

De la puissance des ensembles parfaits de points

Extrait d'une lettre adressée à l'éditeur

G. Cantor

à Halle

... Quant à mon théorème, qui exprime, que les ensembles *parfaits* de points ont tous la même puissance, savoir la puissance du *continu*, je prétends le démontrer, en me bornant d'abord aux ensembles parfaits linéaires ¹, comme il suit. Soit S un ensemble parfait de points quelconque, *qui n'est condensé dans l'étendue d'aucun intervalle*, si petit qu'il soit ; nous admettons, que S est contenu dans l'intervalle $(0 \dots 1)$, dont les points extrêmes 0 et 1 appartiennent à S ; il est évident que tous les autres cas, dans lesquels l'ensemble parfait n'est condensé dans l'étendue d'aucun intervalle, peuvent par projection être réduits à celui-ci.

Or, il existe d'après mes considérations dans Acta mathematica T. 2 page 378 un nombre infini d'intervalles distincts, tout à fait séparés l'un de l'autre, que nous nous représentons rangés suivant leur grandeur, de sorte que les intervalles plus petits viennent après les plus grands ; nous les désignons, dans cet ordre, par :

$$(a_1, \dots b_1), (a_2 \dots b_2), \dots, (a_\nu, \dots b_\nu), \dots; \quad (1)$$

ils sont par rapport à l'ensemble S tels que dans l'intérieur de chacun ne tombe aucun point de S , tandis que leurs points extrêmes a_ν et b_ν en concurrence avec les autres points-limites de l'ensemble de points $\{a_\nu, b_\nu\}$, appartiennent à S et le déterminent ; nous désignons par g l'un quelconque de ces autres points-limites de $\{a_\nu, b_\nu\}$, par $\{g\}$ leur ensemble ; nous avons :

$$S \equiv \{a_\nu\} + \{b_\nu\} + \{g\}. \quad (2)$$

En outre la série (1) d'intervalles est telle que l'espace entre deux d'entre eux $(a_\nu \dots b_\nu)$ et $(a_\mu \dots b_\mu)$ en contient toujours une infinité d'autres et que de plus, $(a_\rho \dots b_\rho)$ étant un quelconque de ces intervalles, il y en a d'autres de la même série (1) qui se rapprochent infiniment soit du point a_ρ , soit du point b_ρ ; car a_ρ et b_ρ comme appartenant comme points à l'ensemble parfait S , en sont aussi des *points-limites*.

Cela établi, je prends un ensemble de la première puissance quelconque :

$$\varphi_1, \varphi_2, \dots, \varphi_\nu, \dots, \quad (3)$$

Acta mathematica. 4. Imprimé le 4 mars 1884.

Transcription en L^AT_EX : Denise Vella-Chemla, juin 2025.

¹M. I. BENDIXSON invité par M. CANTOR à essayer de prouver ce même théorème, en a communiqué une démonstration à la séance du séminaire de l'université de Stockholm, le 21 novembre 1883. Cette démonstration, qui a été trouvée sans que l'auteur ait eu connaissance des recherches que M. CANTOR veut bien me permettre de publier ici, a été présentée à l'Académie royale des sciences de Stockholm, le 12 décembre 1883. Elle se trouve dans Bihang till Svenska Vetenskapsakademiens Handlingar. La démonstration de M. BENDIXSON embrasse le cas d'un ensemble parfait de n dimensions.

L'éditeur.

ensemble de points distincts et placés tous dans l'intervalle $(0 \dots 1)$, *dans toute l'étendue duquel ils sont condensés* ; seulement je suppose que ces points extrêmes 0 et 1 ne se trouvent pas entre les φ_ν .

Pour citer un exemple d'un ensemble tel qu'il nous le faut ici, je rappelle la forme de série, où j'ai mis l'ensemble de tous les nombres rationnels ≥ 0 et ≤ 1 dans Acta mathematica. T. 2 page 319 et où pour notre but il faut supprimer seulement les deux premiers termes, qui y sont 0 et 1.

Mais je tiens à ce que la série (3) soit laissée dans toute sa généralité.

