

M : A.T.H.



MEMOSYNE

UNIVERSITE PARIS VII

Mnémosyne

personnification de la mémoire.

Elle s'unit à Zeus pendant 9 nuits de suite;
de cette union naquirent les neuf Muses.

(Dictionnaire Robert des noms propres)

Illustration de la couverture : **"La mémoire"**
gravure allégorique d'après Gravelot (XVIII ème)

MEMOSYNE

M: *Mathématiques*

A. *Approche par les*

T. *textes*

H. *historiques*



SOMMAIRE

Editorial p.5

Iconographie Sébastien Le Clerc p.6

Bonnes vieilles pages Wantzel p.7

*' Recherche sur les moyens de reconnaître si un problème
de géométrie peut se résoudre avec la règle et le compas'*

Conte du Lundi Euclide a encore quelque chose à dire p.15

Etude Fragments d'une étude des systèmes linéaires. p.19

Dans nos classes

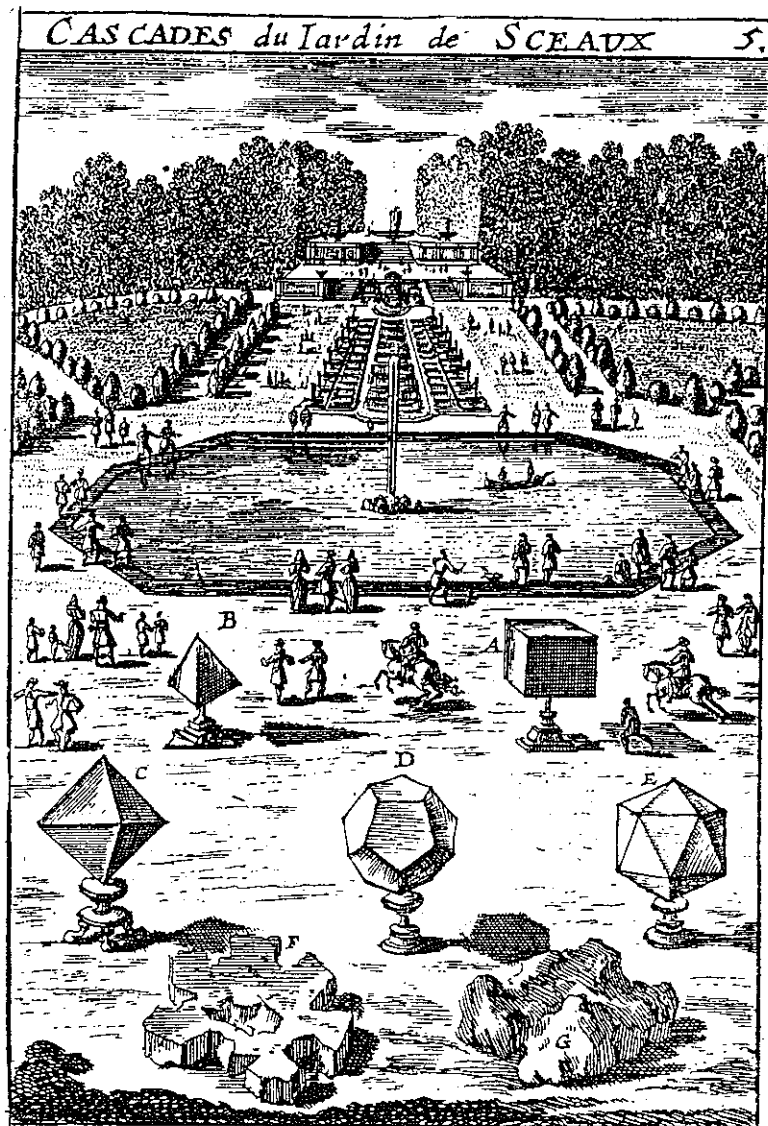
A propos du volume de la pyramide p.59

Notes d'écoute p. 62

Note de lecture p.65

Calendrier

PLANCHE II.



A iij

EDITORIAL

Voici donc le troisième numéro de Mnemosyne . Sa parution coïncide avec le premier anniversaire de la création de cette publication périodique.

Les numéros à venir sont déjà en chantier et même bien avancés pour les prochains :

Le numéro 4 se doit d'être présent à l'Université d'Eté qui se tiendra en juillet à Montpellier.

Quant au numéro 5, il sera disponible pour la rentrée universitaire.

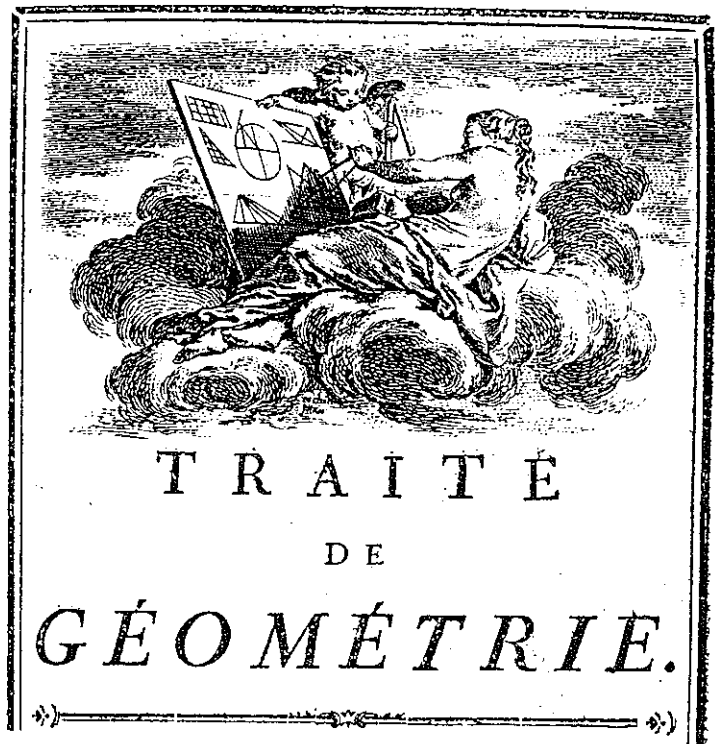
Cette fois, les " *bonnes vieilles pages* " reproduisent un texte souvent cité, mais introuvable car jamais réédité, de Wantzel. C'est une étape importante dans l'histoire de la géométrie car l'auteur s'attache à définir algébriquement l'exigence majeure de la géométrie grecque: n'utiliser que la règle et le compas.

Vous trouverez également dans ce numéro un travail très complet (bien qu'il s'intitule modestement "fragments") sur l'étude des *systèmes d'équations linéaires* qui vous fait parcourir quatre millénaires et bien des régions de notre planète.

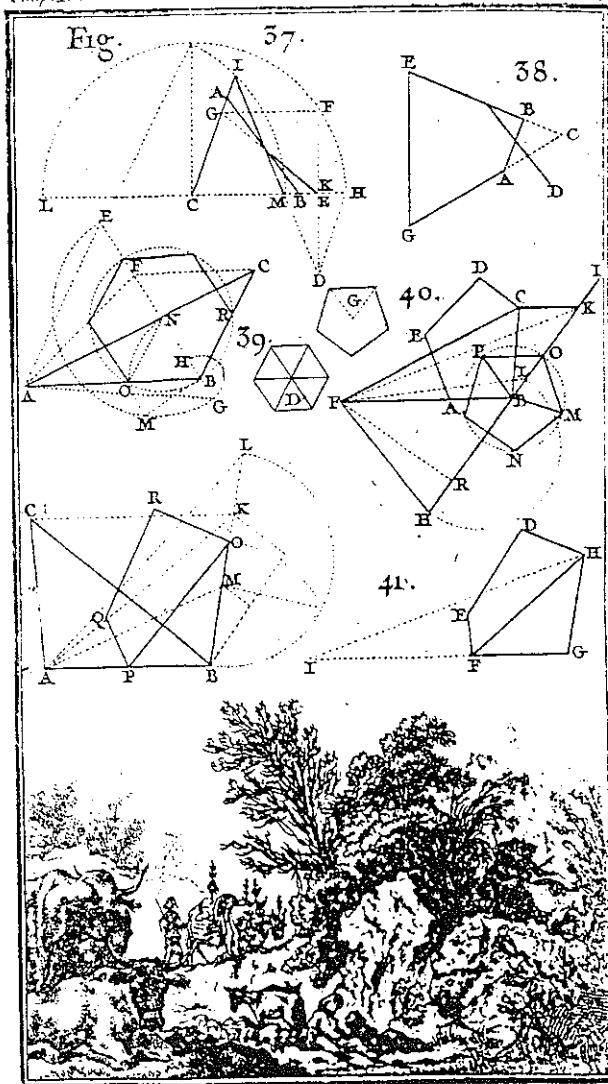
Quant au calcul du *volume de la pyramide*, problème dans lequel un passage à la limite est incontestable, il se prête bien à la réflexion dans les classes. Nous publierons prochainement une étude sur l'histoire de ce problème qui a traversé la géométrie, d'Euclide à Hilbert.

Ainsi les différentes rubriques annoncées dans le premier numéro sont en place; toutes sauf une qui reste désespérément vide: le courrier des lecteurs! Nous avons bien quelques remarques orales mais, prenez votre plume afin de nous faire part de vos suggestions, compléments d'information, remarques voire critiques et de ce que vous désirez trouver dans cette revue.

Jean Luc Verley



Chap. II. Géométrie de le Clerc. Planché 4.



Traité de GEOMETRIE
théorique et pratique,
à l'usage des artistes.

par Sébastien Le Clerc
ed. 1774

BONNES VIEILLES PAGES

Wantzel est encore élève ingénieur à l'Ecole des Ponts et Chaussées quand il publie un article, qui le rendra pour une part célèbre, que nous donnons ci-après *in extenso*, et qui paraît en 1837 dans le journal de Liouville ¹ sous le titre : "*Recherche sur les moyens de reconnaître si un problème de géométrie peut se résoudre avec la règle et le compas.*"

En exceptant le cas de Gauss, que Wantzel cite à la fin de son article, c'est la première vraie tentative pour *algébriser* les problèmes de constructibilité à la règle et au compas et donc la quadrature : c'est à dire par exemple qu'au problème de la construction géométrique d'un carré qui limite la même aire qu'un cercle donné (énoncé traditionnel de la quadrature), Wantzel substitue un problème d'équations : reconnaître le type d'équation que doivent vérifier les coordonnées d'un point pour être constructible à la règle et au compas. Dans cet article, il utilise essentiellement deux arguments mathématiques en sens inverse: la substitution et l'élimination.

Etudiant les équations vérifiées par les abscisses des points d'intersection de cercles et de droites, Wantzel montre par substitution que si un point M (dont on s'intéresse par exemple à l'abscisse : x_n) est constructible à n pas à partir d'un ensemble initial de points à coordonnées rationnelles (posées comme connues et constructibles), alors x_n est nécessairement solution d'une équation algébrique irréductible à coefficients rationnels de degré 2^n . Tout ce qui est constructible est donc algébrique, mais la réciproque est fautive : les constructibles ne sont qu'une partie des algébriques. Ce résultat règle donc (enfin !) le sort de trois problèmes " grecs " de constructibilité restés longtemps en suspens :

la duplication du cube, qui s'algébrique suivant l'équation

$$x^3 - 2a^3 = 0 \quad (a \text{ rationnel})^2 \text{ est donc impossible}$$

tout comme la trisection de l'angle équivalente à

$$x^3 - \frac{3}{4}x + \frac{a}{4} = 0, \text{ ainsi que l'insertion de deux moyennes proportionnelles}$$

(équivalente à : $x^3 = a^2 b$).

Dans la partie IV, Wantzel traite, à la suite de Gauss, la question de la division de la circonférence en parties égales, c'est à dire l'inscriptibilité dans un cercle d'un polygone régulier à N côtés à l'aide de la règle et du compas. Reprenant d'abord et par sa propre méthode, la condition de Gauss: si N est un nombre premier, il doit être de la forme $1 + 2^n$, Wantzel établit ensuite, pour N quelconque, une condition plus générale. Pour intéressants que soient ces résultats, ils laissent évidemment en suspens une question, celle de la quadrature du cercle, puisqu'on ne sait rien d'une éventuelle équation algébrique vérifiée par π , que l'on soupçonne au contraire de n'en vérifier aucune, c'est à dire d'être transcendant³.

Michel SERFATI

¹ Journal de Mathématiques Pures et Appliquées (1837), pages 366-372.

² Sauf quelques exceptions, cette équation est irréductible, de degré trois.

³ Sur cet article de Wantzel, ainsi que la question de la Quadrature du cercle après 1675, cf. M. Serfati: " Quadrature du cercle, Fractions Continues et autres contes." Brochure A.P.M.E.P. " Fragments des Mathématiques, IV "

Recherches sur les moyens de reconnaître si un Problème de Géométrie peut se résoudre avec la règle et le compas ;

PAR M. L. WANTZEL,

Élève-Ingénieur des Ponts-et-Chaussées.

I.

Supposons qu'un problème de Géométrie puisse être résolu par des intersections de lignes droites et de circonférences de cercle : si l'on joint les points ainsi obtenus avec les centres des cercles et avec les points qui déterminent les droites on formera un enchaînement de triangles rectilignes dont les éléments pourront être calculés par les formules de la Trigonométrie ; d'ailleurs ces formules sont des équations algébriques qui ne renferment les côtés et les lignes trigonométriques des angles qu'au premier et au second degré ; ainsi l'inconnue principale du problème s'obtiendra par la résolution d'une série d'équations du second degré dont les coefficients seront fonctions rationnelles des données de la question et des racines des équations précédentes. D'après cela, pour reconnaître si la construction d'un problème de Géométrie peut s'effectuer avec la règle et le compas, il faut chercher s'il est possible de faire dépendre les racines de l'équation à laquelle il conduit de celles d'un système d'équations du second degré composées comme on vient de l'indiquer. Nous traiterons seulement ici le cas où l'équation du problème est algébrique.

II.

Considérons la suite d'équations :

$$(A) \begin{cases} x_1^2 + Ax_1 + B = 0, & x_1^2 + A_1x_1 + B_1 = 0 \dots x_{n-1}^2 + A_{n-1}x_{n-1} + B_{n-1} = 0, \\ & x_n^2 + A_nx_n + B_n = 0, \end{cases}$$

dans lesquelles A et B représentent des fonctions rationnelles des quantités données $p, q, r, \dots$; A_1 et B_1 des fonctions rationnelles de $x_1, p, q, \dots$; et, en général, A_n et B_n des fonctions rationnelles de $x_n, x_{n-1}, \dots, x_1, p, q, \dots$.

Toute fonction rationnelle de x_n telle que A_n ou B_n , prend la forme $\frac{C_{n-1}x_n + D_{n-1}}{E_{n-1}x_n + F_{n-1}}$ si l'on élimine les puissances de x_n supérieures à la pre-

mière au moyen de l'équation $x_n^2 + A_{n-1}x_n + B_{n-1} = 0$, en désignant par C_{n-1} , D_{n-1} , E_{n-1} , F_{n-1} , des fonctions rationnelles de $x_{n-1}, \dots, x_1, p, q, \dots$; elle se ramènera ensuite à la forme $A'_{n-1}x_n + B'_{n-1}$, en multipliant les deux termes de $\frac{C_{n-1}x_n + D_{n-1}}{E_{n-1}x_n + F_{n-1}}$ par $-E_{n-1}(A_{n-1} + D_{n-1}) + F_{n-1}$.

Multiplions l'une par l'autre les deux valeurs que prend le premier membre de la dernière des équations (A) lorsqu'on met successivement à la place de x_{n-1} , dans A_{n-1} et B_{n-1} , les deux racines de l'équation précédente : nous aurons un polynome du quatrième degré en x_n dont les coefficients s'exprimeront en fonction rationnelle de $x_{n-1}, \dots, x_1, p, q, \dots$; remplaçons de même successivement dans ce polynome x_{n-1} par les deux racines de l'équation correspondante, nous obtiendrons deux résultats dont le produit sera un polynome en x_n de degré 2^2 , à coefficient rationnel par rapport à $x_{n-1}, \dots, x_1, p, q, \dots$; et, en continuant de la même manière, nous arriverons à un polynome en x_n de degré 2^n dont les coefficients seront des fonctions rationnelles de $p, q, r, \dots$. Ce polynome égalé à zéro donnera l'équation finale $f(x_n) = 0$ ou $f(x) = 0$, qui renferme toutes les solutions de la question. On peut toujours supposer qu'avant de faire le calcul on a réduit les équations (A) au plus petit nombre possible. Alors une quelconque d'entre elles $x_{n+1}^2 + A_n x_{n+1} + B_n = 0$, ne peut pas être satisfaite par une fonction rationnelle des quantités données et des racines des équations précédentes. Car, s'il en était ainsi, le résultat de la substitution serait une fonction rationnelle de $x_n, \dots, x_1, p, q, \dots$ qu'on peut mettre sous la forme $A'_{n-1}x_n + B'_{n-1}$, et l'on aurait $A'_{n-1}x_n + B'_{n-1} = 0$; on tirerait de cette relation une valeur rationnelle de x_n qui substituée dans l'équation du second degré en x_n conduirait à un résultat de la forme $A'_{n-1}x_{n-1} + B'_{n-1} = 0$. En continuant ainsi, on arriverait à $A'x_1 + B' = 0$, c'est-à-dire que l'équation $x_1^2 + Ax_1 + B = 0$ aurait pour racines des fonctions rationnelles de $p, q, \dots$; le système des équations (A) pourrait donc être remplacé par deux systèmes de $n-1$ équations du second degré, indépendants l'un de l'autre, ce qui est contre la supposition. Si l'une des relations intermédiaires $A'_{n-1}x_{n-1} + B'_{n-1} = 0$, par exemple, était satisfaite identiquement, les deux racines de l'équation $x_n^2 + A_{n-1}x_n + B_{n-1} = 0$ seraient des fonctions rationnelles de $x_{n-1}, \dots, x_1$, pour toutes les valeurs que peuvent prendre ces quantités, en sorte qu'on pourrait supprimer l'équation en x_n et remplacer la racine successivement par ses deux valeurs dans les équations sui-

vantes, ce qui ramènerait encore le système des équations (A) à deux systèmes de $n-1$ équations.

III.

Cela posé, l'équation du degré 2^n , $f(x) = 0$, qui donne toutes les solutions d'un problème susceptible d'être résolu au moyen de n équations du second degré, est nécessairement irréductible, c'est-à-dire qu'elle ne peut avoir de racines communes avec une équation de degré moindre dont les coefficients soient des fonctions rationnelles des données $p, q, \dots$.

En effet, supposons qu'une équation $F(x) = 0$, à coefficients rationnels soit satisfaite par une racine de l'équation $x_n^2 + A_{n-1}x_n + B_{n-1} = 0$, en attribuant certaines valeurs convenables aux quantités $x_{n-1}, x_{n-2}, \dots, x_1$. La fonction rationnelle $F(x_n)$ d'une racine de cette dernière équation peut se ramener à la forme $A'_{n-1}x_n + B'_{n-1}$, en désignant toujours par A'_{n-1} et B'_{n-1} des fonctions rationnelles de $x_{n-1}, \dots, x_1, p, q, \dots$; de même A'_{n-1} et B'_{n-1} peuvent prendre l'un et l'autre la forme $A'_{n-2}x_{n-1} + B'_{n-2}$, et ainsi de suite; on arrivera ainsi à $A'_1x_1 + B'_1$ où A'_1 et B'_1 peuvent être mis sous la forme $A'x_1 + B'$ dans laquelle A' et B' représentent des fonctions rationnelles des données $p, q, \dots$. Puisque $F(x_n) = 0$ pour une des valeurs de x_n , on aura $A'_{n-1}x_n + B'_{n-1} = 0$, et il faudra que A'_{n-1} et B'_{n-1} soient nuls séparément, sans quoi l'équation $x_n^2 + A_{n-1}x_n + B_{n-1} = 0$ serait satisfaite pour la valeur $-\frac{B'_{n-1}}{A'_{n-1}}$ qui est une fonction rationnelle de $x_{n-1}, \dots, x_1, p, q, \dots$, ce qui est impossible; de même, A'_{n-1} et B'_{n-1} étant nuls, A'_{n-2} et B'_{n-2} le seront aussi et ainsi de suite jusqu'à A' et B' qui seront nuls identiquement, puisqu'ils ne renferment que des quantités données. Mais alors A'_1 et B'_1 , qui prennent également la forme $A'x_1 + B'$ quand on met pour x_1 chacune des racines de l'équation $x_1^2 + A_1x_1 + B_1 = 0$, s'annuleront pour ces deux valeurs de x_1 ; pareillement, les coefficients A'_1 et B'_1 peuvent être mis sous la forme $A'_2x_2 + B'_2$, en prenant pour x_2 l'une ou l'autre des racines de l'équation $x_2^2 + A_2x_2 + B_2 = 0$, correspondantes à chacune des valeurs de x_1 , et par conséquent ils s'annuleront pour les quatre valeurs de x_2 et pour les deux valeurs de x_1 qui résultent de la combinaison des deux premières équations (A). On démontrera de même que A'_2 et B'_2 seront nuls en mettant pour x_2 les 2^2 valeurs tirées des trois premières équations (A) conjointement avec les valeurs correspondantes de x_1 et x_2 ;

et continuant de cette manière on conclura que $F(x_*)$ s'annulera pour les 2^m valeurs de x_* auxquelles conduit le système de toutes les équations (A) ou pour les 2^m racines de $f(x) = 0$. Ainsi une équation $F(x) = 0$ à coefficients rationnels ne peut admettre une racine de $f(x) = 0$ sans les admettre toutes; donc l'équation $f(x) = 0$ est irréductible.

IV.

Il résulte immédiatement du théorème précédent que tout problème qui conduit à une équation irréductible dont le degré n'est pas une puissance de 2, ne peut être résolu avec la ligne droite et le cercle. Ainsi la *duplication du cube*, qui dépend de l'équation $x^3 - 2a^3 = 0$ toujours irréductible, ne peut être obtenue par la Géométrie élémentaire. Le problème des *deux moyennes proportionnelles*, qui conduit à l'équation $x^3 - a^3b = 0$ est dans le même cas toutes les fois que le rapport de b à a n'est pas un cube. La *trisection de l'angle* dépend de l'équation $x^3 - \frac{3}{4}x + \frac{1}{4}a = 0$; cette équation est irréductible si elle n'a pas de racine qui soit une fonction rationnelle de a et c'est ce qui arrive tant que x reste algébrique; ainsi le problème ne peut être résolu en général avec la règle et le compas. Il nous semble qu'il n'avait pas encore été démontré rigoureusement que ces problèmes, si célèbres chez les anciens, ne fussent pas susceptibles d'une solution par les constructions géométriques auxquelles ils s'attachaient particulièrement.

