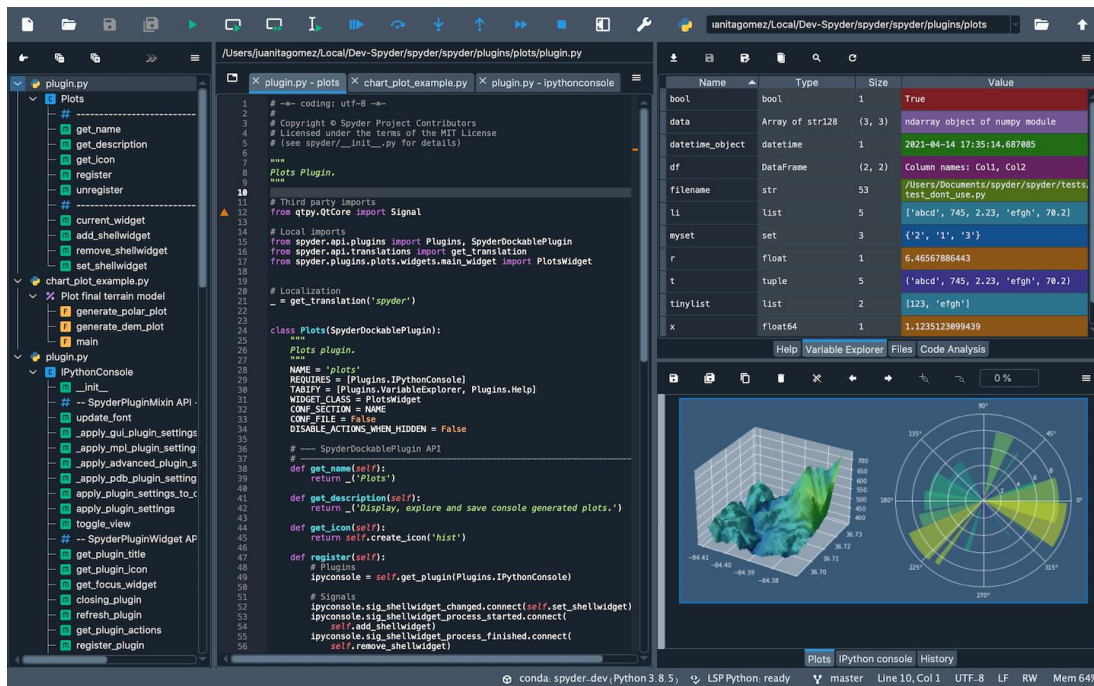




UE 503.1 ANALYSE NUMERIQUE - L3 MFE POLYCOPE DE COURS, TD, TP



Olivier BOTELLA

LEMTA - Université de Lorraine - CNRS
2 avenue de la Forêt de Haye, TSA 60604
54518 Vandoeuvre (France)

Email : Olivier.Botella@univ-lorraine.fr

Lecture web page : <https://arche.univ-lorraine.fr/course/view.php?id=6384>

Interpolation

I) Exemple "naïf"

Les observations scientifiques de la concentration d'ogun dans la ville de Nancy a produit le tableau suivant

ogune (%)	$y_0 = 12.0$	$y_1 = 14.2$	?	$y_3 = 19.8$	?
Année	$x_0 = 1990$	$x_1 = 1994$	$x_2 = 1998$	$x_3 = 2002$	$x_4 = 2006$

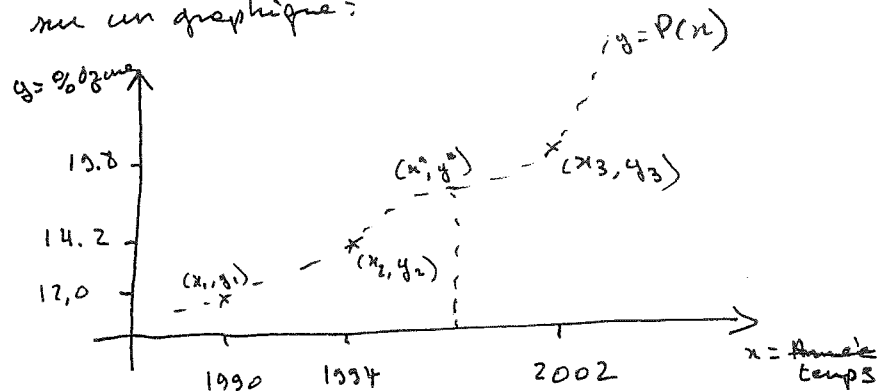
On souhaite maintenant :

* estimer la concentration d'ogun en $x_2 = 1998$
 $\Rightarrow$ Problème d'interpolation.

* prédire la _____ en $x_4 = 2006$
 $\Rightarrow$ Problème d'extrapolation.

La résolution de ces 2 problèmes fait appel à des techniques d'approximation polynomiale.

Naïvement, cela consiste à reporter les données du tableau sur un graphique :



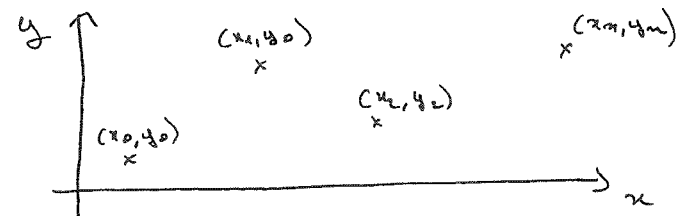
Ensuite, on cherche une fonction "suffisamment simple" qui passe par ces données (puiss. un polynôme $y = P(x)$).
~~Il~~ ^{suffit} reste alors de calculer $y_2 = P(x_2)$ aux points désirés (ici $x_2 = 1998$ ou 2006).

II) Interpolation polynomiale

Rappels sur les polynômes :

- On note IP_N l'ensemble des polyn. de degré $\leq N$.
 (ex.: IP_1 : polynômes de la forme $ax + b$, $ax^2 + bx + c$, ou c)
- $P(x) \in IP_N$ possède au plus N racines complexes.
- IP_N est un espace vectoriel de dimension $N+1$.
 La base canonique est $\{1, x, x^2, \dots, x^N\}$.

II-1 Formulation des problèmes



Supposons qu'on dispose d'un ensemble de $n+1$ données
 $\{(x_i, y_i), i = 0 \dots n\}$
 que l'on écrit sous forme tabulaire

x_0	x_1	...	x_i	...	x_n
y_0	y_1		y_i		y_n

On cherche le polynôme $p(x)$ de degré n tel que

$$p(x_i) = y_i, \quad \text{pour } 0 \leq i \leq n. \quad (1)$$

Théorème (propriété importante)

Le polynôme d'interpolation $p(x)$ est unique.

Exemple:

- par 2 points passe une seule droite ($n=1$)
- par 3 ————— parabole ($n=2$)
- par $n+1$ ————— polyn. de degré n .

Démonstration: Supposons qu'il existe un autre polynôme q qui vérifie (1). Alors, le polynôme

$$r(x) = p(x) - q(x)$$

possède $n+1$ racines $x_0, \dots, x_n$. C'est seulement possible si $r(x) \equiv 0$.

II-2 Une méthode naïve

Pour résoudre le problème d'interpolation (1), une première possibilité est d'écrire $p(x) \in \mathbb{P}_n$ dans sa base canonique:

$$p(x) = a_0 + a_1 x + a_2 x^2 + \dots + a_n x^n = \sum_{j=0}^n a_j x^j.$$

Le problème (1) revient à chercher les $n+1$ coefficients $\{a_j\}$ à partir des $n+1$ données $\{x_i, y_i\}$, soit:

$$a_0 + a_1 x_i + a_2 x_i^2 + \dots + a_n x_i^n = y_i, \quad 0 \leq i \leq n$$

$$\Leftrightarrow \sum_{j=0}^n (x_i)^j a_j = y_i, \quad 0 \leq i \leq n. \quad (2)$$

Le calcul des $\{a_j\}$ par (2) revient à résoudre le système linéaire de taille $(n+1) \times (n+1)$:

$$\underline{A} \underline{X} = \underline{B}, \quad (3)$$

$$\text{avec } \underline{X} = \begin{pmatrix} a_0 \\ a_1 \\ \vdots \\ a_n \end{pmatrix}, \quad \underline{B} = \begin{pmatrix} y_0 \\ y_1 \\ \vdots \\ y_n \end{pmatrix}$$

$$\underline{A} = \begin{bmatrix} 1 & x_0 & x_0^2 & \dots & x_0^n \\ 1 & x_1 & x_1^2 & \dots & x_1^n \\ \vdots & \vdots & \vdots & & \vdots \\ 1 & x_n & x_n^2 & \dots & x_n^n \end{bmatrix}$$

Résoudre "efficacement" un système linéaire n'est pas une tâche triviale (cf: plus tard dans le cours).

Dans la suite de ce cours, nous allons utiliser une technique d'interpolation plus astucieuse, qui utilise la base de Lagrange $\{L_j(x), j=0, \dots, n\}$ plutôt que la base canonique $\{x^j, j=0, \dots, n\}$ pour représenter le polynôme $p(x)$.

11.2 Interpolation de Lagrange

En premier lieu, considérons le problème d'interpolation des données $\{(x_i, y_i), i=0 \dots n\}$ représentées sur le tableau

Tab2

0	0	---	0	1	0	---	0
x_0	x_1		x_{j-1}	x_j	x_{j+1}		x_n

On cherche donc le pol. d'interpolation $L_j(x) \in \mathbb{P}_n$ qui vérifie

$$\text{Pour } i=0, \dots, n : L_j(x_i) = \begin{cases} 1 & \text{si } i=j \\ 0 & \text{si } i \neq j \end{cases} \quad (4)$$

(rq : Notation de Kronecker $L_j(x_i) = \delta_{ij}$).

Ce polynôme de degré n possède n racines $\{x_0, x_1, \dots, x_{j-1}, x_{j+1}, \dots\}$, il est donc forcément de la forme :

$$\begin{aligned} L_j(x) &= d_j (x-x_0)(x-x_1) \dots (x-x_{j-1})(x-x_{j+1}) \dots (x-x_n) \\ &= d_j \prod_{\substack{i=0 \\ i \neq j}}^n (x-x_i), \end{aligned}$$

où d_j est une constante. Pour la déterminer, il suffit de rq :

$$L_j(x_j) = 1 \iff d_j = \frac{1}{\prod_{\substack{i=0 \\ i \neq j}}^n (x_j - x_i)}.$$

Définition : Le polynôme

$$L_j(x) = \prod_{\substack{i=0 \\ i \neq j}}^n \frac{(x-x_i)}{(x_j-x_i)} \quad (5)$$

est appelé le $j^{\text{ème}}$ polynôme de Lagrange.

Maintenant, pour interpréter le problème :

0	0		y_j		0
x_0	x_1		x_j		x_n

l'interpolant s'écrit :

$$p(x) = y_j L_j(x).$$

Pour le problème plus général :

y_0	y_1	---	y_j	---	y_n
x_0	x_1	---	x_j	---	x_n

L'interpolant est une combinaison linéaire des polyn. de Lagrange :

$$p(x) = \sum_{j=0}^n y_j L_j(x). \quad (6)$$

Remarque (Approximation de fonctions)

Un des problèmes de la théorie mathématique de l'approximation de fonctions est de chercher une approximation "simple" (p.ex. polynomiale) d'une fct. $f: \mathbb{R} \rightarrow \mathbb{R}$ quelconque.

Si on suppose que les valeurs $\{f(x_i)\}$ de la fonction sont connues en certains points $\{x_i, i=0, \dots, n\}$, alors le polynôme d'interpolation de Lagrange de $f(x)$ s'écrit

$$\sum_{j=0}^n f(x_j) L_j(x).$$

On note généralement $\Pi_n f$ le polynôme d'interpolation

* L'interpolation de Lagrange est une interpolation compacte, mieux adaptée au calcul numérique que l'interpolation "maître" dans la base $\{x_i\}$, car elle ne nécessite pas la résolution de systèmes linéaires (cf II).

* Lorsque le nombre de données n est grand, l'interpolation polynomiale peut donner de mauvais résultats (cf TD, TP : phénomène de Runge). Bien que le polynôme "adhère" aux données par construction, il peut osciller fortement entre ces points. (Math, cela se traduit par le fait que $\pi_n f$ ne converge pas uniformément vers f .)

Remèdes contre ces oscillations:

- Si on peut choisir les nœuds x_i : ne pas les prendre équirépartis, mais plutôt choisir les points de Tchébychev, qui sont définis dans $U =]a, b[$ par:

$$x_j = \frac{a+b}{2} + \frac{b-a}{2} \cos \left[\frac{(2j+1)\pi}{2n+2} \right], \quad j=0, 1, \dots, n$$

- Utiliser l'interpolation basée sur les polynômes par morceaux.

$$x_i = \frac{a+b}{2} - \frac{(b-a)}{2} \cos(i\pi/n), \quad i=0, \dots, n.$$

III Interpolation par intervalles

Rappel: Espaces fonctionnels $C^k(U)$.

$$C^k(U) = \{ f, \text{ tels que } f \text{ est continue sur } U, \\ \bullet f^{(i)} \text{ est continue sur } U \text{ pour } 1 \leq i \leq k \}$$

exple: $C^0(U)$: ensemble des fonctions continues.

$C^1(U)$: fonctions continues à dérivée première en

III-1 Notion de polynômes par morceaux (piecewise polynomials)

On considère l'intervalle $U =]a, b[$, et la ~~partition~~ partition de des nœuds répartis sur U :

$$U_n = \{ x_i, i=0, \dots, n \}$$

A partir de ces nœuds, on définit les intervalles $U_i =]x_i, x_{i+1}[$ de longueur $h_i = x_{i+1} - x_i$.

de degré k
Déf: Un polynôme par morceaux $P_h(u)$ sur U_n est tel que

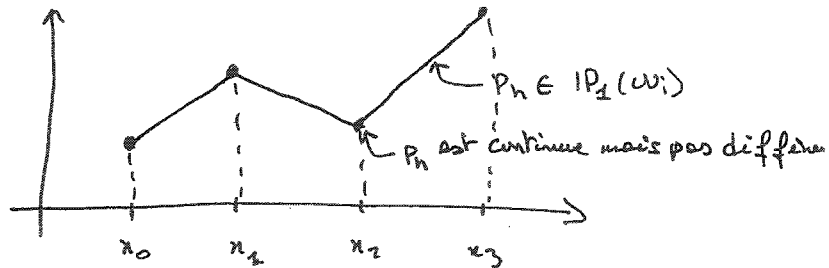
- P_h est un polynôme de $\mathbb{R}_k$ degré k sur U_i
- $P_h \in C^{k-1}(U)$.

rq: En particulier, on remarque que la dérivée $k^{\text{ième}}$ de P_h est discontinue sur les nœuds x_i .

Les polyn. p.m. les plus utilisés sont les polynômes p.m. IP_1 , qui sont tels que:

- $P_h(x) \in IP_1$ sur $\mathcal{U}_i$ (de la forme $ax+b$)
- $P_h(x)$ est seulement continue aux nœuds x_i .

Un exemple de polyn. p.m. est:



III-2 Base de Lagrange des polyn. par morceaux IP_1

Pour pouvoir représenter $P_h(x)$ de façon commode, on définit la base de Lagrange:

$$\{\varphi_j(x), j = 0, \dots, n\},$$

et $P_h(x)$ s'écrit alors:

$$P_h(x) = \sum_{j=0}^n \alpha_j \varphi_j(x) \quad (\alpha_j \in \mathbb{R})$$

Les fonctions de base sont les "fonctions chapeau" de finesse,

$$\varphi_j(x) = \begin{cases} 0 & \text{si } x \leq x_{j-1} \\ \frac{x - x_{j-1}}{x_j - x_{j-1}} & \text{si } x \in \mathcal{U}_{j-1} = [x_{j-1}, x_j] \\ \frac{x_{j+1} - x}{x_{j+1} - x_j} & \text{si } x \in \mathcal{U}_j = [x_j, x_{j+1}] \\ 0 & \text{si } x \geq x_{j+1} \end{cases}$$

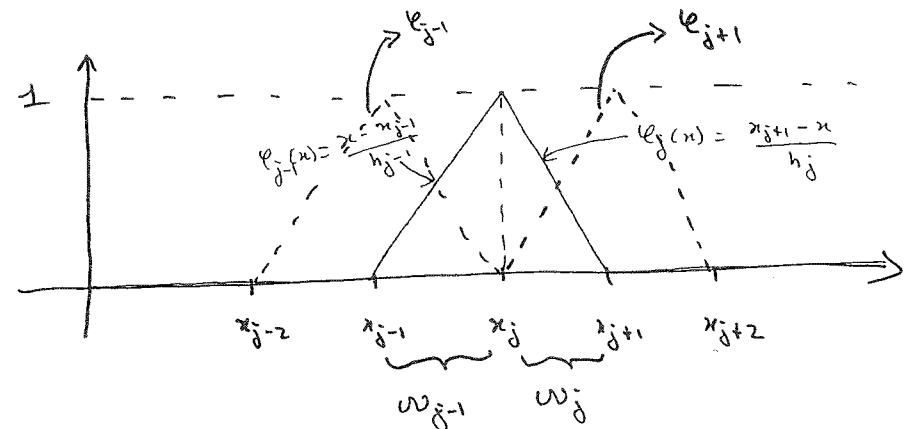


Fig:1 Représentation des fonctions de base $\varphi_j(x)$

* Dans l'intervalle $D_i = [x_i, x_{i+1}]$, on a q. qu seules 2 fonctions de bases ℓ_{i-1} et ℓ_i sont $\neq 0$, donc

$$\begin{aligned} P_h(x) \Big|_{D_i} &= d_i \ell_i(x) + d_{i+1} \ell_{i+1}(x) \\ &= d_i \frac{x_{i+1} - x}{x_{i+1} - x_i} + d_{i+1} \frac{x - x_i}{x_{i+1} - x_i} \end{aligned}$$

On retrouve bien le fait que $P_h(x)$ est un polynôme de degré 1 dans les intervalles D_i .

* Les fonctions $\ell_j(x)$ sont continues aux nœuds x_i , mais pas différentiables.

$$\Rightarrow P_h(x) \in C^0(D).$$

* Il est important de remarquer que $\ell_j(x)$ s'annule sur tous les nœuds sauf x_j :

$$\ell_j(x_i) = \delta_{ij} = \begin{cases} 1 & \text{si } i=j \\ 0 & \text{si } i \neq j \end{cases}$$

III-3 Application à l'interpolation de fonctions

Soit $f(x)$ une fonction à interpoler aux nœuds x_i .

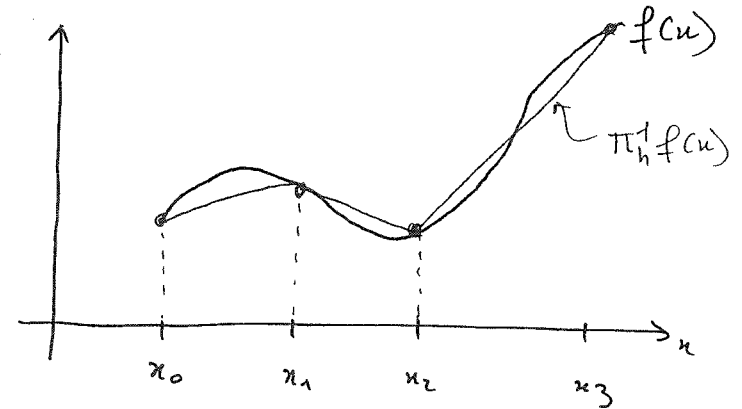
On veut la représenter par un polynôme p.m. IP_1 tel que:

$$\begin{aligned} f(x_i) &= P_h(x_i) \quad \text{pour } i=0, \dots, n \\ &= \sum_{j=0}^n d_j \ell_j(x_i) \\ &= d_i \quad \text{puisque } \ell_j(x_i) = \delta_{ij}. \end{aligned}$$

... de $f(x)$ s'est n'importe lequel :

$$\sum_{j=0}^n f(x_j) \ell_j(x).$$

On note généralement $\Pi_1^h f(x)$ et interpolant.



Ce type d'interpolation est couramment appelée "interpolation polynomiale par morceaux".

Remarques:

* L'avantage de l'interpolation polynomiale par intervalle est de supprimer les oscillations parasites de l'interpol polynomiale globale (cf TP).

* Les polynômes p.m. IP_1 ont de nombreuses applications dans la résolution des eqts de la mécanique (méth. des éléments finis en mécanique des solides).

* Le principal défaut des polyn. p.m. IP_1 est leur manque de régularité (continuité) aux nœuds x_i . Pour y remédier, on utilise généralement l'interpolant $\Pi_h^3 f$ qui a les propriétés:

- $\pi_n^0 f(x_i) = f(x_i)$ pour $i = 0, \dots, n$ (1)

- $\pi_3^h f \in \mathcal{P}_3(u_i)$

- $\pi_3^h f \in C^2(u)$: il est 2 fois différentiable en x_i .

Ce type d'interpolation est aussi connue sous le nom de splines cubiques. Il a été conçu au début des années 1950 dans les bureaux de Design de Citroën. Depuis, l'interpolation par splines a trouvé d'innombrables applications en CAO.

Intégration numérique

I) Exemples et motivation

Dans les problèmes de mécanique, on trouve souvent des intégrales dont le calcul est difficile ou impossible à effectuer analytiquement (pas de primitives).

Par exemple :

$$\int_0^1 e^{-x^2} dx, \quad \int_0^1 \cos x^2 dx.$$

Dans ce chapitre, on va décrire un certain nombre de méthodes de quadrature numérique, pour calculer de façon approchée une intégrale double. Ces méthodes s'appliquent aussi au cas où la fonction est seulement connue en quelques points du domaine d'intégration (forme tabulée).

II) Formules élémentaires

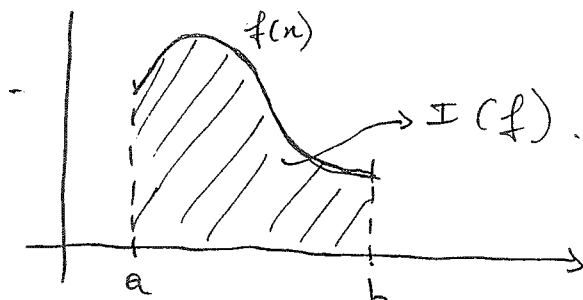
on veut calculer l'intégrale

$$I(f) = \int_a^b f(x) dx. \quad (1)$$

On va ici décrire quelques formules de quadrature, notées $I_q(f)$, telles que

$$I_q(f) \stackrel{?}{=} I(f). \quad (2)$$

↓
signifie: on approche



I) Formule du rectangle (ou du point milieu)

On remplace f par

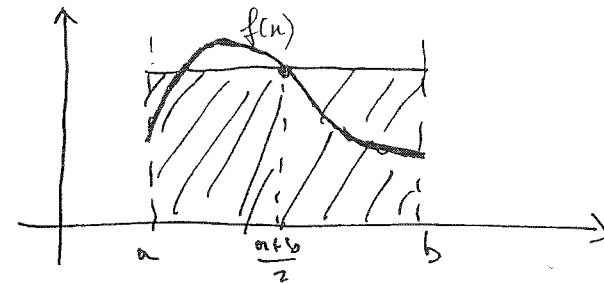
$$\pi_0(f) = f\left(\frac{a+b}{2}\right),$$

la valeur constante prise par f au milieu de l'intervalle $I = [a, b]$. On obtient

$$\int_a^b \pi_0(f) = \int_a^b f\left(\frac{a+b}{2}\right) dx = (b-a) f\left(\frac{a+b}{2}\right).$$

~~On a donc~~ On définit donc

$$I_0(f) = (b-a) f\left(\frac{a+b}{2}\right) \stackrel{?}{=} I(f)$$



2) Formule du trapèze

On remplace $f(u)$ par $\pi_1(f)$, le polynôme d'interpolation de degré 1 aux points a et b :

$$\pi_1(f) = \frac{(u-a)f(b)}{b-a} + \frac{(u-b)f(a)}{(a-b)},$$

par application de l'interpolation de Lagrange.

On définit alors

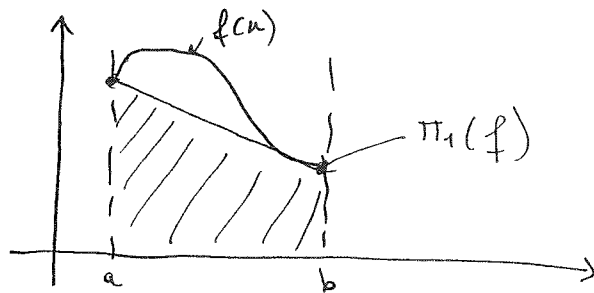
$$I_1(f) = \int_a^b \pi_1(f) du \approx I(f).$$

Calculons $I_1(f)$:

$$\int_a^b \pi_1(f) = \frac{f(b)}{b-a} \int_a^b (u-a) du + \frac{f(a)}{a-b} \int_a^b (u-b) du,$$

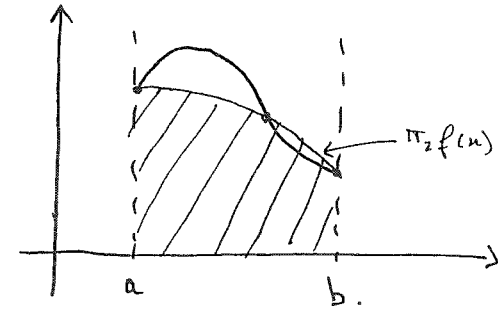
etc..., après calcul:

$$I_1(f) = \frac{(b-a)}{2} f(b) + \frac{(b-a)}{2} f(a) \quad (7)$$



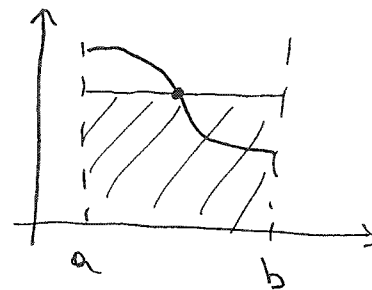
3) Formule de Simpson

On remplace $f(u)$ par $\pi_2 f(u)$ aux points $x_0=a$, $x_1=\frac{a+b}{2}$, $x_2=b$. (déjà vu en TD).

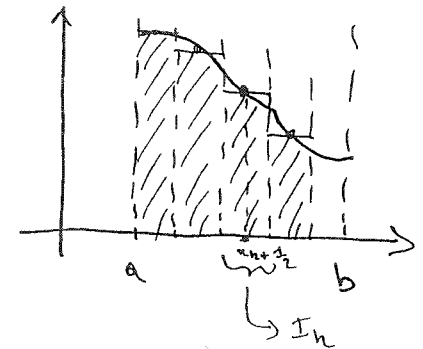


$$I_2(f) = \frac{b-a}{6} [f(a) + 4f(\frac{a+b}{2}) + f(b)] \quad (8)$$

III Formules Composites



formule du rectangle



formule composite du rectangle.

(3)

1) généralités
 Pour obtenir une meilleure approximation numérique, on pourra utiliser des polyg. de Lagrange de degré de + en + grand.
 Néanmoins, on se rendant que l'interpolation de Lagrange peut donner des résultats oscillants: en pratique, on utilise rarement des quadratures utilisant des interpolants de Lagrange de degré ≥ 2 .

Une meilleure stratégie est de construire des formules de quadrature composées. Pour cela, on se coupe le domaine $\Omega =]a, b[$ en n sous-domaines

$$\Omega_k =]x_k, x_{k+1}[, \quad k = 0, \dots, n-1$$

telle que :

- Leur longueur est $h_k = x_{k+1} - x_k$
 - Le milieu de Ω_k est noté $x_{k+\frac{1}{2}} = \frac{x_k + x_{k+1}}{2}$
- (e On note $x_0 = a$, $x_n = b$)

L'intégrale à calculer se décompose comme :

$$\int_a^b f(x) dx = \sum_{k=0}^{n-1} \int_{\Omega_k} f(x) dx \quad (9)$$

Dans chaque sous-domaine, on peut maintenant appliquer une formule de quadrature élémentaire, p. ex. la formule du rectangle :

$$\int_{\Omega_k} f(x) dx \approx h_k f(x_{k+\frac{1}{2}}) \quad (10)$$

On obtient donc la formule de quadrature composée du rectangle :

$$\int_a^b f(x) dx \approx \sum_{k=0}^{n-1} h_k f(x_{k+\frac{1}{2}}) \quad (11)$$

Cas particulier important :

Dans le cas où chaque intervalle est de longueur h constante

On appelle $I_{0,h}$ la formule composée du rectangle, telle

$$I_{0,h} = h \sum_{k=0}^{n-1} f(x_{k+\frac{1}{2}})$$

La longueur h est appelée pas de discrétisation.

2) Autres formules composées

• Formule du Trapèze (notée $I_{1,h}$)

Dans chaque intervalle Ω_k , on applique la formule du trapèze

$$\int_{\Omega_k} f(x) dx \approx \frac{h_k}{2} (f(x_k) + f(x_{k+1}))$$

Pour $h_k = h$ constant, on obtient finalement :

$$I_{1,h} = \frac{h}{2} \left[f_0 + 2 \sum_{k=1}^{n-1} f(x_k) + f(x_n) \right] \quad (12)$$

• Formule de Simpson (notée $I_{2,h}$) :

$$(13) \quad \int_{\Omega_k} f(x) dx \approx \frac{h_k}{6} [f(x_k) + 4f(x_{k+\frac{1}{2}}) + f(x_{k+1})]$$

et donc, pour $h_k = h$, on obtient :

$$(14) \quad I_{2,h} = \frac{h}{3} \left[f_0 + 2 \sum_{k=1}^{n-1} f(x_k) + 4 \sum_{k=1}^{n-1} f(x_{k+\frac{1}{2}}) + f(x_n) \right]$$

IV) Analyse des erreurs

1) Généralités

But de cette partie: Quantifier la précision d'une formule numérique, c'est à dire mesurer l'erreur:

$$E_h = |I(f) - I_h(f)| \quad (15)$$

Cette erreur doit tendre vers zéro lorsque $h \rightarrow 0$.

(Soit, problème: la formule n'est pas constante).

Dans la suite de cette partie, on verra que les formules de quadrature de la partie II sont telles que:

$$E_h \approx C h^q (= O(h^q)) \quad (\text{avec } C: \text{constante})$$

où $q \geq 0$ est l'ordre de la méthode. Plus q est grand, plus l'erreur sera petite $\Rightarrow$ une méthode d'ordre élevé est plus précise qu'une méthode du premier ordre ($q=1$).

Cette analyse des erreurs est effectuée en utilisant les développements de Taylor:

$$f(a+h) = f(a) + h f'(a) + \frac{h^2}{2} f''(a) + \dots + \frac{h^q}{q!} f^{(q)}(a) + O(h^{q+1})$$

Définition (petit "oh" et grand "oh")

intégrer dans la suite du cours.

- $f = O(g)$ s'il existe $k > 0$ t.q.: $\left| \frac{f}{g} \right| \leq k$ lorsque $h \rightarrow 0$.