Voici maintenant ce que j'avance : *l'ensemble de points $\{\varphi_\nu\}$ et l'ensemble d'intervalles $\{(a_\nu \dots b_\nu)\}$ peuvent être associés avec un sens unique l'un à l'autre de sorte que, $(a_\nu \dots b_\nu)$ et $(a_\mu \dots b_\mu)$ étant deux intervalles quelconques appartenant à la série (1), puis φ_{k_ν} et φ_{k_μ} étant les points correspondants de la série (3), on a toujours le nombre φ_{k_ν} plus petit ou plus grand que φ_{k_μ} selon que dans le segment $(0 \dots 1)$ l'intervalle $(a_\nu \dots b_\nu)$ est placé avant l'intervalle $(a_\mu \dots b_\mu)$ ou après lui* ².

Une telle correspondance des deux ensembles $\{\varphi_\nu\}$ et $\{(a_\nu \dots b_\nu)\}$ se peut faire par exemple d'après la règle suivante :

Nous associons à l'intervalle $(a_1 \dots b_1)$ le point φ_1 , à l'intervalle $(a_2 \dots b_2)$ le terme au plus petit indice de la série (3), nous le désignons par φ_{k_2} qui a la même relation par rapport au plus ou moins avec φ_1 , que l'intervalle $(a_2 \dots b_2)$ avec $(a_1 \dots b_1)$ par rapport à leur placement dans le segment $(0 \dots 1)$; de plus nous associons à l'intervalle $(a_3 \dots b_3)$ le terme au plus petit indice, qui a la même relation par rapport au plus ou moins avec φ_1 et avec φ_2 , que l'intervalle $(a_3 \dots b_3)$ avec les intervalles $(a_1 \dots b_1)$ et $(a_2 \dots b_2)$ respectivement par rapport à leur placement dans le segment $(0 \dots 1)$.

Généralement nous associons à l'intervalle $(a_\nu \dots b_\nu)$ le terme au plus petit indice de la série (3), nous le nommerons φ_{k_ν} , tel qu'il a la même relation par rapport au plus ou moins avec tous les points $\varphi_1, \varphi_{k_1}, \dots, \varphi_{k_{\nu-1}}$ dont il a été déjà disposé, que l'intervalle $(a_\nu \dots b_\nu)$ avec les intervalles correspondants $(a_1 \dots b_1), (a_2 \dots b_2), \dots (a_{\nu-1} \dots b_{\nu-1})$ par rapport à leur placement dans le segment $(0 \dots 1)$.

J'avance, que d'après cette règle *les points $\varphi_1, \varphi_2, \dots, \varphi_\nu, \dots$ de la suite (3) seront successivement, quoique selon un ordre différent de la loi de la série (3), associés **tous** à des intervalles distincts de la série (1)* ; car à chaque relation par rapport au plus ou moins entre des points en nombre fini de la série (3) il se trouve plusieurs fois une relation conforme par rapport à la place dans le segment $(0 \dots 1)$ entre des intervalles en même nombre de la série (1) ; cela tient à ce que l'ensemble S est un ensemble parfait qui n'est condensé dans aucun intervalle, quelque petit qu'il soit.

Pour simplifier nous poserons :

$$\varphi_1 = \psi_1 ; \varphi_2 = \psi_2 ; \dots ; \varphi_{k_\nu} = \psi_\nu \dots$$

²Il ne s'agit donc pas ici des places ν et μ qu'occupent ces intervalles dans la série (1).

Par conséquent la série suivante :

$$\psi_1, \psi_2, \dots, \psi_\nu, \dots \quad (4)$$

se compose absolument des mêmes éléments que la série (3) ; les deux séries (3) et (4) ne diffèrent que par rapport à la succession de leurs termes.