La division de la circonférence en parties égales peut toujours se ramener à la résolution de l'équation $x^m - 1 = 0$, dans laquelle m est un nombre premier ou une puissance d'un nombre premier. Lorsque m est premier, l'équation $\frac{x^m - 1}{x - 1} = 0$ du degré $m - 1$ est irréductible, comme M. Gauss l'a fait voir dans ses *Disquisitiones arithmeticae*, section VII; ainsi la division ne peut être effectuée par des constructions géométriques que si $m - 1 = 2^k$. Quand m est de la forme a^k , on peut prouver, en modifiant légèrement la démonstration de M. Gauss que l'équation de degré $(a - 1)a^{k-1}$, obtenue en égalant à zéro le quotient de $x^a - 1$ par $x^{a^{k-1}} - 1$, est irréductible; il faudrait donc que $(a - 1)a^{k-1}$ fût de la forme 2^k en même temps que $a - 1$, ce qui est impossible à moins que $a = 2$. Ainsi, la division de la circonférence en N parties ne peut être effectuée avec la règle et le compas que si les facteurs premiers de N différents de 2 sont de la forme $2^k + 1$ et s'ils entrent seulement à la première puissance dans ce nombre. Ce

principe est annoncé par M. Gauss à la fin de son ouvrage, mais il n'en a pas donné la démonstration.

Si l'on pose $x = k + A' \sqrt[m']{a'} + A'' \sqrt[m'']{a''} + \text{etc.}$ $m', m'' \dots$ étant des puissances de 2, et $k, A', A'' \dots a', a'' \dots$ des nombres commensurables, la valeur de x se construira par la ligne droite et le cercle, en sorte que x ne peut être racine d'une équation irréductible d'un degré m qui ne soit pas une puissance de 2. Par exemple, on ne peut avoir,

$x = A \sqrt[m]{a}$, si $(\sqrt[m]{a})^p$ est irrationnel pour $p < m$; on démontrerait facilement que x ne peut prendre cette valeur lors même que m serait une puissance de 2. Nous retrouvons ainsi plusieurs cas particuliers des théorèmes sur les nombres incommensurables que nous avons établis ailleurs (*).

V.

Supposons qu'un problème ait conduit à une équation de degré 2^n , $F(x) = 0$ et qu'on se soit assuré que cette équation est irréductible; il s'agit de reconnaître si la solution peut s'obtenir au moyen d'une série d'équations du second degré.

Reprenons les équations (A) :

$$(A) \quad \begin{cases} x_1^2 + Ax_1 + B = 0, & x_1^2 + A_1x_1 + B_1 = 0 \dots, \\ x_{n-1}^2 + A_{n-1}x_{n-1} + B_{n-1} = 0, & x_n^2 + A_nx_n + B_n = 0. \end{cases}$$

Il faudra construire l'équation $f(x) = 0$, à coefficients rationnels, qui donne toutes les valeurs de x_n et l'identifier avec l'équation donnée $F(x) = 0$. Pour faire ce calcul on remarque que A_{n-1} et B_{n-1} se ramènent à la forme $a_{n-1}x_{n-1} + a'_{n-1}$ et $b_{n-1}x_{n-1} + b'_{n-1}$, en sorte que l'élimination de x_{n-1} entre les deux dernières équations (A) se fait immédiatement, ce qui donne une équation du quatrième degré en x_n ; on y remplacera ensuite a_{n-1} par $a''_{n-1}x_{n-1} + a'''_{n-1}$, a'_{n-1} par $a''_{n-1}x_{n-1} + a'''_{n-1}$, b_{n-1} par $b''_{n-1}x_{n-1} + b'''_{n-1}$, b'_{n-1} par $b''_{n-1}x_{n-1} + b'''_{n-1}$ et A_{n-1} , B_{n-1} par $a_{n-1}x_{n-1} + a'_{n-1}$, $b_{n-1}x_{n-1} + b'_{n-1}$, puis on éliminera x_{n-1} entre l'équation du 4^e degré déjà obtenue et l'équation $x_{n-1}^2 + A_{n-1}x_{n-1} + B_{n-1} = 0$; et ainsi de suite. Les derniers termes des séries $a_{n-1}, a'_{n-1}, a''_{n-1} \dots, b_{n-1}, b'_{n-1} \dots$, etc., doivent être des fonctions rationnelles des coefficients de $F(x) = 0$; si l'on peut leur assigner des valeurs rationnelles qui satisfassent aux équations de condition obtenues en identifiant, on reproduira les équations (A) dont le système équivaut à l'équation

$F(x) = 0$; si les conditions ne peuvent être vérifiées en donnant des valeurs rationnelles aux indéterminées introduites, le problème ne peut être ramené au second degré.

On peut simplifier ce procédé, en supposant que les racines de chacune des équations (A) donnent le dernier terme de la suivante; ainsi, l'on peut prendre B_{n-1} pour l'inconnue de l'avant-dernière équation, puisque $B_{n-1} = b_{n-1}x_{n-1} + b'_{n-1}$, d'où $x_{n-1} = \frac{B_{n-1} - b'_{n-1}}{b_{n-1}}$; de cette manière les éliminations se font plus rapidement et l'on introduit quatre quantités indéterminées dans l'équation du quatrième degré qui résulte de la première élimination, huit dans l'équation du huitième degré, etc., en sorte que les conditions obtenues en identifiant, sont en même nombre que les quantités à déterminer. Mais on écarte aussi à l'avance le cas où l'une des quantités telle que b_{n-1} serait nulle, et il faut étudier ce cas séparément.

Soit, par exemple, l'équation $x^4 + px^3 + qx^2 + r = 0$. Prenons de suite les équations du second degré sous la forme $x_1^2 + Ax_1 + B = 0$, et $x^2 + (ax_1 + a')x + x_1 = 0$; en éliminant x_1 et identifiant, on aura, $2a_1 - \Lambda a = 0$, $a'^2 - Aaa' - \Lambda + a'B = p$, $2aB - a'A = q$, $B = r$, d'où

$$B = r, \quad a = \frac{2q}{4r - \Lambda^2}, \quad a' = \frac{\Lambda q}{4r - \Lambda^2}, \quad \Lambda^3 + p\Lambda^2 - 4r\Lambda + q^2 - 4rp = 0.$$

Comme B , a et a' sont exprimés rationnellement au moyen de Λ , p , q , r , il faut et il suffit que l'équation du troisième degré en Λ ait pour racine une fonction rationnelle des données. La condition est toujours satisfaite quand $q = 0$, quels que soient p et r , car $\Lambda = -p$ satisfait alors à la dernière équation.

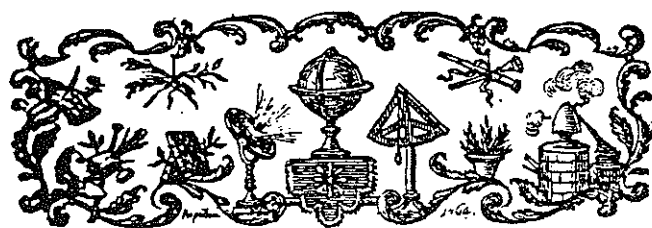
En prenant x_1 pour dernier terme de la deuxième équation du second degré, on a exclu le cas où ce terme serait indépendant de la racine de la première équation; mais en le traitant directement, on ne trouve aucune solution de la question qui ne soit comprise dans les équations ci-dessus.

Ainsi, par un calcul plus ou moins long, on pourra toujours s'assurer si un problème donné est susceptible d'être résolu au moyen d'une série d'équations du second degré, pourvu qu'on sache reconnaître si une équation peut être satisfaite par une fonction rationnelle des données, et si elle est irréductible. Une équation de degré n sera irréductible lorsqu'en cherchant les diviseurs de son premier

nombre de degrés $1, 2, \dots, \frac{n}{2}$, on n'en trouve aucun dont les coefficients soient fonctions rationnelles des quantités données.

La question peut donc toujours être ramenée à rechercher si une équation algébrique $F(x) = 0$ à une seule inconnue peut avoir pour racine une fonction de ce genre. Pour cela, il y a plusieurs cas à considérer. 1° Si les coefficients ne dépendent que de nombres donnés entiers ou fractionnaires, il suffira d'appliquer la méthode des racines commensurables. 2° Il peut arriver que les données représentées par les lettres p, q, r soient susceptibles de prendre une infinité de valeurs, sans que la condition cesse d'être remplie, comme quand elles désignent plusieurs lignes prises arbitrairement; alors, après avoir ramené l'équation $F(x) = 0$ à une forme telle que ses coefficients soient des fractions entières de $p, q, r, \dots$, et que celui du premier terme soit l'unité, on remplacera x par $a_p p^m - a_{p-1} p^{m-1} + \dots + a_0$, et l'on égalera à zéro les coefficients des différentes puissances dans le résultat; les équations obtenues en $a_p, a_{p-1}, \dots$ seront traitées comme l'équation en x , c'est-à-dire qu'on y remplacera ces quantités par des fonctions entières de q , et ainsi de suite jusqu'à ce qu'ayant épuisé toutes les lettres on soit arrivé à des équations numériques qui rentreront dans le premier cas. 3° Lorsque les données sont des nombres irrationnels, ils doivent être racines d'équations algébriques qu'on peut supposer irréductibles; dans ce cas, si l'on remplace x par $a_p p^m + \dots + a_0$ dans $F(x) = 0$, le premier membre de l'équation en p , ainsi obtenue, devra être divisible par celui de l'équation irréductible dont le nombre p est racine; en exprimant que cette division se fait exactement, on arrivera à des équations en $a_p, a_{p-1}, \dots$, que l'on traitera comme l'équation $F(x) = 0$, jusqu'à ce que l'on parvienne à des équations numériques. On doit remarquer que m peut toujours être pris inférieur au degré de l'équation qui donne p .

Ces procédés sont d'une application pénible en général, mais on peut les simplifier et obtenir des résultats plus précis dans certains cas très étendus, que nous étudierons spécialement.



Euclide a encore quelque chose à dire

Gérard HAMON
IREM Rennes

La sortie en 1990 d'une nouvelle traduction par Bernard Vitrac des livres I à IV, (Géométrie plane) du texte de Heiberg (publication de 1883 à 1889)) "EUCLIDE D'ALEXANDRIE - LES ELEMENTS" aux P.U.F., a été un stimulant pour se replonger dans Euclide. Bien que nous utilisions plus ou moins régulièrement le terme de géométrie euclidienne il faut admettre que nous avons bien rarement eu entre les mains une traduction ou un passage important d'une traduction d'Euclide que ce soit comme élève ou comme professeur. Pourtant, les enseignements d'Euclide, antérieurs de 300 ans à ceux du Christ, continuent aujourd'hui encore à régir une bien grande part de la géométrie apprise à l'école, au collège ou au lycée. C'est avec un plaisir certain que nous avons redécouvert sinon découvert l'élaboration pas à pas de cette géométrie, en commençant par les Définitions, 1. "un point est ce dont il n'y a aucune partie".. puis les Demandes, 1 . " Qu'il soit demandé de mener à une ligne droite de tout point à tout point"; les Notions Communes, 1 . " Les choses égales à une même chose sont aussi égales entre elles"... enfin les Propositions, 1 . "Sur une droite limitée donnée construire un triangle équilatéral"... La volonté constante de rigueur, alliée souvent d'une certaine esthétique, caractérise la source d'une partie des mathématiques que nous enseignons aujourd'hui et l'exhibition des modèles qui ont fondé nombre de nos méthodes de raisonnement actuelles ne peuvent laisser indifférent. Aussi il nous a paru utile et rafraîchissant de faire travailler les élèves directement sur Euclide.

Le travail à partir de textes sources doit apparaître comme efficace et surtout ne compliquant pas les choses pour le seul plaisir d'utiliser des textes historiques. Ce souci guidant notre recherche, nous n'avons pas eu à aller très loin pour trouver le thème de notre première activité. Nous nous sommes arrêtés dès la deuxième proposition du livre I, proposition au premier abord qui semblait bien anodine:

"Placer en un point donné une droite égale à une autre droite."

Cela semblait évident et plutôt comme destiné à souscrire à une exigence de rigueur dans le moindre détail. Les compas que nous utilisons régulièrement servent à transporter des longueurs, alors pourquoi diable démontrer cette proposition dès le début comme si c'était très important?

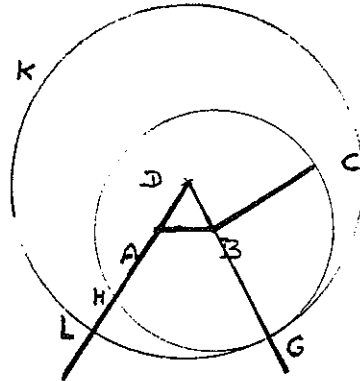
En y regardant de plus près la démonstration d'Euclide ne nous a pas parue banale. En fait c'était là qu'il construisait le compas théorique validant pour toujours les compas matériels utilisés depuis des générations.

Voici la démonstration:

Soit d'une part A le point donné, d'autre part BC, la droite¹ donnée.

Il faut alors placer au point A une droite égale à la droite donnée.

FIG



En effet, que soit jointe la droite AB, du point A jusqu'au point B (Dem. 1), et que, sur elle, soit construit le triangle équilatéral DAB (Prop. 1). Et que les droites AE, BF soient les prolongements en ligne droite de DA, DB (Dem. 2). Et que du centre B et au moyen de l'intervalle BC soit décrit le cercle CGH, et qu'ensuite du centre D et au moyen de l'intervalle DG soit décrit le cercle GKL (Dem. 3).

Or, puisque le point B est le centre du cercle CGH, BC est égale à BG; ensuite, puisque le point D est le centre du cercle GKL, DL est égale à DG (Def. 15), desquelles la partie DA est égale à la partie DB (Def. 20); donc la partie restante AL est égale à la partie restante BG (N.C.1). D'autre part il a été démontré que BC est aussi égale à BG; donc chacune des droites AL, BC est égale à BG; or les choses égales à une même chose sont aussi égales entre elles (N.C.1); et donc AL est égale à BC.

Donc, au point A donné, est placée la droite AL, égale à la droite donnée BC. Ce qu'il fallait faire.

Autour de cette démonstration il nous a paru séduisant de mettre en place une petite axiomatique.

La démonstration de cette proposition II se fait avec

quatre définitions:

- 1- Le point est ce dont il n'y a aucune partie.
- 2- Un segment est une longueur sans largeur.
- 3- Un cercle est une figure plane constituée d'une ligne unique, appelée circonférence, par rapport à laquelle tous les segments menés d'un point unique parmi

¹- Dans Euclide, droite signifie aujourd'hui segment.

ceux qui sont placés à l'intérieur de la figure jusqu'à la circonférence du cercle ont même longueur. Ce point est appelé centre du cercle.

4- Parmi les figures trilatères, le triangle équilatéral est celle qui a les trois côtés égaux ; isocèle, celle qui a seulement deux côtés égaux ; scalène celle qui a ses trois côtés inégaux.

Trois notions communes:

- 1- Les choses égales à une même chose sont égales entre elles;
- 2- Si, à des choses égales, des choses égales sont ajoutées, les tous sont égaux;
- 3- Si, à partir de choses égales, des choses égales sont retranchées, les restes sont égaux.

Trois demandes:

- 1- On peut tracer un segment unique d'un point à un autre point.
- 2- On peut prolonger un segment en ligne droite;
- 3- On ne peut tracer un cercle que si on connaît son centre et un segment rayon de ce cercle.

Enfin avec la Proposition I déjà citée:

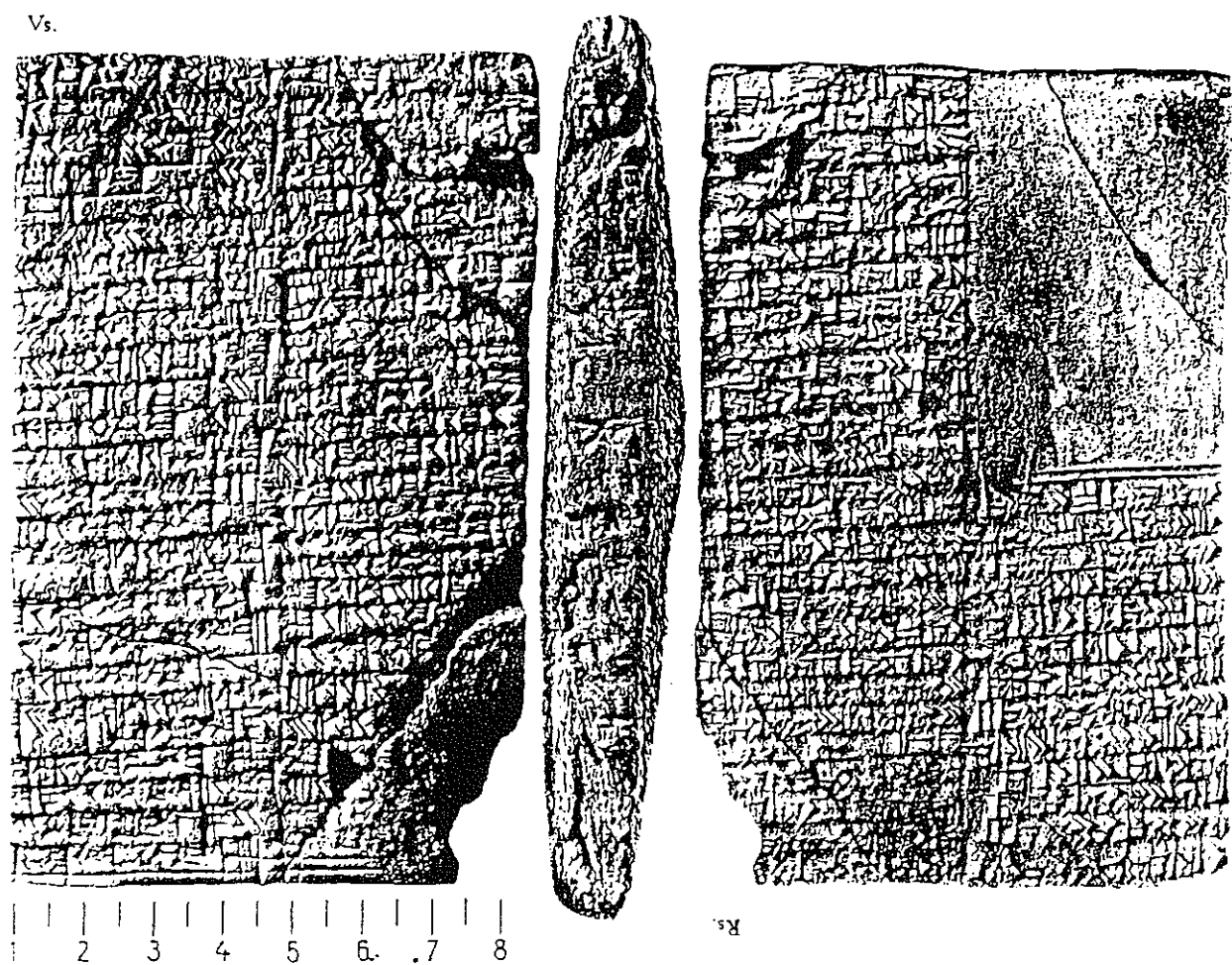
"Sur un segment donné, construire un triangle équilatéral."

L'examen de la démonstration de la proposition I permet de montrer la rigueur d'une démonstration euclidienne et ce que veut dire "démontrer" du point de vue euclidien. Cela fixe l'utilisation usuelle du compas qui n'est pas permise avant la démonstration de la proposition II: **on ne peut tracer un cercle, donc reporter une longueur que si le centre et un point de la circonférence sont connus**. Cette exigence entraîne un travail de recherche intéressant pour construire un losange, une parallèle, un parallélogramme selon des données initiales fixées.

Deux activités ont commencé à être testées, l'une au niveau 4ème 3ème 2de et l'autre au niveau 2de-1èreS². Cela a permis d'aborder les points suivants: le parallélisme et les longueurs, ce qui est attendu lorsqu'en géométrie il est demandé de démontrer; l'examen d'un texte mathématique, d'une construction géométrique.....et aussi une faille dans cette rigueur. En effet Euclide n'a pas défini l'intersection de deux cercles alors que cette notion est nécessaire (cette faille a été remarquée très tardivement par les commentateurs d'Euclide...un élève de seconde l'a remarquée.). En général les élèves s'investissent quel que soit leur niveau car le nombre d'éléments entrant dans les démonstrations et constructions est limité et ceux-ci sont clairement identifiés.

²- Les textes proposés pour ces activités sont intégralement publiés dans la LETTRE D'INFORMATION N°1, IREM de Rennes (Campus de Beaulieu , avenue du général Leclerc 35042 Rennes Cedex.)

O. NEUGEBAUER, *Mathematische Keilschrift-Texte*, 1935,
Springer Verlag, Berlin,



VAT 8389

Fragments d'une histoire des systèmes linéaires

Anne Michel-Pajus

Dans les cours de mathématiques des années 60, les systèmes d'équations linéaires n'avaient rien de bien passionnant. Leur résolution se présentait comme une **application des déterminants**, condensée dans un théorème dit de **Rouché-Fontené**, aussi ennuyeux à énoncer qu'à utiliser. Ensuite est apparue dans les programmes la méthode dite "du **pivot de Gauss**", et je me suis demandé pourquoi l'on attachait le nom de ce prestigieux mathématicien à ce qui ne semblait rien d'autre qu'une résurgence systématisée de la bonne vieille méthode des combinaisons linéaires. En même temps affluaient dans les problèmes de concours des méthodes de résolution par **approximations successives** qui semblaient construites sur mesure pour les ordinateurs. J'ai appris alors que la résolution des très grands systèmes était toujours d'actualité, et loin d'être close.

Sur ces entrefaites j'ai rencontré, dans le cadre de la Commission Inter-Irem d'Histoire et Epistémologie, un petit groupe de gens qui s'intéressaient à l'histoire des algorithmes, et nous avons décidé d'écrire un livre, qui doit paraître l'an prochain [7]. Je tiens à dire tout ce que je dois à ce groupe, et particulièrement à Jean-Luc Chabert et Michel Guillemot.

L'histoire des systèmes linéaires peut se découper en trois phases. Pendant la première, l'objet mathématique "**système linéaire**" n'existe pas en tant que tel. On rencontre seulement des problèmes particuliers, souvent habillés de scénarios aussi improbables que les histoires de robinets du Certificat d'études, assortis d'une solution numérique, et que nous interprétons, nous, comme une résolution de systèmes. Ce qui n'empêche pas les mathématiciens de cette époque de chercher **des méthodes générales** pour résoudre ces problèmes.

La deuxième phase commence après l'invention du symbolisme algébrique, à la fin du 16^{ème} siècle. C'est quand on représente non seulement les inconnues, mais aussi les coefficients, par des lettres, que peut se dégager la structure de "système d'équations". Les efforts portent alors sur la recherche de formules générales donnant les solutions en fonction des coefficients. L'expression, puis la démonstration, de ces formules va donner naissance à la théorie des déterminants, puis des matrices, contribuant ainsi largement à **la naissance de l'algèbre linéaire**.

Au début du 19^{ème} siècle enfin, vont apparaître des méthodes destinées à résoudre numériquement, de façon approchée, des systèmes de grande taille, dont les coefficients sont eux-mêmes donnés avec une certaine approximation. Ces méthodes visent à résoudre des problèmes issus d'autres domaines scientifiques, l'astronomie et la géodésie, considérés d'ailleurs à l'époque comme des branches des mathématiques. La généralisation et la théorisation de ces méthodes va constituer après la deuxième guerre mondiale une part importante de **l'analyse numérique**.

