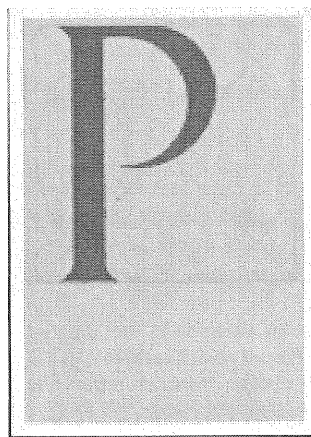


la géométrie grecque, sauvegardée pendant le Moyen-Age grâce aux Arabes et se répandant en Italie à la Renaissance. Il transmet cette connaissance millénaire à Dürer qui à son tour la diffusera au Nord de Alpes ! Un moment historique dont nous sommes témoins en regardant cette peinture !

Références

- [1] N. BOONEN, *Platonische en Archimedische lichamen*, mémoire non-publié, KHLim Lerarenopleiding (Hasselt), 1999.
- [2] P.R. CROMWELL, *Polyhedra*, Cambridge University Press (Cambridge, New York, Melbourne), 1997, ISBN 0-521-66405-5.
- [3] EUCLID, *The thirteen books of the Elements*, Translated with introduction and commentary by Sir Thomas L. Heath, Volume III, Dover Publications (New York), 1956, ISBN 0-486-60090-4.
- [4] J.V. FIELD, *The invention of infinity. Mathematics and Art in the Renaissance*, Oxford University Press (Oxford, New York, Tokyo), 1997, ISBN 0-19-852394-7.
- [5] N. MACKINNON, *The portrait of Fra Luca Pacioli*, *The Mathematical Gazette* 77/479 (1993), p. 129-219.



Une histoire des approximations successives : des équations numériques aux équations fonctionnelles

TOURNES Dominique¹
IUFM de La Réunion (France)

Abstract

Faire un essai, observer attentivement ce qu'il advient et corriger sa stratégie en conséquence : voici sans doute la démarche la plus naturelle qui soit pour résoudre un problème. Tout apprentissage n'est-il pas le fruit d'une suite consciemment organisée d'"essais-erreurs" ? En mathématiques, la *méthode des approximations successives* apparaît comme une formalisation de ce schéma général d'action. Une équation de nature quelconque peut toujours s'écrire sous la forme $x = \varphi(x)$, où x est un élément inconnu appartenant à un ensemble donné et où φ est un opérateur agissant sur cet ensemble. La valeur de x semble irrémédiablement hors d'atteinte puisqu'elle s'exprime en fonction d'elle-même. Comment sortir de ce cercle vicieux ? Imaginons qu'un certain raisonnement ait conduit à considérer comme intéressante une valeur approchée x_1 de x . Il peut être tentant d'appliquer à x_1 l'opérateur φ de manière à obtenir la valeur corrigée $x_2 = \varphi(x_1)$. Le miracle est de constater que, souvent, x_2 est une valeur plus précise que x_1 . Il paraît alors presque naturel de recommencer ce calcul auto-correctif une fois, deux fois, etc., puis, si l'on ose franchir l'obstacle conceptuel, une infinité de fois, afin de parvenir à une solution qui n'a plus besoin d'être corrigée, c'est-à-dire à la solution exacte. Face à une telle conclusion, Gaston Bachelard, qui a étudié la méthode itérative du point de vue épistémologique, n'a pas caché son étonnement : "ainsi, parti de n'importe où, en substituant n'importe quoi, par le seul jeu du système lui-même on arrive à sa solution"². Il reste toutefois à préciser la nature de ce qu'on entend par "solution" : s'agit-il d'une solution donnée *a priori* dont on construit des valeurs approchées aussi précises que l'on veut, ou bien est-ce le processus itératif lui-même qui définit un objet de type nouveau qu'on peut regarder *a posteriori* comme étant une solution ?

La méthode des approximations successives a été pratiquée sous des formes variées depuis deux mille ans, d'abord pour les équations numériques, puis pour des équations plus générales dans d'autres ensembles que les ensembles de nombres. Pendant l'atelier de l'université d'été, nous avons tenté de faire revivre cette longue histoire à travers la lecture de textes originaux. En raison de la durée limitée dont nous disposons, il a fallu se restreindre à la présentation d'un petit nombre de textes, choisis parmi les plus significatifs. Le caractère réducteur d'une telle démarche sera atténué, dans le compte rendu qui suit, par l'insertion de nombreuses références bibliographiques permettant au lecteur intéressé de donner par lui-même davantage d'épaisseur à notre étude³.

¹IUFM de La Réunion, allée des Aigues Marines, Bellepierre, 97487 Saint-Denis Cedex, et REHSEIS-CNRS, 37 rue Jacob, 75006 Paris. Courrier électronique : tournes@univ-reunion.fr.

²BACHELARD, 1973, p. 215.

³Afin de ne pas trop alourdir cette bibliographie, nous ne donnerons pas de références primaires, sauf pour les

1 Itérer les méthodes de fausse position

Dans une première partie, nous allons examiner les tentatives faites avant Newton pour résoudre itérativement les équations numériques¹. Ces tentatives, relativement nombreuses, se rencontrent régulièrement dans les mathématiques grecques, indiennes, arabes et occidentales, et peuvent être placées dans la filiation directe des anciennes méthodes de fausse position.

a) Schéma de l'évolution de l'analyse numérique des équations $f(x) = 0$

La première étape est celle des *équations linéaires*, qui peuvent se mettre sous la forme réduite $ax = b$ (ici, $f(x) = ax - b$). Dans les écrits mathématiques de toutes les civilisations, on rencontre très tôt de nombreux problèmes dont la formalisation moderne conduirait à ce type d'équation. Ces problèmes font l'objet d'un traitement arithmétique consistant à "tester" d'abord une ou deux valeurs numériques particulières de la quantité inconnue, et à les "corriger" ensuite par un raisonnement de proportionnalité pour en déduire la valeur exacte cherchée.

Dans la méthode de *simple fausse position*, on essaye une seule valeur x_1 , conduisant au résultat b_1 . Des égalités $ax = b$ et $ax_1 = b_1$, on tire $\frac{x}{x_1} = \frac{b}{b_1}$, soit $x = \frac{bx_1}{b_1}$. En revanche, dans la méthode de *double fausse position*, on s'appuie sur deux essais x_1 et x_2 , ce qui mène, en désignant par e_1 et e_2 les erreurs commises, aux égalités $ax_1 = b + e_1$ et $ax_2 = b + e_2$. Le plus souvent, on s'arrange pour que l'une des erreurs soit positive et l'autre négative, d'où les expressions parfois rencontrées de "règle du trop et du pas assez", "règle de l'excédent et du déficit", "règle de l'augmentation et de la diminution", etc. En multipliant la première des égalités précédentes par e_2 et la seconde par e_1 , puis en faisant leur différence, on obtient $a(e_2x_1 - e_1x_2) = b(e_2 - e_1)$. En combinant avec $ax = b$, il vient $\frac{x}{e_2x_1 - e_1x_2} = \frac{1}{e_2 - e_1}$, soit enfin $x = \frac{e_2x_1 - e_1x_2}{e_2 - e_1}$.

Il faut bien sûr lire cette présentation des méthodes de fausse position avec les précautions qui s'imposent. La forme algébrique moderne ne reflète en aucun cas les processus mentaux sous-jacents aux textes anciens qui nous sont parvenus; elle masque notamment les difficultés liées, d'une part à l'absence de notations littérales, d'autre part aux contraintes numériques (contraintes variables en fonction du système de numération et des techniques opératoires en usage en un lieu et un temps donnés)².

On peut analyser les méthodes de fausse position d'une autre manière, en faisant appel à la notation fonctionnelle $f(x) = ax - b$ et à une interprétation géométrique (voir Figure 1).

textes étudiés pendant l'atelier. Le lecteur trouvera ces références primaires dans les bibliographies des références secondaires auxquelles nous le renvoyons.

¹Dans le cas des équations numériques, l'histoire de la méthode des approximations successives a donné lieu à une littérature considérable, y compris ces dernières années, ce qui montre que le sujet est loin d'être épuisé. Pour une première vue d'ensemble, se reporter à CHABERT *et al.*, 1994 ou à BAILEY, 1989.

²Pour une histoire détaillée des méthodes de fausse position à travers des textes babyloniens, égyptiens, chinois, indiens, arabes et européens, on pourra se reporter à CHABERT *et al.*, 1994, chap. 3 et à SPIESSER, 1982. Une référence complémentaire intéressante est CHEMLA, 1997; on y découvre comment la méthode de double fausse position, probablement née en Chine vers le début de notre ère, a pu être transmise d'Orient en Occident.

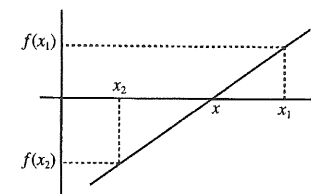


FIGURE 1

Revenons d'abord à la méthode de simple fausse position et à son unique essai x_1 . De l'égalité des rapports $\frac{f(x_1)-0}{x_1-x} = a$, on tire $x = x_1 - \frac{f(x_1)}{a}$. Pour anticiper sur la suite, nous nous permettons d'écrire le résultat sous la forme

$$x = x_1 - \frac{f(x_1)}{f'(x_1)}. \quad (1)$$

De façon analogue, en ce qui concerne la méthode de double fausse position et ses deux essais x_1 et x_2 , l'égalité des rapports $\frac{f(x_1)-f(x_2)}{x_1-x_2} = \frac{0-f(x_2)}{x-x_2}$ conduit à

$$x = x_2 - \frac{f(x_2)}{\frac{f(x_2)-f(x_1)}{x_2-x_1}}. \quad (2)$$

En adoptant les écritures (1) et (2), nous avons seulement voulu montrer que les anciennes techniques arithmétiques de fausse position pouvaient être placées à la source des idées développées ensuite pour la résolution numérique des équations quelconques. En effet, dans un second temps, il est tentant de continuer à appliquer, face à une équation non linéaire $f(x) = 0$, les algorithmes mis au point pour les équations linéaires. C'est le principe même de l'*interpolation linéaire*³ : sur un petit intervalle, on peut faire comme si une grandeur quelconque variait linéairement sans que cela entraîne d'erreur grave.

