

## Une contribution à la théorie des ensembles.

(Mémoire de G. Cantor à Halle.)

(Extrait du Journal de Borchardt, vol. 84.)

Si on peut faire correspondre élément par élément deux ensembles bien définis  $M$  et  $N$  par une opération à sens unique (et, quand on peut le faire d'une manière, on peut le faire aussi de beaucoup d'autres), convenons, pour la suite, de nous exprimer en disant que ces ensembles ont la même puissance, ou encore qu'ils sont équivalents.

Nous appellerons parties intégrantes d'un ensemble tous les autres ensembles  $M'$ , dont les éléments sont en même temps éléments de  $M$ .

Si deux ensembles  $M$  et  $N$  ne sont pas de même puissance, ou bien  $M$  aura la même puissance qu'une partie intégrante de  $N$ , ou bien  $N$  la même qu'une partie intégrante de  $M$  ; dans le premier cas nous disons la puissance de  $M$  plus petite, dans le second cas, nous la disons plus grande que la puissance de  $N$ .

Quand les ensembles à considérer sont finis, c'est-à-dire composés d'un nombre fini d'éléments, la notion de puissance, comme il est facile de le voir, répond alors à celle de nombre dans la signification de dénombrement et par conséquent aussi à celle de nombre entier positif, puisqu'en effet deux ensembles de cette nature n'ont la même puissance que dans l'hypothèse où le nombre de leurs éléments est le même.

Une partie intégrante d'un ensemble fini a toujours une puissance plus petite que l'ensemble lui-même ; ce fait n'a plus lieu dans les ensembles infinis, c'est-à-dire composés d'un nombre infini d'éléments. De cette seule circonstance, qu'un ensemble infini  $M$  est une partie intégrante d'un autre  $N$  ou que l'on peut faire correspondre un à un les éléments de  $M$  à une partie intégrante de  $N$ , par une opération à sens unique, on ne peut aucunement conclure que sa puissance est plus petite que celle de  $N$  ; cette conclusion n'est justifiée que si l'on sait que la puissance de  $M$  n'est pas égale à celle de  $N$  ; de même,  $N$  étant partie intégrante de  $M$  ou tel que ses éléments correspondent un à un à sens unique à ceux d'une partie intégrante de  $M$ , cette circonstance ne suffit pas pour que la puissance de  $M$  soit plus grande que celle de  $N$ .

Pour rappeler un exemple simple, soit  $M$  la série des nombres entiers positifs  $\nu$ ,  $N$  la série des nombres entiers positifs pairs  $2\nu$  ;  $N$  est alors une partie intégrante de  $M$  et néanmoins  $M$  et  $N$  sont de même puissance.

La série des nombres entiers positifs offre, comme il est facile de le montrer, la plus petite de toutes les puissances qui se présentent dans les ensembles infinis. Néanmoins la classe des ensembles qui ont cette plus petite puissance est extraordinairement riche et étendue. À cette classe appartiennent, par exemple, tous les ensembles que M. DEDEKIND appelle corps finis dans ses belles recherches sur les nombres algébriques (cf. leçons de DIRICHLET sur la théorie des nombres, deuxième ou troisième édition, Brunswick 1871 et 1879) ; de même les ensembles que j'ai considérés

---

Acta mathematica, 2, imprimé le 11 juin 1883.

Transcription en Latex : Denise Vella-Chemla, juin 2025.

et que j'ai appelés "systèmes de points de la  $\nu^{\text{ième}}$  espèce (cf. Mathematische Annalen de CLEBSCH et NEUMANN, t. V, p. 129) sont de la première (c'est-à-dire de la plus petite) puissance.

Chaque ensemble se présentant comme une série simplement infinie, avec le terme général  $a_\nu$ , appartient évidemment à cette même classe ; mais de plus, les séries doubles et en général les séries  $n^{\text{ièmes}}$  avec le terme général  $a_{\nu_1, \nu_2, \dots, \nu_n}$  (où  $\nu_1, \nu_2, \dots, \nu_n$  parcourent indépendamment l'un de l'autre tous les nombres entiers positifs) appartiennent aussi à cette classe. J'ai même démontré que l'ensemble  $(\omega)$  de tous les nombres algébriques réels (et on pourrait ajouter : de tous les nombres algébriques complexes). peut se concevoir sous la forme d'une série de terme général  $\omega_\nu$  ; c'est-à-dire que l'ensemble  $(\omega)$ , aussi bien que chacune de ses parties intégrantes infinies, a la puissance de la série de nombres entiers  $1, 2, 3, \dots, \nu, \dots$  (cf. Journal de Borchardt, t. 77, page 258). À l'égard des ensembles de cette première classe, on a les théorèmes suivants, faciles à démontrer :

" $M$  étant un ensemble de la première classe (c'est-à-dire de la puissance de la série des nombres entiers positifs), chaque partie intégrante infinie de  $M$  a la même puissance.

$M', M'', M''' \dots$  étant une série finie ou simplement infinie d'ensembles, dont chacun a la première puissance, l'ensemble  $M$ , qui résulte de la réunion de  $M', M'', M''' \dots$ , a aussi la première puissance."

Nous allons maintenant dans ce qui va suivre examiner au point de vue de leur puissance les ensembles qu'on appelle continus et  $n^{\text{iples}}$ . D'après un théorème, que j'ai démontré dans le § 2 du traité cité (Jour. d. Borchardt, t. 77, page 260) il est certain, que ces ensembles n'appartiennent pas à la première classe, c'est-à-dire qu'ils ont une puissance supérieure à la première.