I Les premières méthodes pratiques

I. Un problème babylonien

Les tablettes babyloniennes nous présentent des problèmes numériques plus ou moins "concrets", suivis d'une liste d'instructions qui conduit à la solution. Les problèmes que nous ramenons à des systèmes linéaires sont souvent résolus par élimination successive des inconnues, parfois assez nombreuses. (Il existerait, d'après [22], un problème à 10 équations et 10 inconnues.)

Le problème de la tablette VAT 8389, gravée vers 1800 av. J.C., ne comporte que deux inconnues: les surfaces de deux champs, connaissant la somme des surfaces (que nous noterons x et y), le rendement de chacune (a et b), et la différence du grain récolté. La traduction est faite à partir de celle de HØYRUP[18]. J'ai ajouté les termes entre crochets et les symboles d'unités pour la rendre plus lisible.

Rappelons que les babyloniens de cette période utilisent une méthode de numération de position et sexagésimale, mais sans symbole qui indiquerait l'absence d'une unité, et différentes mesures de surface et de capacités, sans les préciser non plus. Ici les unités de surface sont le bur et le sar (1 bur = 30x60 sar) et celles de capacité le gur et le sila (1 gur = 5x60 sila). De plus, les instructions mêlent les opérations qui servent à la résolution et les calculs intermédiaires qui servent à exécuter les opérations arithmétiques de base. Les divisions, par exemple, se font par multiplication par l'inverse, que le scribe peut trouver, ou non, dans une table, ainsi que les conversions d'unité. Certaines fractions, comme $\frac{1}{2}$, $\frac{1}{3}$ sont considérées au même titre que les unités.

Tropfke[36], Van der Waerden[38], Høyrup[18], ont proposé des explications à cet enchaînement de calculs. Nous ne les détaillerons pas ici, mais ce premier exemple montre déjà le rôle que jouent en mathématiques les moyens dont on dispose pour les exprimer (numération, notations.).

$$x + y = S$$

$$ax - by = d$$

après conversion en sar et sila
(cf. plus loin) on obtient:

$$a = \frac{40}{60} \quad \text{et} \quad b = \frac{30}{60}$$

$$d = 8(60) + 20 \quad S = 30 \times 60$$

Par 1 bur [de surface d'un premier champ], j'ai récolté 4 gur de grain

Par 1 bur [de surface d'un autre champ], j'ai récolté 3 gur de grain.

Le grain [du premier] excède l'autre de 8;20 [en sila]

J'ai additionné mes champs: 30[x60 sar]

Que sont mes champs?

calcul des rendements en sila/bur

$$4 \text{ gur/bur} = 20 \times 60 \text{ sila/bur}$$

$$3 \text{ gur/bur} = 15 \times 60 \text{ sila/bur}$$

$$S/2 = 15 \times 60 \quad \Rightarrow M_1$$

Pose 30[sar] le bur [de la première parcelle]: pose 20[sila] le grain récolté. Pose 30[sar] le bur [de la] second[e parcelle]: pose 15[sila] le grain récolté. Pose 8;20 ce dont le [premier] grain excède l'autre et pose 30[x60sar] l'accumulation des surfaces des champs.

Et 30 l'accumulation des surfaces des champs fractionnée en deux est 15. Pose 15 et 15 deux fois.

calcul des rendements en sila/sar

$$1 \text{ sar} = \frac{1}{30 \times 60} \text{ bur} = \frac{2}{60 \times 60} \text{ bur}$$

$$20 \times 60 \times \frac{2}{60 \times 60} = \frac{40}{60} = a$$

$$M_1 a = S/2 = 10 \quad \Rightarrow M_2$$

Cherche l'inverse de 30, la [surface de la première] parcelle: 2[soixantièmes]. [Multiplie] 2 par 20, le grain qu'on y récolte: il vient 40, la fausse quantité de grain. [Multiplie] par 15 que tu as posé deux fois: il vient 10. Que ta tête le retienne.

$$15 \times 60 \times \frac{2}{60 \times 60} = \frac{30}{60} = b$$

$$M_1 b = 7 + \frac{30}{60} = 7;30 \Rightarrow M_3$$

$$M_2 - M_1 = 2;30$$

$$d - (M_2 - M_3) = 5;50 \Rightarrow M_4$$

$$a + b = 1;10$$

$$\frac{M_4}{a+b} = 5 \Rightarrow M_5$$

$$x = M_1 + M_5 = 20$$

$$y = M_1 - M_5 = 10$$

Vérification

Cherche l'inverse de 30, la seconde parcelle: 2. [Multiplie] par 15, le grain qu'on y récolte: 30, le grain faux. Multiplie par 15 que tu as posé deux fois: 7;30.

10, que ta tête retient, de quoi excède-t'il 7;30? Il l'excède de 2;30. Enlève cet excès de 2;30 à 8;20 dont un grain excède l'autre tu laisses 5;50. Que ta tête le retienne.

Accumule le 40 [grain faux] et le 30: il vient 1;10. [Je n'en connais pas] l'inverse. Que doit-on poser par 1;10 pour obtenir 5;50 que ta tête retient? Pose 5, par 1;10, cela te donne 5;50.

Des 15 que tu as posé deux fois, soustrais de l'un, ajoute à l'autre ce 5 que tu as posé.

Premièrement, [il vient] 20, deuxièmement: 10. 20 est la superficie du premier champ, 10 est celle du second champ.

Si 20 est la surface du premier champ, 10 la surface du second champ, quel est le grain donné par les deux?

Cherche l'inverse de 30, la parcelle: 2. Multiplie 2 par 20 le grain qu'on y récolte: 40. Multiplie par 20 la surface du champ: 13;20.

[Cherche] l'inverse de 30, la seconde parcelle: 2. [Multiplie] 2 par 15, le grain qu'on y récolte: 30. [Multiplie] 30 par 10, la superficie du champ: il vient 5: c'est le grain des 10, la surface du second champ.

Le grain 13;20 [du premier champ], de combien excède-t'il le grain 5 [du second champ]? [il l'excède de 8;20].

2• Un problème de DIOPHANTE (III^{ème} siècle après J.C.)

Ce problème n'est pas caractéristique des travaux de Diophante, qui s'intéresse généralement à des problèmes qui admettent une infinité de solutions, et dont il cherche les solutions rationnelles strictement positives. Mais il est intéressant par sa limpidité, et l'introduction explicite d'une **inconnue supplémentaire**. Le traducteur contemporain, Ver Eecke[10] a forgé le néologisme d'"arithme" pour la distinguer des autres, appelées "nombres", mais Diophante utilise le même mot.

Trouver trois nombres qui, pris deux à deux, forment des nombres proposés.

Il faut toutefois que la moitié de la somme des nombres proposés soit plus grande que chacun de ces nombres.

Proposons donc que le premier nombre, augmenté du second, forme 20 unités; que le second, augmenté du troisième, forme 30 unités, et que le troisième, augmenté du premier, forme 40 unités.

Posons que la somme des trois nombres est 1 arithme. Dès lors, puisque le premier nombre plus le second forment 20 unités, si nous retranchons 20 unités de 1 arithme, nous aurons comme troisième nombre 1 arithme moins 20 unités. Pour la même raison, le premier nombre sera 1 arithme moins 30 unités, et le second nombre sera 1 arithme moins 40 unités. Il faut encore que la somme des trois nombres devienne égale à 1 arithme. Mais, la somme des trois nombres forme 3 arithmes moins 90 unités. Egalons-les à 1 arithme, et l'arithme devient 45 unités.

Revenons à ce que nous avons posé: le premier nombre sera 15 unités, le troisième sera 25 unités, et la preuve est claire.

SCHEMA DES CALCULS:			
$x+y = a$	avec $a=20$	$x+y+z=S$ (arithme)	$z= S-a$
$y+z = b$	$b=30$		$x= S-b$
$z+x = c$	$c=40$		$y= S-c$
			$x+y+z=3S-(a+b+c)$ $S=3S-(a+b+c)$ $S= \frac{(a+b+c)}{2}$
$x = S - b$		$y = S - c$	$z = S - a$

Des questions qui se traduisent par des systèmes linéaires, énoncées de façon rhétorique, c'est-à-dire sans notations algébriques, se rencontrent fréquemment au Moyen-Age en Inde, dans les pays d'Islam et en Europe, mais on en trouve en Chine bien avant.

3• Les méthodes chinoises

Les plus étonnantes sont celles de la dynastie Han (206 av.J.C., 220 ap.J.C). Les voici décrites par Jean-Claude Martzloff[27]:

"Les techniques chinoises reposent sur la répartition des nombres issus de tel ou tel problème arithmétique en colonnes parallèles (hang), chaque colonne traduisant une condition linéaire imposée à un ensemble d'inconnues, appelées les "choses" (wu). Ces colonnes de nombres (représentées concrètement par des groupements de baguettes affectées aux diverses unités décimales des nombres) sont posées au sens propre du mot sur une "surface à calculer" et laissent donc apparaître des configurations de nombres correspondant à ce que nous appellerions "la matrice du système augmentée des seconds membres"...A partir de là, les procédés de résolution utilisent tout un arsenal se composant d'opérations telles que la "multiplication partout"(bian cheng), multiplication d'une colonne de nombres par un même facteur, ou bien la "réduction directe"(zhi chu), c'est-à-dire la soustraction terme à terme des coefficients de deux colonnes en vue de parvenir à l'élimination du coefficient de l'une des inconnues. L'un des procédés chinois les plus remarquables consiste à réduire la "matrice" du système à la forme triangulaire, puis à calculer les inconnues par substitutions successives..." (on trouve des exemples d'applications jusqu'à 6 inconnues).

4• Des méthodes de fausse position

Les chinois utilisent aussi la "*ying bu zu shu*" (règle du trop ou du pas assez). Cette méthode occupe un chapitre entier de la bible arithmétique chinoise, le *Jiuzhang suanshu* (*Procédés calculatoires en neuf chapitres*) . Il s'agit d'un cas particulier de la méthode de double fausse position, que nous détaillerons sur un exemple de résolution d'un problème à 3 inconnues, exposé par Clavius¹. Il est possible de lire comme des méthodes de fausse position (simple) la résolution de certains problèmes égyptiens et babyloniens. Utilisées aussi par les Indiens (Bhaskara, au XII^{ème}), ces méthodes sont introduites par les Arabes en Europe, où on les retrouve à partir du douzième siècle (chez Ben Ezra, par exemple, et Fibonacci)[33].

¹Jésuite du 16^{ème} siècle, qui oeuvra beaucoup pour l'enseignement des mathématiques dans les collèges de Jésuites.

Cette méthode restera en usage dans l'enseignement jusqu'au XX^{ème} siècle. Elle présente l'avantage de ne pas opérer de manière générale sur une inconnue, dans le cadre algébrique, ce qui pose des difficultés conceptuelles, mais de rester dans le cadre arithmétique, en travaillant sur des valeurs numériques particulières.

Dans la méthode de double fausse position, appliquée à ce que nous notons

$$(S_1) \quad ax_1 + bx_2 = c_1$$

$$(S_2) \quad cx_1 + dx_2 = c_2$$

on part d'une valeur de x_1 (posée ou supposée), et on en déduit x_2 à l'aide de (S_1) ; on calcule alors le résidu $e = cx_1 + dx_2 - c_2$ à l'aide de (S_2) .

Un calcul immédiat montre que $e = x_1(c - \frac{ad}{b}) + \frac{dc_1}{b} - c_2$. Nous voyons que:

- e est une fonction affine de x_1
- $e = 0$ correspond à la valeur exacte de x_1 .

Il suffit donc de choisir deux valeurs de x_1 (x'_1 et x''_1), de calculer les valeurs correspondantes de e (e' et e''), et d'utiliser une interpolation linéaire pour obtenir la valeur exacte de x_1

La Règle (nommée **Regula Falsi** au Moyen-Age) donne

$$x_1 = \frac{x'_1 e'' - x''_1 e'}{e'' - e'}$$

Cette méthode et sa justification se généralisent à un système d'ordre n . Soit le système (supposé de Cramer):

$$(S_1) \quad a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n = c_1$$

$$(S_2) \quad a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n = c_2$$

.

$$(S_{n-1}) \quad a_{n-11}x_1 + a_{n-12}x_2 + \dots + a_{n-1n}x_n = c_{n-1}$$

$$(S_n) \quad a_{n1}x_1 + a_{n2}x_2 + \dots + a_{nn}x_n = c_n$$

Pour tout n -uplet $(x_1, x_2, \dots, x_n)$ vérifiant les $(n-1)$ premières équations, on appelle résidu $e = a_{n1}x_1 + \dots + a_{nn}x_n - c_n$.

Si l'on fixe x_1 par exemple, on en déduit $(x_2, \dots, x_n)$ comme solution du système

$$(S_1) \quad a_{12}x_2 + \dots + a_{1n}x_n = c_1 - a_{11}x_1$$

$$(S_2) \quad a_{22}x_2 + \dots + a_{2n}x_n = c_2 - a_{21}x_1$$

.

$$(S_{n-1}) \quad a_{n-12}x_2 + \dots + a_{n-1n}x_n = c_{n-1} - a_{n-11}x_1$$

$x_2, \dots, x_n$ sont donc des fonctions affines de x_1 , et donc e également.

Ceci est évidemment valable pour chaque x_i . Il suffit donc de disposer pour chaque variable de 2 couples de valeurs $(x'_i, e(x'_i))$ et $(x''_i, e(x''_i))$ pour obtenir la valeur exacte de x_i à l'aide de la regula falsi.

Il est intéressant de noter qu'il n'est pas nécessaire de recommencer tout le processus pour chaque variable. En effet, les calculs intermédiaires exécutés pour la première variable fournissent 2 n -uplets solutions des $(n-1)$ premières équations et les 2 valeurs correspondantes du résidu.

CLAVIUS *Aritmetica practica*, Rome 1586 (traduction Maryvonne Spiesser et Anne Michel-Pajus[33])

On cherche trois nombres tels que : le premier augmenté de 73 égale le double des deux autres; le second augmenté de 73 égale le triple des deux autres, le troisième augmenté de 73 le quadruple des deux autres.

Tu poses le premier nombre égal à 1 (...). Puisque 1 avec 73 fait 74, lequel doit selon la question proposée être le double des deux autres, il faut que la somme des deux autres soit 37. Et parce que le second avec 73 doit faire un nombre triple du premier(qui est 1) et du troisième ensemble, il faudra partager (comme dans la question enseignée précédemment) le nombre 37 en deux parties, dont la première avec 73 fait un nombre triple du nombre composé de la seconde partie et du 1. Ainsi, avant de résoudre la question proposée, il est nécessaire d'en résoudre une autre qui apparaît dans cette opération.

Tu poses donc que la première partie de 37 est 2 et par suite la seconde est 35. La première partie avec 73 fait 75 et la seconde avec 1 fait 36, dont le triple n'est pas 75, mais 108. Donc nous avons manqué la vérité de 33 unités.

Tu poses alors la première partie égale à 5 et donc la seconde égale à 32; la première avec 73 fait 78 et la seconde avec 1 fait 33, duquel le triple n'est pas 78, mais 99. Donc on a manqué de nouveau la vérité de 21 unités.

$$\begin{array}{r} 2 \quad 5 \\ 35 \text{ M} \quad \text{M} 32 \\ 33 \quad 21 \\ \hline 12 \text{ diviseur} \end{array}$$

On fait alors, selon le précepte de la règle du faux et on trouvera que la première partie est 10 1/4, et la deuxième partie est 26 3/4. (...)

N.B. **M** signifie "manque", **P** signifie "(dé)passé"

Résumons ce passage et le suite avec des notations algébriques:

Système à résoudre:			
$x+73=2(y+z)$ $y+73=3(z+x)$ $z+73=4(x+y)$			
Si $x=1$	alors	$y+z = 37$ $y+73 = 3(z+1)$	
Si $y = 2$	alors $z=35$ (e' = -33)	33 Manque	
Si $y=5$	alors $z=32$ (e'' = -21)	21 Manque	
$\begin{array}{r} 2 \quad 5 \\ 35 \text{ M} \quad \text{M} 32 \\ 33 \quad 21 \\ \hline 12 \text{ diviseur} \end{array}$		$y=10 \quad 1/4$ $z=26 \quad 3/4$	
Si $x=3$	alors	$y+z = 38$ $y+73 = 3(z+3)$	
Si $y = 2$	alors $z=36$ (e' = -42)	42 Manque	
Si $y=23$	alors $z=15$ (e'' = 42)	42 Pass	
$\begin{array}{r} 2 \quad 23 \\ 36 \text{ M} \quad \text{P} 15 \\ 42 \quad 42 \\ \hline 84 \text{ diviseur} \end{array}$		$y=12 \quad 1/2$ $z=25 \quad 1/2$	
$z+73=99 \quad 3/4$		$4(x+y)=45$	
54 3/4 Passe			
$z+73=98 \quad 1/2$		$4(x+y)=62$	
36 1/2 Pass			
$\begin{array}{r} 1 \quad 3 \\ 10 \quad 1/4 \quad \text{P} \quad 12 \quad 1/2 \\ 26 \quad 3/4 \quad \text{P} \quad 25 \quad 1/2 \\ 54 \quad 3/4 \quad \text{P} \quad 36 \quad 1/2 \\ \hline 18 \quad 1/4 \text{ diviseur} \end{array}$			
$x = 7$ $y = 17$ $z = 23$			

II Les formules générales

C'est à la fin du 16^{ème} qu'à la suite de Viète (*l'Introduction à l'Art Analytique* est publiée en 1591), les mathématiciens désignent par des lettres, non seulement les inconnues (voyelles) et leurs puissances, mais les coefficients indéterminés (consonnes). Ce mouvement se poursuit chez les mathématiciens du 17^{ème}, en particulier avec Descartes qui désire utiliser l'algèbre en Géométrie, et il semble donc que ce soit à la fin du 17^{ème} qu'apparaît l'objet mathématique que nous appelons système linéaire, c'est à dire un système d'équations général, avec des coefficients littéraux. Les notations adoptées ne sont pas très pratiques, et ne mettent pas toujours en évidence la structure des formules. Leibniz, génial précurseur, restera longtemps ignoré.

1• Leibniz

Il faut attendre la publication par Gerhardt, entre 1849 et 1863, de nombreux manuscrits de Leibniz pour découvrir ses travaux sur les équations. Depuis cette date, les chercheurs continuent à étudier ses innombrables manuscrits. J'ai traduit le texte suivant d'après une édition de 1980([23] p 5-6). Il est daté de juin 1678. On y trouve les formules généralement attribuées à Cramer, avec des "déterminants" développés selon une colonne, celle qui est justement différente au numérateur et au dénominateur.

Pour retrouver nos notations habituelles, il suffit de considérer les doubles chiffres comme les doubles indices des coefficients, et les chiffres simples (2,3,4,5) comme les indices des inconnues. Ainsi, $12,2 + 13,3 + 14,4 + 15,5 - A$ égal à 0 se traduit par :

$$a_{12}x_2 + a_{13}x_3 + a_{14}x_4 + a_{15}x_5 - A = 0$$

$$12,23,34,45 \text{ par } a_{12} \cdot a_{23} \cdot a_{34} \cdot a_{45}$$

La dernière formule du texte ci-après correspond à

$$x_5 = \frac{\begin{vmatrix} a_{12} & a_{13} & a_{14} & A \\ a_{22} & a_{23} & a_{24} & B \\ a_{32} & a_{33} & a_{34} & C \\ a_{42} & a_{43} & a_{44} & D \end{vmatrix}}{\begin{vmatrix} a_{12} & a_{13} & a_{14} & a_{15} \\ a_{22} & a_{23} & a_{24} & a_{25} \\ a_{32} & a_{33} & a_{34} & a_{35} \\ a_{42} & a_{43} & a_{44} & a_{45} \end{vmatrix}}$$

lorsque l'on développe les déterminants suivant la dernière colonne.

Specimen d'Analyse nouvelle par laquelle on évite les erreurs, l'esprit est conduit comme par la main, et l'on trouve facilement des développements.

(...)

L'une des remarques qui peuvent produire en calcul de très grands effets est celle-ci: de même que les théorèmes remarquables, et les développements, se dévoilent facilement et comme spontanément, de même peuvent-ils s'écrire le plus souvent sans calculs, pour autant que l'on ait effleuré les prémisses.

Assurément, si quelqu'un, dans le présent exemple qui comporte quatre lettres [inconnues] données, reliées par un nombre égal d'équations simples, cherche la valeur de l'une d'elles et choisit indistinctement comme il est d'usage les lettres a b x y comme il lui plaît, il sera confronté à une horrible confusion et à un labeur immense. Avec notre manière, la chose est accomplie presque sans tracas, comme ce sera évident.

La règle qui résume donc cet art de la caractéristique est que tous les caractères qui se trouvent dans l'objet désigné expriment ce qui, au moyen de nombres, apportera le maximum de facilité de copie et de calcul. Et de même cela sera d'un grand usage en géométrie, pour exprimer l'emplacement. [rature]

12,2 + 13,3 + 14,4 + 15,5 - A égal à 0
 22,2 + 23,3 + 24,4 + 25,5 - B égal à 0
 32,2 + 33,3 + 34,4 + 35,5 - C égal à 0
 42,2 + 43,3 + 44,4 + 45,5 - D égal à 0

Le dénominateur de la fraction exprimant la valeur de [l'inconnue] 5 est:

12,23,34,45 12,23,35,44 12,24,33,45 12,24,35,43 12,25,33,44 13,25,34,43
 - 13,22,34,45 + 13,22,35,44 - 13,24,32,45 + 13,24,35,42 - 13,25,32,44 + 13,25,34,42
 14, 23, 32,45 14,23,35,42 14,22,35,43 14,22,33,45 14,25,32,43 14,25,33,42
 15,22,33,44 15,22,34,43 15,23,32,44 15,23,34,42 15,24,32,43 15,24,33,42

que nous écrivons très facilement au moyen de nombres comme l'on voit, puisque'il faut prendre tour à tour parmi les quantités ou plutôt les nombres , 1.2.3.4. pour les premiers caractères, et pour les suivants,d'abord 2.3.4.5./ 2.3.5.4./ 2.4.3.5./ 2.4.5.3./ 2.5.3.4./ 2.5.4.3./ avec 2 placé en premier,puis 3.2.4.5./3.2.5.4./3.4.2.5./3.4.5.2./3.5.2.4./3.5.4.2./ et ainsi de suite en suivant les transpositions possibles des nombres 2.3.4.5. Où il est clair aussi que l'on obtient la ligne inférieure à partir de la supérieure en gardant les mêmes nombres excepté deux qui s'échangent, comme dans

2.3.4.5

3.2.4.5.