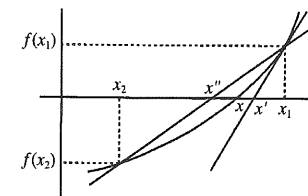


FIGURE 2

³Cette idée se rencontre très tôt chez les astronomes babyloniens, qui dressent des éphémérides des corps célestes et ont ensuite à interpoler entre deux valeurs consécutives de leurs tables. Voir NEUGEBAUER, 1957 et surtout NEUGEBAUER, 1975. Une fois décrites en termes modernes, la méthode d'interpolation linéaire entre deux valeurs d'une table (une valeur par excès et une valeur par défaut) et la méthode de double fausse position (règle de l'excédent et du déficit) semblent voisines, voire identiques. Pourtant, elles sont probablement le fruit de deux traditions historiques longtemps indépendantes, la première se développant pour les besoins des calculs astronomiques, et la seconde en réponse aux problèmes concrets des comptables ou des marchands.

Si l'on connaît seulement un point $(x_1, f(x_1))$, on peut remplacer la courbe par sa tangente en ce point (voir Figure 2); la valeur exacte x de la racine est alors approchée par

$$x \approx x' = x_1 - \frac{f(x_1)}{f'(x_1)}. \quad (1')$$

De façon analogue, si l'on connaît deux points $(x_1, f(x_1))$ et $(x_2, f(x_2))$, on peut remplacer la courbe par la sécante joignant ces deux points, d'où la seconde approximation

$$x \approx x'' = x_2 - \frac{f(x_2) - f(x_1)}{f'(x_2) - f'(x_1)}. \quad (2')$$

Ainsi, par ces processus de correction de la (ou des) valeur(s) fausse(s) dont on était parti, on obtient une nouvelle valeur approchée meilleure –en général– que la (ou les) précédente(s). Ces méthodes de fausse position généralisée (au sens où elles ne sont plus appliquées à des problèmes du premier degré) apparaissent, sous une forme ou sous une autre, dans de nombreux textes de diverses civilisations et de diverses époques, mais, presque toujours, on s'y contente d'une seule correction. Ce qui est plus rare –et c'est ce qui va nous intéresser dans cet article–, ce sont les textes dans lesquels on trouve l'idée d'*itérer la correction*, autrement dit le fondement d'un algorithme permettant d'obtenir des valeurs approchées aussi précises que l'on veut.

L'itération de la formule de correction (1') conduit à la *méthode de Newton-Raphson*, méthode à convergence rapide d'ordre 2, définie par

$$\begin{cases} x_1 \\ x_{n+1} \end{cases} = \begin{cases} x_1 \\ x_n - \frac{f(x_n)}{f'(x_n)} \end{cases}.$$

La formule de correction (2'), quant à elle, peut s'itérer de deux façons. Selon que l'on adopte une récurrence simple ou double, on obtient la *méthode de Lagrange* (d'ordre 1) ou la *méthode de la sécante* (d'ordre le nombre d'or $(1 + \sqrt{5})/2$), définies respectivement par

$$\begin{cases} x_1, x_2 \\ x_{n+1} \end{cases} = \begin{cases} x_1, x_2 \\ x_n - \frac{f(x_n)}{f'(x_n) - f'(x_1)} \end{cases}, \quad \begin{cases} x_1, x_2 \\ x_{n+1} \end{cases} = \begin{cases} x_1, x_2 \\ x_n - \frac{f(x_n)}{f'(x_n) - f'(x_{n-1})} \end{cases}.$$

Outre ces trois schémas issus des anciennes techniques de fausse position ou d'interpolation, il peut apparaître des schémas plus généraux nés de la transformation, par un procédé quelconque, de l'équation à résoudre $f(x) = 0$ en une équation à point fixe $g(x) = x$. Le processus d'approximation est alors décrit par la suite récurrente

$$\begin{cases} x_1 \\ x_{n+1} \end{cases} = \begin{cases} x_1 \\ g(x_n) \end{cases}.$$

Nous allons maintenant analyser quelques textes illustrant l'émergence historique des techniques précédentes. Ces textes sont empruntés aux mathématiques grecques, arabes et indiennes.

b) Texte 1 : Héron d'Alexandrie, I^{er} siècle

Héron vivait à Alexandrie, au I^{er} siècle de notre ère. Ses principaux écrits, retrouvés à Constantinople en 1896, s'intitulent *Les Métriques* et *La Dioptrique*. Ce sont des manuels de

mathématiques destinés à des mécaniciens et des ingénieurs. L'extrait suivant des *Métriques*⁴, qui présente le calcul approché de $\sqrt{720}$, est incontournable : c'est, semble-t-il, le premier texte connu dans lequel on rencontre le schéma des approximations successives.

Puisque 720 n'a pas de côté rationnel, nous extrairons le côté avec une très petite différence de la façon suivante. Comme le premier nombre carré plus grand que 720 est 729 qui a pour côté 27, divise 720 par 27; cela fait $26\frac{2}{3}$; ajoute 27, cela fait $53\frac{2}{3}$; prends-en la moitié, cela fait $26\frac{1}{3}$. Ainsi donc le côté de 720 sera très proche de $26\frac{1}{3}$. En fait, $26\frac{1}{3}$ multiplié par lui-même donne $720\frac{1}{36}$, de sorte que la différence est $\frac{1}{36}$. Si nous voulons rendre cette différence inférieure encore à $\frac{1}{36}$, nous mettrons $720\frac{1}{36}$ trouvé tout à l'heure à la place de 729 et, en procédant de la même façon, nous trouverons que la différence est beaucoup plus petite que $\frac{1}{36}$.

Les calculs de Héron peuvent être résumés ainsi :

$$\begin{aligned} \frac{1}{2} \left(\frac{720}{27} + 27 \right) &= 26\frac{1}{3} \approx 26,83333333; \\ \frac{1}{2} \left(\frac{720}{26\frac{1}{3}} + 26\frac{1}{3} \right) &= 26\frac{1609}{1932} \approx 26,83281574. \end{aligned}$$

Sachant que $\sqrt{720} \approx 26,83281573$, on constate que la méthode est excellente. Bien que Héron s'en tienne à une seule itération, sans doute parce que la précision obtenue est d'ores et déjà suffisante pour les besoins de la pratique, le texte contient implicitement l'idée que l'on pourrait continuer, avec, à chaque étape, une erreur beaucoup plus petite que précédemment. Il serait abusif d'y voir un algorithme infini au sens moderne du terme mais il s'agit malgré tout d'un processus "potentiellement infini" en ce sens qu'il peut être appliqué un nombre fini de fois aussi grand que nécessaire, jusqu'à ce que l'on obtienne la précision souhaitée. L'algorithme s'écrit

$$\begin{cases} x_1 = 27 \\ x_{n+1} = \frac{1}{2} \left(\frac{720}{x_n} + x_n \right) \end{cases} = x_n - \frac{x_n^2 - 720}{2x_n}.$$

Si l'on considère que l'équation à résoudre est $f(x) = 0$, avec $f(x) = x^2 - 720$, on observe que $x_{n+1} = x_n - f(x_n)/f'(x_n)$. Le procédé de Héron apparaît donc à nos yeux comme un cas particulier de la méthode de Newton mais, bien entendu, Héron ne connaissait pas la notion de dérivée ! Nombreuses sont les hypothèses sur le raisonnement qui a pu conduire à la mise au point d'un tel algorithme⁵. Notons a une valeur approchée par excès de $\sqrt{A}$, et désignons par e l'erreur commise. Si a est très proche de $\sqrt{A}$, e est petit, d'où $A = (a - e)^2 = a^2 - 2ae + e^2 \approx a^2 - 2ae$. On obtient alors $e = \frac{a^2 - A}{2a}$, soit $\sqrt{A} \approx a - \frac{a^2 - A}{2a} = \frac{1}{2} \left(\frac{A}{a} + a \right)$, et on retrouve la formule du texte. Ce raisonnement par linéarisation, équivalent à la méthode moderne de Newton, est des plus séduisants mais on ne peut pas savoir s'il traduit correctement la pensée de Héron. Connaissant le goût des Grecs pour la théorie des "médietés", voici une autre possibilité : a est une valeur approchée par excès de $\sqrt{A}$, donc $\frac{A}{a}$ en est une valeur approchée par défaut, et $\sqrt{A}$ est la moyenne géométrique de ces deux valeurs ($\sqrt{A} = \sqrt{\frac{A}{a} \times a}$); il est alors tentant de prendre comme nouvelle approximation la moyenne arithmétique $\frac{1}{2} \left(\frac{A}{a} + a \right)$.

⁴Texte grec dans HÉRON D'ALEXANDRIE, 1903, Livre I des *Métriques*, p. 18-20. Nous avons repris la traduction française de J.-P. Levet (CHABERT *et al.*, 1994, p. 231).

⁵On trouvera une analyse fine de la méthode de Héron dans CAVEING, 1998, chap. 1. Caveing soutient que cette méthode était connue des Babyloniens et des Grecs bien avant le premier siècle, et qu'elle a pu jouer un rôle dans la découverte de l'irrationalité.

La méthode de Héron se retrouve ensuite régulièrement, sous diverses formes, dans des textes alexandrins, arabes, byzantins et européens. Chez les mathématiciens arabes⁶, qui ont approfondi les algorithmes numériques grecs et indiens, on rencontre, à côté de l'approximation de Héron $\sqrt{A} \approx a + \frac{A-a^2}{2a}$, une autre approximation, qualifiée de "conventionnelle" : $\sqrt{A} \approx a + \frac{A-a^2}{2a+1}$. Cette dernière formule correspond à une interpolation linéaire entre deux entiers consécutifs a et $a+1$, ainsi qu'on le voit sur l'égalité $a + \frac{A-a^2}{2a+1} = a - \frac{a^2-A}{(a+1)-a}$. Au XII^{ème} siècle, al-Samaw'al, dans son *Traité d'arithmétique*, généralise le résultat au calcul d'une racine n^e en proposant d'itérer la formule $\sqrt[n]{A} \approx a + \frac{A-a^n}{\sum_{k=0}^{n-1} C_n^k a^k}$, formule que nous pouvons aussi écrire sous la forme $\sqrt[n]{A} = a - \frac{a^n-A}{(a+1)^n-a^n}$. Chez al-Samaw'al, l'idée des approximations successives est extrêmement claire, puisqu'il écrit : "Nous pouvons ainsi obtenir des quantités rationnelles dont le nombre est infini et dont chacune est plus proche que celle qui la précède de la quantité que l'on cherche à approcher".

c) Texte 2 : Paramesvara, XV^{ème} siècle

Faisons maintenant une incursion dans l'astronomie indienne. Comme les Arabes, les Indiens ont hérité des techniques d'interpolation de l'astronomie babylonienne et grecque (nombre de leurs travaux se situent dans la lignée de l'*Almageste* de Ptolémée). Dans les traités astronomiques indiens, il apparaît dès le V^{ème} siècle des techniques itératives pour résoudre des équations, avec des calculs reposant souvent sur des procédés de simple ou double fausse position⁷. Le texte suivant⁸ est dû à Paramesvara (environ 1380-1460), un élève de Mādhava, astronome renommé de Kerala. C'est un extrait du *Siddhāntadīpikā*, commentaire d'un commentaire (sic) du *Mahābhāskariya* de Bhāskara I. L'objectif est de calculer le sinus d'un angle donné.

