

Cours d'analyse numérique

SMI-S4

Introduction

L'objet de l'analyse numérique est de concevoir et d'étudier des méthodes de résolution de certains problèmes mathématiques, en général issus de problèmes réels, et dont on cherche à calculer la solution à l'aide d'un ordinateur.

Exemples de problèmes à résoudre:

- Systèmes linéaires ou non linéaires.
 - Optimisation.
 - Equations différentielles ou aux dérivées partielles.
- etc. . .

Références pour le cours:

(1) P. G. **Ciarlet**

Introduction à l'analyse numérique matricielle et à l'optimisation.

(2) A. **Quarteroni**

Méthodes numériques pour le calcul scientifique.

(3) M. **Sibony** & J. Cl. **Mardon**

Analyse numérique 1: Systèmes linéaires et non linéaires.

Analyse numérique 2: Approximations et équations différentielles.

1 Résolution de systèmes linéaires

Objectifs:

On note par $\mathcal{M}_N(\mathbb{R})$ l'ensemble des matrices carrées d'ordre N .

Soit $A \in \mathcal{M}_N(\mathbb{R})$ une matrice inversible et $b \in \mathbb{R}^N$.

On cherche à résoudre le système linéaire:

trouver $x \in \mathbb{R}^N$, tel que $Ax = b$.

Nous allons voir deux types de méthodes:

- Méthodes directes,
- Méthodes itératives.

Selon le type et la taille d'une matrice, on utilise une méthode directe ou une méthode itérative.

1.1 Quelques rappels d'algèbre linéaire

1.1.1 Matrices

- La transposée, d'une matrice $A = (a_{i,j})_{1 \leq i,j \leq N}$, est une matrice

$A^T = (a'_{i,j})_{1 \leq i,j \leq N}$, avec $a'_{i,j} = a_{j,i}$ pour $1 \leq i, j \leq N$.

On a évidemment les propriétés suivantes:

$$\begin{aligned}(A^T)^T &= A; \\ (A+B)^T &= A^T + B^T; \\ (AB)^T &= B^T A^T; \\ (\alpha A)^T &= \alpha A^T, \forall \alpha \in \mathbb{R}; \\ \det(A^T) &= \det(A). \\ \text{Si } A \text{ est inversible } (A^T)^{-1} &= (A^{-1})^T.\end{aligned}$$

- La matrice adjointe de A est $A^* = \overline{A^T}$. on a alors les propriétés:

$$\begin{aligned}(A^*)^* &= A; \\ (A+B)^* &= A^* + B^*; \\ (AB)^* &= B^* A^*; \\ (\alpha A)^* &= \overline{\alpha} A^*, \forall \alpha \in \mathbb{C}; \\ \det(A^*) &= \det(A); \\ \text{Si } A \text{ est inversible } (A^*)^{-1} &= (A^{-1})^*.\end{aligned}$$

- Une matrice $A \in \mathcal{M}_N(\mathbb{R})$ est *symétrique* si $A^T = A$.

- Une matrice $A \in \mathcal{M}_N(\mathbb{C})$ est *hermitienne* si $A^* = A$.

- Une matrice $A \in \mathcal{M}_N(\mathbb{R})$ est *orthogonale* si $A^T A = A A^T = I_N$.
(c'est à dire $A^T = A^{-1}$)

- Une matrice $A \in \mathcal{M}_N(\mathbb{C})$ est *unitaire* si $A^* A = A A^* = I_N$.
(c'est à dire $A^* = A^{-1}$).

On dit que A est *normale* si $A^* A = A A^*$.

- Si $A = (a_{i,j})_{1 \leq i,j \leq N}$ la trace de A est définie par $tr(A) = \sum_{i=1}^N a_{i,i}$.

1.1.2 Normes matricielles

Soit E un espace vectoriel sur $\mathbb{R}$.

On définit $\|\cdot\|$ une norme vectorielle sur E , c'est-à-dire une application: $E \rightarrow \mathbb{R}^+$ qui vérifie les propriétés suivantes:

- (i) $\forall x \in E, \|x\| = 0 \iff x = 0$.
- (ii) $\forall \lambda \in \mathbb{R}, \forall x \in E, \|\lambda x\| = |\lambda| \cdot \|x\|$.
- (iii) $\forall x, y \in E, \|x + y\| \leq \|x\| + \|y\|$.

Exemples de normes vectorielles sur $\mathbb{R}^N$:

Soit $x = (x_1, x_2, \dots, x_N) \in \mathbb{R}^N$

- 1) $x \rightarrow \|x\|_1 = |x_1| + |x_2| + \dots + |x_N|$
- 2) $x \rightarrow \|x\|_2 = \sqrt{|x_1|^2 + |x_2|^2 + \dots + |x_N|^2}$
- 3) $x \rightarrow \|x\|_\infty = \max(|x_1|, |x_2|, \dots, |x_N|)$

sont trois normes vectorielles sur $\mathbb{R}^N$.

Définition 2:

Soit $\|\cdot\|$ une norme vectorielle sur $E = \mathbb{R}^N$.

On appelle **une** norme matricielle induite, par cette norme vectorielle, sur $\mathcal{M}_N(\mathbb{R})$, qu'on note encore par $\|\cdot\|$:

$$A \rightarrow \|A\| = \sup\{\|Ax\|, x \in \mathbb{R}^N, \|x\| = 1\}.$$

Remarque:

on dit aussi que c'est une norme matricielle *subordonnée* ou associée à la norme vectorielle.

Proposition 1 Si $\|\cdot\|$ est une norme matricielle induite (par une norme vectorielle $\|\cdot\|$ sur $E = \mathbb{R}^N$) sur $\mathcal{M}_N(\mathbb{R})$, alors on a les propriétés suivantes:

- (1) $\forall x \in \mathbb{R}^N$ et $\forall A \in \mathcal{M}_N(\mathbb{R})$,
 $\|Ax\| \leq \|A\| \cdot \|x\|$.
- (2) $\|A\| = \max\{\|Ax\|, x \in \mathbb{R}^N, \|x\| = 1\}$.
- (3) $\|A\| = \max\{\frac{\|Ax\|}{\|x\|}, x \in \mathbb{R}^N, x \neq 0\}$.

Démonstration

(1) Soit $x \in \mathbb{R}^N$ non nul,

$$y = \frac{1}{\|x\|}x \implies \|y\| = 1$$

$$\implies \|Ay\| \leq \|A\| \quad (\text{par définition})$$

$$\iff \|A \frac{1}{\|x\|}x\| \leq \|A\|$$

$$\iff \frac{\|Ax\|}{\|x\|} \leq \|A\|$$

$$\iff \|Ax\| \leq \|A\| \cdot \|x\|.$$

Si $x = 0$ alors $Ax = 0$ et l'inégalité

$$\|Ax\| \leq \|A\| \cdot \|x\| \text{ est encore vérifiée.}$$

(2) $\|A\| = \sup\{\|Ax\|, x \in \mathbb{R}^N, \|x\| = 1\}$

Comme l'application $x \rightarrow \|x\|$ est continue sur $\{x \in \mathbb{R}^N, \|x\| = 1\}$

qui est un compact de $\mathbb{R}^N$,

$\exists x_0 \in \{x \in \mathbb{R}^N, \|x\| = 1\}$ tel que:

$$\|A\| = \|Ax_0\| = \max\{\|Ax\|, x \in \mathbb{R}^N, \|x\| = 1\}.$$

(3) Si x est non nul,

$$\frac{\|Ax\|}{\|x\|} = \|A(\frac{x}{\|x\|})\|$$

et $\frac{x}{\|x\|} \in \{x \in \mathbb{R}^N, \|x\| = 1\}$.

Proposition

Pour toute matrice $A = (a_{i,j})_{1 \leq i,j \leq N}$, on a

$$1) \|A\|_1 = \sup_{\|x\|_1=1} \frac{\|Ax\|_1}{\|x\|_1} = \max_{1 \leq j \leq N} \sum_{i=1}^N |a_{i,j}|.$$

$$2) \|A\|_\infty = \sup_{\|x\|_\infty=1} \frac{\|Ax\|_\infty}{\|x\|_\infty} = \max_{1 \leq i \leq N} \sum_{j=1}^N |a_{i,j}|.$$

Exemple

Si $A = \begin{pmatrix} 1 & -2 \\ 3 & 0 \end{pmatrix};$

$\|A\|_1 = \max(4, 2) = 4$ et $\|A\|_\infty = 3$.