Les recherches de RIEMANN, de HELMOLTZ et d'autres après eux sur les hypothèses qui servent de fondement à la géométrie, partent, comme on sait, de la notion d'un ensemble continu, d'étendue  $n$ , et en font consister le caractère essentiel en ce que leurs éléments dépendent de  $n$  variables réelles, continues, indépendantes l'une de l'autre, en sorte qu'à chaque élément de l'ensemble appartient un système de valeurs  $x_1, x_2, \dots, x_n$  admissible, et réciproquement à chaque système de valeurs  $x_1, x_2, \dots, x_n$  admissible appartient un certain élément de l'ensemble.

Le plus souvent, comme il résulte de la suite de ces recherches, on suppose en outre tacitement que la correspondance des éléments de l'ensemble et du système de valeurs, posée comme base, est continue, en sorte qu'à chaque changement infiniment petit du système de valeurs répond un changement infiniment petit de l'élément correspondant de l'ensemble et réciproquement, à chaque changement infiniment petit des éléments de l'ensemble, correspond un changement semblable des valeurs de ses coordonnées.

Quant à savoir si cette supposition suffit, ou s'il faut la compléter par des conditions encore plus spéciales, pour pouvoir regarder comme bien fondée et incontestable l'idée que ces auteurs se sont faite de l'ensemble  $n^{\text{iple}}$  et continu, c'est une question que nous passerons d'abord sous silence <sup>1</sup> ; nous avons seulement à montrer ici, que si on laisse cette supposition de côté, (ce qui arrive

---

<sup>1</sup>La réponse à cette question à laquelle nous reviendrons dans une autre circonstance, ne me paraît donner lieu, à aucune difficulté sérieuse.

bien souvent dans les traités de ces auteurs), si par rapport à la correspondance entre l'ensemble et ses coordonnées on n'admet aucune limitation, ce caractère, considéré par les auteurs comme essentiel (d'après lequel un ensemble  $n^{\text{ple}}$  continu est tel qu'on peut en déterminer les éléments par coordonnées réelles, continues, indépendantes l'une de l'autre) devient absolument sans valeur.

Comme notre travail le montrera, on peut même déterminer les éléments d'un ensemble continu d'étendue  $n$  par une seule coordonnée réelle et continue au moyen d'une opération à sens unique. Il suit de là, que si on ne fait aucune supposition par rapport à la nature de la correspondance, le nombre des coordonnées réelles continues et indépendantes qui peuvent servir à la détermination à sens unique des éléments d'un ensemble continu d'étendue  $n$ , peut être tout nombre donné  $m$  et que par conséquent on ne peut le considérer comme caractère invariable d'un ensemble donné.

En me posant la question de savoir si un ensemble continu de  $n$  dimensions peut être relié au moyen d'une opération à sens unique à un ensemble continu d'une seule dimension, de telle sorte qu'à chaque élément de l'un d'eux réponde un élément, et un seulement, de l'autre, il s'est trouvé qu'une telle correspondance existe toujours.

D'après cela une surface continue peut être rapportée complètement par une opération à sens unique à une ligne continue ; la même chose est vraie des corps continus et des ensembles continus géométriques à un nombre quelconque de dimensions.

En appliquant l'expression introduite plus haut, nous pouvons donc dire que la puissance d'un ensemble continu d'étendue  $n$  et choisi à volonté est égale à la puissance d'un ensemble d'étendue simple, comme par exemple, d'un segment de droite continu et limité.

## § 1.

Comme deux ensembles continus, d'un nombre égal de dimensions, peuvent, au moyen de fonctions analytiques, se rapporter l'un à l'autre complètement et à sens unique, par rapport au but que nous poursuivons (et qui est de montrer qu'on peut joindre d'une façon complète et à sens unique des ensembles continus qui n'ont pas le même nombre de dimensions) tout se ramène, comme on l'entrevoit facilement, à la démonstration du théorème suivant :

(A). *“Soient  $x_1, x_2, \dots, x_n, n$  grandeurs réelles, variables, indépendantes l'une de l'autre, dont chacune peut prendre toutes les valeurs  $\geq 0$  et  $\leq 1$ , et soit  $t$  une autre variable comprise dans les mêmes limites ( $0 \leq t \leq 1$ ), on peut faire correspondre cette grandeur  $t$  au système des grandeurs  $x_1, x_2, \dots, x_n$  de telle sorte qu'à chaque valeur déterminée de  $t$  appartienne un système de valeurs déterminées  $x_1, x_2, \dots, x_n$  et vice versa à chaque système de valeurs déterminées  $x_1, x_2, \dots, x_n$  une certaine valeur de  $t$ .”*

Comme conséquence de ce théorème se présente cet autre que nous avons en vue :

(B). *“On peut faire correspondre d'une façon complète et à sens unique un ensemble continu à  $n$  dimensions à un ensemble continu d'une seule dimension ; deux ensembles continus l'un de  $n$ , l'autre de  $m$  dimensions,  $n$  étant  $\lesseqgtr$ , ont la même puissance ; les éléments d'un ensemble continu à*

*n dimensions peuvent être déterminés à sens unique par une seule coordonnée t continue et réelle; mais ils peuvent aussi être déterminés à sens unique par un système de m coordonnées continues  $t_1, t_2, \dots, t_m$ .*”

## § 2.

Pour démontrer (A) nous partons de ce théorème connu que tout nombre irrationnel  $e \geqslant_1^0$  peut être représenté d’une manière complètement déterminée, sous la forme d’une fraction continue infinie :

$$\frac{1}{\alpha_1 + \frac{1}{\alpha_2 + \frac{1}{\alpha_\nu + \ddots}}} = (\alpha_1, \alpha_2, \dots, \alpha_\nu, \dots)$$

où les  $\alpha_\nu$  sont des nombres positifs entiers rationnels.