Le sinus peut être obtenu à partir de l'arc au moyen d'une méthode dite "sans différence".

On doit prendre la moitié de l'arc donné, puis la moitié du résultat, et ainsi de suite. Il faut continuer à diviser par deux jusqu'à ce que le résultat soit plus petit que 1/300 du quadrant. Alors la corde de l'arc divisé est égale à cet arc.

En raison de la petitesse du sinus-verse, le sinus aussi est presque égal à cet arc. Le cosinus est calculé au moyen du sinus, et le sinus-verse à partir du cosinus. De proche en proche, la racine carrée du carré de la corde diminué du carré du sinus-verse est le sinus corrigé et, à partir de là, le cosinus est calculé puis, à partir du cosinus, le sinus-verse.

La racine carrée de la différence des carrés de la corde et du sinus-verse est le sinus corrigé; et, en ayant fait ceci sans cesse, on peut rendre le résultat indiscernable.

Le sinus itéré, multiplié par deux, est la corde précise du double de cet arc. Et le sinus est obtenu ainsi :

Si on suppose que le sinus est la corde diminuée de quelque chose (ou si on choisit que le sinus est la corde diminuée par la différence entre la corde et son arc, multipliée par 30 et divisée par 11), le

⁶Se reporter à RASHED, 1984, p. 115-120 ou à RASHED, 1997, p. 61-67.

⁷Pour des approximations successives analogues dans l'astronomie occidentale, en particulier au XIV^{ème} siècle, on pourra consulter MANCHA, 1998. Dans GLUSHKOV, 1976 est défendue la thèse selon laquelle Leonardo Fibonacci aurait lui aussi conduit certains calculs en itérant la méthode de double fausse position.

⁸Ce texte est étudié dans PLOFKE, 1996, où on le trouvera en sanskrit. La traduction française est de nous, à partir de la traduction anglaise de Plofker.

cosinus et le sinus-verse sont obtenus à partir de là. Et le sinus est la racine carrée de la différence des carrés de la corde et du sinus-verse. On doit corriger ceci itérativement. Ce sinus, multiplié par deux, est la corde du double de l'arc.

De cette façon on peut obtenir la corde de l'arc successivement doublé et son sinus corrigé itérativement, jusqu'à ce qu'il en résulte le sinus désiré de l'arc donné. De la même façon, le cosinus se déduit de l'arc complémentaire de la même manière.

Cette correction itérative entraîne souvent beaucoup de répétitions. On explique donc une certaine méthode pour accélérer la correction :

La racine carrée de la somme des carrés du sinus et du sinus-verse est la corde. On peut déduire la corde de chaque sinus et sinus-verse.

À chaque étape, la différence entre les deux cordes précédemment obtenues est le "diviseur", et le "multiplicateur" est la différence entre les deux précédents sinus-verses.

La différence entre la corde donnée et la corde produite à l'étape précédente est le "multiplicande". Le résultat obtenu est successivement ajouté ou retranché du sinus-verse précédemment obtenu. La règle est la suivante : quand la corde résultant du sinus donné est plus grande que la corde précédemment obtenue, elle est ajoutée; sinon, retranchée. Ceci donne la correction du sinus-verse, et le sinus en est déduit comme auparavant.

Quand on a calculé la corde et le sinus-verse et ainsi de suite par la méthode décrite, on doit corriger itérativement. Alors la correction est rapidement réalisée.

Introduisons quelques notations. Les lignes trigonométriques indiennes correspondent à un cercle de rayon R donné; nous les noterons avec une majuscule pour les distinguer des fonctions modernes associées à un cercle de rayon unité. Les fonctions sinus, cosinus, sinus-verse et corde, dont il est question dans le texte, seront donc définies respectivement (voir Figure 3) par $\text{Sin } \theta = R \sin \theta$, $\text{Cos } \theta = R \cos \theta$, $\text{Vers } \theta = R - \text{Cos } \theta = R(1 - \cos \theta)$ et $\text{Cord } \theta = \sqrt{\text{Sin}^2 \theta + \text{Vers}^2 \theta} = 2 \text{Sin } \frac{\theta}{2}$.

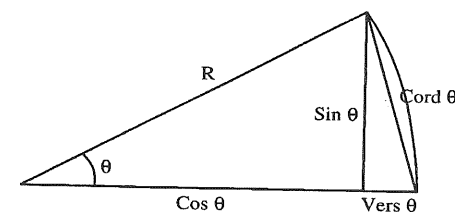


FIGURE 3

Soit à calculer $\text{Sin } \theta$ pour un angle θ donné. On détermine d'abord un entier p tel que $\frac{\theta}{2^p} \leq \frac{90^\circ}{300} = 18'$. On peut considérer que la valeur de $\text{Cord } \frac{\theta}{2^p}$ est connue puisque, pour un angle aussi petit, la corde s'identifie à l'arc (on vérifie effectivement que $\frac{\pi}{600}$ et $2 \sin(\frac{\pi}{1200})$ coïncident jusqu'à la huitième décimale). On calcule alors $\text{Sin } \frac{\theta}{2^p}$ à partir de $\text{Cord } \frac{\theta}{2^p}$ par un processus itératif que nous décrirons plus loin. Une fois calculée la valeur de $\text{Sin } \frac{\theta}{2^p}$, la relation $\text{Cord } \frac{\theta}{2^{p-1}} = 2 \text{Sin } \frac{\theta}{2^p}$ permet d'obtenir la valeur de $\text{Cord } \frac{\theta}{2^{p-1}}$. On applique à nouveau le

processus itératif pour déterminer $\text{Sin } \frac{\theta}{2^{p-1}}$ à partir de $\text{Cord } \frac{\theta}{2^{p-1}}$, et on recommence : de proche en proche, on obtient les valeurs de $\text{Sin } \frac{\theta}{2^{p-2}}, \dots, \text{Sin } \frac{\theta}{2}$ et $\text{Sin } \theta$.

À chaque étape, le problème est de déterminer le sinus d'un angle donné α à partir de sa corde. Cela se fait par le schéma de calcul $\text{Cos } \alpha = \sqrt{R^2 - \text{Sin}^2 \alpha}$, $\text{Vers } \alpha = R - \text{Cos } \alpha$ et $\text{Sin } \alpha = \sqrt{\text{Cord}^2 \alpha - \text{Vers}^2 \alpha}$, autrement dit par l'équation à point fixe

$$\text{Sin } \alpha = \sqrt{\text{Cord}^2 \alpha - \left(R - \sqrt{R^2 - \text{Sin}^2 \alpha}\right)^2}.$$

Notons C la corde connue de α . Si l'on tient compte de la valeur initiale proposée par le texte⁹, la valeur de $\text{Sin } \alpha$ est finalement calculée à l'aide des approximations successives

$$\begin{cases} u_1 &= C - \frac{30}{11}(R\alpha - C) \\ u_{n+1} &= \sqrt{C^2 - \left(R - \sqrt{R^2 - u_n^2}\right)^2}. \end{cases}$$

Il est aisé de vérifier que la méthode converge pour un angle α inférieur à 60° . Lorsque l'angle s'approche de 60° , la convergence est de plus en plus lente, ce qui amène notre astronome indien à proposer une méthode d'accélération de convergence. La méthode accélérée porte cette fois sur le calcul de $\text{Vers } \alpha$ (ce qui équivaut au calcul de $\text{Sin } \alpha$ puisque $\text{Cord } \alpha = \sqrt{\text{Sin}^2 \alpha + \text{Vers}^2 \alpha}$). Si on note Vers_n et Cord_n les valeurs du sinus-verse et de la corde calculées à chaque étape, l'algorithme présenté par le texte peut s'écrire¹⁰

$$\text{Vers}_{n+1} = \text{Vers}_n + \frac{(C - \text{Cord}_n)(\text{Vers}_n - \text{Vers}_{n-1})}{\text{Cord}_n - \text{Cord}_{n-1}} = \text{Vers}_n - \frac{(C - \text{Cord}_n)}{\frac{(C - \text{Cord}_n) - (C - \text{Cord}_{n-1})}{\text{Vers}_n - \text{Vers}_{n-1}}}.$$

Remarquons que

$$C - \text{Cord } \alpha = C - \sqrt{\text{Sin}^2 \alpha + \text{Vers}^2 \alpha} = C - \sqrt{R^2 - (R - \text{Vers } \alpha)^2 + \text{Vers}^2 \alpha} = C - \sqrt{2R\text{Vers } \alpha}.$$

En introduisant la fonction $f(x) = C - \sqrt{2Rx}$ et en réutilisant les deux valeurs initiales issues du premier algorithme, l'algorithme accéléré s'écrit donc

$$\begin{cases} v_1 &= \sqrt{C^2 - u_1^2}, v_2 = \sqrt{C^2 - u_2^2} \\ v_{n+1} &= v_n - \frac{f(v_n)}{\frac{f(v_n) - f(v_{n-1})}{v_n - v_{n-1}}}. \end{cases}$$

On reconnaît exactement la méthode de la sécante, dont la convergence est effectivement plus rapide que la méthode ordinaire des approximations successives. En guise d'illustration, nous avons testé les deux algorithmes de Paramēśvara pour un angle de 45° . À la précision 10^{-6} , les valeurs correctes sont $\text{sin } 45^\circ \approx 0,707107$ et $\text{Vers } 45^\circ \approx 0,292893$. Voici maintenant les valeurs fournies par les deux algorithmes du texte :

⁹Cette valeur initiale n'est autre, sans doute, qu'une ancienne formule empirique reliant la corde et le sinus. Comme elle fournit déjà une assez bonne valeur approchée, il est naturel de s'en servir pour initialiser le processus itératif. Dans le cas particulier du premier calcul (celui de $\text{Sin } \frac{\theta}{2^p}$), la corde n'est pas connue mais nous avons vu qu'on pouvait considérer, sans perte significative de précision, que $C = R\alpha$; la valeur initiale devient donc, tout simplement, $u_1 = R\alpha$.