1.1.3 Rayon spectral

Définition:

Soit $A \in \mathcal{M}_N(\mathbb{R})$ et $\lambda \in \mathbb{C}$.

On dit que λ est une *valeur propre* de A , s'il existe un vecteur $u \in \mathbb{R}^N, u \neq 0$ et tel que: $Au = \lambda u$.

On dit alors que u est un vecteur propre de A associé à λ .

Proposition 2 On a l'équivalence suivante:

$$\lambda \text{ est une valeur propre de } A \iff \det(A - \lambda I) = 0.$$

Les valeurs propres de A sont donc les racines du polynôme

$P_A(x) = \det(A - xI)$, appelé le *polynôme caractéristique* de A .

C'est un polynôme de degré N , il a donc N racines dans $\mathbb{C}$ (non nécessairement distinctes).

La matrice A a donc N valeurs propres dans $\mathbb{C}$ (non nécessairement distinctes): $\lambda_1, \lambda_2, \dots, \lambda_N$.

On montre qu'on a : $\boxed{\det(A) = \prod_{i=1}^N \lambda_i \text{ et } \operatorname{tr}(A) = \sum_{i=1}^N \lambda_i}$

On montre qu'on a :

Propriété:

Si $A \in \mathcal{M}_N(\mathbb{R})$ est une matrice *symétrique*

alors toutes ses *valeurs propres* sont *réelles*

et il existe une matrice *orthogonale* P ($P^{-1} = P^T$) telle que:

$$P^{-1}AP = \begin{pmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \ddots & & \vdots \\ \vdots & & \ddots & \\ 0 & 0 & \dots & \lambda_N \end{pmatrix} = D$$

où les λ_i sont les valeurs propres de A .

A est donc semblable à la matrice diagonale D .

Définition: On dit que deux matrices A et B sont semblables s'il existe une matrice inversible P telle que $A = P^{-1}BP$.

Exemple:

$$A = \begin{pmatrix} 0 & 1 & 1 \\ 1 & 0 & 1 \\ 1 & 1 & 0 \end{pmatrix}$$

$$x^3 - 3x - 2 = (x - 2)(x + 1)^2,$$

Les valeurs propres de A sont: $2, -1, -1$.

Les vecteurs propres sont:

$$\text{Pour } \lambda_1 = -1 : \alpha \begin{pmatrix} -1 \\ 0 \\ 1 \end{pmatrix} + \beta \begin{pmatrix} 1 \\ -2 \\ 1 \end{pmatrix},$$

avec $\alpha, \beta \in \mathbb{R}, (\alpha, \beta) \neq (0, 0)$.

$$\text{Pour } \lambda_2 = 2 : \alpha \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}, \text{ avec } \alpha \in \mathbb{R}^*.$$

$$\text{Si } P = \begin{pmatrix} -1 & 1 & 1 \\ 0 & -2 & 1 \\ 1 & 1 & 1 \end{pmatrix}, \text{ on a } P^{-1} = \begin{pmatrix} -\frac{1}{2} & 0 & \frac{1}{2} \\ \frac{1}{6} & -\frac{1}{3} & \frac{1}{6} \\ \frac{1}{3} & \frac{1}{3} & \frac{1}{3} \end{pmatrix}$$

$$P^{-1}AP = \begin{pmatrix} -\frac{1}{2} & 0 & \frac{1}{2} \\ \frac{1}{6} & -\frac{1}{3} & \frac{1}{6} \\ \frac{1}{3} & \frac{1}{3} & \frac{1}{3} \end{pmatrix} \begin{pmatrix} 0 & 1 & 1 \\ 1 & 0 & 1 \\ 1 & 1 & 0 \end{pmatrix} \begin{pmatrix} -1 & 1 & 1 \\ 0 & -2 & 1 \\ 1 & 1 & 1 \end{pmatrix}$$

$$= \begin{pmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 2 \end{pmatrix}$$

$$\text{Remarque: Si on prend } P' = \begin{pmatrix} -\frac{1}{\sqrt{2}} & \frac{1}{\sqrt{6}} & \frac{1}{\sqrt{3}} \\ 0 & -\frac{2}{\sqrt{6}} & \frac{1}{\sqrt{3}} \\ \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{6}} & \frac{1}{\sqrt{3}} \end{pmatrix},$$

son inverse: $P^{-1} = P^T$ et on a:

$$P'^T A P' = \begin{pmatrix} -\frac{1}{\sqrt{2}} & 0 & \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{6}} & -\frac{2}{\sqrt{6}} & \frac{1}{\sqrt{6}} \\ \frac{1}{\sqrt{3}} & \frac{1}{\sqrt{3}} & \frac{1}{\sqrt{3}} \end{pmatrix} \begin{pmatrix} 0 & 1 & 1 \\ 1 & 0 & 1 \\ 1 & 1 & 0 \end{pmatrix} \begin{pmatrix} -\frac{1}{\sqrt{2}} & \frac{1}{\sqrt{6}} & \frac{1}{\sqrt{3}} \\ 0 & -\frac{2}{\sqrt{6}} & \frac{1}{\sqrt{3}} \\ \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{6}} & \frac{1}{\sqrt{3}} \end{pmatrix}$$

$$= \begin{pmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 2 \end{pmatrix}.$$

Définition

Le rayon spectral de A est

$$\rho(A) = \max\{|\lambda|, \lambda \in \mathbb{C}, \lambda \text{ valeur propre de } A\}.$$

Proposition 3 Soit $A \in \mathcal{M}_N(\mathbb{R})$, alors pour toute norme matricielle, (induite ou non) on a:

$$\rho(A) \leq \|A\|$$

Démonstration

Soit λ une valeur propre de A telle que : $\rho(A) = |\lambda|$.

$\exists p \in \mathbb{R}^N$ non nul, tel que $Ap = \lambda p$.

p vecteur non nul $\implies \exists q \in \mathbb{R}^N$ tel que $pq^T \neq 0$.

$$\begin{aligned} \implies \rho(A) \|pq^T\| &= |\lambda| \cdot \|pq^T\| = \|(\lambda p)q^T\| \\ &= \|(Ap)q^T\| = \|A(pq^T)\| \\ &\leq \|A\| \cdot \|pq^T\| \end{aligned}$$

$$\implies \rho(A) \leq \|A\|.$$

On montre aussi le résultat suivant:

Proposition 4 $\forall \epsilon > 0$ et $A \in \mathcal{M}_N(\mathbb{R})$, il existe au moins une norme matricielle subordonnée telle que:

$$\|A\| \leq \rho(A) + \epsilon.$$

On montre aussi la proposition suivante:

Proposition 5 Soit $A = (a_{i,j})_{1 \leq i,j \leq N} \in \mathcal{M}_N(\mathbb{R})$, alors on a

$$\|A\|_2 = \sup_{\|x\|_2=1} \frac{\|Ax\|_2}{\|x\|_2} = \sqrt{\rho(A^T A)} = \sqrt{\rho(AA^T)}.$$

Théorème

On désigne par I_N la matrice identité de $\mathcal{M}_N(\mathbb{R})$.

1) Soit $\|\cdot\|$ une norme matricielle induite sur $\mathcal{M}_N(\mathbb{R})$.

Si $A \in \mathcal{M}_N(\mathbb{R})$ est telle que $\|A\| < 1$, alors

la matrice $I_N + A$ est inversible et

$$\|(I_N + A) - 1\| \leq \frac{1}{1 - \|A\|}.$$

2) Si $I_N + A$ est singulière, alors $\|A\| \geq 1$, pour toute norme matricielle sur $\mathcal{M}_N(\mathbb{R})$.

Démonstration:

1) $(I_N + A)x = 0$

$$\implies \|x\| = \|Ax\| \leq \|A\| \cdot \|x\|$$

Comme $\|A\| < 1$, on a $x = 0$.

