtenir certains des résultats publiés par M. Baire.

La condition suffisante pour qu'une fonction soit de classe α au plus que donne le théorème XIV, condition qui est évidemment nécessaire,

peut être considérée comme la généralisation du théorème XV, donnée par M. Baire, malgré la différence des énoncés. Il n'est pas difficile, d'ailleurs, de donner un énoncé analogue à celui du théorème XV et convenant à tous les cas. Pour arriver à un tel énoncé, je reprends la définition de la continuité en un point P , quand on néglige les ensembles d'une famille F . Dans l'application de cette définition, on peut faire deux conventions distinctes :

1° Celle qui a été faite jusqu'ici et d'après laquelle, quelle que soit P , on néglige un ensemble quelconque de la famille F , *contenant ou non* P ;

2° Celle d'après laquelle, pour définir la continuité en P , on néglige un ensemble de la famille F *ne contenant pas* P .

Ces deux conventions sont évidemment distinctes; la fonction $\chi(x)$, page 140, est partout continue, quand on néglige les ensembles dénombrables avec la convention 1°, elle ne l'est que pour les valeurs irrationnelles de x avec la convention 2°; la continuité, quand on néglige les ensembles finis (contenant un nombre fini de points), est la continuité ordinaire si l'on fait la convention 2°, avec la convention 1° c'est la continuité ordinaire dans l'ensemble des points autres que celui que l'on considère. La convention 1° permet de définir ce que l'on pourrait appeler *la continuité autour de* P , la convention 2° permet de définir ce que nous appellerons maintenant *la continuité en* P ⁽¹⁾.

Avec la convention nouvelle que nous venons de faire, le théorème XVI n'est plus démontré; mais il est facile de compléter la démonstration précédente.

Remarquons d'abord que les points d'un ensemble M qui ne sont pas *intérieurs* à un intervalle dans lequel M est partout de seconde catégorie forment un ensemble de première catégorie, d'après la décomposition d'un ensemble quelconque en ensembles de première caté-

(1) Je rappelle (*voir* la note de la page 185) que le langage que j'ai adopté est différent de celui choisi par M. Baire. J'ai seulement traduit dans un langage voisin de celui que j'adopte définitivement une propriété donnée par M. Baire. J'insiste d'ailleurs sur la différence entre les deux conventions citées, parce que l'une d'elles seulement conduit à un énoncé général fournissant immédiatement le théorème XV comme cas particulier.

gorie et en ensembles partout de seconde catégorie. Ceci posé dans la démonstration du théorème XVI, à chaque nombre ε nous avons fait correspondre une infinité dénombrable d'ensembles M . Les points Q de E , qui ne sont pas intérieurs à un intervalle dans lequel l'un des (E, M) contenant Q est partout de seconde catégorie sur E , forment un ensemble $E(\varepsilon)$ de première catégorie sur E . L'ensemble $e = E\left(\frac{1}{1}\right) + E\left(\frac{1}{2}\right) + E\left(\frac{1}{3}\right) + \dots$ est donc de première catégorie sur E et il existe des points de E n'appartenant pas à e .

Soit P l'un de ces points, j'appelle M_k l'un des M , correspondant à $\varepsilon = \frac{1}{k}$, et partout de seconde catégorie dans un certain intervalle contenant P à son intérieur. J'appelle I_k cet intervalle. Il est évident que l'on peut se servir des intervalles I_k et des ensembles M_k comme de ceux qui ont été définis précédemment; le théorème XVI est donc démontré ⁽¹⁾.

Supposons maintenant que, au lieu d'invoquer le théorème IV, pour la démonstration du théorème XVI, nous nous soyons servi du théorème X, les ensembles M sont alors de rang α au plus ⁽²⁾, ainsi que les (E, M) , puisque E est de rang 1. Soit A l'un des ensembles (E, M) de première catégorie sur E ; il est la somme des ensembles $B_1, B_2, \dots$, partout non denses sur E , donc aussi la somme des ensembles $(A, B_i + B'_i)$, $(A, B_2 + B'_2), \dots$ partout non denses sur E . Or $B_i + B'_i$ est fermé, donc de rang 1, A est de rang α au plus, $(A, B_i + B'_i)$ est de rang α au plus. Et nous pouvons supposer que tous ceux des (E, M) de première catégorie que nous rencontrerons seront remplacés par une infinité dénombrable d'ensembles de rang α

(1) La nature particulière des ensembles M_k n'intervient pas, c'est parce que la propriété démontrée n'est qu'un cas particulier de la suivante que l'on vérifiera sans peine : Si, en négligeant les ensembles de première catégorie sur un ensemble parfait E , une fonction est ponctuellement discontinue sur E lorsque l'on adopte l'une des conventions 1^o et 2^o, elle l'est aussi lorsque l'on adopte l'autre convention.

Il n'y a donc en réalité aucune différence entre les deux sens de l'énoncé XVI

(2) Au lieu de supposer f constante, à ε près, sur chaque M , on peut supposer que f est, à ε près, de classe inférieure à α . Cela permet de donner au théorème XVII une autre forme.

au plus, partout non dense sur E. Dès lors le point P jouit de la propriété qui intervient dans la définition suivante :

Une fonction est dite continue (α) sur l'ensemble parfait E, au point P de E, si, à tout nombre positif ε , on peut faire correspondre un intervalle contenant P à son intérieur, dans lequel E peut être considéré comme la somme d'une infinité dénombrable d'ensembles de rang α au plus, sur chacun desquels f est constante à ε près et qui sont tous, sauf l'un d'eux contenant P, partout non denses sur E.

En un tel point P, d'après X, f est de classe α au plus sur E ; de plus f est continue quand on néglige les ensembles de première catégorie sur E. La réciproque est vraie ; je ne m'en servirai pas. Remarquons encore que la continuité (1) est la continuité ordinaire, et cela grâce à l'emploi de la convention 2°.

Lorsque E sera le dérivé de l'ensemble des points où f est continue (α) sur E, nous dirons que f est ponctuellement discontinue (α) sur E.

Avec ces définitions, en tenant compte du théorème XIV, nous pouvons énoncer une proposition contenant les théorèmes XV et XVI comme cas particuliers.

XVII. *Pour qu'une fonction soit de classe α au plus, il faut et il suffit qu'elle soit ponctuellement discontinue (α) sur tout ensemble parfait.*

VII. — Relations entre différentes familles de fonctions.

Fonctions définies analytiquement. — On sait que, un ensemble E de points $(x_1, x_2, \dots, x_n)$ étant donné, on appelle *projection de cet ensemble sur la variété* $x_{i+1} = x_{i+2} = \dots = x_n = 0$ l'ensemble e de tous les systèmes de valeurs associées $(x_1, x_2, \dots, x_i)$. Je vais démontrer que, si E est mesurable B, sa projection l'est aussi.

Cela est évident si E est un intervalle, car alors e en est un aussi. Or tout ensemble mesurable B se déduit d'intervalles par l'application

sieurs déterminations. J'indique seulement la nature de ces propositions.

Supposons qu'on dise qu'une fonction à plusieurs déterminations est mesurable B quand, quels que soient a , b et l'entier n , l'ensemble des points, où n déterminations, et n seulement, satisfont à l'inégalité $a \leq f \leq b$, est mesurable B. On s'assurera facilement qu'il est nécessaire et suffisant qu'une fonction à un nombre fini ou à une infinité dénombrable de déterminations soit mesurable B pour qu'on puisse définir simultanément toutes ces déterminations par une relation analytique entre la fonction et les variables. Bien entendu, une telle relation ne définira la fonction qu'implicitement si, comme je l'ai supposé jusqu'à présent, on n'admet dans les expressions analytiques que les seuls signes $+$, $\times$, $\lim$. Mais, si l'on y adjoint un signe d'opération non uniforme, il n'en est plus nécessairement ainsi. Par exemple, si l'on emploie le signe $\sqrt[n]{}$, on pourra représenter explicitement, par une seule expression analytique, toutes les déterminations d'une fonction mesurable B pourvu que le nombre de ces déterminations soit toujours une puissance de 2 ou qu'il y ait une infinité dénombrable de déterminations. C'est, en un certain sens, une généralisation du théorème XVIII.

Pour démontrer ce théorème nous avons utilisé une remarque sur les projections des ensembles. En la précisant et en la généralisant nous pourrions en déduire de nouvelles conséquences.