¹⁰La notation algébrique moderne permet de donner une seule formule là où le texte indien doit énoncer deux règles selon le "signe" de la correction à effectuer.

$$\begin{aligned} u_1 &= 0,710736; u_2 = 0,705585; u_3 = 0,707734; u_4 = 0,706846; \\ u_5 &= 0,707215; u_6 = 0,707062; u_7 = 0,707125; u_8 = 0,707099; \\ u_9 &= 0,707110; u_{10} = 0,707105; u_{11} = 0,707107. \end{aligned}$$

$$v_1 = 0,283973; v_2 = 0,296541; v_3 = 0,292921; v_4 = 0,292893.$$

En définitive, ce texte témoigne du haut niveau atteint par les astronomes indiens dans le domaine du calcul numérique. On appréciera notamment la clarté et la précision du langage utilisé par l'auteur pour décrire, en l'absence de toute notation symbolique, des algorithmes très élaborés.

d) Texte 3 : Qādī-Zāda, XV^{ème} siècle

Chez les Arabes, on a retrouvé des traces de calculs itératifs dès le IX^{ème} siècle : Habash al-Hāsib, un contemporain d'al-Khwārizmī, traite par approximations successives une équation transcendante de la forme $t = x - e \sin x$, où t et e sont donnés (en 1618, Kepler a lui aussi résolu par un procédé itératif une équation analogue pour établir ses célèbres *Tables rudolphines*¹¹).

Toujours chez les Arabes, le mathématicien al-Kāshī, qui vivait à Samarkand dans la première moitié du XV^{ème} siècle, a développé un procédé original d'itération pour le calcul de $\sin 1^\circ$. Depuis Ptolémée, ce calcul était une étape délicate et incontournable dans la confection des tables trigonométriques¹². Al-Kāshī part de la formule $\sin 3\alpha = 3 \sin \alpha - 4 \sin^3 \alpha$, formule qui fait apparaître $\sin 1^\circ$ comme solution de l'équation $3x - 4x^3 = \sin 3^\circ$. L'équation, mise sous la forme $x = \frac{q+x^3}{p}$, fait ensuite l'objet d'un traitement consistant à "extraire" un à un les chiffres successifs (en numération sexagésimale) de la racine. L'ouvrage d'al-Kāshī, intitulé *Sur la corde et le sinus*, n'a pas été retrouvé, mais la méthode de calcul de $\sin 1^\circ$ nous a cependant été transmise car elle est décrite par Mīram Shalabī (un astronome ayant vécu en Turquie au début du XVI^{ème} siècle) dans ses *Règles des opérations et correction des tables*. Par chance, on a trouvé récemment un traité de Qādī-Zāda (le grand-père de Shalabī), intitulé *Traité sur la détermination du sinus d'un degré*, qui donne de la méthode une description issue plus directement de l'enseignement d'al-Kāshī. En voici un extrait¹³.

On divise d'abord quelques chiffres du nombre par le nombre des choses. On forme le cube du quotient et on l'ajoute au reste du nombre. On divise encore une fois leur somme. On forme le cube de la somme des deux quotients et on ajoute son excès par rapport au cube obtenu d'abord au reste de la somme. On divise ensuite leur deuxième somme encore une fois. On forme le cube de la somme des quotients et on ajoute l'excès de ce cube par rapport au cube de la somme des deux quotients au reste de la deuxième somme. On divise ensuite encore une fois la troisième somme et on procède comme auparavant. L'opération est terminée dès que l'on arrive à ce que l'on peut négliger.

¹¹On trouvera le texte de Kepler et une analyse dans CHABERT *et al.*, 1994, p. 250-254.

¹²Les formules d'addition et de duplication permettent de calculer directement $\sin 3^\circ$ à partir des lignes trigonométriques connues de 30° et de 36° : il suffit de remarquer que $3^\circ = (36^\circ - 30^\circ)/2$. Par contre, il n'existe pas de méthode analogue pour $\sin 1^\circ$; dans ce cas, on doit nécessairement faire appel à un procédé de calcul approché.

¹³La traduction française est tirée de YOUSCHKEVITCH, 1976, p. 176. Pour d'autres études sur la méthode d'al-Kāshī, voir AABOE, 1954 et RIAHI, 1995.

Dans l'équation $x = \frac{q+x^3}{p}$, ou $px = q + x^3$, q est le "nombre" et p le "nombre des choses". La quantité inconnue x , c'est-à-dire les "choses", est cherchée sous la forme $a + b + c + \dots$, où $a, b, c, \dots$ sont les chiffres successifs de x en unités sexagésimales¹⁴. On peut alors expliquer le texte de la façon suivante :

- Première étape : x étant très petit (c'est le sinus de 1°), son cube est négligeable et on peut considérer que $x = a + \dots = \frac{q+x^3}{p} \approx \frac{q}{p}$. Le chiffre a s'obtient donc comme le premier chiffre de la division de q par p .
- Deuxième étape : sachant maintenant que $x = a$, on peut écrire de façon plus précise que $x = a + b + \dots = \frac{q+x^3}{p} = \frac{q+a^3}{p}$, d'où $b + \dots = \frac{q+a^3}{p} - a = \frac{(q-ap)+a^3}{p}$. Le chiffre b est ainsi le premier chiffre de la division de $(q - ap) + a^3$ par p (les parenthèses servent seulement à indiquer que la quantité $q - ap$ n'est pas recalculée ici, puisque c'est le "reste du nombre", autrement dit le reste de la division effectuée lors de la première étape).
- Troisième étape : on recommence le processus avec la nouvelle approximation $x \approx a + b$. De façon analogue, il vient $x = a + b + c + \dots = \frac{q+x^3}{p} = \frac{q+(a+b)^3}{p}$, d'où l'on déduit que $c + \dots \approx \frac{q+(a+b)^3}{p} - a - b = \frac{(q-ap+a^3-bp)+(a+b)^3-a^3}{p}$, ce qui permet de calculer le chiffre c (à nouveau, on remarque que la quantité $q - ap + a^3 - bp$ est le reste de la division effectuée lors de l'étape précédente).

Il est loisible de continuer ainsi pour obtenir autant de chiffres que nécessaire. Abstraction faite des problèmes techniques de calcul liés à l'utilisation du système sexagésimal, l'algorithme d'al-Kāshī peut se résumer par la suite récurrente

$$\begin{cases} x_1 &= 0 \\ x_{n+1} &= \frac{q+x_n^3}{p} \end{cases}$$

Cette idée d'"extraire" l'un après l'autre les chiffres successifs de la racine d'une équation, un peu comme on le fait dans les procédés classiques d'extraction des racines carrées ou cubiques, n'est pas vraiment nouvelle, au XV^{ème} siècle, dans les mathématiques arabes. On la rencontre déjà trois siècles plus tôt dans le *Traité des équations* de Sharaf al-Dīn al-Tūsī. Ce dernier ouvrage contient en particulier un procédé de résolution numérique des équations du troisième degré¹⁵ qui s'apparente assez étroitement à celui qui sera exposé plus tard par Newton.

2 Le génie de Newton : des équations numériques aux équations fonctionnelles

Ainsi qu'il vient d'être dit, Newton n'est pas l'inventeur de la méthode qui porte son nom. On sait maintenant que son œuvre dans le domaine de la résolution numérique des équations est l'aboutissement d'une longue tradition, commençant au moins au XII^{ème} siècle avec Sharaf

¹⁴Il s'agit là d'un abus de notation destiné à alléger les écritures. Il faut comprendre que, si l'écriture réelle en base soixante de $\sin 1^\circ$ est $0; \alpha, \beta, \gamma, \dots$, alors $a = \alpha.60^{-1}$, $b = \beta.60^{-2}$, $c = \gamma.60^{-3}$, ...

¹⁵Sur les algorithmes numériques de Sharaf al-Dīn al-Tūsī, la référence qui s'impose est RASHED, 1984, chap. III. Dans HOUZEL, 1995, on trouvera un complément intéressant permettant de bien comprendre en quoi les calculs du mathématicien arabe préfigurent ceux de Newton.

al-Dīn al-Tūsī, et se poursuivant ensuite chez les mathématiciens arabes puis européens, avec notamment Viète (1600), Kepler (1618), Harriot (1631) et Oughtred (1652). D'ailleurs, dans une lettre à Oldenbourg du 26 juin 1676, Newton indique lui-même que, pour sa part, il s'est inspiré des travaux de Viète et d'Oughtred¹⁶.

a) Une méthode universelle pour toutes les équations

Un exposé détaillé de la méthode de Newton se trouve dans *La méthode des fluxions et des suites infinies*¹⁷, traité achevé dès 1671. On peut y lire ceci :

Il faut que nous entrions dans un détail un peu plus grand pour expliquer comment on doit réduire les racines de ces équations à des suites infinies, car ce que les géomètres nous ont donné sur les équations en nombres est extrêmement embarrassé et chargé d'opérations superflues, de sorte qu'on ne peut prendre sur cela un bon modèle pour faire les mêmes opérations en espèces. Je ferai donc voir d'abord comment se doit faire en nombres la réduction des équations affectées, et ensuite j'appliquerai la méthode aux espèces.

L'extrait est des plus explicites quant aux objectifs de l'auteur : en ce qui concerne les "équations en nombres" (les équations numériques), Newton ne prétend pas apporter quelque chose de fondamentalement nouveau, si ce n'est clarifier et simplifier la démarche "embarrassée" de ses prédécesseurs; par contre, il se propose d'étendre le procédé des approximations successives aux "équations en espèces", c'est-à-dire aux équations littérales.

Le but annoncé est des plus ambitieux : disposer d'une méthode universelle pour résoudre toutes les équations, qu'il s'agisse d'équations numériques $f(x) = 0$, d'équations fonctionnelles simples $f(x, y) = 0$ reliant deux quantités inconnues, ou même d'équations fonctionnelles plus complexes $f(x, y, \dot{x}, \dot{y}) = 0$ faisant intervenir, en outre, les "fluxions" des quantités inconnues¹⁸ (dans les deux derniers cas, il s'agit bien d'une équation fonctionnelle dans la mesure où la résolution consiste à tenter d'exprimer l'une des quantités inconnues en "fonction" de l'autre). L'idée de base de la technique utilisée est d'effectuer les opérations littérales de l'algèbre à l'aide des suites¹⁹ infinies, de la même façon qu'on effectue les opérations numériques de l'arithmétique au moyen des fractions décimales. Au développement décimal d'un nombre correspond formellement le développement en série d'une expression littérale. Une réflexion en profondeur sur les anciennes méthodes d'extraction des chiffres successifs de la racine d'une équation numérique (que ce soit en base 10 ou en base 60) a conduit Newton à imaginer un processus analogue pour les équations fonctionnelles : la fonction inconnue est cherchée sous la forme d'une série de puissances de la variable indépendante, et les termes successifs de la série sont extraits un à un.