$\implies I_N + A$ est inversible et on a

$$\text{Comme } (I_N + A)^{-1} = I_N - A(I_N + A)^{-1}$$

on a $\|(I_N + A)^{-1}\| \leq 1 + \|A\| \cdot \|(I_N + A)^{-1}\|$

$$\implies (1 - \|A\|) \cdot \|(I_N + A)^{-1}\| \leq 1$$

$$\implies \|(I_N + A)^{-1}\| \leq \frac{1}{1 - \|A\|}.$$

2) $I_N + A$ singulière $\iff \det(I_N + A) = 0$

$$\iff \lambda = -1 \text{ est une valeur propre de } A$$

$$\implies 1 \leq \rho(A) \leq \|A\|.$$

1.1.4 Produit scalaire et matrices définies positives

Définition Si E est un espace vectoriel sur $\mathbb{K} = \mathbb{R}$ ou $\mathbb{C}$, un produit scalaire sur E est une application de

$$E \times E \rightarrow \mathbb{K}$$

$$(x, y) \mapsto \langle x, y \rangle$$

qui vérifie les propriétés suivantes:

(i) $\langle x, x \rangle \geq 0$, $\forall x \in E$ et $\langle x, x \rangle = 0 \implies x = 0$.

(ii) $\langle \alpha(x + y), z \rangle = \alpha \langle x, z \rangle + \alpha \langle y, z \rangle$, $\forall \alpha \in \mathbb{K}$, $\forall x, y, z \in E$

(iii) $\langle y, x \rangle = \overline{\langle x, y \rangle}, \forall x, y \in E$.

Exemple: $\mathbb{K} = \mathbb{R}, E = \mathbb{R}^N$, pour $x, y \in E$,
 $x^T = (x_1, \dots, x_N)$ et $y^T = (y_1, \dots, y_N)$,

L'application $(x, y) \rightarrow \langle x, y \rangle = y^T x = \sum_{i=1}^N x_i y_i$

définit un produit scalaire sur $\mathbb{R}^N$.

Définition

Une matrice $A \in \mathcal{M}_N(\mathbb{R})$ est définie positive si

$$(Ax, x) > 0, \forall x \in \mathbb{R}^N.$$

Si l'inégalité stricte est remplacée par une inégalité au sens large (≥ 0), on dit que la matrice est semi-définie positive.

Définition: les mineurs principaux dominants de $A = (a_{ij})_{1 \leq i, j \leq N}$ sont les N déterminants des sous-matrices de A :

$$(a_{ij})_{1 \leq i, j \leq k}, (1 \leq k \leq N).$$

Proposition 6 Soit $A \in \mathcal{M}_N(\mathbb{R})$ une matrice symétrique.

A est définie positive si et seulement si l'une des propriétés suivantes est satisfaite:

- (1) $(Ax, x) > 0, \forall x \in \mathbb{R}^N, x \neq 0$.
- (2) Les valeurs propres de A sont > 0 .
- (3) Les mineurs principaux dominants de A sont tous > 0 .

1.2 Méthodes directes

Une méthode directe de résolution d'un système linéaire est une méthode qui calcule x , la solution exacte du système, après un nombre fini d'opérations élémentaires $(+, -, \times, /)$.

1.2.1 Méthode de Gauss et factorisation LU

Méthode de Gauss Soit $A \in \mathcal{M}_N(\mathbb{R})$ une matrice inversible et $b \in \mathbb{R}^N$.

pb: trouver $x \in \mathbb{R}^N, Ax = b$.

Méthode de Gauss:

On pose $A^{(1)} = A$ et $b^{(1)} = b$.

On construit une suite de matrices et de vecteurs:

$$A^{(1)} \rightarrow A^{(2)} \rightarrow \dots \rightarrow A^{(k)} \rightarrow \dots \rightarrow A^{(N)}.$$

$$b^{(1)} \rightarrow b^{(2)} \rightarrow \dots \rightarrow b^{(k)} \rightarrow \dots \rightarrow b^{(N)}.$$

Les systèmes linéaires obtenus sont équivalents:

$$Ax = b \iff A^{(k)}x = b^{(k)} \iff A^{(N)}x = b^{(N)},$$

La matrice $A^{(N)}$ est triangulaire supérieure.

Si aucun pivot n'est nul:

Passage de $A^{(k)}$ à $A^{(k+1)}$, si $a_{k,k}^{(k)} \neq 0$:

- Pour $i \leq k$ et pour $j = 1, 2, \dots, N$
 $a_{i,j}^{(k+1)} = a_{i,j}^{(k)}$ et $b_i^{(k+1)} = b_i^{(k)}$.

- Pour $i > k$:

$$a_{i,j}^{(k+1)} = a_{i,j}^{(k)} - \frac{a_{i,k}^{(k)}}{a_{k,k}^{(k)}} a_{k,j}^{(k)},$$

$$b_i^{(k+1)} = b_i^{(k)} - \frac{a_{i,k}^{(k)}}{a_{k,k}^{(k)}} b_k^{(k)},$$

pour $i = j, \dots, N$.

$$A^{(k)} = \begin{pmatrix} a_{1,1}^{(1)} & a_{1,2}^{(1)} & \dots & \dots & \dots & \dots & a_{1,N}^{(1)} \\ 0 & a_{2,2}^{(2)} & a_{2,3}^{(2)} & \dots & \dots & \dots & a_{2,N}^{(2)} \\ \vdots & \ddots & \ddots & & & & \vdots \\ \vdots & & 0 & a_{k,k}^{(k)} & \dots & \dots & a_{k,N}^{(k)} \\ \vdots & \vdots & a_{k+1,k}^{(k)} & \ddots & & & \vdots \\ \vdots & \vdots & \vdots & & \ddots & & \vdots \\ 0 & \dots & 0 & a_{N,k}^{(k)} & \dots & \dots & a_{N,N}^{(k)} \end{pmatrix}, 1 \leq k \leq N.$$

La matrice $A^{(N)}$ est alors sous la forme:

$$A^{(N)} = \begin{pmatrix} a_{1,1}^{(1)} & a_{1,2}^{(1)} & \dots & \dots & \dots & \dots & a_{1,N}^{(1)} \\ 0 & a_{2,2}^{(2)} & a_{2,3}^{(2)} & \dots & \dots & \dots & a_{2,N}^{(2)} \\ \vdots & \ddots & \ddots & & & & \vdots \\ \vdots & & \ddots & a_{k,k}^{(k)} & \dots & \dots & a_{k,N}^{(k)} \\ \vdots & \vdots & \vdots & \ddots & \ddots & \ddots & \vdots \\ \vdots & \vdots & \vdots & \vdots & \ddots & \ddots & \vdots \\ 0 & \dots & \dots & \dots & \dots & 0 & a_{N,N}^{(N)} \end{pmatrix}$$

La résolution du système linéaire $A^{(N)}x = b^{(N)}$ se fait en remontant et on calcule successivement: $x_N, x_{N-1}, \dots$ et x_1 :

$$\begin{cases} x_N = \frac{b_N^{(N)}}{a_{N,N}^{(N)}}, \\ x_i = \frac{1}{a_{i,i}^{(N)}} (b_i^{(N)} - \sum_{j=i+1}^N a_{i,j}^{(N)} x_j), \quad i = N-1, \dots, 1. \end{cases}$$

Remarques: $\det(A) = \det(A^{(k)}) = \det(A^{(N)})$.

Cas d'un pivot nul:

Si à l'étape k , le pivot $a_{k,k} = 0$, on peut permuter la ligne k avec une ligne i_0 , telle que $i_0 > k$ et $a_{i_0,k} \neq 0$ et on continue la méthode de Gauss.

Coût de la méthode de Gauss pour résoudre le système linéaire $Ax = b$

1) Pour l'élimination: $A \rightarrow A^{(N)}$

Nombre d'additions:

$$(N-1)^2 + (N-2)^2 + \dots + 1^2 = \frac{N(N-1)(2N-1)}{6};$$

$$\text{Nombre de multiplications: } \frac{N(N-1)(2N-1)}{6};$$

$$\text{Nombre de divisions: } (N-1) + (N-2) + \dots + 1 = \frac{N(N-1)}{2}.$$

2) Pour passer de $b \rightarrow b^{(N)}$:

$(N-1) + (N-2) + \dots + 1 = \frac{N(N-1)}{2}$ additions
et $\frac{N(N-1)}{2}$ multiplications.