À chaque nombre irrationnel  $e \geqslant_1^0$  appartient une série déterminée infinie de nombres entiers positifs  $\alpha_\nu$  et réciproquement chaque série semblable détermine un certain nombre irrationnel  $e \geqslant_1^0$ .

Soient maintenant  $e_1, e_2, \dots, e_n$ ,  $n$  grandeurs variables indépendantes l’une de l’autre, dont chacune peut prendre toutes les valeurs numériques irrationnelles de l’intervalle  $(0 \dots 1)$ , qu’on pose alors :

$$\begin{aligned} e_1 &= (\alpha_{1,1}, \alpha_{1,2}, \dots, \alpha_{1,\nu}, \dots) \\ &\dots\dots\dots \\ e_\mu &= (\alpha_{\mu,1}, \alpha_{\mu,2}, \dots, \alpha_{\mu,\nu}, \dots) \\ &\dots\dots\dots \\ e_n &= (\alpha_{n,1}, \alpha_{n,2}, \dots, \alpha_{n,\nu}, \dots) \end{aligned}$$

ces  $n$  nombres irrationnels déterminent à sens unique un  $\overline{n+1}^{\text{ième}}$  nombre irrationnel  $d \geqslant_1^0$  :

$$d = (\beta_1, \beta_2, \dots, \beta_\nu, \dots),$$

si on établit entre les nombres  $\alpha$  et  $\beta$  le rapport suivant :

$$\beta_{\nu-1n+\mu} = \alpha_{\mu,\nu} \quad \begin{cases} \mu = 1, 2, \dots, n \\ \nu = 1, 2, \dots, \infty \end{cases} \quad (1)$$

Mais aussi de l’autre côté si on part d’un nombre irrationnel  $d \geqslant_1^0$ , il détermine la série des  $\beta_\nu$  et au moyen de (1) les séries des  $\alpha_{\mu,\nu}$  ; par conséquent  $d$  détermine complètement et à sens unique le système des  $n$  nombres irrationnels  $e_1, e_2, \dots, e_n$ . De cette considération résulte tout d’abord le théorème suivant :

(C). “Soient  $e_1, e_2, \dots, e_n$   $n$  grandeurs variables indépendantes l’une de l’autre, dont chacune peut prendre toutes les valeurs numériques irrationnelles de l’intervalle  $(0 \dots 1)$  et soit  $d$  une autre variable irrationnelle dans les mêmes limites, on peut faire correspondre d’une manière complète et à sens unique cette grandeur  $d$  et le système des  $n$  grandeurs  $e_1, e_2, \dots, e_n$ .”

### § 3.

Après avoir démontré, dans le paragraphe précédent, le théorème (C), nous avons maintenant à démontrer le théorème suivant :

(D). *Une grandeur variable  $e$ , qui peut prendre toutes les valeurs numériques irrationnelles de l'intervalle  $(0 \dots 1)$ , peut se joindre à sens unique à une variable  $x$  comportant toutes les valeurs réelles, c'est-à-dire rationnelles et irrationnelles, qui sont  $\geq 0$  et  $\leq 1$ , en sorte qu'à chaque valeur irrationnelle de  $e \geq_1^0$  corresponde une valeur réelle de  $x \leq_1^0$  et une seulement et que réciproquement à chaque valeur réelle de  $x$  corresponde une certaine valeur irrationnelle de  $e$ ".*

Car une fois ce théorème démontré, qu'on se représente (en l'appliquant), comme correspondantes aux  $n + 1$  grandeurs variables désignées dans le § 2 par  $e_1, e_2 \dots e_n$ , et  $d$ , les autres variables  $x_1, x_2, \dots, x_n$  et  $t$ , reliées aux premières par une opération à sens unique, chacune des dernières variables pouvant prendre sans restriction toutes les valeurs réelles  $\geq 0$  et  $\leq 1$ . Comme nous avons établi une correspondance complète et à sens unique entre la variable  $d$  et le système des  $n$  variables  $e_1, e_2, \dots e_n$  dans le § 2, on obtient de cette manière une association complète, déterminée et à sens unique de la variable continue  $t$  et du système des  $n$  variables continues  $x_1, x_2 \dots x_n$ , ce qui démontrera la vérité du théorème (A).

Nous n'aurons donc plus à nous occuper dans la suite que de la démonstration du théorème (D); qu'on nous permette d'employer, pour plus de brièveté, un formalisme simple que nous allons d'abord faire connaître.

Nous appellerons ensemble linéaire de nombres réels tout ensemble bien défini de nombres réels, distincts les uns des autres, c'est-à-dire inégaux en sorte qu'un seul et même nombre ne se présente pas plus d'une fois comme élément dans un ensemble linéaire.

Les variables réelles, qui se présentent dans le cours de ce travail, sont toutes de telle nature que le champ de chacune d'elles, c'est-à-dire l'ensemble des valeurs qu'elle peut prendre est un ensemble linéaire donné ; nous n'appuierons donc plus dans la suite sur cette supposition que partout nous ferons tacitement.