- $f = o(g)$ si:

$$\left| \frac{f}{g} \right| \rightarrow 0 \text{ lorsque } h \rightarrow 0.$$

2) Analyse des erreurs pour la formule du rectangle:

Sur l'intervalle $[x_n, x_{n+1}]$, la formule du rectangle (pas h constant) s'écrit:

$$\int_{x_n}^{x_{n+1}} f(u) du \approx h f(x_{n+\frac{1}{2}}).$$

On utilise un D.L. de $f(u)$ au point $x_{n+\frac{1}{2}}$:

$$f(u) = f(x_{n+\frac{1}{2}}) + (u - x_{n+\frac{1}{2}}) f'(x_{n+\frac{1}{2}}) + \frac{(u - x_{n+\frac{1}{2}})^2}{2} f''(x_{n+\frac{1}{2}}) + \dots$$

On intègre ce D.L. dans $[x_n, x_{n+1}]$:

$$\begin{aligned} \int_{x_n}^{x_{n+1}} f(u) du &= \int_{x_n}^{x_{n+1}} f(x_{n+\frac{1}{2}}) du + f'(x_{n+\frac{1}{2}}) \underbrace{\int_{x_n}^{x_{n+1}} (u - x_{n+\frac{1}{2}}) du}_{=0 \text{ (fonction impaire sur } [x_n, x_{n+1}])} \\ &\quad + \frac{f''(x_{n+\frac{1}{2}})}{2} \int_{x_n}^{x_{n+1}} (u - x_{n+\frac{1}{2}})^2 du + \dots \end{aligned}$$

On remarque que les termes de puissance impaire ont une intégrale nulle, et il reste donc:

$$\underbrace{\int_{x_n}^{x_{n+1}} f(u) du}_{\text{intégrale exacte}} = \underbrace{h f(x_{n+\frac{1}{2}})}_{\text{formule du rectangle}} + \frac{h^3}{24} f''(x_{n+\frac{1}{2}}) + \frac{h^5}{1920} f^{(4)}(x_{n+\frac{1}{2}}) + \dots \quad (\text{terme d'erreur } E_h \text{ sur } [x_n, x_{n+1}])$$

~~Si f admet des dérivées d'ordre ≥ 2~~

- Si toutes les dérivées d'ordre ≥ 2 de $f(u)$ sont nulles, alors $E_h = 0 \quad \forall h$. La formule du rectangle est exacte pour les polynômes de degré 1.

- Le terme dominant de l'erreur est telle que $E_h = 0$.
Donc la formule du rectangle est précise au troisième ordre sur un intervalle I_h .

Pour obtenir l'ordre d'une formule composite, il faut sommer ⁽¹⁶⁾ sur tous les intervalles :

$$\underbrace{\int_a^b f(u) du}_{I(f)} = \underbrace{\sum_{h=0}^{n-1} h f(x_{h+\frac{1}{2}})}_{I_{1,h}(f)} + \underbrace{\sum_{h=0}^{n-1} \frac{h^3}{24} f''(x_{h+\frac{1}{2}})}_{E_h} + \dots$$

Pour simplifier la démonstration, on écrit que le terme dominant de l'erreur est :

$$E_h = \sum_{h=0}^{n-1} O(h^3) = h^2 O(h^2) \sum_{h=0}^{n-1} h = (b-a) O(h^2) = O(h^2).$$

La formule du rectangle est précise au second ordre sur l'intervalle tout entier $D = [a, b]$.

3) Formule du trapèze

$$\int_{x_n}^{x_{n+1}} f(u) du \approx \frac{h}{2} (f(x_n) + f(x_{n+1}))$$

Pour déterminer l'ordre de cette formule, on va utiliser au maximum ce qu'on a vu précédemment.

On considère les D.L. suivants (autour de $x_{n+\frac{1}{2}}$) :

$$f(u) = f(x_{n+\frac{1}{2}}) - \frac{h}{2} f'(x_n) + \frac{1}{2} \left(-\frac{h}{2}\right)^2 f''(x_n) + \frac{1}{6} \left(-\frac{h}{2}\right)^3 f'''(x_n) + O(h^4)$$

$$f(x_{n+\frac{1}{2}}) = f(x_{n+\frac{1}{2}}) + \frac{h}{2} f'_{+\frac{1}{2}}(x_n) + \frac{1}{2} \left(\frac{h}{2}\right)^2 f''_{+\frac{1}{2}}(x_n) + \dots$$

La $\frac{1}{2}$ somme de ces 2 expressions conduit à :

$$f(x_{n+\frac{1}{2}}) = \frac{1}{2} (f(x_n) + f(x_{n+1})) - \frac{1}{8} h^2 f''(x_{n+\frac{1}{2}}) - \frac{1}{384} h^4 f^{(4)}(x_{n+\frac{1}{2}}) + \dots$$

On injecte cette expression dans l'estimation d'erreur (16) du rectangle, pour obtenir :

$$\int_{x_n}^{x_{n+1}} f(u) du = \frac{h}{2} (f(x_n) + f(x_{n+1})) - \frac{h^3}{12} f''(x_{n+\frac{1}{2}}) - \frac{h^5}{480} f^{(4)}(x_{n+\frac{1}{2}}) + \dots$$

Conclusions:

- Elle est précise au 3^{ème} ordre
- Son degré d'exactitude est 1
- Dans le cas de la formule composite, on a même comme particulier que :

$$E_h = (b-a) O(h^2),$$

3) Formule de Simpson

- cf. +0 et ci-dessous.

(10)

IV) Résumé

Pour un intervalle $I_h = [x_n, x_{n+1}]$, toutes les formules d'quadrature peuvent se mettre sous la forme :

$$\int_{x_n}^{x_{n+1}} f(u) du \approx I(f) = \sum_{i=0} d_i f(\xi_i)$$

avec :

- les $\{\xi_i\}$ sont les nœuds de la formule
- Les $\{d_i\}$ ——— poids ———

(ex: Rectangle: $m=0$, $d_i = h_n$, $\xi_i = x_{n+\frac{1}{2}}$).

Pour obtenir une formule composite, on somme sur les intervalles I_k , $k=0, \dots, n-1$.

Formule	Poids	Nœuds	r	Q	Q_{com}
Rectangle	h_n	$x_{n+\frac{1}{2}}$	1	3	
Trapeze	$\{\frac{h_n}{2}, \frac{h_n}{2}\}$	$\{x_n, x_{n+1}\}$	1	3	
Simpson	$\{\frac{h_n}{6}, \frac{2}{3}h_n, \frac{h_n}{6}\}$	$\{x_n, x_{n+\frac{1}{2}}, x_{n+1}\}$	3	5	

- r : degré d'exactitude
(intégr. exactement des polynômes de degré r ou plus)
- Q : ordre de précision.

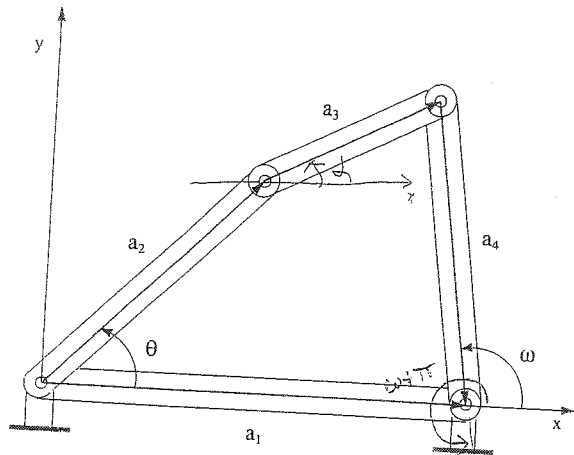


Fig. 1: Configuration de 4 tiges articulées.
On remarque que la tige $\vec{a}_1$ est immobile (ses 2 extrémités sont en appui solide), tandis que les 3 autres sont mobiles.

① Exemple et motivation

On considère le système mécanique de la fig. 1: on cherche à savoir déterminer l'angle θ entre la tige $\vec{a}_1$ et $\vec{a}_2$ lorsque la tige $\vec{a}_4$ est perpendiculaire avec un angle $\omega + \pi = \text{mes}(\vec{ou}, \vec{a}_4)$ avec l'horizontale.

Dans une première étape, on met en équation le problème:
On part de l'identité vectorielle:

$$\vec{a}_1 = \vec{a}_2 + \vec{a}_3 + \vec{a}_4,$$

qu'on projette sur les axes (Ox) et (Oy) , soit:

$$a_1 = a_2 \cos \theta + a_3 \cos \delta + a_4 \cos(\omega + \pi) \quad (1a)$$

$$0 = a_2 \sin \theta + a_3 \sin \delta + a_4 \sin(\omega + \pi) \quad (1b)$$

avec $\delta = \text{mes}(\vec{ou}, \vec{a}_3)$

système (1), ce qui entraîne:

$$(a_1 - a_2 \cos \theta + a_4 \cos \omega)^2 + (a_4 \sin \omega - a_2 \sin \theta)^2 = a_3^2$$

Avec un peu de trigonométrie, cette équation devient:

$$\underbrace{\frac{a_1}{a_2} \cos \omega - \frac{a_1}{a_4} \cos \theta - \cos(\omega - \theta) + \frac{a_1^2 + a_2^2 - a_3^2 + a_4^2}{2a_2 a_4}}_{f_w(\theta)} = 0$$

Le problème de mécanique s'est donc sous la forme:

$$\text{Trouver } \theta \in \mathbb{R} \text{ tel que: } f_w(\theta) = 0, \quad (2)$$

$\hookrightarrow [-\pi, \pi]$

où $w \in \mathbb{R}$ est un paramètre donné. On remarque que la fonction $f_w: \mathbb{R} \rightarrow \mathbb{R}$ est non linéaire:

$$\forall \alpha, \beta \quad f_w(\alpha \theta_1 + \beta \theta_2) \neq \alpha f_w(\theta_1) + \beta f_w(\theta_2).$$

Pour w quelconque, on ne peut trouver une solution analytique au problème (2). Il faut donc utiliser des méthodes numériques de résolution d'équation (ou système d'éqs) non-linéaire.

I) Fonctions linéaires

1/2

Une équation linéaire :

$$f(u) = 0. \quad (1)$$

est une équation qui vérifie la superposition des solutions :

si x_1 et x_2 vérifient (1), alors leur combinaison linéaire $\alpha x_1 + \beta x_2$ ($\alpha, \beta \in \mathbb{R}$) vérifie (1) tout autant de (1) :

$$(2) \quad f(\alpha x_1 + \beta x_2) = \alpha \underbrace{f(x_1)}_{=0} + \beta \underbrace{f(x_2)}_{=0} = 0.$$

Exemples de fct linéaires :

en 1D : $f(x) = 3x$.

en nD : $\underline{x} = \begin{pmatrix} x_1 \\ y \end{pmatrix}$, $\begin{cases} 3x_1 + 2y = 1 \\ -x_1 + y = 4 \end{cases}$ (Système linéaire).

Propriétés importantes :

— Par le principe de superposition de solutions, on peut toujours calculer le solution analytique de α équations linéaires, par des méthodes systématiques (ex : cours d'EDO du 1er et 2nd ordre).

— Malheureusement, la plupart des équations de la physique sont non-linéaires.

II) Fonc / équations linéaires:

2/2

Une fonction ~~linéaire~~ non linéaire ne vérifie pas la superposition de solutions (2).

Exemple : $f(x) = x^2$, $f(x) = \sin(x)$

on a bien $f(x_1 + x_2) \neq f(x_1) + f(x_2)$.

En conséquence, il n'existe pas de méthodes systématiques pour résoudre les équations non-linéaires, sauf dans des cas très particuliers, p. ex :

- polynômes du 2nd ordre : $ax^2 + bx + c = 0$ et 3^{em} ordre, pour les ordres impaires 2.
- EDO : méthodes de Riccati, etc...

Pour résoudre des eq^{ts} non-linéaires issus de modèles physiques complexes (niveau Master) il faut utiliser des méthodes numériques.

II) Généralités

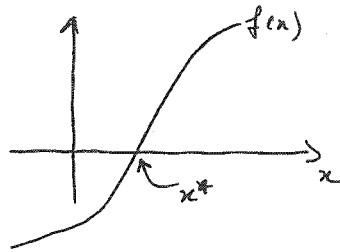
II-1 Positionnement du problème

Soit $f: \mathbb{R} \rightarrow \mathbb{R}$ une fonction continue. On cherche les zéros x^* , tel que:

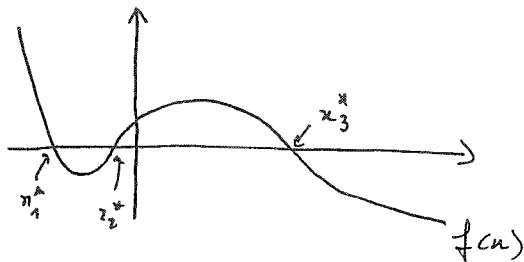
$$f(x^*) = 0.$$

II-2 Exemples rencontrés

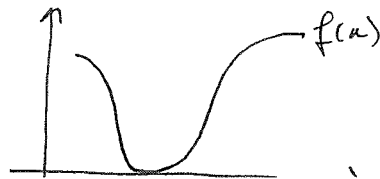
a) Le plus simple: un seul zéro



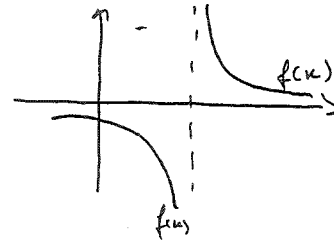
b) + plusieurs zéros



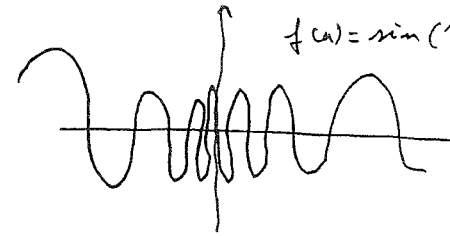
c) Zéros multiples



d) Exemples pathologiques (à éviter)



fonctions discontinues



fonction avec de nombreuses racines infinies.

II-3 Stratégie à suivre:

a) Localiser grossièrement (ex: graphiquement) le(s) zéro(s) x_0 cette évaluation qui

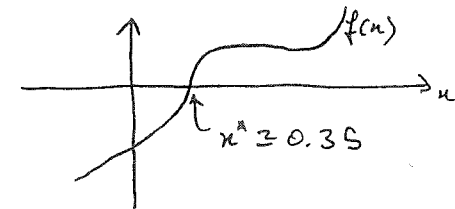
b) Construire à partir de x_0 une suite $\{x_1, x_2, \dots\}$ telle que:

$$\lim_{n \rightarrow +\infty} x_n = x^*$$

II-4 Exemple test:

Dans la suite des cours, on va considérer l'exemple:

Trouver le zéro de $f(x) = \sin(2x) - 1 + x$ (1)



On voit graphiq^t que $x_0 \approx 0.35$ est unique. On peut donc raisonnable^t utiliser $x_0 \in [0, 0.5]$

il reste maintenant à définir la suite:

III méthode de la dichotomie (ou bisection)

NL

Une 1^{re} méthode purement "géométrique" pour définir $\{x_n, n \geq 0\}$ est de remarquer que si on trouve un intervalle $I_0 = [a_0, b_0]$ t. q :

$$f(a_0)f(b_0) < 0 \quad (3)$$

alors f change de signe dans I_0 , et donc $x_n \in I_0$ puisqu' f est continue.

On peut alors subdiviser I_0 en intervalles plus petits (faire la dichotomie)

petits $I_n = [a_n, b_n]$, en utilisant le test (3), de façon à encadrer la valeur x_n par des intervalles de + en plus petits: ainsi le centre $x_n = \frac{a_n + b_n}{2} \rightarrow x_n$ lorsque $n \rightarrow \infty$

L'algo de la bisection prend la forme suivante:

1) Initialisation: On choisit $I_0 = [a_0, b_0]$, tel que (3) est v.
 $x_0 = \frac{a_0 + b_0}{2}$.

(Pour le pbme test (3), on peut choisir $I_0 = [0, 1] \Rightarrow x_0 = \frac{1}{2}$).

L'erreur commise s'écrit:

$$e_0 = |x_0 - x_n| \leq \frac{b_0 - a_0}{2}$$

2) Détermination de $I_1 = [a_1, b_1]$ (et donc $x_1 = \frac{a_1 + b_1}{2}$)

Suivant le signe du produit:

$$f(x_0)f(a)$$

Cas 1 $f(x_0)f(a) < 0$: alors nécessairement $x^* \in [a, x_0]$.

et on définit $I_1 = [a, x_0]$ dont le centre est:

$$x_1 = \text{milieu}(I_1) = \frac{a + x_0}{2}$$

Cas 2: $f(x_0)f(a) > 0$: nécessairement $x^* \notin I_1 = [x_0, b]$, NL

- dont le centre est:

$$x_1 = \text{milieu}(I_1) = \frac{x_0 + b}{2}$$

3) Test de convergence: On vérifie si $f(x_1) = 0$.

Si ce n'est pas le cas, on sait quand même que l'erreur est telle que:

$$e_1 = |x_1 - x_0| \leq \frac{b - a}{4}$$

et on revient à l'itération 2) pour définir x_2 (et I_2).

De proche en proche, on définit la suite $\{x_n, n \geq 0\}$ qui est telle que:

$$e_n = |x_n - x_n| \leq \frac{|x_{n-1} - x_0|}{2} \leq \frac{b - a}{2^{n+1}} \quad (4)$$

Donc $e_n \rightarrow 0$ lorsque $n \rightarrow \infty$. Clairement, si on a trouvé un bon intervalle de départ I_0 , la méthode de bisection converge toujours.

Comme on le verra en TD, la convergence de la méthode est très lente. Dans la suite de ce cours, on va définir des méth. plus performantes.

$$|\phi(x_n) - \phi(x_*)| \leq C |x^n - x^*|, \quad 0 \leq C < 1$$

Les funct. d'iteration qui verifient cette inegalite sont appelees fonctions contractantes. C'est en general difficile a prouver qu'une fonction est contractante. A la place, nous allons utiliser le thm suivant pour obtenir un critere pratique assurant la convergence d'un algo. du point fixe :

Thm (Convergence). Soit ϕ une fonction d'iteration contractante. Si ϕ verifie la condition de stabilite :

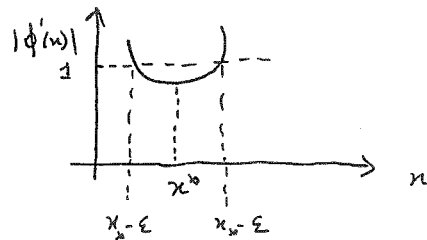
$$|\phi'(x_*)| < 1, \quad (11)$$

alors pour x_0 "suffisamment" proche de x_* , l'iteration du point fixe va converger.

Idee de la dem.

$$x^{n+1} - x_* = \phi(x_n) - \phi(x_*) = \phi'(\eta)(x_n - x_*)$$

pour $\eta \in [x_n, x_*]$, en utilisant un D.L. avec reste epoque en $x = x_*$.



on suppose x_n proche de x_* , donc si $|\phi'(x_*)| < 1$, on a $|\phi'(\eta)| = C < 1$

On a donc :

$$|x^{n+1} - x_*| \leq \max_{x \in [x_n, x_*]} |\phi'(x)| |x^n - x_*|$$

soit $|\phi'(x_n)| < 1$, alors par continuite :

$$\exists \epsilon > 0 \text{ tq } |\phi'(x)| < 1, \quad \forall x \in]x^* - \epsilon, x^* + \epsilon[.$$

Si x_n est suffisamment proche de x^* , alors

$$[x_n, x_*] \subset]x^* - \epsilon, x^* + \epsilon[,$$

et donc

$$\max_{x \in [x_n, x_*]} |\phi'(x)| = C < 1, \text{ soit :}$$

$$|x^{n+1} - x_*| \leq C |x^n - x_*| \text{ avec } 0 < C < 1.$$

□

Définition (l'ordre d'une methode)

Une methode est convergente d'ordre q si elle verifie :

$$|x_{n+1} - x_*| \leq C |x_n - x_*|^q$$

- Si $q=1$ et $0 < C < 1$: Convergence lineaire.
- Si $q=2$: Convergence quadratique.
- Si $q=3$: ——— cubique.

Remarques :

* D'après Eq. (4), la methode de bisection est convergente au 1er ordre.

* D'après le thm pcdt :

— si $0 < |\phi'(x_*)| < 1$: une methode du point fixe a une cv. lineaire.

— si $\phi'(x_*) = 0$: on verra que la cv est quadratique.

Page (8) du Cours :

Définition (Stabilité)

Une fonction d'itération Φ est stable si :

$$|\Phi'(x_*)| < 1 \quad (11)$$

Théorème (Convergence)

Si Φ est stable et contractante, alors pour x_0 "suffisamment" proche de x_* , la méthode du point fixe va converger.

Démo simplifiée

On fait un DL en x_* :

$$\underbrace{\Phi(x_{n+1})}_{x_{n+1}} \approx \underbrace{\Phi(x_*)}_{x_*} + \Phi'(x_*) (x_n - x_*)$$

$$\Rightarrow x_{n+1} - x_* \approx \Phi'(x_*) (x_n - x_*)$$

$$\Rightarrow E_{n+1} \approx |\Phi'(x_*)| E_n \quad \text{avec } E_n = |x_n - x_*|$$

$$\text{Si on pose } C \equiv |\Phi'(x_*)| < 1$$

alors la suite géométrique

$$E_{n+1} \approx C E_n = C^2 E_{n-1} = \dots = C^{n+1} E_0$$

converge vers 0.

Synonymes: Newton-Raphson, meth. des tangentes.

On revient au problème original :

$$f(x) = 0$$

dont la solution est x_* .

On initialise l'algorithme par la condition initiale x_0 et, si elle est suffisamment proche de x_* , alors l'approximation linéaire suivante est valide : (D.L. en x_0)

$$f(x) \approx L(x) = f(x_0) + f'(x_0)(x - x_0). \quad (12)$$

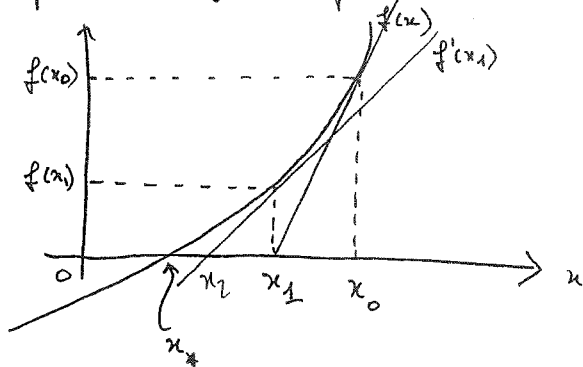
On définit x_1 comme la racine de $L(x) = 0$, soit :

$$x_1 = x_0 - \frac{f(x_0)}{f'(x_0)}.$$

De proche en proche, on définit l'algorithme Newton :

$$\begin{cases} \bullet x_0 \in \mathbb{R} \text{ donnée, tq } f'(x_0) \neq 0. \\ \bullet x_{n+1} = x_n - \frac{f(x_n)}{f'(x_n)}. \end{cases} \quad (13)$$

Interprétation géométrique :



Selon (12) : x_{n+1} est l'intersection entre

- l'axe Ox
- la droite passant par $f(x_n)$ de coeff. directeur $f'(x_n)$.

D'après (13), la méthode de Newton est une méthode du point fixe avec :

$$\phi(x) = x - \frac{f(x)}{f'(x)}. \quad (14)$$

• Vérifions sa consistance :

$$x^* = x^* - \frac{f(x^*)}{f'(x^*)} \Leftrightarrow f(x^*) = 0 \text{ si } f'(x^*) \neq 0.$$

On en déduit (provisoirement) que :

- 1) l'algo de Newton est consistant. D'après le thm¹, il converge si $|\phi'(x_n)| < 1$.
- 2) Si $f'(x_*) = 0$ (x^* : zéro double) $\Rightarrow$ problème de convergence : il faut utiliser des meth. de Newton modifiées (cf T.D).

• Examinons maintenant la cv de l'algo de Newton dans le cas d'un zéro simple ($f(x^*) = 0, f'(x_*) \neq 0$).

Après calcul, on a :

$$\phi'(x) = \frac{f(x)f''(x)}{[f'(x)]^2} \Rightarrow \phi'(x_*) = 0.$$

Donc la meth. de Newton converge, au moins linéairement. Quel est son ordre réel ? On reprend les grandes lignes de la démon du thm¹ : En utilisant un D.L. avec reste exacte en $x = x_*$:

$$x_{n+1} - x_* = \phi(x_n) - \phi(x_*) = \underbrace{\phi'(x_*)}_{=0} (x_n - x_*) + \frac{\phi''(\eta)}{2} (x_n - x_*)^2$$

pour $\eta \in [x_n, x_*]$.

IV-1 Généralités

On revient au pbme test (1), mt:

$$\text{Trouver } x^* \text{ tq. } \underbrace{\sin(2x^*) - 1 + x^*}_{f(x^*)} = 0. \quad (5)$$

On peut écrire ce pbme comme:

$$\text{Trouver } x^* \text{ tq. } x^* = \underbrace{1 - \sin(2x^*)}_{\phi_1(x^*)} \quad (6)$$

On peut aussi écrire:

$$\begin{aligned} \text{Trouver } x^* \text{ tq. } \sin(2x^*) &= 1 - x^* \\ \Leftrightarrow x^* &= \underbrace{\frac{1}{2} \arcsin(1 - x^*)}_{\phi_2(x^*)} \quad (7) \end{aligned}$$

Dans les 2 cas, il semble naturelle de formuler la suite:

$$x^{n+1} = \phi_i(x^n)$$

Si la suite converge, on aura trouvé x_* tel que:

$$x_* = \phi_i(x_*) \Leftrightarrow f(x_*) = 0.$$

On aura alors transformé la recherche du zéro de f en la recherche du point fixe de $\phi_i(x)$.

Définition: Une méthode qui cherche les racines du problème (5) avec les itérations

$$\begin{cases} x_0 \text{ donnée} \\ x_{n+1} = \phi(x_n) \end{cases} \quad (8)$$

est appelée méthode du point fixe. La fonction $\phi(x)$ est appelée fonction d'itération.

forme en x a vu pour l'exemple test (1), il y a d'innombrables façons de définir $\phi(x)$.
Le plus simple est de poser:

$$\phi(x) = x - f(x) \quad (9)$$

c'est la méthode de Picard. Plus généralement, on pose:

$$\phi(x) = x + q f(x) \quad (10)$$

avec $q \in \mathbb{R}^*$. On peut même choisir $q = q(x)$ tant que $q(x)$ ne s'annule pas.

Pour une fonction d'itération quelconque $\phi(x)$, il faut s'assurer de la consistance avec le problème de base (1):

Définition (consistance)

On dit que la méth. du point fixe (8) est consistante avec le pbme (1) si:

$$x = \phi(x) \Leftrightarrow f(x) = 0.$$

exple: Les itérations (6) et (7) sont consistantes avec le pbme test (1), mais pas $\phi_3(x) = \frac{1}{2} [1 - \sin(2x)]$

IV-2

Analyse de la convergence

Une fois que la consistance d'une méth. du point fixe est établie, il faut s'assurer de sa convergence: Il faut que l'erreur $e_n = |x_n - x_*|$ diminue à chaque itération, mt:

$$e^{n+1} = |x^{n+1} - x_*| \leq c |x^n - x_*|, \text{ avec } \boxed{0 \leq c < 1}$$

$$\eta \in [x_n, x_{n+1}]$$

$$|x_{n+1} - x_*| \leq C |x_n - x_*|^2 \quad (15)$$

La meth. de Newton a une conv. quadratique pour un zéro simple.

IV-4 Résumé

Méthode des point fixe : $x^{n+1} = \phi(x_n)$ pour résoudre $f(x) = 0$ sur l'intervalle I .

1) Etudier l'existence : $x_* = \phi(x_*) \Leftrightarrow f(x_*) = 0$.

2) Etudier la stabilité : Pour x_0 "suffisamment" proche de x_* , on a :

$$\begin{cases} \text{Si } 0 < |\phi'(x_*)| < 1 : & \text{cv d'ordre 1} \\ \text{estimation d'erreur : } e_{n+1} \leq C e_n, & \text{avec } C = \max_{x \in I} |\phi'(x)| \\ \text{Si } \phi'(x_*) = 0 : & \text{cv d'ordre 2} \\ e_{n+1} \leq C e_n^2, & \text{avec } C = \frac{1}{2} \max_{x \in I} |\phi''(x)| \end{cases}$$

Cas particuliers :

• Méthode de Picard :

$$\phi(x) = x - f(x)$$

• Méthode de Newton-Corde : $\phi(x) = x - \frac{f(x)}{f'(x_0)}$

• Méthode de Newton :

$$\phi(x) = x - \frac{f(x)}{f'(x)}$$

rq : ordre 2 pour racine simple ($f(x_*) = 0, f'(x_*) \neq 0$).

• — 1 — double ($C \notin \text{TD}$).
($f(x_*) = 0, f'(x_*) = 0, f''(x_*) \neq 0$).

On va généraliser les algos. sur deux cas 1D ou cas multi-D, par ex :

$$\begin{cases} \text{Trouver les } (x, y) \text{ zéros de} \\ \text{syst. N-2 :} \end{cases} \begin{cases} x^2 + y^2 - 1 = 0 & (1) \\ 2x + y - 1 = 0 & (2) \end{cases}$$

Graphiquement, on voit 2 racines $x_1^* = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$ et $x_2^* = \begin{pmatrix} 4/5 \\ -3/5 \end{pmatrix}$ qui sont l'intersection du cercle (Eq (1)) et de la droite (Eq (2)).

V-I Rappels de maths

a) Analyse vectorielle / matricielle

$$\text{Soit } x = \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} \in \mathbb{R}^n.$$

On définit le produit scalaire :

$$(x, y) = \sum_{i=1}^n x_i y_i$$

qui induit une norme vectorielle euclidienne :

$$\|x\|_2 = \sqrt{(x, x)}^{1/2}$$

Soit $\underline{A} \in \mathbb{R}^{n, n}$ ^{siq. n, n colonnes} (matrice carrée de taille $n \times n$)

On définit la norme matricielle euclidienne :

$$\|\underline{A}\|_2 = \sup_{\substack{x \in \mathbb{R}^n \\ x \neq 0}} \frac{\|\underline{A}x\|_2}{\|x\|_2} \quad (16)$$

Cette norme matricielle est peu pratique à utiliser. A la place, nous allons plutôt utiliser la p-norm :

$$\|\underline{A}\|_2 \geq \max_{i=1, \dots, n} |\lambda_i| \quad (17)$$

où λ_i sont les val. propres de $\underline{A}$.

$$\rho(\underline{A}) = \max_{i=1, \dots, n} |\lambda_i|, \quad (18)$$

Il y a une autre relation très utile: D'après (16), on a

$$\|\underline{A} \underline{x}\|_2 \leq \|\underline{A}\|_2 \|\underline{x}\|_2 \quad (19)$$

b) calcul différentiel

* fonction $f: \mathbb{R} \rightarrow \mathbb{R}$

Formule de Taylor: $f(x+h) = f(x) + h \underbrace{f'(x)}_{\text{Dérivée 1ère}} + \dots$

* fonction $f: \mathbb{R}^n \rightarrow \mathbb{R}$

$\underline{x} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}$ et $f(\underline{x})$ est un réel.