¹⁶Sur la façon dont Newton a étudié les travaux de ses prédécesseurs, et sur ses tentatives antérieures à 1671, on pourra consulter YPMA, 1995.

¹⁷NEWTON, 1671. Dans la suite, nous nous référerons à la traduction française du Marquis de Buffon.

¹⁸Dans la conception de Newton, les fluxions $\dot{x}$ et $\dot{y}$ représentent les variations instantanées des grandeurs variables x et y par rapport à une grandeur de référence variant uniformément, et qu'on peut supposer être le temps t . Avec les notations de Leibniz, on aurait $\dot{x} = dx/dt$ et $\dot{y} = dy/dt$. Le quotient $\dot{y}/\dot{x}$ représente alors la dérivée dy/dx de y par rapport à x , notée aujourd'hui y' . Lorsqu'on suppose que x est la variable de référence, il vient tout simplement $\dot{x} = 1$ et $\dot{y}/\dot{x} = y'$.

¹⁹Au XVII^{ème} siècle, le mot "suite" est employé avec le sens actuel de "série".

Afin d'illustrer cette généralisation audacieuse, nous allons commenter trois extraits de *La méthode des fluxions*²⁰, traitant tour à tour d'une équation numérique, d'une équation fonctionnelle ordinaire et d'une équation différentielle.

b) Textes 4, 5, 6 : Newton, 1671

Le premier exemple est une équation numérique du troisième degré. Dans le texte original, les calculs sont disposés dans un tableau que nous n'avons pas reproduit ici.

Soit l'équation $y^3 - 2y - 5 = 0$ à réduire en suite infinie; prenez un nombre comme 2, qui ne diffère pas d'une de ses dixièmes parties de la vraie valeur de la racine, et faites $2 + p = y$, substituez $2 + p$ pour y dans l'équation donnée, et vous aurez $p^3 + 6p^2 + 10p - 1 = 0$, dont il faut chercher la racine pour l'ajouter au quotient; rejetez $p^3 + 6p^2$ à cause de sa petitesse, il restera $10p - 1 = 0$, ou $p = 0,1$; ce qui est très près de la vraie valeur de p ; c'est pourquoi, l'écrivant au quotient, je fais $0,1 + q = p$, et substituant comme auparavant, j'ai $q^3 + 6,3q^2 + 11,23q + 0,061 = 0$; négligeant les deux premiers termes, il reste $11,23q + 0,061 = 0$, ou $q = -0,0054$ à peu près [...]. J'écris donc $-0,0054$ dans le quotient, mais au-dessous parce que ce terme est négatif, et supposant $-0,0054 + r = q$, je substitue comme auparavant, et je continue ainsi l'opération aussi longtemps qu'il convient [...].

L'équation s'écrit $f(y) = 0$, avec $f(y) = y^3 - 2y - 5$. Newton prend comme première approximation $y_1 = 2$, sans doute après avoir constaté que $f(2) < 0$, que $f(3) > 0$, et donc que la racine se trouve entre 2 et 3. L'idée générale est alors de chercher la racine sous la forme d'une "suite infinie" $y = 2 + p + q + r + \dots$, dont chaque terme est petit par rapport au précédent. Pour trouver la seconde approximation $y_2 = 2 + p$, on écrit que $f(2 + p) = 0$, égalité qui devient $10p - 1 = 0$ après suppression des termes négligeables en p^2 et en p^3 . Ainsi que le remarquera Joseph Raphson²¹ en 1715, cette linéarisation de l'équation revient à déterminer p par l'égalité $f(2) + pf'(2) = 0$, de sorte que $y_2 = 2 + 0,1 = 2 - \frac{f(2)}{f'(2)}$. De même, si on pose ensuite $y_3 = 2,1 + q$, l'équation obtenue par Newton pour la détermination de q n'est autre que $f(2,1) + qf'(2,1) = 0$. Grâce à la remarque de Raphson, le processus assez lourd décrit par le texte (à chaque étape, changement d'inconnue et linéarisation de la nouvelle équation) peut être ramené au schéma itératif simplifié utilisé aujourd'hui :

$$\begin{cases} y_1 &= 2 \\ y_{n+1} &= y_n - \frac{f(y_n)}{f'(y_n)}. \end{cases}$$

Pendant longtemps, les deux méthodes, celle de Newton et celle de Raphson, ont été exposées comme si elles étaient distinctes. C'est Lagrange, en 1798, qui mettra définitivement en évidence le fait que "ces deux méthodes ne sont au fond que la même présentée différemment". Depuis, on a pris l'habitude de parler de la "méthode de Newton-Raphson"²².

²⁰Ces extraits se trouvent respectivement p. 7, p. 12 et p. 34.

²¹Sur les travaux de Raphson, voir CHABERT *et al.*, 1994, p. 201-204.

²²Dans YPMA, 1995, on apprend que Thomas Simpson a joué un rôle non négligeable dans la mise au point de la version définitive de la méthode. Alors que Newton et Raphson conduisaient des raisonnements purement algébriques, Simpson a, le premier, fait intervenir explicitement le calcul des fluxions. Cependant, Lagrange, dans son *Traité de la résolution des équations numériques* (publié en 1798), n'a fait référence qu'à Newton et à Raphson. De même, Fourier, dans son *Analyse des équations indéterminées* (publiée en 1831), a seulement parlé de "la méthode newtonienne". On comprend donc pourquoi le nom de Simpson a ensuite été oublié.

Sans nous attarder davantage sur cette méthode de résolution numérique des équations qui, comme tout le montre depuis le début de cet article, est l'aboutissement de deux millénaires d'efforts continus²³, passons à ce qui constitue l'apport original et exceptionnel de Newton : la résolution approchée des équations fonctionnelles. L'extrait suivant de *La méthode des fluxions* explique tout d'abord comment il faut procéder dans le cas d'une équation fonctionnelle simple, c'est-à-dire sans fluxions (comme pour l'extrait précédent, nous n'avons pas reproduit le tableau dans lequel sont organisés les calculs).

Soit donc pour exemple l'équation $y^3 + a^2y + axy - 2a^3 - x^3 = 0$, dont il faille tirer la racine; je choisis, comme je l'ai dit, les termes $y^3 + a^2y - 2a^3$, qui étant égaux à zéro me donnent $y - a = 0$, ainsi j'écris $+a$ dans le quotient; mais, parce que $+a$ n'est pas la valeur complète de y , je fais $a + p = y$, et je substitue dans l'équation cette valeur de y , et parmi les termes $p^3 + 3ap^2 + axp$, etc. qui en résultent, je choisis de la même façon les termes $+4a^2p + a^2x$, qui étant égaux à zéro donnent $p = -\frac{1}{4}x$, j'écris donc $-\frac{1}{4}x$ au quotient; mais parce que $-\frac{1}{4}x$ n'est pas la valeur exacte de p , je fais $-\frac{1}{4}x + q = p$, et substituant cette valeur je choisis dans les termes $q^3 - \frac{3}{4}xq^2 + 3aq^2$, etc. qui en résultent, les termes $4a^2q - \frac{1}{16}ax^2$ qui étant égaux à zéro donnent $q = \frac{ax}{64a}$ que j'écris au quotient, mais comme $\frac{ax}{64a}$ n'est pas la valeur exacte de q , je fais $\frac{ax}{64a} + r = q$, et je substitue comme ci-dessus, ce qui se peut continuer aussi longtemps qu'on voudra [...].

Cette fois, l'équation à résoudre s'écrit $f(x, y) = 0$, avec $f(x, y) = y^3 + a^2y + axy - 2a^3 - x^3$. Newton cherche la solution sous la forme d'une série de puissances de la variable indépendante x . Les calculs du texte s'interprètent plus facilement si on note cette série $y = \alpha + \beta x + \gamma x^2 + \dots$. En substituant dans l'équation et en négligeant les termes de degré supérieur ou égal à 1, il vient $\alpha^3 + a^2\alpha - 2a^3 = (\alpha - a)(\alpha^2 + a\alpha + 2a^2) = 0$, d'où $\alpha = a$. La première approximation est donc $y_1 = a$. Newton reprend alors, avec les mêmes notations, la démarche mise au point pour les équations numériques : il effectue le changement d'inconnue $y = a + p$, ce qui conduit à la nouvelle équation $p^3 + 3ap^2 + (4a^2 + ax)p + a^2x - x^3 = 0$. En remplaçant p par $\beta x + \gamma x^2 + \dots$ et en négligeant les termes de degré supérieur ou égal à 2, on aboutit à $4a^2\beta x + a^2x = 0$, c'est-à-dire à $\beta = -\frac{1}{4}$. À partir de la seconde approximation $y_2 = y_1 - \frac{1}{4}x$, une nouvelle itération conduit à $y_3 = y_2 + \frac{1}{64a}x^2$, et ainsi de suite.

Soit proposée l'équation $\frac{y}{x} = 1 - 3x + y + x^2 + xy$; j'écris de suite et de gauche à droite dans une ligne au-dessus les termes $1 - 3x + x^2$, qui ne sont pas affectés de la quantité relative y , et j'écris le reste y et xy dans une colonne à main gauche.

	+1	-3x	+x ²	
+y	*	+x	-xx	$+\frac{1}{3}x^3 - \frac{1}{6}x^4 + \frac{1}{30}x^5$, &c.
+xy	*	*	+xx	$-x^3 + \frac{1}{3}x^4 - \frac{1}{6}x^5 + \frac{1}{30}x^6$, &c.
Somme	1	-2x	+xx	$-\frac{2}{3}x^3 + \frac{1}{6}x^4 - \frac{4}{30}x^5$, &c.
y =	x	-xx	$+\frac{1}{3}x^3 - \frac{1}{6}x^4 + \frac{1}{30}x^5 - \frac{1}{45}x^6$, &c.	