De deux variables de cette nature  $a$  et  $b$  nous dirons qu'elles n'ont aucune liaison, si aucune des valeurs que peut prendre  $a$  n'est égale à une valeur de  $b$  c'est-à-dire si les deux ensembles de valeurs que peuvent prendre les variables  $a$  et  $b$  n'ont pas d'éléments communs, on dit que  $a$  et  $b$  sont sans liaison <sup>2</sup>.

Si on a une série finie ou infinie  $a', a'', a''', \dots, a^{(\nu)}, \dots$  de variables bien définies ou de constantes telles que  $a^{(\nu)}$  et  $a^{(\mu)}$  n'ont aucune liaison entre eux, on peut définir une variable  $a$  par ce caractère que son champ se compose de l'ensemble des champs de  $a', a'', \dots a^{(\nu)}, \dots$  ; réciproquement une variable donnée  $a$  peut se décomposer d'après les modes les plus divers en d'autres  $a', a'', \dots$  qui

---

<sup>2</sup>Deux ensembles  $M$  et  $N$  ou bien n'ont aucune liaison, s'ils n'ont aucun élément qui leur soit commun ; ou bien ils sont reliés par un troisième ensemble déterminé  $P$  c'est-à-dire par l'ensemble de leurs éléments communs. Je désigne l'ensemble  $P$  par  $\mathcal{D}(M, N)$ .

n'ont aucune liaison deux à deux ; dans ces deux cas nous exprimons le rapport de la variable  $a$  aux variables  $a', a'', \dots, a^{(\nu)}, \dots$  par la formule suivante :

$$a \equiv \{a', a'', \dots, a^{(\nu)}, \dots\}$$

Cette formule exprime à la fois 1° que toute valeur que pourrait prendre une des variables  $a^{(\nu)}$  est aussi une valeur qui convient à la variable  $a$  ; 2° que toute valeur que peut recevoir  $a$  peut être prise aussi par une des grandeurs  $a^{(\nu)}$ , et par une seulement. Pour expliquer cette formule, soit par exemple  $\varphi$  une variable qui peut prendre toutes les valeurs rationnelles  $\geq 0$  et  $\leq 1$ ,  $e$  une variable qui peut prendre toutes les valeurs irrationnelles de l'intervalle  $(0 \dots 1)$  et enfin  $x$  une variable qui peut prendre toutes les valeurs réelles, rationnelles et irrationnelles  $\geq 0$  et  $\leq 1$ , on a :

$$x \equiv \{\varphi, e\}.$$

Soit  $a$  et  $b$  deux grandeurs variables de telle nature qu'on puisse les joindre l'une à l'autre d'une façon complète et à sens unique, en d'autres termes si le champ de l'une et de l'autre à la même puissance, nous dirons  $a$  et  $b$  équivalentes l'une à l'autre et nous l'exprimerons par une des deux formules  $a \sim b$  ou  $b \sim a$ . D'après cette définition de l'équivalence de deux grandeurs variables il suit immédiatement que  $a \sim a$  ; et que, si  $a \sim b$  et  $b \sim c$ , on a toujours aussi  $a \sim c$ .

Dans la suite du travail le théorème ci-dessous dont nous pouvons omettre la démonstration à cause de sa simplicité, trouvera son application en divers endroits :

(E). “Soit  $a', a'', \dots, a^{(\nu)}$  une série finie ou infinie de variables ou de constantes qui n'ont aucune liaison deux à deux,  $b', b'', \dots, b^{(\nu)}, \dots$  une autre série de la même nature, si à chaque variable  $a^{(\nu)}$  de la première série répond une variable déterminée  $b^{(\nu)}$  de la seconde et si ces variables correspondantes sont constamment équivalentes l'une à l'autre, c'est-à-dire que  $a^{(\nu)} \sim b^{(\nu)}$ , on aura toujours aussi :  $a \sim b$ , si

$$a \equiv \{a', a'', \dots, a^{(\nu)}, \dots\}$$

et

$$b \equiv \{b', b'', \dots, b^{(\nu)}, \dots\}. \quad "$$

#### § 4.

Au point où nous en sommes arrivés de notre travail, il n'y a plus qu'à démontrer le théorème (D) du § 3. Pour cela prenons comme point de départ qu'on peut écrire tous les nombres rationnels qui sont  $\geq 0$  et  $\leq 1$ , sous la forme d'une série simplement infinie :

$$\varphi_1, \varphi_2, \varphi_3, \dots, \varphi_\nu, \dots$$

avec un terme général  $\varphi_\nu$ .

On peut le prouver de la façon la plus simple, comme il suit :  $\frac{p}{q}$  étant la forme irréductible pour un nombre rationnel  $\geq 0$  et  $\leq 1$ , où par conséquent  $p$  et  $q$  sont des nombres entiers non négatifs avec

le plus grand commun diviseur 1, qu'on pose  $p + q = N$ . Dès lors à chaque nombre  $\frac{p}{q}$  appartient une valeur déterminée, entière et positive de  $N$ , et réciproquement à cette valeur de  $N$  appartient toujours un nombre fini de quantités  $\frac{p}{q}$ . Si on imagine maintenant les nombres  $\frac{p}{q}$  rangés dans un ordre tel que ceux qui appartiennent à des valeurs plus petites de  $N$  précèdent ceux pour lesquels  $N$  a une valeur plus grande, et que de plus les nombres  $\frac{p}{q}$ , pour lesquels  $N$  a la même valeur se suivent les uns les autres par ordre de grandeur, les plus grands après les plus petits, chacun des nombres  $\frac{p}{q}$  vient occuper une place parfaitement déterminée dans une série simplement infinie, dont le terme général sera désigné par  $\varphi_\nu$ . Mais cette proposition peut aussi se tirer comme conclusion de ce j'ai dit ailleurs, que l'ensemble  $(\omega)$  de tous les nombres réels algébriques peut se mettre sous la forme d'une série infinie :

$$\omega_1, \omega_2, \dots, \omega_\nu, \dots$$

de terme général  $\omega_\nu$  ; cette propriété de l'ensemble  $(\omega)$  se transmet en effet à l'ensemble de tous les nombres rationnels  $\geq 0$  et  $\leq 1$ , parce que ce dernier ensemble est une partie intégrante du premier  $(\omega)$ .