3) Pour la remontée:

$\frac{N(N-1)}{2}$ additions et $\frac{N(N-1)}{2}$ multiplications et N divisions.

Au total, l'ordre du nombre d'opérations élémentaires nécessaires à la méthode de Gauss est:

$\frac{N^3}{3}$ additions, $\frac{N^3}{3}$ multiplications et $\frac{N^2}{2}$ divisions.

Comme le calcul direct d'un déterminant nécessite:

$N! - 1$ additions et $(N-1)N!$ multiplications,

la méthode de Cramer va nécessiter

$(N+1)!$ additions, $(N+2)!$ multiplications et N divisions

Par exemple pour $N = 10$:

La méthode de Gauss nécessite : 700 opérations,

et la méthode de Cramer: 500 000 000 opérations!

$(11! + 12! = 39\,916\,800 + 479\,001\,600 = 5.1892 \times 10^8)$.

Stratégies de pivot Plusieurs stratégies de choix de i_0 sont possibles:

- Stratégie du pivot partiel:

On peut choisir i_0 tel que:

$$|a_{i_0,k}| = \max\{|a_{i,k}|, i = k, \dots, N\}.$$

- Stratégie du pivot total:

On choisit i_0, j_0 tels que:

$$|a_{i_0,j_0}| = \max\{|a_{i,j}|, i, j = k, \dots, N\},$$

et on permute

$$\begin{cases} \text{la ligne } k \text{ avec la ligne } i_0 \\ \text{et} \\ \text{la colonne } k \text{ avec la colonne } j_0 \end{cases}$$

Factorisation LU Supposons que dans l'élimination de Gauss on n'utilise aucune stratégie de pivotage et que tous les pivots $a_{k,k}^{(k)} \neq 0$.

Dans ce cas le passage de $A^{(k)} \rightarrow A^{(k+1)}$ ($1 \leq k \leq N-1$) revient à multiplier à gauche la matrice $A^{(k)}$ par la matrice $N \times N$:

$$E^{(k)} = \begin{pmatrix} 1 & 0 & \dots & \dots & \dots & \dots & 0 \\ 0 & 1 & 0 & \dots & \dots & \dots & \\ \vdots & \ddots & \ddots & & & & \vdots \\ \vdots & & 0 & 1 & \dots & \dots & 0 \\ \vdots & & \vdots & -l_{k+1,k} & \ddots & & \vdots \\ \vdots & & \vdots & \vdots & 0 & \ddots & \vdots \\ 0 & \dots & 0 & -l_{N,k} & 0 & \dots & 1 \end{pmatrix}$$

avec $l_{i,k} = \frac{a_{i,k}^{(k)}}{a_{k,k}^{(k)}}$, pour $k+1 \leq i \leq N$.

La matrice $A^{(N)}$, qui est triangulaire supérieure, est alors égale à

$A^{(N)} = MA$
 avec $M = E^{(N-1)}E^{(N-2)}\dots E^{(1)}$
 M est le produit de matrices triangulaires inférieures, donc M est aussi
 triangulaire inférieure, on a $\det(M) = \prod_{i=1}^{N-1} \det(E^{(i)}) = 1$ et l'inverse de M est
 aussi triangulaire inférieure.

En posant $U = A^{(N)}$ et $L = M^{-1}$, on a $A = LU$.

$$(E^{(k)})^{-1} = \begin{pmatrix} 1 & 0 & \dots & \dots & \dots & \dots & 0 \\ 0 & 1 & 0 & \dots & \dots & \dots & \\ \vdots & \ddots & \ddots & \ddots & & & \vdots \\ \vdots & & 0 & 1 & \ddots & \dots & 0 \\ \vdots & & \vdots & l_{k+1,k} & \ddots & \ddots & \vdots \\ \vdots & & \vdots & \vdots & 0 & \ddots & 0 \\ 0 & \dots & 0 & l_{N,k} & 0 & \dots & 1 \end{pmatrix}$$

$$L = M^{-1} = (E^{(1)})^{-1}(E^{(2)})^{-1}\dots(E^{(N-1)})^{-1}$$

$$L = \begin{pmatrix} 1 & 0 & \dots & \dots & \dots & \dots & 0 \\ l_{2,1} & 1 & 0 & \dots & \dots & \dots & \\ \vdots & \ddots & \ddots & & & & \vdots \\ \vdots & & \ddots & 1 & \dots & \dots & 0 \\ \vdots & & \vdots & l_{k+1,k} & \ddots & & \vdots \\ \vdots & & \vdots & \vdots & & \ddots & \vdots \\ l_{N,1} & \dots & \dots & l_{N,k} & \dots & \dots & l_{N,N-1} & 1 \end{pmatrix}$$

Exemple:

$$A = \begin{pmatrix} 5 & 2 & 1 \\ 5 & -6 & 2 \\ -4 & 2 & 1 \end{pmatrix} = A^{(1)}$$

$$E^{(1)} = \begin{pmatrix} 1 & 0 & 0 \\ -1 & 1 & 0 \\ \frac{4}{5} & 0 & 1 \end{pmatrix}; A^{(2)} = \begin{pmatrix} 5 & 2 & 1 \\ 0 & -8 & 1 \\ 0 & \frac{18}{5} & \frac{9}{5} \end{pmatrix}$$

$$E^{(2)} = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & \frac{9}{20} & 1 \end{pmatrix}; A^{(3)} = \begin{pmatrix} 5 & 2 & 1 \\ 0 & -8 & 1 \\ 0 & 0 & \frac{9}{4} \end{pmatrix} = U$$

$$L = \begin{pmatrix} 1 & 0 & 0 \\ 1 & 1 & 0 \\ -\frac{4}{5} & -\frac{9}{20} & 1 \end{pmatrix} \text{ et } A = LU$$

Théorème

Soit $A = (a_{i,j})_{1 \leq i,j \leq N}$ une matrice carrée d'ordre N telle que les N sous-matrices de A :

$$\begin{pmatrix} a_{11} & \cdots & \cdots & a_{1k} \\ \vdots & & & \vdots \\ \vdots & & & \vdots \\ a_{k1} & \cdots & \cdots & a_{kk} \end{pmatrix}, 1 \leq k \leq N$$

soient inversibles,

alors il existe une matrice triangulaire inférieure $L = (l_{ii})_{1 \leq i \leq N}$, avec $l_{ii} = 1$ ($1 \leq i \leq N$), et une matrice triangulaire supérieure U telles que $A = LU$.

De plus cette factorisation est unique.

Démonstration

Il suffit de montrer qu'aucun pivot n'est nul.

1.2.2 Matrice symétrique définie positive: Méthode de Cholesky

Dans le cas où la matrice A est symétrique définie positive, la condition du théorème précédent est satisfaite et elle admet une factorisation $A = LU$, avec

$$L = \begin{pmatrix} 1 & 0 & 0 \\ l_{2,1} & \ddots & \vdots \\ \vdots & & 0 \\ l_{N,1} & l_{N,N-1} & 1 \end{pmatrix} \text{ et } U = \begin{pmatrix} u_{1,1} & \cdots & u_{1,N} \\ 0 & \ddots & \vdots \\ \vdots & & \vdots \\ 0 & 0 & u_{N,N} \end{pmatrix}$$

Comme la matrice A est définie positive, on a

$$\prod_{i=1}^k u_{i,i} = \Delta_k > 0, k = 1, 2, \dots, N,$$

(les Δ_k sont les mineurs principaux dominants de A)

et $u_{i,i} > 0$, pour $i = 1, 2, \dots, N$.

$$\text{Si on pose } D = \begin{pmatrix} \sqrt{u_{1,1}} & 0 & \cdots & 0 \\ 0 & & \ddots & \vdots \\ \vdots & & & 0 \\ 0 & 0 & \sqrt{u_{N,N}} \end{pmatrix}$$

Cette matrice est inversible et son inverse est

$$D^{-1} = \begin{pmatrix} \frac{1}{\sqrt{u_{1,1}}} & 0 & \cdots & 0 \\ 0 & & \ddots & \vdots \\ \vdots & & & 0 \\ 0 & 0 & \frac{1}{\sqrt{u_{N,N}}} \end{pmatrix}$$

$$A = LU = (LD)(D^{-1}U) = RB^T$$

La matrice $R = LD$ est triangulaire inférieure et $B^T = D^{-1}U$ est triangulaire supérieure et elles sont toutes les deux inversibles.