Formule de Taylor: $f(\underline{x} + \underline{h}) = f(\underline{x}) + \underbrace{(\nabla f(\underline{x}), \underline{h})}_{\substack{\text{différentielle} \\ \text{au sens notation } D(f)(\underline{x})}} + \dots \quad (20)$

$$\text{avec } \nabla f(\underline{x}) = \begin{pmatrix} \frac{\partial f}{\partial x_1}(\underline{x}) \\ \vdots \\ \frac{\partial f}{\partial x_n}(\underline{x}) \end{pmatrix} \quad (21)$$

exple: Dans $\mathbb{R}^2$, $\underline{x} = \begin{pmatrix} x \\ y \end{pmatrix}$, $f(\underline{x}) = f(x, y) = x^2 + y^2 - 1$

$$\nabla f(\underline{x}) = \begin{pmatrix} \frac{\partial f}{\partial x} \\ \frac{\partial f}{\partial y} \end{pmatrix} = \begin{pmatrix} 2x \\ 2y \end{pmatrix}$$

* fonction $f: \mathbb{R}^n \rightarrow \mathbb{R}^p$

$$\underline{x} \in \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} \text{ et } \underline{f}(\underline{x}) = \begin{pmatrix} f_1(\underline{x}) \\ \vdots \\ f_p(\underline{x}) \end{pmatrix} \in \mathbb{R}^p \quad (13)$$

Formule de Taylor: $\underline{f}(\underline{x} + \underline{h}) = \underline{f}(\underline{x}) + [\underline{\nabla} \underline{f}(\underline{x})] \underline{h} + \dots \quad (22)$

La différentielle est cette fois une matrice $\in \mathbb{R}^{p, n}$

$$[\underline{\nabla} \underline{f}(\underline{x})] = \begin{pmatrix} \nabla f_1(\underline{x}) \\ \nabla f_2(\underline{x}) \\ \vdots \\ \nabla f_p(\underline{x}) \end{pmatrix} = \begin{pmatrix} \frac{\partial f_1}{\partial x_1} & \frac{\partial f_1}{\partial x_2} & \dots & \frac{\partial f_1}{\partial x_n} \\ \frac{\partial f_2}{\partial x_1} & \frac{\partial f_2}{\partial x_2} & \dots & \frac{\partial f_2}{\partial x_n} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{\partial f_p}{\partial x_1} & \dots & \dots & \frac{\partial f_p}{\partial x_n} \end{pmatrix} = \left(\frac{\partial f_i}{\partial x_j}(\underline{x}) \right)_{i,j} \quad (23)$$

Rq: En math, on appelle cette matrice la matrice jacobien (ou simplement le jacobien), notations équivalentes:

$$[\underline{\nabla} \underline{f}(\underline{x})], \quad J_{\underline{f}}(\underline{x}), \quad D \underline{f}(\underline{x})$$

exemple: dans $\mathbb{R}^2$:

$$\underline{f}(\underline{x}) = \begin{cases} f_1(x, y) = x^2 + y^2 - 1 \\ f_2(x, y) = 2x + y - 1 \end{cases}$$

$$\underline{\nabla} \underline{f}(\underline{x}) = \begin{pmatrix} \frac{\partial f_1}{\partial x} & \frac{\partial f_1}{\partial y} \\ \frac{\partial f_2}{\partial x} & \frac{\partial f_2}{\partial y} \end{pmatrix} = \begin{pmatrix} 2x & 2y \\ 2 & 1 \end{pmatrix}$$

V-3 Généralisation des algos. aux systèmes multi-D

a) Positionnement du problème

On cherche le(s) zéro(s) $\underline{x}^*$ du système:

$$\underline{f}(\underline{x}) = 0 \quad \text{avec } \underline{f}: \mathbb{R}^n \rightarrow \mathbb{R}^n.$$

b) Itérations du point fixe

Comme en 1D, l'algo. du point fixe s'écrit:

$$\begin{cases} \bullet \underline{x}^0 \in \mathbb{R}^n \text{ donné} \\ \bullet \underline{x}^{k+1} = \underline{\Phi}(\underline{x}^k) \end{cases}$$

Il faut donc s'assurer de sa consistance:

$$\underline{x}^* = \underline{\Phi}(\underline{x}^*) \Leftrightarrow \underline{f}(\underline{x}^*) = 0 \quad (25)$$

et aussi sa stabilité, il faut que:

$$\rho(\underline{\nabla} \underline{\Phi}(\underline{x}^*)) \leq 1 \quad (\text{car: } |\phi'(x^*)| \leq 1) \quad (26)$$

Demo: (simplifiée)

$$\begin{aligned} \underline{x}^{k+1} - \underline{x}^* &= \underline{\Phi}(\underline{x}^k) - \underline{\Phi}(\underline{x}^*) \\ &\approx [\underline{\nabla} \underline{\Phi}(\underline{x}^*)](\underline{x}^k - \underline{x}^*) \quad \text{par D.L.} \end{aligned}$$

ce qui implique:

$$\|\underline{x}^{k+1} - \underline{x}^*\| \leq \underbrace{\|\underline{\nabla} \underline{\Phi}(\underline{x}^*)\|}_{\geq \rho(\underline{\nabla} \underline{\Phi}(\underline{x}^*))} \|\underline{x}^k - \underline{x}^*\|$$

dnc c.v. linéaire si $\rho(\underline{\nabla} \underline{\Phi}(\underline{x}^*)) \leq 1$.

anch: $\exists$ point fixe + stab. $\Rightarrow$ point fixe c.v. linéairement (pour $\underline{x}_0$ suffisamment proche de $\underline{x}^*$ (i.e. $\|\underline{x}_0 - \underline{x}^*\|$ pet

c) Algorithme de Newton

Pour établir cet algo., on reprend le principe vu en 1D:

Par un D.L. de Taylor en $\underline{x}_k$, on établit l'approx. linéaire

$$\underline{f}(\underline{x}_{k+1}) \approx L(\underline{x}) = \underline{f}(\underline{x}_k) + [\underline{\nabla} \underline{f}(\underline{x}_k)](\underline{x} - \underline{x}_k)$$

On définit $\underline{x}_{k+1}$ une solution du système linéaire:

$$L(\underline{x}_{k+1}) = 0$$

soit:

$$[\underline{\nabla} \underline{f}(\underline{x}_k)](\underline{x}_{k+1} - \underline{x}_k) = -\underline{f}(\underline{x}_k) \quad (27)$$

Ainsi, pour calculer $\underline{x}_{k+1}$, il faut résoudre un système linéaire. L'algo. de Newton s'écrit de façon pratique:

- $\underline{x}_0 \in \mathbb{R}^n$ donné, t.q. $[\underline{\nabla} \underline{f}(\underline{x}_0)]$ est inversible.
- Itérations, $k \geq 0$:

- Calculer la matrice $\underline{A}_k = \underline{\nabla} \underline{f}(\underline{x}_k)$

- Inverser le syst. linéaire $\underline{A}_k \delta \underline{x} = -\underline{f}(\underline{x}_k)$

- Effectuer: $\underline{x}_{k+1} = \underline{x}_k + \delta \underline{x}$.

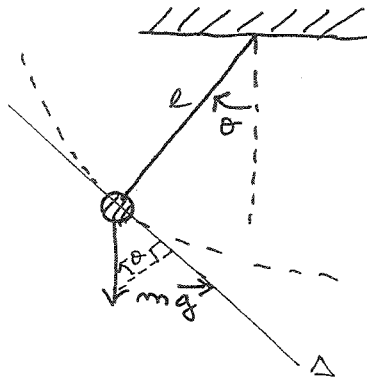
Thm (Admis) Si $[\underline{\nabla} \underline{f}(\underline{x}_0)]$ est inversible dans, pour $\underline{x}_0$ suffisamment proche de $\underline{x}^*$, la meth. de Newton admet une c.v. quadratique.

Equations Différentielles Ordinaires

EDO

1) Motivation : Le pendule oscillant

Mouvement d'un pendule de masse m attaché à une tige de longueur l :



$\theta(t)$: angle du pendule à l'instant t .

Le mouvement est régi par le principe fondamental de la dynamique :

$$m \vec{a} = m \vec{a}$$

On projette cette eq^r sur la tangente (Δ) à la trajectoire :

$$m g \sin \theta = m \underbrace{\ddot{\theta}}_{\text{accél. tangentielle}} l$$

L'eq^r du mv^t est donc l'Eq^r diff. ordinaire (EDO) :

$$\ddot{\theta} = - \omega_0^2 \sin \theta, \text{ avec } \omega_0 = \sqrt{\frac{g}{l}} \quad (1)$$

Pour définir complètement le mv^t, on lui adjoint les conditions initiales à $t=0$:

$$\begin{cases} \theta(0) = \theta_0 : \text{position initiale de la masse} \\ \dot{\theta}(0) = \dot{\theta}_0 : \text{vitesse initiale} \end{cases} \quad (2)$$

Le système (1)-(2) est un problème de Cauchy pour les EDO. Malgré son apparence simple, il est fortement non-linéaire et il n'existe pas de solution analytique à ce problème. Ce chapitre de cours est dédié à l'approximation numérique de ce type de problèmes.

Remarque : Le système (1)-(2) est une EDO du 2nd ordre. Il est possible de le transformer en système d'EDO du premier ordre :

$$(3a) \quad \dot{\underline{y}}(t) = \underline{F}(\underline{y}, t), \quad \underline{y} = \begin{pmatrix} y_1 \\ y_2 \end{pmatrix}, \quad \underline{F} : \mathbb{R}^2 \rightarrow \mathbb{R}^2$$

$$(3b) \quad \underline{y}(0) = \underline{y}_0$$

en posant $y_1 = \theta$ et en définissant la variable auxiliaire $y_2 = \dot{\theta}$. De cette façon, l'EDO (1) devient :

$$\begin{cases} \dot{y}_1 = y_2 \\ \dot{y}_2 = - \omega_0^2 \sin y_1 \end{cases}$$

qui est de la forme (3a) avec :

$$\underline{F}(\underline{y}) = \begin{cases} f_1(y_1, y_2) = y_2 \\ f_2(y_1, y_2) = - \omega_0^2 \sin y_1 \end{cases}$$

De même, la cond. initiale (2) s'écrit :

$$\underline{y}(0) = \begin{pmatrix} y_1(0) \\ y_2(0) \end{pmatrix} = \begin{pmatrix} \theta_0 \\ \dot{\theta}_0 \end{pmatrix}$$

II) Généralités

(Notation: $y(t) = \frac{dy}{dt}(t)$)

On s'intéresse aux équations diff. ordinaires (E.O.) du 1^{er} ordre, de la forme:

$$\dot{y}(t) = f(t, y(t)), \quad (1)$$

où $f: \mathbb{R}^2 \rightarrow \mathbb{R}$ est une fonction donnée.

~~Dep. Chap. II, on va se consacrer à la résolution des E.D.O. d'ordre plus élevé en systèmes d'E.D.O. en premier lieu. Exemples de problèmes.~~

Exemple 1: On considère $f(y, t) = 3y - 3t$, l'E.D.O.

s'écrit donc:

$$\dot{y}(t) = 3y(t) - 3t \quad (2)$$

dont les solutions ont de la forme $y_d = (d - \frac{1}{3})e^{3t} + t + \frac{1}{3}$.
Cette E.D.O. est linéaire, car on peut superposer 2 solutions y_{d1} et y_{d2} .

Exemple 2 $f(y, t) = 1 - y^2$, donc

$$\dot{y}(t) = 1 - y^2 \quad (3)$$

C'est un exemple d'E.D.O. non-linéaire. De plus cette E.D.O. est dite autonome ($f(y, t)$ ne dépend que de y)

Pour déterminer complètement la solution, on lui adjoint une condition initiale, p.ex:

$$y(t=0) = y_0. \quad (4)$$

Le système (1), (4) est appelé 'problème de Cauchy', soit:

$$(5) \quad \begin{cases} \dot{y}(t) = f(t, y(t)) & \text{pour } t \in [0, T] \\ y(0) = y_0. \end{cases}$$

EDP.
15/1

III) Schémas numériques de base

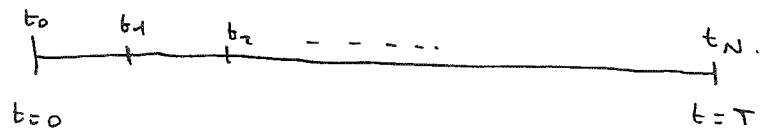
1) Construction

Il existe de nombreuses méthodes de construction. On va donc ici des méthodes basées sur des formules de quadrature numérique.

On effectue préalablement une subdivision du domaine de calcul $D = [0, T]$, de la forme

$$t_n = nh, \quad n = 0, \dots, N,$$

avec $h = \frac{T}{N}$ le pas de temps.



On intègre l'E.D.O (1) entre t_n et t_{n+1} :

$$\int_{t_n}^{t_{n+1}} \dot{y} dt = \int_{t_n}^{t_{n+1}} f(y(t), t) dt,$$

$$\text{soit } y^{n+1} - y^n = \int_{t_n}^{t_{n+1}} f(y(t), t) dt. \quad (6)$$

(notation: $y^n = y(t_n)$).

Il faut maintenant évaluer le RHS de (6) par quadrature numérique.

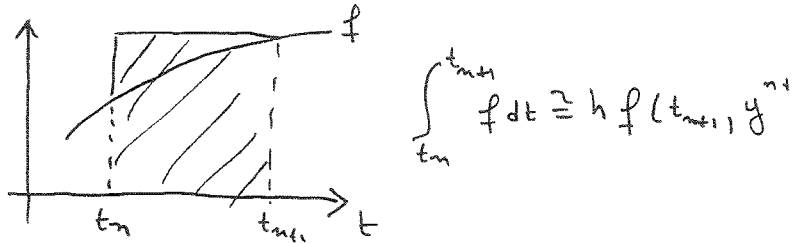
La méthode la plus simple consiste à approximer l'intégrale en utilisant la valeur au point t_n :



$$\int_{t_n}^{t_{n+1}} f dt \approx h f(t_n, y^n)$$

$$\frac{y^{n+1} - y^n}{h} = f(t_n, y^n). \quad (7)$$

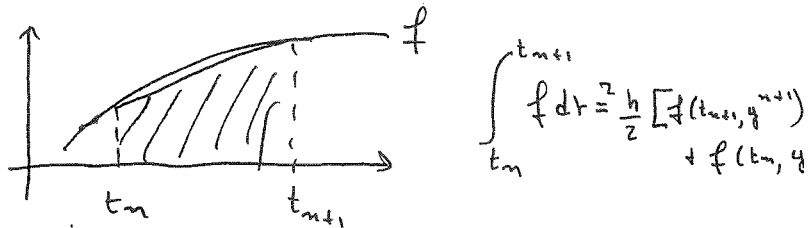
On peut aussi utiliser les valeurs au point final t_{n+1} :



On obtient alors la méthode d'Euler rétrograde EN1:

$$\frac{y^{n+1} - y^n}{h} = f(t_{n+1}, y^{n+1}) \quad (8)$$

On peut aussi utiliser la méthode du trapèze:



On obtient alors la méthode de Crank-Nicholson:

$$\frac{y^{n+1} - y^n}{h} = \frac{1}{2} [f(t_{n+1}, y^{n+1}) + f(t_n, y^n)]$$

Des schémas plus sophistiqués peuvent être construits avec les quadrature numériques avec des ordres supérieurs. On les verra en temps voulu en TD.

A chacun de ces schémas, on lui adjoint la C.I. (4):

$$y^0 = y_0 \quad (10)$$

III-2) Utilisation

Concrètement, tous ces schémas sont utilisés par le principe d'avancement en temps à partir de la C.I.:

$$y^0 = y(0).$$

Par exemple, l'application du schéma EP1 pour $n=1$ donne:

$$y^1 = y_0 + h f(t_0, y_0)$$

qui est une approximation de $y(t_1)$, puis:

$$y^2 = y_1 + h f(t_1, y_1)$$

etc... jusqu'à $t_N = T$. Le schéma EP1 est qualifié d'explicite car il n'y a pas d'éq^e à résoudre pour calculer les y^{n+1} .

Examinons l'utilisation du schéma EN1. Au temps $t=t_1$ on a:

$$y^1 = y_0 + h f(t_1, y^1),$$

et il n'y a pas de formule de récurrence simple qui donne facilement y^1 lorsque f n'est pas linéaire. Par exemple pour l'ex. 2 (Eq. (3)), on en obtient:

$$y^1 + h(y^1)^2 = y_0 + h.$$

On est amené à résoudre une eq^e non linéaire pour calculer y^1 . Pour cette raison, le schéma EN1 est qualifié d'implicite.

Rq: Dans le cas où f est linéaire, i.e. l'exemple (1), on a $f = 3y - 3t$ et le schéma EN1 donne:

$$y^1 = y_0 + h(3y^1 - 3t_1)$$

$$\text{soit } (1-3h)y^1 = y_0 - 3t_1.$$

TD ...

III Analyse des schémas

Le schéma explicite EP1 semble préférable à ER1 car il nécessite moins d'efforts pour avancer la solution en temps.

Néanmoins, il y a un prix à payer à cette facilité de calcul : la stabilité du schéma.

III-1 Stabilité

défi : c'est la propriété qui permet d'obtenir une solution numérique dont l'erreur ne s'amplifie pas lorsque $t \rightarrow +\infty$.

Considérons par exemple le problème modèle, ~~ou pour~~ $\lambda < 0$:

$$\begin{cases} \dot{y} = \lambda y, & t \in [0, +\infty[. \\ y(0) = y_0. \end{cases} \quad (11)$$

La solution exacte est

$$y(t) = y_0 e^{\lambda t},$$

qui a la ppte de s'amortir au cours du temps :

$$\lim_{t \rightarrow \infty} y(t) = 0$$

puisque $\lambda \leq 0$.

Lorsqu'on applique EP1 à (13), on obtient :

$$\begin{aligned} y^n &= (1 + \lambda h) y^{n-1} \\ &= (1 + \lambda h)^n y_0 \\ &= \sigma^n y_0 \end{aligned}$$

avec $\sigma = 1 + \lambda h$. Le réel σ est appelé facteur d'ampli

Pour que la solution numérique y^n reste bornée lorsque $n \rightarrow \infty$, on doit avoir nécessairement :

$$|\sigma| \leq 1,$$

ce qui conduit à

$$-1 \leq 1 + \lambda h \leq 1$$

$$0 < h < -\frac{2}{\lambda} = \frac{2}{|\lambda|}.$$

Pour obtenir une solution bornée, le schéma EP1 doit être utilisé avec un pas de temps h vérifiant la condition

$$h \leq \frac{2}{|\lambda|}. \quad (12)$$

On dit alors que le schéma est conditionnellement stable.

Que se passe-t-il pour le schéma ER1 ? Lorsque'il est appliqué au problème test (13), on obtient :

$$y^{n+1} = y^n + \lambda h y^{n+1}$$

soit

$$y^n = \sigma^n y_0 \quad \text{avec} \quad \sigma = \frac{1}{1 - \lambda h}.$$

La condition de stabilité s'écrit :

$$|\sigma| \leq 1 \Leftrightarrow |1 - \lambda h| \geq 1, \quad (13)$$

qui est toujours vérifiée pour tout $h > 0$. Le schéma ER1 est stable sans condition.

III-2 - Stabilité dans le cas général

Dans le cas d'un rhs quelconque $f(t, y)$, une analyse de stabilité analogue peut être réalisée, mais quelques restrictions.

Rappel: Formule de Taylor pour fonctions à 2 variables (forme symbolique)

$$\begin{aligned} f(x+h, y+k) &= f \sum_{n \geq 0} \frac{1}{n!} \left(h \frac{\partial}{\partial x} + k \frac{\partial}{\partial y} \right)^n f(x, y) \\ &= f(x, y) + \left(h \frac{\partial}{\partial x} + k \frac{\partial}{\partial y} \right) f(x, y) + \frac{1}{2!} \left(h \frac{\partial}{\partial x} + k \frac{\partial}{\partial y} \right)^2 f(x, y) \\ &\quad + \dots \end{aligned}$$

Première restriction: l'analyse de stabilité s'effectue pour un E.D.O linéaire.

Il faut donc linéariser $f(y, t)$ autour du point (y_n, t_n) .
En ne gardant que les termes linéaire du D.L. de f autour de (y_n, t_n) , on obtient :

$$\begin{aligned} f(t, y) &\approx f(y_n, t_n) + (t - t_n) \frac{\partial f}{\partial t}(y_n, t_n) + (y - y_n) \frac{\partial f}{\partial y}(y_n, t_n) \\ &= \lambda y + \alpha + \beta t \end{aligned}$$

Avec :

$$\lambda = \frac{\partial f}{\partial y}(y_n, t_n)$$

$$\alpha = f(y_n, t_n) - t_n \frac{\partial f}{\partial t}(y_n, t_n) - y_n \frac{\partial f}{\partial y}(y_n, t_n)$$

$$\beta = \frac{\partial f}{\partial t}(y_n, t_n)$$

l'E.D.O linéarisé est donc :

$$- \dot{y} = \lambda y + \alpha + \beta t.$$

La solution s'écrit :

$$y(t) = C_0 e^{\lambda t} - \left(\frac{\alpha}{\lambda} + \frac{\beta}{\lambda^2} + \frac{\beta}{\lambda} t \right),$$

la partie la plus dangereuse de l'instabilité est la partie exponentielle, qui est solution du problème test :

$$\dot{y} = \lambda y.$$

2^{ème} restriction: l'analyse de stabilité concerne la partie exponentielle de la solution : On ne retient donc que la partie autonome de l'E.D.O linéarisée.

Donc, au temps (t_n, y_n) , la stabilité du schéma EP1 est :

$$h \leq \frac{2}{-\lambda}, \quad \lambda = \frac{\partial f}{\partial y}(y_n, t_n).$$

Lors d'une interrogation en temps dans $t \in [0, T]$, il faut donc que :

$$h \leq \frac{2}{-\frac{\partial f}{\partial y}(y_n, t_n)}, \quad \forall n = 0, \dots, N.$$

3- stabilité - Résumé

Considérons le pbme de Cauchy

$$\begin{cases} \dot{y} = f(t, y(t)) \\ y(0) = y_0 \end{cases} \quad (IS)$$

à résoudre par schéma numérique. On peut définir

3 classes de schémas ou 3 méthodes.

• Schémas stables sans conditions :

La sol. numérique reste bornée (ie ne croît pas vers $\pm \infty$) quelle que soit la valeur du pas de temps h .

Un schéma implicite (ex: EP1) fait généralement partie de cette classe. Le prix à payer est de résoudre une équation (qui peut être non-linéaire) pour calculer y^{n+1} .

• Schémas instables : La solution numérique explose quelle que soit la valeur de h . $\Rightarrow$ schéma instable.

• Schémas conditionnellement stables :

Avec certains choix du paramètre h , la solution reste bornée. Les schémas explicites (ex: EP1) appartiennent généralement à cette catégorie.

Pour déterminer dans quelle classe appartient un schéma, il faut réaliser une analyse de stabilité sur le problème modèle :

$$\dot{y} = \lambda y,$$

qui est + simple que l'EDO originale (15), tout en gardant ses caractéristiques principales (p.ex: en effectuant une linéarisation de (15)).

Il s'agit d'analyser la stabilité sur le problème modèle sous une condition nécessaire de stabilité, qui élimine les instabilités de forme exponentielle (les plus dangereuses).

Précision

En plus Outre la propriété de stabilité du schéma, on cherche à connaître sa propriété de précision (ou convergence), c.à.d savoir si à un instant $t_n = nh$ la solution numérique y^n vérifie :

$$|y^n - y(t_n)| \rightarrow 0 \text{ lorsque } h \rightarrow 0.$$

On veut aussi déterminer l'ordre de précision q du schéma, telle que

$$|y^n - y(t_n)| = O(h^q).$$

Il y a essentiellement 2 méthodes basées pour déterminer l'ordre de précision q :

* méthode basée sur les D.L. de Taylor.

* ——— sur le problème modèle (11).

On va les expliciter pour le schéma EP1.

a) Précision basée sur le problème modèle

La solution exacte du problème (11) est :

$$y(t_n) = e^{\lambda t_n} y_0 = e^{\lambda nh} y_0 = (e^{\lambda h})^n y_0.$$

Le D.L. de $e^{\lambda h}$ est :

$$e^{\lambda h} = 1 + \lambda h + \frac{(\lambda h)^2}{2} + \frac{(\lambda h)^3}{3!} + \dots \quad (16)$$

Par contre, la solution numérique donnée par EP1 est :

$$y^n = \sigma^n y^0,$$

avec le facteur d'amplification :

$$\sigma = 1 + \lambda h.$$

On observe que σ représente le D.L. du facteur d'amplification exact $e^{\lambda h}$ jusqu'à l'ordre 1 inclus.

à l'ordre q son facteur d'amplification représente le D.L. de $e^{\lambda h}$ jusqu'à l'ordre $\frac{(\lambda h)^q}{q!}$ inclus. $\times$

On a donc la propriété fondamentale:

Théorème: (Admis)

Consistance à l'ordre q + stabilité $\Rightarrow$

$$\Rightarrow \begin{cases} \text{convergence à l'ordre } q: \\ |y^n - y(t_n)| = O(h^q). \end{cases}$$

B) Consistance basée sur les D.L. de Taylor.

Le schéma EP1 s'écrit:

$$\frac{y^{n+1} - y^n}{h} = f(y^n, t_n). \quad (17)$$

On effectue le D.L. de y^{n+1} au temps t_n :

$$y^{n+1} = y^n + h \dot{y}^n + \frac{h^2}{2} \ddot{y}^n + O(h^3)$$

que l'on reporte dans le schéma (17), on obtient:

$$\underbrace{\ddot{y}^n = f(y^n, t_n)}_{\text{E.D.O au temps } t^n} + \underbrace{\frac{h}{2} \ddot{y}^n}_{\substack{\hookrightarrow \tau_n: \text{reste ou erreur de} \\ \text{truncation}}} + O(h)$$

Pour ce schéma, l'erreur de truncation est:

$$\tau_h = \frac{h}{2} \ddot{y}^n = O(h)$$

Le schéma EP1 est donc consistant à l'ordre 1.

$\times$ Pour le schéma EP1, en résumé:

- il est stable pour $h < 2/|\lambda|$
- Il est consistant à l'ordre 1.

Il est donc convergent lorsqu'on fixe $h < 2/|\lambda|$.

$\times$ Pour le schéma ER1

$$\frac{y^{n+1} - y^n}{h} = f(y^{n+1}, t_{n+1}).$$

son facteur d'amplification est

$$\sigma = \frac{1}{1 - \lambda h} = 1 + \lambda h + (\lambda h)^2 + \dots$$

Il représente donc le D.L. de $e^{\lambda h}$ jusqu'à l'ordre 1:

il est consistant à l'ordre 1.

$\Rightarrow$ Il est convergent $\forall h > 0$.

IV-1 Motivations

Dans la pratique, on utilise des schémas plus sophistiqués que ER1 et EP1, qui ont l'avantage d'être plus précis (on veut au moins $q \geq 2$) et posséder une stabilité équivalente ou parfois supérieure.

Nous allons ici présenter les méthodes les plus populaires: RK.

Considérons la forme intégrale (6) de l'EDO:

$$y^{n+1} - y^n = \int_{t_n}^{t_{n+1}} f(t, y(t)) dt.$$

Pour augmenter la précision, une stratégie est de considérer des formules de quadrature + précises, p. ex. la formule du trapèze donne:

$$y^{n+1} - y^n = \frac{h}{2} [f(t_n, y^n) + f(t_{n+1}, y^{n+1})] \quad (18).$$

C'est le schéma C-N, en $O(h^2)$, stable sans condition (cf. TD).

Malheureusement, il est implicite.

Essayons-le puis rendre explicite: On peut essayer de "prédire" y^{n+1} par y^* en appliquant ER1:

$$y^* = y^n + h f(t_n, y^n) \quad (19)$$

~~et de corriger~~ La prédiction $y^* \approx y(t_{n+1})$ est au 1^{er} ordre.

On peut maintenant essayer de corriger en la reportant dans C-N, soit:

$$y^{n+1} = y^n + \frac{h}{2} [f(t_n, y^n) + f(t_{n+1}, y^*)] \quad (20).$$

Le schéma ^{explicite} à 2 étapes (19)-(20) est connu sous le nom de schéma "prédicteur-correcteur" (ou sch. de Heun), et on montrera qu'il est en $O(h^2)$.

EDO (14)

Le schéma (19)-(20) peut s'écrire de façon équivalente sous la forme:

$$k_1 = h f(t_n, y^n) \quad (21)$$

$$k_2 = h f(t_n + h, y^n + k_1) \quad (22)$$

$$y_{n+1} = y_n + \frac{1}{2} [k_1 + k_2]. \quad (23)$$

Si on le compare au schéma EP1, qui s'écrit:

$$y_{n+1} = y_n + k_1$$

le schéma prédicteur-correcteur augmente la précision de EP1 en évaluant f à un point supplémentaire dans $[t_n, t_{n+1}]$.

Cette construction est typique des méthodes de Runge-Kutta, où de judicieuses évaluations de f en q points de $[t_n, t_{n+1}]$ donne un schéma d'ordre q .

IV-2 Les schémas R-K d'ordre 2

La forme générale du schéma R-K ^{explicite} à 2 pas est:

$$k_1 = h f(y_n, t_n) \quad (24)$$

$$k_2 = h f(t_n + \alpha h, y_n + \beta k_1) \quad (25)$$

$$y^{n+1} = y^n + \gamma_1 k_1 + \gamma_2 k_2, \quad (26)$$

où $\gamma_1, \gamma_2, \alpha$ et β sont des coefficients à déterminer. (on voit bien que le schéma pred/cor. (21)-(23) est un cas particulier). Les coefficients sont déterminés de façon que le schéma soit le + précis possible.

Composons le D.L. de ce schéma à celui de la solution exacte au temps $t = t_n$.

La solution exacte donne:

$$y_{n+1} = y_n + h \dot{y}_n + \frac{h^2}{2} \ddot{y}_n + \dots \quad (27)$$

De plus,

$$\dot{y} = f(y, t) \Rightarrow \dot{y}_n = f(y_n, t_n) = f_n \text{ (notation)}$$

et:

$$\ddot{y} = \frac{d}{dt} \dot{y} = \frac{d}{dt} f(t, y(t)) = \frac{\partial f}{\partial t} + f \frac{\partial f}{\partial y} \Rightarrow \ddot{y}_n = \frac{\partial}{\partial t} f_n + f_n \frac{\partial}{\partial y} f_n$$

Le D.L. de la sol. exacte s'écrit donc:

$$y_{n+1} = y_n + h f_n + \frac{h^2}{2} \left(\frac{\partial}{\partial t} f_n + f_n \frac{\partial}{\partial y} f_n \right) + \dots \quad (28)$$

Effectuons maintenant le D.L. du schéma de R-K (24)-(26).