Relations entre les fonctions de plusieurs variables et les fonctions d'une seule variable. — Considérons la transformation T

$$(T) \quad X_1 = X_1(x_1, x_2, \dots, x_n), \quad \dots, \quad X_p = X_p(x_1, x_2, \dots, x_n),$$

définie à l'aide des fonctions X_i continues dans le domaine d de l'espace $x_1, x_2, \dots, x_n$; et supposons qu'elle transforme le domaine d en un domaine D de l'espace $X_1, X_2, \dots, X_p$. Ces transformations sont bien connues si $n = p$; le passage d'un ensemble à sa projection définie par

$$X_1 = x_1, \quad X_2 = x_2, \quad \dots, \quad X_p = x_p$$

fournit un exemple de ces transformations pour le cas $n > p$. M. Peano a donné le premier un exemple du cas $n < p$; il a indi-

qué en effet comment on pouvait construire une courbe remplissant un intervalle I (*Math. Annalen*, t. XXXVI) ⁽¹⁾, si cette courbe est définie par

$$X_1 = X_1(t), \quad X_2 = X_2(t), \quad \dots, \quad X_p = X_p(t);$$

ces formules définissent une transformation de la nature considérée entre l'intervalle I (à p dimensions) et un intervalle (à une dimension) de l'axe des t .

Soit a un point de d , T lui fait correspondre un point A bien déterminé, le transformé de a . A est en général le transformé de plusieurs points que j'appelle *les homologues de a*. Un ensemble e de d étant donné, le transformé E de e est l'ensemble des transformés des points de e ; je désignerai par e , l'ensemble de tous les points dont les transformés font partie de E; en d'autres termes e , contient les points de e et leurs homologues; je dirai que e , est l'ensemble complet correspondant à E. Cherchons des relations entre les classes des ensembles e , E, e_1 .

Supposons que e soit fermé, c'est-à-dire soit F de classe 0. Soit M un point de D limite des points $M_1, M_2, \dots$ de E, soient $m_1, m_2, \dots$ des points de e qui ont respectivement pour transformés $M_1, M_2, \dots$. Soient enfin m un de leurs points limites, et $m_k, m_{k_2}, \dots$ des points de la suite $m_1, m_2, \dots$, tendant vers m , je dis que m admet M pour transformé. En effet les fonctions $X_1, X_2, \dots, X_p$, qui figurent dans la transformation T, étant continues, $X_i(m)$ est la limite de $X_i(m_{k_1})$; or $X_i(m_{k_1})$ étant la coordonnée X_i de M_{k_1} , $X_i(m)$ sera la coordonnée X_i de M et M est bien le transformé de m . L'ensemble E est fermé.

Donc, si e est F de classe 0, E est F de classe 0. Il faut bien remarquer que cela n'est plus vrai en général si l'on remplace F par O; en

⁽¹⁾ Voir aussi une note de M. Hilbert (*Math. Ann.*, t. XXXVIII) et la page 44 de mes *Leçons sur l'intégration*. — La transformation qui a servi au début du paragraphe I pour définir un domaine à partir d'un intervalle, transforme une courbe remplissant l'intervalle en une courbe remplissant le domaine. Tout domaine peut donc être rempli par une courbe. Ces courbes permettraient, si l'on y voyait avantage, de ne s'occuper que des transformations T où $n = 1$.

d'autres termes, si e est ouvert il n'en résulte pas que E soit aussi ouvert. De l'étude des ensembles F on ne peut pas conclure pour les ensembles O parce que, si e et g sont complémentaires par rapport à d , il n'est pas vrai en général que E et G soient complémentaires par rapport à D ; la somme $E + G$ est bien identique à D , mais E et G ont en général des points communs.

Supposons que e soit O de classe 0 ou 1. Alors e est la somme d'une infinité dénombrable d'ensembles fermés et dont les transformés sont par suite aussi fermés. Mais la somme de ces transformés est identique à E , donc E est O au plus de classe 1.

Supposons que e soit F de classe 1 ou 2. Alors e est la partie commune à une infinité dénombrable d'ensembles O de classe 1 au plus *et que l'on peut supposer contenus les uns dans les autres*. Dans ces conditions leurs transformés, qui sont O de classe 1 au plus, admettent E pour partie commune; donc E est F de classe 2 au plus.

En continuant ainsi on voit que, *si e est F de classe $2n$ au plus, E est F de classe $2n$ au plus; si e est O de classe $2n + 1$ au plus, E est O de classe $2n + 1$ au plus*. En particulier, si e est F ou O de classe inférieure à ω , E est aussi F ou O de classe inférieure à ω .

Si e est de rang ω , il est la partie commune à une infinité dénombrable d'ensembles F des classes inférieures à ω , que l'on peut supposer contenus les uns dans les autres, et dont les transformés, qui sont F des classes inférieures à ω , ont pour partie commune E . E est donc de rang ω .

Si e est le complément d'un ensemble de rang ω , il est la somme d'une infinité dénombrable d'ensembles O des classes inférieures à ω , lesquels ont pour transformés des ensembles O des classes inférieures à ω ; leur somme E est donc le complémentaire d'un ensemble de rang ω .

Si e est F de classe ω , il est la partie commune à une infinité dénombrable d'ensembles complémentaires d'ensembles de rang ω , contenus les uns dans les autres. Alors E est la partie commune des transformés de ces ensembles, donc E est de classe ω au plus.

Si e est O de classe ω , il est la somme d'une infinité dénombrable d'ensembles de rang ω , donc E est O de classe ω au plus.

A partir de là, en raisonnant par récurrence, on verra que, *si e*

est F (ou O) de classe $\alpha \geq \omega$, E est F (ou O) de classe α au plus ⁽¹⁾.

Pour le passage de e_i à E les conclusions précédentes s'appliquent, mais on peut aller plus loin. Soit g_i le complémentaire de e_i par rapport à O , g_i est évidemment l'ensemble complet correspondant au complémentaire G de E par rapport à D . Si e_i est F de classe $\alpha \geq \omega$ ou de classe finie et paire, E est, on le sait, un ensemble F de classe α au plus. Si e_i est F de classe finie et impaire α , g_i est O de classe impaire α , donc G est, on le sait, un ensemble O de classe α au plus et par suite E est F de classe α au plus. Donc, dans tous les cas, E est au plus de la même classe que e_i ⁽²⁾.

Ce résultat peut s'interpréter autrement. Soit $F(X_1, X_2, \dots, X_p)$ une fonction, la transformation T lui fait correspondre une fonction

$$f(x_1, x_2, \dots, x_n) \equiv F[X_1(x_1, \dots, x_n), \dots, X_p(x_1, \dots, x_n)].$$

Comme $E(a \leq f \leq b)$ est le correspondant complet de $E(a \leq F \leq b)$, f est au moins de la classe F .

Mais, au sujet de la fonction $F[X_1(x_1, \dots, x_n), \dots]$ composée à l'aide des fonctions continues $X_1(x_1, \dots, x_n), \dots, X_p(x_1, \dots, x_n)$ et de la fonction $F(X_1, X_2, \dots, X_p)$, nous pouvons démontrer une propriété analogue au théorème I, savoir que la fonction $f(x_1, x_2, \dots, x_n)$ est au plus de la classe de la fonction $F(X_1, X_2, \dots, X_p)$. Dire que F est d'une classe déterminée, c'est dire en effet que l'on peut construire cette fonction à l'aide d'un certain procédé opératoire à partir de fonctions continues $\Phi(X_1, X_2, \dots, X_p)$. Pour substituer dans F les valeurs de $X_1, X_2, \dots, X_p$ données par (T) , il suffit de faire la substitution dans les Φ , ce qui donne des fonctions continues $\varphi(x_1, x_2, \dots, x_n)$, puis d'employer, à partir de ces nouvelles fonctions, le même procédé opératoire qu'avant la substitution.

Comme il n'est pas démontré qu'après la substitution ce procédé

(1) En appliquant ces résultats au cas où E est une projection de e , on pourrait préciser la nature des fonctions définies implicitement par des relations obtenues en égalant à zéro des fonctions de classes connues.

(2) Je reviens ici au langage abrégé dans lequel *ensemble* F de classe α est remplacé par *ensemble de classe* α .

opérateur est le plus simple possible, nous pouvons seulement conclure que f est au plus de la classe de F .