Comme A est symétrique, on a $A^T = A$.

$$\Rightarrow RB^T = BR^T \Rightarrow B^{-1}R = R^T(B^T)^{-1}$$

De plus $B^{-1}R$ est une matrice triangulaire inférieure et $R^T(B^T)^{-1}$ est une matrice triangulaire supérieure, ce qui implique que

$$B^{-1}R = R^T(B^T)^{-1} = \text{matrice diagonale}$$

Or les éléments diagonaux: $(B^{-1}R)_{i,i} = 1$

$$\implies B^{-1}R = R^T(B^T)^{-1} = I_N$$

$$\implies B = R \text{ et } A = RR^T.$$

Comme $R = LD = (r_{i,j}) \implies r_{i,i} = \sqrt{u_{i,i}} > 0$, pour $1 \leq i \leq N$.

Unicité de R :

Supposons que $A = RR^T = BB^T$, avec $R = (r_{i,j})$ et $B = (b_{i,j})$ des matrices triangulaires inférieures avec des éléments diagonaux $r_{ii} > 0$ et $b_{i,i} > 0$, pour $1 \leq i \leq N$.

$$RR^T = BB^T \implies R^T(B^T)^{-1} = R^{-1}B$$

$$\implies R^T(B^T)^{-1} = R^{-1}B = D \text{ matrice diagonale,}$$

$$\text{et } \frac{r_{i,i}}{b_{i,i}} = \frac{b_{i,i}}{r_{i,i}} \implies r_{i,i}^2 = b_{i,i}^2, 1 \leq i \leq N \implies r_{i,i} = b_{i,i}, 1 \leq i \leq N.$$

$$\text{et alors } d_{ii} = \frac{b_{i,i}}{r_{i,i}} = 1, 1 \leq i \leq N \text{ et } R = B.$$

On a donc montré le théorème suivant:

Théorème

Une matrice A est symétrique définie positive *si et seulement si* il existe une matrice triangulaire inférieure inversible R telle que: $A = RR^T$.

Si on impose que les éléments diagonaux de R , $r_{ii} > 0, 1 \leq i \leq N$, cette décomposition est unique.

Calcul effectif de la matrice R :

$$A = RR^T, R = (r_{i,j}) \implies a_{i,j} = \sum_{k=1}^i r_{i,k}r_{j,k} \text{ pour } 1 \leq i, j \leq N$$

$\boxed{i=1}$ on détermine la 1^{ère} colonne de R :

$$\rightarrow j = 1 : a_{1,1} = r_{1,1}^2 \implies r_{1,1} = \sqrt{a_{1,1}};$$

$$\rightarrow j = 2 : a_{1,2} = r_{1,1}r_{2,1} \implies r_{2,1} = \frac{a_{1,2}}{r_{1,1}};$$

$\vdots$

$$\rightarrow j = N : a_{1,N} = r_{1,1}r_{N,1} \implies r_{N,1} = \frac{a_{1,N}}{r_{1,1}};$$

De proche en proche, on détermine la $i^{\text{ème}}$ colonne de R ($i=2, \dots, N$):

$\boxed{i}$

$$\rightarrow j = i : a_{i,i} = r_{i,1}^2 + \dots + r_{i,i-1}^2$$

$$\implies r_{i,i} = \sqrt{a_{i,i} - (r_{i,1}^2 + \dots + r_{i,i-1}^2)};$$

$$\rightarrow j = i+1 : a_{i,i+1} = r_{i,1}r_{i+1,1} + \dots + r_{i,i}r_{i+1,i}$$

$$\implies r_{i+1,i} = \frac{a_{i,i+1} - (r_{i,1}r_{i+1,1} + \dots + r_{i,i-1}r_{i+1,i-1})}{r_{i,i}};$$

$\vdots$

$$\rightarrow j = N : a_{i,N} = r_{i,1}r_{N,1} + \dots + r_{i,i}r_{N,i}$$

$$\implies r_{N,i} = \frac{a_{i,N} - (r_{i,1}r_{N,1} + \dots + r_{i,i-1}r_{N,i-1})}{r_{i,i}}.$$

Applications de la factorisation de Cholesky:

1) Résoudre le système linéaire $Ax = b$ se ramène à résoudre successivement les 2 systèmes linéaires à matrices triangulaires:

$$Ax = b \iff \begin{cases} Ry = b \\ R^T x = y \end{cases}$$

2) Calcul du déterminant de A : $\det(A) = (\prod_{i=1}^N r_{i,i})^2$.

L'ordre du nombre d'opérations pour résoudre un système linéaire, par la méthode de Cholesky:

$\frac{N^3}{6}$ additions;

$\frac{N^3}{6}$ multiplications;

$\frac{N^2}{6}$ divisions;

N extractions de racines carrées.

Exemple:

$A = \begin{pmatrix} 1 & \frac{1}{2} & \frac{1}{3} \\ \frac{1}{2} & \frac{1}{3} & \frac{1}{4} \\ \frac{1}{3} & \frac{1}{4} & \frac{1}{5} \end{pmatrix}$; la factorisation de Cholesky de A :

$$A = \begin{pmatrix} 1 & 0 & 0 \\ \frac{1}{2} & \frac{\sqrt{3}}{6} & 0 \\ \frac{1}{3} & \frac{\sqrt{3}}{6} & \frac{\sqrt{5}}{30} \end{pmatrix} \begin{pmatrix} 1 & \frac{1}{2} & \frac{1}{3} \\ 0 & \frac{\sqrt{3}}{6} & \frac{\sqrt{3}}{6} \\ 0 & 0 & \frac{\sqrt{5}}{30} \end{pmatrix} = RR^T.$$

1) $\det(A) = (\frac{\sqrt{3}}{6} \frac{\sqrt{5}}{30})^2 = \frac{1}{2160}$.

2) Pour résoudre $Ax = b$; avec $b = \begin{pmatrix} 1 \\ 1 \\ 1 \end{pmatrix}$,

$$Ry = b \implies y_1 = 1, y_2 = \sqrt{3} \text{ et } y_3 = \sqrt{5}$$

$$R^T x = y \implies x = \begin{pmatrix} 3 \\ -24 \\ 30 \end{pmatrix}.$$

1.3 Conditionnement d'un système linéaire

Soit $\|\cdot\|$ une norme matricielle subordonnée à une norme vectorielle.

Définition Soit $A \in \mathcal{M}_n(\mathbb{K})$ une matrice inversible.

On définit le nombre

$$\text{cond}(A) = \|A\| \|A^{-1}\|$$

appelé le conditionnement de la matrice A , relativement à la norme matricielle considérée.

Théorème:

Soit $A \in \mathcal{M}_n(\mathbb{K})$ une matrice inversible et x un vecteur tel que:

$$Ax = b,$$

1) Soit $x + \delta x$ la solution de

$$A(x + \delta x) = b + \delta b$$

On suppose $b \neq 0$, alors on a :

$$\frac{\|\delta x\|}{\|x\|} \leq \text{cond}(A) \frac{\|\delta b\|}{\|b\|} \quad (1)$$

et c'est la meilleure possible (c'est à dire: $\exists b \neq 0$ et $\exists \delta b \neq 0$ tels que (1) devienne égalité).

2) Soit $x + \Delta x$ solution de

$$(A + \Delta A)(x + \Delta x) = b$$

On suppose $b \neq 0$, alors on a :

$$\frac{\|\Delta x\|}{\|x + \Delta x\|} \leq \text{cond}(A) \frac{\|\Delta A\|}{\|A\|} \quad (2)$$

et c'est la meilleure possible

(c'est à dire: $\exists b \neq 0$ et $\exists \delta b \neq 0$ tels que (2) devienne égalité).