évidemment la ligne inférieure est obtenue à partir de la supérieure en changeant 2 en 3 et inversement. Ceci étant entendu pour les caractères qui suivent ou plutôt qui sont à droite. D'où, en ordonnant, on aura le numérateur et le dénominateur ou plutôt la fraction:

- 12,23,34	D	+ 12,23,44	C	- 12,33,44	B	+ 22,33,44	A
+ 12,24,33		- 12,24,43		+ 12,34,43		- 22,34,43	
+ 13,22,34		- 13,22,44		+ 13,32,44		- 23,32,44	
- 13,24,32		+ 13,24,42		- 13,34,42		+ 23,34,42	
- 14,22,33		+ 14,22,43		- 14,32,43		+ 24,32,43	
+ 14,23,32		- 14,23,42		+ 14,33,42		- 24,33,42	

[var] 5 égale

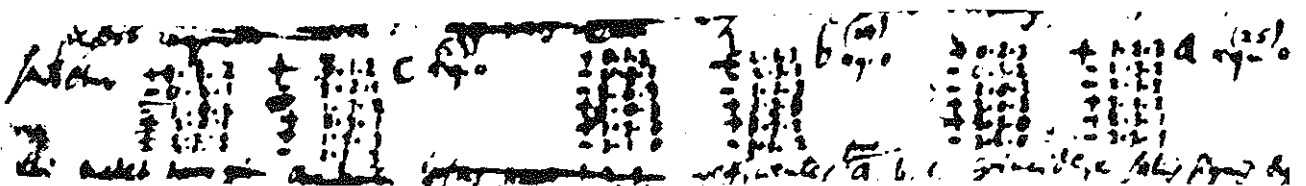
+	45	+	35	-	25	+	15
-		-		+		-	
+		+		-		+	
-		-		+		-	
+		+		-		+	
-		-		+		-	

Où il apparaît également facilement la façon dont il faut écrire. Si je n'avais pas utilisé des nombres au lieu de lettres, je n'aurais pas pu l'écrire ainsi facilement. (...)



Fragment de manuscrit
LH XXXV 14,1 B1. 261
(Stück 8, Zeile 1-116)

Gottfried Wilhelm Leibniz



2. MACLAURIN

Nous avons dit que, faute de publication, les travaux de Leibniz resteront longtemps ignorés, et c'est en 1748 qu'est publié, à titre posthume, l'ouvrage de Colin Maclaurin, "A treatise of algebra" datant sans doute de 1729. Voici les "théorèmes généraux pour trouver les valeurs des inconnues", dans une traduction de 1753. [26]

86

TRAITE' D'ALGEBRE: THEOREMES GENERAUX

Pour trouver les valeurs des inconnues.

204. J'appellerai coefficients du même ordre, les quantités qui accompagnent la même inconnue dans différentes équations, ou celles qui n'en accompagnent aucune.

205. J'appellerai coefficients opposés, ceux qui accompagnent différentes inconnues dans différentes équations.

THEOREME PREMIER.

206. Deux équations & deux inconnues étant données, chacune de ces inconnues est égale à une fraction dont le numérateur est la différence des produits des coefficients opposés des ordres où cette inconnue ne se trouve point, & le dénominateur est la différence des produits des coefficients opposés des deux inconnues.

DEMONSTRATION.

Soit $ax + by = c$

$$dx + ey = f$$

je dis que $y = \frac{af - dc}{ae - db}$

& $x = \frac{ce - bf}{ae - db}$

car par la première équation

$$x = \frac{c - by}{a}$$

par la deuxième ... $x = \frac{f - ey}{d}$

donc $\frac{c - by}{a} = \frac{f - ey}{d}$

$$\left. \begin{array}{l} cd - dby = af - aey \\ aey - dby = af - cd \end{array} \right| y = \frac{af - cd}{ae - db}$$

substituant cette valeur de y dans une des valeurs de x , on trouvera

$$x = \frac{ce - bf}{ae - db}$$

THÉOREME II.

207. Trois équations & trois inconnues étant données, chaque inconnue fera égale à une fraction dont le numérateur contiendra tous les produits qu'on peut faire de trois coefficients opposés, pris dans les ordres où cette inconnue ne se trouve point, & le dénominateur contiendra les différens produits qu'on peut former de trois coefficients opposés, pris dans les ordres qui renferment les trois inconnues.

DEMONSTRATION.

Soit les trois inconnues x, y, z , & les trois équations

$$ax + by + cz = m$$

$$dx + ey + fz = n$$

$$gx + hy + kz = p$$

Ne regardant d'abord comme inconnue que x & y , & n'ayant égard qu'aux deux premières équations, on aura par le Théorème précédent

$$y = \frac{an - afz - dm + dcz}{ae - db}$$

n'ayant égard qu'à la première & à la dernière

$$y = \frac{ap - akz - gm + gcz}{ah - gb}$$

$$\text{donc...} \frac{an - afz - dm + dcz}{ae - db} = \frac{ap - akz - gm + gcz}{ah - gb}$$

& ayant fait évanouir les fractions, détruit les termes semblables, & divisé les deux membres de l'équation par a , on trouvera

$$aekz - afhz + cdhz - bdkz + bfgz - cegz = aep - ahn + dhm - dbp + gbn - gem$$

$$\text{\& par conséquent } z = \frac{aep - ahn + dhm - dbp + gbn - gem}{aek - afh + cdh - bdk + bfg - ceg}$$

pour trouver la valeur de y , au lieu de substituer celle de z , ce qui feroit une opération fort embarrassante, on recommence entièrement l'opération, dans laquelle on traite y comme on a traité z dans la précédente, & on trouve

$$y = \frac{afp - aku + dkm - dep + cgn - fgm}{aek - afh + cdh - bdk + bfg - ceg}$$

S'il manquoit quelque termes dans quelqu'une des trois équations données, on trouveroit des valeurs plus simples des inconnues; je suppose, par exemple, $f=0, k=0$; alors le terme fz s'évanouira dans la seconde équation, & kz dans la troisième, & l'on aura

$$z = \frac{aep - ahn + dhm - dbp + gbn - gem}{cdh - ceg}$$

$$y = \frac{cgn - dep}{cdh - ceg}$$

3. CRAMER

Cramer est amené à rechercher ces formules dans le cadre de ses travaux de géométrie, en particulier la détermination des coniques passant par cinq points donnés. Il les publie dans *l'Introduction à l'analyse des lignes courbes algébriques* en 1750[8].

Ses notations sont un peu meilleures que celles de Maclaurin, mais pas aussi pratiques que celles de Leibniz ! Il démontre les résultats pour les systèmes de 2 et 3 équations puis généralise par induction la formule pour n équations. Il étudie ensuite des systèmes impossibles ou indéterminés.

Soient plusieurs inconnues z, y, x, v, ϕ . & autant d'équations

$$\begin{aligned} A' &= Z'z + \Gamma'y + X'x + V'v + \phi \\ A'' &= Z''z + \Gamma''y + X''x + V''v + \phi \\ A''' &= Z'''z + \Gamma'''y + X'''x + V'''v + \phi \\ A^{(4)} &= Z^{(4)}z + \Gamma^{(4)}y + X^{(4)}x + V^{(4)}v + \phi \end{aligned}$$

où les lettres $A', A'', A''', A^{(4)}, \phi$ ne marquent pas, comme à l'ordinaire, les puissances d' A , mais le premier membre, supposé connu, de la première, seconde, troisième, quatrième &c. équation. De même $Z', Z'', Z''', Z^{(4)}, \phi$ sont les coefficients de z ; $\Gamma', \Gamma'', \Gamma''', \Gamma^{(4)}, \phi$ ceux de y ; $X', X'', X''', X^{(4)}, \phi$ ceux de x ; $V', V'', V''', V^{(4)}, \phi$ ceux de v ; &c. dans la première, seconde, &c. équation.

Cette Notation supposée, s'il n'y a qu'une équation & qu'une inconnue z ; on aura $z = \frac{A'}{Z'}$. S'il y a deux équations & deux inconnues z & y ; on trouvera $z = \frac{A' \Gamma'' - A'' \Gamma'}{Z' \Gamma'' - Z'' \Gamma'}$, & $y = \frac{Z' A'' - Z'' A'}{Z' \Gamma'' - Z'' \Gamma'}$. S'il y a trois équations & trois inconnues z, y , & x ; on trouvera

$$\begin{aligned} z &= \frac{A' \Gamma'' X''' - A'' \Gamma' X''' - A''' \Gamma X'' + A' \Gamma'' X'' + A'' \Gamma' X'' - A''' \Gamma X'}{Z' \Gamma'' X''' - Z'' \Gamma' X''' - Z''' \Gamma X'' + Z' \Gamma'' X'' + Z'' \Gamma' X'' - Z''' \Gamma X'} \\ y &= \frac{Z' A'' X''' - Z'' A' X''' - Z''' A X'' + Z' A'' X'' + Z'' A' X'' - Z''' A X'}{Z' \Gamma'' X''' - Z'' \Gamma' X''' - Z''' \Gamma X'' + Z' \Gamma'' X'' + Z'' \Gamma' X'' - Z''' \Gamma X'} \\ x &= \frac{Z' \Gamma'' X''' - Z'' \Gamma' X''' - Z''' \Gamma X'' + Z' \Gamma'' X'' + Z'' \Gamma' X'' - Z''' \Gamma X'}{Z' \Gamma'' X''' - Z'' \Gamma' X''' - Z''' \Gamma X'' + Z' \Gamma'' X'' + Z'' \Gamma' X'' - Z''' \Gamma X'} \end{aligned}$$

L'examen de ces Formules fournit cette Règle générale. Le nombre des équations & des inconnues étant n , on trouvera la valeur de chaque inconnue en formant n fractions dont le dénominateur commun a autant de termes qu'il y a de divers arrangements de n choses différentes. Chaque terme est composé des lettres $ZYXV$ &c. toujours écrites dans le même ordre, mais auxquelles on distribue, comme exposants, les n premiers chiffres rangés en toutes les manières possibles. Ainsi, lorsqu'on a trois inconnues, le dénominateur a $[1 \times 2 \times 3 =]$ 6 termes, composés des trois lettres ZYX , qui reçoivent successivement les exposants 123, 132, 213, 231, 312, 321. On donne à ces termes les signes + ou —, selon la Règle suivante. Quand un exposant est suivi dans le même terme, médiatement ou immédiatement, d'un exposant plus petit que lui, j'appellerai cela un *dérangement*. Qu'on compte, pour chaque terme, le nombre des dérangements: s'il est pair ou nul, le terme aura le signe +; s'il est impair, le terme aura le signe —. Par ex. dans le terme $Z^1Y^2X^3$ il n'y a aucun dérangement: ce terme aura donc le signe +. Le terme $Z^1Y^3X^2$ a aussi le signe +, parce qu'il a deux dérangements, 3 avant 1 & 2 avant 3. Mais le terme $Z^2Y^1X^3$, qui a trois dérangements, 3 avant 2, 2 avant 1, & 1 avant 3, aura le signe —.

Le dénominateur commun étant ainsi formé, on aura la valeur de z en donnant à ce dénominateur le numérateur qui se forme en changeant, dans tous les termes, Z en A . Et la valeur d' y est la fraction qui a le même dénominateur & pour numérateur la quantité qui résulte quand on change Y en A , dans tous les termes du dénominateur. Et on trouve d'une manière semblable la valeur des autres inconnues.

Généralement parlant, le Problème est déterminé. Mais il peut y avoir des Cas particuliers, où il reste indéterminé; & d'autres où il devient impossible. C'est lorsque le dénominateur commun se trouve égal à zéro

Ces résultats sont immédiatement diffusés et enseignés. "Cette méthode était tellement en faveur que les examens aux écoles des services publics ne roulaient, pour ainsi dire, que sur elle; on était admis ou rejeté suivant qu'on la possédait bien ou mal" aurait dit Gergonne¹.

¹Recteur de l'Académie de Montpellier dans la première moitié du 19^{ème} siècle (Muir I, p14)

4. CAUCHY

Malgré de nombreuses tentatives d'expression ou de démonstration des formules générales, comme celle de Bézout (1764 [2]), Vandermonde (1771[37]), Laplace (1772[21]), ce n'est qu'en 1812, dans le fameux Mémoire de Cauchy [3] "*Sur les fonctions qui ne peuvent obtenir que 2 valeurs,égales et de signes contraires par suite des transpositions opérées entre les variables qu'elles renferment*", que figure la première démonstration générale des formules pour n inconnues. Cauchy y baptise les déterminants (en empruntant le mot à Gauss chez qui il désigne le discriminant d'une forme quadratique) et y inaugure aussi le double indice et la disposition en tableau (sans les barres verticales, qui seront ajoutées par Cayley en 1841).

Le déterminant de ce que nous appelons la matrice A, de terme général $[a_{i,j}]$, est noté $S(+a_{1,1}a_{2,2}a_{3,3}\dots a_{n,n})$. Après avoir démontré les formules de développement suivant une rangée, et défini le cofacteur b_{ij} de a_{ij} , qu'il appelle quantité adjointe, Cauchy démontre les formules (appelées équations(10) dans l'extrait suivant), et que nous résumons par

$$\sum_{k=1}^n a_{i,k} b_{j,k} = \delta_{ij} \det(A)$$

La résolution des systèmes n'est plus qu'une application des déterminants et n'occupe que 4 pages sur 82.

la forme Les équations dont il s'agit seront de

[illegible]

Cela posé, les relations établies par les équations (10) entre les termes du système $(a_{i,n})$ et ceux du système adjoint $(b_{i,n})$ fournissent un moyen facile d'obtenir la valeur d'une inconnue prise à volonté. Ainsi, par exemple, si l'on multiplie la première des équations (20) par $b_{1,1}$, la deuxième par $b_{2,1}$, ... enfin la dernière par $b_{n,1}$, et qu'on les ajoute entre elles en ayant égard aux équations (10), on aura, pour déterminer la valeur de x_1 , l'équation suivante :

$$D_n x_1 = m_1 b_{1,1} + m_2 b_{2,1} + \dots + m_n b_{n,1}.$$

En général, si l'on multiplie la première des équations (20) par $b_{1,\mu}$, la deuxième par $b_{2,\mu}$, ..., enfin la dernière par $b_{n,\mu}$, et qu'ensuite on les ajoute entre elles, on aura, en vertu des équations (10),

$$(21) \quad D_n x_\mu = m_1 b_{1,\mu} + m_2 b_{2,\mu} + \dots + m_n b_{n,\mu}.$$

Par suite, la valeur générale de l'une quelconque des inconnues sera

$$x_\mu = \frac{m_1 b_{1,\mu} + m_2 b_{2,\mu} + \dots + m_n b_{n,\mu}}{D_n} = \frac{m_1 b_{1,\mu} + m_2 b_{2,\mu} + \dots + m_n b_{n,\mu}}{a_{1,\mu} b_{1,\mu} + a_{2,\mu} b_{2,\mu} + \dots + a_{n,\mu} b_{n,\mu}}.$$

Cette valeur se présente donc sous la forme d'une fraction qui a pour dénominateur le déterminant

$$D_n = S(\pm a_{1,1} a_{2,2} \dots a_{n,n})$$

et pour numérateur ce que devient ce déterminant quand on y remplace les coefficients de l'inconnue pris dans les équations (20) par les seconds membres de ces mêmes équations.

Les compléments sur la résolution exacte des systèmes vont préciser les cas d'indétermination et d'impossibilité. Le théorème complet est connu en France sous le nom de théorème de ROUCHE-FONTENE. Ces deux professeurs, dont le premier, polytechnicien, était professeur au Lycée Charlemagne, et le second, agrégé de mathématiques, au Collège Rollin, ont publié plusieurs notes entre 1875 et 1880 sur le sujet. On doit à Eugène Rouché les dénominations de déterminant principal et déterminants caractéristiques. Mais Muir[29] montre que C.L.DODGSON, plus connu sous le nom de Lewis Carroll, avait déjà publié ces résultats en 1867[11].

Notons que ces recherches vont conduire Frobenius, en 1879, à dégager le concept de rang. La résolution des systèmes linéaires joue ainsi un rôle fondamental dans l'émergence des concepts de l'algèbre linéaire.[12]

III Les méthodes de résolution approchées

1. LA METHODE DES MOINDRES CARRÉS

Nous avons vu que les mathématiciens disposent donc à partir de la deuxième moitié du XVIII^{ème} de méthodes tout à fait efficaces pour résoudre les systèmes à 2,3, à la rigueur 4, inconnues, mais les recherches dans d'autres branches des mathématiques, telles l'astronomie ou la géodésie, vont conduire à des systèmes **d'une nature toute différente**. Non seulement le nombre d'inconnues est beaucoup plus grand, mais les coefficients des équations sont incertains, puisque résultant de mesures physiques. L'on dispose en revanche d'un nombre pratiquement illimité d'équations...Tous les systèmes obtenus ainsi sont incompatibles ! comment obtenir malgré tout une solution acceptable?

La réponse est la méthode des moindres carrés.

Elle est publiée pour la première fois par Legendre[25] en appendice de ses *Nouvelles méthodes pour la Détermination des orbites des Comètes*, en 1805, mais Gauss l'a utilisée déjà au moins dès 1801, pour déterminer à la stupéfaction générale l'orbite de Cérès, malgré des observations très courtes. La revendication de paternité de la méthode provoquera d'ailleurs des controverses assez violentes ! [6]

Dans le cadre conceptuel actuel, le principe est très simple. Notons $MX+N=0$, ou $MX=(-N)$ le système à résoudre. Il est incompatible en général, car $(-N)$ n'appartient pas à l'image de M . On va choisir comme solution un vecteur X_0 tel que la distance de MX_0 à $(-N)$ soit minimale; dans une structure euclidienne, c'est donc que MX_0 est la projection orthogonale de $(-N)$ sur l'image de M . On a donc pour tout vecteur X , $\langle MX_0 + N, MX \rangle = \langle {}^tMMX_0 + {}^tMN, X \rangle = 0$. Ce qui équivaut à: X_0 est solution de ${}^tMMX_0 = -{}^tMN$. On se ramène ainsi à un système qui a autant d'équations que d'inconnues, de matrice $A = {}^tMM$, symétrique, positive et généralement définie positive.

Ce n'est évidemment pas ainsi que les mathématiciens du XIX^{ème} justifient la méthode. Prenons les notations de Gauss:

La transposition matricielle du système à résoudre est $N + MX = 0$, et, à tout vecteur X , on associe le vecteur $W = N + MX$, avec

$$M = \begin{pmatrix} a & b & c & \dots & \dots \\ a' & b' & c' & \dots & \dots \\ a'' & b'' & c'' & \dots & \dots \\ \dots & \dots & \dots & \dots & \dots \end{pmatrix} \quad X = \begin{pmatrix} p \\ p' \\ r \\ \dots \end{pmatrix} \quad N = \begin{pmatrix} n \\ n' \\ n'' \\ \dots \end{pmatrix} \quad W = \begin{pmatrix} w \\ w' \\ w'' \\ \dots \end{pmatrix}$$

w, w', w'' sont appelés les erreurs ou les écarts. La quantité à minimiser est donc $W = w^2 + w'^2 + w''^2 \dots$, qui est fonction des inconnues $p, q, r \dots$. Legendre (en 1805) obtient le nouveau système en annulant les dérivées partielles de W par rapport à chacune des inconnues. Gauss donne plusieurs justifications, entre 1809 et 1821 en s'appuyant sur un raisonnement de probabilité. On peut aussi utiliser une autre norme pour la distance, c'est ce que fait Euler, en 1749, qui minimise le maximum de la valeur absolue des écarts, ou Laplace, en 1799, qui minimise la somme des valeurs absolues des écarts. Mais la méthode des moindres carrés aboutit à des calculs plus agréables, car elle conserve la linéarité. Les nouvelles équations obtenues sont désignées par le terme d'équations normales.

2. Le "Pivot" de Gauss

C'est dans le *De ellipticis Palladis* (1810) [14] que j'ai trouvé le texte qui se rapproche le plus de ce que nous désignons ainsi¹. Après le succès remporté avec Cérès, Gauss se penche sur l'orbite de Pallas, découverte par Olbers. Avec 6 inconnues et 11 équations, il obtient 6 équations normales par la méthode des moindres carrés. Pour éviter les "pénibles" calculs de l'élimination classique, Gauss propose une autre méthode, qui consiste en fait à exécuter ce que nous appelons la réduction en carrés de Gauss de la forme quadratique W, avec un carré supplémentaire qui correspond aux termes constants. Il obtient le minimum en annulant tous les carrés, ce qui donne les n équations d'un système triangulaire qui coïncide avec celui obtenu par la méthode dite "du pivot".

(Traduction de Joseph Bertrand
1855)

Traduction matricielle

$$N + AX = W$$

$$\Omega = \|AX + N\|^2$$

$$\frac{\partial \Omega}{\partial p} = \frac{\partial \Omega}{\partial q} = \dots = 0$$

$${}^tAN + {}^tAAX = 0$$

Dans l'impossibilité où nous sommes de satisfaire exactement aux onze équations proposées, c'est-à-dire d'annuler tous les seconds membres, nous chercherons à rendre la somme de leurs carrés aussi petite que possible.

On aperçoit facilement que si l'on considère les fonctions linéaires

$$\begin{aligned} n + np + bq + cr + ds + \dots &= w, \\ n' + n'p + b'q + c'r + d's + \dots &= w', \\ n'' + n''p + b''q + c''r + d''s + \dots &= w'', \\ &\dots\dots\dots \end{aligned}$$

les équations qu'il faut résoudre pour rendre

$$\Omega = w^2 + w'^2 + w''^2 + \dots$$

un minimum, sont

$$\begin{aligned} n w + n' w' + n'' w'' + \dots &= 0, \\ b w + b' w' + b'' w'' + \dots &= 0, \\ c w + c' w' + c'' w'' + \dots &= 0, \\ &\dots\dots\dots \end{aligned}$$

ou, en posant, pour abréger,

$$\begin{aligned} n n + n' n' + n'' n'' + \dots &= (nn), \\ n' + n'' + n''' + \dots &= (na), \\ ab + a' b' + a'' b'' + \dots &= (ab), \\ &\dots\dots\dots \\ b^2 + b'^2 + b''^2 + \dots &= (bb), \\ bc + b' c' + b'' c'' + \dots &= (bc), \\ &\dots\dots\dots \end{aligned}$$

p, q, r, s, etc., devront se déterminer par les équations suivantes :

$$\begin{aligned} (an) + (na)p + (ab)q + (ac)r + \dots &= 0, \\ (bn) + (nb)p + (bb)q + (bc)r + \dots &= 0, \\ (cn) + (nc)p + (bc)q + (cc)r + \dots &= 0, \\ &\dots\dots\dots \end{aligned}$$

¹Gauss formule des remarques à ce sujet à la fin de la Théorie du Mouvement des corps célestes de 1809 et reprend sa méthode dans la deuxième partie de la Théorie de la Combinaison des Observations parue en 1823

L'élimination, très-pénible lorsque le nombre des inconnues est considérable, peut se simplifier notablement de la manière suivante. Outre les coefficients (an) , (aa) , etc., [dont le nombre est $\frac{1}{2}(i^2 + 3i)$, si le nombre des inconnues est i], supposons que l'on ait calculé la somme

$$n^2 + n'^2 + n''^2 + \dots = (nn);$$

on voit facilement que l'on a

$$\begin{aligned}\Omega &= \langle AX+N, AX+N \rangle \\ &= \|N\|^2 + 2 \langle AN, X \rangle + X^t A A X\end{aligned}$$

$$\begin{aligned}\Omega &= (nn) + 2(an)p + 2(bn)q + 2(cn)r + \dots \\ &\quad + (aa)p^2 + 2(ab)pq + 2(ac)pr + \dots \\ &\quad + (bb)q^2 + 2(bc)qr + 2(bd)qs + \dots \\ &\quad + (cc)r^2 + 2(cd)rs + \dots;\end{aligned}$$

et, en désignant

$$(an) + (aa)p + (ab)q + \dots$$

par Λ , tous les termes de $\frac{\Lambda^2}{(aa)}$ qui contiennent le facteur p , se trouvent dans l'expression Ω , et, par suite,

$$\Omega - \frac{\Lambda^2}{(aa)}$$

est une fonction indépendante de p . C'est pourquoi, en posant

* erreur du traducteur : lire:

$$\begin{aligned}(bn) &= \frac{(an)(ab)}{(aa)} \\ (cn) &= \frac{(an)(ac)}{(aa)}\end{aligned}$$

$$\begin{aligned}(nn) - \frac{(an)^2}{(aa)} &= (nn, 1), \\ (bn) - \frac{(an)(bn)}{(aa)}^* &= (bn, 1), \\ (cn) - \frac{(an)(cn)}{(aa)}^* &= (cn, 1), \\ &\dots\dots\dots \\ (bb) - \frac{(ab)^2}{(aa)} &= (bb, 1), \\ (bc) - \frac{(ab)(ac)}{(aa)} &= (bc, 1), \\ (bd) - \frac{(ab)(ad)}{(aa)} &= (bd, 1), \\ &\dots\dots\dots\end{aligned}$$

on aura

$$\begin{aligned}\Omega - \frac{\Lambda^2}{(aa)} &= (nn, 1) + 2(bn, 1)q + 2(cn, 1)r + 2(dn, 1)s \dots \\ &\quad + (bb, 1)q^2 + 2(bc, 1)qr + 2(bd, 1)qs \\ &\quad + (cc, 1)r^2 + 2(cd, 1)rs \\ &\quad + \dots\dots\dots\end{aligned}$$

nous désignerons cette fonction par Ω' .