Ce résultat, rapproché du précédent, montre que :

La transformation T fait correspondre à tout ensemble E contenu dans D et de classe α un ensemble complet e contenu dans d et exactement de classe α .

Et aussi que :

La transformation T fait correspondre à toute fonction $F(X_1, \dots, X_p)$ des variables $X_1, X_2, \dots, X_p$, qui est de classe α , une fonction

$$F[X_1(x_1, x_2, \dots, x_n), X_2(x_1, x_2, \dots, x_n), \dots, X_p(x_1, x_2, \dots, x_n)]$$

des variables $x_1, x_2, \dots, x_n$, qui est exactement de classe α .

Le fait que D est un domaine n'est presque jamais intervenu ; supprimons cette condition, alors D est un ensemble parfait. La démonstration du premier des deux énoncés précédent reste applicable, occupons-nous du second. Convenons de dire qu'une fonction est de classe α sur D s'il existe une fonction partout définie, qui soit de classe α , qui se réduise sur D à la fonction donnée et s'il n'en existe pas de classe inférieure à α . Alors il est évident que, si F est de classe α , f est de classe α au plus.

Si f est de classe 0 , la fonction F , d'où dérive f , et qui n'est connue que sur D , est continue d'après les résultats qui précèdent ; mais puisqu'on sait qu'il existe une fonction continue, partout définie, se réduisant sur D à une fonction continue donnée, F est, sur D , de classe 0 au plus.

Si f est de classe $\alpha \geq 1$, la fonction F , définie sur D , est telle que $E(\alpha \leq F \leq b)$ est toujours de classe α au plus, mais l'ensemble complémentaire de D , étant de classe 1 , est aussi de classe α au plus, Donc la fonction égale à F sur D , à 0 pour les autres points est de classe α au plus, F est, sur D , de classe α au plus.

Les deux énoncés trouvés sont donc exacts dans tous les cas où d est

un domaine ⁽¹⁾. La seule application que j'en veux faire est la suivante : Considérons une courbe C définie à l'aide de fonctions continues d'un paramètre t variant dans un intervalle d ; soit

$$X_1 = X_1(t), \quad X_2 = X_2(t), \quad \dots, \quad X_p = X_p(t)$$

cette courbe. A toute fonction $F(X_1, X_2, \dots, X_p)$ définie sur elle, correspond une fonction $f(t) \equiv F[X_1(t), X_2(t), \dots, X_p(t)]$ définie dans d . La classe de cette fonction $f(t)$ est ce que nous appellerons *la classe de F sur la courbe C*. Il est évident que cette classe est bien attachée à la courbe C et qu'elle ne dépend pas de la représentation paramétrique choisie, c'est-à-dire qu'elle reste la même si, dans les fonctions $X_1(t), X_2(t), \dots, X_p(t)$, on substitue une fonction $t(\theta)$ continue et variant toujours dans le même sens. Mais l'extension que nous avons donnée aux énoncés précédents permet d'aller plus loin : la classe de F sur la courbe C étant la classe de F sur l'ensemble parfait D des points de C reste la même si l'on substitue à C une autre courbe Γ assujettie à la seule condition d'avoir un même ensemble de points que C . Et cela revient à substituer, dans les coordonnées des points de C , des fonctions $t(\theta)$ qui sont discontinues, à plusieurs déterminations, l'ensemble de ces déterminations pouvant même avoir la puissance du continu pour certaines valeurs de θ .

Si l'on ne s'occupe pas de ces généralisations, on ne peut conclure pour toutes les courbes, mais on peut toujours conclure pour les courbes qui remplissent un domaine, et cela suffit pour légitimer le théorème XIX :

XIX. Pour qu'une fonction soit de classe α dans un domaine, il faut et il suffit qu'elle soit de classe α au plus pour chaque courbe du domaine et qu'elle soit effectivement de classe α pour une courbe du domaine.

Cette courbe particulière peut d'ailleurs être prise indépendamment

⁽¹⁾ Ils sont vrais aussi si d est un ensemble parfait, mais cette généralisation nous est tout à fait inutile.

Lorsque d est un domaine, D est un ensemble parfait, mais il faut bien remarquer que ce n'est pas un ensemble parfait quelconque.

de la fonction considérée, il suffit de prendre une courbe remplissant le domaine (¹).

Il est bien évident que, si l'on veut déduire la classe d'une fonction dans un domaine de la classe de cette fonction sur une courbe, toujours la même, il faut que cette courbe passe par tous les points du domaine. Mais on peut préférer à l'emploi de cette courbe, unique mais nécessairement compliquée, l'emploi d'une famille de courbes plus simples.

Je ne sais pas s'il est possible de nommer une telle famille, mais je vais montrer par un exemple simple que la famille des courbes analytiques ne répond pas à la question (²).

Considérons trois circonférences tangentes intérieurement en un point A. Soient C_1 la plus grande, C_2 la circonférence moyenne, C_3 la plus petite. Soit $f(x, y)$ une fonction, nulle sur C_1 , C_3 , à l'extérieur de C_1 , à l'intérieur de C_3 ; égale à 1 sur C_2 sauf en A; linéaire sur chaque rayon de C_3 , entre C_1 et C_2 d'une part, entre C_2 et C_3 d'autre

(¹) J'ai énoncé ce théorème dans une Note des *Comptes rendus* (27 mars 1899) et j'en ai déduit qu'on pouvait appliquer à tous les cas l'énoncé du théorème XV, démontré par M. Baire pour les fonctions d'une variable. Ma démonstration primitive était très différente comme forme de celle du texte; en réalité, les propriétés utilisées étaient les mêmes. J'employais encore les relations entre un ensemble, son transformé et l'ensemble complet correspondant. Seulement, au lieu de considérer une courbe quelconque remplissant le domaine, je ne me servais que des courbes de Peano dont l'ensemble des points multiples est particulièrement simple.

(²) Il suffit qu'une fonction soit de classe $\alpha = 0$ sur toute droite $x = \text{const.}$ et sur toute courbe $y = f(x)$ pour qu'elle soit de classe $\alpha = 0$ dans tout le plan (x, y) . Je ne sais pas si cette propriété est encore exacte pour $\alpha > 0$; si elle l'était, on aurait là une famille simple répondant à la question. Cette famille serait d'autant plus intéressante qu'elle ne comprend que les courbes qu'on étudie ordinairement, l'étude d'une courbe $y = f(x)$ étant souvent beaucoup plus facile que celle d'une courbe quelconque. Par exemple, M. Baire a démontré qu'une fonction $F(x, y)$, continue en x et en y , définit sur chacune de ces courbes une fonction de classe 1; il est fort possible que les méthodes de M. Baire puissent être étendues au cas des courbes quelconques, il n'en est pas moins intéressant de se demander si l'on a le droit de conclure du résultat de M. Baire que F est de classe 1, comme cela sera démontré plus loin.

part. Il est évident que $f(x, y)$ est discontinue en A, et pourtant elle est continue sur toute droite du plan.

Faisons une construction analogue en remplaçant C_1, C_2, C_3 par des courbes qui, au voisinage de A, sont transcendantes et ont un contact d'ordre infini. On peut supposer, par exemple, que ces arcs de courbes ont pour équations

$$y = f_1(x), \quad y = f_2(x), \quad y = f_3(x),$$

et que l'on a

$$f_1(0) = f_2(0) = f_3(0) = a_0,$$

$$f_1^{(p)}(0) = f_2^{(p)}(0) = f_3^{(p)}(0) = a_p p!,$$

$f_i^{(p)}$ désignant la dérivée $p^{\text{ième}}$ de f_i , les a_0 et les a_p étant choisis à l'avance de façon que la série $a_0 + \sum a_p x^p$ ne soit pas convergente pour $x \neq 0$. Alors la fonction $f(x, y)$ est continue sur toute courbe analytique, et cependant elle est discontinue en A. Bien entendu, une série uniformément convergente de telles fonctions convenablement choisies donnerait *une fonction continue sur tout arc analytique et cependant discontinue dans tout domaine*.