Le D.L. de k_2 en (t_n, y_n) s'écrit:

$$\begin{aligned} k_2 &= h \left[f_n + \alpha h \frac{\partial}{\partial t} f_n + \beta k_1 \frac{\partial}{\partial y} f_n \right] \\ &= h f_n + \alpha h^2 \frac{\partial}{\partial t} f_n + \beta h^2 f_n \frac{\partial}{\partial y} f_n, \end{aligned}$$

On en déduit que le D.L. du schéma (26) s'écrit:

$$y_{n+1} = y_n + h (\delta_1 f_n + \delta_2 f_n) + h^2 \delta_2 \left(\beta f_n \frac{\partial}{\partial y} f_n + \alpha \frac{\partial}{\partial t} f_n \right) + \dots \quad (29)$$

En comparant (28) et (29), une condition nécessaire ^{fm} que le schéma R-K soit en $O(h^2)$ est que:

$$\begin{cases} \delta_1 + \delta_2 = 1 \\ \delta_2 \beta = \frac{1}{2} \\ \delta_2 \alpha = \frac{1}{2} \end{cases}$$

Il y a 4 coeff. pour 3 eq's (non linéaires). On utilise, p.ex., α comme paramètre pour écrire:

$$\delta_2 = \frac{1}{2\alpha}, \quad \beta = \alpha, \quad \delta_1 = 1 - \frac{1}{2\alpha}.$$

Finalement, le schéma Rk2-2 (24)-(26) s'écrit:

$$k_1 = h f(y_n, t_n) \quad (29)$$

$$k_2 = h f(t_n + \alpha h, y_n + \alpha k_1) \quad (30)$$

$$y^{n+1} = y^n + \left(1 - \frac{1}{2\alpha}\right) k_1 + \frac{1}{2\alpha} k_2. \quad (31)$$

Dans la pratique, il y a 2 formes couramment utilisées:

* $\alpha = 1$ correspond au schéma de Heun (19)-

* $\alpha = \frac{1}{2}$ correspond au schéma "pred./correct":

$$y_{n+\frac{1}{2}}^* = y_n + \frac{h}{2} f(y_n, t_n) \quad (32)$$

$$y_{n+1} = y_n + h f(y_{n+\frac{1}{2}}^*, t_{n+\frac{1}{2}}) \quad (33)$$

Effectuons l'analyse du schéma Rk2-2 pour le problème modèle $\dot{y} = \lambda y$: on obtient

$$k_1 = \lambda h y^n$$

$$k_2 = \dots = \lambda h y^n + \alpha \lambda^2 h^2 y^n$$

$$\begin{aligned} y^{n+1} &= y^n + \left(1 - \frac{1}{2\alpha}\right) \lambda h y^n + \frac{1}{2\alpha} (\lambda h y^n + \alpha \lambda^2 h^2 y^n) \\ &= y^n \underbrace{\left(1 + \lambda h + \frac{\lambda^2 h^2}{2}\right)}_{\delta} \end{aligned}$$

Le facteur d'amplification représente les 2 premiers termes du D.L. de $e^{\lambda h}$, ce qui confirme la précision en $O(h^2)$ de Rk2-2

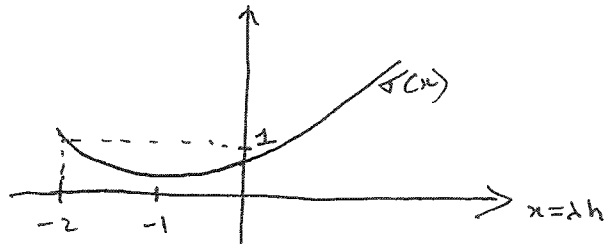
La stabilité requiert que

$$-1 \leq \gamma(\lambda h) \leq 1.$$

Pour déterminer cette condition, le plus simple est de tracer le graphe de

$$x \mapsto \gamma(x) = 1 + x + \frac{x^2}{2}$$

et d'examiner les valeurs de $x < 0$ pour lesquelles $|\gamma(x)| \leq 1$.



La stabilité requiert:

$$\begin{aligned} -2 &\leq x \leq 0 \\ -2 &\leq \lambda h \leq 0 \\ \Rightarrow \boxed{h \leq -\frac{2}{\lambda}} \end{aligned}$$

c'est la même condition de stabilité que EP1.

Concl. Les schémas RK2-k sont:

- * explicites
- * précis en $O(h^2)$
- * Aussi stable que EP1.

Rq: Par un procédé similaire (augmenter le nombre d'étapes $k=1, k=2, k=3$, etc...) il est possible de construire des schémas RK-k explicites de + en + précis. Le grand intérêt de ces schémas RK-k est que leur stabilité s'améliore lorsque la précision $\rightarrow$. En TD on étudiera le schéma RK4

5) Systèmes d'E.D.O.

On peut appliquer les schémas numériques scalaires des sections précédentes aux systèmes d'E.D.O, de la forme:

$$(34) \quad \dot{\underline{X}}(t) = \underline{N}(t, \underline{X}(t)), \quad \underline{X}(0) = \underline{X}_0$$

avec $\underline{X}(t) = \begin{pmatrix} x_1(t) \\ \vdots \\ x_n(t) \end{pmatrix} \in \mathbb{R}^n$.

Pour la clarté de ce chapitre, on va considérer un système d'E.D.O linéaire à coefficients constants, tel que $\underline{N}$ a la forme:

$$(35) \quad \underline{N}(t, \underline{X}) = \underline{A} \underline{X} + \underline{F}(t),$$

où $\underline{A}$ est une matrice de $\mathbb{R}^{n,n}$ à coefficients constants.

Le cas où $\underline{N}$ a une forme générale (i.e, non linéaire) sera effectué dans l'exo 5 du TD 3.

L'application des schémas numériques de la partie II et IV est évidente, par exemple le schéma EP1 s'écrit:

$$\frac{1}{h}(\underline{X}^{n+1} - \underline{X}^n) = \underline{A} \underline{X}^n + \underline{F}(t_n), \quad (36a)$$

qui s'écrit sous forme itérative:

$$\underline{X}^{n+1} = (\underline{I} + h \underline{A}) \underline{X}^n + h \underline{F}(t_n) \quad (36b)$$

L'Analyse de ce schéma reprend les grandes lignes de ce qui a été fait au chap. III pour les E.D.O scalaires. En particulier, l'analyse de stabilité pour le système (34)-(35)

$$\dot{X} = AX \quad (37)$$

Afin de pouvoir appliquer les résultats de la Sect. III-2, il faut diagonaliser la matrice $\underline{A}$, c-a-d :

$$\underline{A} = \underline{S} \underline{D} \underline{S}^{-1} \quad (38)$$

où $\underline{D}$ est la matrice diagonale des valeurs propres de $\underline{A}$, et $\underline{S}$ la matrice de passage (ie, la $i^{\text{ème}}$ colonne de $\underline{S}$ contient le vecteur propre associé à la $i^{\text{ème}}$ val. propre de $\underline{A}$).

En effectuant le changement de base $\underline{Y} = \underline{S}^{-1} \underline{X}$, le syst. (3) s'écrit :

$$\underline{S}^{-1} \dot{\underline{X}} = \underline{S}^{-1} \underline{A} \underline{S} \underline{S}^{-1} \underline{X} + \underline{S}^{-1} \underline{F},$$

Pour la variable $\underline{Y}$, ce système prend la forme diagonale

$$\dot{\underline{Y}} = \underline{D} \underline{Y} + \underline{G}, \quad \underline{G} = \underline{S}^{-1} \underline{F}$$

qui s'écrit composante par composante :

$$\text{Pour } i=1, \dots, n: \quad \dot{y}_i = \lambda_i y_i + g_i(t), \quad (39)$$

où λ_i est une v.p. de $\underline{A}$. L'avantage de ce changement de variable est d'un lien au système d'EDO totalement découplé (39).

On peut appliquer le schéma EP1 à cette équation, i

$$y_i^{n+1} = (1 + \lambda_i h) y_i^n + h g_i(t). \quad (40)$$

D'après l'analyse de stabilité de la section III-2, ce schéma est stable si :

$$h \leq \frac{2}{-\lambda_i}, \quad i=1, \dots, n. \quad (41)$$

Ainsi, pour assurer la stabilité du schéma EP1 (36) il faut que toutes les conditions de stabilité (41) soient vérifiées, i.e :

$$h \leq \frac{2}{\max_i \{-\lambda_i\}}$$

Résumé: Pour effectuer l'analyse de stabilité d'un schéma numérique au système d'EDO :

$$\dot{X} = AX + F(t),$$

il faut :

- 1) Calculer les v.p. λ_i de $\underline{A}$
- 2) Effectuer l'étude de stabilité du schéma numérique appliqué à l'EDO scalaire :

$$\dot{y} = \lambda_i y, \quad i=1, \dots, n$$

- 3) La cond. de stabilité est alors donnée par la plus restrictive des cond. de stabilité pour chaque EDO scalaire.

Remarques

EDD 22

1) Même si $\underline{A}$ est à coefficients réels, ses v.p. λ_i ne sont pas nécessairement réelles (il faudrait en outre que $\underline{A}$ soit symétrique). On peut énoncer que :

- Si $\lambda_i \leq 0$, on utilise l'analyse de stab. du Chap III-1.
- Si $\lambda_i > 0$, on ne peut rien dire sur la stabilité du schéma (cf chap III-2).
- Si $\lambda_i \in \mathbb{C}$, avec une partie réelle négative, la "mode d'emploi" de l'analyse de stabilité est donné dans l'exo 5. du TD 3.

2) Le cas général $\dot{X} = N(t, X)$ où N est non-linéaire se ramène au cas (34)-(35) après linéarisation de $N(t, X)$ (d'une façon analogue à ce qui a été fait dans la sect. III-2 et l'exo 2 du TD 3 pour une EDO scalaire).

Un exemple est donné dans l'exo 5.

Appendix A: Methods for Ordinary Differential Equations (ODEs)

We consider the initial value problem:

$$(1a) \quad \dot{y} = f(t, y(t)), \quad t \in [0, T] = \mathcal{U}$$

$$(1b) \quad y(t=0) = y_0$$

often called a Cauchy problem.

I) Building of basic numerical schemes

the basic methods used in CFD (Euler scheme, Crank-Nicolson method) will be built here by means of numerical quadrature.

First, the time domain $\mathcal{U} = [0, T]$ is subdivided as

$$\text{as } \{t_n = n \Delta t, \quad n = 0, \dots, N\}$$

where $\Delta t = \frac{T}{N}$ is the time-step.

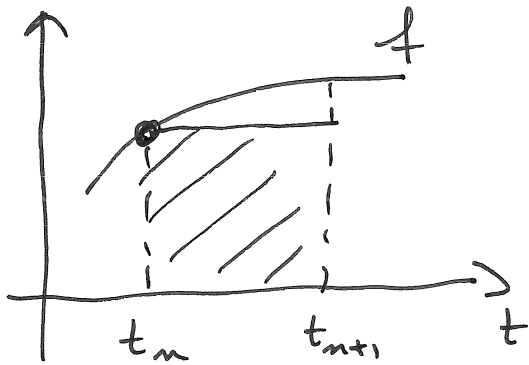
The EDO (1a) is first integrated in $[t_n, t_{n+1}]$ to yield:

$$y(t_{n+1}) - y(t_n) = \int_{t_n}^{t_{n+1}} f(t, y(t)) dt \quad (2)$$

~~Now~~ No discrete approximation has yet been performed. Now, the RHS of (2) will be approximated with numerical quadrature based on Lagrange polynomials.

I-1) Forward Euler method (FE1) (2)

The simplest method stems from the approximation of (2) by using values at time t_n :



$$\int_{t_n}^{t_{n+1}} f dt \approx \Delta t f(t_n, y^n)$$

where we used the compact notation:

$$\underbrace{y^n}_{\text{approx. solution.}} \approx \underbrace{y(t_n)}_{\text{exact solution}}$$

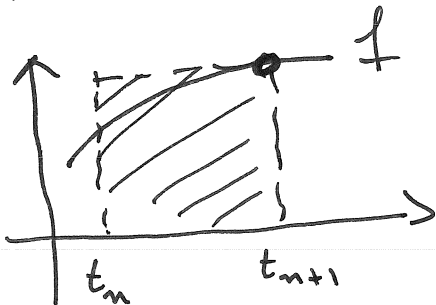
the forward Euler scheme reads then:

$$\frac{y^{n+1} - y^n}{\Delta t} = f(t_n, y^n) \quad (3)$$

Its acronym is FE1, because we shall see later that the method is first order in time.

I-2) Backward Euler Method (BE1)

If the RHS of (2) is approximated with end point t_{n+1} :



$$\int_{t_n}^{t_{n+1}} f dt \approx \Delta t f(t_{n+1}, y^{n+1})$$

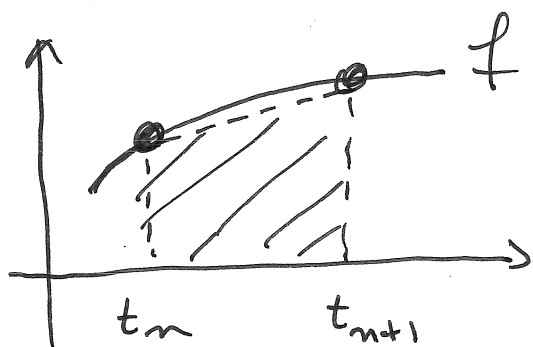
we obtain the backward Euler method (BE1):

$$\frac{y^{n+1} - y^n}{\Delta t} = f(t_{n+1}, y^{n+1}) \quad (4)$$

I-3] Trapezoidal method

(3)

The RHS of (2) can be approximated by the Trapezoidal method (use of Lagrange polynomial of order 2)



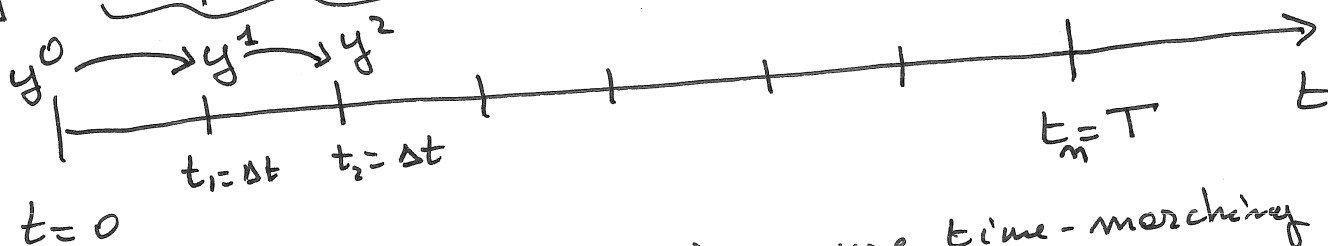
$$\int_{t_n}^{t_{n+1}} f dt \approx \frac{\Delta t}{2} [f(t_{n+1}, y^{n+1}) + f(t_n, y^n)]$$

The resulting trapezoidal method:

$$\frac{y^{n+1} - y^n}{\Delta t} = \frac{1}{2} [f(t_{n+1}, y^{n+1}) + f(t_n, y^n)] \quad (5)$$

is often called the Crank-Nicolson method (CN) when used for time integration of PDEs.

II] Use of time integration scheme. Implicit or explicit method



For ~~integrating~~ solving ODE (1), we use time-marching by starting with the initial condition at $t=0$:

$$y^0 = y(0),$$

Then, y^1 at $t=t_1$ is calculated with FE1, for example, written in its iterative form:

$$\underbrace{y^1}_{\text{Unknown at } t=t_1} = \underbrace{y^0 + \Delta t f(t_0, y^0)}_{\text{known at } t=t_1} \quad (6)$$

(4)

This expression is ready to be programmed in any language, such as FORTRAN or MATLAB. Since there is no equation to solve to obtain y^1 from values at previous time-level, the FE1 scheme is called explicit. Furthermore, Eq. (6) involves the solution at only time levels to and t_1 , thus FE1 is called a two-level scheme.

By contrast, the BE1 scheme leads to the iterative equation:

$$y^1 - \Delta t f(t_1, y^1) = y^0, \quad (7)$$

the solution of an equation is needed to obtain y^1 , resulting in more computational work. The BE1 is called an implicit scheme, and is also a two-level scheme.

The CN (trapezoidal) scheme is an implicit two-level scheme.

III) Truncation error of time-integrated scheme.

The local truncation error of a scheme (LTE) is the difference between the numerical scheme and the ODE at time t_n . It can be summarized as:

$$\text{Scheme at } t_n = \text{ODE at } t_n + \tau^n$$

where τ^n is the LTE at time t_n . It is usually performed by using Taylor expansion of the solution, i.e.

$$y^{n+1} = y^n + \Delta t \dot{y}^n + \frac{\Delta t^2}{2} \ddot{y}^n + O(\Delta t^3) \quad (8)$$

Let us calculate the LTE of backward-Euler scheme (5) by introducing this Taylor expansion in (3):

$$\underbrace{\frac{y^{n+1} - y^n}{\Delta t} - f(t_n, y_n)}_{\text{FE1 scheme at } t_n} = \frac{y^n + \Delta t \dot{y}^n + \frac{\Delta t^2}{2} \ddot{y}^n + O(\Delta t^3) - y^n}{\Delta t} - f(t_n, y_n)$$

$$= \underbrace{\dot{y}^n - f(t_n, y_n)}_{\text{ODE at } t_n} + \tau^n$$

where the LTE is $\tau^n = \frac{\Delta t}{2} \ddot{y}^n + O(\Delta t^2)$.

Hence, the FE1 scheme is first-order accurate.

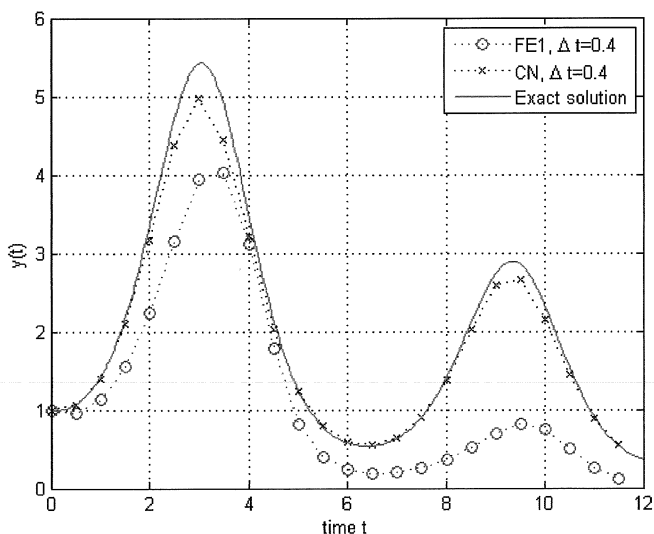
In a similar fashion, we can show the the LTE of

BE1 is:

$$\tau^n = -\frac{\Delta t}{2} \ddot{y}^n + O(\Delta t^2)$$

and the LTE of CN scheme is:

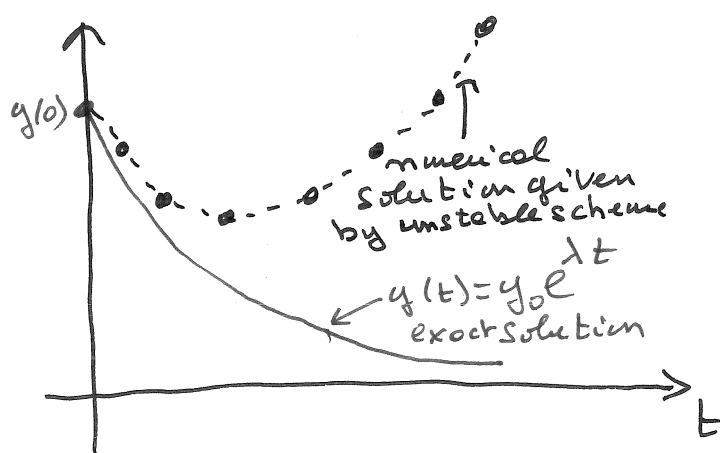
$$\tau^n = \frac{\Delta t^2}{12} \ddot{y}^n + O(\Delta t^3)$$



IV

Stability of a numerical scheme

(6)



The stability (also called absolute stability) is the property of the scheme to yield a numerical solution who does not diverge to ∞ when t is large.

Let us consider the model problem for $\lambda < 0$:

$$\begin{cases} \dot{y} = \lambda y, & t \in [0, +\infty[\\ y(0) = y_0 \end{cases} \quad (9)$$

Its exact solution is:

$$y(t) = y_0 e^{\lambda t},$$

and since $\lambda < 0$ the solution decays with time ($\lim_{t \rightarrow \infty} y(t) = 0$).

When BE1 is applied to model problem (9), the iterative equation reads:

$$y^{n+1} = y^n + \Delta t \lambda y^n \Rightarrow y^n = \sigma y^{n-1} = \dots = (\sigma)^n y^0$$

where $\sigma = 1 + \lambda \Delta t$ is the amplification factor of the scheme. In order to have a solution that does not blow up in time, we must have:

$$|\sigma| < 1 \quad (10)$$

$$\text{which leads to } 0 < \Delta t < -\frac{2}{\lambda}.$$

Hence, the BE1 scheme is stable if we use a time-step that verifies the stability condition:

$$\Delta t < \frac{2}{-\lambda} \quad (11)$$

We say that BE1 is conditionally stable.

- If we use BE1 to solve model problem (9), the amplification factor is $\sigma = \frac{1}{1 - \lambda \Delta t}$,

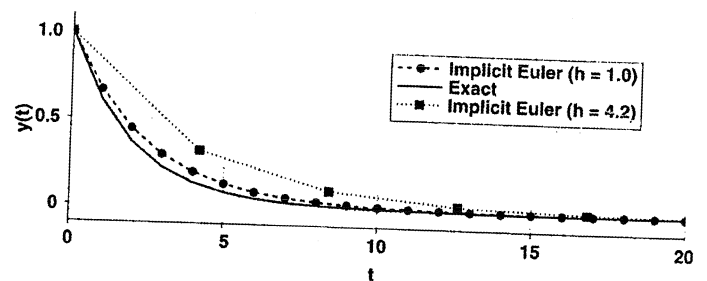
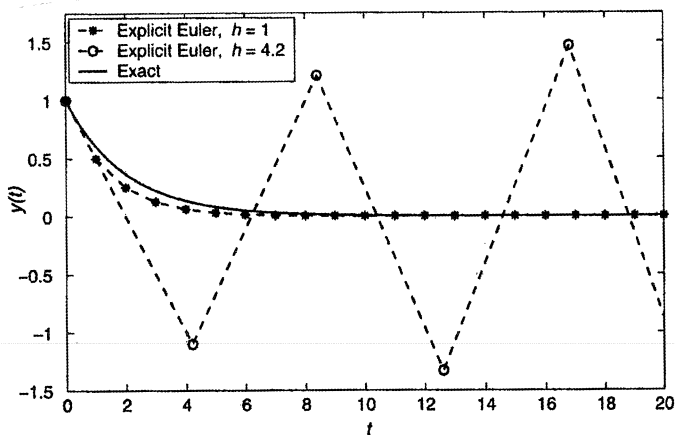
and stability condition $|\sigma| < 1$ leads to $|1 - \lambda \Delta t| > 1$, which is always verified for $\Delta t > 0$. Hence, the BE1 is stable with no conditions.

- The amplification factor for CN scheme is

$$\sigma = \frac{1 + \lambda \Delta t/2}{1 - \lambda \Delta t/2}$$

and stability condition (10) is verified for all $\Delta t > 0$. As the BE1 scheme, the CN scheme is always stable.

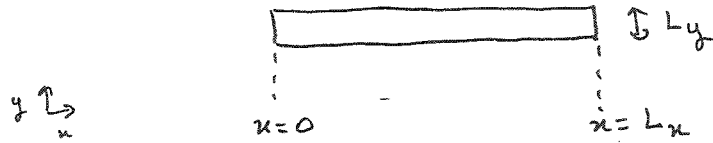
Application: We consider the model problem $\dot{y} + \frac{1}{2}y = 0$. the stability condition of BE1 is $\Delta t < 4$, while BE1 and CN are stable for all $\Delta t > 0$. Fig. 1 shows computations with FE with various values of Δt , the one with $\Delta t = 4.2$ is unstable. Fig. 2 shows that BE1 is stable for all values of Δt .



Résolution de Systèmes linéaires I : Généralités et méthodes directes

I) Motivation

On souhaite connaître l'équilibre thermique d'un barreau



tel que :

- Le barreau est long : $L_x \gg L_y$
- On impose une source de chaleur uniforme répartie dans le barreau : $\dot{\omega}(x, y)$ (en W/m^3)
- Aux extrémités $x=0$ et $x=L_x$ on impose la température nulle : $T=0$.

On cherche à connaître la distribution de température $T(x, y)$ dans ce barreau. Pour cela, on considère que dans un volume élémentaire dV du barreau, il existe :

- $\vec{q}(x, y)$: flux de chaleur au travers de la surface $P = dV$ (en W/m^2)
- $\dot{\omega}(x, y)$, la source de chaleur.

Pour que le volume de dV soit en équilibre thermique, il faut que le flux de chaleur au travers de P soit égal à la quantité de chaleur reçue :

$$\int_{dV} \vec{q} \cdot \vec{n} dV = \int_{dV} \dot{\omega} dV$$

mit, en utilisant le théorème de la divergence.

$$\int_V (\nabla \cdot \vec{q} - \dot{\omega}) dV = 0. \quad (1)$$

La loi de Fourier énonce que $\vec{q} = -k \nabla T$, où k est la conductivité thermique du métal, ce qui conduit à

$$\int_V [-k \nabla^2 T - \dot{\omega}] dV = 0,$$

en supposant que k est une constante. Puisque cette intégrale doit être nulle pour tout volume V du barreau, ainsi localement :

$$-k \nabla^2 T = \dot{\omega}. \quad (1)$$

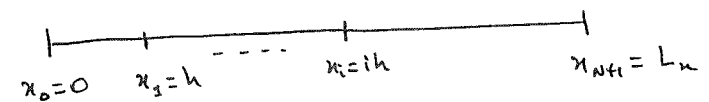
Cette équation aux dérivées partielles est connue sous le nom d'équation de Laplace.

On la simplifie en considérant que $L_x \gg L_y \Rightarrow T(x, y) = T(x)$ et on obtient le problème aux limites (avec $f = \dot{\omega}/k$)

$$(2a) \quad -\frac{d^2 T}{dx^2} = f(x) \quad \text{pour } x \in]0, L_x[$$

$$(2b) \quad T(0) = 0, \quad T(L_x) = 0.$$

Il est possible de résoudre cette équation par la méthode des différences finies. Pour cela, on introduit le maillage du barreau :



qui est tel que $x_i = ih$ pour $i = 0, \dots, N+1$ avec $h = \frac{L_x}{N+1}$

On cherche maintenant à calculer $T_i = T(x_i)$ en ces points du maillage.

conclusion $\bar{\sigma}$:

$$T_0 = 0, \quad T_{N+1} = 0.$$

Aux points $\{x_i, i=1, \dots, N\}$ situés à l'intérieur du maillage, l'Eq. (2a) s'écrit :

$$-T_i'' = f_i, \quad \text{pour } i=1, \dots, N. \quad (3)$$

Pour obtenir une approximation de T_i'' , on utilise la formule du TD1 :

$$T_i'' \approx \frac{T_{i-1} - 2T_i + T_{i+1}}{h^2}. \quad (4)$$

En introduisant cette approx. dans Eq. (3), on obtient le système linéaire de taille $N \times N$:

en $i=1$: $-\left(\overset{\substack{\text{valeur} \\ \text{d'après les c.l.}}}{T_0} - 2T_1 + T_2\right) = f_1 \quad (5a)$

pour $i=1, \dots, N-1$: $-\left(T_{i-1} - 2T_i + T_{i+1}\right) = f_i \quad (5b)$

pour $i=N$: $-\left(T_{N-1} - 2T_N + \overset{\substack{\text{valeur} \\ \text{d'après les c.l.}}}{T_{N+1}}\right) = f_N. \quad (5c)$

Sous forme matricielle, il s'écrit comme :

$$\underline{K}_n \underline{T} = \underline{F} \quad (6)$$

avec :

$$\underline{T} = \begin{pmatrix} T_1 \\ \vdots \\ T_N \end{pmatrix}, \quad \underline{F} = \begin{pmatrix} f_1 \\ \vdots \\ f_N \end{pmatrix},$$

$$\underline{K}_n = \begin{bmatrix} 2 & -1 & & & \\ -1 & 2 & -1 & & \\ & \ddots & \ddots & \ddots & \\ & & -1 & 2 & -1 \\ & & & -1 & 2 \end{bmatrix} \quad (7)$$

Conclusion : La discrétisation des équations ^{linéaires} de la mécanique donne lieu à des systèmes linéaires, que'il faut maintenant résoudre. Pour des systèmes de grande taille ($n \gg 1$), il faut utiliser des méthodes qui sont :

- * efficaces (i.e., rapides) ^{lesquelles sont utilisées sur un ordinateur}
- * peu sensibles aux erreurs d'arrondi
- * robustes : elle doivent fournir un résultat pour un grand nombre de problèmes différents.

II Généralités sur les systèmes linéaires

Un système linéaire d'ordre n s'écrit sous forme algébrique :

$$\text{pour } i=1, \dots, n \quad \sum_{j=1}^n A_{ij} x_j = b_i \quad (8)$$

(i : indice des lignes, j : indice des colonnes).

Il s'écrit aussi sous la forme matricielle :

$$\underline{A} \underline{x} = \underline{b}$$

avec $\underline{x}$ et $\underline{b}$ des vecteurs de taille n , et $\underline{A}$ une matrice de taille $n \times n$ telle que :

$$\underline{A} = (A_{ij})_{i,j=1, \dots, n} = \begin{bmatrix} A_{11} & A_{12} & \dots & A_{1n} \\ A_{21} & & & \\ \vdots & & & \end{bmatrix}$$

* La matrice $\underline{A}$ est triangulaire supérieure si elle est telle que $A_{ij} = 0$ pour $i > j$, soit pour $n=3$:

$$\underline{A} = \begin{bmatrix} A_{11} & A_{12} & A_{13} \\ 0 & A_{22} & A_{23} \\ 0 & 0 & A_{33} \end{bmatrix}$$

* Elle est dite triangulaire inférieure si elle est telle que $A_{ij} = 0$ pour $j > i$, soit p. ex. $n=3$:

$$\underline{A} = \begin{bmatrix} A_{11} & 0 & 0 \\ A_{21} & A_{22} & 0 \\ A_{31} & A_{32} & A_{33} \end{bmatrix}$$

* Elle est dite symétrique si $A_{ij} = A_{ji}$, par exemple:

$$\begin{bmatrix} 1 & 2 & 3 \\ 2 & 0 & 4 \\ 3 & 4 & 2 \end{bmatrix}$$

* Un type de matrice très courant sont les matrices tridiagonales, de la forme:

$$\underline{A} = \begin{bmatrix} b_1 & c_1 & 0 & 0 \\ a_2 & b_2 & c_2 & 0 \\ 0 & a_3 & b_3 & c_3 \\ 0 & 0 & a_4 & b_4 \end{bmatrix} \quad \text{pour } n=4.$$

Dans ce type de matrices, seuls les éléments diagonaux (les b_i), sur la 1^{ère} diagonale supérieure (les c_i) et inférieure (les a_i) peuvent être non nuls.

On rencontre souvent ce type de matrices dans la méthode des différences finies. En mécanique, un exemple classé est la ~~matrice~~ matrice K_n (Eq. (7)) pour laquelle, quel que soit n , la matrice a seulement 3 éléments non nuls par ligne. De plus, on remarque que K_n est symétrique.

III Méthodes de résolution directe.

On souhaite donc résoudre le système d'ordre n :

$$\underline{A} \underline{X} = \underline{B} \quad (C_0).$$

Dans ce cours, on considérera toujours des systèmes non singuliers ($\det \underline{A} \neq 0$), pour lesquels la solution unique de (C₀) s'écrit:

$$\underline{X} = \underline{A}^{-1} \underline{B}.$$

Pour des raisons liées aux erreurs d'arrondi (cf chap on ne doit jamaïs résoudre un syst. lin. sur ordinateur en calculant $\underline{A}^{-1}$ et en le multipliant avec $\underline{B}$).

On utilise plus généralement le procédé d'élimination de Gauss. C'est l'exemple le plus populaire de méthodes de résolution directe, qui

fournissent la solution en un nombre fini d'étape. (À l'opposé: les méthodes de résolution itératives, qui fournissent la solution en un nombre infini (théoriquement) d'étapes, seront vu dans le chap. suivant).