De même, en posant

$$(bn, 1) + (bb, 1)q + (bc, 1)r + \dots = B,$$

la différence

$$\Omega' = \frac{B^2}{(bb, 1)}$$

sera indépendante de q ; nous la représenterons par Ω'' .

En posant, de même,

$$(nn, 1) - \frac{(bn, 1)^2}{(bb, 1)} = (nn, 2),$$

$$(cn, 1) - \frac{(bn, 1)(bc, 1)}{(bb, 1)} = (cn, 2),$$

$$(cr, 1) - \frac{(bc, 1)^2}{(bb, 1)} = (cc, 2),$$

.....

et

$$(cn, 2) + (cc, 2)r + (cd, 2)s + \dots = C,$$

la différence

$$\Omega'' = \frac{C^2}{(cc, 2)}$$

sera une fonction indépendante de r .

En continuant ainsi, nous formerons une suite d'expressions $\Omega, \Omega', \Omega'',$ etc., dont la dernière sera indépendante des diverses inconnues, et représentée par (nn, μ) , si μ désigne le nombre de ces inconnues; nous aurons alors

$$\Omega = \frac{A^2}{(aa)} + \frac{B^2}{(bb, 1)} + \frac{C^2}{(cc, 2)} + \frac{D^2}{(dd, 3)} + \dots + (nn, \mu).$$

On prouvera facilement que Ω étant une somme de carrés

$$\omega^2 + \omega'^2 + \omega''^2 + \dots,$$

et ne pouvant devenir négative, les diviseurs $(aa), (bb, 1), (cc, 2),$ etc., sont tous positifs. (Nous supprimons, pour abréger, le détail de la démonstration.) D'après cela, la valeur minimum de Ω correspond évidemment aux valeurs des inconnues, pour lesquelles

$$A = 0, \quad B = 0, \quad C = 0, \dots,$$

et, en commençant à résoudre le système par la dernière équation, qui ne contient qu'une inconnue, on trouvera les valeurs de $p, q, r, s,$ etc., sans avoir aucune élimination à effectuer. La méthode donne, en même temps, la valeur minimum de Ω , qui est (nn, μ) .

Les problèmes posés par les arrondis des calculs successifs seront abordés par Moulton en 1913[28], qui donne le premier exemple connu de matrice "mal conditionnée". La première définition du "conditionnement d'une matrice" sera introduite par Todd en 1950[35].

3. La méthode indirecte de GAUSS

Vers 1817, Gauss achève ses travaux d'Astronomie et se consacre à la Géodésie. Son programme de triangulation du Hanovre va l'occuper jusqu'en 1825. Il utilise une chaîne de 26 triangles. Les systèmes deviennent monstrueux, et de plus, il est fatigué par les travaux d'arpentage sur le terrain. Il invente alors une méthode qui supporte les erreurs de calcul et qu'il peut pratiquer, comme il l'écrit à son ami Gerling, le 26 décembre 1823, "à moitié endormi ou en pensant à autre chose".

Elle peut être considérée comme l'ancêtre des méthodes que Southwell a décrites entre 1935 et 1949, et baptisées : **méthodes de relaxation**. Il ne semble pas que ce soit pour leurs qualités relaxantes, mais par analogie avec le problème mécanique qu'il étudie, un cadre chargé de poids:

Les méthodes indirectes de résolution des systèmes linéaires

Nous reprenons la notation matricielle du système : $MX + B = 0$, avec

$$M = \begin{pmatrix} m_1^1 & m_1^2 & \dots & m_1^n \\ m_2^1 & m_2^2 & \dots & m_2^n \\ \dots & \dots & \dots & \dots \\ m_n^1 & m_n^2 & \dots & m_n^n \end{pmatrix} \quad X = \begin{pmatrix} x_1 \\ x_2 \\ \dots \\ x_n \end{pmatrix} \quad B = \begin{pmatrix} b_1 \\ b_2 \\ \dots \\ b_n \end{pmatrix}$$

Pour un vecteur Y donné, on appelle résidu relatif au vecteur Y le vecteur $W = MY + B$.

Les méthodes consistent en la construction par récurrence des suites $X^{(p)}$ et $W^{(p)} = M X^{(p)} + B$, $X^{(0)}$ étant choisi arbitrairement. Il est clair que, si M est inversible, et si $W^{(p)}$ tend vers le vecteur nul, $X^{(p)}$ a une limite, qui est la solution du système.

Les méthodes de relaxation:

Pour passer de $X^{(p)}$ à $X^{(p+1)}$, il faut choisir une équation, soit la i-ème:

$$\sum_{j=1}^n m_i^j x_j + b_i = 0$$

et une composante: la k-ème, qui sera la seule à être modifiée,

de sorte que la i-ème composante du résidu $W^{(p+1)}$ soit nulle.

$X^{(p+1)}$ est donc définie par:

$$\forall j \neq k \dots x_j^{(p+1)} = x_j^{(p)}$$

$$x_k^{(p+1)} = \frac{1}{m_i^k} \left[b_i - \sum_{j \neq k} m_i^j x_j^{(p)} \right] = x_k^{(p)} - \frac{w_i^{(p)}}{m_i^k}$$

Il faut évidemment que m_i^k soit non nul. Dans les méthodes que nous allons voir, l'auteur a toujours choisi $k=i$, ce qui exige des coefficients diagonaux non nuls.

Ces méthodes se différencient donc par le critère de choix de i.

Pour GAUSS[16], i est tel que $\frac{|w_i^{(p)}|}{m_i^i}$ soit maximal

Pour SEIDEL[31], i est tel que $\frac{|w_i^{(p)}|^2}{m_i^i}$ soit maximal

Pour SOUTHWELL[32], i est tel que $|w_i^{(p)}|$ soit maximal.

Les méthodes itératives :

Ce sont des méthodes de relaxation cycliques : à chaque pas, i augmente d'une unité (modulo n). On peut alors exprimer (comme c'est clairement expliqué dans le texte de Nekrasov pour la méthode de Seidel) $X^{(p+1)}$ en fonction de $X^{(p)}$ sous forme matricielle:

si $M = T + A$, la suite est définie par $TX^{(p+1)} = -AX^{(p)} - B$, ou $X^{(p+1)} = -T^{-1}AX^{(p)} - B$.

Les méthodes diffèrent par le choix de T. Pour Jacobi, T est la diagonale de M, pour la version cyclique de Seidel, T est le triangle inférieur de M.

On voit facilement sur cette expression que si la suite $X^{(p)}$ converge, sa limite est solution du système.

Dans sa lettre à Gerling - que nous présentons intégralement en annexe, comme les autres textes de cette troisième partie inédits ou difficiles à trouver - Gauss donne seulement un exemple d'application de sa méthode à la résolution d'un système numérique de 4 équations normales à 4 inconnues, obtenues à partir de la méthode des moindres carrés.

À la suite des techniques de compensation de la géodésie, la matrice du système a tous ses éléments diagonaux positifs, et opposés à la somme des autres éléments de la même rangée, et le vecteur cherché est supposé très petit.

Gauss part du vecteur nul, et à chaque étape modifie la composante dont l'indice donne un maximum pour la valeur absolue de la modification. Il s'arrête quand cette modification est inférieure à 0,5. Notons que le système est à coefficients entiers, et qu'il arrondit toujours à l'entier le plus proche, qui mesure le millième de seconde. Les vecteurs résidus, disposés en colonne, ont une norme strictement décroissante. Le processus s'arrête forcément en un nombre fini d'étapes. On peut remarquer aussi que Gauss utilise le fait que la somme des composantes du résidu est toujours nulle pour vérifier l'exactitude des calculs. Ceci fait penser aux invariants de boucle de nos algorithmes.

Bien que Gerling publie en 1843 un ouvrage [17] où il détaille cette méthode avec de nombreux exemples, il semble qu'elle ne se diffuse pas dans le milieu mathématique. La plus ancienne référence que j'ai trouvée se trouve dans WHITTAKER et ROBINSON (1924)[39]

4• La méthode de JACOBI (1845)

C'est dans un article de Jacobi que nous trouvons la première description systématique d'une **méthode itérative**. Dans le cadre de ses recherches sur les perturbations séculaires des planètes, il utilise la méthode des moindres carrés, et pour simplifier les calculs, remarque que dans ce cas "les coefficients diagonaux sont prépondérants parce qu'ils sont sommes de carrés, tandis que les autres résultent de l'addition de nombres positifs et négatifs qui s'annulent en partie."

Pour le cas où la diagonale n'est pas assez dominante, Jacobi propose une méthode de transformation qui revient en fait à diagonaliser la matrice, puisqu'elle donne les coefficients hors diagonale aussi petits que l'on veut. "Mais à partir d'un certain point, laissé à l'appréciation du calculateur, il sera avantageux d'introduire la méthode d'approximation. Si on le fait trop tôt, cette méthode elle-même mettra en évidence les coefficients qui compromettent son succès et qui doivent par conséquent être éliminés par de nouvelles transformations. "Il donne également des invariants de contrôle. Il sent bien que l'aspect dominant de la diagonale joue un rôle dans la convergence mais il ne donne pas de justifications de celle-ci.

"La difficulté de résoudre exactement un grand nombre d'équations linéaires (comme c'est souvent le cas) quand on applique la méthode des moindres carrés, a conduit à penser à des méthodes d'approximation. Il en est une qui se présente d'elle même quand, dans les différentes équations, ce n'est jamais la même variable qui est affectée d'un coefficient nettement supérieur aux autres. Soient en effet les équations:

$$(00)x + (01)x_1 + (02)x_2 + \dots = (0m)$$

$$(10)x + (11)x_1 + (12)x_2 + \dots = (1m)$$

$$(20)x + (21)x_1 + (22)x_2 + \dots = (2m)$$

etc. etc. etc.

dans lesquelles tous les coefficients (ik) sont très petits par rapport à ceux (jj) de la diagonale; on obtiendra une valeur approchée des inconnues x, x_1, x_2 , etc. par:

$$(00)x = (0m), (11)x_1 = (1m), (22)x_2 = (2m) \text{ etc...}$$

Si on désigne ces valeurs par a, a_1, a_2 etc..., on obtient leurs premières corrections $\Delta, \Delta_1, \Delta_2$ etc. par:

$$(00) \Delta = -((01)a_1 + (02)a_2 + \dots)$$

$$(11) \Delta_1 = -((10)a + (12)a_2 + \dots) \text{ etc....}$$

et, si on pose en général:

$$x = a + \Delta + \Delta^2 + \Delta^3 \dots$$

$$x_1 = a_1 + \Delta_1 + \Delta_1^2 + \Delta_1^3 \dots$$

$$x_2 = a_2 + \Delta_2 + \Delta_2^2 + \Delta_2^3 \dots \text{ etc}$$

où les indices supérieurs désignent les corrections successives, de plus en plus petites, on déduira des Δ^{i+1} des Δ^i par les égalités:

$$(00)\Delta^{i+1} = -((01) \Delta_1^i + (02) \Delta_2^i + \dots),$$

$$(11)\Delta_1^{i+1} = -((10) \Delta^i + (12) \Delta_2^i + \dots),$$

$$(22)\Delta_2^{i+1} = -((20) \Delta^i + (21) \Delta_1^i + (23) \Delta_3^i + \dots),$$

etc. etc.

Dans les équations auxquelles conduit la méthode des moindres carrés, les coefficients diagonaux sont effectivement prépondérants parce qu'ils sont des sommes de carrés, tandis que les autres résultent de l'addition de nombres positifs et négatifs qui s'annulent en partie....."

[traduction : Colette BLOCH]

Appelons D la diagonale de la matrice M du système (avec $M = D+C$).

Puisqu'elle est prépondérante, on obtient une première solution approchée $X^{(0)}$ en résolvant $D X^{(0)} = -N$, puis on ajoute $D^{(1)}$, vérifiant $DD^{(1)} = CX^{(0)}$ On obtient ainsi $X^{(i+1)} = X^{(i)} + D^{(i+1)}$, puis par récurrence, $DX^{(i+1)} = CX^{(i)} - N$.

5. La méthode de SEIDEL (1874)

Seidel est un élève de Jacobi et il effectue beaucoup de calculs pour lui. En particulier il doit résoudre des systèmes à 72 inconnues pour calculer les valeurs les plus probables des log des luminosités des étoiles. Pour accélérer le procédé, il revient à la méthode de Gauss (celle de Pallas, car il semble ignorer l'autre, pourtant publiée par Gerling en 1845) et reprend la quantité W à minimiser.

$$\begin{cases} a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n + b_1 = w_1 \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n + b_2 = w_2 \\ \dots \\ a_{n1}x_1 + a_{n2}x_2 + \dots + a_{nn}x_n + b_n = w_n \end{cases}$$

Première étape de la réduction de Gauss:

$$\Omega = \frac{w_1^2}{a_{11}} + f(x_2, x_3, \dots, x_n)$$

Remplaçons x_1 par $x_1 + Dx_1$, avec $Dx_1 = -\frac{W_1}{a_{11}}$:

W_1 devient $W'_1 = 0$

$$W_i \text{ devient } W'_i = W_i + a_{1i} Dx_1 \text{ pour } i \neq 1$$

Ω diminue de $\frac{W_1^2}{a_{11}}$ et devient Ω'

Maintenant, choisissons k tel que $\frac{W_k^2}{a_{kk}}$ soit maximal :

$$\Omega' = \frac{w_k^2}{a_{kk}} + f'(x_1, \dots, x_{k-1}, x_{k+1}, \dots, x_n)$$

En remplaçant x_k par $x_k + Dx_k$, avec $Dx_k = \frac{-W'_k}{a_{kk}}$.

W'_i devient $W''_i = W'_i + a_{ki} D x_k$ pour $i \neq k$, et Ω diminue de $\frac{W_k'^2}{a_{kk}}$

La variante cyclique de Gauss-Seidel a été étudiée en 1885 par Nekrasov[30], qui la baptise méthode de Seidel, démontre sa convergence et relie la vitesse de convergence aux valeurs propres de la matrice. La traduction du texte de Nekrasov est en annexe (annexe 2), comme celle de Seidel (annexe 3).

6• La méthode de CHOLESKY (1923)

Ce polytechnicien, né en 1875, était chargé de recherches géodésiques. Il a participé en particulier à la triangulation de la Crète. Tué au combat en 1918, il avait imaginé une méthode très performante de résolution des équations normales, que son ami le Commandant Benoit a publiée en 1923.

Soit à résoudre: (1) $AX+K=0$
et soit (2) $X = {}^tAL$ (l'équation corrélative)
(1) et (2) $\Rightarrow$ (3) $A {}^tAL + K = 0$

Cholesky détermine par identification les relations qui permettent de calculer les coefficients de la matrice T triangulaire inférieure telle que $T {}^tT = A {}^tA$.

On peut alors considérer que (3) provient de :

$$(1') TY+K=0 \text{ et}$$

$$(2') Y = {}^tTL$$

Il reste à résoudre successivement:

$$(1') TY+K=0 \text{ qui donne } Y \text{ par un système triangulaire}$$

$$(2') Y = {}^tTL \text{ qui donne } L$$

$$(2) X = {}^tAL$$

Si l'on part de $BL+K=0$, où B est une matrice symétrique définie positive, il faut déterminer T , et l'on obtient Y par (1'), puis L par (2').

Cholesky semble avoir ignoré les méthodes itératives. On n'en trouve d'ailleurs aucune trace dans les cours de géodésie de l'Ecole Polytechnique (donnés par Callandreau, puis Poincaré à l'époque), alors qu'il y a trente lignes de conseils sur l'art de disposer ses calculs.

On peut y voir une marque du mépris dans lequel ont longtemps été tenues les Mathématiques Appliquées en France. La méthode des moindres carrés y a d'ailleurs été également oubliée jusqu'après la Guerre de 1870 [21]. A ma connaissance, la méthode de Seidel a été publiée en français pour la première fois dans l'Encyclopédie de Molk, traduite de l'Encyclopédie allemande, entre 1911 et 1916.

Les méthodes itératives, comme les Mathématiques Appliquées en général, connaîtront un grand développement aux Etats-Unis pendant la deuxième guerre mondiale. Un de leurs promoteurs est G.E. Forsythe, et c'est un de ses articles [13] qui m'a donné les pistes pour cette dernière partie. Il s'intitule "Solving linear algebraic equations can be interesting".

J'espère vous en avoir convaincu(e) !

Bibliographie

- [1] BENOIT Ct.,1923 Note sur une méthode de résolution des équations normales provenant de l'application de la méthode des moindres carrés à un système d'équations linéaires en nombre inférieur à celui des inconnues(Procédé du Ct Cholesky), *Bulletin Géodésique*, 2, p 67-77.
- [2] BEZOUT E.,1764.*Histoire de l'Académie Royale des Sciences de Paris*.
- [3] CAUCHY L.A.,1815.Mémoire sur les fonctions qui ne peuvent obtenir que deux valeurs égales et de signes contraires par suite des transpositions opérées entre les variables qu'elles renferment. *Journal de l'Ecole Polytechnique*, 10,p 29-112. Oeuvres 2^e série ,p 191-169.
- [4] CAUCHY L.A.,1847.Méthode générale pour la résolution des systèmes d'équations simultanées. *Comptes Rendus de l'Académie des Sciences de Paris*,25,p 536-538.
- [5] CAYLEY A.,1858. A memoir on the theory of Matrices. *Philosophical Transactions*,148,p 17-37.
- [6] CHABERT J.L.,1989. Gauss et la méthode des moindres carrés, *Revue d'Histoire des Sciences*,T42, p 5-26.
- [7] CHABERT et alii,1993. *Histoire d'algorithmes: du caillou à la puce*, Belin,Paris.
- [8] CRAMER G., 1750. *Introduction à l'analyse des lignes courbes algébriques*. Genève.
- [9] DAHAN-DALMEDICO A. et PEIFFER J.,1986. *Une histoire des mathématiques*. Points Sciences, Le Seuil, Paris.
- [10] DIOPHANTE d'Alexandrie, *Les six livres d'arithmétique...* Trad Paul Ver Eecke, Bruges 1926. Réimpression Blanchard Paris, 1959.
- [11] DODGSON C.L.(Lewis Carroll) ,1867 *An elementary theory of déterminants*.
- [12] DORIER J.L.,1990. *Analyse historique de l'émergence des concepts élémentaires d'algèbre linéaire*, Cahier Didirem 7,IREM Paris VII.
- [13] FORSYTHE G.E.,1953,Solving linear algebraic equations can be intersting. *Bulletin of the American Mathematical Society*, 59,299-329
- [14] GAUSS C.F.,1810 Disquisitio de elementis ellipticis Palladis. Mémoire présenté le 25 Novembre 1810 à la Soc Roy Sciences de Göttingen. *Oeuvres* VI,1874,3-24
- [15] GAUSS C.F.,1821 Theoria Combinationis observationum erroribus minimis obnoxiae. 1^o partie présentée à la Soc Roy des Sciences de Göttingen en 1821, 2^o en 1823. Supplément en 1826. *Oeuvres*,IV,1873-, 1-108. (Trad. J. Bertrand. Méthodes des moindres carrés.Paris 1855)
- [16] GAUSS C.F.1823. Lettre du 26 Décembre 1823. *Briefwechsel zwischen Carl Friedrich Gauss und Christian Ludwig Gerling*. Berlin 1927
- [17] GERLING C.L.,1843. *Die Ausgleichungs-Rechnungen des practischen Geometrie, oder die Methode der kleinsten Quadrate mit Anschwenchungen für geödatische Aufgaben* (Appendice à Application du Calcul des compensations à la géométrie pratique ,ou la méthode des moindres carrés appliquée aux problèmes géodésiques. Appendice intitulé : Sur l'élimination des équations normales)
- [18] HØYRUP J.,1990. Algebra and naive geometry. An Investigation of Some Basic Aspects of Old Babylonian Mathematical Thought(II) *Altorientalische Forshungen* 17,1990,2.
- [19] HOUSEHOLDER A.S.,1953, *Principles of Numerical analysis*, New York

- [20] JACOBI C.G., 1845. Über eine neue Auflösungsart der bei der Methode der kleinsten Quadrate vorkommenden linearen Gleichungen. *Astronomische Nachrichten*, 22, p 297-306, *Gesammelte Werke*, 3, 1884, p 469-478.
- [21] JOZEAU M.F., 1992. *Comparaison de méthodes géodésiques en France et en Allemagne dans la première moitié du XIXe siècle*, Univ. Paris XIII.
- [22] KLINE M. 1972. *Mathematical thought from Ancient to Modern Times*, Oxford University Press.
- [23] KNOBLOCH E., 1980, *Der Beginn der Determinantentheorie, Leibnizens...* Hildersheim, Gerstenberg Verlag.
- [24] LAPLACE P.S., 1772. Recherches sur le calcul intégral et sur le système du monde. *Mémoires de l'Académie des Sciences de Paris. Oeuvres*, vol VIII.
- [25] LEGENDRE A.M. 1805. *Nouvelles méthodes pour la détermination des orbites des comètes*, Paris.
- [26] MACLAURIN C., 1748. *A treatise on algebra*, Londres. Trad fr. Paris 1753
- [27] MARTZLOFF J.C., 1987, *Histoire des mathématiques chinoises*, Masson, Paris.
- [28] MOULTON F.R., 1913. On the solutions of linear equations having small determinants, *The American Mathematical Monthly*, XX, p 242-249.
- [29] MUIR T., 1906-1923. *The theory of determinants in the historical order of development*, Londres.
- [30] NEKRASOV P.A. 1885. Détermination des inconnues par la méthode des moindres carrés quand le nombre d'inconnues est très grand. *Matematicheskij Sbornik* T 12, 189-203 (en russe)
- [31] SEIDEL L., 1874. Ueber ein Verfahren die Gleichungen, auf welche die Methode der kleinsten Quadrate führt, sowie lineare Gleichungen überhaupt, durch successive Annäherung aufzulösen. *Communication à la Section math-physique de L'Académie Royale des Sciences de Berlin, séance du 7 février 1874*.
- [32] SOUTHWELL R.V., 1940. *Relaxation methods in engineering science, a treatise on approximate computation*, Oxford University Press.
- [33] SPIESSER M., 1982. *Equations du premier degré. Méthode de fausse position*. IREM de Toulouse.
- [34] TEMPLE G., 1939 (ou 1949), The general theory of relaxation methods applied to linear systems, *Proceedings of the Royal Society of London*
- [35] TODD J., 1949, The condition of a certain matrix, *Proceedings of the Cambridge Philos.Soc.*, vol 46, p 116-118.
- [36] TROPFKE J., 1902-1903. *Geschichte der Elementarmathematik*. Leipzig.
- [37] VANDERMONDE A.T., 1772. Mémoire sur l'élimination. *Histoire de l'Académie Royale des Sciences de Paris*.
- [38] VAN DER WAERDEN B.L., 1983. *Geometry and algebra in ancient civilisations*, Springer Verlag.
- [39] WHITTAKER E.T. et ROBINSON G., 1924. *The calculus of observations* (a treatise on numerical mathematics).