Et d'abord je multiplie le premier terme 1 par la quantité corrélatrice x , et divisant le produit $1x$ ou x par le nombre 1 des dimensions, j'écris $\frac{x}{1}$ ou simplement x dans le quotient au-dessous; puis au

²³On pourrait signaler aussi que, peu après Newton, son compatriote Edmond Halley a publié, en 1694, une méthode encore plus performante se traduisant par le schéma itératif $x_{n+1} = x_n - \frac{2f(x_n)f'(x_n)}{2f'(x_n)^2 - f(x_n)f''(x_n)}$. La méthode de Halley revient à effectuer un développement de Taylor à l'ordre 2 alors que celle de Newton se limitait à l'ordre 1 (par suite, l'ordre de convergence est 3 au lieu de 2). Mais on ne sait pas vraiment comment Halley a pu raisonner pour aboutir à un tel résultat : voir SCAVO et THOO, 1995, où l'on trouvera aussi un grand nombre de références antérieures sur la méthode de Halley.

lieu de y substituant cette valeur dans les termes $+y$ et $+xy$ de la colonne à main gauche, j'ai $+x$ et $+xx$, que j'écris vis-à-vis et à main droite; ensuite je prends dans ce qui reste les termes les plus bas $-3x$ et $+x$, dont la somme $-2x$ multipliés par x devient $-2xx$, qui divisé par le nombre 2 des dimensions donne $-xx$ pour le second terme de la valeur de y dans le quotient. Prenant donc ce terme et le substituant au lieu de y , j'ai $-xx$ et $-x^3$ qu'il faut ajouter respectivement aux termes $+x$ et $+xx$, écrits vis-à-vis de y et yx . Je prends de même les plus bas termes $+xx - xx + xx$ de la somme xx desquels etc. je tire le troisième terme $+\frac{1}{3}x^3$ de la valeur de y , et après l'avoir substitué etc. je tire des plus bas termes $\frac{1}{3}x^3$ et $-x^3$ le quatrième terme $-\frac{1}{6}x^4$. Ce que l'on peut continuer aussi longtemps qu'on le jugera à propos.

Ce procédé pour extraire la fluente y , terme après terme, est tout à fait analogue à celui utilisé plus haut pour extraire une racine d'une équation numérique. Nous pouvons encore l'analyser comme un schéma itératif, en interprétant le calcul de Newton de la façon suivante : prenons $y_1 = 0$ comme première approximation de l'équation différentielle $y' = 1 - 3x + x^2 + (1+x)y$; en substituant dans le second membre, on obtient $y' = 1 - 3x + x^2$, puis, en négligeant les termes d'ordre supérieur ou égal à 1, $y' = 1$, d'où la nouvelle approximation $y_2 = x$; en substituant à nouveau, on aboutit à l'équation $y' = 1 - 2x + 2x^2$, qui, par suppression des termes d'ordre supérieur ou égal à 2, s'écrit plus simplement $y' = 1 - 2x$, d'où l'approximation $y_3 = x - x^2$; et ainsi de suite. Finalement, en faisant abstraction de la suppression, à chaque étape, des termes non significatifs de degré élevé, l'algorithme peut s'écrire

$$\begin{cases} y_1 &= 0 \\ y_{n+1} &= \int_0^x f(x, y_n) dx. \end{cases}$$

Ce schéma d'intégrations successives conduit à un développement en série de la solution qui s'annule en 0. Le problème de la constante d'intégration n'est pas directement évoqué.

3 Vers le théorème du point fixe de Banach

Newton ayant ainsi ouvert la voie, on va assister pendant les XVIII^{ème} et XIX^{ème} siècles à une extension progressive de la méthode des approximations successives aux équations fonctionnelles les plus diverses : équations différentielles ordinaires, équations aux dérivées partielles, équations aux différences finies, équations intégrales, équations intégral-différentielles, etc.

a) Les approximations successives des astronomes

C'est en astronomie que l'idée originale de Newton devait connaître ses applications les plus ambitieuses. En 1747, à peu près simultanément, Clairaut, d'Alembert et Euler réussissent à mettre en équations le redoutable problème des trois corps et en donnent les premières solutions approchées analytiques sous forme de développements en séries trigonométriques²⁴. D'Alembert et Euler, dans la plupart de leurs recherches, s'inspirent directement de Newton : ils déterminent un à un les termes des séries grâce à des substitutions successives dans les équations différentielles. Cette démarche sera généralement retenue par la suite, ainsi qu'en témoigne Laplace, en 1776, dans la seconde partie de ses *Recherches sur le calcul intégral et sur le système du monde* :

²⁴Vu leur complexité, nous ne pouvons pas aborder ces calculs ici. On pourra consulter TOURNÈS, 1996, chap. IV et TATON et WILSON, 1995.

Ayant ainsi une première valeur approchée des variables, on substitua dans les équations différentielles, au lieu de chaque variable, cette valeur, plus une très petite indéterminée dont on négligea le carré et les puissances supérieures, et, en continuant d'opérer ainsi, on eut une seconde, une troisième, etc. valeur approchée. Cette méthode, analogue à celle de Newton pour déterminer par approximation les racines des équations numériques, se présenta naturellement aux géomètres qui résolurent les premiers le problème des trois corps.

Clairaut, quant à lui, adopte un point de vue assez différent²⁵. Dans ses recherches sur la théorie de la Lune (1750-1765) et sur le retour de la comète de Halley (1757-1760), il aboutit à une équation différentielle de la forme $\frac{d^2r}{dv^2} = F(v, r, \frac{dr}{dv})$, analogue à celle de d'Alembert. Mais, au lieu de procéder directement à des substitutions successives dans cette équation différentielle, il la transforme d'abord en une équation intégrale de la forme $r = G(v, r, \frac{dr}{dv})$. Ce n'est qu'après qu'il part d'une première valeur supposée r_0 de la fonction inconnue r , qu'il obtient une valeur corrigée de cette même fonction en substituant r_0 dans le second membre, soit $r_1 = G(v, r_0, \frac{dr_0}{dv})$, et ainsi de suite. Cette façon de remplacer explicitement l'équation différentielle par une équation intégrale équivalente constitue, en quelque sorte, une anticipation de la méthode de PICARD dont nous parlerons plus loin.

Ainsi, sans entrer dans les détails, on peut dire que la méthode des approximations successives a été abondamment pratiquée pendant la seconde moitié du XVIII^{ème} siècle et au début du XIX^{ème} siècle, essentiellement par les astronomes. Le texte que nous allons étudier maintenant permettra de faire le point sur la situation à l'issue de cette longue période de recherches empiriques.

b) Texte 7 : SCHMIDTEN, 1821

Dans un mémoire paru en 1821²⁶, Henri Gerner SCHMIDTEN offre un panorama de la méthode des approximations successives et de ses possibles applications dans les diverses branches de l'analyse. Selon lui, la méthode des approximations successives, connue depuis longtemps, s'applique à tous les types d'équations :

Depuis longtemps on se sert du principe des substitutions successives, comme d'une méthode d'approximation, fondée sur des valeurs particulières des quantités qui entrent dans l'équation proposée; et on l'a employée, faute de méthodes plus rigoureuses; c'est pourquoi je me suis surtout attaché à l'exposer sous un point de vue qui doit la faire considérer comme la seule méthode générale qui existe pour l'intégration des équations.

Pour illustrer cette philosophie toute newtonienne, SCHMIDTEN propose de nombreux exemples, en considérant tour à tour des équations différentielles à deux variables, des équations aux différences finies à deux variables, des équations aux différences mêlées à deux variables et des équations linéaires aux différences partielles. Ces exemples, traités de façon purement

²⁵Voir TOURNÈS, 1996, (chap. IV et V) et TOURNÈS, 1998. Nous montrons notamment en quoi la méthode de Clairaut est à l'origine des quadratures mécaniques par approximations successives pratiquées par les astronomes, et des méthodes multipas pour l'intégration numérique des équations différentielles ordinaires. Par ailleurs, dans GREENBERG, 1988, on verra que Clairaut savait aussi appliquer la méthode itérative à des équations numériques : plus précisément, Greenberg explique comment Clairaut résout ainsi un système de quatre équations numériques à quatre inconnues.

²⁶SCHMIDTEN, 1820/21. L'extrait sélectionné se trouve p. 270-271.

formelle, donnent lieu à des développements calculatoires dans la tradition de l'école combinatoire allemande. Auparavant, le début du mémoire contient un exposé abstrait de la méthode des substitutions successives, d'une grande beauté formelle en même temps qu'étonnamment moderne :

Soit, en effet, y une fonction d'un certain nombre de variables indépendantes, donnée par l'équation différentielle

$$\varphi \cdot y = f \cdot y;$$

$\varphi \cdot y$ étant une fonction qui contient les coefficients différentiels ou aux différences de l'ordre le plus élevé qui soient dans l'équation proposée, et $f \cdot y$ étant une autre fonction quelconque des variables indépendantes [contenant] des coefficients différentiels ou aux différences; on aura l'équation intégrale

$$y = X + \frac{1}{\varphi} f \cdot y;$$

$1/\varphi$ signifiant la fonction inverse de φ ; et X étant la fonction la plus générale qui satisfasse à l'équation $\varphi \cdot X = 0$.

Au moyen de cette relation implicite, on trouvera facilement la valeur explicite de y , par des substitutions successives; ce sera

$$y = X + \frac{1}{\varphi} f(X + \frac{1}{\varphi} f(X + \frac{1}{\varphi} f(X + \dots$$

Maintenant, il se peut que chaque substitution rapproche cette série de la véritable valeur de y ; mais il se peut aussi qu'elle l'en éloigne; et alors on devra donner une autre forme à la série; ce qui est toujours possible d'autant de manière différentes qu'il y en aura de partager l'équation entre les deux termes $\varphi \cdot y$ et $f \cdot y$.

On voit cependant que la valeur de y restera, en général, très compliquée, à moins que $\varphi \cdot y$ et $f \cdot y$ ne soient linéaires par rapport à y , ce qui embrasse déjà une classe d'équation très étendue et très importante : celle des équations linéaires.

On a, dans ce cas,

$$y = X + \frac{1}{\varphi} f \cdot X + \frac{1}{\varphi} f \left(\frac{1}{\varphi} f \cdot X \right) + \frac{1}{\varphi} f \left(\frac{1}{\varphi} f \left(\frac{1}{\varphi} f \cdot X \right) \right) + \dots;$$

et je me propose d'en exposer les principales conséquences [...].

Ce texte appelle trois groupes de remarques :

1) SCHMIDTEN raisonne directement sur une équation dont l'inconnue est une fonction. À cette fonction sont appliqués des opérateurs sur lesquels on peut effectuer, à nouveau, des opérations d'un niveau supérieur (par exemple, prendre l'inverse, c'est-à-dire appliquer l'opérateur réciproque). Le concept d'équation fonctionnelle paraît donc tout à fait clair en ce début du XIX^{ème} siècle, soit plus tôt que ce que l'on dit parfois. Si, en l'absence de toute notion topologique et de toute considération de convergence, il serait abusif de parler déjà d'analyse fonctionnelle, par contre le terme d'algèbre fonctionnelle semble parfaitement approprié.