Soit maintenant  $e$  la variable qui se présente dans le théorème (D) et qui peut prendre toutes les valeurs numériques réelles de l'intervalle  $(0 \dots 1)$ , à l'exception des nombres  $\varphi_\nu$ .

Qu'on prenne ensuite dans l'intervalle  $(0 \dots 1)$  une série quelconque infinie de nombres  $\varepsilon_\nu$  irrationnels et soumise aux conditions qu'en général on a  $\varepsilon_\nu < \varepsilon_{\nu+1}$ , et que  $\lim_{\nu=\infty} \varepsilon_\nu = 1$  ; soit par exemple :

$$\varepsilon_\nu = 1 - \frac{\sqrt{2}}{2^\nu}.$$

Qu'on désigne par  $f$  une variable, qui peut prendre toutes les valeurs réelles de l'intervalle  $(0 \dots 1)$ , à l'exception des valeurs  $\varepsilon_\nu$ , par  $g$  une autre variable, qui peut prendre toutes les valeurs réelles de l'intervalle  $(0..1)$ , à l'exception des  $\varepsilon_\nu$  et des  $\varphi_\nu$ .

Nous disons que :

$$e \sim f.$$

En effet d'après la notation du § 3 on a :

$$e \equiv \{g, \varepsilon_\nu\}.$$

et :

$$f \equiv \{g, \varphi_\nu\}.$$

Mais on a :  $g \sim g$  ;  $\varepsilon_\nu \sim \varphi_\nu$  ; nous concluons donc d'après (E) que  $e \sim f$ .

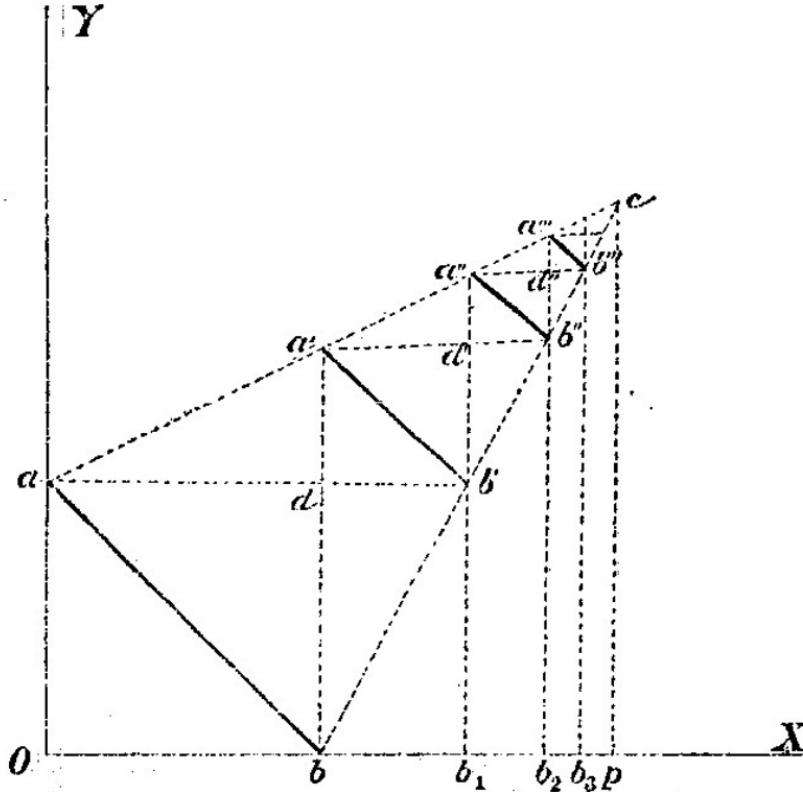
Le théorème à démontrer (D) est donc ramené au théorème suivant :

(F). “Une variable  $f$  qui peut prendre toutes les valeurs de l'intervalle  $(0 \dots 1)$ , à l'exception des valeurs d'une série donnée  $\varepsilon_\nu$ , soumise aux conditions que  $\varepsilon_\nu < \varepsilon_{\nu+1}$  et que  $\lim_{\nu=\infty} \varepsilon_\nu = 1$ , peut se joindre d'une façon complète et à sens unique à une variable  $x$  qui peut prendre toutes les valeurs  $\geq 0$  et  $\leq 1$  ; en d'autres termes, on a  $f \sim x$ ”.

§ 5.

Nous appuyons la démonstration de (F) sur les théorèmes suivants (G), (H), (J) :

(G). “Soit  $y$  une variable, qui peut prendre toutes les valeurs de l'intervalle  $(0 \dots 1)$  à l'exception seulement de 0,  $x$  une variable qui comporte toutes les valeurs de l'intervalle  $(0 \dots 1)$  sans exception, on a :  $y \sim x$ .”.