ANNEXE 1

Lettre de Gauss à Gerling du 26 décembre 1823,

Briefwechsel zwischen Carl Friedrich Gauss und Christian Ludwig Gerling, Berlin, 1927

Traduit par A. Michel-Pajus

Göttingen, le 26 décembre 1823,

Ma lettre est parvenue trop tard à la poste et elle m'est revenue. Aussi je la décachette pour ajouter des indications pratiques relatives à l'élimination. A vrai dire, il y a là beaucoup de petits avantages particuliers que l'on n'apprend qu'à l'usage.

Je prends par exemple vos mesures d'Orber-Reisig¹.

Je prends d'abord [pour la direction de]² $1 = 0$

d'où, avec 1.3,

$$3 = 77^{\circ}57'53",107$$

(je choisis celui-ci car 1.3 a plus de poids que 1.2); donc avec

$$\begin{array}{rcl} 13 & 1.2 & 2 = 26^{\circ}44' 7",423 \\ 50 & 2.3 & 2 = \quad \quad 6, 507 \end{array} \quad \left| \quad 2 = 26^{\circ}44' 6",696 \right.$$

finalemeut avec

$$\begin{array}{rcl} 26 & 1.4 & 4 = 136^{\circ}21'13",481 \\ 6 & 2.4 & 4 = \quad \quad 8, 529 \\ 78 & 3.4 & 4 = \quad \quad 11,268 \end{array} \quad \left| \quad 4 = 136^{\circ}21'11",641 \right.$$

Pour améliorer encore l'approximation, j'obtiens, avec

$$\begin{array}{rcl} 13 & 1.2 & 1 = -0",727 \\ 28 & 1.3 & 1 = 0 \\ 26 & 1.4 & 1 = -1,840 \end{array} \quad \left| \quad 1 = -0",855 \right.$$

Puisque l'on ne s'occupe que de positions relatives, je peux ajouter +0",855 aux quatre valeurs et je pose

$$\begin{aligned} 1 &= 0^{\circ} 0' 0",000 + a \\ 2 &= 26 44 7, 551 + b \\ 3 &= 77 57 53,962 + c \\ 4 &= 136 21 12,496 + d \end{aligned}$$

Dans la méthode indirecte, il est très avantageux d'ajouter une modification à chaque direction. Vous pouvez vous en convaincre aisément en calculant sur le même exemple sans cette astuce, en outre, vous vous privez de la grande commodité d'avoir toujours un contrôle au moyen de la somme des termes absolus = 0. Maintenant, je

¹Les mesures d'angle données par Gerling (trouvées sur une feuille dans les papiers de Gauss) étaient, sachant que 1 représente Berger Warte, 2 Johannisberg, 3 Taufstein, 4 Milseburg

Répétitions	Angle
13	1.2 = 26°44' 7",423
28	1.3 = 77 57 53,107
26	1.4 = 136 21 13,481
50	2.3 = 51 13 46,600
6	2.4 = 109 37 1,833
78	3.4 = 58 23 18,161

(Note de Kruger qui a préparé le volume IX des *Werke*)

² ajouté par Kruger

forme les équations normales suivant le schéma suivant (en fait, si les termes sont trop nombreux, je sépare les négatifs des positifs¹):

$$\begin{array}{llll} ab - 1664 & ba + 1661 & ca + 23940 & da - 25610 \\ ac - 23940 & bc + 9450 & cb - 9450 & db + 8672 \\ ad + 25610 & bd - 18672 & cd - 29094 & dc + 29094 \end{array}$$

Les équations de condition sont donc :

$$\begin{array}{l} 0 = + 6 + 67a - 13b - 28c - 26d \\ 0 = - 7558 - 13a + 69b - 50c - 6d \\ 0 = -14604 - 28a - 50b + 156c - 78d \\ 0 = +22156 - 26a - 6b - 78c + 110d \\ \text{Somme} = 0 \end{array}$$

Pour l'élimination indirecte, je remarque que, si 3 des quantités a,b,c,d sont prises égales à 0, la quatrième prend la valeur la plus grande quand d est choisie comme quatrième. Naturellement chaque quantité doit être déterminée à partir de sa propre équation, et par conséquent, d à partir de la quatrième. Je pose donc $d = -201$ et substitue cette valeur. Les termes constants deviennent: +5232, -6352, +1074, +46; le reste est inchangé.

Maintenant c'est au tour de b, je trouve $b = +92$, je substitue et je trouve les termes absolus : +4036, -4, -3526, -506. Je continue ainsi jusqu'à ce qu'il n'y ait plus rien à corriger. Mais de tout ce calcul je n'écris en réalité que le tableau suivant :

	d = -201	b = +92	a = -60	c = +12	a = +5	b = -2	a = -1
+6	+5232	+4036	+16	-320	+15	+41	-26
-7558	-6352	-4	+776	+176	+111	-27	-14
-14604	+1074	-3526	-1846	+26	-114	-14	+14
+22156	+46	-506	+1054	+118	-12	0	+26

Dans la mesure où je ne pousse le calcul qu'au 2000ème de seconde près, je vois qu'il n'y a maintenant plus rien à corriger. J'en rassemble:

a = -60	b = +92	c = +12	d = -201
+5	-2		
-1			
-56	+90	+12	-201

en ajoutant à chacun la correction +56, j'obtiens:

a = 0	b = +146	c = +68	d = -145
-------	----------	---------	----------

ainsi les valeurs [des directions] sont:

1	0° 0' 0",000
2	26 44 7,697
3	77 57 54,030
4	136 21 12,351

Presque chaque soir, je fais une nouvelle édition du tableau, qu'il est facile d'améliorer n'importe où. Contre la monotonie du travail d'arpentage, c'est toujours une plaisante distraction; on peut aussi voir immédiatement si quelque chose de douteux s'est glissé, ce qui reste à obtenir, etc. Je vous recommande cette méthode comme modèle. Il ne vous arrivera presque plus jamais de pratiquer une élimination directe, du moins quand vous avez plus de deux inconnues. La procédure indirecte peut se faire à moitié endormi, ou en pensant à autre chose.

¹(Note de Krüger) l'unité des constantes est le millièème de seconde

ANNEXE 2

DETERMINATION DES INCONNUES PAR LA METHODE DES MOINDRES CARRÉS QUAND LE NOMBRE DES INCONNUES EST TRÈS GRAND.

P.A. Nekrassov

Extrait de *Matematicheskij Sbornik*, p 189-203, Tome 12, Moscou, 1885.

Traduit du Russe par Jacques SIMON, Professeur de Mathématiques Spéciales.

§ 1. L'astronomie et la géodésie présentent des situations où il faut déterminer par la méthode des moindres carrés un grand nombre d'inconnues. Dans ces cas la résolution par un système normal d'équations servant à définir les inconnues, présente d'énormes difficultés. Ainsi un système normal contient parfois jusqu'à 70 inconnues. Les déterminants au moyen desquels il faut exprimer ces inconnues contiennent chacun 70 lignes et sont constitués d'un nombre énorme de termes exprimé par le produit 1.2.3. ... 70. L'évaluation d'un tel déterminant est impensable sans extrêmes difficultés¹. Pour contourner ces difficultés, l'astronomie fait un calcul approché des inconnues. Parmi les procédés de ce genre l'un des plus efficaces est la méthode de Ludwig Seidel, qui consiste à faire le calcul approché successivement des solutions cherchées, et dont la description détaillée est donnée par l'auteur dans son mémoire appelé: "Über ein Verfahren die Gleichungen ,.... München"

Examinant la méthode de Seidel, à la demande du respecté astronome V.K. Tséraski, qui est concerné par les difficultés évoquées, j'ai remarqué que dans son mémoire, Seidel ne s'intéresse pas au problème, important dans ses implications pratiques, de la rapidité d'approximation des solutions. Pour combler cette lacune, je montrerai que dans les cas favorables la méthode de Seidel converge assez vite mais il est des cas où elle converge lentement et même infiniment lentement.

§2. Les fondements de la méthode de Seidel.

Soit $x, y, z, \dots, t$ formant m inconnues. Posant:

$$f_i = a_i x + b_i y + c_i z + \dots + p_i t - q_i$$

Supposons que les équations obtenues à la suite d'observations sont:

$$f_1 = 0, f_2 = 0, f_3 = 0, \dots, f_m = 0,$$

où $m \geq \mu$. Le système normal d'équations déduites de la méthode des moindres carrés, est, comme il est connu:

$$(aa)x + (ab)y + (ac)z + \dots + (ap)t - (aq) = 0 \quad (1)$$

$$(ba)x + (bb)y + (bc)z + \dots + (bp)t - (bq) = 0$$

$$(ca)x + (cb)y + (cc)z + \dots + (cp)t - (cq) = 0$$

$$\dots \dots \dots (pa)x + (pb)y + (pc)z + \dots + (pp)t - (pq) = 0$$

¹ Souvent grâce à une certaine symétrie et grâce aux observations déduites de la pratique, on peut alléger de tels calculs de déterminants dont dépend la résolution des équations. J'ai indiqué comment procéder dans une communication faite à la Société des Amateurs des Sciences Naturelles et à la Société Mathématique. Je développerai cette méthode dans l'un des articles suivants.

$$\text{où: } (kl) = k_1 l_1 + k_2 l_2 + \dots + k_m l_m$$

Les grandeurs $x, y, z, \dots, t$ qui satisfont à l'équation (1), comme il est bien connu, rendent minimum le polynôme:

$$Q = \frac{1}{2}(f^2_1 + f^2_2 + \dots + f^2_m) \quad (2)$$

En ordonnant le deuxième membre par rapport à x , on obtient:

$$Q = \frac{1}{2}((aa)x^2 + 2(ab)xy + 2(ac)xz + \dots + 2(ap)xt - 2(aq)x) + R \quad (2')$$

où R est un polynôme ne contenant pas x . L'égalité (2') peut encore s'écrire ainsi:

$$Q = \frac{1}{2(aa)}\{(aa)x + (ab)y + (ac)z + \dots + (ap)t - (aq)\}^2 + P \quad (3)$$

où P est un polynôme ne contenant pas x .

Soit $x_o, \dots, t_o$ un système de valeurs particulières remplaçant $x, \dots, t$, et ne vérifiant pas les équations (1). Alors la grandeur Q correspondant à ces valeurs sera:

$$Q_o = \frac{1}{2(aa)}\{(aa)x_o + (ab)y_o + (ac)z_o + \dots + (ap)t_o - (aq)\}^2 + P_o \quad (3')$$

Soit x_1 une grandeur vérifiant l'équation:

$$(aa)x_1 + (ab)y_o + (ac)z_o + \dots + (ap)t_o - (aq) = 0$$

Il est clair que pour le système de valeurs particulières $x_1, y_o, \dots, t_o$ remplaçant $x, y, \dots, t$ la grandeur Q prendra une valeur Q_1 , égale à P_o , et en outre nous aurons: $Q_1 < Q_o$.

Ainsi pour les valeurs $x_1, y_o, \dots, t_o$ la grandeur Q est plus proche de son minimum, que pour les valeurs $x_o, y_o, \dots, t_o$ des variables $x, y, \dots, t$.

Soit y_1 une grandeur satisfaisant à l'équation:

$$(ba)x_1 + (bb)y_1 + (bc)z_o + \dots + (bp)t_o - (bq) = 0$$

Par la méthode indiquée, il est facile de se convaincre que Q qui peut se présenter sous la forme:

$$Q = \frac{1}{2(bb)}\{(ba)x + (bb)y + (bc)z + \dots + (bp)t - (bq)\}^2 + L$$

pour les valeurs $x_1, y_1, z_o, \dots, t_o$ des variables $x, y, \dots, t$ sera plus proche de son minimum que la grandeur Q correspondant aux valeurs $x_1, y_o, \dots, t_o$.

Poursuivant ainsi par la même méthode on peut successivement remplacer toutes les grandeurs $x_o, y_o, \dots, t_o$ par $x_1, y_1, \dots, t_1$ puis de la même façon par $x_2, y_2, \dots, t_2$ etc... A la limite après un nombre infini d'opérations Q tendra vers son minimum et le

En éliminant les grandeurs B, C, ..., P, on obtient pour un α déterminé l'égalité suivante:

$$\begin{vmatrix} (aa)\alpha & (ab) & (ac) & \dots & (ap) \\ (ba)\alpha & (bb)\alpha & (bc) & \dots & (bp) \\ (ca)\alpha & (cb)\alpha & (cc)\alpha & \dots & (cp) \\ \dots & \dots & \dots & \dots & \dots \\ (pa) & (pb) & (pc) & \dots & (pp) \end{vmatrix} = 0 \quad (7)$$

Cette équation de degré $\mu-1$ par rapport à a admet $\mu-1$ racines $a_1, \dots, a_{\mu-1}$. A chacune d'entre elles correspond un système de grandeurs B, C, ..., P, déterminées à l'aide des égalités (6). Il y a $\mu-1$ intégrales particulières de la forme (5) avec une constante arbitraire A. A partir de ces intégrales particulières il est facile de former l'intégrale générale qui sera:

$$\begin{aligned} \xi_n &= A_1 \alpha_1^n + A_2 \alpha_2^n + \dots + A_{\mu-1} \alpha_{\mu-1}^n \\ \eta_n &= A_1 B_1 \alpha_1^n + A_2 B_2 \alpha_2^n + \dots + A_{\mu-1} B_{\mu-1} \alpha_{\mu-1}^n \\ \zeta_n &= A_1 C_1 \alpha_1^n + A_2 C_2 \alpha_2^n + \dots + A_{\mu-1} C_{\mu-1} \alpha_{\mu-1}^n \\ &\dots \dots \dots \\ \tau_n &= A_1 P_1 \alpha_1^n + A_2 P_2 \alpha_2^n + \dots + A_{\mu-1} P_{\mu-1} \alpha_{\mu-1}^n \end{aligned} \quad (8)$$

où $A_1, \dots, A_{\mu-1}$ sont des constantes arbitraires d'intégration, $\alpha_i, B_i, C_i, \dots, P_i$ sont des grandeurs définies par les égalités (6), puis α_i satisfait à l'égalité (7). Les constantes arbitraires $A_1, \dots, A_{\mu-1}$ sont liées aux erreurs initiales $h_0, x_0, \dots, t_0$ et aux valeurs choisies arbitrairement $y_0, \dots, t_0$ par les égalités:

$$\begin{aligned} \eta_0 &= y_0 - y' = A_1 B_1 + A_2 B_2 + \dots + A_{\mu-1} B_{\mu-1} \\ \zeta_0 &= z_0 - z' = A_1 C_1 + A_2 C_2 + \dots + A_{\mu-1} C_{\mu-1} \\ &\dots \dots \dots \\ \tau_n &= t_0 - t' = A_1 P_1 + A_2 P_2 + \dots + A_{\mu-1} P_{\mu-1} \end{aligned} \quad (9)$$

Si l'équation (6) a des racines non toutes distinctes, les seconds membres des égalités (8) doivent être modifiés de façon évidente, mais les conclusions ci-dessous restent valables.

§ 4. Rapidité de convergence de la méthode de Seidel

Comme l'erreur sur les grandeurs $x_n, y_n, \dots, t_n$ doit, comme il a été prouvé au §2, tendre vers 0 quand n tend vers $+\infty$, alors les racines de l'équation (7) doivent avoir un module inférieur à 1. Seulement à cette condition les seconds membres des égalités (8) tendront vers 0 quand n augmente indéfiniment.

Soit α la racine de (7) de plus grand module, nous l'appellerons la racine principale. Par rapport à la racine α , les erreurs $\xi_n, \dots$ déterminées par les égalités (8) seront $A\alpha^n, AB\alpha^n, AC\alpha^n, \dots$. La grandeur A, dépendant, comme le montrent les équations (9) de $y_0 - y', \dots, t_0 - t'$, ne peut être donnée à l'avance car $y', \dots, t'$ sont supposées inconnues. Ainsi pour un choix arbitraire de $y_0, \dots, t_0$, A ne sera pas nulle en

général et c'est pourquoi les valeurs $A\alpha^n, AB\alpha^n \dots$ des erreurs, quand n augmentera, diminueront lentement en comparaison avec les autres erreurs. Dans la méthode de Seidel la convergence sera rapide si la racine principale a un module petit. Mais si la racine principale a un module voisin de 1, ce qui n'est pas rare, alors la convergence vers les solutions de (1) sera lente sauf si, mais c'est peu probable, le choix des valeurs initiales $y_0 \dots t_0$ entraîne $A = 0$.

De façon générale la rapidité de convergence de la méthode de Seidel est entièrement liée à la valeur de la racine principale de l'équation (7). Si toutes les solutions de (7), y compris la principale sont petites la méthode convergera vite, sinon elle convergera lentement. Nous allons donner un exemple de chacun de ces cas.

Exemple I. Soit $m = \mu = 3$

$$\begin{aligned} f_1 &= x - 3y + z - 3 = 0 \\ f_2 &= 2x + y - 4z - 10 = 0 \\ f_3 &= x + 1,01y + 7,1z - 9,1 = 0 \end{aligned}$$

Le système normal en fonctions de ces données sera:

$$\begin{aligned} 6x + 0,01y + 0,1z - 12,1 &= 0 \\ 0,01x + 11,0201y + 0,171z - 0,191 &= 0 \\ 0,1x + 0,171y + 67,41z - 67,61 &= 0 \end{aligned}$$

L'équation (7) prendra la forme:

$$4457,189646\alpha^2 - 0,28657801\alpha + 0,000171 = 0$$

Il y aura deux racines imaginaires conjuguées. Le module de chacune d'elles sera $\rho = 0,0000196$. La méthode de Seidel convergera très rapidement.

En prenant par exemple $y_0 = 10$ et $z_0 = 3$, on trouve:

$$x_1 = 1,95, \quad y_1 = -0,0309888, \quad z_1 = 1,000152$$

Ces grandeurs sont déjà relativement proches des solutions exactes qui sont 2, 0 et 1.

Appliquant la méthode de Seidel une deuxième fois on obtiendra les résultats suivants presque exacts:

$$x_2 = 2,0000491, \quad y_2 = -0,00000240, \quad z_2 = 0,99999993$$

Exemple II. Soit $m = \mu = 3$,

$$\begin{aligned} f_1 &= x + y + z - 6 = 0 & 6x + 5y + 4z - 28 &= 0 \\ f_2 &= 2x + y + z - 7 = 0 & 5x + 6y + 4z - 29 &= 0 \\ f_3 &= x + 2y + z - 8 = 0 & 4x + 4y + 3z - 21 &= 0 \end{aligned}$$

L'équation (7) prendra la forme: $108\alpha^2 - 187\alpha + 80 = 0$

Les racines de cette équation sont: $\alpha_1 = 0,959 \dots$ et $\alpha_2 = 0,872 \dots$

La racine principale est proche de 1. La convergence sera lente.

Si on pose par exemple $y_0 = 3, z_0 = 4$, on trouve :

$$\begin{aligned} x_1 &= -0,5, & y_1 &= 2,58333, & z_1 &= 4,22222 \\ x_2 &= -0,300925, & y_2 &= 2,26929, & z_2 &= 4,37557 \\ x_3 &= -0,141455, & y_3 &= 2,04095, & z_3 &= 4,47638 \end{aligned}$$

Ces valeurs sont sensiblement différentes des valeurs exactes: 1, 2, 3.

Le système d'équations $f_1 = 0, f_2 = 0, f_3 = 0, \dots, f_m = 0,$ (10)

peut être choisi tel que la racine principale de l'équation (7) soit aussi proche que l'on veut de 1. En effet soit $m - \mu + p$ relations linéaires entre les fonctions $f_1, \dots, f_m$ de la forme : $\lambda_1 f_1 + \lambda_2 f_2 + \dots + \lambda_m f_m$, où les coefficients $\lambda_1 \dots \lambda_m$ sont des constantes. Si $p \geq 1$ le système (10) sera indéterminé. En même temps le système (1) sera aussi indéterminé et en outre on aura :

$$\begin{vmatrix} (aa) & (ab) & (ac) & \dots & (ap) \\ (ba) & (bb) & (bc) & \dots & (bp) \\ (ca) & (cb) & (cc) & \dots & (cp) \\ \dots & \dots & \dots & \dots & \dots \\ (pa) & (pb) & (pc) & \dots & (pp) \end{vmatrix} = 0 \quad (11)$$

Par suite l'équation (7) aura une racine α égale à 1

Donnons aux grandeurs $a_1, \dots, b_1, \dots, p_1, \dots$ un petit accroissement de sorte que le système (1) devienne déterminé. En même temps le déterminant premier membre de (11) sera non nul bien que très petit et l'équation (7) aura une racine très proche de 1. On peut obtenir de tels systèmes avec racine principale de module proche de 1 par d'autres méthodes. On peut faire en sorte que cette racine soit proche de -1, ou soit un nombre imaginaire de module 1. Par suite les systèmes de cette sorte peuvent se rencontrer très souvent. Dans de telles situations qu'on peut observer effectivement, la méthode de Seidel se montre non adéquate.

Remarquons encore que l'équation (7) peut se présenter sous la forme :

$$(aa)(bb)(cc) \dots (pp) \alpha^{\mu-1} + \dots \pm (ab)(bc)(cd) \dots (pa) = 0 \quad (7')$$

Il en résulte :

$$\alpha_1 \alpha_2 \dots \alpha_{\mu-1} = \pm \frac{(ab)(bc) \dots (pa)}{(aa)(bb) \dots (pp)}$$

Cette égalité montre que $\text{mod}(\alpha) > K \quad (12)$

où α est la racine principale de l'équation (7) et

$$K = \text{mod} \sqrt[\mu-1]{\frac{(ab)(bc)(cd) \dots (pa)}{(aa)(bb)(cc) \dots (pp)}} \quad (13)$$

Si la valeur de K n'est pas assez petite, alors il est clair que la méthode de Seidel convergera lentement.