Avant d'aborder le cas général, considérons l'exemple du syst. triang. sup. d'ordre 3 défini par :

$$\underline{A} = \begin{bmatrix} 4 & 8 & 12 \\ 0 & 2 & 4 \\ 0 & 0 & 2 \end{bmatrix}, \quad \underline{b} = \begin{pmatrix} 4 \\ 2 \\ 4 \end{pmatrix}$$

qui s'écrit :

$$\begin{cases} 4x_1 + 8x_2 + 12x_3 = 4 \\ 2x_2 + 4x_3 = 2 \\ 2x_3 = 4 \end{cases} \quad (10)$$

On la résout facilement en observant que la dernière ligne donne $x_3 = 2$, celle d'en dessus donne alors :

$$x_2 = \frac{1}{2} [2 - 4x_3] = -3$$

et la première ligne donne enfin :

$$x_1 = \frac{1}{4} [4 - 8x_2 - 12x_3] = 1.$$

Cette méthode est l'algorithme de substitution rétrograde ("backward sweep"). Pour un système d'ordre n , cet algo s'écrit :

$$x_n = b_n / A_{nn}, \quad (11a)$$

et, pour $i = n-1$ jusqu'à 1 :

$$x_i = \frac{1}{A_{ii}} \left(b_i - \sum_{j=i+1}^n A_{ij} x_j \right) \quad (11b)$$

Cet algorithme se programme en 5 lignes sur un ordi., et il est au cœur de la méthode de substitution de Gauss pour les matrices A quelconques.

III - 2) Méthode d'élimination de Gauss (MEG)

La MEG a pour but de transformer le système $\underline{A}x = \underline{b}$ en un système équivalent (ie, ayant la même solution) et la forme

$$\underline{U}x = \underline{b}' \quad (12)$$

où U (ie: Upper) est une matrice triang. & supérieure. Le système (12) peut alors être aisément résolu par l'algo. de substitution rétrograde de la section précédente. Pour définir la MEG, prenons l'exemple du système d'ordre 3

$$\underline{A} = \begin{bmatrix} 4 & 8 & 12 \\ 3 & 8 & 13 \\ 2 & 9 & 18 \end{bmatrix}, \quad \underline{b} = \begin{pmatrix} 4 \\ 5 \\ 11 \end{pmatrix}$$

qui s'écrit :

$$\begin{cases} 4x_1 + 8x_2 + 12x_3 = 4 & r_1^{(1)} \\ \boxed{3}x_1 + 8x_2 + 13x_3 = 5 & r_2^{(1)} \\ \boxed{2}x_1 + \boxed{9}x_2 + 18x_3 = 11 & r_3^{(1)} \end{cases} \quad (13)$$

et on notera $r_i^{(h)}$ la i ème ligne du système obtenue à la h ème étape de la MEG (pour le syst. (13), on a $h=1$).

1ère partie de la MEG : mise sous forme triangulaire supérieure

Afin de rendre le syst. (13) triang. sup., il faut annuler les coefficients encadrés dans (13) en effectuant des opérations algébriques sur les lignes $r_i^{(1)}$ de ce syst.

En effectuant les opérations suivantes :

$$(14) \quad \begin{cases} r_1^{(2)} = r_1^{(1)} \\ r_2^{(2)} = r_2^{(1)} - m_{21} r_1^{(1)}, \text{ avec } m_{21} = 3/4 \end{cases}$$

Le système (13) se transforme en :

$$\begin{cases} -4x_1 + 8x_2 + 12x_3 = 4 & r_1^{(2)} \\ 2x_2 + 4x_3 = 2 & r_2^{(2)} \\ \boxed{5}x_2 + 12x_3 = 9 & r_3^{(2)} \end{cases} \quad (15)$$

Pour annuler le coefficient encadré dans (15), il faut faire les opérations

$$\begin{aligned} r_1^{(3)} &= r_1^{(2)} \\ r_2^{(3)} &= r_2^{(2)} \\ r_3^{(3)} &= r_3^{(2)} - m_{32} r_2^{(2)}, \quad m_{32} = 5/2. \end{aligned} \quad (16)$$

On obtient alors le système triangulaire supérieure

$$\begin{cases} 4x_1 + 8x_2 + 12x_3 = 4 \\ 2x_2 + 4x_3 = 2 \\ 2x_3 = 4 \end{cases} \quad (17)$$

2^{ème} Partie de la MEG : Résolution par substitution rétrograde

Le syst. triang. sup. peut être aisément résolu par l'algo de la section III-1. Dans notre exemple, on aura reconnu que (17) est l'exemple de la section précédente, dont la solution est :

$$x_1 = 1, \quad x_2 = -3, \quad x_3 = 2.$$

III.3 Factorisation LU

Il est parfois utile de réinterpréter la MEG comme une méthode de factorisation de la matrice A sous la forme

$$\underline{\underline{A}} = \underline{\underline{L}} \underline{\underline{U}}, \quad (18)$$

telle que U est la matrice triang. sup. donnée par la première partie de la MEG ; pour notre exemple :

$$\underline{\underline{U}} = \begin{bmatrix} 4 & 8 & 12 \\ 0 & 2 & 4 \\ 0 & 0 & 2 \end{bmatrix}$$

et L (Lower) est la matrice triang. inférieure de la forme :

$$\underline{\underline{L}} = \begin{bmatrix} 1 & 0 & 0 \\ m_{21} & 1 & 0 \\ m_{31} & m_{32} & 1 \end{bmatrix}$$

où les m_{ij} , $i > j$, sont les multiplicateurs utilisés par les opérations élémentaires (14) et (16) de la première partie de la MEG. Pour notre exemple :

$$\underline{\underline{L}} = \begin{bmatrix} 1 & 0 & 0 \\ 3/4 & 1 & 0 \\ 1/2 & 5/2 & 1 \end{bmatrix}.$$

Lorsque la factorisation LU de la matrice A est connue, il est aisé de résoudre le système $\underline{\underline{L}} \underline{\underline{U}} \underline{x} = \underline{b}$ par les 2 étapes successives :

puis -

$$\underline{U} \underline{x} = \underline{y} \quad (19b)$$

Puisque ces 2 systèmes sont triangulaires, on les résout aisément par un algorithme de substitution rétrograde (pour (19b)) et de substitution progressive (pour (19a)).

III-4 Complexité Algorithmique

Afin d'évaluer l'efficacité (i.e. le temps de calcul) d'un algorithme de résolution numérique, il est important de connaître sa complexité algorithmique, i.e. le nombre d'opérations arithmétiques scalaires (multiplications, additions, etc.) qu'il nécessite. Ces opérations arithmétiques sont souvent appelées flops (floating points scalar operations).

Par exemple, comptons le nombre de flops que nécessite le multiplicateur matrice-vecteur $\underline{A} \underline{x}$, qui s'écrit :

$$\text{pour } i = 1, \dots, n \quad \sum_{j=1}^n A_{ij} x_j = A_{i1} x_1 + A_{i2} x_2 + \dots + A_{in} x_n$$

Pour chaque ligne i , on fait donc n multiplications et $n-1$ additions, soit $2n-1$ flops. La complexité de l'algo est donc :

$$n(2n-1) = O(n^2) \text{ flops.}$$

Ainsi, lorsqu'on double la taille n de la matrice, le nombre d'opérations est multiplié par 4. Le coût de cet algorithme peut être réduit dans le cas ~~de~~ particulier des matrices tri-diagonales (p.ex. la matrice de diff. finies $K_n(\tau)$) l'opération (20) devient alors

$$\text{pour } i = 1, \dots, n \quad \sum_{j=1}^n A_{ij} x_j = a_i x_{i-1} + b_i x_i + c_i x_{i+1}$$

et elle nécessite seulement $5n = O(n)$ opérations. Le coût de cet algorithme varie seulement linéairement avec la taille n de la matrice.

Il est crucial de connaître la complexité d'un algorithme de résolution de syst. linéaire, car une méthode qui semble séduisante sur le plan théorique peut se révéler inutile en pratique. Prenons l'exemple des formules de Cramer qui donne la solution du système $A\underline{x} = \underline{b}$ sous la forme

$$\underline{x}_j = \frac{\Delta_j}{\det(A)}, \quad j = 1, \dots, n,$$

où Δ_j est le déterminant de la matrice obtenue en remplaçant la $j^{\text{ème}}$ colonne de $\underline{A}$ par le second membre $\underline{b}$.

La complexité algorithmique de cette formule est représentée par le calcul (par récurrence) des déterminants, et on peut montrer que'elle est de $(n+1)!$ flops, ce qui est gigantesque: pour résoudre un système linéaire d'ordre $n=50$, un P.C. actuel nécessiterait plusieurs milliards d'années de temps de calcul ("temps CPU") $\Rightarrow$ La méthode de Cramer n'est jamais utilisée en pratique.

Pour la MEG, on peut facilement montrer que:

- La mise sous forme triang. supérieure nécessite $O(n^3)$ flop
- La substitution rétrograde nécessite n^2 flops.

Soit un total de $O(n^3) + O(n^2)$ flops.

Ainsi, pour résoudre un problème de taille $n=50$, la MEG nécessite environ 10000 flops, ce qui prend 10^{-4} s sur un PC de 200 MHz: bien mieux que la meth. de Cramer!

Il faut aussi retenir que le coût dominant de la MEG est la mise sous forme triang. sup., soit $O(n^3)$ opérations.

Ainsi, lorsque la taille de la matrice est multipliée par 2, le coût algorithmique est multiplié par 8. Pour les systèmes de grandes tailles, le temps de calcul de la MEG peut se révéler prohibitif. Les méthodes itératives (cf chap. suivant) sont alors utilisées dans l'espoir de diminuer les temps de calcul de la MEG.

): Un P.C. à 200 MHz calcule environ 10^9 flops par secondes, soit $3 \cdot 10^{18}$ flops par années. Pour réaliser $51! \approx 10^{67}$ opérations, le P.C. devrait calculer pendant $5 \cdot 10^{49}$ années.

Lorsqu'on représente des nombres réels sur un ordinateur (ou une calculatrice), il résulte généralement des erreurs d'arrondis. Par exemple, la représentation à 4 chiffres significatifs du rationnel $b = \frac{1}{3}$ est :

$$b + \delta b = 0,3333$$

ce qui génère une erreur relative d'arrondi $\frac{\delta b}{b} = 10^{-4}$.

Supposons que de telles erreurs d'arrondi sont introduites dans le second-membre du système linéaire $\underline{A}\underline{x} = \underline{B}$ qui devient alors :

$$\underline{A}\underline{x} = \underline{B} + \underline{\delta B}.$$

Dans ce cas, les opérations arithmétiques (additions, mult., ... de sa résolution par méthode de Gauss vont propager les erreurs d'arrondis jusqu'à la solution calculée $\underline{x}$ +

Dans certains cas "critiques", on observe qu'une faible perturbation ($\frac{\|\delta B\|}{\|B\|} \ll 1$) va s'amplifier jusqu'à

entraîner de larges erreurs ($\frac{\|\delta x\|}{\|x\|} \gg 1$) sur la

solution calculée, ce qui est inacceptable : On dit alors que le système est mal conditionné.

Ce paragraphe a pour but de :

- Définir le nombre de condition $\chi(\underline{A})$ d'un système linéaire.
- Estimer les erreurs de la méthode de Gauss pour les systèmes mal conditionnés.
- Vous convaincre que de nombreux problèmes de la mécanique sont mal conditionnés.

10-1 Exemple de systèmes mal conditionnés

Considérons le système linéaire d'ordre 2 :

$$\underline{A}\underline{x} = \underline{B}, \quad (20)$$

avec $\underline{A} = \begin{bmatrix} 1 & 1 \\ 1 & \frac{10001}{10000} \end{bmatrix}$ et $\underline{B} = \begin{pmatrix} 2 \\ 2 \end{pmatrix}$. (20b)

L'unique solution de (20) est $\underline{x} = \begin{pmatrix} 2 \\ 0 \end{pmatrix}$.

Considérons maintenant une petite perturbation de ce système linéaire, soit :

$$\underline{A}\underline{x}' = \underline{B} + \underline{\delta B}, \quad (21)$$

avec $\underline{\delta B} = \begin{pmatrix} 0 \\ 10^{-4} \end{pmatrix}$. La solution de ce système est $\underline{x}' = \underline{x} + \underline{\delta x}$.

Pour cet exemple, on remarque avec étonnement qu'un changement relatif du second membre, de l'ordre de :

$$\frac{\|\delta B\|}{\|B\|} = \frac{10^{-4}}{2\sqrt{2}} \approx 0.04\%,$$

produit un changement relatif dans la solution :

$$\frac{\|\delta x\|}{\|x\|} = \frac{\sqrt{2}}{2} \approx 71\%,$$

soit une amplification de $\frac{\|\delta x\|/\|x\|}{\|\delta B\|/\|B\|} \approx 2 \cdot 10^4$ de la

perturbation du second-membre. Que se passe-t-il ?

explication géométrique

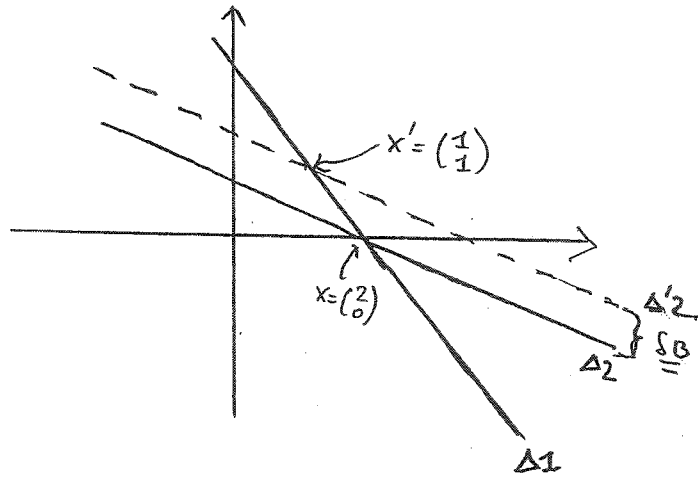
On peut faire une première estimation qualitative de ce phénomène en interprétant la solution du système soit :

$$\begin{cases} x + y = 2 \\ \dots \end{cases} \quad (21)$$

Comme l'intersection $X = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$ des deux droites Δ_1 et Δ_2 .
 Ces 2 droites sont quasiment parallèles : une petite chose
 translation de Δ_2 , par exemple :

$$x + (1.0001)y = 2 + 10^{-4} \quad (\Delta_2')$$

va produire un grand changement sur l'intersection
 de ces deux droites, qui est maintenant $X' = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$.



IV-2 Nombre de Conditionnement

On souhaite estimer l'erreur relative $\frac{\| \delta X \|}{\| X \|}$ que l'on
 commet sur la solution, en fonction des données du
 problème ($\underline{A}$ et $\underline{B}$), et de la perturbation $\delta \underline{B}$. Pour cela,
 on dispose du théorème :

Théorème On appelle $\chi(\underline{A}) \geq 1$ le nombre de condition-
nement du système (20) défini par :

$$\chi(\underline{A}) = \| \underline{A} \| \| \underline{A}^{-1} \|, \quad (22)$$

où $\| \cdot \|$ est une norme matricielle. On dispose de
 l'estimation d'erreur :

$$0 \leq \frac{\frac{\| \delta X \|}{\| X \|}}{\frac{\| \delta B \|}{\| B \|}} \leq \chi(\underline{A}) \quad (23).$$

Démonstration : On rappelle qu'une norme matricielle $\| \cdot \|$
 est définie par :

$$\| \underline{A} \| = \sup_{\substack{X \in \mathbb{R}^n \\ X \neq 0}} \frac{\| \underline{A} X \|}{\| X \|}.$$

En particulier, on en déduit que pour tout $X \in \mathbb{R}^n$, on a
 l'inégalité bien utile :

$$\| \underline{A} X \| \leq \| \underline{A} \| \| X \|. \quad (24)$$

A partir de la relation $\underline{B} = \underline{A} X$, on déduit de cette inégalité q

$$\| \underline{B} \| \leq \| \underline{A} \| \| X \|,$$

et de $X = \underline{A}^{-1} \underline{B}$, on déduit que :

$$\| X \| \leq \| \underline{A}^{-1} \| \| \underline{B} \|,$$

ces 2 relations conduisent à l'encadrement :

$$\frac{1}{\| \underline{A}^{-1} \| \| \underline{B} \|} \leq \frac{1}{\| X \|} \leq \frac{\| \underline{A} \|}{\| \underline{B} \|} \quad (25).$$

D'une manière analogue, à partir des relations $\delta \underline{B} = \underline{A} \delta X$
 et $\delta X = \underline{A}^{-1} \delta \underline{B}$ on déduit l'encadrement :

$$\frac{\| \delta \underline{B} \|}{\| \underline{A} \|} \leq \| \delta X \| \leq \| \underline{A}^{-1} \| \| \delta \underline{B} \| \quad (26).$$

A partir de (25) et (26), on obtient finalement :

$$\frac{1}{\chi(\underline{A})} \frac{\| \delta \underline{B} \|}{\| \underline{B} \|} \leq \frac{\| \delta X \|}{\| X \|} \leq \chi(\underline{A}) \frac{\| \delta \underline{B} \|}{\| \underline{B} \|},$$

avec $\chi(\underline{A}) = \| \underline{A} \| \| \underline{A}^{-1} \|$, d'où on déduit (22).

ou est commun à deux (ou) matrices que la norme relative (sur la solution) peut être aussi bien "innofensive" $\frac{\| \delta x \|}{\| x \|} \geq 0$, que potentiellement dangereuse $\frac{\| \delta x \|}{\| x \|} \geq 2$ lorsque $\chi(A)$ est grand.

Comment calculer "efficacement" le nombre de conditionnement $\chi(A) = \|A\| \|A^{-1}\|$ sans calculer les normes matricielles. Pour un certain type de matrices, très courant en ma'-conique, il suffit de connaître les val. propres de A pour calculer $\chi(A)$:

Propriété: Une matrice symétrique A est dite définie positive (SDP) si toutes ses valeurs propres sont strictement positives. Le conditionnement d'une matrice SDP est simplement

$$\chi(A) = \frac{\lambda_{\max}}{\lambda_{\min}} \geq 1,$$

où $\lambda_{\max}$ et $\lambda_{\min}$ sont respectivement la plus grande et plus petite valeur propre de A .

Calculons $\chi(A)$ pour l'exemple (20). La matrice A est symétrique, ses val. propres sont $\lambda_1 \approx 5 \cdot 10^5$ et $\lambda_2 \approx 2$, elle est donc SPD et son conditionnement est:

$$\chi(A) = \frac{\lambda_2}{\lambda_1} \approx 4 \cdot 10^4.$$

L'estimation d'erreur (2) devient alors:

$$0 \leq \frac{\| \delta x \| / \| x \|}{\| \delta B \| / \| B \|} \leq 4 \cdot 10^4$$

et on avait observé que $\| \delta x \| / \| x \| \approx 2 \cdot 10^4 = \frac{1}{2} \chi(A)$

pour l'exemple (20). Hâti, sans faire de choix particulier sur la perturbation δB , on observe une erreur relative proche du cas extrême $\chi(A) \approx 4 \cdot 10^4$.

Désormais, on qualifie les systèmes linéaires pour lesquels $\chi(A)$ est grand ~~comme~~ de systèmes mal conditionnés qui sont (potentiellement) très sensibles aux erreurs d'arr.

Les systèmes linéaires issus de la mécanique sont notoirement connus pour être mal conditionnés. Par exemple, la TD4 nous montre que les val. propres de la matrice de diff. finies symétrique K_n (Eq. (7)) sont:

$$\lambda_k = 2 - 2 \cos \frac{k\pi}{n+1}, \quad k = 1, \dots, n.$$

La matrice K_n est donc SDP, et son conditionnement est:

$$\chi(K_n) = \frac{\lambda_{\max}}{\lambda_{\min}} = \frac{1 + \cos \frac{\pi}{n+1}}{1 - \cos \frac{\pi}{n+1}},$$

et, lorsque n est grand, un développement limité montre que

$$\chi(K_n) \approx \frac{4}{\pi^2} (n+1)^2.$$

Cette matrice est donc mal conditionnée pour des problèmes de grande taille.

IV - 3 Conditionnement et erreurs d'arrondi des méthodes directes

Même si on ne connaît pas la solution exacte d'un système linéaire, il est possible d'évaluer l'importance des erreurs d'arrondis dues par une méthode directe.

Supposons que pour résoudre le système linéaire $\underline{A}\underline{x} = \underline{B}$ on ait utilisé une méthode directe (p.ex. la ME 6) pour obtenir une solution $\underline{x}'$, qui est entachée d'erreurs d'arrondis : alors forcément la quantité

$$\underline{A}\underline{x}' - \underline{B} \neq 0.$$

On définit alors le résidu comme :

$$\underline{r} = \underline{B} - \underline{A}\underline{x}', \quad (27)$$

et forcément $\underline{x}'$ est la solution exacte du système :

$$\underline{A}\underline{x}' = \underline{B} + \underline{r},$$

où $\underline{r}$ correspond à la petite perturbation $\delta \underline{B}$ du syst. (21).
L'estimation d'erreur (23) appliquée à ce système donne, dans le cas extrême :

$$\frac{\|\underline{x}' - \underline{x}\|}{\|\underline{x}\|} \geq \chi(\underline{A}) \frac{\|\underline{r}\|}{\|\underline{B}\|}$$

Ce qui montre que l'erreur relative croît avec le nombre de conditionnement : lorsque $\chi(\underline{A})$ est suffisamment grand les erreurs sur la solution peuvent être du même ordre que la solution elle-même : Très mauvais !

Exemple On considère le système linéaire $\underline{A}\underline{x} = \underline{B}$ où $\underline{A}$ est la matrice de Hilbert d'ordre n définie par :

$$\underline{A}_{ij} = \frac{1}{i+j-1}, \quad i, j = 1, \dots, n. \quad (28)$$

On montre que son conditionnement se comporte comme $\chi(\underline{A})^2$ c'est donc un système très mal conditionné, même pour des systèmes de taille modeste ($n \geq 10$). La Fig 1 montre que l'erreur obtenue par méthode directe croît avec le nombre de conditionnement $\chi(\underline{A})$, et qu'elle devient inacceptable lorsque $\chi(\underline{A})$ est grand ($n \geq 12$). Par contre, l'erreur donnée par la méthode itérative reste acceptable.

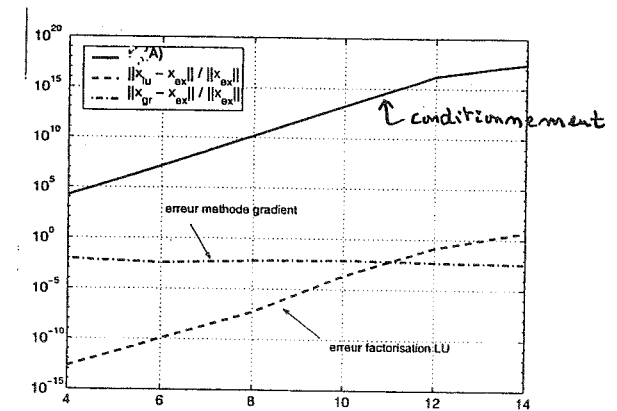


Fig. 1 : Résolution du système de Hilbert (28) pour diverses valeurs de n : Erreur relative entre la solution exacte $\underline{x}_{ex}$ et la solution donnée par :

- Méthode directe (fact. LU) : $\underline{x}_{LU}$
- Méthode itérative ("gradient") : $\underline{x}_{gr}$

Résolution de systèmes linéaires II:

- Méthodes itératives

I Généralités

I-1 Définition du problème

On veut résoudre le système linéaire d'ordre n

$$\underline{A} \underline{x} = \underline{F}, \quad (1)$$

où $\underline{A}$ est "régulière" ($\det \underline{A} \neq 0$). Son unique solution est notée $\underline{x}^*$. Pour cela, on souhaite définir une suite $\{\underline{x}^h, h \geq 0\}$ telle que:

$$\lim_{h \rightarrow +\infty} \underline{x}^h = \underline{x}^*.$$

I-2 Motivation

Le développement de méthodes itératives est basé sur les 2 principales carences des méthodes directes, qui nous avons identifiées dans le chapitre précédent:

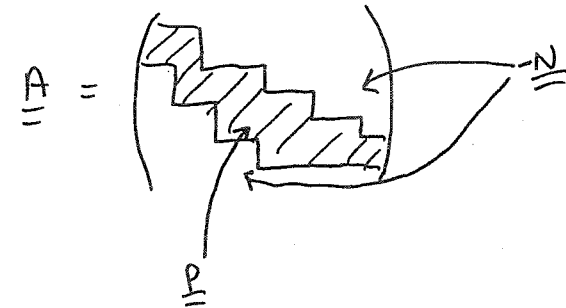
* Pour une matrice $\underline{A}$ quelconque, l'élimination de Gauss (ou méthode LU) nécessite $O(n^3)$ opérations arithmétiques (flops), ce qui peut être prohibitif pour des problèmes de grande taille.

En contraste, le passage $\underline{x}^h \rightarrow \underline{x}^{h+1}$ pour une méth. itérative nécessite seulement $O(n^2)$ flops (essentiellement l'évaluation du résidu $\underline{r}^h = \underline{b} - \underline{A} \underline{x}^h$).

Une méthode itérative peut donc devenir compétitive si elle converge en un nombre d'itérations indépendant de n .

* Si $\underline{A}$ est mal conditionnée, une méthode directe peut entraîner de larges erreurs d'arrondis. Ce phénomène est beaucoup moins sensible avec les méthodes itératives (cf. Fig 1 du chap. précédent) car, en particulier, on peut contrôler la valeur du résidu $\underline{r}_h$ de la solution.

II Principes de base: les méth. de décomposition ("Splitting")


$$\underline{A} = \underline{P} - \underline{N}$$

L'idée de base est de décomposer la matrice $\underline{A}$ en:

$$\underline{A} = \underline{P} - \underline{N}, \quad (2)$$

le problème (1) s'écrit alors de façon équivalente:

$$\underline{P} \underline{x} = \underline{N} \underline{x} + \underline{F}. \quad (3)$$

On définit alors la méthode itérative:

$$\underline{P} \underline{x}^{h+1} = \underline{N} \underline{x}^h + \underline{F} \quad (4)$$

La matrice $\underline{P}$, appelée matrice de préconditionnement

- $\underline{P}$ est inversible (i.e., $\underline{P}^{-1}$ existe)
- $\underline{P}$ est "facile" à inverser.

Sous ces hypothèses, on peut formellement écrire l'itération sous la forme :

$$\underline{X}^{k+1} = \underbrace{\underline{P}^{-1}\underline{N}}_{\underline{B}} \underline{X}^k + \underline{P}^{-1}\underline{F} \quad (5)$$

On peut remarquer que cette itération correspond à une méthode du point fixe $\underline{x}^{k+1} = \phi(\underline{x}^k)$. (cf cours sur les syst. non-linéaires).

Par analogie à ce type de méthodes, la matrice $\underline{B} = \underline{P}^{-1}\underline{N}$ est appelée matrice d'itération.

Il est aisé de vérifier que la méthode (4) est consistante

$$\begin{aligned} \underline{P}\underline{X}^* &= \underline{N}\underline{X}^* + \underline{F} \Leftrightarrow (\underline{P} - \underline{N})\underline{X}^* = \underline{F} \\ &\Leftrightarrow \underline{A}\underline{X}^* = \underline{F} \end{aligned}$$

Examinons maintenant la convergence de la méthode (4).

Définition (Convergence): Une méthode itérative est convergente si l'erreur à la $k^{\text{ième}}$ itération $\underline{E}^k = \underline{X}^k - \underline{X}^*$ vérifie :

$$\lim_{k \rightarrow +\infty} \|\underline{E}^k\| = 0$$

En combinant (4) et (3), on obtient :

$$\underline{X}^k - \underline{X}^* = \underline{B} (\underline{X}^{k-1} - \underline{X}^*)$$

d'où on déduit l'inégalité :

$$\|\underline{E}^k\| \leq \|\underline{B}\| \|\underline{E}^{k-1}\|$$

on obtient finalement :

$$\|\underline{E}^k\| \leq [\rho(\underline{B})]^k \|\underline{E}^0\|.$$

Ainsi, l'itération (4), qui est consistante, converge si et seulement si est elle est stable :

$$\boxed{\rho(\underline{B}) < 1.} \quad (6)$$

Rq: Contrairement aux meth. itér. pour syst. non lin., on peut remarquer que l'itération (4) converge pour toute solution initiale $\underline{X}^0 \in \mathbb{R}^n$.

III Les méthodes itératives de base

Dans cette partie, on va présenter les 3 méthodes "historiques" de Jacobi, Gauss-Seidel, relaxation (S.O.R.), qui reposent toutes sur la décomposition (2) :

$$\underline{A} = \underline{P} - \underline{N}.$$

Elles diffèrent l'une de l'autre par le choix du préconditionneur $\underline{P}$ qu'il faut inverser.

Leurs performances dépendent des valeurs propres de la matrice d'itération $\underline{B} = \underline{P}^{-1}\underline{N}$, plus particulièrement de son rayon spectral :

$$\rho(\underline{B}) = \max_{i=1, \dots, n} |\lambda_i(\underline{B})|.$$

En effet, d'après l'estimation d'erreur :

$$\|\underline{E}^k\| \leq \rho(\underline{B}) \|\underline{E}^{k-1}\|,$$

on déduit que plus $\rho(\underline{B})$ sera petit, plus la convergence sera rapide.

On comparera les performances de ces méthodes sur le problème test d'ordre 2 :

$$\underline{A} \underline{x} = \underline{F}, \quad \text{avec} \quad \underline{A} = \begin{bmatrix} 2 & -1 \\ -1 & 2 \end{bmatrix} \quad (7)$$

qui correspond à la matrice de différence-finies K_n avec du chapitre précédent. Le cas n quelconque sera examiné en TD.