Démonstration:

$$1) \rightarrow Ax = b \implies \|b\| = \|Ax\| \leq \|A\| \cdot \|x\| \implies \|x\| \geq \frac{\|b\|}{\|A\|}$$

$$\text{et } A(x + \delta x) = b + \delta b \implies A(\delta x) = \delta b$$

Donc

$$\frac{\|\delta x\|}{\|x\|} = \frac{\|A^{-1}\delta b\|}{\|x\|} \leq \frac{\|A^{-1}\| \cdot \|\delta b\|}{\|x\|} \leq \|A^{-1}\| \cdot \|\delta b\| \cdot \frac{\|A\|}{\|b\|}$$

$$\implies \frac{\|\delta x\|}{\|x\|} \leq \text{cond}(A) \frac{\|\delta b\|}{\|b\|}.$$

→ Si on choisit

$$\text{un vecteur } y \text{ tel que: } \|Ay\| = \|A\| \cdot \|y\|$$

et

$$\delta b \text{ tel que: } \|A^{-1}\delta b\| = \|A^{-1}\| \cdot \|\delta b\|,$$

on a l'égalité dans (1).

$$2) \rightarrow (A + \Delta A)(x + \Delta x) = b$$

$$\implies \Delta x = -A^{-1}\Delta A(x + \Delta x)$$

$$\implies \|\Delta x\| \leq \|A^{-1}\| \cdot \|\Delta A\| \cdot \|x + \Delta x\|$$

$$\implies \frac{\|\Delta x\|}{\|x + \Delta x\|} = \frac{\|A^{-1}\Delta A(x + \Delta x)\|}{\|x + \Delta x\|} \leq \|A^{-1}\| \cdot \|\Delta A\| = \text{cond}(A) \frac{\|\Delta A\|}{\|A\|}.$$

→ Si on choisit;

$$\text{un vecteur } y \neq 0, \text{ tel que: } \|A^{-1}y\| = \|A^{-1}\| \cdot \|y\|$$

et

$$\text{un scalaire } \beta \neq 0,$$

et on prend alors:

$$b = (A + \beta I)y, \Delta A = \beta I \text{ et } \Delta x = -\beta A^{-1}y,$$

$$\implies x + \Delta x = y, Ax = b \text{ et } (A + \Delta A)(x + \Delta x) = (A + \beta I)y = b.$$

$$\implies \|\Delta x\| \leq |\beta| \cdot \|A^{-1}y\| = |\beta| \cdot \|A^{-1}\| \cdot \|y\| = \|\Delta A\| \cdot \|A^{-1}\| \cdot \|x + \Delta x\|.$$

De plus si β n'est pas une valeur propre de A , alors $b \neq 0$.

Propriétés Soit $A \in \mathcal{M}_N(\mathbb{R})$

$$1) \text{ Si } \Delta A \in \mathcal{M}_N(\mathbb{R}), \text{ alors } \|\Delta A\| < \frac{1}{\|A^{-1}\|}$$

$$\implies \frac{\|\Delta x\|}{\|x\|} \leq \text{cond}(A) \cdot \frac{\|\Delta A\|}{\|A\|} \cdot \frac{1}{1 - \|A^{-1}\| \cdot \|\Delta A\|}.$$

$$2) \text{ (i) } \text{cond}(A) \geq 1,$$

$$\text{(ii) } \text{cond}(A) = \text{cond}(A^{-1}),$$

$$\text{(iii) } \text{cond}(\alpha A) = \text{cond}(A), \forall \alpha \text{ scalaire } \neq 0.$$

$$3) \text{ On désigne par } \text{cond}_2(A)$$

le conditionnement de la matrice A , relatif à la norme euclidienne.

$$\text{(i) } \text{cond}_2(A) = \sqrt{\frac{\mu_N(A)}{\mu_1(A)}},$$

où $\mu_1(A)$ et $\mu_N(A)$ sont respectivement la plus petite et la plus grande des valeurs propres de la matrice $A^T A$.

(on rappelle que A inversible $\implies A^T A$ symétrique définie positive)

$$\text{(ii) Si } A \text{ est une matrice normale } (AA^T = A^T A),$$

$$\text{cond}_2(A) = \frac{\max_i |\lambda_i(A)|}{\min_i |\lambda_i(A)|},$$

où les $\lambda_i(A)$ sont les valeurs propres de la matrice A .

(iii) Si A est unitaire ou orthogonale, $\text{cond}_2(A) = 1$.

(iv) U orthogonale ($U^T U = U U^T = I_N$)

$\Rightarrow \text{cond}_2(A) = \text{cond}_2(AU) = \text{cond}_2(UA) = \text{cond}_2(U^T AU)$.

1.4 Méthodes itératives

On va voir un type de méthodes itératives de résolution du système linéaire $Ax = b$ sous la forme:

$$(3) \quad \begin{cases} x^{(0)} \text{ vecteur arbitraire,} \\ x^{(k+1)} = Bx^{(k)} + c, k \geq 0 \end{cases}$$

lorsque

$$Ax = b \iff x = Bx + c,$$

la matrice B et le vecteur c sont en fonction de A et b .

Définition

La méthode itérative (3) est convergente si

$$\lim_{k \rightarrow +\infty} x^{(k)} = x, \forall x^{(0)}.$$

Remarque:

Si on pose $e^{(k)} = x^{(k)} - x$, pour $k = 0, 1, \dots$

Comme $x = Bx + c$ et $x^{(k+1)} = Bx^{(k)} + c$, on a

$$e^{(k)} = Bx^{(k-1)} - Bx = Be^{(k-1)} = \dots = B^k e^{(0)}$$

$$\lim_{k \rightarrow +\infty} x^{(k)} = x \iff \lim_{k \rightarrow +\infty} e^{(k)} = 0 \iff \lim_{k \rightarrow +\infty} B^k e^{(0)} = 0$$

Donc

La méthode itérative (3) est convergente si

$$\lim_{k \rightarrow +\infty} B^k v = 0, \forall v$$

ce qui équivaut à

$$\lim_{k \rightarrow +\infty} \|B^k v\| = 0, \forall v$$

pour toute norme vectorielle $\|\cdot\|$.

1.4.1 Convergence des méthodes itératives

Théorème

Les propositions suivantes sont équivalentes:

(1) la méthode itérative (3) est convergente;

(2) $\rho(B) < 1$;

(3) $\|B\| < 1$ pour au moins une norme matricielle subordonnée $\|\cdot\|$.

Démonstration

(1) $\Rightarrow$ (2)

Supposons $\rho(B) \geq 1$, $\rho(B) = |\lambda|$, donc $\exists$ un vecteur p :

$$p \neq 0, Bp = \lambda p \text{ et } |\lambda| \geq 1$$

$$\Rightarrow \forall k \geq 0, \|B^k p\| = \|\lambda^k p\| = |\lambda|^k \|p\| = |\lambda|^k \|p\| \geq \|p\|$$

ce qui contredit $\lim_{k \rightarrow +\infty} B^k p = 0$.

(2) $\implies$ (3)

On utilise la proposition 4: $\forall \epsilon > 0$, il existe au moins une norme matricielle subordonnée telle que:

$$\|B\| \leq \rho(B) + \epsilon.$$

(3) $\implies$ (1)

Soit $\|\cdot\|$ une norme matricielle subordonnée.

On utilise alors la propriété:

$$\begin{aligned} \forall k \geq 0, \forall v, \quad \|B^k v\| &\leq \|B^k\| \cdot \|v\| \leq \|B\|^k \cdot \|v\| \\ \|B\| < 1 &\implies \lim_{k \rightarrow +\infty} \|B\|^k = 0 \implies \lim_{k \rightarrow +\infty} \|B^k v\| = 0. \end{aligned}$$

1.4.2 Méthode de Jacobi, de Gauss-Seidel et de relaxation

Ces méthodes sont des cas particuliers de la méthode suivante:

$A = M - N$ avec M inversible et assez simple. On aurait alors:

$$\begin{aligned} Ax = b &\iff Mx = Nx + b \\ &\iff x = M^{-1}Nx + M^{-1}b \\ &\iff x = Bx + c, \end{aligned}$$

avec $B = M^{-1}N$ et $c = M^{-1}b$.