Remarquons enfin que la méthode de Seidel convergera rapidement quand les grandeurs $(aa), \dots, (pp)$ seront très grandes par rapport aux autres coefficients liés aux inconnues dans les équations (1). En effet dans ce cas le coefficient du premier terme de l'égalité (7') sera très grand par rapport aux autres termes de cette égalité et par suite toutes les racines de (7') et de (7) seront très petites.

L'ordre de recherche des valeurs approchées par la méthode de Seidel peut être différent de celui qui a été indiqué au §3. Par suite la rapidité de convergence peut visiblement être quelque peu améliorée. Mais dans le cas où la racine principale de (7) est très proche de 1, un tel procédé ne rend pas la méthode de Seidel suffisamment rapide pour être utilement utilisé dans la pratique. Pour prouver cette affirmation je vais examiner un cas où l'ordre de calcul le plus favorable donne une convergence très lente. etc... Les grandeurs $M_1, \dots, M_m$ étant calculées, choisissons la plus grande. Supposons que ce soit :

$$M_3 = \frac{N_3^2}{2(cc)}$$

De façon évidente la première approximation optimum concernera z_0 . L'ayant réalisée, on aura :

$$z_1 = z_0 - \frac{N_3}{(cc)}$$

Supposons que les premiers membres de (1) après remplacement des grandeurs $x, y, \dots, z$ par $x_0, \dots, z_0$ prennent les valeurs $N'_1, \dots, N'_\mu$. Ensuite nous poserons:

$$M'_1 = \frac{(N'_1)^2}{2(aa)} \quad , \quad M'_2 = \frac{(N'_2)^2}{2(bb)} \quad , \dots \quad , \quad M'_\mu = \frac{(N'_\mu)^2}{2(pp)}$$

La deuxième approximation devra concerner celle des grandeurs $x, y, z_1, \dots, z_0$ qui correspond à la plus grande des grandeurs:

$$M'_1, \dots, M'_\mu$$

etc² En effectuant ainsi les calculs, la méthode de Seidel permettra d'approcher les inconnues un peu plus rapidement, mais le gain en rapidité n'est pas grand. En effet si le module de la racine principale de l'équation (7) est très proche de 1, alors après quelques opérations par la méthode de Seidel, les grandeurs $M_1, \dots, M_m$ deviennent trop peu significatives et la plus grande est de très peu inférieure à Q, sans tenir compte des erreurs significatives sur les solutions. Nous allons éclaircir cela par un exemple numérique.

Exemple: Soit $m = \mu = 3$ et

$$f_1 = 10x + y + z - 15 = 0$$

$$f_2 = 20x + 2y + z - 27 = 0$$

$$f_3 = 31x + 3y + 2z - 43 = 0$$

Le système normal en fonctions de ces données sera:

$$\begin{aligned} 1461x + 143y + 92z - 2023 &= 0 \\ 143x + 14y + 9z - 198 &= 0 \\ 92x + 9y + 6z - 128 &= 0 \end{aligned} \quad (14)$$

L'équation (7) prendra la forme: $122724 \alpha^2 - 241127 \alpha + 118404 = 0$

Cette équation admet une racine très proche de 1.

Les solutions des équations (14) sont: $x = 1, y = 2, z = 3$.

Soit $x_0 = 0,9, \dots, y_0 = 1,9, \dots, z_0 = 2,9$.

On trouve: $N_1 = -169,6 \quad N_2 = -16,6 \quad N_3 = -10,7$
 $M_1 = 9,843997 \quad M_2 = 9,8411 \quad M_3 = 9,5408$

La première approximation doit concerner la grandeur x_0 , on aura:

$$x_1 = x_0 + 0,1160849 = 1,0160849.$$

Ensuite nous trouvons: $N'_1 = 0,0000389 \quad N'_2 = 0,0001407 \quad N'_3 = -0,0201892$
 $M'_1 = 0 \quad M'_2 = 0,0000000007 \quad M'_3 = 0,000034$

La deuxième approximation doit concerner z_0 , après quoi Q deviendra très petit.

A partir de z_0 on obtiendra: $z_1 = 2,9033649$.

Ensuite nous trouvons: $N''_1 = 0,3096097 \quad N''_2 = 0,0304248 \quad N''_3 = 0,0000002$
 $M''_1 = 0,00003281 \quad M''_2 = 0,00003306 \quad M''_3 = 0$

La troisième approximation concernera y_0 , après quoi Q sera très petit.

Pour y_0 , on aura: $y_1 - y_0 = -0,0021732$. Cette approximation éloigne même la valeur approchée y de sa valeur exacte.

Ensuite nous trouvons: $N'''_1 = -0,0011579 \quad N'''_2 = 0 \quad N'''_3 = -0,0195586$
 $M'''_1 = 46 \cdot 10^{-11} \quad M'''_2 = 0 \quad M'''_3 = 0,0000318$

On obtiendra alors: $z_2 - z_1 = 0,0032598$.

La petitesse de l'approximation par rapport à l'importance de l'erreur montre qu'il faut procéder à un calcul long pour obtenir une approximation qui n'est que de 0,01.

² Dans la pratique un tel ordre de calcul possède l'inconvénient de rendre le calcul plus compliqué puisque en plus de l'ordre il convient de calculer une série de valeurs prises par la lettre M.

ANNEXE 3

SUR UN PROCÉDE POUR RESOUDRE PAR APPROXIMATIONS SUCCESSIONS LES EQUATIONS (DITES "NORMALES") FOURNIES PAR LA METHODE DES MOINDRES CARRÉS OU, EN GENERAL, TOUT SYSTEME D'EQUATIONS LINÉAIRES.

par Ludwig SEIDEL

(traduction Colette BLOCH)

Communication à la section math-physique de l'Académie Royale des Sciences, Berlin,

Séance du 7 février 1874.

1

En sus des hypothèses communément admises dans la théorie de l'ajustement des résultats d'observations, supposons que chacun de ces résultats s'exprime par les équations linéaires suivantes en fonction des inconnues $x, y, z, \dots$:

$$\begin{aligned} & a_1x + b_1y + c_1z + \dots + n_1 = 0 \\ (A) \quad & a_2x + b_2y + c_2z + \dots + n_2 = 0 \\ & a_3x + b_3y + c_3z + \dots + n_3 = 0 \quad \text{etc.} \end{aligned}$$

où le nombre des observations (et par conséquent celui des équations A) est plus grand que le nombre des inconnues $x, y, z, \dots$, mais où chaque équation est altérée par les erreurs d'observation (ou pourrait l'être) et où les seconds membres ne seraient donc en général pas nuls, même si on pouvait substituer à $x, y, z, \dots$ leurs valeurs exactes.

On suppose également que les poids éventuellement différents des différentes observations sont déjà inclus dans la forme des équations, selon la règle habituelle, si bien qu'on n'a plus aucune raison à priori de s'attendre à une erreur plus grande dans telle équation que dans telle autre, c'est-à-dire à une plus grande valeur absolue de l'expression qui devrait être nulle de par l'équation.

On suppose encore que les hypothèses relatives aux conditions d'observation, qui justifient rationnellement l'application de la "méthode des moindres carrés", sont adéquates ou du moins qu'on peut les considérer comme telles. Il s'ensuit que le système le plus probable⁽¹⁾ de valeurs de $x, y, z, \dots$ sera celui qui vérifie la condition suivante: rendre minimale la somme des carrés des premiers membres des équations (A).

Employons, comme il est d'usage dans cette théorie, les crochets [] comme signe de sommation, si bien que, par exemple

$$[aa] = a_1^2 + a_2^2 + a_3^2 + \dots$$

$$[ab] = [ba] = a_1 b_1 + a_2 b_2 + a_3 b_3 + \dots$$

chaque somme a autant de termes qu'il y a d'observations), alors la somme des carrés s'écrit

$$\begin{aligned} Q = & [aa]x^2 + [bb]y^2 + [cc]z^2 + \dots \\ & + 2[ab]xy + 2[ac]xz + 2[bc]yz + \dots \\ & + 2[an]x + 2[bn]y + 2[cn]z + \dots \\ & + [nn] \end{aligned}$$

et elle est minimale quand les valeurs des inconnues vérifient les équations normales suivantes:

$$\begin{aligned} (B) \quad & [aa]x + [ab]y + [ac]z + \dots + [an] = 0 \\ & [ab]x + [bb]y + [bc]z + \dots + [bn] = 0 \\ & [ac]x + [cb]y + [cc]z + \dots + [cn] = 0 \quad \text{etc.} \end{aligned}$$

Si simple que soit, par sa nature mathématique, le problème de calculer un nombre quelconque d'inconnues à partir d'un même nombre d'équations linéaires, il n'en reste pas moins très pénible quand le nombre des inconnues est important, et dans de tels cas (par exemple la résolution d'un réseau géodésique quelque peu étendu) on se voit réduit à sacrifier le calcul systématique exact, à former des systèmes partiels d'inconnues et à les recoller tant bien que mal après les avoir résolus un par un. J'ignore si on a jamais intégralement calculé un complexe de plus de soixante-dix inconnues. Le nombre de 70 est atteint dans le réseau de triangulation de la Prusse orientale (et de plus dans ce cas les inconnues sont liées par 31 conditions impératives, circonstance qui augmente encore la difficulté de la résolution classique). Personnellement j'ai eu à faire à 72 inconnues pour calculer les valeurs les plus probables des logarithmes des luminosités des étoiles qui intervenaient dans mon réseau photométrique.

La méthode habituelle de résolution, donnée par Gauss, consiste, comme on sait, à exprimer une inconnue en fonction des autres en la tirant de l'équation où elle a pour coefficient une somme de carrés, autrement dit où elle figure sur la diagonale principale du système, ainsi par exemple la valeur de x sera tirée de la première équation dans (B) et portée dans les autres équations, ce qui fournit le premier système transformé des équations normales :

$$\begin{aligned} (C) \quad & [bb.1]y + [bc.1]z + \dots + [bn.1] = 0 \\ & [bc.1]y + [cc.1]z + \dots + [cn.1] = 0 \quad \text{etc.} \end{aligned}$$

qui partage avec le système (B) d'origine la propriété de symétrie des coefficients par rapport à la diagonale et dans lequel on a posé:

$$[bc.1] = [cb.1] = [bc] - \frac{[ab][ac]}{[aa]}$$

On procède de même à la substitution d'une deuxième inconnue, etc. et on obtient des systèmes transformés dont chacun contient une inconnue de moins que le précédent jusqu'à la dernière inconnue qui figure seule et chacune des autres est calculée par l'équation qui a servi à son élimination en remontant dans l'ordre inverse.

Jacobi a élaboré un autre procédé et l'a appliqué aux 7 équations établies par LeVerrier pour calculer une partie des irrégularités séculaires du système planétaire d'après Laplace; encore étudiant j'ai eu l'honneur de faire pour lui

les calculs numériques. Par ce procédé on arrive à annuler les plus gros coefficients non situés sur la diagonale, grâce à une substitution linéaire convenable (correspondant à une rotation d'axes orthogonaux) qui introduit, à la place de deux inconnues à fort coefficient, deux nouvelles inconnues, en préservant la symétrie de l'ensemble ; il est vrai que dans la suite du calcul les nouvelles substitutions réintroduisent un coefficient non nul, mais la somme des carrés des coefficients hors diagonale diminue constamment au profit des coefficients diagonaux et, par ce procédé dont la convergence est rigoureusement démontrée, on s'approche autant qu'on veut du but final où (en dehors des termes constants) il ne reste que les termes diagonaux qui fournissent immédiatement les dernières inconnues. Cette méthode est parfaitement adaptée aux difficultés particulières du cas auquel Jacobi l'a appliquée (où les coefficients diagonaux étaient eux-mêmes fonctions linéaires de variables complémentaires), mais par ailleurs elle ne me paraît en aucune façon plus avantageuse que la méthode généralement employée, dans les cas simples qui se présentent habituellement ; et je doute quelle ait jamais été appliquée jusqu'à présent à un autre cas.

J'ai ouvert une troisième voie dans mon travail photométrique cité plus haut; le choix de cette méthode était particulièrement indiqué par la forme simple des équations expérimentales. Dans le présent article j'ai le dessein d'exposer et de démontrer la méthode de résolution sous la forme où elle est universellement applicable et de dégager en même temps quelques particularités qui lui sont propres.

2

Supposons qu' on ait pris d'abord pour les inconnues $x, y, z, \dots$ un ensemble de valeurs quelconques qui ne vérifie pas encore le système le plus probable des "équations normales" (B) mais qui donne

$$[aa]x + [ab]y + \dots [an] = N_1$$

$$[ab]x + [bb]y + \dots [bn] = N_2 \quad \text{etc.} \dots$$

Par identification, la somme des carrés des écarts peut s'écrire:

$$\begin{aligned} Q &= \frac{1}{[aa]}([aa]x + [ab]y + [ac]z + \dots + [an])^2 \\ &+ [bb.1]y^2 + [cc.1]z^2 + \dots \\ &+ 2[bc.1]yz + \dots + 2[bn.1]y + 2[cn.1]z + \dots + [nn.1] \\ &= \frac{1}{[aa]}N_1^2 + [bb.1]y^2 + \dots + [nn.1] \end{aligned}$$

Sous cette forme l'inconnue x ne figure que dans le premier terme (c'est à dire dans N_1); d'où l'on voit aussitôt qu'on diminuera Q de la quantité

$$\frac{1}{[aa]} N_1^2$$

et, sans changer les valeurs prises pour $y, z, \dots$, on modifie la valeur de x de telle sorte que l'expression qui valait N_1 devienne nulle.

$$\text{Ce sera réalisé en corrigeant } x \text{ de } \Delta x = - \frac{N_1}{[aa]}$$

et la valeur ainsi améliorée $x + \Delta x$ est celle qui, pour la première inconnue, s'accorde le mieux aux valeurs choisies pour les autres inconnues et serait alors la meilleure si les valeurs provisoires des autres variables étaient déjà les

vraies valeurs cherchées.

La modification de x , qui remplace N_1 par $N'_1 = 0$ modifie du même coup les valeurs de $N_2, N_3, \dots$ en

$$N'_2 = N_2 + [ab] \Delta x$$

$$N'_3 = N_3 + [ab] \Delta x \quad \text{etc.}$$

Si, au lieu de garder $y, z, \dots$ et de corriger x de $-\frac{N_1}{[aa]}$, on avait gardé $x, z, \dots$

et corrigé y de $-\frac{N_2}{[bb]}$, la somme Q aurait diminué de $\frac{N_2^2}{[bb]}$ la somme Q aurait

diminué de $\frac{N_3^2}{[cc]}$ si on avait corrigé z seul (de $-\frac{N_3}{[cc]}$.)

Une deuxième étape va donc consister à corriger y de $\Delta y = -\frac{N'_2}{[bb]}$

$y + \Delta y$ sera la meilleure valeur s'accordant au système de valeurs $x + \Delta x, z, \dots$

et réduisant Q de $\frac{N'^2_2}{[bb]}$ alors que Q avait déjà diminué de $\frac{N^2_1}{[aa]}$

Cette modification de y remplace le système de valeurs $N'_1 = 0, N'_2, N'_3, \dots$

par $N''_1, N''_2 = 0, N''_3, \dots$ où

$$N''_1 = N'_1 + [ab] \Delta y = [ab] \Delta y$$

$$N''_2 = N'_2 + [bb] \Delta y = 0$$

$$N''_3 = N'_3 + [bc] \Delta y \quad \text{etc.}$$

Si maintenant on corrigeait z de telle façon que sa nouvelle valeur $z + \Delta z$ s'accorde le mieux possible aux valeurs $x + \Delta x, y + \Delta y$ et aux valeurs initiales des inconnues suivantes, on diminuerait Q de

$$\frac{N''^2_3}{[cc]} ;$$

par contre Q diminuera de $\frac{N''^2_1}{[aa]}$ si on revient à x (car $x + \Delta x$ n'est plus la meilleure

valeur convenant à $y + \Delta y, z, \dots$) et qu'on lui apporte une deuxième correction $-\frac{N''_1}{[aa]}$.

Si donc partant d'un système quelconque de valeurs initiales, on les modifie successivement en les prenant dans un ordre arbitraire (et il n'est pas nécessaire de parcourir tout le cycle avant de revenir à une inconnue dont la valeur a déjà été corrigée), tout en prenant soin de déterminer toujours chaque correction de façon à vérifier celle des équations normales où la variable concernée occupe la position "diagonale", alors on diminue pas à pas la somme des carrés des écarts (et chaque fois d'une quantité immédiatement connue, de la forme

$\frac{N_2}{[aa]}$ ou $[aa] \Delta x^2$), et ceci aussi longtemps qu'on peut la diminuer.

Car les amoindrissements de Q et les corrections à apporter aux inconnues (la dernière de la forme $-\frac{N}{[aa]}$) ne deviennent négligeables que lorsque tous les N simultanément ont décré jusqu'à des valeurs pratiquement nulles.

Quand ce but est atteint, toutes les équations normales sont vérifiées et, par améliorations successives, les inconnues ont pris les valeurs les plus vraisemblables.

Il faut bien remarquer que la preuve de la décroissance de Q et aussi de la convergence de ce processus d'approximation repose entièrement sur la condition que pour chaque variable on calcule les améliorations successives par rapport à l'équation où elle figure en terme diagonal; car si, dans un système quelconque de n équations linéaires à n inconnues, on prétendait partir d'un ensemble de valeurs arbitraire pour les inconnues et apporter des améliorations successives en déterminant la correction pour x à partir d'une quelconque des équations, puis la correction pour y à partir d'une autre, arbitrairement choisie, etc....., - on ne pourrait pas prouver qu'en général on se rapproche indéfiniment de la solution du système - il pourrait bien plutôt arriver (comme on s'en convaincra facilement) que les valeurs successives des inconnues tendent vers l'infini ou oscillent indéfiniment entre deux valeurs distinctes finies.

Mais tout le système, de n équations linéaires à n inconnues peut être mis sous la forme normale

$$\begin{aligned} (B) \quad & [aa] x + [ab] y + [ac] z + \dots + [an] = 0 \\ & [ab] x + [bb] y + [bc] z + \dots + [bn] = 0 \\ & [ac] x + [cb] y + [cc] z + \dots + [cn] = 0 \quad \text{etc.....} \end{aligned}$$

à partir des équations

$$\begin{aligned} (A) \quad & a_1 x + b_1 y + c_1 z + \dots + n_1 = 0 \\ & a_2 x + b_2 y + c_2 z + \dots + n_2 = 0 \\ & a_3 x + b_3 y + c_3 z + \dots + n_3 = 0 \quad \text{etc.....} \end{aligned}$$

et sous la forme (B) on peut le résoudre par notre méthode, comme on le fait pour les équations normales déduites d'observations expérimentales. L'oscillation sans fin, ou la croissance sans borne des valeurs des variables sont, dans les deux cas, exclues a priori, puisqu'il est prouvé que chaque étape rapproche du but (diminuer la somme des carrés) et qu'on cesse de s'en rapprocher de façon appréciable seulement quand toutes les équations (B) sont vérifiées à e près et que toutes les inconnues ont alors atteint leurs valeurs définitives. La somme Q existe naturellement tout autant pour un système qui doit être résolu exactement que pour le système déduit d'observations empiriques; la seule différence est que dans le premier cas la valeur minimale qu'elle obtient est 0.

(1)La supériorité de ce système le plus probable tient à ce que la probabilité pour que les vraies valeurs des inconnues soient dans des intervalles de longueur donnée autour de valeurs calculées est plus grande avec ce système qu'avec tout autre (pour les mêmes amplitudes d'intervalles). Ce ne paraît pas toujours être compris avec la netteté désirable.

Le volume de la pyramide

Michèle Grégoire

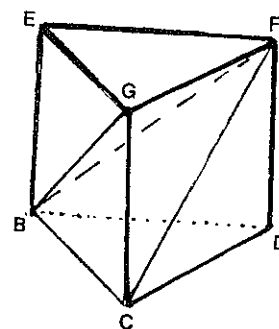
Mesurer de façon élémentaire un solide consiste à le comparer à un volume simple pris pour unité. La pyramide triangulaire est un solide simple dont le calcul du volume n'est pas immédiat. On peut, à son propos, tenter de faire faire aux élèves le lien entre la notion commune de volume et des notions plus abstraites liées à la mesure d'un volume: limites, somme des termes d'une suite, intégrale...

Un des exemples proposés est un problème destiné à une activité de type travaux dirigés pour des élèves de première ou de terminale scientifiques, qui commence par des constructions mettant en évidence qu'un prisme peut être décomposé en trois pyramides. Dans certains cas particuliers, il est évident que ces pyramides ont même volume, lorsque ces pyramides sont superposables ou symétriques. Pour démontrer que, dans le cas général, les trois pyramides ont le même volume, le problème introduit la lecture d'un extrait des *Eléments de Géométrie* de Legendre (12^{ème} édition, 1823). La propriété est de nature infinitésimale, mais le raisonnement de Legendre évite tout recours à l'infini, grâce à un raisonnement par l'absurde. La preuve de Legendre fait sentir ce qui débouchera sur la formulation technique précise du concept de limite qu'on n'expose plus aux élèves de lycée. Elle permet aussi d'aborder en dernière partie une autre méthode, qui annonce l'usage des sommes de Riemann pour définir le volume de la pyramide.

Problème

I Constructions

1) On découpe le prisme droit BCDEFG de hauteur a et de base BCD triangle rectangle isocèle de côté a , en trois pyramides comme l'indique la figure 1. Dessiner les patrons des trois pyramides. Les construire et les assembler pour former le prisme. Comparer les patrons et les trois pyramides. Quelle relation entre les volumes du prisme et des pyramides peut-on conjecturer ?



2) Faire les mêmes constructions avec le prisme droit B'C'D'E'F'G', de base B'C'D', où $B'C = a$, $C'D' = 2a$, $B'C'D' = 60^\circ$, $D'F' = 3a$. La conjecture ci-dessus s'applique-t-elle dans ce cas? Montrer que les pyramides ont deux à deux même hauteur et des bases d'aires égales.

II Lecture de la proposition XVII du livre 6 des *Eléments de Géométrie* de Legendre.

Les questions posées se réfèrent aux différents paragraphes du texte joint.

§ 1. Préciser les hypothèses et la conclusion recherchée. Quel mode de raisonnement emploie Legendre? Le prisme de base ABC et de hauteur Ax. est plus loin dans le texte désigné par ABCX. Que représente-t-il?

§ 2. Posons $n = \frac{AT}{k}$. Préciser la transformation de l'espace par laquelle ABC a pour image DEF, et celle qui transforme *abc* en *def*. Quel prisme inscrit dans *sabc* a même volume que DEFG?

§ 3. Prouver la propriété énoncée. On notera V et v les volumes de SABC et de *sabc*.

§ 4. Combien de prismes ont été construits dans chaque pyramide? On notera, en partant du sommet de chaque pyramide, $P_1, P_2, \dots, P_i, \dots$, les prismes construits autour de SABC, et $p_1, p_2, \dots, p_i, \dots$, les prismes inscrits dans *sabc*. Ecrire avec ces notations les inégalités énoncées par Legendre. Comment conclut-il?