La série (4) de points ψ_ν a donc ce rapport remarquable avec la série (1) d'intervalles, que toutes les fois que ψ_ν est plus petit ou plus grand que ψ_μ , aussi a_ν et b_ν sont respectivement plus petits ou plus grands que a_μ et b_μ . Et je rappelle de nouveau que l'ensemble $\{\psi_\nu\}$, puisqu'il coïncide avec l'ensemble donné $\{\varphi_\nu\}$, à part la succession des termes, est condensé dans toute l'étendue du segment $(0 \dots 1)$ et que les points extrêmes de celui-ci, 0 et 1, n'appartiennent pas à cet ensemble.

Les conséquences d'une telle association des deux ensembles $\{\psi_\nu\}$ et $\{(a_\nu \dots b_\nu)\}$ sont maintenant, comme il est facile de le démontrer, les suivantes :

Si $(a_{\lambda_1} \dots b_{\lambda_1}), (a_{\lambda_2} \dots b_{\lambda_2}), \dots, (a_{\lambda_\nu} \dots b_{\lambda_\nu}), \dots$ est une série quelconque d'intervalles appartenants à la série (1), qui convergent infiniment soit vers le point a_ρ , soit vers le point b_ρ , alors la série correspondante de points $\psi_{\lambda_1}, \psi_{\lambda_2}, \dots, \psi_{\lambda_\nu}, \dots$, appartenant tous à la série (4), converge infiniment vers le point ψ_ρ , et réciproquement.

Si $(a_{\lambda_1} \dots b_{\lambda_1}), (a_{\lambda_2} \dots b_{\lambda_2}), \dots, (a_{\lambda_\nu} \dots b_{\lambda_\nu}), \dots$ est une série quelconque de la même espèce, mais telle que ses termes convergent infiniment vers un point g de l'ensemble S (voir la formule (2) et la signification de g), alors la série correspondante $\psi_{\lambda_1}, \psi_{\lambda_2}, \dots, \psi_{\lambda_\nu}, \dots$ à son tour converge infiniment vers un point déterminé du segment $(0 \dots 1)$, qui ne coïncide avec aucun point de la série (3) ou (4) et qui de plus est entièrement déterminé par g ; nous désignerons ce point correspondant à g par h ; réciproquement soit h un point quelconque du segment $(0 \dots 1)$, qui n'appartient pas à la série (3) ou (4) il détermine un point g de l'ensemble S différent des points a_ν et b_ν ; en sorte que les deux nombres variables g et h sont des fonctions à sens unique l'une de l'autre et que les ensembles $\{g\}$ et $\{h\}$ par suite sont certainement de la même puissance.

De là suit la démonstration du théorème en question.

Car nous avons d'après la formule (2) :

$$S \equiv \{a_\nu\} + \{b_\nu\} + \{g\}.$$

Puis il est évident que :

$$(0 \dots 1) \equiv \{\varphi_{2\nu}\} + \{\varphi_{2\nu-1}\} + \{h\}.$$

Mais comme on a les formules suivantes :

$$\{a_\nu\} \sim \{\varphi_{2\nu}\} ; \{b_\nu\} \sim \{\varphi_{2\nu-1}\} ; \text{ et } \{g\} \sim \{h\}$$

on conclut d'après le théorème (E) des Acta Mathematica T. 2 p. 318 la formule :

$$S \sim (0 \dots 1)$$

c'est-à-dire l'ensemble parfait S a la même puissance que le segment continu $(0 \dots 1)$; ce qui était à démontrer.

... Cette démonstration a l'avantage de nous dévoiler une grande classe remarquable de fonctions *continues* d'une variable réelle x , dont les propriétés donnent lieu à des recherches intéressantes, soit en les considérant d'après la définition, qui se rattache à notre développement, *soit en tâchant de les mettre sous la forme de séries trigonométriques, qui certainement leur sont conformes, parce que ces fonctions continues ne jouissent pas d'un nombre infini de maxima et minima.*

En effet nous pouvons établir dans l'intervalle $(0 \dots 1)$ une fonction $\psi(x)$ satisfaisant aux conditions suivantes :