Les fonctions ainsi construites sont de classe 1; il est facile de le voir de bien des manières et cela résultera d'un théorème qui va être démontré; je vais cependant le démontrer pour la première fonction $f(x, y)$ construite, ce qui me conduira à une série intéressante. De A comme centre, je décris la circonférence Γ_n de rayon $\frac{1}{n}$; soit φ_n une fonction continue égale à f en tout point où $f = 0$, égale à f à l'extérieur de Γ_n et comprise entre 0 et 1; f est la limite de φ_n , donc f est de classe 1, mais, de plus, la suite obtenue qui n'est pas partout uniformément convergente est cependant uniformément convergente sur toute droite. Cette suite peut évidemment être remplacée par une série de polynômes, et ce qui a été dit de la première des fonctions $f(x, y)$ peut l'être des autres, de sorte qu'il existe des séries de polynômes uniformément convergentes sur tout arc analytique sans être uniformément convergentes dans aucune aire.

Toutes les fonctions $f(x, y)$ que nous venons de construire sont de classe 1; avant de rattacher cela à un théorème général, je montre

que, contrairement à ce que pourraient faire penser ce théorème et les exemples précédents, il ne suffit pas de connaître la classe d'une fonction sur toute droite du plan pour connaître une limite supérieure de la classe de cette fonction. En effet, une fonction partout nulle, sauf peut-être pour les points d'une circonférence, est de la classe 1 au plus sur toute droite, elle peut être cependant de classe quelconque, ou même échapper à tout mode de représentation analytique.

Fonctions de plusieurs variables continues par rapport à chacune d'elles.

XX. *Une fonction de n variables, continue par rapport à chacune d'elles, est de classe $n - 1$ au plus.*

Soit $f(x_1, x_2, \dots, x_n)$ une telle fonction et supposons-la définie dans un intervalle, ce qui ne restreint pas la généralité, puisque, si f n'est définie que dans D , il est possible de la définir à l'extérieur de D en respectant les continuités; cela suppose toutefois que f est continue par rapport aux variables en tous les points frontières de D , mais cela est admis implicitement dans l'énoncé.

Supposons donc que f est définie pour $a \leq x_i \leq b$ et divisons l'intervalle (a, b) en n parties égales à l'aide des points $a_0 = a, a_1, \dots, a_n = b$. Soit φ_n la fonction égale à f quand x a l'une de ces valeurs a_i et variant linéairement quand, $x_2, x_3, \dots, x_n$ restant constantes, x_1 varie de a_i à a_{i+1} . φ_n est une fonction continue par rapport aux ensembles $(x_1, x_2), (x_1, x_3), \dots, (x_1, x_n)$ et de plus φ_n tend vers f , quand n augmente indéfiniment.

Opérons sur φ_n comme sur f en faisant jouer à x_2 le rôle de x_1 ; nous verrons que φ_n est la limite d'une suite convergente de fonctions continues en $(x_1, x_2, x_3), (x_1, x_2, x_4), \dots, (x_1, x_2, x_n)$. Nous opérerons sur ces nouvelles fonctions comme sur f et φ_n en faisant jouer à x_3 le rôle de x_1 et x_2 , et ainsi de suite. Au bout de $n - 1$ opérations, nous arriverons à des fonctions continues par rapport à l'ensemble des variables et à partir desquelles f s'obtient par $n - 1$ passages successifs à la limite; donc f est de classe $n - 1$ au plus ⁽¹⁾.

(1) J'ai déjà donné cette démonstration dans une Note du *Bulletin des Sciences mathématiques* (*Sur l'approximation des fonctions*, novembre 1898).

On peut se demander, il est vrai, si la limite supérieure trouvée pour la classe peut être effectivement atteinte. La réponse est affirmative; nous allons démontrer, en effet, que :

XXI. Si $f(t)$ est une fonction de classe n , il existe une fonction $\varphi(x_1, x_2, \dots, x_{n+1})$ continue par rapport à chacune de ses $n+1$ variables et telle que $f(t)$ soit identique à $\varphi(t, t, \dots, t)$.

En d'autres termes, $\varphi(x_1, x_2, \dots, x_{n+1})$ se réduit à f sur une certaine droite.

1° $n = 1$. — On a alors

$$f(t) = \lim_{n \rightarrow \infty} f_n(t),$$

les $f_n(t)$ étant continues. Soient $\varepsilon_1, \varepsilon_2, \dots$ des nombres décroissants et tendant vers zéro, et soit η_n un nombre tel que, dans un intervalle quelconque de longueur η_n , l'oscillation de f_n soit inférieure à ε_n ; je suppose, de plus, que les η_n sont choisis décroissant avec $\frac{1}{n}$ et tendant vers zéro.

Nous définissons φ par les égalités

$$\varphi(x_1, x_2) = \varphi(x_2, x_1), \quad \varphi(t, t) = f(t)$$

et

$$\varphi(x_1, t) = f_{n+1}(t) + [f_n(t) - f_{n+1}(t)] \frac{x_1 - t - \eta_{n+2}}{\eta_{n+1} - \eta_{n+2}},$$

quand

$$t + \eta_{n+2} \leq x_1 \leq t + \eta_{n+1}.$$

Il est évident que, pour $x_1 \geq x_2$, φ est continue en x_1 ; montrons que φ est continue en x_2 . Donnons à x_1 une valeur constante, lorsque l'on a

$$x_1 - \eta_{n+1} \leq x_2 \leq x_1 - \eta_{n+2},$$

φ est comprise entre le plus petit et le plus grand des nombres suivants :

$$f_n + \varepsilon_n, \quad f_n - \varepsilon_n, \quad f_{n+1} + \varepsilon_{n+1}, \quad f_{n+1} - \varepsilon_{n+1},$$

et, puisque tous ces nombres tendent vers f , φ , qui est évidemment continue en x_2 pour $x_2 < x_1$, l'est aussi pour $x_2 = x_1$.

Ce raisonnement deviendrait plus simple encore si l'on traduisait en langage géométrique la construction indiquée analytiquement ⁽¹⁾.

2° n quelconque. — On peut maintenant passer d'une valeur de n à la valeur immédiatement supérieure à l'aide d'un raisonnement que je me contenterai d'exposer dans le cas particulier du passage de $n = 1$ à $n = 2$. Cela me permettra d'employer le langage géométrique et de bien mettre en évidence la propriété des domaines à plus de deux dimensions qui est fondamentale dans ce raisonnement.

Par la droite D , $x = y = z$, faisons passer les trois plans Dx , Dy , Dz qui contiennent respectivement Ox , Oy , Oz . Ces trois plans divisent l'espace en six domaines indéfinis, limités chacun par deux faces d'un dièdre. Entre ces systèmes de deux faces, menons, par D , deux plans; par exemple, traçons les six plans Π passant par D et faisant 20° avec l'un des plans Dx , Dy , Dz . Les plans Π , Dx , Dy , Dz divisent l'espace en dix-huit dièdres de 20° d'angle; les plans Π seuls divisent l'espace en douze dièdres : les uns, les dièdres A de 40° d'angle, sont tels qu'ils contiennent les parallèles à l'une des directions de l'un des axes, menées par les points de D ; les autres, les dièdres B de 20° d'angle, ne contiennent aucune de ces parallèles.

Si l'on veut construire une fonction $\varphi(x, y, z)$ continue en x , en y et en z et égale sur D à une fonction donnée, et si l'on sait résoudre cette question pour chaque domaine A , on saura évidemment la résoudre pour tout l'espace; en effet, si l'on donne à $\varphi(x, y, z)$ les valeurs déjà connues, quand (x, y, z) est point d'un domaine A , il sera facile de définir $\varphi(x, y, z)$ dans les domaines B tout en respectant les continuités, parce qu'il n'y aura plus à s'occuper de la continuité aux points de D .

Il suffit donc de résoudre la question pour un domaine A et, dans un tel domaine, on n'a à se préoccuper, pour les points de D , que de la continuité par rapport à une seule des variables. Cette simplification résulte, comme on vient de le voir, de la possibilité d'une certaine

(¹) En 1899, M. Baire m'avait communiqué une démonstration due à M. Volterra et s'appliquant au cas de $n = 1$.

division de l'espace; une division analogue est possible pour les espaces à plus de trois dimensions, elle est impossible pour l'espace à deux dimensions.

Pour démontrer le théorème dans le cas de $n = 2$, il suffit de définir la fonction φ , dans un domaine A , l'un des dièdres précédemment défini ou tout autre domaine possédant la propriété remarquable indiquée. Supposons que ce domaine contienne les parallèles menées par les points de D à la direction positive de Ox .