III Application.

Soit SABC une pyramide quelconque et ABCQPS le prisme de base ABC et d'arête AS. Montrer que ce prisme se décompose en trois pyramides de même volume que SABC.

IV Une autre méthode pour calculer le volume V de la pyramide SABC.

En s'inspirant de la figure de Legendre, construire une pyramide de hauteur h coupée par des plans équidistants et parallèles à la base ABC qui définissent des prismes intérieurs $p_1, p_2, \dots$, et des prismes extérieurs $P_1, P_2, \dots$ à cette pyramide. Soit $\frac{h}{n}$ la hauteur commune de ces prismes.

1) Calculer le volume des prismes p_1 puis p_2 , puis p_i . En déduire le volume v_n de la réunion des prismes intérieurs à SABC.

2) Montrer que $1^2 + 2^2 + \dots + n^2 = \frac{n(n+1)(2n+1)}{6}$.

3) Calculer le volume V_n des prismes extérieurs à la pyramide.

4) Expliquer pourquoi pour tout entier n, $v_n < V < V_n$. Quelle est la limite de la suite (v_n) et la limite de la suite (V_n) ? Que peut-on en déduire pour V?

Avant d'aborder le problème avec les élèves, il faut leur rappeler que, son objet étant de démontrer la formule $V = \frac{1}{3} Bh$, il ne faut pas utiliser cette formule dans le cours du problème. Les constructions de patrons et de solides seront peut-être faites à la maison plutôt qu'en classe. La lecture du texte, tout à fait compréhensible par les élèves, est très utile pour continuer le problème. Il est souhaitable pour certains élèves de mener cette lecture en classe.

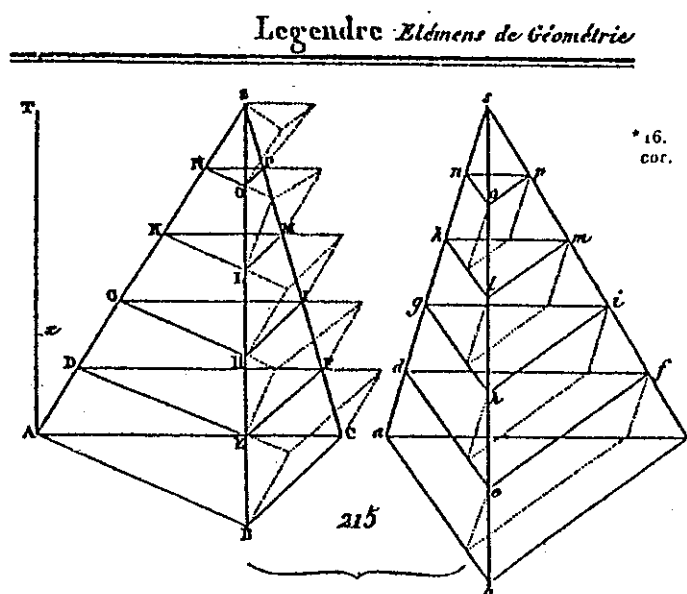
Deux pyramides triangulaires qui ont des bases équivalentes et des hauteurs égales, sont équivalentes.

Fig. 215. Soient $SABC$, $sabc$ les deux pyramides dont les bases ABC , abc , que nous supposons placées sur un même plan, sont équivalentes et qui ont même hauteur TA ; si ces pyramides ne sont pas équivalentes, soit $sabc$ la plus petite et soit Ax la hauteur d'un prisme qui étant construit sur la base ABC , serait égal à leur différence.

Divisez la hauteur commune AT en parties égales plus petites que Ax , et soit k une de ces parties; par les points de division de la hauteur, faites passer des plans parallèles au plan des bases; les sections faites par chacun de ces plans dans les deux pyramides, seront équivalentes*, telles que DEF et def , GHI et ghi , etc. Cela posé, sur les triangles ABC , DEF , GHI , etc., pris pour bases, construisez des prismes extérieurs qui aient pour arêtes les parties AD , DG , GK , etc., du côté SA ; de même sur les triangles def , ghi , kln , etc., pris pour bases, construisez dans la seconde pyramide des prismes intérieurs qui aient pour arêtes les parties correspondantes du côté sa ; tous ces prismes partiels auront pour hauteur commune k .

La somme des prismes extérieurs de la pyramide $SABC$ est plus grande que cette pyramide, la somme des prismes intérieurs de la pyramide $sabc$ est plus petite que cette pyramide; donc par ces deux raisons la différence entre les deux sommes de prismes devra être plus grande que la différence entre les deux pyramides.

Or à partir des bases ABC , abc , le second prisme extérieur $DEFG$ est équivalent au premier prisme intérieur $defa$, puisque leurs bases DEF , def , sont équivalentes et qu'ils ont une même hauteur k ; sont équivalents par la même raison le troisième prisme extérieur $GHIK$ et le second intérieur $ghid$, le quatrième extérieur et le troisième intérieur, ainsi de suite jusqu'au dernier des uns et des autres. Donc tous les prismes extérieurs de la pyramide $SABC$, à l'exception du premier $ABCD$, ont leurs équivalents dans les prismes intérieurs de la pyramide $sabc$. Donc le prisme $ABCD$ est la différence entre la somme des prismes extérieurs de la pyramide $SABC$ et la somme des prismes intérieurs de la pyramide $sabc$; mais la différence de ces deux sommes est plus grande que la différence des deux pyramides; donc il faudrait que le prisme $ABCD$ fût plus grand que le prisme $ABCX$; or au contraire il est plus petit, puisqu'ils ont une même base ABC , et que la hauteur k du premier est moindre que la hauteur Ax du second. Donc l'hypothèse d'où l'on est parti ne saurait avoir lieu; donc les deux pyramides $SABC$, $sabc$, de bases équivalentes et de hauteurs égales, sont équivalentes.



NOTE D'ECOUTE

A propos de l'Analyse non standard

Compte-rendu de la conférence de Pierre Cartier au séminaire de l'U.P.S.

Remarquant que parler en termes d'infinitésimaux, même seulement sur un mode métaphorique, pouvait rendre des services, P. Cartier a d'abord brossé à grands traits une histoire du traitement des infinitésimaux du 17^{ème} au 20^{ème} siècles, signalé leur abandon en mathématiques mais pas en physique au 19^{ème}, présenté l'apport d'Abraham Robinson et précisé ses propres conceptions sur le rôle que pourrait jouer l'analyse non standard.

P. Cartier a rappelé de façon très rapide et un peu schématique que Cavalieri et Pascal mesuraient une figure à l'aide de ses indivisibles - indivisibles sans épaisseur mais d'écartement constant - . L'aire d'une figure, désignée par l'expression "tous les indivisibles de la figure" est en quelque sorte la somme de ces indivisibles et notée par le symbole *omn*. (du latin *omnia*, tous les ...). Leibniz a précisé qu'il considérait une aire curviligne comme une somme de petits rectangles d'épaisseur infiniment petite, désignée par dx , et introduit le symbole $\int f(x)dx$, qui représente littéralement une somme infinie d'éléments infinitésimaux (P. Cartier a rappelé que le symbole de l'intégrale définie n'a été introduit que beaucoup plus tard, par Fourier). Leibniz traite l'infiniment petit dx presque comme un nombre ordinaire. Dans cette utilisation des infiniment petits il y a cependant une contradiction récurrente. Par exemple, en confondant l'aire curviligne avec la somme des aires de petits rectangles de largeur infiniment petite, Leibniz utilise des infiniment petits dans le cours du calcul, mais les néglige à la fin de ce calcul. Lazare Carnot a, selon P. Cartier, bien analysé les contradictions et les problèmes posés par les infiniment petits dans ses *Réflexions sur la métaphysique du calcul infinitésimal* : pour les classiques, le signe d'égalité recouvre deux significations différentes, l'égalité stricte et l'égalité approchée (à une erreur négligeable près). D'Alembert voudra tout ramener à la notion de limite qu'il n'élucidera cependant pas. On considère que Cauchy a introduit la rigueur dans l'édifice de l'analyse, mais son langage manipule encore des infiniment petits et des infiniment grands. Cauchy a élevé au rang de principe un mode de raisonnement utilisé avant lui, en énonçant le critère de convergence qui porte son nom, qu'il énonce en termes d'infiniment grands et de quantités négligeables. Weierstrass codifie ensuite la notion de limite en termes d'épsilon et d'alpha. Il y a permanence du sens mais glissement du langage; il ne s'agit plus, à

partir de Weierstrass d'entier n infiniment grands, mais d'écrire une condition du type $|x_n| < \varepsilon$ vérifiée pour tous les n supérieurs à un certain n_0 .

Pierre Cartier nous a parlé d'une expérience d'enseignement à un jeune homme quasi-autodidacte qui avait entrepris de façon tardive de passer un bac G. Pierre Cartier avait réussi à résumer le programme d'analyse en quelques règles concises et faciles à utiliser. Il avait proposé de compléter les nombres usuels par des nombres illimités et défini l'égalité approchée de deux nombres sous la forme

$$a \approx b \Leftrightarrow (a = b \text{ (égalité stricte) ou } a \neq b \text{ et } \frac{1}{a - b} \text{ illimité})$$

En fait il faut toute une hiérarchie d'éléments négligeables et de quasi-égalités. En effet, lorsqu'on écrit $\frac{d(x^2)}{dx} \approx 2x$, on néglige un infiniment petit du premier ordre et lorsqu'on écrit $dx^2 \approx 2x dx$, on néglige un infiniment petit du second ordre. Les difficultés apparaissent véritablement lorsqu'on veut définir par exemple l'intégrale, où intervient une somme infinie de termes infiniment petits et des infiniment petits d'ordres différents. Un garde-fou devient nécessaire pour éviter les contradictions.

Le conférencier a alors présenté Abraham Robinson (1918 - 1974), le fondateur de l'analyse non standard: chassé d'Allemagne vers 1933 par le nazisme, il se réfugie en Israël où il sera l'élève du grand logicien A. Fraenkel. Un séjour d'étude à Paris est vite interrompu par la guerre; il fuit en Angleterre, où il se spécialise en aérodynamique et participe par ses recherches à l'effort de guerre. Après 1945, il est ingénieur consultant chez Boeing où il contribue à l'invention de l'aile en delta et au développement de l'aviation supersonique. Les deux tiers de son oeuvre mathématique concerneront l'aérodynamique, mais il recommence également ses recherches en logique et il est l'un des fondateurs de la théorie des modèles. Au tout début des années soixante, il applique les résultats de cette théorie à l'invention de l'analyse non standard. S'appuyant sur un résultat de Skolem, selon lequel il y a plus d'un modèle possible pour l'arithmétique, Robinson définit un modèle N^* qui contient l'ensemble N des entiers usuels et dans lequel les propriétés vraies dans N restent vraies, mais qui contient également d'autres objets, les hyper-entiers, qui sont "au-delà" des entiers. Ces hyper-entiers ω sont plus grands que tous les entiers, et vérifient par exemple les énoncés $\omega^2 > \omega$ et $2\omega > \omega$ La propriété "il y a un entier supérieur à tous les entiers" est fausse aussi dans N^* . Les démonstrations de Robinson sont assises sur une théorie logique très solide, mais aride et coûteuse. Cette théorie donne un statut rigoureux aux infinitésimaux.

Après la mort de Robinson, une école américaine où figure en particulier E. Nelson, développe une variante de la théorie de Robinson, la théorie des ensembles internes. Un peu plus tard, Reeb et l'école française vont faire un effort de simplification et tenter de rendre l'outil plus accessible, dans le but de simplifier l'exposé de l'analyse et sa pédagogie.

Pierre Cartier tient que, au cours du 20ème siècle, l'analyse mathématique et l'analyse numérique se sont progressivement séparées en deux enseignements parallèles qui ne se recoupent pas. L'analyse mathématique souvent ne s'applique pas au concret. Un point de vue influencé par l'analyse non standard permettrait de lutter contre le divorce actuel entre ces deux domaines. Cartier donne l'exemple du théorème d'Hadamard, selon lequel le déterminant d'une matrice dont les éléments dépendent continûment d'un paramètre λ est lui-même une fonction continue de λ . Cependant, sur un exemple de matrice 20×20 , P. Cartier a pu observer des variations très brusques qui infirment toute tentative d'interpolation. La fonction étudiée, est "macroscopiquement" non-continue, même si elle est "microscopiquement" continue. P. Cartier en conclut que la notion de continuité n'est pas absolue, qu'il faudrait définir des notions de continuité dépendant de l'ordre de grandeur considéré. Pour lui, l'idée importante de la théorie des modèles introduisant la non unicité du modèle, met en question l'apparent excès d'unité des théories mathématiques régnantes, trop rigides et univoques. Les méthodes infinitésimales sont l'occasion de réfléchir à ces questions, d'introduire une relativisation du point de vue. Ajoutons que Pierre Cartier prépare un ouvrage sur le calcul intégral traité dans le cadre de l'analyse non standard, où il s'attache à alléger le contenu logique nécessaire.

Michèle Grégoire



Abraham Robinson

NOTE DE LECTURE

Comme nous l'avions annoncé dans le numéro 2 de Mnémosyne, voici, sous la plume de Maryvonne Spiesser, une analyse d'un premier ouvrage général d'histoire des mathématiques.

Histoire des Mathématiques

Jean-Paul COLLETTE

Tome 1 - Editions du Renouveau Pédagogique Inc. (ERPI) - Montréal, 1973 (228 pages)

Tome 2 - Vuibert. Paris - ERPI. Montréal, 1979. (359 pages)

Composition de l'ouvrage

C'est à la suite d'un cours d'histoire des mathématiques donné à l'université du Québec que J.P. Collette a nourri le projet de ce livre. Il s'agit avant tout pour l'auteur de *"présenter un manuel d'histoire des mathématiques plutôt qu'un traité, afin d'exposer surtout les notions historiques communément acceptées par les historiens et de faciliter d'autant la lecture de son contenu"* [préfaces des T1 et T2]. Cet ouvrage est destiné en particulier aux enseignants et aux étudiants désireux de "prendre contact" avec l'histoire des mathématiques.

L'auteur a fait le choix d'une progression chronologique. Le tome 1 couvre toute la période qui va de la préhistoire (premières activités de comptage chez l'homme primitif) à l'aube du XVII^{ème} siècle. Il est divisé en 11 chapitres, abordant les civilisations babylonienne et égyptienne, les mathématiques en Grèce, en Chine, en Inde, en Islam, puis en Europe au Moyen-Age et à la Renaissance, pour conclure avec les grands mathématiciens qui furent à la charnière des XVI^{ème} et XVII^{ème} siècle et ouvrirent la voie aux mathématiques nouvelles.

Le deuxième tome est également découpé en 11 chapitres regroupés en trois grandes parties correspondant respectivement aux XVII^{ème}, XVIII^{ème} et XIX^{ème} siècles. Après une présentation générale du siècle, J.P. Collette cite les principaux mathématiciens, et décrit leurs contributions. Quelques paragraphes sont consacrés aux grandes théories et au développement des différentes branches des mathématiques (exemples : théorie des ensembles, la résolubilité des équations, l'analyse vectorielle, ...).

A la fin de chaque tome, on trouve une liste des illustrations, un index des principaux auteurs (avec leurs dates) et des sujets traités (séparé en deux index distincts dans le tome 2).

Enfin, 32 sujets de recherche sont proposés dans le premier tome, accompagnés d'une courte bibliographie pour seule aide. Ce sont soit des recherches sur un thème épistémologique, soit, mais plus rarement, sur une œuvre ou un auteur. En voici deux exemples :

- 31. L'origine historique de l'Induction mathématique**¹
- [Tome 1 p. 216]
- Bussey W.H. 'The origin of Mathematical Induction,'
The American Mathematical Monthly, 24 (1917): 199-207.
- Cajori Florian. 'Analysis of Ueber das Wesen des Mathematik by A. Voss'.
Bulletin of the American Mathematical Society, 15 (1909) :405-409
- Cajori Florian. 'Fermat's Method on infinite Descent'
The American Mathematical Monthly, 2(1918): 333-337.
- Cajori Florian. 'Origin of the Name Mathematical induction'.
The American Mathematical Monthly, 25 (1918): 197-201.
- Hara Kōkichi. 'Pascal et l'induction mathématique'. *Revue d'histoire des sciences et leurs applications*, 15 (1962); 287-302.
- Rabinovitch, Nachum L. 'Rabbi Levi Ben Gershon and the Origins of Mathematical Induction'. *Archive for History of Exact Sciences*, 6 (1970); 237-248.
- Sierpinski W. 'L'induction incomplète dans la théorie des nombres'.
Scripta Mathematica, 28 (1962): 5-13.
- Vacca. 'Maurolycus, the First Discoverer of the Principle of Mathematical Induction'. *Bulletin on the American Mathematical Society*, 16 (1910) ;70-73.
- 23. L'"Arénaire" d'Archimède**
- Le but de cet ouvrage, son contenu mathématique et ses principales caractéristiques.
- [Tome 1 p. 213]
- Archimedes. 'The Sand Reckoner' in *The World of Mathematics*, J.R. Newman, ed. Vol 1, New York, Simon and Schuster, 1956, p. 420-429.
- Heath, Sir Thomas L., ed. *The Works of Archimedes with the Method of Archimedes*. New York, Dover, 1953, p. 221-232.
- Miller, Leland. 'Large Numbers' in *Historical Topics for Mathematics Classroom*, The National Council of Teachers of Mathematics, ed. 31st Yearbook, Washington, D.C., N.C.T.M., 1969, p. 82-83
- Read, Cecil B. 'Archimedes and his Sandreckoner'.
School Science and Mathematics, 61 (1961): 81-84.

Composition des chapitres

Chaque chapitre comporte un résumé qui rassemble les points importants et les principales idées traitées. Dans le premier tome, il s'intitule "En bref" et vient en conclusion ; dans le tome suivant, c'est en introduction.

On trouve également à la fin du chapitre :

- une bibliographie²

¹ou récurrence,

²dépassée maintenant, puisqu'elle date de 1979 au plus tard. Une grande majorité des ouvrages cités est en langue anglaise.

- des exercices (non corrigés !) proposant, soit une réflexion sur les idées maîtresses qui ont été développées, soit une démonstration ou une application d'un résultat. Voici, par exemple, quelques-uns des exercices proposés à la suite du chapitre 5 du tome I, intitulé "de Platon à Euclide" :

ex. 1 : *Quelles ont été les contributions mathématiques de Platon et son rôle dans l'évolution des mathématiques grecques ?*

ex. 7 : *Trouver les 5ème, 6ème et 7ème nombres parfaits par la méthode d'Euclide.*

ex 10 : *Platon, dans son Théétète, affirme que Théodore a prouvé l'irrationalité de $\sqrt{3}$. Pouvez-vous en établir une preuve rigoureuse ?* [Tome 1, p. 81-82]

Tous les chapitres sont divisés en paragraphes se rapportant parfois à un thème particulier [ex : le développement de la trigonométrie pendant la Renaissance, (T.1, ch. 10)], mais le plus souvent, à un auteur et à son oeuvre³. La présentation de l'ouvrage est claire et la lecture aisée. Les titres des paragraphes, les noms des mathématiciens, sont bien mis en évidence. Comme nous l'avons précisé plus haut, ce manuel est à visée pédagogique. Destiné à permettre une première approche de l'histoire des mathématiques, c'est en quelque sorte une étape vers une étude plus approfondie d'un auteur ou d'un sujet. Le découpage chronologique permet effectivement un accès facile à l'oeuvre d'un mathématicien.

J.P. Collette s'est fortement inspiré de l'ouvrage de l'américain Carl B. Boyer, *A History of Mathematics*, publié en 1968. Que ce soit dans les objectifs, ou dans la conception pratique du manuel, la ressemblance est grande.

Présenter en 500 pages une histoire chronologique des mathématiques qui, depuis l'antiquité grecque seulement, court sur plus de 25 siècles, ne permet pas de s'attarder beaucoup sur chaque sujet. Collette fait un simple panorama historique. Les auteurs sont brièvement replacés dans leur époque, la description des oeuvres est appuyée par quelques courtes citations. On ne trouvera pas ici de réflexion épistémologique, d'analyse de l'évolution d'un concept mathématique. On ne trouvera pas non plus, ou presque pas, de textes originaux conséquents, et cela nous expose à tous les dangers d'une lecture indirecte déjà interprétée (que ce soit d'ailleurs par Collette ou non : il précise bien qu'il va exposer des "*notions historiques communément acceptées par les historiens*").

En conclusion, l'Histoire des Mathématiques de Collette est un ouvrage bien présenté, pratique à consulter. Il n'est pas un outil permettant de suivre l'évolution d'une idée, d'une théorie ou d'une branche des mathématiques, mais il représente une source d'information intéressante pour un premier accès à un auteur ou à une période de l'histoire.

Maryvonne Spiesser

IREM de Toulouse

³ se reporter à l'ouvrage T1, ch. 4 p. 51 et T.2 ch. 2 p. 43 par exemple.

Comité de rédaction:

<i>Philippe BRIN</i>	<i>Lycée Technique E. Branly Créteil Animateur à l'IREM Paris VII</i>
<i>Michèle GREGOIRE</i>	<i>Lycée Lavoisier Paris Animateur à l'IREM Paris VII</i>
<i>Maryvonne HALLEZ</i>	<i>Collège P. Bert Paris Animateur à l'IREM Paris VII</i>
<i>Marie Françoise JOZEAU</i>	<i>Lycée G. de Nerval Luzarches Animateur à l'IREM Paris VII</i>
<i>Michèle LACOMBE</i>	<i>Lycée J. Monod Clamart Animateur à l'IREM Paris VII</i>
<i>Anne MICHEL-PAJUS</i>	<i>Lycée C. Bernard Paris Animateur à l'IREM Paris VII</i>
<i>Michel SERFATI</i>	<i>Lycée Technique Raspail Paris Animateur à l'IREM Paris VII</i>
<i>Jean Luc VERLEY</i>	<i>Université Paris VII IREM Paris VII</i>

*avec la participation de Gérard HAMON
IREM de RENNES*

*et celle de Maryvonne SPIESSER
IREM de TOULOUSE*

<u>Courrier à adresser à :</u> Groupe M: A.T.H. IREM de l'université Paris VII Tour 55-56 3ème étage 75005 PARIS
--

*Pour échanger expériences et réflexion à propos de l'histoire et de
l'enseignement des mathématiques*

M. *Mathématiques*

A. *Approche par les*

T. *textes*

H. *historiques*

SOMMAIRE

Editorial

Bonnes vieilles pages

Wantzel

*' Recherche sur les moyens de reconnaître si un problème
de géométrie peut se résoudre avec la règle et le compas.'*

Iconographies

Conte du Lundi

Euclide a encore quelque chose à dire

Etude

Fragments d'une étude des systèmes linéaires

Dans nos classes

A propos du volume de la pyramide

Note d'écoute

Note de lecture

Calendrier

Editeur : IREM

Directeur responsable de la publication : R. DOUADY

Dépôt légal : Avril 1993

ISBN : 2-86612-085-X

IREM Université Paris VII Denis Diderot
Tour 56-55 - 3^{ème} étage, Case 7018
2, place Jussieu 75 251 Paris Cedex 05
Tel : 01 44 27 53 83