III-1 Décompositions de la matrice $\underline{A}$

Toutes les méthodes itératives de base reposent sur la décomposition :

$$\underline{A} = \begin{bmatrix} \diagup & \diagdown \\ -\underline{L} & \underline{D} \end{bmatrix}, \quad \text{soit} \quad \underline{A} = \underline{D} - \underline{L} - \underline{U} \quad (8)$$

tel que :

* $\underline{D}$ ("diagonal") est composée des éléments diagonaux de

$$\underline{D}_{i,j} = \begin{cases} A_{ii} & \text{pour } i = 1, \dots, n \\ 0 & \text{pour } i \neq j \end{cases}$$

* $\underline{L}$ ("Lower") : éléments triangulaires inférieurs de $\underline{A}$:

$$\underline{L}_{i,j} = \begin{cases} -A_{ij} & \text{si } i > j \text{ (ie, pour } j=1, \dots, i-1) \\ 0 & \text{si } i \leq j \end{cases}$$

* $\underline{U}$ ("Upper") : éléments triangulaires supérieurs de $\underline{A}$:

$$\underline{U}_{i,j} = \begin{cases} -A_{ij} & \text{si } i < j \text{ (ie, pour } j=i+1, \dots, n) \\ 0 & \text{si } i \geq j \end{cases}$$

Exemple : Pour la matrice diff. finies^{AM} d'ordre 3

$$\underline{A} = \begin{bmatrix} 2 & -1 & 0 \\ -1 & 2 & -1 \\ 0 & -1 & 2 \end{bmatrix},$$

la décomposition (8) correspond à :

$$\underline{D} = \begin{bmatrix} 2 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 2 \end{bmatrix}, \quad \underline{L} = \begin{bmatrix} 0 & 0 & 0 \\ 1 & 0 & 0 \\ 0 & 1 & 0 \end{bmatrix}, \quad \underline{U} = \begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \end{bmatrix}$$

III-2 La méthode de Jacobi

La méthode la plus simple consiste à choisir le préconditionneur

$$\underline{P}_J = \underline{D}, \quad (9)$$

car dans ce cas le préconditionneur est trivial à inverser :

$$\underline{D}^{-1} = \begin{cases} \frac{1}{A_{ii}} & \text{pour } i=j \\ 0 & \text{sinon} \end{cases}$$

Ce choix entraîne $\underline{N}_J = \underline{P}_J^{-1} \underline{A} = \underline{L} + \underline{U}$, et l'itération (4) devient :

$$\underline{x}^{h+1} = \underline{D}^{-1} (\underline{L} + \underline{U}) \underline{x}^h + \underline{D}^{-1} \underline{F}. \quad (10a)$$

En posant $\underline{x}^h = (x_i^h, i=1, \dots, n)$ et $\underline{F} = (f_i, i=1, \dots, n)$, elle s'écrit pour les composantes scalaires :

$$x_i^{h+1} = \frac{1}{A_{ii}} \left(f_i - \sum_{\substack{j=1 \\ j \neq i}}^n A_{ij} x_j^h \right), \quad i=1, \dots, n. \quad (10)$$

Cette méthode est connue sous le nom de méth. de Jacobi

Le préconditionneur est :

$$\underline{P}_S = \begin{bmatrix} 2 & 0 \\ 0 & 2 \end{bmatrix} \Rightarrow \underline{P}_S^{-1} = \begin{bmatrix} 1/2 & 0 \\ 0 & 1/2 \end{bmatrix}$$

d'où on déduit :

$$\underline{N}_S = \underline{P}_S - \underline{A} = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \Rightarrow \underline{B}_S = \underline{P}_S^{-1} \underline{N}_S = \begin{bmatrix} 0 & 1/2 \\ 1/2 & 0 \end{bmatrix}$$

L'itération (10a) s'écrit donc :

$$\underline{X}^{k+1} = \begin{bmatrix} 0 & 1/2 \\ 1/2 & 0 \end{bmatrix} \underline{X}^k + \begin{pmatrix} f_{1/2} \\ f_{2/2} \end{pmatrix}. \quad (11)$$

Les valeurs propres de $\underline{B}_S$ sont $1/2$ et $-1/2$, et le facteur de réduction est donc :

$$\rho(\underline{B}_S) = \frac{1}{2}.$$

La méthode de Jacobi converge et, de plus :

$$\|\underline{E}^k\| \leq \frac{1}{2} \|\underline{E}^{k-1}\|,$$

Ce qui montre que l'erreur est au moins divisée par 2 à chaque itération.

Pour ce problème simple, la méthode de Jacobi se révèle efficace. Comme on le verra en TD, son efficacité chute considérablement pour des systèmes mal conditionnés.

Des améliorations s'imposent !

III-3 La méthode de Gauss-Seidel

Afin de présenter cette méthode, écrivons l'itération de Jacobi (11) pour les composantes :

$$2 x_1^{k+1} = x_2^k + f_1$$

puis,

$$2 x_2^{k+1} = x_1^k + f_2.$$

Une amélioration est d'utiliser dans cette dernière ligne la dernière (et meilleure) estimation de x_1 , soit x_1^{k+1} , ce qui donne :

$$2 x_1^{k+1} = x_2^k + f_1$$

puis,

$$2 x_2^{k+1} = x_1^{k+1} + f_2.$$

En regroupant à gauche l'itér $k+1$, cette nouvelle méthode s'écrit matriciellement :

$$\begin{bmatrix} 2 & 0 \\ -1 & 2 \end{bmatrix} \begin{pmatrix} x_1^{k+1} \\ x_2^{k+1} \end{pmatrix} = \begin{bmatrix} 0 & 1 \\ 0 & 1 \end{bmatrix} \begin{pmatrix} x_1^k \\ x_2^k \end{pmatrix} + \begin{pmatrix} f_1 \\ f_2 \end{pmatrix}$$

Ce qui correspond à l'itération :

$$(-\underline{L} + \underline{D}) \underline{X}^{k+1} = \underline{U} \underline{X}^k + \underline{F}. \quad (12)$$

C'est la méthode de Gauss-Seidel. Elle est de la forme (4) avec :

$$\underline{P}_{GS} = -\underline{L} + \underline{D}, \quad \underline{N}_{GS} = \underline{P}_{GS} - \underline{A} = \underline{U}. \quad (13)$$

Pour les composantes, elle s'écrit :

$$\text{Pour } i=1, \dots, n: x_i^{k+1} = \frac{1}{A_{ii}} \left[f_i - \sum_{j=1}^{i-1} A_{ij} x_j^{k+1} - \sum_{j=i+1}^n A_{ij} x_j^k \right]$$

Application au problème test (7)

$$\underline{P}_{GS} = \begin{bmatrix} 2 & 0 \\ -1 & 2 \end{bmatrix}, \quad \underline{P}_{GS}^{-1} = \begin{bmatrix} 1/2 & 0 \\ 1/4 & 1/2 \end{bmatrix}, \quad \underline{N}_{GS} = \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix}$$

$$\Rightarrow \underline{B}_{GS} = \begin{bmatrix} 0 & 1/2 \\ 0 & 1/4 \end{bmatrix}$$

Les val. propres de $\underline{B}_{GS}$ sont 0 et $1/4$, donc :

$$\rho(\underline{B}_{GS}) = [\rho(\underline{B}_{GS})]^2 = \frac{1}{4}.$$

La méthode de Gauss-Seidel dérive l'ensemble par 4 à chaque itération : on en déduit qu'une itération de G-S équivaut à une itération de Jacobi ; pour un coût algorithmique raisonnable : la méthode de Gauss-Seidel est donc plus performante.

III La méthode de relaxation (ou S.O.R.)

(S.O.R. : "Successive Over Relaxation")

On rappelle que le préconditionneur de la méthode de Gauss-Seidel est (cf. Eq. (13)) :

$$\underline{P} = \underline{D} - \underline{L}.$$

Une généralisation de la méthode G.S. est de considérer maintenant le préconditionneur :

$$\underline{P}_w = \frac{1}{w} \underline{D} - \underline{L}, \quad \text{avec } w \in \mathbb{R}. \quad (15)$$

Le paramètre de relaxation (ou d'accélération) w devra être choisi de manière à obtenir le facteur de réduction ρ_w le plus petit possible. On remarquera que le cas $w=1$ correspond à la méthode de Gauss-Seidel.

Pour obtenir une méthode consistante, on définit $\underline{N}_w$ par

$$\underline{N}_w = \underline{P}_w - \underline{A} = \underline{P}_w - (\underline{D} - \underline{L} - \underline{U}) = \left(\frac{1-w}{w}\right) \underline{D} + \underline{U}. \quad (16)$$

Avec ces définitions de $\underline{P}_w$ et $\underline{N}_w$, l'itération (4) s'écrit :

$$\left(\frac{1}{w} \underline{P} - \underline{L}\right) \underline{x}^{h+1} = \left[\left(\frac{1-w}{w}\right) \underline{D} + \underline{U}\right] \underline{x}^h + \underline{F}.$$

Après multiplication de cette équation par w , on obtient pour la composante :

$$\text{pour } i=1, \dots, n: \quad x_i^{h+1} = (1-w)x_i^h + \frac{w}{A_{ii}} \left(f_i - \sum_{j=1}^{i-1} A_{ij} x_j^{h+1} - \sum_{j=i+1}^n A_{ij} x_j^h \right)$$

Comment choisir le paramètre w ? cela se fait en calculant, pour chaque problème $\underline{A}\underline{x} = \underline{F}$, la matrice d'itération $\underline{B}(w) = \underline{P}_w^{-1} \underline{N}_w$ et en cherchant la valeur de w qui minimise son rayon spectral ρ_w . On présente le calcul de w pour le problème test (7) plus loin dans le cours.

Le théorème suivant détermine les valeurs "permissibles" de w :

Théorème : Le facteur de réduction ρ_w de la méthode S.O. vérifie :

$$|1-w| \leq \rho_w.$$

Ainsi, pour que la méthode converge (i.e. $\rho_w < 1$), il faut nécessairement que :

$$0 < w < 2.$$

Démonstration : On note λ_i les valeurs propres de la matrice d'itération $\underline{B}(w) = \underline{P}_w^{-1} \underline{N}_w$. Par définition du rayon spectral, on a :

$$\rho_w = \max_{i=1, \dots, n} |\lambda_i| \geq \sqrt[n]{\prod_{i=1}^n |\lambda_i|} \quad (16)$$

à une part, on voit que :

$$\prod_{i=1}^n |\lambda_i| = |\det(\underline{\underline{B}}(\omega))| = \frac{|\det(\underline{\underline{N}}\omega)|}{|\det(\underline{\underline{P}}\omega)|}$$

Calculons ces déterminants :

- $\underline{\underline{P}}\omega = \frac{1}{\omega} \underline{\underline{D}} - \underline{\underline{L}}$ est triangulaire inférieure, de diagonal $\frac{1}{\omega} \underline{\underline{D}}$, donc :

$$\det(\underline{\underline{P}}\omega) = \frac{\det(\underline{\underline{D}})}{\omega^n}$$

- De même, pour $\underline{\underline{N}}\omega = \frac{1-\omega}{\omega} \underline{\underline{D}} + \underline{\underline{U}}$ triang. supérieure :

$$\det(\underline{\underline{N}}\omega) = \left(\frac{1-\omega}{\omega}\right)^n \det(\underline{\underline{D}})$$

Ainsi, l'inégalité (16) devient :

$$\rho_\omega \geq \sqrt[n]{|\det(\underline{\underline{B}}(\omega))|} = |1-\omega|$$

Pour converger, il faut que $\rho_\omega < 1$, donc nécessairement ω doit être choisi tel que :

$$0 < \omega < 2$$

Application au problème test (7)

Nous allons maintenant déterminer le choix "optimal" de $\omega \in]0, 2[$ pour le problème test, qui correspond à minimiser le facteur de réduction ρ_ω de la méthode S.O.R. Avec ce choix de ω , la méthode S.O.R. présentera un bien meilleur taux de convergence que la méthode de Gauss-Seidel (et donc de Jacobi).

Pour la matrice $\underline{A} = \begin{bmatrix} 2 & -1 \\ -1 & 2 \end{bmatrix}$, on obtient : (12)

$$P_w = \frac{1}{w} \underline{D} - \underline{L} = \begin{bmatrix} 2/w & 0 \\ -1 & 2/w \end{bmatrix}, \quad N_w = \left(\frac{1-w}{w} \right) \underline{D} + \underline{U} = \begin{bmatrix} \frac{2(1-w)}{w} & 1 \\ 0 & \frac{2(1-w)}{w} \end{bmatrix}$$

L'itération S.O.R s'écrit :

$$\begin{bmatrix} 2/w & 0 \\ -1 & 2/w \end{bmatrix} \begin{pmatrix} x_1^{k+1} \\ x_2^{k+1} \end{pmatrix} = \begin{bmatrix} \frac{2(1-w)}{w} & 1 \\ 0 & \frac{2(1-w)}{w} \end{bmatrix} \begin{pmatrix} x_1^k \\ x_2^k \end{pmatrix} + \begin{pmatrix} f_1 \\ f_2 \end{pmatrix}$$

La matrice d'itération de la méthode s'écrit :

$$B_w = P_w^{-1} N_w = \begin{bmatrix} 1-w & w/2 \\ \frac{w(1-w)}{2} & \frac{1}{4} w^2 + 1-w \end{bmatrix}$$

Les valeurs propres de la matrice B_w sont :

$$\lambda_{+,-} = 1-w + \frac{w^2}{8} \pm \frac{w}{8} \sqrt{p(w)}$$

$$\text{avec } p(w) = w^2 - 16w + 16 = (w - w_0)(w - w_1)$$

$$\text{tel que } w_1 = 8 + 4\sqrt{3} > 2 \text{ et } w_0 = 8 - 4\sqrt{3} \approx 1,07.$$

On remarque que $p(w)$ change de signe dans $]0, 2[$: les valeurs propres λ_+ et λ_- peuvent donc être complexes.

Ainsi, pour calculer $\rho_w = \max \{ |\lambda_+|, |\lambda_-| \}$, il faut considérer 2 cas :

1^{er} cas : Pour $w \in]0, w_0[$, on a $p(w) > 0$, les valeurs propres sont donc réelles et on trouve facilement :

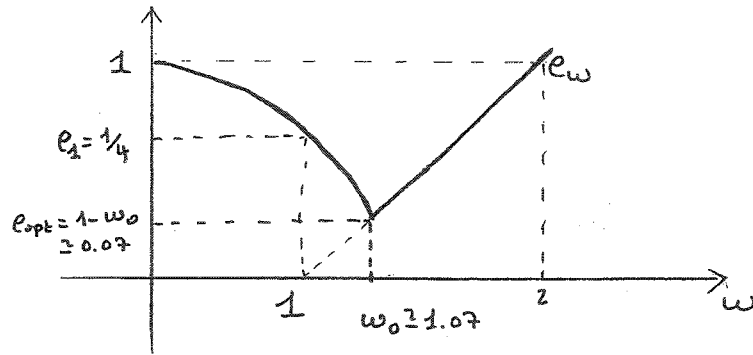
$$\rho_w = 1 - w + \frac{w^2}{8} + \frac{w}{8} \sqrt{p(w)}$$

propres sont des complexes conjugués. On obtient :

$$\rho_w = |\lambda_+| = |\lambda_-| = \sqrt{\left(1 - w + \frac{w^2}{8}\right)^2 + \frac{w^2}{64}(-p(w))}$$

$$= w - 1 \quad \text{après calcul.}$$

On en déduit le graphe de ρ_w (pas à l'échelle !)



En premier lieu, on retient bien que :

* Pour $w \leq 0$ et $w \geq 2$, on a $\rho_w \geq 1 \Rightarrow$ la méthode diverge

* Pour $w = 1$, on retient le facteur de réduction de la méthode de Gauss-Seidel :

$$\rho_1 = \rho_{GS} = \frac{1}{4}$$

Graphiquement, on observe que le minimum de ρ_w est obtenu en $w = w_0 \approx 1.07$. Le facteur de réduction optimal de la méthode S.O.R est donc :

$$\rho_{opt} = \rho_{w_0} = w_0 - 1 \approx 0.07.$$

De plus, on remarque que :

$$\rho_{opt} \approx (\rho_{GS})^2$$

ce qui montre qu'une itération de S.O.R avec $w = w_{opt}$:

- * équivaut à 2 itérations de G.-S.
- * équivaut à 4 itérations de Jacobi.

En conclusion de ce paragraphe :

* Par un choix intelligent du paramètre w , et donc des valeurs propres de la matrice d'itération, on peut bâtir une méthode itérative performante.

* Dans le cas d'une matrice $\underline{A}$ générale, le calcul de w_{opt} peut être très compliqué. Heureusement, pour les systèmes lin. issus de la mécanique, il existe des théorèmes mathématiques pour déterminer w_{opt} facilement (cf T.D.).

IV Résultats de convergence pour les méthodes itératives de bas

Supposons que vous souhaitiez résoudre le système $\underline{A}\underline{x} = \underline{F}$, avec $\underline{A}$ quelconque, en utilisant la meth. de Jacobi, de G. ou S.O.R. Afin de s'assurer que la méthode que vous choisissez soit convergente, il faudrait calculer son facteur de réduction :

$$\rho(\underline{B}) = \rho(\underline{P}^{-1}\underline{N}) = \rho(\underline{I} - \underline{P}^{-1}\underline{A})$$

et vérifier qu'il est strictement inférieur à 1. Le calcul est, en général, très fastidieux. Heureusement, pour certaines matrices spéciales (que l'on retrouve en mécanique), nous disposons de théorèmes bien utiles pour arriver à déterminer le facteur de réduction de ces méthodes itératives.

La première classe de matrices spéciales regroupe les matrices SDP (rappel: une matrice symétrique est dite définie positive si ses val. propres sont strictement positives).

On dispose d'un critère bien pratique pour déterminer si une matrice est SDP:

Critère de Sylvester: Une matrice $\underline{A}$ symétrique est définie positive si

$$\det(\underline{A}_k) > 0, \quad k=1, \dots, n$$

où $\underline{A}_k = (\underline{A}_{ij})_{i,j=1, \dots, k}$ est la sous matrice principale d'ordre k de $\underline{A}$.

Exemple: Considérons la matrice $\underline{A} = \begin{bmatrix} 2 & -1 & 0 \\ -1 & 2 & -1 \\ 0 & -1 & 2 \end{bmatrix}$.

On calcule ses sous-matrices principales:

$$\underline{A}_1 = [2] \Rightarrow \det(\underline{A}_1) = 2 > 0$$

$$\underline{A}_2 = \begin{bmatrix} 2 & -1 \\ -1 & 2 \end{bmatrix} \Rightarrow \det(\underline{A}_2) = 3 > 0$$

$$\underline{A}_3 = \underline{A} \Rightarrow \det(\underline{A}_3) = 4 > 0.$$

Donc, d'après le critère de Sylvester, elle est SDP (on savait déjà, cf. TD).

Pour les matrices SDP, on dispose du théorème:

Théorème: Si $\underline{A}$ est SDP alors la méthode de G-S et la méthode S.O.R (pour $w \in]0, 2[$) convergent.

Démon: (A suivre)

Rq: Il n'y a aucun résultat de convergence pour les matrices SDP avec la méthode de Jacobi. Dans certains cas, cette méthode converge (p.ex. le plomb test 7) dans d'autre, elle diverge (p.ex. la matrice de Hilbert $H_{ij} = \frac{1}{i+j-1}$ d'ordre $n=3$, qui est SDP, vérifie $\rho_3 > 1$).

Une autre classe intéressante de matrices regroupe les matrices à diagonale dominante:

Définition: Une matrice $\underline{A}$ est dite diagonale dominante stricte par ligne si:

$$|A_{ii}| > \sum_{\substack{j=1 \\ j \neq i}}^n |A_{ij}|, \quad \forall i=1, \dots, r$$

Exemple: Considérons la matrice d'ordre 3:

$$\underline{A} = \begin{bmatrix} -3 & 2 & 0 \\ 1 & 4 & -2 \\ 1 & -1 & 3 \end{bmatrix}$$

pour la ligne 1: $|-3| > 2$, ligne 2: $|4| > |-2| + |1|$, ligne 3: $|3| > |1| + |-1|$
donc cette matrice est strictement diagonale dominante par ligne.

Théorème: Si $\underline{A}$ est diagonale dominante stricte par ligne, alors les méthodes de Jacobi et Gauss-Seidel convergent, et la méthode S.O.R converge si $w \in]0, 1[$.

La méthode de Richardson

Dans les sections précédentes, on a établi que les méthodes itératives de base (Jacobi, G-S, S.O.R.) s'écrivent sous la forme (4):

$$\underline{P} \underline{X}^{h+1} = \underline{N} \underline{X}^h + \underline{F}, \quad (17)$$

avec $\underline{N} = \underline{P}^{-1} \underline{A}$, et la matrice d'itération s'écrit:

$$\underline{B}_1 = \underline{P}^{-1} \underline{N} = \underline{I} - \underline{P}^{-1} \underline{A}.$$

Il est intéressant d'écrire (17) sous la forme:

$$\underline{P} \underline{X}^{h+1} = \underline{P} \underline{X}^h + \underline{R}^h, \quad (18)$$

avec $\underline{R}^h$ le résidu de l'équation à la $h^{\text{ème}}$ itération, défini par:

$$\underline{R}^h = \underline{F} - \underline{A} \underline{X}^h. \quad (19)$$

Une généralisation de (18) consiste à introduire un paramètre d'accélération $\alpha \neq 0$, tel que:

$$\underline{P} \underline{X}^{h+1} = \underline{P} \underline{X}^h + \alpha \underline{R}^h. \quad (20)$$

Cette itération est appelée méthode de Richardson.

Sa matrice d'itération a la forme:

$$\underline{B}_2 = \underline{I} - \alpha \underline{P}^{-1} \underline{A} \quad (21)$$

Le paramètre $\alpha \neq 0$ doit être choisi de manière à accélérer la convergence de l'itération (18), soit:

$$\rho(\underline{B}_2) < \rho(\underline{B}_1).$$

Comme on l'a vu en TD, il est possible de calculer la valeur optimale de α à partir des val. propres de $\underline{P}^{-1} \underline{A}$:

Théorème: Si $\underline{P}^{-1} \underline{A}$ possède des val. propres positives strictement, notées $\lambda_1 > \lambda_2 > \dots > \lambda_n > 0$, alors le paramètre α_{opt} qui minimise $\rho(\underline{B}_2)$ est:

$$\alpha_{\text{opt}} = \frac{2}{\lambda_1 + \lambda_n}$$

et la méthode de Richardson a pour facteur de réduction

$$\rho_{\text{opt}} = \frac{\lambda_1 - \lambda_n}{\lambda_1 + \lambda_n} < 1.$$

La méthode (21), avec $\alpha = \alpha_{\text{opt}}$, appelée méthode de Richardson à pas fixe, est un premier exemple d'accélérateurs de convergence.

Dans les années suivantes, vous verrez que la plupart des méthodes itératives modernes (ie, utilisées dans les logiciels commerciaux de calcul scientifique pour la mécanique), sont la combinaison de:

* Un préconditionneur $\underline{P}$: (ex: Jacobi, S.O.R., ou multigrilles)

* Un accélérateur de convergence: basé sur des améliorations de la méthode de Richardson (méthode du Gradient Conjugué, G.M.R.E.S, etc...)

On va utiliser ce théorème pour accélérer la méthode S.O.R., définie par $P_{SOR}(w_0) = \begin{bmatrix} \frac{2}{w_0} & 0 \\ -1 & \frac{2}{w_0} \end{bmatrix}$

et $w_0 = 8 - 4\sqrt{3}$ la valeur optimale du chap. IV.

Après calcul, on trouve que:

$$P^{-1}A = \begin{bmatrix} w_0 & -\frac{1}{2}w_0 \\ \frac{1}{2}w_0(w_0-1) & w_0(1-\frac{w_0}{4}) \end{bmatrix}$$

et les val. propres de cette matrice sont doubles:

$$\lambda_1 = \lambda_2 = \left(1 - \frac{1}{8}w_0 + \sqrt{16 - 16w_0 + w_0^2}\right)w_0.$$

Pour la valeur optimale $w_0 = 8 - 4\sqrt{3}$, on trouve:

$$\lambda_1 = \lambda_2 = 4\sqrt{3} - 6 > 0.$$

Les v.p. de $P^{-1}A$ sont positives, on peut appliquer le théorème pour trouver le paramètre d'accélération optimal:

$$\omega = \frac{2}{\lambda_1 + \lambda_2} = \frac{2\sqrt{3} + 3}{6} \approx 1.077.$$

Le facteur de réduction est alors:

$$\rho_2 = \frac{\lambda_1 - \lambda_2}{\lambda_1 + \lambda_2} = 0$$

ce qui est ^{encore} meilleur que $\rho_1 = w_0 - 1 \approx 0.97$ de la méthode S.O.R.

R.G.: L'application de ^{et accélération de CV} la méthode de Richardson est mieux adaptée rapidement ses limites: il faut que les v.p. de $P^{-1}A$ soient positives, ce qui n'est pas le cas pour le problème test avec $n \geq 2$. Dans les années soixante, on venait des accélérations de CV plus sophistiquées (méthodes du gradient).

V-1 Critère d'arrêt

Théoriquement, les méthodes itératives donnent la solution exacte d'un système linéaire lorsque $k \rightarrow \infty$. En pratique, on se satisfait d'arrêter l'algorithme lorsque ~~l'erreur relative est~~ nous jugeons que l'erreur relative est suffisamment petite, par exemple lorsque :

$$\frac{\|X^k - X\|}{\|X\|} < \delta \quad (22)$$

avec $\delta = 10^{-2}$ (i.e., 1% d'erreur relative).

Puisqu'on ne connaît pas la solution exacte, on effectue généralement ce test d'arrêt sur le résidu :

$$\frac{\|R^k\|}{\|F\|} < \epsilon, \quad (23)$$

où $\epsilon > 0$ est à choisir. Afin de bien choisir la tolérance nous allons relier ce test d'arrêt à (22). D'après la définition du résidu (19), on obtient :

$$A^T R^k = A^T F - X^k = X - X^k,$$

d'où on déduit l'inégalité :

$$\frac{\|X - X^k\|}{\|X\|} \leq \|A^T\| \frac{\|R^k\|}{\|F\|} \quad (24)$$

D'autre part, on tire de $F = AX$ l'inégalité :

$$\|F\| \leq \|A\| \|X\|,$$

et l'inégalité (24) devient finalement

$$\frac{\|X - X^k\|}{\|X\|} \leq \kappa(A) \frac{\|R^k\|}{\|F\|} \quad (25)$$

ou $\kappa(A)$ est le nombre de conditionnement $\kappa(A) = \frac{\|A\|}{\|A^{-1}\|}$.
Ainsi, pour obtenir une erreur relative sur la solution de l'ordre de $\delta = 1\%$, il suffit d'utiliser le critère d'arrêt avec :

$$\epsilon = \frac{\delta}{\kappa(A)}. \quad (26)$$

Contrairement aux méthodes directes, une méthode itérative permet de contrôler les erreurs d'arrondis ~~pour~~ après de la solution numérique (cf Fig. 1 pour les méthodes directes).

V-2 Algorithme pratique

On souhaite maintenant programmer l'itération de Richardson (20) sur un ordinateur. Pour cela, il est commode de la réécrire sous la forme :

$$P \underline{S}X = \underline{d} \underline{R}^k, \text{ avec } \underline{S}X = X^{k+1} - X^k$$

On l'implémente par l'algorithme suivant :

• Initialisation : calculer le résidu $\underline{R}^0 = F - A\underline{X}^0$

- On choisit le nombre d'itérations maximales, par exemple : $k_{\max} = 100$.
- On choisit la tolérance ϵ , par exemple avec le critère (26).
- On choisit une cond. initiale $X^0 \in \mathbb{R}^n$.
- On calcule le résidu initial $R^0 = F - AX^0$.

• Itérations : Pour $k = 0, \dots, k_{\max} - 1$:

- Résoudre $P \underline{S}X = \underline{d} \underline{R}^k$
- Calculer $X^{k+1} = X^k + \underline{S}X$
- Calculer $\underline{R}^{k+1} = F - AX^{k+1}$
- Test d'arrêt : on stoppe l'algo si $\frac{\|R^{k+1}\|}{\|F\|} < \epsilon$

Analyse Numérique - Licence de Mécanique

TD 1 - Interpolation - Intégration numérique

Exercice 1. (*Interpolation de Lagrange - Phénomène de Runge - Applications à l'intégration et différentiation numérique*)

Nous cherchons à interpoler une fonction $f(x)$ par un polynôme de degré $n = 2$ aux points $x_0 = -1$, $x_1 = 0$ et $x_2 = 1$.

- 1) Calculez les éléments $\{L_j(x)\}$ de la base de Lagrange. Représentez-les sur un graphe.
- 2) On veut interpoler une fonction quelconque $f(x)$. Ecrivez la forme générale du polynôme d'interpolation de Lagrange, et ordonnez-le en puissances de x .
- 3) Interpolez les fonctions :

$$f(x) = 1, \quad f(x) = x, \quad f(x) = x^2, \quad f(x) = x^3.$$

Que remarquez-vous. Pouviez-vous prévoir théoriquement ces résultats ?

- 4) Interpolez la fonction $f(x) = e^x$ dans $\Omega =]-1, 1[$ et représentez sur le même graphe cette fonction et son interpolant $\pi_2 f(x)$. Estimez l'erreur relative d'interpolation commise au point $x = 0.5$.
- 5) Même question pour la fonction de Runge définie dans $\Omega =]-1, 1[$ par :

$$f(x) = \frac{1}{1 + 100x^2}.$$

(Les interpolants de degré supérieur à 2 sont représentés sur la figure 1.)

- 6) Utilisez la question 2) pour bâtir une formule approchée permettant de calculer l'intégrale :

$$I(f) = \int_{-1}^1 f(t) \, dt.$$

Donnez une signification géométrique à cette formule. Utilisez-la pour calculer une valeur approchée de $I(e^x)$ et comparez avec la solution exacte.

- 7) A l'aide du changement de variables

$$x = \varphi(t) \equiv a + (b - a) \frac{t + 1}{2},$$

utilisez la question précédente pour établir une approximation de l'intégrale

$$\int_a^b g(x) \, dx.$$

(Cette formule de quadrature, utilisant une interpolation parabolique, est connue sous le nom de *formule de Simpson*).

- 8) On souhaite maintenant développer une approximation discrète de $f''(x)$ en un point $x = x_i$ qui utilise les valeurs de f aux points équidistants :

$$x_{i-1} = x_i - h, \quad x_i, \quad x_{i+1} = x_i + h, \quad \text{avec } h > 0.$$

Pour cela, construisez le polynôme d'interpolation $\pi_2 f(x)$ passant par ces trois points, dérivez-le deux fois, et montrez que vous obtenez la *formule aux différences* :

$$f''(x_i) \cong \frac{f(x_{i-1}) - 2f(x_i) + f(x_{i+1}))}{h^2}.$$

Pourquoi est-elle communément appelée formule centrée ?

Exercice 2. (Estimation de l'ordre des erreurs numériques)

Soit I_h une formule de quadrature composite de pas h , appliquée à l'intégration de $I(f)$. On suppose que I_h est une formule d'ordre q , et donc l'erreur commise avec un pas d'intégration h est telle que :

$$E_h = |I_h - I| \sim Ch^q, \quad (1)$$

lorsque h est suffisamment petit. Dans de nombreuses situations, on est amené à vérifier *a posteriori* que l'ordre de précision q est atteint, par exemple pour vérifier qu'il n'y a pas d'erreur dans un programme informatique. Cet exercice a pour but d'évaluer l'ordre de précision d'une méthode à partir des résultats numériques obtenus sur un problème test.

- 1) Le rapport entre l'erreur commise avec un pas h et un pas réduit de moitié est appelé facteur de réduction (FdR). Montrez que pour une méthode d'ordre q , le FdR s'écrit :

$$\text{FdR} = 2^q.$$

En déduire l'intérêt que présente une méthode d'ordre élevé (p.ex. $q = 4$) par rapport à une méthode basique (p.ex. $q = 2$).

- 2) Application : pour calculer l'intégrale

$$I(f) = \int_0^{2\pi} x e^x \cos 2x \, dx,$$

on a utilisé deux formules de quadratures pour des valeurs décroissantes de h . Les résultats sont reportés dans le tableau 1. Complétez ce tableau en calculant le FdR de chacune des 2 méthodes, et déduisez-en leur ordre de précision q . L'une de ces méthodes est la quadrature du rectangle, l'autre la quadrature de Simpson : pouvez-vous les identifier ?

h	Erreur, I_h^A	FdR	Erreur, I_h^B	FdR
$2\pi/10$	1.83×10^{-2}	—	4.66×10^{-4}	—
$2\pi/20$	4.26×10^{-3}		3.05×10^{-5}	
$2\pi/40$	1.05×10^{-3}		1.92×10^{-6}	
$2\pi/80$	2.60×10^{-4}		1.20×10^{-7}	

TAB. 1 – Erreur E_h commises par les formules de quadrature I_h^A et I_h^B .

Exercice 3. (Extrapolation de Richardson - Quadrature de Romberg)

Pour le calcul approché de l'intégrale

$$I = \int_a^b f(x) \, dx,$$

on va développer des solutions numériques de haute précision en se basant sur le principe de l'extrapolation de Richardson.