Lorsque x est compris dans l'un quelconque des intervalles $(a_\nu \dots b_\nu)$ c'est-à-dire pour $a_\nu \leq x \leq b_\nu$, $\psi(x)$ est égale à ψ_ν ; lorsque x reçoit une valeur g qui s'obtient comme limite d'une série d'intervalles $(a_{\lambda_1} \dots b_{\lambda_1}), \dots (a_{\lambda_\nu} \dots b_{\lambda_\nu}), \dots$ alors on définit :

$$\psi(g) = h = \lim_{\nu=\infty} \psi_{\lambda_\nu}. \quad (5)$$

Certes, la fonction $\psi(x)$, d'après ce que nous avons vu, est une fonction continue, monotone ³ de la variable continue x ; lorsque x croît de 0 à 1, $\psi(x)$ varie d'une manière continue sans diminuer de 0 à 1 ; son image géométrique se compose d'un ensemble scalariforme de segments droits, tous parallèles à l'axe des x et de certains points interposés, qui font, que cette courbe devient un continu. Un cas particulier de ces fonctions est déjà compris dans un exemple, que j'ai mentionné dans Acta mathematica T. 2, page 407. En posant :

$$z = \frac{c_1}{3} + \frac{c_2}{3^2} + \dots + \frac{c_\rho}{3^\rho} + \dots, \quad (6)$$

où les coefficients c_μ peuvent prendre à volonté les deux valeurs 0 et 2 et où la série peut être composée d'un nombre fini ou infini de membres, l'ensemble $\{z\}$ est un ensemble *parfait* S , situé dans l'intervalle $(0 \dots 1)$, les points extrêmes 0 et 1 appartiennent à cet ensemble $\{z\}$; de plus l'ensemble $\{z\} = S$ est ici tel, qu'il n'est condensé dans l'étendue d'aucun intervalle, si petit qu'il soit ; enfin on peut aussi s'assurer, que cet ensemble $S = \{z\}$ a une *grandeur* $\mathfrak{I}(S)$ (*notion* que j'expliquerai à l'instant) égale à zéro.

Ici les points, que nous avons désignés par b_ν résultent de la formule (6) pour z en prenant $c_\rho = 0$ à partir d'un certain ρ plus grand que 1, en sorte que tous les b_ν sont compris dans la formule :

$$b_\nu = \frac{c_1}{3} + \frac{c_2}{3^2} + \dots + \frac{c_{\mu-1}}{3^{\mu-1}} + \frac{2}{3^\mu}. \quad (7)$$

Les points a_ν résultent de la même formule pour z , en prenant c_ρ à partir d'un certain ρ toujours égal à 2, en sorte qu'en vertu de l'équation :

$$1 = \frac{2}{3} + \frac{2}{3^2} + \frac{2}{3^3} + \dots$$

³C'est une expression introduite par M. CH. NEUMANN (voir *Ueber die nach Kreis-, Kugel- und Cylinder-functionen fortschreitenden Entwicklungen.* Leipzig 1881, p. 26).

on a, en prenant $c_\mu = 0, c_{\mu+1} = c_{\mu+2} = \dots = 2$,

$$a_\nu = \frac{c_1}{3} + \frac{c_2}{3^2} + \dots + \frac{c_{\mu-1}}{3^{\mu-1}} + \frac{1}{3^\mu}. \quad (8)$$

Joignons maintenant la variable z à une autre y , définie par la formule :

$$y = \frac{1}{2} \left(\frac{c_1}{2} + \frac{c_2}{2^2} + \dots + \frac{c_\rho}{2^\rho} + \dots \right) \quad (9)$$

dans laquelle nous convenons, que les coefficients c_ρ ont la même valeur que dans (6).

Par cette liaison y devient évidemment une fonction de z , que nous appelons $\psi(z)$. Remarquons maintenant que les deux valeurs de $\psi(z)$ pour $z = a_\nu$, et pour $z = b_\nu$ deviennent égales, savoir :

$$\psi(a_\nu) = \psi(b_\nu) = \frac{1}{2} \left(\frac{c_1}{2^1} + \frac{c_2}{2^2} + \dots + \frac{c_{\mu-1}}{2^{\mu-1}} + \frac{2}{2^\mu} \right)$$

De là résulte une fonction continue et monotone $\psi(x)$ de la variable continue x , définie de la manière suivante :

Pour $a_\nu < x < b_\nu$, on pose : $\psi(x) = \psi(a_\nu) = \psi(b_\nu)$, et pour $x = z$, on a $\psi(x) = y = \psi(z)$.