Méthode de Jacobi: En posant $A = D - (E + F)$,

$M = D$ est la diagonale de A et $N = E + F$.

$$Ax = b \iff Dx = (E + F)x + b.$$

On suppose que D est inversible, c'est à dire $a_{ii} \neq 0, 1 \leq i \leq N$.

La matrice

$$J = D^{-1}(E + F) = I_N - D^{-1}A$$

est appelée la matrice de Jacobi.

$$\begin{cases} x^{(0)} \text{ donné,} \\ Dx^{(k+1)} = (E + F)x^{(k)} + b, k \geq 0 \end{cases}$$

A chaque étape, on calcule les N composantes $x_1^{(k+1)}, \dots, x_N^{(k+1)}$ du vecteur $x^{(k+1)}$:

$$\begin{cases} a_{1,1}x_1^{(k+1)} = -a_{1,2}x_2^{(k)} - \dots - a_{1,N}x_N^{(k)} + b_1 \\ a_{2,2}x_2^{(k+1)} = -a_{2,1}x_1^{(k)} - a_{2,3}x_3^{(k)} \dots - a_{2,N}x_N^{(k)} + b_2 \\ \vdots \\ a_{N,N}x_N^{(k+1)} = -a_{N,1}x_1^{(k)} - \dots - a_{N,N-1}x_{N-1}^{(k)} + b_N \end{cases}$$

Méthode de Gauss-Seidel M est la partie triangulaire inférieure de A : $M = D - E$ et $N = F$.

On pourrait améliorer la méthode précédente en utilisant les quantités déjà calculées, on calcule successivement les N composantes

$x_1^{(k+1)}, \dots, x_N^{(k+1)}$ du vecteur $x^{(k+1)}$:

$$\begin{cases} a_{1,1}x_1^{(k+1)} = -a_{1,2}x_2^{(k)} - \dots - a_{1,N}x_N^{(k)} + b_1 \\ a_{2,2}x_2^{(k+1)} = -a_{2,1}x_1^{(k+1)} - a_{2,3}x_3^{(k)} \dots - a_{2,N}x_N^{(k)} + b_2 \\ \vdots \\ a_{N,N}x_N^{(k+1)} = -a_{N,1}x_1^{(k+1)} - \dots - a_{N,N-1}x_{N-1}^{(k+1)} + b_N \end{cases}$$

ce qui revient à écrire:

$$Dx^{(k+1)} = Ex^{(k+1)} + Fx^{(k)} + b$$

ou encore

$$(D - E)x^{(k+1)} = Fx^{(k)} + b$$

$$x^{(k+1)} = (D - E)^{-1}Fx^{(k)} + (D - E)^{-1}b$$

$\mathcal{L}_1 = (D - E)^{-1}F$ est la matrice de Gauss-Seidel.

Elle est inversible si $a_{ii} \neq 0, 1 \leq i \leq N$.

Méthode de relaxation Pour $\omega \neq 0$, en posant $M = \frac{1}{\omega}D - E$, $N = (\frac{1-\omega}{\omega})D + F$, on a

$$\begin{cases} A = M - N = \{\frac{1}{\omega}D - E\} - \{(\frac{1-\omega}{\omega})D + F\} \\ a_{1,1}x_1^{(k+1)} = a_{1,1}x_1^{(k)} - \omega\{a_{1,1}x_1^{(k)} + a_{1,2}x_2^{(k)} + \dots + a_{1,N}x_N^{(k)} + b_1\} \\ a_{2,2}x_2^{(k+1)} = a_{2,2}x_2^{(k)} - \omega\{a_{2,1}x_1^{(k+1)} + a_{2,2}x_2^{(k)} \dots + a_{2,N}x_N^{(k)} + b_2\} \\ \vdots \\ a_{N,N}x_N^{(k+1)} = a_{N,N}x_N^{(k)} - \omega\{a_{N,1}x_1^{(k+1)} + \dots + a_{N,N-1}x_{N-1}^{(k+1)} + a_{N,N}x_N^{(k)} + b_N\} \end{cases}$$

ce qui revient à écrire:

$$\{\frac{1}{\omega}D - E\}x^{(k+1)} = \{(\frac{1-\omega}{\omega})D + F\}x^{(k)} + b$$

La matrice de relaxation est

$$\mathcal{L}_\omega = \{\frac{1}{\omega}D - E\}^{-1}\{(\frac{1-\omega}{\omega})D + F\} = (D - \omega E)^{-1}\{(1 - \omega)D + \omega F\}$$

Exemples: Etude de la convergence des méthodes de Jacobi et de Gauss-Seidel dans le cas où la matrice du système linéaire est:

$$1) A_1 = \begin{pmatrix} -4 & 1 & 0 \\ 1 & -4 & 1 \\ 0 & 1 & -4 \end{pmatrix};$$

Les deux méthodes convergent.

$$2) A_2 = \begin{pmatrix} 2 & 1 & 1 \\ 1 & 2 & 1 \\ 1 & 1 & 2 \end{pmatrix};$$

la méthode de Jacobi diverge et La méthode de Gauss-Seidel converge.

$$3) A = \begin{pmatrix} 1 & 2 & -2 \\ 1 & 1 & 1 \\ 2 & 2 & 1 \end{pmatrix}.$$

la méthode de Jacobi converge et la méthode de Gauss-Seidel diverge.

1.4.3 Convergence des méthodes de Jacobi, de Gauss-Seidel et relaxation

Théorème

Soit A une matrice symétrique définie positive.

On suppose que

$$A = M - N, \text{ avec } M \text{ inversible.}$$

Si la matrice symétrique

$$M^T + N$$

est définie positive,

alors

$$\rho(M^{-1}N) < 1.$$

Démonstration

1) $M^T + N$ est symétrique:

$$A = M - N \implies M^T + N = A^T + N^T + N = A + N^T + N = M + N^T.$$

2) Comme A est symétrique définie positive, l'application

$$v \in \mathbb{R}^N \rightarrow \|v\| = (v^T A v)^{\frac{1}{2}}$$

définit une norme sur $\mathbb{R}^N$.

On considère alors la norme matricielle $\|\cdot\|$ induite par cette norme vectorielle.

$$\|M^{-1}N\| = \|I_N - M^{-1}A\| = \sup_{\|v\|=1} \|v - M^{-1}Av\|$$

Soit v un vecteur tel que $\|v\| = 1$. On pose $w = M^{-1}Av$,

$$\begin{aligned} \|v - w\|^2 &= (v - w)^T A (v - w) = 1 - v^T A w - w^T A v + w^T A w \\ &= 1 - w^T M^T w - w^T M w + w^T A w \\ &= 1 - w^T (M^T + N) w. \end{aligned}$$

$$v \neq 0 \implies w = M^{-1}Av \neq 0 \implies w^T (M^T + N) w > 0$$

Donc si la matrice symétrique $M^T + N$ est définie positive, on a

$$\|v - w\|^2 = 1 - w^T (M^T + N) w < 1$$

Exemple:

$$A = \begin{pmatrix} 2 & -1 & 0 & 0 \\ -1 & \ddots & \ddots & 0 \\ 0 & \ddots & \ddots & -1 \\ 0 & 0 & -1 & 2 \end{pmatrix}$$

1) Pour tout $u \in \mathbb{R}^N$,

$$u^T A u = u_1^2 + u_N^2 + \sum_{i=2}^N (u_i - u_{i-1})^2$$

$\implies$ la matrice symétrique A est définie positive.

2) Pour la méthode de Jacobi:

$$M = D \text{ et } N = E + F:$$

$A = M - N$ est symétrique définie positive, car

$$\begin{aligned} u^T (M^T + N) u &= u_1^2 + u_N^2 + \sum_{i=2}^N (u_i + u_{i-1})^2 \\ &> 0, \text{ si } u \neq 0. \end{aligned}$$

$$M^T + N \text{ est définie positive} \implies \rho(M^{-1}N) < 1.$$

Proposition 7 (CS) Soit A une matrice symétrique définie positive.

$0 < \omega < 2 \implies$ la méthode de relaxation converge.