La démonstration de ce théorème (G) se fait de la manière la plus simple en considérant la courbe ci-avant, dont les abscisses à partir de 0 représentent la grandeur  $x$ , et les ordonnées la grandeur  $y$ . Cette courbe est composée d'un nombre infini de segments de droites  $\overline{ab}, \overline{a'b'}, \dots, \overline{a^{(\nu)}b^{(\nu)}}, \dots$  (qui sont parallèles entre eux et deviennent infiniment petits quand  $\nu$  croît à l'infini) et du point isolé  $c$ , dont ces segments se rapprochent asymptotiquement. Mais les points extrêmes  $a, a', \dots, a^{(\nu)}$  devant être considérés comme faisant partie de la courbe, au contraire les points extrêmes  $b, b', \dots, b^{(\nu)}, \dots$  doivent être regardés comme en dehors de cette courbe. Les longueurs représentées dans la figure sont :

$$\overline{Op} = \overline{pc} = 1 ; \overline{Ob} = \overline{bp} = \overline{Oa} = \frac{1}{2} ;$$

$$\overline{a^{(\nu)}d^{(\nu)}} = \overline{d^{(\nu)}b^{(\nu)}} = \overline{b_{\nu-1}b_{\nu}} = \frac{1}{2^{\nu+1}}.$$

On se convainc que, tandis que l'abscisse  $x$  prend toutes les valeurs de 0 à 1, l'ordonnée  $y$  les prend aussi toutes, à l'exception seulement de la valeur 0.

Le théorème (G) étant ainsi démontré, on obtient en appliquant les formules de transformation :  $y = \frac{z - \alpha}{\beta - \alpha}$  ;  $y = \frac{u - \alpha}{\beta - \alpha}$  la généralisation de (G) :

(H). “Une variable  $z$ , qui peut prendre toutes les valeurs d’un intervalle  $(\alpha \dots \beta)$ , à  $\alpha \geq \beta$ , à l’exception de la seule valeur  $\alpha$ , est équivalente à une variable  $u$  qui peut prendre toutes les valeurs du même intervalle  $(\alpha \dots \beta)$  sans exception.”

De là nous arrivons immédiatement au théorème suivant :

(1). “ $\omega$  étant une variable susceptible de prendre toutes les valeurs de l’intervalle  $(\alpha \dots \beta)$  à l’exception des deux valeurs extrêmes  $\alpha$  et  $\beta$ ,  $u$  étant la même variable que dans (H), on a  $\omega \sim u$ ”.

En effet : soit  $\gamma$  une valeur quelconque entre  $\alpha$  et  $\beta$  ; qu’on introduise comme auxiliaires quatre nouvelles variables  $\omega'$ ,  $\omega''$ ,  $u''$  et  $z$ .

Supposons que  $z$  soit la même variable que dans (H), que  $\omega'$  prenne toutes les valeurs de l’intervalle  $(\alpha \dots \gamma)$  à l’exception des deux valeurs extrêmes  $\alpha$  et  $\gamma$  ; qu’on donne à  $\omega''$  toutes les valeurs de l’intervalle  $(\gamma \dots \beta)$  à l’exception de la seule valeur extrême  $\beta$  ; soit enfin  $u''$  une variable susceptible de prendre toutes les valeurs de l’intervalle  $(\gamma \dots \beta)$  y compris les valeurs extrêmes. On a alors :

$$\begin{aligned}\omega &\equiv \{\omega', \omega''\} \\ z &\equiv \{\omega', u''\}.\end{aligned}$$

Mais par suite de (H) on a  $\omega'' \sim u''$  ; nous concluons donc que  $\omega \sim z$ . Mais d’après (H) on a aussi :  $z \sim u$  ; par conséquent on a encore :  $\omega \sim u$ , ce qui démontre le théorème (J).

Nous pouvons maintenant démontrer le théorème (F) comme il suit :

En renvoyant à la signification des variables  $f$  et  $x$  dans l’énoncé de (F), nous introduisons certaines variables auxiliaires :

$$f', f'', \dots, f^{(\nu)}, \dots$$

et

$$x'', x^{\text{IV}}, \dots, x^{(2\nu)}, \dots$$

Soient :  $f'$  une variable qui comporte toutes les valeurs de l’intervalle  $(0 \dots \varepsilon_1)$  à l’exception de la seule valeur extrême ;  $f^{(\nu)}$  pour  $\nu > 1$  une variable susceptible de prendre toutes les valeurs de l’intervalle  $(\varepsilon_{\nu-1} \dots \varepsilon_\nu)$  à l’exception des deux valeurs extrêmes  $\varepsilon_{\nu-1}$  et  $\varepsilon_\nu$  ;  $x^{(2\nu)}$  une variable qui comporte toutes les valeurs de l’intervalle  $(\varepsilon_{2\nu-1} \dots \varepsilon_{2\nu})$  sans exception.