M. L. SCHEEFFER à Berlin a observé, que cette fonction $\psi(x)$, ainsi que beaucoup d'autres, est en contradiction avec un théorème de M. HARNACK (v. Math. Annalen Bd. 19, page 241, Lehrs. 5). En effet cette fonction $\psi(x)$ a sa dérivée $\psi'(x)$ égale à zéro pour toutes les valeurs de x , à l'exception de celles que nous avons nommées z ; celles-ci constituent un ensemble parfait $\{z\}$, dont la grandeur $\mathfrak{J}(\{z\})$ est égale à zéro. Mais M. SCHEEFFER m'a aussi dit, qu'il pouvait remplacer ce théorème par un autre, qui serait exempt de doute ; j'espère qu'il publiera bientôt dans les Acta ses recherches sur ce sujet aussi bien que sur diverses autres questions intéressantes, dont il s'occupe.

... Dans ce qui précède j'ai démontré, que tous les ensembles parfaits et linéaires de points, qui ne sont condensés dans aucune partie du segment, dans lequel ils sont placés, si petite qu'elle soit, sont de la même puissance que le continu linéaire.

Prenons maintenant un ensemble parfait et linéaire de points S quelconque, placé dans l'intervalle $(-\omega \dots +\omega)$ je dis qu'également cet ensemble S a la puissance du continu $(0 \dots 1)$.

En effet, comme nous avons déjà traité le cas, où l'ensemble S n'est condensé dans aucune partie continue du segment $(-\omega \dots +\omega)$, prenons un intervalle quelconque $(c \dots d)$, dans l'intérieur duquel S soit condensé partout. Tous les points de $(c \dots d)$ appartiendront aussi à S , parce que S est un ensemble parfait.

L'ensemble de points $(c \dots d)$ est un système partiel de S et S un système partiel du segment $(-\omega \dots + \omega)$. Comme l'ensemble $(c \dots d)$ a la même puissance que l'ensemble $(-\omega \dots + \omega)$, on en conclut aussi, que S a la même puissance que $(-\omega \dots + \omega)$, c'est-à-dire la puissance de $(0 \dots 1)$; car on a le théorème général :

Étant donné un ensemble bien défini M d'une puissance quelconque, un ensemble partiel M' pris dans M et un ensemble partiel M'' pris dans M' , si le dernier système M'' possède la même puissance que le premier M , l'ensemble moyen M' est aussi toujours de la même puissance que M et M'' . (Voir Acta mathematica, T. 2, page 392).

Lorsqu'un ensemble P est tel, que son premier ensemble dérivé $P^{(1)}$ en est diviseur, je nomme P un *ensemble fermé*.

Chaque ensemble fermé P d'une puissance supérieure à la première se décompose, comme nous le savons, d'une seule manière en un ensemble R de la première puissance et en un ensemble parfait S . On en conclut au moyen des théorèmes obtenus, le suivant : *Tous les ensembles fermés de points se divisent en deux classes, les uns sont de la première puissance, les autres ont la puissance du continu arithmétique*. Dans une prochaine communication je montrerai que cette division en deux classes a aussi lieu pour les ensembles de points non fermés. Par là nous arriverons à l'aide des principes du § 13 de mon mémoire dans Acta mathematica T. 2, page 390, à la détermination de la *puissance du continu arithmétique*, en démontrant qu'elle coïncide avec celle de la *deuxième classe des nombres* (II).

Il y a une *notion de volume* ou de *grandeur*, qui se rapporte à tout ensemble P , situé dans un espace plan G_n à n dimensions, que cet ensemble P soit continu ou non.

Dans le cas où P se réduit à un ensemble continu à n dimensions, ou à un système de tels ensembles, cette notion se confond avec la notion ordinaire de volume.