L'extrapolation de Richardson est une technique générale permettant d'obtenir une estimation numérique précise à partir de deux (ou plusieurs) solutions données par une méthode basique, de faible précision. La méthode de base que nous utiliserons est la quadrature du trapèze, notée $I_{1,h}$. L'ingrédient essentiel de cette extrapolation repose sur la connaissance de la forme de l'erreur de troncature de la méthode de base, ici :

$$I = I_{1,h} + c_1 h^2 + c_2 h^4 + c_3 h^6 + \dots \quad (1)$$

Il est alors possible de combiner deux calculs obtenus avec la quadrature du trapèze pour éliminer le terme en $\mathcal{O}(h^2)$ et obtenir une estimation d'ordre supérieur.

- 1) Ecrivez l'erreur de troncature de la formule du trapèze pour le pas $h/2$. Effectuez une combinaison linéaire de cette équation avec (1) pour éliminer le terme d'erreur du second-ordre. Montrez qu'on définit alors une nouvelle formule de quadrature, notée $I_{2,h}$, de la forme :

$$I_{2,h} = \frac{4I_{1,h/2} - I_{1,h}}{3}, \quad (2)$$

dont l'erreur de troncature s'écrit

$$I = I_{2,h} - \frac{c_2}{4}h^4 - \frac{5}{16}c_3h^6 + \dots \quad (3)$$

Calculez l'expression analytique de cette nouvelle formule de quadrature, et montrez qu'on retrouve la quadrature de Simpson. En déduire que la formule de quadrature de Simpson est précise au quatrième ordre.

- 2) Appliquez de nouveau le principe de l'itération de Richardson à $I_{2,h}$ pour obtenir la formule du sixième ordre :

$$I_{3,h} = \frac{16I_{2,h/2} - I_{2,h}}{15}. \quad (4)$$

Le principe d'extrapolation peut être répété de nouveau pour éliminer le terme en $\mathcal{O}(h^6)$, etc ... Cet algorithme est alors connu sous le nom d'intégration de Romberg.

- 3) Application : utilisez la quadrature de Romberg pour calculer l'intégrale

$$I = \int_1^\pi \frac{\sin x}{2x^3} dx = 0.198557\dots \quad (5)$$

La première colonne du tableau 2 donne les résultats de l'intégrale du trapèze sur un nombre croissants d'intervalles n . Complétez ce tableau en utilisant les extrapolations développées dans les questions précédentes. Vérifiez la convergence des approximations vers le résultat exact.

n			
2	$I_{1,h} = 0.278173$		
4	$I_{1,h/2} = 0.220713$	$I_{2,h}$	
8	$I_{1,h/4} = 0.204304$	$I_{2,h/2}$	$I_{3,h}$
16	$I_{1,h/8} = 0.200009$	$I_{2,h/4}$	$I_{3,h/2}$

TAB. 2 – Résultats de l'approximation de l'intégrale (5).

- 4) Ecrivez la procédure *Maple* `trapeze=proc(a,b,n,f)` qui calcule l'intégrale $\int_a^b f(x)dx$ par la formule du trapèze composite avec n sous intervalles de taille uniforme¹. Utilisez-la avec Maple pour retrouver les résultats de la question précédente.

¹On rappelle que, d'après le cours, cette formule s'écrit :

$$I_h = \frac{h}{2} \sum_{k=0}^{n-1} [f(x_k) + f(x_{k+1})],$$

avec $h = (b - a)/n$ et $x_k = a + kh$.

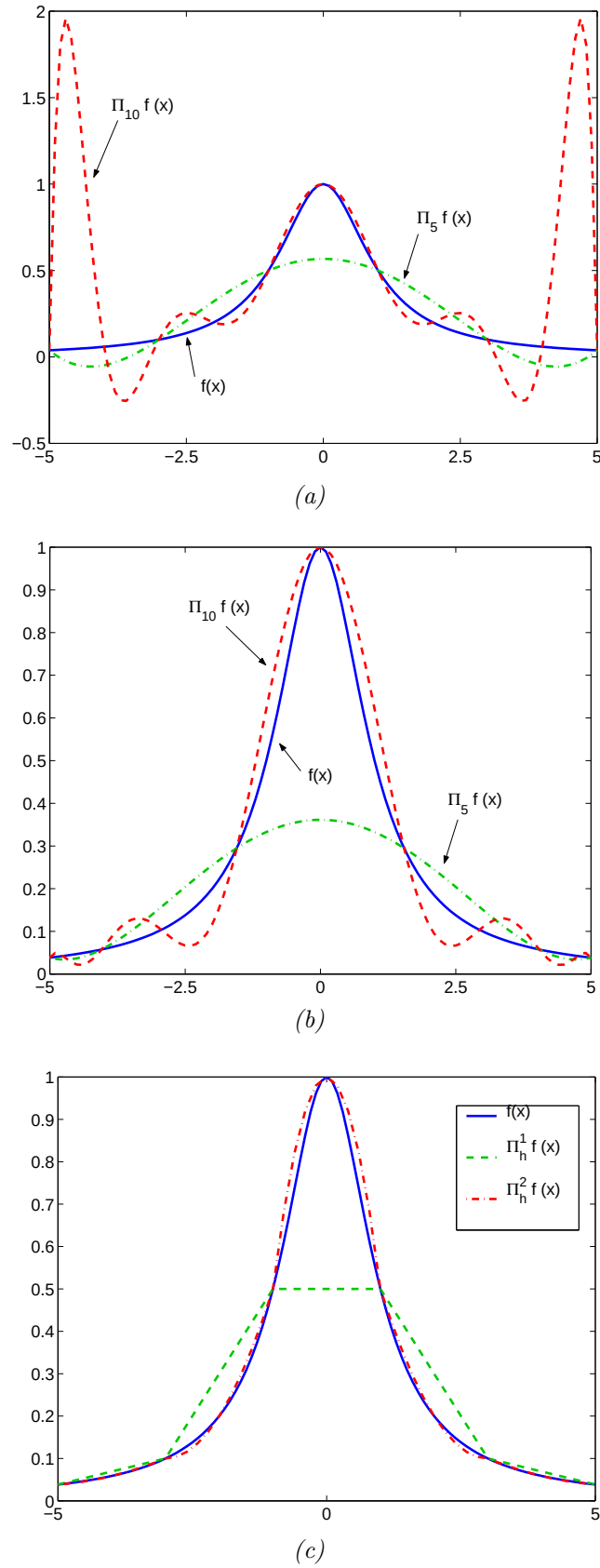


FIG. 1 – Interpolation de la fonction de Runge $f(x) = \frac{1}{1+20x^2}$ dans $\Omega =]-5, 5[$. Dans (a), l'interpolant global $\pi_n f$ qui utilise $n+1$ nœuds équidistants présente des oscillations qui augmentent avec le degré n du polynôme. Dans (b), les oscillations sont fortement atténuées lorsque les points de Tchebychev ($x_i = -5 \cos i\pi/n$, $i = 0, \dots, n$) sont utilisés. Dans (c), les oscillations ont disparues en utilisant l'interpolant $\pi_n^h f$ sur 5 sous-intervalles de longueur $h = 2$.

ANALYSE NUMERIQUE EX1 du TD1

Les Graphes

>

> **L0:=x->x^2/2-x/2;**

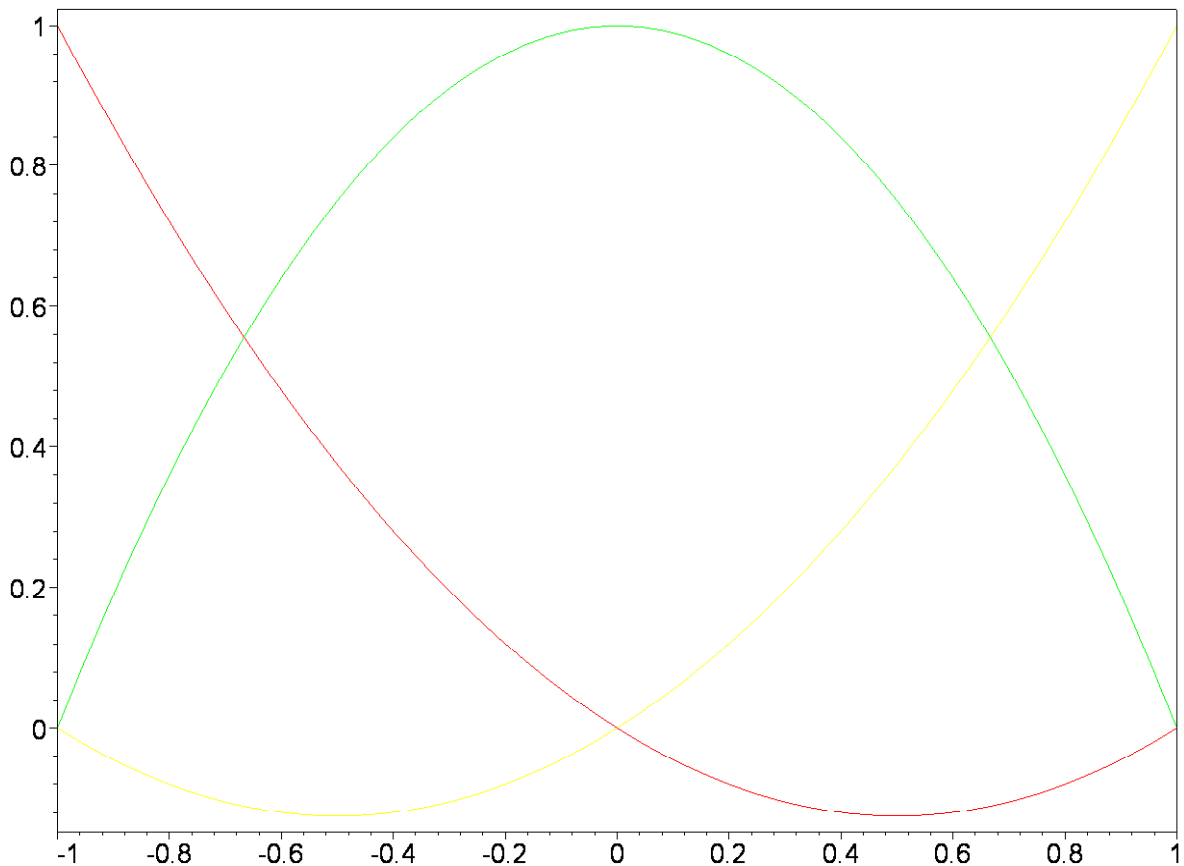
> **L1:=x->1-x**2;**

$$L1 := x \rightarrow 1 - x^2$$

> **L2:=x->x^2/2+x/2;**

$$L2 := x \rightarrow \frac{1}{2}x^2 + \frac{1}{2}x$$

> **plot({L0,L1,L2}, -1..1);**



Elements de la base de Lagrange (question 1)

> **a0:=exp(-1.);a1:=exp(0.);a2:=exp(1.);**

$$a0 := .3678794412$$

$$a1 := 1.$$

$$a2 := 2.718281828$$

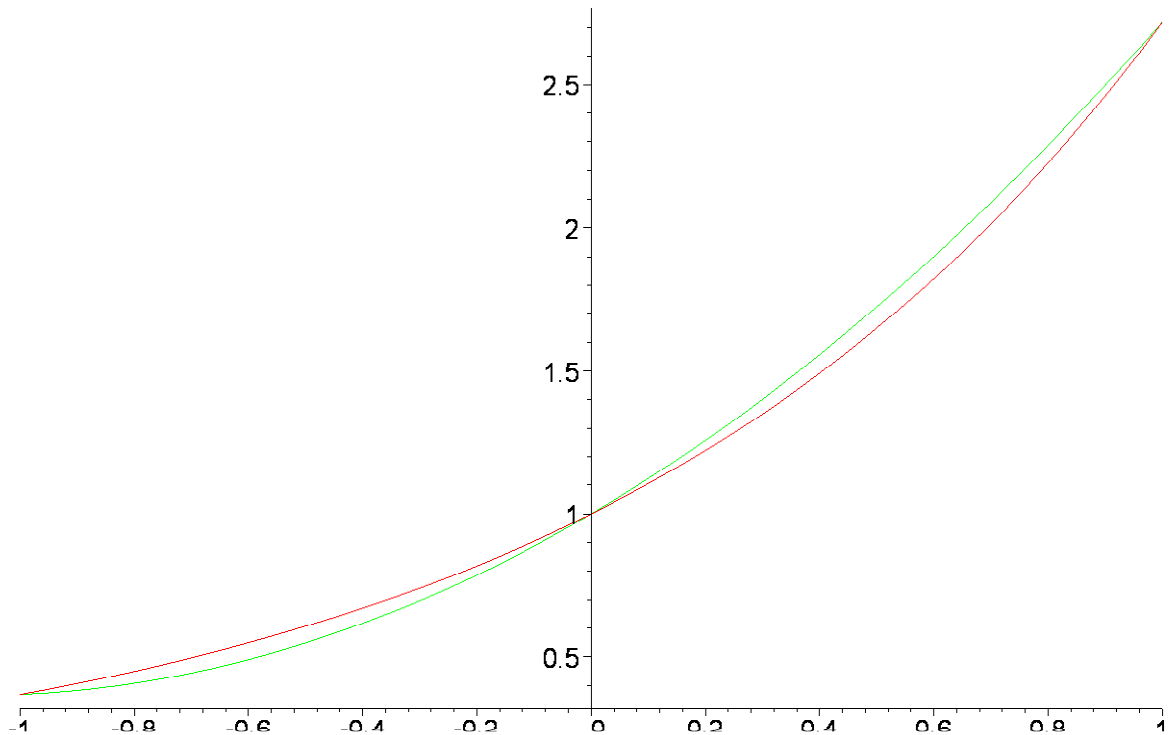
> **pol:=x->a0*L0(x)+a1*L1(x)+a2*L2(x);**

$$pol := x \rightarrow a0 L0(x) + a1 L1(x) + a2 L2(x)$$

> **(pol(0.5)-exp(0.5))/exp(0.5);**

$$.04527720138$$


```
> plot({exp,pol}, -1..1);
```



function exp(x) et son interpolant de Lagrange (Q. 4)

```
> runge:=x->1/(1+100*(x-0.)^2);a0:=runge(-1.);a1:=runge(0.);a2:=runge(1.);
```

$$runge := x \rightarrow \frac{1}{1 + 100 x^2}$$

$$a0 := .009900990099$$

$$a1 := 1.$$

$$a2 := .009900990099$$

```
> pol:=x->a0*L0(x)+a1*L1(x)+a2*L2(x);
```

```
> pol(x);
```

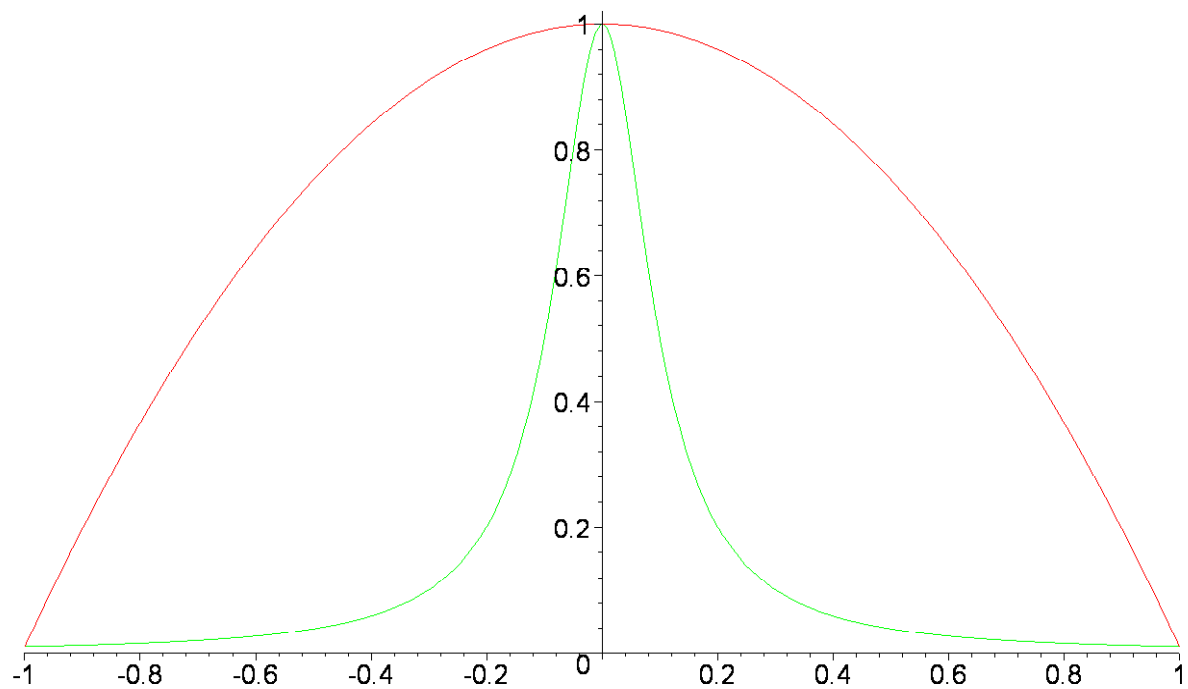
$$pol := x \rightarrow a0 L0(x) + a1 L1(x) + a2 L2(x)$$

$$-.9900990099 x^2 + 1.$$

```
> (pol(0.5)-runge(0.5))/runge(0.5);
```

$$18.56435643$$

```
> plot({runge,pol}, -1..1);
```



Function de Runge et son interpolant de Lagrange (Q. 5)

>

Analyse Numérique - Licence de Mécanique

TD 2 - Equations Non-Linéaires

Exercice 1. (*Quelques itérations du point fixe*)

On veut résoudre le problème

$$\text{Trouver la racine } x_\star \text{ de la fonction } f(x) = \sin 2x - 1 + x, \quad (1)$$

par diverses méthodes itératives. Le graphe de f est représenté sur la figure 1 : la racine $x_\star \simeq 0.35$ est unique.

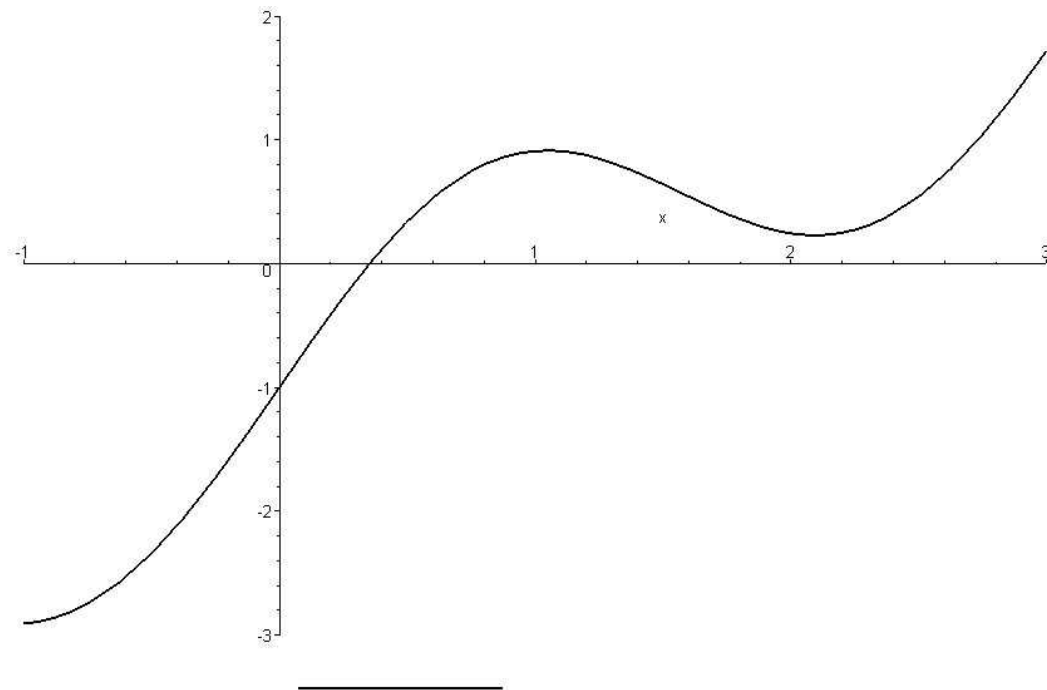


Figure 1: Graphe de la fonction $f(x) = \sin 2x - 1 + x$.

- 1) La méthode de bisection va-t-elle converger ? Proposez un intervalle I_0 pour l'initialiser.
- 2) Pour $i = 1$ et 2 , on considère les méthodes itératives

$$\begin{cases} x_0 \in \mathbb{R}, \\ x_{k+1} = \phi_i(x_k), \end{cases}$$

avec $\phi_1(x) = 1 - \sin(2x)$, et $\phi_2(x) = \frac{1}{2} \arcsin(1 - x)$ pour $x \in]0, 2[$.

- a) A quelle classe de méthodes appartiennent-elles ? Sont-elles consistantes avec le problème (1) ?
 - b) En utilisant la figure 2, construisez géométriquement les 4 premiers itérés de chacune des méthodes à partir de $x_0 = 0$. Que remarquez-vous ?
 - c) Etudiez la stabilité de ces méthodes (conseil : utilisez la figure 3). Pouvez-vous prédire les comportements observés dans la question précédente ?
- 3) Ecrivez la méthode de Newton pour résoudre (1). Pour ce problème, quel est l'ordre de convergence de cette méthode ? En partant de $x_0 = 0$, construisez géométriquement les premiers itérés de la méthode de Newton. Même question pour $x_0 = 0.9$.

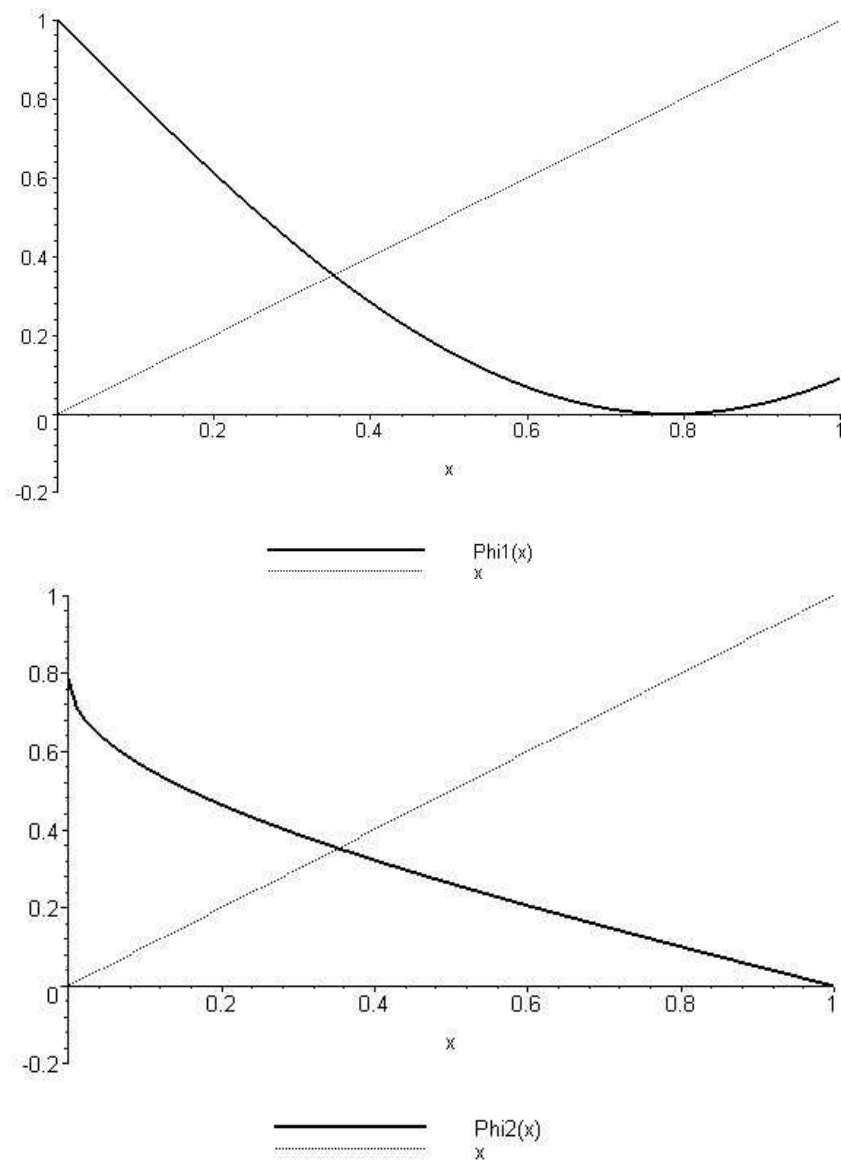


Figure 2: Graphe des fonctions d'itération $\phi_1(x)$ et $\phi_2(x)$.

- 4) Dans l'algorithme de Newton, la partie la plus coûteuse en temps de calcul est l'évaluation de $f'(x_k)$ à chaque itération¹. Pour cela, il peut se révéler judicieux de remplacer $f'(x_k)$ par une approximation bien choisie : on obtient alors des méthodes de Newton modifiées ou quasi-Newton. Le choix

$$\forall k \geq 1, \quad f'(x_k) \cong f'(x_0),$$

correspond à la *méthode de la corde*, ainsi appelée car elle utilise une approximation linéaire de $f(x_{k+1})$ avec toujours la même pente $f'(x_0)$.

- Ecrivez l'algorithme de la corde. A quelle classe de méthodes appartient-elle ? Cet algorithme est-il consistant avec le problème (1) ?
- Etudiez la stabilité de cet algorithme. En déduire que sa convergence est au plus d'ordre 1. Que se passe-t-il pour $x_0 = 0$ et $x_0 = 0.9$? Construisez les premiers itérés pour ces conditions initiales.

¹C'est encore plus vrai pour des systèmes d'équations non-linéaires dans $\mathbb{R}^n$.

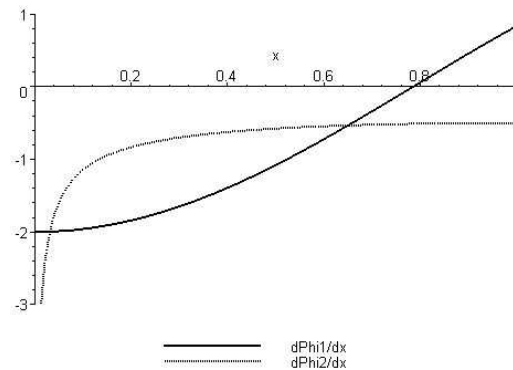


Figure 3: Graphe de la dérivée première des fonctions d'itération.

Exercice 2. (Performance des méthode de bisection et de Newton)

On cherche le(s) zéro(s) de $f(x) = \ln(x+1) - 2$ sur l'intervalle $I_0 = [4, 12]$.

- 1) Montrer que cette fonction admet un zéro x_* dans I_0 , et qu'il est unique.
- 2) En utilisant la méthode de la bisection dans I_0 , estimer le nombre minimal d'itérations nécessaire pour calculer x_* avec la tolérance $\epsilon = 10^{-1}$ et 10^{-10} .
- 3) On considère la méthode du point fixe définie par sa fonction d'itération :

$$\phi(x) = x - (x+1)(\ln(x+1) - 2). \quad (1)$$

- a) Vérifiez la consistance et la stabilité de cette méthode. La reconnaissez-vous ?
- b) Trouvez une constante $C > 0$ telle que

$$|x_{k+1} - x_*| \leq C|x_k - x_*|^2.$$

- c) Supposons que la méthode (1) soit initiée par $x_0 \in I_0$ tel que $|x_0 - x_*| \leq 10^{-1}$. En utilisant la question précédente, trouvez le nombre minimal d'itérations nécessaire pour obtenir une erreur inférieure à $\epsilon = 10^{-10}$.
- 4) On applique maintenant la méthode de la bisection pour trouver une approximation $\tilde{x}$ de la racine, telle que $|\tilde{x} - x_*| \leq 10^{-1}$, et ensuite la méthode du point fixe (1) à partir de la condition initiale $x_0 = \tilde{x}$. Estimer, dans ce cas, le nombre minimal d'itérations nécessaire pour obtenir une erreur inférieure à $\epsilon = 10^{-10}$.

Exercice 3. (Méthode de Newton modifiée pour les racines doubles)

Soit $f(x)$ une fonction admettant une racine double x_* .

- 1) En écrivant f sous la forme

$$f(x) = (x - x_*)^2 g(x), \quad g(x_*) \neq 0,$$

vérifiez que la convergence de la méthode de Newton est seulement linéaire.

- 2) On considère maintenant la méthode de Newton modifiée :

$$x_{k+1} = x_k - 2 \frac{f(x_k)}{f'(x_k)}.$$

Etudiez la consistance et la stabilité de cet algorithme. Montrez qu'il converge au second ordre.

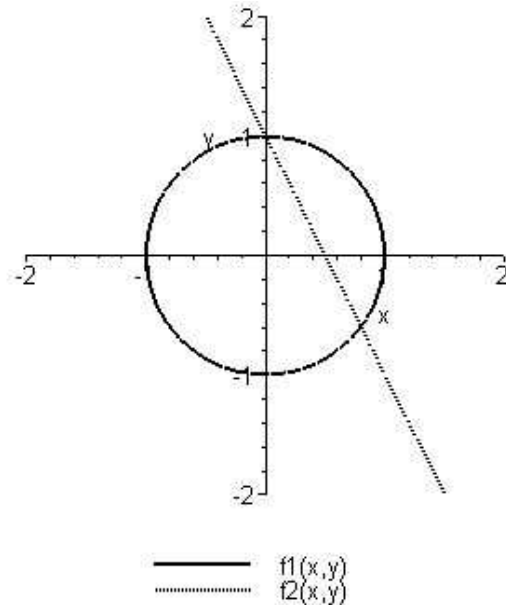


Figure 4: Courbes $f_1(x, y) = 0$ et $f_2(x, y) = 0$. Leurs intersections sont les solutions du système (1).

Exercice 4. (*Système d'équations non linéaires*)

On pose $\mathbf{x} = (x, y)$, et on considère le système non-linéaire dans $\mathbb{R}^2$:

$$\mathbf{f}(\mathbf{x}) = 0, \quad \text{avec} \quad \mathbf{f}(\mathbf{x}) = \begin{pmatrix} f_1(x, y) = x^2 + y^2 - 1 \\ f_2(x, y) = 2x + y - 1 \end{pmatrix}, \quad (1)$$

dont les solutions sont $\mathbf{x}_1^* = (0, 1)$ et $\mathbf{x}_2^* = (4/5, -3/5)$, cf. Fig. 4.

- 1) Pour résoudre (1), on considère la méthode du point fixe définie par sa fonction d'itération :

$$\phi(\mathbf{x}) = \begin{pmatrix} \phi_1(x, y) = \frac{1-y}{2} \\ \phi_2(x, y) = \sqrt{1-x^2} \end{pmatrix}. \quad (2)$$

Vérifiez la consistance et stabilité de cet algorithme lorsqu'on l'applique au calcul de $\mathbf{x}_1^*$. Pouvez-vous l'utiliser pour calculer $\mathbf{x}_2^*$?

- 2) Proposez une modification de $\phi(\mathbf{x})$ permettant de calculer $\mathbf{x}_2^*$. Montrez qu'elle est convergente pour $\mathbf{x}_0$ suffisamment proche de cette racine.
- 3) Décrivez la méthode de Newton-corde pour résoudre le système (1) en partant de $\mathbf{x}_0 = (1/2, 0)$.

Exercice 1. (Stabilité des schémas d'Euler)

On veut résoudre numériquement le problème de Cauchy

$$\begin{cases} \dot{y}(t) + 0.5y(t) = 0, & t \in [0, 20], \\ y(0) = 1, \end{cases}$$

avec les méthodes d'Euler.