2) Pour SCHMIDTEN, le problème de l'existence ne se pose pas. Lorsqu'il parle de la "véritable valeur de y ", il est clair que cette valeur existe et qu'il s'agit seulement d'en trouver un développement en série. De plus, les substitutions successives fournissent toujours la "valeur explicite de y ", que le développement soit convergent ou divergent. On est en plein dans la tradition eulérienne : les algorithmes infinis restent recevables même quand ils ne débouchent pas sur la possibilité d'un calcul numérique. Une idée supplémentaire très importante se greffe là-dessus : il y a une infinité de façons de transformer une équation en une équation à point fixe et, sans que l'on sache expliquer pourquoi, certaines équations à point fixe conduisent à des séries convergentes et d'autres non.

3) En dehors des équations linéaires, les calculs sont inextricables; plus précisément, la loi de formation du développement en série reste implicite. Pour une équation non linéaire, on peut toujours calculer l'un après l'autre les premiers termes du développement mais on ne disposera jamais que d'un nombre fini de ces termes : le développement reste potentiel. Or, pour aborder valablement le problème de la convergence, il faut pouvoir raisonner sur l'infinité des termes de la série. On voit mal comment y parvenir sans disposer d'une expression explicite, directe ou récurrente, de tous les termes. Si SCHMIDTEN se restreint aux équations linéaires, c'est parce que, pour lui, dans l'optique de l'analyse algébrique, l'essentiel est de caractériser complètement la loi de formation du développement. Plus tard, lorsqu'on s'intéressera à la convergence, il sera également naturel, pour la raison que nous avons indiquée, de s'en tenir aux équations linéaires.

Ce texte de SCHMIDTEN de 1821 apporte la preuve que, au début du XIX^{ème} siècle, la technique empirique des substitutions successives avait déjà revêtu l'aspect d'un problème de point fixe très général, assez bien théorisé dans son aspect formel. Restait à expliquer ce qui, dans certains cas, provoquait le succès numérique du procédé, c'est-à-dire déterminer des conditions suffisantes assurant la convergence des séries obtenues.

c) Le problème de la convergence

Jusqu'à présent, nous n'avons rencontré dans aucun des textes étudiés ce qu'en termes modernes nous appelons le *problème de la convergence*. Pour tous les auteurs cités, y compris SCHMIDTEN, ce problème ne se pose pas. Dans leur esprit, il va de soi que l'équation à résoudre, qu'elle soit numérique ou fonctionnelle, admet une solution, et qu'il s'agit seulement d'en calculer des valeurs approchées. SCHMIDTEN souligne clairement que, dans la mise en œuvre des substitutions successives, on constate parfois que les valeurs successives se rapprochent de la solution, et parfois qu'elles s'en éloignent, mais la réflexion n'est pas poussée plus avant.

Dans le cas des équations numériques, c'est Fourier qui, semble-t-il, aborde pour la première fois la question de la convergence de la méthode de Newton dans un article de 1818, mais c'est Cauchy qui en réalise le premier traitement correct dans son *Cours d'Analyse* de 1821²⁷.

En ce qui concerne les équations différentielles ordinaires, on peut mentionner des recherches de Liouville en 1830, sur la théorie de la chaleur, et en 1837-38, sur ce qui deviendra la "théorie de Sturm-Liouville". Au cours de ces travaux, Liouville résout par approximations successives

²⁷ Sur ces travaux de Fourier et Cauchy, voir CHABERT *et al.*, 1994, p. 210-216.

des équations différentielles linéaires du second ordre, avec une preuve quasiment complète de la convergence. Liouville démontre non seulement la convergence de la série qui définit la solution, mais aussi la convergence de la série des dérivées, ce qui justifie la dérivation terme à terme. La seule faiblesse est l'absence d'une notion explicite de convergence uniforme, mais cela ne compromet pas fondamentalement la validité de la preuve dans la mesure où n'interviennent que des séries entières. On trouve des raisonnements analogues à ceux de Liouville dans divers travaux de Cauchy à partir de 1830, de sorte qu'il est difficile de savoir lequel des deux hommes a eu le premier l'idée de ce type de démonstration²⁸.

D'autres travaux sur la méthode des approximations successives sont publiés par Caqué (1864), Fuchs (1870) et Peano (1887), toujours dans le cas des équations différentielles ordinaires linéaires. Parallèlement, la méthode itérative est également employée pour résoudre des équations aux dérivées partielles linéaires : en 1877, Neumann aborde ainsi l'équation de Laplace et, en 1885, Schwarz fait de même pour l'équation des membranes vibrantes²⁹. C'est dans ces travaux de Neumann et de Schwarz que PICARD trouve son inspiration pour publier, en 1890³⁰, une preuve de la méthode des approximations successives s'appliquant, pour la première fois, à des équations *non linéaires*. Le mémoire de Picard débute par ces mots :

Considérons une équation du second ordre aux dérivées partielles de la forme [...]. On peut, pour intégrer cette équation, avec des conditions aux limites déterminées, procéder de la manière suivante par approximations successives.

Ce n'est qu'au bout d'une cinquantaine de pages que PICARD, presque incidemment, s'intéresse aux équations différentielles ordinaires : "Les méthodes d'approximation dont nous venons de faire usage peuvent évidemment être employées dans le cas des équations différentielles ordinaires [...]". Ainsi donc, c'est par le biais des équations aux dérivées partielles que PICARD aboutit finalement à une nouvelle démonstration du théorème d'existence de Cauchy-Lipschitz. Par cette démonstration générale, il parachève l'œuvre entreprise auparavant par Liouville, Cauchy et PEANO. De nos jours, du moins en ce qui concerne les équations différentielles, l'expression "méthode de Picard" est devenue synonyme de "méthode des approximations successives".

d) Texte 8 : PEANO

Plutôt qu'un texte de PICARD, nous avons choisi d'étudier un extrait du mémoire de PEANO de 1887 sur l'intégration par séries des équations différentielles linéaires. Laissant de côté le mémoire original en italien, nous nous référons à la traduction française parue l'année suivante avec de légères modifications³¹. L'originalité de PEANO est double : la méthode itérative est utilisée pour intégrer un système d'équations linéaires du premier ordre (alors qu'on s'était surtout intéressé, jusque là, à une seule équation linéaire d'ordre supérieur) et, pour la première fois, un système d'équations est traité comme une *équation unique*. PEANO considère le système d'équations différentielles linéaires homogènes

$$\begin{aligned} \frac{dx_1}{dt} &= r_{11}x_1 + \dots + r_{1n}x_n \\ &\dots \dots \dots \\ \frac{dx_n}{dt} &= r_{n1}x_1 + \dots + r_{nn}x_n \end{aligned}$$

où les r_{ij} sont des fonctions réelles de la variable t , continues dans l'intervalle (p, q) . S'inspirant des travaux de Grassmann, Hamilton, Cayley et Sylvester³², il ramène le système à la seule équation

$$\frac{dx}{dt} = Rx,$$

en introduisant le "nombre complexe" (le vecteur) $x = [x_1, \dots, x_n]$ et la "substitution" (la transformation linéaire) représentée par la "matrice"

$$R = \begin{Bmatrix} r_{11} & \dots & r_{1n} \\ \dots & \dots & \dots \\ r_{n1} & \dots & r_{nn} \end{Bmatrix} = \{r_{ij}\}.$$

La moitié de l'article est naturellement consacrée à une étude préliminaire des propriétés de ces objets, qui nous sont aujourd'hui si familiers mais qui étaient alors peu connus des mathématiciens. Pour comprendre la suite, il suffira de savoir que PEANO définit les "modules" (les normes) d'un nombre complexe et d'une substitution par les formules :

$$\begin{aligned} \text{mod.} x &= \sqrt{x_1^2 + x_2^2 + \dots + x_n^2}, \\ \text{mod.} R &= \text{maximum de } \frac{\text{mod.}(Rx)}{\text{mod.} x}. \end{aligned}$$

Une fois ces notions et notations introduites, la démonstration de PEANO est un modèle de rigueur, de concision et d'élégance :

Le système d'équations données se réduira à l'équation unique

$$\frac{dx}{dt} = Rx. \quad (1)$$

Soit a un complexe constant quelconque. Posons :

$$a' = \int Ra \, dt, \quad a'' = \int Ra' \, dt, \quad a''' = \int Ra'' \, dt, \dots,$$

où les intégrales s'étendent de t_0 à t . On doit prouver que la série

$$x = a + a' + a'' + \dots \quad (2)$$

²⁸Une analyse détaillée de ces recherches de Liouville et de Cauchy se trouve dans TOURNÈS, 1996, p. 268-285.

²⁹Voir DIEUDONNÉ, 1981 et ARCHIBALD, 1996.

³⁰PICARD, 1890.

³¹PEANO, 1888.

³²C'est vers 1843-45 que Cayley et Grassmann, voulant étendre le langage géométrique, considèrent les systèmes de n nombres en tant qu'éléments d'un nouvel espace. PEANO lui-même s'intéresse à la question en logicien et donne une définition générale, proche de la nôtre, des espaces vectoriels et des applications linéaires sur le corps des nombres réels dans *Calcolo geometrico secondo l'Ausdehnungslehre di Grassmann, preceduto dalle operazioni della logica deduttiva*, publié en 1888.

est convergente, et que sa somme x est une fonction de t (qui, évidemment, a la valeur a pour $t = t_0$) satisfaisant à l'équation (1). En effet puisque les r_{ij} sont des fonctions continues de t dans l'intervalle (p, q) , $\text{mod.}R$ sera aussi une fonction continue de t , et soit m son maximum. En supposant, pour simplifier, $t > t_0$, on a :

$$\begin{aligned}\text{mod.}a' &\leq \int \text{mod.}(Ra) dr \leq \int \text{mod.}R \cdot \text{mod.}a \cdot dt \leq m \cdot \text{mod.}a \cdot (t - t_0), \\ \text{mod.}a'' &\leq \int \text{mod.}(Ra') dt \leq \int \text{mod.}R \cdot \text{mod.}a' \cdot dt \leq \frac{1}{2}m^2 \cdot \text{mod.}a \cdot (t - t_0)^2, \\ &\text{etc.} \\ \text{mod.}a^{(n)} &\leq \frac{1}{n!}m^n \cdot \text{mod.}a \cdot (t - t_0)^n.\end{aligned}$$

Donc la série (2) est uniformément convergente, car les modules des termes sont moindres que des quantités constantes qui forment une série convergente. En différenciant les termes de (2) on obtient la série $0 + Ra + Ra' + \dots$, qui converge uniformément vers Rx . Donc x satisfait bien à l'équation proposée.