Démonstration

On applique le résultat précédent.

On a $A = M - N$

avec $M = \frac{1}{\omega}D - E$ et $N = (\frac{1-\omega}{\omega})D + F$

$$M^T + N = \frac{1}{\omega}D - E^T + (\frac{1-\omega}{\omega})D + F$$

$$= \frac{1}{\omega}D + (\frac{1-\omega}{\omega})D$$

$$= (\frac{2-\omega}{\omega})D$$

$$0 < \omega < 2 \implies \frac{2-\omega}{\omega} > 0$$

$\implies (\frac{2-\omega}{\omega})D$ est définie positive.

Proposition 8 Soit A une matrice, alors
pour tout $\omega \neq 0$, $\rho(\mathcal{L}_\omega) \geq |\omega - 1|$.

Démonstration

$$\prod_{i=1}^N \lambda_i(\mathcal{L}_\omega) = \det(\mathcal{L}_\omega) = \frac{\det(\frac{1-\omega}{\omega}D + F)}{\det(\frac{1}{\omega}D - E)} = (1 - \omega)^N$$

$$\implies \rho(\mathcal{L}_\omega) \geq \left| \prod_{i=1}^N \lambda_i(\mathcal{L}_\omega) \right|^{\frac{1}{N}} = |1 - \omega|.$$

Corollaire

Si A est une matrice symétrique définie positive,

alors

$0 < \omega < 2 \iff$ la méthode de relaxation converge.

1.4.4 Autres méthodes itératives:

Méthode du gradient $\begin{cases} x^{(0)} \text{ arbitraire} \\ x^{(k+1)} = x^{(k)} - \alpha(Ax^{(k)} - b), k \geq 0 \end{cases}$

où α est une constante fixée.

Si on pose $r^{(k)} = Ax^{(k)} - b, k \geq 0$ on a

$$x^{(k+1)} = x^{(k)} - \alpha(Ax^{(k)} - b)$$

$$= x^{(k)} - \alpha Ar^{(k)}$$

$$\implies r^{(k+1)} = r^{(k)} - \alpha Ar^{(k)}$$

$$= (I - \alpha A)r^{(k)}.$$

Une condition nécessaire et suffisante

de convergence est:

$$\rho(I - \alpha A) < 1.$$

Proposition 9 Si la matrice A est symétrique définie positive alors la méthode du gradient à pas fixe converge si et seulement si

$$0 < \alpha < \frac{2}{\lambda_n}$$

Avec λ_n la plus grande valeur propre de A .

De plus $\alpha = \frac{2}{\lambda_1 + \lambda_n}$ est une valeur optimale ($I - \alpha A$ a le plus petit rayon spectral)

Méthode du gradient à pas optimal $\begin{cases} x^{(0)} \text{ arbitraire} \\ x^{(k+1)} = x^{(k)} - \alpha_k (Ax^{(k)} - b), k \geq 0 \end{cases}$

où α_k est choisi tel que:

$$r^{(k+1)} \perp r^{(k)}$$

avec $r^{(k)} = Ax^{(k)} - b, k \geq 0$.

Donc $r^{(k+1)} = r^{(k)} - \alpha_k \cdot r^{(k)}$.

$$r^{(k+1)} \perp r^{(k)} \iff \langle r^{(k)}, r^{(k)} - \alpha_k \cdot Ar^{(k)} \rangle = 0$$

$$\iff \alpha_k = \frac{\|r^{(k)}\|_2^2}{\langle r^{(k)}, Ar^{(k)} \rangle}.$$

L'algorithme est:

(0) $k = 0$

$x^{(0)}$ arbitraire

(1) calcul de $r^{(k)} = Ax^{(k)} - b$

(2) Si $r^{(k)} = 0$, c'est terminé.

(3) Si $r^{(k)} \neq 0$, on calcule:

$$\alpha_k = \frac{\|r^{(k)}\|_2^2}{\langle r^{(k)}, Ar^{(k)} \rangle}$$

et

$$x^{(k+1)} = x^{(k)} - \alpha_k \cdot r^{(k)}.$$

(3) $k := k + 1$ et aller à (1).

On montre le résultat suivant:

Proposition

Si A est une matrice symétrique définie positive, alors la méthode du gradient à pas optimal converge.

Exemples:

$$A = \begin{pmatrix} 2 & 1 \\ 1 & 3 \end{pmatrix} \text{ et } b = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$$

$$\text{La solution de } Ax = b \text{ est } x = \begin{pmatrix} \frac{3}{5} \\ \frac{-1}{5} \end{pmatrix}$$

Les valeurs propres de

$$A = \begin{pmatrix} 2 & 1 \\ 1 & 3 \end{pmatrix}$$

sont

$$\lambda_1 = \frac{5}{2} - \frac{1}{2}\sqrt{5} = 1.3820$$

et

$$\lambda_2 = \frac{1}{2}\sqrt{5} + \frac{5}{2} = 3.618$$

Donc A est symétrique définie positive.

$$\text{On choisit } x^{(0)} = \begin{pmatrix} 1 \\ \frac{1}{2} \end{pmatrix}$$

et on calcule le premier itéré pour:

1) la méthode de Jacobi:

$$x^{(1)} = \begin{pmatrix} \frac{1}{4} \\ \frac{-1}{3} \end{pmatrix}$$

2) la méthode de Gauss-Seidel:

$$x^{(1)} = \begin{pmatrix} \frac{1}{4} \\ \frac{-1}{12} \end{pmatrix}$$

3) la méthode du gradient à pas fixe

$$r^{(0)} = \begin{pmatrix} 2 & 1 \\ 1 & 3 \end{pmatrix} \begin{pmatrix} 1 \\ \frac{1}{2} \end{pmatrix} - \begin{pmatrix} 1 \\ 0 \end{pmatrix} \\ = \begin{pmatrix} \frac{3}{2} \\ \frac{5}{2} \end{pmatrix}$$

Si on choisit α tel que

$$0 < \alpha < \frac{2}{\lambda_2} = \frac{2}{3.618} = 0.552\,79,$$

$$\alpha = \frac{2}{\lambda_2 + \lambda_1} = \frac{2}{5}$$

la méthode converge.

Si on prend $\alpha = \frac{1}{2}$:

$$x^{(1)} = \begin{pmatrix} 1 \\ \frac{1}{2} \end{pmatrix} - \frac{1}{2} \begin{pmatrix} \frac{3}{2} \\ \frac{5}{2} \end{pmatrix} \\ = \begin{pmatrix} \frac{1}{4} \\ -\frac{3}{4} \end{pmatrix}$$

Si on prend $\alpha = \frac{2}{5}$

$$x^{(1)} = \begin{pmatrix} 1 \\ \frac{1}{2} \end{pmatrix} - \frac{2}{5} \begin{pmatrix} \frac{3}{2} \\ \frac{5}{2} \end{pmatrix} \\ = \begin{pmatrix} \frac{2}{5} \\ -\frac{1}{2} \end{pmatrix}$$

4) la méthode du gradient à pas optimal:

$$r^{(0)} = \begin{pmatrix} \frac{3}{2} \\ \frac{5}{2} \end{pmatrix}$$

$$\|r^{(0)}\|_2^2 = \begin{pmatrix} \frac{3}{2} & \frac{5}{2} \end{pmatrix} \begin{pmatrix} \frac{3}{2} \\ \frac{5}{2} \end{pmatrix} = \frac{17}{2}$$

$$\begin{pmatrix} \frac{3}{2} & \frac{5}{2} \end{pmatrix} \begin{pmatrix} 2 & 1 \\ 1 & 3 \end{pmatrix} \begin{pmatrix} \frac{3}{2} \\ \frac{5}{2} \end{pmatrix} = \frac{123}{4}$$

$$\lambda_0 = \frac{17}{2} \frac{4}{124} = \frac{17}{62}$$

$$x^{(1)} = \begin{pmatrix} 1 \\ \frac{1}{2} \end{pmatrix} - \frac{17}{62} \begin{pmatrix} \frac{3}{2} \\ \frac{5}{2} \end{pmatrix} \\ = \begin{pmatrix} \frac{73}{124} \\ -\frac{23}{124} \end{pmatrix}$$