Soit $f(t) = \lim_{n \rightarrow \infty} f_n(t, t)$ la fonction donnée; $f_n(u, v)$ étant une fonction continue par rapport à chacune des deux variables u, v . Désignons par $P_1, P_2, \dots$ des plans parallèles entre eux et à D , se rapprochant constamment de D , dont la distance à D tend vers zéro, coupant le domaine A et non parallèles à Ox . On pourra prendre pour P_n

$$2x - y - z = \frac{1}{n}.$$

La fonction $\varphi(x, y, z)$ est maintenant définie par les conditions suivantes : 1° $\varphi(x, y, z)$ est égale à $f(x)$ sur D , c'est-à-dire que l'on a $\varphi(x, x, x) = f(x)$; 2° au point de P_n , dont les coordonnées sont $\frac{1}{2}(y + z + \frac{1}{n})$, y, z , la fonction $\varphi(x, y, z)$ est égale à $f_n(y, z)$; 3° $\varphi(x, y, z)$ est linéaire en x quand on se déplace sur une parallèle à l'axe des x entre deux plans P_i consécutifs.

Ces conditions déterminent évidemment, en tout point de A , une fonction φ remplissant les conditions imposées. Le théorème est donc établi pour $n = 2$, et l'on voit que le raisonnement se généralise facilement.

Des théorèmes XX et XXI il résulte qu'il y a identité entre les fonctions de classe n et celles qu'on obtient en donnant des valeurs égales aux variables d'une fonction de $n + 1$ variables continues par rapport à chacune d'elles. Cette identité fournit un nouveau moyen de rattacher les fonctions de classe n aux fonctions continues, qui pourrait être pris pour base d'une étude des fonctions de classe finie, et l'on verra facilement qu'en l'employant certains des théorèmes précédents s'obtiennent très facilement. M. Baire est parvenu, pour le cas de $n = 1$, à la démonstration de l'identité qui vient d'être indiquée en

cherchant simultanément les conditions nécessaires et suffisantes pour qu'une fonction soit de première classe et pour qu'elle puisse être obtenue en faisant $x = y$ dans une fonction $f(x, y)$ continue en x et en y ; il a trouvé comme réponse à ces deux questions des conditions identiques, celles qui sont fournies par le théorème XV (¹).

VIII. — Exemples de fonctions.

Je vais maintenant démontrer l'existence de fonctions de toute classe et l'existence de fonctions échappant à tout mode de représentation analytique.

Un théorème de M. G. Cantor est fondamental pour notre objet : l'ensemble des fonctions a une puissance supérieure au continu. En se servant de ce théorème, page 71 de sa thèse, M. Baire a pu affirmer l'existence de fonctions qui n'appartiennent à aucune classe. Je vais reprendre ces points en précisant davantage. J'essaierai de ne jamais parler d'une fonction sans la définir effectivement; je me place ainsi à un point de vue fort analogue à celui choisi par M. Borel dans ses *Leçons sur la théorie des fonctions*.

Un objet est défini ou donné quand on a prononcé un nombre fini de mots s'appliquant à cet objet et à celui-là seulement; c'est-à-dire quand on a nommé une propriété caractéristique de l'objet. Pour donner une fonction $f(x_1, x_2, \dots, x_n)$ on nomme généralement une propriété appartenant à tous les ensembles de nombres

$$x_1, x_2, \dots, x_n, f(x_1, x_2, \dots, x_n)$$

(¹) Il a été question précédemment de fonctions continues sur toutes les droites; on pourrait se demander si la classe de ces fonctions peut atteindre la même limite que lorsqu'il s'agit simplement de fonctions continues par rapport à chacune des variables; la réponse est affirmative.

En modifiant légèrement les constructions employées pour le théorème XXI, on construira une fonction de $n + 1$ variables, continue sur toute droite et se réduisant, sur une courbe C, à une fonction arbitrairement donnée de classe n . La courbe C ne peut pas être prise arbitrairement; pour $n = 1$, on pourra prendre un arc de cercle; pour $n = 2$, un arc d'hélice; et, d'une manière générale, on pourra prendre pour C un arc de courbe rencontrant, en un nombre fini de points, toute variété linéaire à n dimensions.

et à ceux-là seulement; mais cela n'est nullement nécessaire, on peut nommer d'autres propriétés caractéristiques de cette fonction. C'est ce que l'on fait, par exemple, lorsque, une fonction $f(x)$ étant définie d'une manière quelconque, on dit que $F(x)$ est celle des fonctions primitives de $f(x)$ qui s'annule pour $x = 0$ ⁽¹⁾. C'est nommer une fonction que de dire qu'elle est égale à *zéro* ou *un* suivant que la constante d'Euler C est rationnelle ou non.

Il ne faudrait d'ailleurs pas croire qu'une fonction est nécessairement mieux définie quand on donne une propriété caractéristique de l'ensemble γ , $x_1, x_2, \dots, x_n$, car une telle propriété ne permet pas en général de calculer γ . Par exemple, la fonction $\chi(x)$, page 140, qui admet même une représentation analytique connue, n'est pas connue pour $x = C$, bien que l'on sache calculer C avec autant de décimales que l'on veut ⁽²⁾ et, si nous la connaissons pour $x = \pi$, ce n'est pas son expression analytique qui nous la fait connaître.

On ne devra donc pas s'étonner si, dans la suite, je considère comme parfaitement définies et données des fonctions que je ne saurai calculer pour aucune valeur des variables; je ne connaîtrai en général rien de plus, sur ces fonctions, que leur définition; cela me permettra cependant de donner un sens beaucoup plus précis à l'énoncé de M. Cantor. Je reprends la démonstration du théorème de M. Cantor.

Convenons de dire que, une certaine famille de fonctions étant donnée, on sait réaliser une application de cette famille sur le continu lorsque, à toute fonction de la famille, on sait faire correspondre un

(1) Je rappelle que, si f n'est pas bornée, on ne connaît aucun procédé général permettant de reconnaître si F existe et donnant sa valeur.

(2) Il y aurait lieu de rechercher quelles sont les expressions analytiques qui fournissent réellement des procédés de calcul. La réponse à cette question dépendra évidemment de la nature des opérations que nous considérerons comme réellement effectuelles.

D'une façon générale un calcul est illusoire s'il suppose que l'on effectue successivement deux passages à la limite, à moins que le second ne soit relatif à une suite uniformément convergente. Or, c'est un tel calcul que l'on aurait à effectuer pour calculer $\chi(C)$, C étant supposée donnée par la suite de ses chiffres décimaux, c'est-à-dire par une série.

ensemble de valeurs de la variable t , comprise entre 0 et 1, de manière qu'à une valeur de t corresponde une fonction au plus.

Supposons qu'une telle application soit possible pour une famille $\mathcal{F}$ de fonctions d'une variable. Définissons une fonction $F(t)$ par la condition d'être partout nulle sauf pour les valeurs de t auxquelles correspondent des fonctions de $\mathcal{F}$ qui s'annulent pour cette valeur particulière t de la variable; alors nous prendrons $F(t) = 1$. La fonction F n'est évidemment identique à aucune des fonctions de la famille $\mathcal{F}$ ⁽¹⁾.

C'est ce raisonnement que l'on traduit ordinairement en disant qu'il *existe* une fonction ne faisant pas partie de $\mathcal{F}$, ou encore que la puissance de l'ensemble des fonctions est supérieure à celle de l'ensemble des fonctions de $\mathcal{F}$. Pour être précis, nous devons seulement conclure que : *Ou bien il est impossible de nommer une application de $\mathcal{F}$ sur le continu, ou bien il est possible de nommer une fonction F ne faisant pas partie de $\mathcal{F}$.*

Ainsi, si l'on peut nommer une application, on peut nommer une fonction F ; si l'on sait seulement qu'il existe une application, on doit conclure qu'il existe une fonction F . En général, on démontre l'existence d'un objet en donnant le moyen de le nommer; il n'en est cependant pas toujours ainsi, j'en donnerai plus loin un exemple en faisant cependant des réserves ⁽²⁾.

M. Borel, dans la Note III de ses *Leçons sur la théorie des fonctions*, cite une classe étendue de familles de fonctions pour lesquelles on peut nommer une application sur le continu : les familles de fonctions définies par des conditions dénombrables. Un ensemble dénombrable de telles familles constitue évidemment une nouvelle famille de fonctions définies par des conditions dénombrables. La famille des fonctions continues est une famille de fonctions définies par des conditions dénombrables, et l'on peut, de bien des manières, *citer* une application de cette famille sur le continu. Partant d'une de ces appli-

⁽¹⁾ On pourrait raisonner de même pour les fonctions de plusieurs variables, mais les théorèmes de M. Cantor et ceux qui précèdent nous permettent de nous occuper uniquement des fonctions d'une seule variable; c'est ce qui sera fait dans la suite; je supposerai même qu'il s'agit de fonctions définies dans $(0, 1)$.