Il s'agit, pour la première fois, d'une démonstration d'existence ne laissant rien à désirer : convergence de la série en tout point de l'intervalle de régularité, justification correcte de la dérivation terme à terme pour prouver que la série est bien solution de l'équation. Ce travail de PEANO apparaît ainsi comme l'ultime avatar des idées semées vers 1830 par Liouville et Cauchy. La seule possibilité de dépasser la perfection atteinte par le mathématicien italien réside désormais dans l'abandon de l'hypothèse de linéarité, ce qui sera réalisé par Picard en 1890, ainsi que nous l'avons souligné plus haut.

e) Texte 9 : BANACH, 1920

Après PICARD, les approximations successives deviennent une méthode générale de résolution des équations fonctionnelles. La démarche de PICARD trouve son aboutissement dans la thèse³³ que soutient BANACH en 1920 et qui contient un théorème abstrait de point fixe :

THÉORÈME 6. Si

- 1° $U(X)$ est une opération continue dans E , le contre-domaine de $U(X)$ étant contenu dans E ;
- 2° il existe un nombre $0 < M < 1$ qui pour tout X' et X'' remplit l'inégalité

$$\|U(X') - U(X'')\| \leq M \|X' - X''\|,$$

il existe un élément X tel que $X = U(X)$.

Démonstration. Y désignant un élément choisi d'une façon arbitraire, soit $\{X_n\}$ une suite qui satisfait aux conditions :

$$X_1 = Y \text{ et pour tout } n, \quad X_{n+1} = U(X_n).$$

Nous allons démontrer que la suite $\{X_n\}$ converge suivant la norme vers un certain élément X . On observera dans ce but que l'on a pour tout $n > 1$:

³³BANACH, 1920.

$$\|X_{n+1} - X_n\| = \|U(X_n) - U(X_{n-1})\| \leq M \|X_n - X_{n-1}\|,$$

d'où

$$\|X_{n+1} - X_n\| \leq M^{n-1} \|X_2 - X_1\|.$$

On a par hypothèse $M < 1$; la série $\sum_{n=1}^{\infty} \|X_{n+1} - X_n\|$ est donc convergente, ce qui implique que la série $X_1 + \sum_{n=1}^{\infty} (X_{n+1} - X_n)$ converge suivant la norme vers un certain élément X .

Or,

$$X_1 + \sum_{n=1}^{n_1} (X_{n+1} - X_n) = X_{n_1}$$

donc

$$\overline{\lim}_{n \rightarrow \infty} X_n = X.$$

$U(X)$ étant continu, on a $\overline{\lim}_{n \rightarrow \infty} U(X_n) = U(X)$ et comme

$$X_n = U(X_{n-1}),$$

on trouve

$$\overline{\lim}_{n \rightarrow \infty} X_n = \overline{\lim}_{n \rightarrow \infty} U(X_{n-1})$$

et finalement

$$X = U(X).$$

c.q.f.d.

La preuve, identique à celle que l'on trouve aujourd'hui dans les traités d'analyse, n'appelle pas de remarque particulière. Par la considération d'espaces abstraits munis de normes adéquates, la résolution de l'équation fonctionnelle la plus complexe devient ainsi formellement identique à la résolution de l'équation numérique la plus simple. En quelque sorte, le théorème de BANACH détermine sous quelles conditions les procédés formels très généraux décrits plus haut par SCHMIDTEN peuvent être considérés comme recevables en analyse. À lui seul, ce théorème résume et justifie tous les algorithmes d'approximations successives utilisés pendant près de deux mille ans³⁴.

Bibliographie

- AABOE, A. Al-Kāshī's iteration method for the determination of $\sin 1^\circ$, *Scripta Mathematica*, **20** (1954), p. 24-29.
 ARCHIBALD, T. From attraction theory to existence proofs : The evolution of potential-theoretic methods in the study of boundary-value problems, 1860-1890, *Revue d'histoire des mathématiques*, **2** (1996), p. 67-93.
 BACHELARD, G. *Essai sur la connaissance approchée*, Paris : Vrin, 1973.
 BAILEY, D.F. A historical survey of solution by functional iteration, *Mathematics Magazine*, **62** (1989), p. 155-166.
 BANACH, S. *Sur les opérations dans les ensembles abstraits et leur application aux équations*

³⁴Pour suivre les développements de la méthode des approximations successives après la thèse de Banach, on pourra consulter KANTOROVITCH, 1939 et MAHWIN, 1994.

- intégrales, Thèse présentée en juin 1920 à l'Université de Léopol pour obtenir le grade de docteur en philosophie. Publiée dans *Fundamenta Mathematicae*, 3 (1922), p. 133-181.
- CAVEING, M. *La constitution du type mathématique de l'idéalité dans la pensée grecque*. Vol. 3 : *L'irrationalité dans les mathématiques grecques jusqu'à Euclide*, Villeneuve d'Ascq : Presses Universitaires du Septentrion, 1998.
- CHABERT, J.-L. et al. *Histoire d'algorithmes. Du caillou à la puce*, Paris : Belin, 1994. Trad. angl. par C. Weeks, *A history of algorithms. From the pebble to the microchip*, New York : Springer, à paraître.
- CHEMLA, K. Reflections on the world-wide history of the rule of false double position, or : How a loop was closed, *Centaurus*, 39 (1997), p. 97-120.
- DIEUDONNÉ, J. *History of functional analysis*, Amsterdam : North-Holland, 1981.
- GLUSHKOV, S. On approximation methods of Leonardo Fibonacci, *Historia Mathematica*, 3 (1976), p. 291-296.
- GREENBERG, J.L. Breaking a "vicious circle" : Unscrambling A.-C. Clairaut's iterative method of 1743, *Historia Mathematica*, 15 (1988), p. 228-239.
- HÉRON D'ALEXANDRIE *Heronis Alexandrini Opera quæ supersunt omnia*. Vol. III : *Metrica, Dioptra*, Ed. H. Schöne, Leipzig : Teubner, 1903.
- HOUZEL, C. Sharaf al-Dīn al-Tūsī et le polygone de Newton, *Arabic Sciences and Philosophy*, 5 (1995), p. 239-262.
- KANTOROVITCH, L. The method of successive approximations for functional equations, *Acta Mathematica*, 71 (1939), p. 63-97.
- KENNEDY, E. S. An early method of successive approximations, *Centaurus*, 13 (1969), p. 248-250.
- KENNEDY, E. S. & TRANSUE, W. R. A medieval iterative algorithm, *American Mathematical Monthly*, 63-2 (1956), p. 80-83.
- MANCHA, J.L. Heuristic reasoning : Approximation procedures in Levi ben Gerson's Astronomy, *Archive for History of Exact Sciences*, 52 (1998), p. 13-50.
- MAWHIN, J. Boundary value problems for non linear ordinary differential equations : from successive approximations to topology, dans J.-P. Pier (Ed.), *Development of mathematics 1900-1950*, Basel-Boston-Berlin : Birkhäuser, 1994.
- NEUGEBAUER, O. *The exact sciences in Antiquity*, 2^e Ed., Providence : Brown University Press, 1957; rééd. New York : Dover, 1969.
- A history of ancient mathematical astronomy*, 3 vol., New York-Heidelberg-Berlin : Springer, 1975.
- NEWTON, I. *Methodus fluxionum et serierum infinitarum*, 1671. Trad. angl. par J. Colson, *The method of fluxions and infinite series*, Londres, 1736. Trad. fr. par M. de Buffon, *La méthode des fluxions et des suites infinies*, Paris, 1740; rééd. Paris : Blanchard, 1994.
- PEANO, G. Integrazione per serie delle equazioni differenziali lineari, *Atti della Reale Accademia delle Scienze di Torino*, 22 (1887), p. 437-446. Intégration par séries des équations différentielles linéaires, *Mathematische Annalen*, 32 (1888), p. 450-456.
- PICARD, É. Mémoire sur la théorie des équations aux dérivées partielles et la méthode des approximations successives, *Journal de mathématiques pures et appliquées* (4^e série), 6 (1890), p. 145-210.
- PLOFKER, K. An example of the secant method of iterative approximation in a fifteenth century sanskrit text, *Historia Mathematica*, 23 (1996), p. 246-256.
- RASHED, R. *Entre arithmétique et algèbre. Recherches sur l'histoire des mathématiques arabes*, Paris : Les Belles Lettres, 1984. Trad. angl., *The development of arabic mathematics : Between arithmetic and algebra*, Dordrecht : Kluwer Academic Publishers (Boston Studies

- in the Philosophy of Science, vol. 156), 1994.
- Histoire des sciences arabes*. Vol. 2 : *Mathématiques et physique*, sous la direction de R. Rashed, Paris : Seuil, 1997.
- RIABI, F. An early iterative method for the determination of $\sin 1^\circ$, *The College Mathematics Journal*, 26-1 (1995), p. 16-21.
- SCAVO, T.R. & THOO, J. B. On the geometry of Halley's method, *American Mathematical Monthly*, 102 (1995), p. 417-426.
- SCHMIDTEN, H.G. Mémoire sur l'intégration des équations linéaires, *Annales de mathématiques pures et appliquées*, 11 (1820/21), p. 269-316.
- SPIESSER, M. *Équations du premier degré. Méthode de "fausse position"*, Toulouse : IREM, 1982.
- TATON, R. et WILSON, C. (Eds.) *Planetary astronomy from the Renaissance to the rise of astrophysics. Part B : The eighteenth and nineteenth centuries*, Cambridge : University Press, 1995.
- TOURNÈS, D. *L'intégration approchée des équations différentielles ordinaires (1671-1914)*, Thèse de doctorat de l'université Paris 7-Denis Diderot (juin 1996), Villeneuve d'Ascq : Presses Universitaires du Septentrion, 1997.
- L'origine des méthodes multipas pour l'intégration numérique des équations différentielles ordinaires, *Revue d'histoire des mathématiques*, 4 (1998), p. 5-72.
- YOUSCHKEVITCH, A.P. *Les mathématiques arabes (VIII^{ème}-XV^{ème} siècles)*, trad. par M. Cazenave et K. Jaouiche, Paris : Vrin, 1976.
- YPMMA, T.J. Historical development of the Newton-Raphson method, *SIAM Review*, 37 (1995), p. 531-551.