Si on joint encore aux variables  $f', f'', \dots, f^{(\nu)}, \dots$  la quantité constante 1, toutes grandeurs prises ensemble ont le même champ que  $f$ , c’est-à-dire qu’on a :

$$f \equiv \{f', f'', \dots, f^{(\nu)}, \dots, 1\}$$

De même on se convainc que :

$$x \equiv \{f', x'', f''', x^{\text{IV}}, \dots, f^{(2\nu-1)}, x^{(2\nu)}, \dots, 1\}.$$

Mais par suite du théorème (J) on a :

$$f^{(2\nu)} \sim x^{(2\nu)} ; \text{ puis } : f^{(2\nu-1)} \sim x^{(2\nu-1)} ; 1 \sim 1 ;$$

d'où à cause du théorème (E) § 3 :

$$f \sim x.$$

## § 6.

Je vais maintenant donner une démonstration beaucoup plus courte pour le théorème (D) ; si je ne me suis pas borné à celle-là, cela est venu de ce que les théorèmes auxiliaires (F), (G), (H), (I), qui ont servi à une démonstration plus compliquée, ont de l'intérêt en eux-mêmes.

Nous désignons par  $x$ , comme plus haut, une variable qui peut prendre toutes les valeurs réelles de l'intervalle  $(0 \dots 1)$ , y compris les valeurs extrêmes ; soit  $e$  une variable qui ne comporte que les valeurs irrationnelles de l'intervalle  $(0 \dots 1)$  ; il faut démontrer que  $x \sim e$ .

Nous nous représentons, comme dans le § 4, les nombres rationnels  $\geq 0$  et  $\leq 1$  sous forme de série, avec le terme général  $\varphi_\nu$ , où  $\nu$  a à parcourir la série des nombres 1, 2, 3, ... Nous prenons ensuite dans l'intervalle  $(0 \dots 1)$  une série infinie quelconque de nombres irrationnels distincts entre eux ; soit  $\eta_\nu$  le terme général de cette série (par exemple :  $\eta_\nu = \frac{\sqrt{2}}{2^\nu}$ ).

Qu'on désigne par  $h$  une variable susceptible de prendre toutes les valeurs de l'intervalle  $(0 \dots 1)$  à l'exception des  $\varphi_\nu$  aussi bien que des  $\eta_\nu$ .

D'après le formalisme adopté dans le § 3 on alors :

$$x \equiv \{h, \eta_\nu, \varphi_\nu\} \tag{2}$$

et

$$e \equiv \{h, \eta_\nu\}.$$

Nous pouvons aussi écrire la dernière formule comme il suit :

$$e \equiv \{h, \eta_{2\nu-1}, \eta_{2\nu}\}. \tag{3}$$

Si maintenant nous remarquons que :

$$h \sim h ; \eta_\nu \sim \eta_{2\nu-1} ; \varphi_\nu \sim \eta_{2\nu}$$

et si nous appliquons aux deux formules (1) et (2) le théorème (E) § 3, nous obtenons  $x \sim e$  ;  
c. q. f. d.

## § 7.

La pensée viendrait naturellement de choisir, pour la démonstration de (A) la forme de représentation des fractions décimales infinies au lieu des fractions continues que nous avons employées ; il semblerait que cette méthode nous aurait conduit plus promptement au but ; mais au contraire elle entraîne avec elle une difficulté sur laquelle je veux attirer l'attention ici ; et c'est la raison qui m'a fait renoncer dans ce travail à l'emploi des fractions décimales.

Si on a par exemple deux variables  $x_1$  et  $x_2$  et qu'on pose :

$$x_1 = \frac{\alpha_1}{10} + \frac{\alpha_2}{10^2} + \dots + \frac{\alpha_\nu}{10^\nu} + \dots$$

$$x_2 = \frac{\beta_1}{10} + \frac{\beta_2}{10^2} + \dots + \frac{\beta_\nu}{10^\nu} + \dots$$

en supposant que les nombres  $\alpha_\nu, \beta_\nu$  deviennent des nombres entiers  $\geq 0$  et  $\leq 9$  et ne prennent pas constamment, à partir d'un certain  $\nu$ , la valeur 0 (excepté lorsque  $x_1$  ou  $x_2$  est égal à 0), ces expressions de  $x_1, x_2$  seront déterminées, dans tous les cas avec une signification unique, c'est-à-dire  $x_1$  et  $x_2$ , déterminent les séries infinies de nombres  $\alpha_\nu$  et  $\beta_\nu$ , et réciproquement.

Si maintenant on tire de  $x_1$  et  $x_2$  un nombre :

$$t = \frac{\gamma_1}{10} + \frac{\gamma_2}{10^2} + \dots + \frac{\gamma_\nu}{10^\nu} + \dots,$$

en posant :

$$\gamma_{2\nu-1} = \alpha_\nu ; \gamma_{2\nu} = \beta_\nu,$$

pour

$$\nu = 1, 2, \dots, \infty,$$

il s'établit un rapport à sens unique entre le système  $x_1, x_2$  et la variable  $t$  ; car un seul système de valeurs  $x_1, x_2$  conduit à une valeur donnée de  $t$ .

Mais la variable  $t$ , et c'est la particularité à remarquer ici, ne prend pas toutes les valeurs de l'intervalle  $(0 \dots 1)$ , elle a une variabilité restreinte, tandis que  $x_1$  et  $x_2$  ne sont soumis à aucune restriction dans ce même intervalle. En effet, toutes les valeurs de la somme de la série :

$$\frac{\gamma_1}{10} + \frac{\gamma_2}{10^2} + \dots + \frac{\gamma_\nu}{10^\nu} + \dots,$$

où, à partir d'un certain  $\nu > 1$ , tous les  $\gamma_{2\nu-1}$  ou tous les  $\gamma_{2\nu}$  ont la valeur zéro, doivent être considérées comme au dehors des limites de variabilité de  $t$ , parce qu'elles ramèneraient à des représentations, par fractions décimales, de  $x_1$  ou  $x_2$ , qui sont finies et par conséquent inadmissibles.