⁽²⁾ Page 213. Voir aussi la note de la page 176.

cations et remarquant qu'une fonction de classe α est définie par une suite dénombrable convergente de fonctions de classe inférieure à α , on en déduira facilement une application particulière des fonctions de classe α au plus sur le continu. Du théorème de Cantor il résulte alors *qu'on peut nommer une fonction qui n'est ni de la classe α , ni d'une classe inférieure* (¹).

Nous ne pouvons pas conclure de là l'existence de fonctions de toute classe, car nous ne savons pas si la fonction ainsi définie est représentable analytiquement et de classe supérieure à α ou si elle n'est pas représentable analytiquement. Nous allons voir qu'on peut arriver à nommer une fonction représentable analytiquement et de classe supérieure à α .

Je démontre d'abord une propriété simple des symboles de classe. *Si l'on a des symboles de classe $\alpha_1, \alpha_2, \dots$ décroissants, le nombre de ces symboles est certainement fini.* Supposons, en effet, qu'on puisse avoir une suite de symboles décroissants $\alpha_1, \alpha_2, \dots$, et considérons l'ensemble des symboles de classe inférieurs à tous les α_i , cet ensemble E est nécessairement dénombrable puisqu'il ne contient que des symboles inférieurs à α_1 , donc il existe un symbole β qui est le premier venant après tous ceux de E. β ne faisant pas partie de E est égal ou supérieur à l'un des α_i . Il ne peut être supérieur à l'un des α_i puisque tous les symboles inférieurs à β font partie de E; donc β égale l'un des α_i ; si $\beta = \alpha_i$, la suite ne contient que i symboles, le théorème est donc démontré (²).

Supposons donné un symbole de classe α . Cela veut dire qu'on a, par un procédé quelconque, défini les symboles de classe inférieure à α . Rangeons ces symboles, qui sont en nombre *fini* ou *dénombrable*, en suite simplement infinie $\alpha_1, \alpha_2, \dots$, le même symbole pouvant revenir

(¹) Cette application particulière des propriétés des familles de fonctions définies par des conditions dénombrables, que M. Borel a données dans son Livre cité, avait été faite par M. Borel lui-même, qui m'a communiqué son raisonnement, analogue à celui du texte.

(²) Dans ce théorème il s'agit, bien entendu, de symboles ordonnés de manière qu'après chaque symbole il y ait un premier symbole déterminé et que l'ensemble ait un premier symbole; ce théorème ne veut pas dire qu'avant chaque classe il n'y a qu'un nombre fini de classes.

plusieurs fois (¹). Soit f une fonction de classe α , on peut considérer f comme la limite d'une suite convergente $f_1, f_2, \dots$. Si α est de première espèce, je supposerai que toutes ces fonctions sont de classe $\alpha - 1$; si α est de seconde espèce, je supposerai que f_n est de classe α_{r_n} , α_{r_1} étant α , et α_{r_n} étant le premier des symboles $\alpha_{r_{n-1}+1}, \alpha_{r_{n-1}+2}, \dots$, qui surpasse $\alpha_{r_{n-1}}$ (²).

Je peux considérer f_n comme la limite d'une suite convergente $f_{n,1}, f_{n,2}, \dots$. Si α_{r_n} est de première espèce, je supposerai que toutes ces fonctions sont de la classe $\alpha_{r_n} - 1$; si α_{r_n} est de seconde espèce, je supposerai que $f_{n,p}$ est de la classe $\alpha_{r_{n,p}}$, $\alpha_{r_{n,1}}$ étant le premier des symboles $\alpha_1, \alpha_2, \dots$, inférieur à α_{r_n} et $\alpha_{r_{n,p}}$ étant le premier des symboles $\alpha_{r_{n,p-1}+1}, \alpha_{r_{n,p-1}+2}, \dots$, qui surpasse $\alpha_{r_{n,p-1}}$ sans égaler α_{r_n} .

Les fonctions $f_{n,p}$ seront supposées définies comme limites de fonctions $f_{n,p,q}$ dont les classes seront fixées par le même procédé, et ainsi de suite. Nous pourrions toujours déduire, d'une fonction à k indices, des fonctions à $k + 1$ indices ayant les k premiers indices communs avec la fonction primitive; nous serons arrêtés cependant quand nous arriverons à une fonction de classe 0. Considérons une suite de fonctions déduites les unes des autres $f_n, f_{n,p}, f_{n,p,q}, \dots$; les classes correspondantes, qui sont inférieures à α , vont en décroissant, donc la suite ne comprend qu'un nombre fini de fonctions. Toutes les fonctions que nous définissons ont donc un nombre fini d'indices; elles forment, par suite, un ensemble dénombrable. Il faut remarquer que tout cortège d'indices ne correspond pas à une fonction; mais, un cortège d'indices étant donné, nous savons reconnaître s'il correspond à une fonction et, si oui, quelle est la classe de cette fonction.

En particulier, nous savons reconnaître quels cortèges d'indices correspondent aux fonctions de classe 0. D'après la façon dont les classes

(¹) Si l'on voit une difficulté dans le fait que je n'indique pas d'une façon générale la loi de formation de la suite $\alpha_1, \alpha_2, \dots$, et je ne puis l'indiquer puisque je ne sais pas nommer le symbole α le plus général, on pourra considérer que ce qui suit n'est applicable qu'aux symboles α qu'on aura nommés d'une façon spéciale : en nommant la loi suivant laquelle est formée la suite $\alpha_1, \alpha_2, \dots$.

(²) Tout cela suppose pour un instant l'existence effective des classes inférieures à $\alpha + 1$, et même l'existence de fonctions bornées appartenant à ces classes.

ont été définies, si $f_{1,p}$, où I est un certain cortège d'indices, est de classe 0, la suite $f_{1,1}, f_{1,2}, \dots$ est une suite convergente de fonctions de classe 0; or une telle suite pouvant toujours être remplacée par une suite, convergeant vers la même limite, formée de polynômes dont le $p^{\text{ième}}$ est de degré p et à coefficients inférieurs à 2^p , je peux supposer que $f_{1,p}$ est un polynôme de degré p et à coefficients inférieurs à 2^p . Je poserai donc

$$(1) \quad f_{1,p} = 2^p(a_{1,p,1} + a_{1,p,2}x + \dots + a_{1,p,p+1}x^{p+1}),$$

les a seront compris entre 0 et 1. Il suffit évidemment de se donner ces a et leurs indices pour que la fonction f soit déterminée.

Je vais définir les nombres a comme fonctions d'un paramètre t , pour cela je range d'abord en suite simplement infinie les cortèges d'indices correspondant aux a . Par exemple, j'adopte le procédé suivant : la suite $I(n, p, \dots, t)$ sera placée avant la suite $I_1(n_1, p_1, \dots, t_1)$ si la somme $n + p + \dots + t$ est inférieure à la somme $n_1 + p_1 + \dots + t_1$, ou si, ces deux sommes étant égales et les indices de I étant égaux à ceux de même rang dans I_1 jusqu'au $p^{\text{ième}}$ rang, le $(p+1)^{\text{ième}}$ indice de I est inférieur au $(p+1)^{\text{ième}}$ indice de I_1 (¹). Soit $I_1, I_2, \dots$ cette suite de cortège d'indices.