Lorsque P est un continu à un nombre de dimensions plus petit que n la valeur du volume devient zéro ; la même chose arrive lorsque P est tel que $P^{(1)}$ a la première puissance et encore dans divers autres cas. Mais ce qui, au premier moment, paraîtra peut être étonnant, c'est que ce volume, je le désigne par $\mathfrak{J}(P)$, a quelquefois une valeur différente de zéro pour des ensembles P contenus dans G_n de l'espèce de ceux, qui ne sont condensés dans aucune partie continue à n dimensions de G_n , si petite qu'elle soit.

J'arrive à cette notion générale de *volume* ou de *grandeur* $\mathfrak{J}(P)$ d'un ensemble *quelconque* P contenu dans G_n en prenant *chaque* point p , qui appartient à P ou à $P^{(1)}$, pour centre d'une sphère pleine à n dimensions au rayon ρ , que nous appellerons $K(p, \rho)$. Le plus petit multiple de toutes ces sphères pleines $K(p, \rho)$ (voir la définition du plus petit multiple, Acta mathematica T. 2, page 357) savoir :

$$\mathfrak{M}[K(p, \rho)],$$

(où ρ est une constante) constitue pour chaque valeur de ρ un ensemble qui se compose de pièces continues à n dimensions et dont le volume se détermine d'après les règles connues au moyen d'une

intégrale n -tuple.

Soit $f(\rho)$ la valeur de cette intégrale ; $f(\rho)$ est une fonction continue de ρ , qui diminue avec ρ ; la limite de $f(\rho)$, lorsque ρ converge vers zéro, me sert de définition du volume $\mathfrak{I}(P)$; en sorte, que nous avons :

$$\mathfrak{I}(P) = \lim_{\rho=0} f(\rho). \quad (10)$$

Je fais remarquer expressément que cette valeur du *volume* ou de la *grandeur* d'un ensemble quelconque P contenu dans un espace continu plan G_n à n dimensions est absolument dépendante de l'espace plan G_n même, duquel P est considéré comme une partie composante, et particulièrement du nombre n ; de sorte que, si l'on considère *le même ensemble* P comme une partie constituante d'un autre espace continu plan H_m la valeur du volume de P par rapport à l'espace H_m est en général différente de celle, qui se rapporte au même ensemble P , considéré comme partie constitutive de G_n .

Un carré p. e. dont le côté est égal à l'unité, a sa *grandeur* égale à zéro lorsqu'il est considéré comme partie constituante de l'espace à trois dimensions, mais il a la *grandeur* égale à 1, lorsqu'on le regarde comme partie d'un plan à deux dimensions. Cette notion générale de *volume* ou de *grandeur* m'est indispensable dans les recherches sur les *dimensions* des *ensembles continus*, que j'ai promises dans Acta mathematica T. 2, page 407 et que je vous enverrai plus tard pour votre journal.

En nous bornant ici aux ensembles *linéaires* de points, compris dans l'intervalle $(0 \dots 1)$, le *volume* ou la *grandeur* d'un tel ensemble P se détermine facilement en suivant la méthode exposée dans Acta mathematica T. 2, page 378, où nous avons considéré des intervalles, désignés par $(c_\nu \dots d_\nu)$ et liés d'après une loi manifeste à P et $P^{(1)}$ ou, comme je l'ai exprimé là, à $\mathfrak{M}(P, P^{(1)})$. Nous y avons posé :

$$\sum (d_\nu - c_\nu) = \sigma$$

où σ est une quantité déterminée positive ≤ 1 . Or dans notre cas on se convaincra facilement, que l'on a :

$$\mathfrak{I}(P) = 1 - \sigma \quad (11)$$

... Les ensembles linéaires parfaits de points S , qui ne sont condensés dans aucun intervalle, si petit qu'il soit, ont en général une *grandeur* $\mathfrak{I}(S)$ différente de zéro, mais il peuvent aussi avoir une *grandeur* $\mathfrak{I}(S)$ égale à zéro.