## § 8.

Le travail que nous avons en vue étant terminé dans les paragraphes précédents, quelques remarques plus générales pourront trouver place ici, comme conclusion.

Le principe (A) et par suite le principe (B) peuvent être généralisés, de sorte que des ensembles continus d'un nombre infiniment grand de dimensions ont la même puissance que les ensembles continus d'une seule dimension ; toutefois cette généralisation est essentiellement liée à l'hypothèse que les dimensions infiniment nombreuses forment elles-mêmes un ensemble de la première classe ou puissance. Au lieu du théorème (A) on a le suivant :

(A'). *“Soit  $x_1, x_2, \dots, x_\mu$ , une série simplement infinie de grandeurs variables, réelles, indépendantes l'une de l'autre, dont chacune peut prendre toutes les valeurs  $\geq 0$  et  $\leq 1$ , et soit  $t$  une autre variable avec les mêmes limites ( $0 \leq t \leq 1$ ), on peut faire correspondre par une opération à sens unique cette grandeur  $t$  au système des  $x_1, x_2, \dots, x_\mu, \dots$ , qui sont en nombre infini.”*

Ce théorème (A') se ramène, à l'aide du théorème (D) § 3, au suivant :

(C'). *“Soit  $e_1, e_2, \dots, e_\mu, \dots$  une série infinie de grandeurs variables indépendantes l'une de l'autre, dont chacune peut prendre toutes les valeurs numériques irrationnelles de l'intervalle  $(0 \dots 1)$  et soit  $d$  une autre variable irrationnelle avec les mêmes limites, on peut joindre cette grandeur  $d$  par une opération à sens unique au système des grandeurs en nombre infini :  $e_1, e_2, \dots, e_\mu, \dots$ ”*

La démonstration de (C') se fait de la manière la plus simple, en appliquant le développement de la fraction continue et en posant, comme dans le § 2 :

$$e_\mu = (\alpha_{\mu,1}, \alpha_{\mu,2}, \dots, \alpha_{\mu,\nu}, \dots)$$

pour  $\mu = 1, 2, \dots, \infty$

$$d = (\beta_1, \beta_2, \dots, \beta_\lambda, \dots)$$

et en établissant entre les nombres entiers positifs  $\alpha$  et  $\beta$  le rapport :

$$\alpha_{\mu,\nu} = \beta_\lambda,$$

où :

$$\lambda = \mu + \frac{(\mu + \nu - 1)(\mu + \nu - 2)}{2}$$

En effet la fonction  $\frac{(\mu + \nu - 1)(\mu + \nu - 2)}{2}$ , comme il est facile de le montrer, jouit de cette propriété remarquable de représenter tous les nombres entiers positifs, et chacun d'eux une fois seulement, quand  $\mu$  et  $\nu$  y prennent également, indépendamment l'un de l'autre toutes les valeurs positives entières.

Le théorème (A') paraît indiquer le terme jusqu'auquel on peut généraliser le théorème (A) et les conséquences qui en découlent. Et maintenant que nous avons ainsi démontré, pour un champ extraordinairement riche et étendu d'ensembles, la propriété de pouvoir se joindre à sens complet et unique à une droite continue ou à une partie de cette droite (en entendant par parties d'une

ligne tous les ensembles de points qui y sont contenus) la question se pose de savoir comment se comportent, au point de vue de leur puissance, les différentes parties d'une ligne droite continue c'est-à-dire les différents ensembles de points qu'on y peut imaginer en nombre infini.

Si nous dépouillons ce problème de sa forme géométrique et si, comme il a déjà été expliqué au § 3, nous entendons par ensemble linéaire de nombres réels tout ensemble imaginable de quantités réelles distinctes entre elles et en nombre infini, la question se pose ainsi : en quelles classes se divisent les ensembles linéaires, et quel est le nombre de ces classes, si on groupe dans des classes différentes les ensembles de différente puissance, et dans la même les ensembles de même puissance ? Par un procédé d'induction, dans la description duquel nous n'entrerons pas davantage, on est amené à ce théorème, que le nombre des classes d'ensembles obtenus d'après ce mode de groupement est un nombre fini et qu'il est égal à deux.

D'après cela les ensembles linéaires comprendraient deux classes <sup>3</sup> dont la première contient tous les ensembles susceptibles d'être ramenés à la forme : *functio ipsius*  $\nu$  (où  $\nu$  parcourt tous les nombres entiers positifs) ; tandis que la seconde classe embrasse tous les ensembles réductibles à la forme : *functio ipsius*  $x$  (où  $x$  peut prendre toutes les valeurs réelles  $\geq 0$  et  $\leq 1$ ).

Il n'y aurait donc dans les ensembles linéaires infinis, et par conséquent aussi dans tous les autres qui s'y ramènent par une opération à sens complet et unique, que deux espèces de puissances, répondant à ces deux classes ; nous remettons à plus tard la solution exacte de cette question.

Halle a. S. 11 Juillet 1877.

---

<sup>3</sup>Que ces deux classes soient distinctes en réalité, c'est une conséquence immédiate du théorème démontré dans le 2 du travail cité plus haut (Journ. d. Math. pures appliquées t. 77, p. 260), d'après lequel, si on a une série régulière infinie  $u_1, u_2, \dots, u_\nu, \dots$ , on peut toujours trouver dans chaque intervalle donné  $(\alpha \dots \beta)$  des nombres  $\nu$  qui ne se présentent pas dans la série proposée.