Je suppose que t , compris entre 0 et 1, peut être écrit sous la forme

$$t = \frac{\theta_1}{3} + \frac{\theta_2}{3^2} + \dots + \frac{\theta_n}{3^n} + \dots,$$

tous les θ étant égaux à 0 ou 2. Alors t fait partie d'un certain ensemble parfait Z . Je prends

$$\begin{aligned} 2a_{1,1} &= \frac{\theta_1}{2} + \frac{\theta_3}{2^2} + \dots + \frac{\theta_{2n+1}}{2^n} + \dots, \\ 2a_{1,2} &= \frac{\theta_{2,1}}{2} + \frac{\theta_{2,3}}{2^2} + \dots + \frac{\theta_{2(2n+1)}}{2^n} + \dots, \\ &\dots\dots\dots \\ 2a_{1,p} &= \frac{\theta_{2p-1,1}}{2} + \frac{\theta_{2p-1,3}}{2^2} + \dots + \frac{\theta_{2p-1,2n+1}}{2^n} + \dots, \\ &\dots\dots\dots \end{aligned}$$

Soit t une valeur ne faisant pas partie de Z . Puisque Z est parfait,

(¹) Pour faire cette comparaison on ajoutera, si cela est nécessaire, un certain nombre de zéros à la suite des indices de I ou de I_1 .

t fait partie d'un intervalle (t_0, t_1) dont les extrémités appartiennent à Z et dont aucun autre point ne fait partie de Z . Soient $\alpha_{l_p}^0, \alpha_{l_p}^1$ les valeurs de α_{l_p} correspondant à t_0 et t_1 ; dans (t_0, t_1) je prends

$$\alpha_{l_p} = \alpha_{l_p}^0 + \frac{\alpha_{l_p}^1 - \alpha_{l_p}^0}{t_1 - t_0} (t - t_0).$$

On verra facilement que les α_{l_p} ainsi définies sont des fonctions continues de t ⁽¹⁾.

A partir de ces α_{l_p} comme coefficients, on peut former par les égalités telles que (1) des polynômes en x , je continuerai à les désigner par les mêmes symboles f que précédemment, mais j'indiquerai qu'il s'agit maintenant de fonctions de deux variables t, x . Ainsi, si f_1 désignait un polynôme en x , $f_1(t, x)$ désignera une fonction continue en (t, x) . A partir de ces fonctions continues $f_1(t, x)$, comme auparavant à partir des fonctions continues f_1 , par des passages à la limite nous formerons des expressions $f_1(t, x)$ ayant de moins en moins d'indices et correspondant aux fonctions f_1 . Nous arriverons ainsi en particulier à une expression $f(t, x)$ correspondant à f . Les expressions $f_1(t, x)$, $f(t, x)$ n'auront, en général, aucun sens, car les passages à la limite qu'indiquent ces expressions ne s'appliquent plus nécessairement à des suites convergentes; mais, quelle que soit la fonction f de classe α , ou même de classe inférieure à α , il existe une infinité de valeurs de t , et même une infinité de valeurs de t faisant partie de Z , telles que $f(t, x)$ soit identique, quel que soit x , à f .

Nous pouvons donc appliquer le procédé de M. Cantor aux fonctions $f(t, x)$; posons $\varphi(x) = 0$, sauf quand $f(x, x)$ a un sens et est égal à 0, auquel cas nous prendrons $\varphi(x) = 1$. Il est évident que $\varphi(x)$ est représentable analytiquement et n'est pas de classe égale ou inférieure à α . Donc il existe des fonctions de toute classe; c'est-à-dire qu'il y a contradiction à admettre à la fois les axiomes ordinaires de l'analyse plus cette propriété : toute suite convergente de fonctions des classes égales ou inférieures à α a pour limite une fonction de classe égale ou inférieure à α , puisque alors $\varphi(x)$ serait certainement

(1) On pourra consulter sur ce point la page 44 de mes *Leçons sur l'intégration*. Les considérations du texte définissent en somme une courbe remplissant tout un domaine de l'espace à une infinité dénombrable de dimensions.

de classe égale ou inférieure à α . Mais on peut aller plus loin. Cherchons la classe de $\varphi(x)$.

Cette classe est supérieure à α . D'autre part, si f_1 est de classe 0, $f_1(t, x)$ est de classe 0; de là on déduit que, si f_1 est de classe β , $f_1(t, x)$ est au plus de classe β ; donc $f(t, x)$ est au plus de classe α , donc $f(x, x)$ est au plus de classe α ⁽¹⁾. Les deux ensembles

$$E[f(x, x) \neq 0] = E[\varphi(x) = 0], \quad E[f(x, x) = 0] = E[\varphi(x) = 1],$$

étant respectivement O et F de classe α au plus, sont F de classe $\alpha + 1$ au plus; $\varphi(x)$ est au plus de classe $\alpha + 1$. Donc $\varphi(x)$ est de classe $\alpha + 1$.

Soient β un symbole de classe quelconque et $\beta_1, \beta_2, \dots$ les symboles inférieurs rangés en suite simplement infinie. Si β est de première espèce, barrons dans la suite le symbole $\beta - 1$ et opérons sur la suite restante comme sur la suite $\alpha_1, \alpha_2, \dots$, cela nous donnera une fonction $\varphi(x)$ de classe β . Si β est de seconde espèce, barrons dans la suite des β_i tous ceux qui sont de seconde espèce, et soit $\gamma_1, \gamma_2, \dots$ la nouvelle suite. Soit γ_i un symbole quelconque de cette suite, barrons tous ceux des γ_j qui sont égaux ou supérieurs à γ_i ; cela donne une suite $\delta_1, \delta_2, \dots$ ⁽²⁾ à partir de laquelle, opérant comme sur $\alpha_1, \alpha_2, \dots$, nous obtiendrons une fonction $\varphi_i(x)$ de classe γ_i ; je représente cette fonction par $\varphi_i(x)$. Je pose alors $\varphi(x) = 0$ si $x = \frac{1}{2^i}$ quel que soit l'entier i et, pour $\frac{1}{2^{i+1}} < x < \frac{1}{2^i}$, $\varphi(x) = \varphi_i(2^i x - 1)$. Il est alors évident que $\varphi(x)$ est de classe β . Nous avons donc nommé une fonction de classe donnée β , je l'appelle $\varphi_\beta(x)$ ⁽³⁾.

Pour nommer maintenant une fonction échappant à tout mode de représentation analytique il nous suffira d'établir une certaine correspondance entre les symboles de classe et les points d'un segment. J'indique d'abord un raisonnement vague d'où l'on est peut-être en droit de conclure qu'il existe des fonctions non représentables analytique-

(1) Ici il est fait usage de la classification des fonctions non partout définies, mais cela n'entraîne aucune difficulté (voir p. 164).

(2) Si cette suite était finie et s'écrivait $\delta_1, \delta_2, \dots, \delta_q$, on la remplacerait par la suite infinie $\delta_1, \delta_2, \dots, \delta_q, \delta_1, \delta_2, \dots, \delta_q, \delta_1, \delta_2, \dots$.

(3) Cette fonction n'est déterminée que lorsqu'on connaît la suite $\beta_1, \beta_2, \dots$.

ment. Le raisonnement de M. Cantor montre que l'ensemble des fonctions a une puissance supérieure à celle du continu, montrons donc que l'ensemble E des fonctions représentables analytiquement n'a pas une puissance supérieure au continu. Cela est vrai de l'ensemble E_α des fonctions de classe α au plus, puisque celles-ci dépendent de l'infinité dénombrable des constantes a_r . L'ensemble E est la somme des E_α ; d'un théorème de M. Cantor⁽¹⁾ il résulte que, si l'ensemble des α n'a pas une puissance supérieure au continu, E n'a pas une puissance supérieure au continu. Raisonnons ainsi : au symbole 1 faisons correspondre un point t_1 ; au symbole 2, un point t_2 ; ...; au symbole ω , un point t_ω ; ...; au symbole α , un point t_α , et ainsi de suite. Nous ne serons jamais arrêté, car, jusqu'à α , nous n'avons employé qu'une infinité dénombrable de points t_α , donc il reste des points et nous pourrions isoler l'un d'eux $t_{\alpha+1}$. De là nous concluons que l'ensemble des symboles a une puissance au plus égale à celle du continu; il existe des fonctions non représentables analytiquement.

Mais, si l'on se reporte à ce que j'ai dit sur le théorème de Cantor, et si l'on remarque que nous avons prétendu démontrer l'existence d'une application de l'ensemble des symboles sur le continu, sans en nommer aucune, on aura peut-être quelque doute sur la valeur du raisonnement précédent. On doit même se demander, à mon avis, s'il a un sens quelconque et s'il est légitime de parler d'un nombre infini d'opérations ou de choix sans les déterminer en donnant la loi suivant laquelle ils sont faits. Je reprends donc la question.