Quant à ceux, pour lesquels $\mathfrak{I}(S)$ est différent de zéro, ils peuvent être réduits par composition (addition) et à ceux pour lesquels $\mathfrak{I}(S) = 0$ et à de tels ensembles parfaits, qui non seulement sont d'une *grandeur* différente de zéro, mais dont toutes les parties *parfaites*, que l'on obtient en se bornant à des intervalles partiels de $(0 \dots 1)$, ont à leur tour une *grandeur* différente de zéro.

Pour *cette dernière classe* d'ensembles parfaits linéaires il y a une démonstration très simple du théorème démontré plus haut, que leur puissance est celle du continu.

En effet prenons un tel ensemble parfait S dans l'intervalle $(0 \dots 1)$ et supposons que les points extrêmes 0 et 1 appartiennent à S ; nous établissons d'abord la série (1) d'intervalles $(a_\nu \dots b_\nu)$, appartenant dans le sens expliqué à l'ensemble parfait S .

Soit x une grandeur quelconque > 0 et ≤ 1 , nous désignons par S_x l'ensemble, qui est constitué par tous les points de S , qui sont situés dans l'intervalle $(0 \dots x)$ et définissons une fonction $\varphi(x)$ par les conditions suivantes :

$$\varphi(0) = 0, \quad \varphi(x) = \mathfrak{I}(S_x) \quad \text{pour } x > 0 \text{ et } \leq 1.$$

Cette fonction $\varphi(x)$ est, comme on le voit sans peine, continue et monotone dans l'intervalle $(0 \dots 1)$; pour la valeur $x = 1$ elle prend la valeur $\varphi(1) = \mathfrak{I}(S) = c$, différente de zéro d'après l'hypothèse faite par rapport à S . De plus, dans chacun des intervalles $(a_\nu \dots b_\nu)$, c'est-à-dire pour $a_\nu \leq x \leq b_\nu$ elle conserve une valeur constante $\varphi(x) = \varphi(a_\nu) = \varphi(b_\nu)$; lorsque x est plus petit que a_ν on a toujours $\varphi(x) < \varphi(a_\nu)$ lorsque x est plus grand que b_ν on a $\varphi(x) > \varphi(b_\nu)$; *cela tient à ce que nous avons supposé un ensemble S tel que tout ensemble partiel parfait, que l'on obtient en se bornant à des intervalles partiels de $(0 \dots 1)$ est à son tour d'une grandeur différente de zéro.*

La fonction continue $\varphi(x)$ prend toutes les valeurs entre 0 et c ; elle prend chaque valeur entre celles qui sont égales à $\varphi(a_\nu) = \varphi(b_\nu)$ un nombre infini de fois, savoir pour tous les x , qui sont $\leq a_\nu$ et $\leq b_\nu$; mais elle ne prend *qu'une seule* fois chaque valeur h de l'intervalle $(0 \dots c)$, qui est différente des valeurs $\varphi(a_\nu) = \varphi(b_\nu)$, pour une valeur *distincte* g de x , où g diffère de toutes les valeurs appartenant aux intervalles $(a_\nu \dots b_\nu)$, soit des valeurs extrêmes a_ν et b_ν soit des intermédiaires.

Et puisque à chacune de ces valeurs g de x il appartient une certaine valeur $h = \varphi(g)$, différente des valeurs $\varphi(a_\nu) = \varphi(b_\nu)$, et vice versa, on a comme dans notre première démonstration :

$$\{g\} \sim \{h\}$$

d'où l'on conclut comme plus haut, que la puissance de S est celle du continu $(0 \dots c)$.

... Après avoir obtenu ces résultats je suis revenu à mes recherches sur les séries trigonométriques, que j'ai publiées il y a maintenant treize ans et que j'avais laissées de côté depuis longtemps ; non seulement je suis parvenu à démontrer, que le théorème Acta mathematica T. 2 page 348 reste juste, lorsque le système de points, que j'y ai désigné par P , est tel, que son ensemble dérivé $P^{(1)}$ a la première puissance, mais je possède maintenant même quelques résultats pour le cas où $P^{(1)}$ est d'une puissance plus grande que la première ; je vous les enverrai une autre fois.

Halle, 15 Novembre 1883.