Je suppose les nombres rationnels compris entre 0 et 1 rangés dans un certain ordre que je ne précise pas, mais que je suppose précisé; soit $z_1, z_2, \dots$ la suite considérée. Je prends une valeur de t comprise entre 0 et 1 et je l'écris dans le système de numération de base 2, en employant seulement, lorsque cela est possible, un nombre fini de fois le chiffre 1; l'expression de t est bien déterminée dès que t est donnée

$$t = \frac{0_1}{2} + \frac{0_2}{2^2} + \dots$$

Barrons dans la suite $z_1, z_2, \dots$ tous les z_i qui correspondent aux

(¹) *Acta mathematica*, t. II. — BOREL, *Leçons sur la théorie des fonctions*, p. 16 et 17.

indices i des chiffres θ_i qui sont nuls; soient $z'_1, z'_2, \dots$ les z conservés. Il se peut, et il en sera toujours ainsi si les z' sont en nombre fini, qu'on puisse trouver un symbole de classe α jouissant de la propriété suivante : on peut établir une correspondance entre les z' et tous les symboles β inférieurs à α de manière qu'à un z' corresponde un seul β , et inversement, et que, si à β et β_1 correspondent $z'(\beta)$ et $z'(\beta_1)$, on ait $z'(\beta) < z'(\beta_1)$ si β est plus petit que β_1 . Alors α sera dit le symbole correspondant à t . De plus, à tout symbole β ($\beta < \alpha$) correspond un entier i bien déterminé, l'indice du z' correspondant; si donc on range les β en leur donnant comme rang ces entiers i on aura une suite $\alpha_1, \alpha_2, \dots$ formée des symboles inférieurs à α ⁽¹⁾. A la valeur de t considéré nous pouvons, par suite, faire correspondre une fonction $\varphi_\alpha(x)$ de classe α ; il suffit d'employer le procédé précédemment indiqué; cette fonction $\varphi_\alpha(x)$, qui n'est pas déterminée dès que α est donné, mais qui l'est dès que t est donné, si à t correspond un symbole α , je l'appellerai $\varphi(t, x)$. Si à t ne correspond pas de symbole α , je poserai $\varphi(t, x) = 0$.

Je dis que la fonction $\varphi(t, x)$ échappe à toute représentation analytique; cela sera évidemment démontré quand il sera prouvé que tout symbole α correspond à une valeur t , de sorte que, quel que soit α , il existe une valeur de t telle que $\varphi(t, x)$, en tant que fonction de x , est de classe α . Je vais examiner seulement le cas le plus difficile $\alpha \geq \omega$, dans lequel chaque α correspond à une infinité non dénombrable de valeurs de t .

Soient α un symbole et $\alpha_1, \alpha_2, \dots$ les symboles inférieurs à α , pris chacun une seule fois, et rangés en suite simplement infinie. Je prends arbitrairement une série convergente de somme inférieure à 0,5 et dont les termes $\varepsilon_1, \varepsilon_2, \dots$ sont positifs. Soit β un symbole inférieur à α , il a un certain rang i dans la suite $\alpha_1, \alpha_2, \dots$; à β je fais correspondre ε_i . Soit S_β la somme, inférieure à 0,5 — ε_i , de tous les ε qui correspondent

(1) J'admets dans tout ceci que, t étant donné, α , s'il existe, est déterminé et que la correspondance entre les β et les z' ne peut être établie que d'une façon. La démonstration de ces faits est immédiate, je ne la donne pas pour abrégé. On pourra consulter sur ce point le Mémoire de M. Cantor *Sur les fondements de la théorie des ensembles transfinis* (traduction de F. Marotte, Paris, Hermann, 1899) ou la Note qui termine mes *Leçons sur l'intégration et la recherche des fonctions primitives*.

à des symboles inférieurs à β , posons $l_\beta = 2S_\beta + \varepsilon_i$, $l'_\beta = 2S_\beta + 2\varepsilon_i$; soit enfin y_β celui des nombres rationnels compris entre l_β et l'_β qui, écrit sous forme irréductible, fournit pour la somme de ses deux termes le résultat le plus petit possible. La suite des y_β forme une suite telle que $z'_1, z'_2, \dots$. Dans la suite des z' , y_β occupe le $i^{\text{ième}}$ rang, dans la suite des z il occupe le rang I ($I \geq i$). Je pose

$$t = \frac{\theta_1}{2} + \frac{\theta_2}{2^2} + \dots + \frac{\theta_n}{2^n} + \dots,$$

en convenant de prendre $\theta_k = 1$ si K est l'un des rangs I correspondant aux $\beta < \alpha$ et $\theta_k = 0$ dans le cas contraire. Il est évident alors qu'à cette valeur de t correspond α ; il reste cependant à démontrer que l'expression choisie pour t n'est pas l'une de celles exclues; c'est-à-dire que tous les θ d'indice assez grand ne sont pas égaux à 1. S'il en était ainsi l'ensemble des y_β ne différerait de celui des z_β que par la suppression d'un nombre fini de termes, or cela est impossible puisque tous les nombres rationnels compris entre y_1 et y_2 ne font pas partie des y_β .

La fonction $\varphi(t, x)$ non représentable analytiquement nous fournira autant de *fonctions d'une variable non représentables analytiquement* que nous le voudrons; il suffira, d'après le théorème XIX, de considérer la fonction d'une variable définie par $\varphi(t, x)$ sur une courbe remplissant le domaine $0 \leq t \leq 1$, $0 \leq x \leq 1$.

Ainsi on peut nommer des fonctions non représentables analytiquement, aussi ne faudrait-il pas confondre l'étude qui vient d'être entreprise avec celle des *fonctions qu'on peut nommer*, étude qu'il serait intéressant d'aborder.

J'ajoute quelques remarques qui se déduisent immédiatement des considérations précédentes.

On a vu qu'à tout symbole α on pouvait faire correspondre une infinité de valeurs t ($0 < t < 1$) et qu'une valeur t correspondait à un symbole au plus; on peut traduire cela en disant que *l'ensemble des symboles (ou des nombres transfinis) a au plus la puissance du continu*, propriété qui, comme je l'ai dit, ne me paraît pas rigoureusement démontrée par les raisonnements que j'ai rappelés.

Soit $\psi(t)$ l'une des fonctions non représentables analytiquement que nous avons appris à former; elle ne prend que les valeurs 0 ou 1, donc

les deux ensembles $E[\psi(t) = 0]$, $E[\psi(t) = 1]$ sont tous deux non mesurables B. Ainsi, *nous savons nommer des ensembles non mesurables B.*

Soit E un ensemble non mesurable B formé de valeurs de x comprises entre 0 et 1. Je puis toujours écrire

$$x = \frac{\theta_1}{2} + \frac{\theta_2}{2^2} + \dots + \frac{\theta_n}{2^n} + \dots,$$

où les θ sont égaux à 0 ou 1; lorsque cela sera possible de deux manières, je prendrai celle qui n'emploie qu'un nombre fini de fois le chiffre 1. Je pose

$$t = \frac{2\theta_1}{3} + \frac{2\theta_2}{3^2} + \dots + \frac{2\theta_n}{3^n} + \dots$$

La correspondance entre x et t est exprimable analytiquement, donc, à tout ensemble de valeurs de x (ou de t) mesurable B, elle fait correspondre un ensemble en t (ou x) mesurable B; par suite à E correspond un ensemble E_1 non mesurable B. Or, cet ensemble étant formé de points de l'ensemble Z de la page 210, lequel est mesurable B et de mesure nulle, E_1 est mesurable et de mesure nulle ⁽¹⁾. Donc, *nous savons nommer un ensemble mesurable qui n'est pas mesurable B.*

Mais la question beaucoup plus intéressante : peut-on nommer un ensemble non mesurable? reste entière ⁽²⁾.

⁽¹⁾ Une fonction bornée nulle partout, sauf aux points de E_1 , est intégrable au sens de Riemann et cependant échappe à tout mode de représentation analytique.

⁽²⁾ En corrigeant les épreuves, je signale deux Ouvrages relatifs aux questions traitées ici et qui sont parus depuis l'époque où je rédigeais ce Mémoire (mai 1904). Ce sont les *Leçons sur les fonctions discontinues* de M. René Baire, et les *Leçons sur les fonctions de variables réelles et les développements en séries de polynômes* de M. Émile Borel (tous deux chez Gauthier-Villars, Paris, 1905).

Dans la note III de ce dernier Ouvrage se trouve le raisonnement dont je parlais dans la note I de la page 208. Dans la note II du même Ouvrage j'ai démontré le théorème XV par les méthodes du texte; dans quelques pages publiées dans le *Bulletin de la Société mathématique de 1904* j'ai repris cette démonstration sous une autre forme.