- 1) Énoncez la condition de stabilité du schéma d'Euler progressif (EP1). Pour les pas de temps $h = 4.2, 3$, et 1 , établissez la formule de récurrence donnant la solution numérique à l'instant t_n en fonction de la condition initiale. Pour ces valeurs de h , représentez la solution numérique obtenue (complétez la figure 1). Commentez ces résultats.

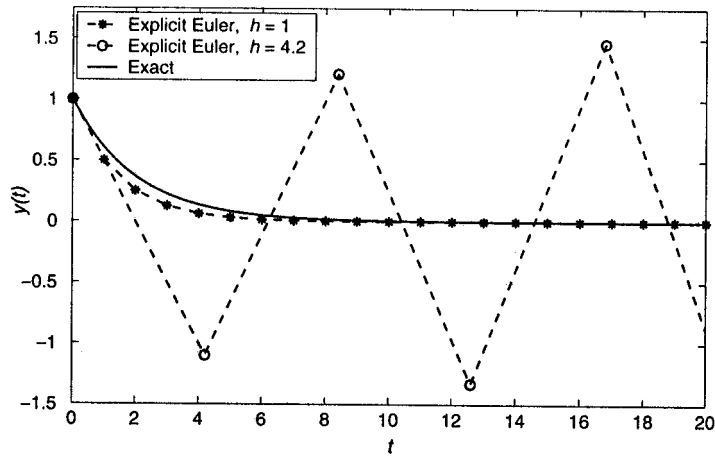


Figure 1: Solution numérique obtenue avec le schéma EP1.

- 2) Même question pour le schéma d'Euler rétrograde (ER1).

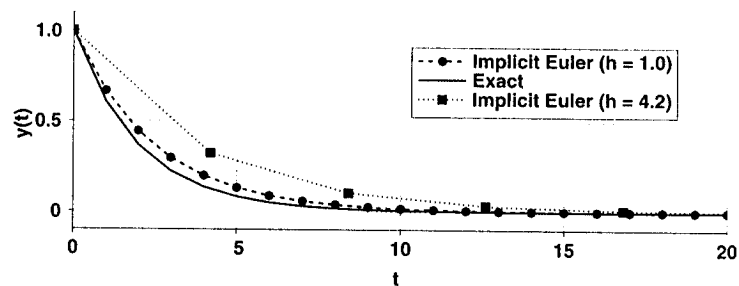


Figure 2: Solution numérique obtenue avec le schéma ER1.

Exercice 2. (*Résolution d'une E.D.O. non-linéaire*)

On considère le problème

$$\begin{cases} \dot{y}(t) = -t y^2(t), & t \in [0, 1], \\ y(0) = 2, \end{cases}$$

dont la solution exacte est $y(t) = 2/(1+t^2)$.

- 1) Cette E.D.O. est-elle linéaire ? Pour le schéma EP1, établir la relation de récurrence donnant y^{n+1} en fonction de y^n . Pour $h = 1/10$, calculez la solution aux temps t_1 et t_2 .
- 2) Déterminez une condition nécessaire de stabilité pour EP1. Vérifiez que la méthode est stable pour le pas de temps utilisé dans la question précédente.
- 3) Appliquez ER1 à la résolution de cette EDO. Montrez que ce schéma se met sous la forme de récurrence :

$$G(t_{n+1}, y^{n+1}) = H(t_n, y^n),$$

et explicitez les fonctions G et H . Pouvez-vous facilement calculer y^{n+1} ? Pourquoi ? Décrivez une méthode itérative permettant de calculer y^{n+1} .

Exercice 3. (*Schéma de Crank-Nicolson*)

On considère le problème de Cauchy :

$$\dot{y}(t) = f(t, y(t)), \quad t \in [0, T] \quad (1a)$$

$$y(0) = y_0. \quad (1b)$$

- 1) En appliquant la quadrature du trapèze entre les instants t_n et $t_{n+1} = t_n + h$, montrez qu'on définit le schéma :

$$\frac{y^{n+1} - y^n}{h} = \frac{1}{2} [f(t_n, y^n) + f(t_{n+1}, y^{n+1})]. \quad (2)$$

qui est connu sous le nom de *schéma de Crank-Nicolson (C-N)*. Cette méthode est-elle explicite ou implicite ?

- 2) En utilisant des développements de Taylor, montrez que le schéma ER1 est précis au premier ordre, tandis que C-N est au second-ordre.
- 3) On applique le schéma C-N au problème test

$$\dot{y}(t) = \lambda y(t), \quad \lambda < 0.$$

Calculez son facteur d'amplification. Vérifiez que la méthode est bien au second-ordre. Etudiez la stabilité de la solution numérique.

- 4) On a utilisé deux schémas X et Y pour intégrer numériquement une E.D.O. de la forme (1). Les résultats sont reportés dans le tableau 1. L'une de ces méthodes est ER1, l'autre C-N : pouvez-vous les identifier ?

h	Schéma X	Schéma Y
1	3.05×10^{-2}	5.22×10^{-3}
0.5	1.46×10^{-2}	1.03×10^{-3}
0.25	7.16×10^{-3}	2.50×10^{-4}
0.125	3.55×10^{-3}	6.19×10^{-5}
0.0625	1.77×10^{-3}	1.54×10^{-5}

Table 1: Erreurs commises à l'instant $t = t^*$ pour différentes valeurs du pas de temps.

Exercice 4. (Système d'EDO linéaires : l'oscillateur mécanique)

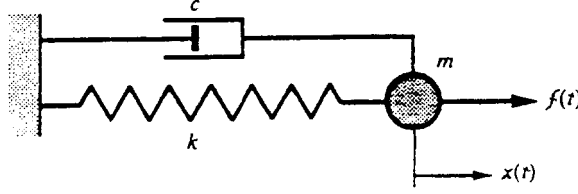


Figure 3: Schéma de l'oscillateur mécanique linéaire.

On considère l'oscillateur mécanique représenté sur la figure 3, formé de :

- Masse m indéformable telle que $x(t)$ représente son déplacement horizontal.
- Ressort sans masse de raideur $k > 0$, fournissant une force élastique proportionnelle et opposée au sens du déplacement, soit : $-kx(t)$.
- Amortisseur fournissant une force de freinage proportionnelle et opposée à la vitesse, soit : $-C\dot{x}(t)$, avec $C \geq 0$ la constante d'amortissement visqueux.
- Une force extérieure $f(t)$ agissant sur la masse m .

Au temps $t = 0$, la masse placée en x_0 est lâchée à la vitesse v_0 .

- 1) Par application du principe fondamental de la dynamique, montrez que le déplacement de la masse est régi par l'EDO du second-ordre :

$$\ddot{x}(t) + 2\gamma \dot{x}(t) + \omega_*^2 x(t) = \frac{f(t)}{m}, \quad (1)$$

avec $\gamma = C/(2m)$ le coefficient d'amortissement et $\omega_* = \sqrt{k/m}$ la pulsation propre de l'oscillateur.

- 2) Calculez l'énergie mécanique $E(t)$ du système. Dans le cas où le système n'est pas forcé (*i.e.* $f(t) \equiv 0$), utilisez l'EDO (1) pour montrer l'identité :

$$\frac{dE}{dt} = -C [\dot{x}(t)]^2.$$

En déduire le comportement du système lorsqu'il est freiné ou non freiné.

Dans la suite de l'exercice, les coefficients de raideur et d'amortissement sont choisis tels que $\gamma > \omega_* > 0$.

- 3) Formulez l'EDO (1) sous la forme du problème de Cauchy :

$$\begin{cases} \dot{X}(t) = \mathcal{A}X(t) + F(t), \\ X(0) = X_0. \end{cases} \quad (2)$$

Appliquez le schéma EP1 à la résolution de ce système.

- 4) Cette question concerne la stabilité du schéma de la question précédente.
 - a) A partir de la diagonalisation $\mathcal{A} = \mathcal{S}D\mathcal{S}^{-1}$, rappelez les étapes principales de l'étude de stabilité du schéma.
 - b) Calculez les valeurs propres λ_1, λ_2 de $\mathcal{A}$ en fonction de γ et ω_* , et déterminez leur signe.
 - c) En déduire la condition de stabilité du schéma. Application numérique : déterminez cette condition dans le cas $\gamma = 5, \omega_* = 4$.

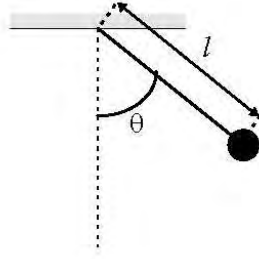


Figure 4: Le pendule oscillant. La position de la masse est repérée par l'angle θ . L'origine $\theta = 0$ est fixée à la position de repos de la masse.

- 5) Appliquez les schémas ER1 et CN à la résolution de (2) et énoncez leur condition de stabilité.

Exercice 5. (*Pendule oscillant - Système d'équations non-linéaires - Diagrammes de stabilité - Schéma RK4*)

Considérons le pendule constitué d'une masse m attachée à une tige rigide de longueur l , qui se meut dans un plan vertical (cf. Fig. 4).

- 1) A partir de la loi de conservation de l'énergie mécanique, montrez que le mouvement de la masse est régi par l'EDO :

$$\ddot{\theta} = -\omega_{\star}^2 \sin \theta, \quad (1)$$

avec $\omega_{\star} = \sqrt{g/l}$, où g est la constante de gravitation. Mettez cette équation sous la forme du problème de Cauchy

$$\dot{X}(t) = F(t, X(t)), \quad (2)$$

avec

$$X(t) = \begin{pmatrix} x_1(t) \\ x_2(t) \end{pmatrix}, \quad F = \begin{pmatrix} f_1(t, x_1, x_2) \\ f_2(t, x_1, x_2) \end{pmatrix}.$$

Par consistance avec la suite de l'énoncé, on aura défini x_1 comme l'angle du pendule et x_2 sa vitesse angulaire.

- 2) On veut résoudre le système non-linéaire (2) par une méthode numérique. D'après l'exercice précédent, rappelez les étapes principales de l'analyse de stabilité appliquée à (2).
- 3) Calculez la matrice jacobienne de F , soit :

$$\mathcal{A} = \frac{\partial F}{\partial X} = \left(\frac{\partial f_i}{\partial x_j} \right)_{i,j}.$$

Montrez que ses valeurs propres s'écrivent :

$$\lambda = \pm \sqrt{-\cos x_1} \omega_{\star}.$$

- 4) Dédurre de la question précédente que pour des oscillations de faible amplitude ($\theta \simeq 0$), l'équation du pendule s'écrit

$$\ddot{\theta} = -\omega_{\star}^2 \theta, \quad (3)$$

et calculez ses valeurs propres. A quel système physique correspond cette EDO ?

- 5) On souhaite utiliser les schémas classiques (EP1, ...) pour résoudre (1) ou (3). Pouvez-vous appliquer les conditions de stabilité usuelles ?

Les systèmes d'EDO à valeurs propres complexes sont souvent qualifiés de systèmes oscillants. Leur approximation numérique motive l'étude de la stabilité des schémas pour l'EDO (scalaire) modèle

$$\dot{y}(t) = \lambda y(t), \quad \lambda = \lambda_R + i \lambda_I \in \mathbb{C}, \quad \lambda_R \leq 0. \quad (4)$$

6) Montrez que le schéma EP1 appliqué à (4) est stable si le pas de temps h vérifie

$$\sqrt{(1 + \lambda_R h)^2 + (\lambda_I h)^2} < 1. \quad (5)$$

Vérifiez que pour $\lambda \in \mathbb{R}^{-,*}$, on retrouve la condition de stabilité usuelle. Qu'en est-il dans le cas purement complexe ? Conseillerez-vous d'utiliser EP1 pour résoudre les systèmes oscillants ?

7) Effectuez l'analyse de stabilité des schémas ER1 et C-N.

Peut-on utiliser des schémas explicites pour résoudre des systèmes oscillants ? Tournons-nous maintenant vers les schémas RK2- α (cf. cours), ainsi que le schéma RK4 défini par

$$K_1 = hf(t_n, y^n), \quad (6a)$$

$$K_2 = hf(t_n + \frac{h}{2}, y^n + \frac{1}{2}K_1), \quad (6b)$$

$$K_3 = hf(t_n + \frac{h}{2}, y^n + \frac{1}{2}K_2), \quad (6c)$$

$$K_4 = hf(t_n + h, y^n + K_3), \quad (6d)$$

$$y^{n+1} = y^n + \frac{1}{6}(K_1 + 2K_2 + 2K_3 + K_4). \quad (6e)$$

8) A quelle classe de méthodes ce schéma appartient-il ? Est-il explicite ou implicite ?

9) Calculez son facteur d'amplification σ et montrez que le schéma RK4 est précis au quatrième ordre.

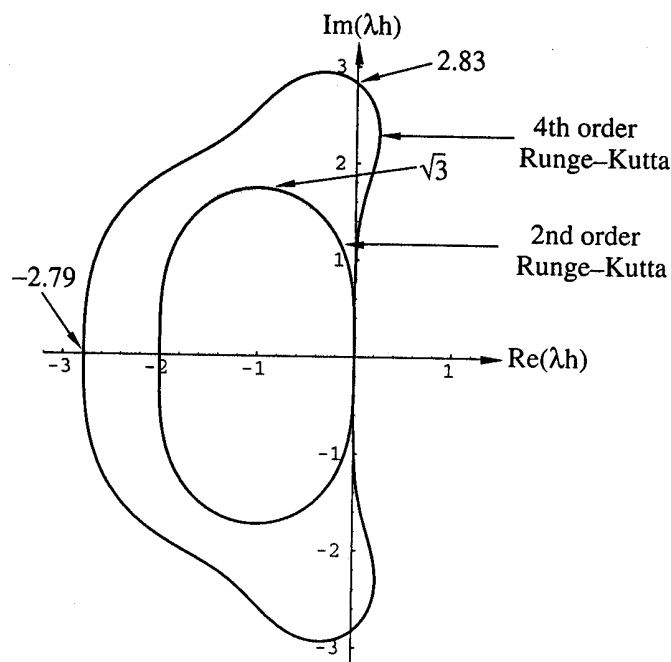


Figure 5: Diagramme de stabilité des schémas de Runge-Kutta. Le schéma est stable pour les valeurs de λh situées à l'intérieur de la courbe fermée. On peut remarquer que la stabilité des méthodes R-K croît avec leur ordre de précision.

Afin d'évaluer et comparer la stabilité des schémas numériques, il est très courant de représenter leur région de stabilité dans le plan complexe $(\lambda_R h, \lambda_I h)$, c'est-à-dire l'ensemble des complexes λh qui vérifient

$$|\sigma(\lambda h)| < 1.$$

Les zones de stabilité de RK2- α et de RK4 sont représentées¹ sur la figure 5.

- 10) Complétez la figure 5 en traçant la région de stabilité de EP1 (aidez-vous de la question 6).
- 11) A, l'aide des diagrammes de la figure 5, déterminez graphiquement les conditions de stabilité de RK2- α et RK4 dans les cas purement réel et complexe.
- 12) Confirmez que RK2- α est instable pour les problèmes oscillants en calculant $|\sigma(i\lambda_I)|$.
- 13) Quel schéma conseillez vous pour résoudre les systèmes (1) et (3) ? Pour chacun des systèmes, déterminez ses conditions de stabilité lorsque $g = 9.81 \text{ m/sec}^2$ et $l = 0.6 \text{ m}$.

Schémas	Type	Ordre q	$\sigma(\lambda h)$	Stabilité pour $\dot{y} = \lambda y$	
				$\lambda \in \mathbb{R}^{-,*}$	$\lambda \in i\mathbb{R}^*$
EP1	Expl.	1	$1 + \lambda h$	$h < \frac{2}{-\lambda}$	
ER1	Impl.	1	$\frac{1}{1 - \lambda h}$	Oui	Oui
CN					
RK2- α ($\alpha = \frac{1}{2}$: Préd./Corr., $\alpha = 1$: Heun)	Expl.	2	$1 + \lambda h + \frac{(\lambda h)^2}{2!}$	$h < \frac{2}{-\lambda}$	
RK4					

Table 2: Principales propriétés des schémas numériques pour EDO.

Exercice 6. Utilisez les exercices précédents pour compléter le tableau 2.

¹Ces diagrammes sont réalisés en traçant la frontière de ces courbes, correspondant aux points vérifiant

$$\sigma(\lambda h) - e^{i\varphi} = 0,$$

en faisant varier φ dans $[0, 2\pi[$. Les diagrammes de stabilité compliqués, comme celui de RK4, nécessitent l'utilisation de méthodes numériques pour extraire les racines d'une fonction à valeur complexe.

Analyse Numérique - Licence de Mécanique

TD 4 - Système d'Equations Linéaires

Exercice 1. (*Algorithme de Thomas pour systèmes tridiagonaux*)

On considère le système linéaire $\mathcal{A}X = F$ d'ordre n , où $\mathcal{A}$ est la matrice tridiagonale définie par

$$\mathcal{A} = \begin{bmatrix} b_1 & c_1 & & & \\ a_2 & b_2 & c_2 & & \\ & \ddots & \ddots & \ddots & \\ & & a_{n-1} & b_{n-1} & c_{n-1} \\ & & & a_n & b_n \end{bmatrix}.$$

Le but de cet exercice est de construire la méthode de résolution directe connue sous le nom d'*algorithme de Thomas*, qui est un cas particulier de la méthode d'élimination de Gauss pour systèmes tridiagonaux.

- 1) On cherche la décomposition $\mathcal{A} = \mathcal{L}\mathcal{U}$ avec $\mathcal{L}$ et $\mathcal{U}$ des matrices bidiagonales de la forme :

$$\mathcal{L} = \begin{bmatrix} 1 & 0 & & & \\ \alpha_2 & 1 & & & \\ & \ddots & \ddots & \ddots & \\ & & \alpha_{n-1} & 1 & 0 \\ & & & \alpha_n & 1 \end{bmatrix}, \quad \mathcal{U} = \begin{bmatrix} \beta_1 & c_1 & & & \\ 0 & \beta_2 & c_2 & & \\ & \ddots & \ddots & \ddots & \\ & & 0 & \beta_{n-1} & c_{n-1} \\ & & & 0 & \beta_n \end{bmatrix},$$

où les coefficients $\{\alpha_i, i = 2, \dots, n\}$ et $\{\beta_i, i = 1, \dots, n\}$ sont à déterminer.

- a) En identifiant les éléments de la matrice $\mathcal{A}$ avec ceux du produit $\mathcal{L}\mathcal{U}$, montrez que les α_i et β_i s'écrivent : $\beta_1 = b_1$ et

$$\alpha_i = \frac{a_i}{\beta_{i-1}}, \quad \beta_i = b_i - \alpha_i c_{i-1}, \quad \text{pour } i = 2, \dots, n.$$

- b) En posant $X = (x_i, i = 1, \dots, n)^T$ et $F = (f_i, i = 1, \dots, n)^T$, résolvez le système par la méthode LU . Plus précisément, on vous demande d'établir la formule de récurrence progressive pour résoudre le système bidiagonal inférieur, puis la formule de récurrence rétrograde pour résoudre le système bidiagonal supérieur.

- 2) On appelle *algorithme de Thomas* les formules de récurrence de la question précédente : elles sont destinées à être programmées sur un ordinateur. Vérifiez qu'elles sont correctes en les appliquant au système d'ordre 3 défini par

$$\mathcal{A} = \begin{bmatrix} 2 & -1 & 0 \\ -1 & 2 & -1 \\ 0 & -1 & 1 \end{bmatrix}, \quad F = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix},$$

dont la solution est $X = (1, 1, 1)^T$.

- 3) Donnez le nombre d'opérations algébriques élémentaires (flops) de l'algorithme de Thomas. Comparez sa complexité algorithmique avec celle d'un algorithme de résolution directe pour système quelconque (e.g., la méthode de Gauss).

Exercice 2. On considère le système linéaire $\mathcal{A}X = B$ avec

$$\mathcal{A} = \begin{bmatrix} 1 & 2 \\ 2 & 5 \end{bmatrix}, \quad B = \begin{pmatrix} 0 \\ -1 \end{pmatrix}.$$

- 1) Montrez que le système admet une solution unique, et qu'il est symétrique défini positif (SPD). Que pouvez-vous dire sur la convergence des méthodes de Jacobi et de Gauss-Seidel (GS) ?
- 2) Construisez la matrice d'itération $\mathcal{B}_J$ de la méthode de Jacobi, et étudiez la convergence de cette méthode.
- 3) Ecrivez la méthode GS pour résoudre ce système et calculez son facteur de réduction ρ_{GS} .
- 4) Comparez la vitesse de convergence des deux méthodes. En utilisant la relation

$$\|E_k\|_2 \leq \rho \|E_{k-1}\|_2,$$

où E_k désigne l'erreur à la $k^{\text{ième}}$ itération, établissez pour chaque méthode le nombre minimum d'itérations nécessaire pour faire décroître l'erreur initiale $\|E_0\|_2$ d'un facteur de 1000.

Exercice 3. (Méthode de Richardson à pas stationnaire)

On considère le système linéaire d'ordre n :

$$\mathcal{A}X = F. \quad (1)$$

On note $R_k = F - \mathcal{A}X_k$ le résidu de cette équation à la $k^{\text{ième}}$ itération, et on définit la méthode de Richardson

$$\mathcal{P}X_{k+1} = \mathcal{P}X_k + \alpha R_k,$$

où $\mathcal{P}$ une matrice inversible, et α un paramètre d'accélération réel à déterminer de façon judicieuse.

- 1) Etudiez la consistance de la méthode suivant la valeur de α .
- 2) Ecrivez la méthode de Richardson sous la forme

$$X_{k+1} = \mathcal{B}_\alpha X_k + F_\alpha,$$

et explicitez $\mathcal{B}_\alpha$ et F_α .

- 3) On note $\mathcal{A} = \mathcal{D} - \mathcal{L} - \mathcal{U}$ le "splitting" habituel de la matrice $\mathcal{A}$. Pour $\alpha = 1$ et $\mathcal{P} = \mathcal{D} - \mathcal{L}$, la méthode de Richardson correspond-elle à une méthode connue ?
- 4) On se place dans le cas $n = 2$, et on suppose que les valeurs propres de $\mathcal{P}^{-1}\mathcal{A}$, notées λ_i , sont strictement positives, soit : $\lambda_1 > \lambda_2 > 0$.

- a) Calculez les valeurs propres de $\mathcal{B}_\alpha$ et montrez que son rayon spectral est strictement inférieur à 1 si

$$\rho_i(\alpha) = |1 - \alpha\lambda_i| < 1, \quad \text{pour } i = 1, 2.$$

- b) En déduire que la méthode de Richardson converge si

$$0 < \alpha < 2/\lambda_1.$$

- c) Tracez les $\rho_i(\alpha)$ sur un même graphe. Montrez graphiquement que le facteur de réduction est minimal lorsque le paramètre d'accélération est choisi tel que $\alpha = \alpha_{R,\text{opt}}$ avec

$$\alpha_{R,\text{opt}} = \frac{2}{\lambda_1 + \lambda_2}. \quad (2)$$

Calculez le facteur de réduction $\rho_{R,\text{opt}}$ correspondant.

- 5) Application : pour le système (1) avec la matrice

$$\mathcal{A} = \begin{bmatrix} 1 & 1 \\ 1 & 2 \end{bmatrix},$$

on va comparer les performances de la méthode de Gauss-Seidel et de la méthode de Richardson avec $\mathcal{P} = \mathcal{D} - \mathcal{L}$ et le pas $\alpha_{R,\text{opt}}$ donné par (2).

- Calculez les valeurs propres de $\mathcal{P}^{-1}\mathcal{A}$. Déduisez-en le facteur de réduction de chacune des 2 méthodes.
- Pour ce système linéaire, écrivez les deux méthodes et comparez leur coût numérique et leur vitesse de convergence.
- Au vu des résultats de cette partie, justifiez l'assertion : “La méthode de Richardson est un accélérateur de la convergence d’une méthode itérative de base”.

Exercice 4. (Performance des méthodes itératives pour les systèmes mal conditionnés)

Soient a, b, c trois réels. On définit $\mathcal{A} = \text{Tridiag}_N[a, b, c]$ la matrice tridiagonale de taille $N \times N$ telle que :

$$\mathcal{A} = \begin{bmatrix} b & c & & & \\ a & b & c & & \\ & \ddots & \ddots & \ddots & \\ & & a & b & c \\ & & & a & b \end{bmatrix}.$$

On notera D_N son déterminant.

- Montrez que $D_1 = b$, $D_2 = b^2 - ac$, et que pour $N \geq 3$:

$$D_N = b D_{N-1} - ac D_{N-2}.$$

- Utiliser la question précédente pour montrer que :

$$\text{pour } N \geq 1, \quad D_N = \frac{\sin(N+1)\theta}{\sin \theta} r^N, \quad (1)$$

avec $r = \sqrt{ac}$, et θ tel que $2r \cos \theta = b$. (on rappelle les formules : $\sin(x \pm y) = \sin x \cos y \pm \cos x \sin y$, $\cos(x \pm y) = \cos x \cos y \mp \sin x \sin y$, $\sin(x - y) + \sin(x + y) = 2 \sin x \cos y$.)

- Montrez que les valeurs propres de $\mathcal{A}$ sont :

$$\lambda_k = b - 2\sqrt{ac} \cos \frac{k\pi}{N+1}, \quad k = 1, 2, \dots, N.$$

(Suggestion : calculez les racines du déterminant de $\text{Tridiag}_N[a, b - \lambda, c]$ en utilisant (1).)

On considère maintenant la matrice $\mathcal{K} = \text{Tridiag}_N[-1, 2, -1]$, et on pose $h = 1/(N+1)$. On souhaite comparer les performances des méthodes itératives classiques pour résoudre le système linéaire

$$\mathcal{K}X = F, \quad (2)$$

où F est un RHS donné.

- Déduire de la question précédente que les valeurs propres de $\mathcal{K}$ s'écrivent

$$\lambda_k = 2 - 2 \cos k\pi h, \quad \text{pour } k = 1, \dots, N, \quad (3)$$

et que cette matrice est SDP. Que pouvez-vous dire sur la convergence des méthodes de Jacobi, Gauss-Seidel et *S.O.R.* appliquées au système (2) ?

- En remarquant que $\cos N\pi h = -\cos \pi h$, montrez que le nombre de conditionnement de $\mathcal{K}$ s'écrit

$$\chi(\mathcal{K}) = \frac{1 + \cos \pi h}{1 - \cos \pi h}.$$

En déduire que le système (2) est mal conditionné lorsque N est grand.

- 6) On note $\mathcal{D}$ la partie diagonale de $\mathcal{K}$. En calculant les valeurs propres de la matrice d'itération $\mathcal{B}_J = \mathcal{I} - \mathcal{D}^{-1}\mathcal{K}$ de la méthode de Jacobi, montrez que son facteur de réduction s'écrit :

$$\rho_J = \cos \pi h, \quad (4)$$

et que $\rho_J \simeq 1 - (\pi h)^2/2$ lorsque h est petit. Commentez les performances de la méthode de Jacobi pour les systèmes mal conditionnés.

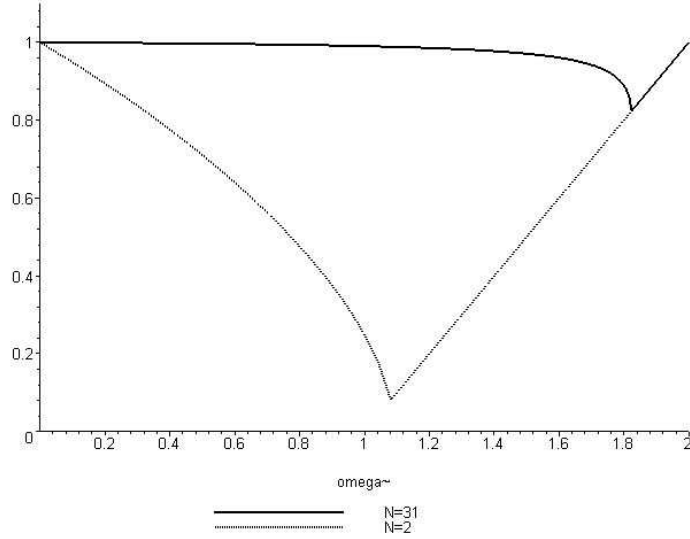


Figure 1: Facteur de réduction de la méthode *S.O.R.* appliqué à la résolution du système (2) pour $N = 2$ ($\rho_J = 1/2$) et $N = 31$ ($\rho_J \simeq 0.995$).

- 7) Un fameux théorème dû à Young (1950) énonce que si les valeurs propres de $\mathcal{B}_J$ sont réelles et que $\rho_J < 1$, alors la méthode *S.O.R.* est convergente pour le paramètre de relaxation $\omega \in]0, 2[$, avec le facteur de réduction

$$\rho_{S.O.R.}(\omega) = \begin{cases} 1 - \omega + \frac{\omega^2 \rho_J^2}{2} + \omega \rho_J \sqrt{1 - \omega + \frac{\omega^2 \rho_J^2}{4}} & \text{pour } 0 < \omega \leq \omega_{\text{opt}}, \\ \omega - 1 & \text{pour } \omega_{\text{opt}} \leq \omega < 2, \end{cases} \quad (5a)$$

tel que

$$\omega_{\text{opt}} = \frac{2}{1 + \sqrt{1 - \rho_J^2}} > 1. \quad (5b)$$

Le graphe de $\rho_{S.O.R.}(\omega)$ en fonction de ω est représenté sur la figure 1. On peut montrer que pour toute valeur de $\rho_J < 1$, le minimum de $\rho_{S.O.R.}(\omega)$ est obtenu en $\omega = \omega_{\text{opt}}$.

- a) Montrez que le facteur de réduction de la méthode de Gauss-Seidel (G.S) s'écrit :

$$\rho_{GS} = \rho_J^2, \quad (6)$$

et que $\rho_{GS} \simeq 1 - (\pi h)^2$ lorsque h est petit. Comparez les performances de G.S. avec celle de Jacobi pour la résolution du système (2).

- b) Pour ce même système, montrez que le paramètre de relaxation et le facteur de réduction optimaux de la méthode *S.O.R.* sont :

$$\omega_{\text{opt}} = \frac{2}{1 + \sin \pi h}, \quad \rho_{S.O.R.}(\omega_{\text{opt}}) = \frac{1 - \sin \pi h}{1 + \sin \pi h}, \quad (7)$$

et que $\rho_{\text{S.O.R}}(\omega_{\text{opt}}) \simeq 1 - 2\pi h$ lorsque h est petit. Comparez avec les méthodes précédentes.

- 8) Pour le système (2) avec $N = 31$, calculez les facteurs de réduction des trois méthodes. En utilisant la relation

$$\|E_k\|_2 \leq \rho \|E_{k-1}\|_2,$$

où E_k désigne l'erreur à la $k^{\text{ième}}$ itération, établissez pour chaque méthode le nombre minimum d'itérations nécessaire pour faire décroître l'erreur initiale $\|E_0\|_2$ d'un facteur de 1000. Au vu de ces résultats, quelle méthode itérative recommandez-vous ?

Analyse Numérique - Licence 3 Mécanique des fluides et Energie

Enoncés des TP - Analyse numérique avec PYTHON

olivier.botella@univ-lorraine.fr

1 Présentation des TP

Les travaux pratiques d'analyse numérique s'effectuent dans le langage de programmation Python et sur l'environnement de développement **Spyder**¹. Ils se déroulent sur 6 séances encadrées de 2h. Pour réaliser ces TP, vous devez télécharger sur la page ARCHE du cours² l'archive `tp_python_analyse_num_l3mfe.zip` et l'extraire sur votre PC. Cette archive contient les programmes Python à utiliser et compléter lors des séances.

1.1 L'environnement de développement intégré (IDE) Spyder

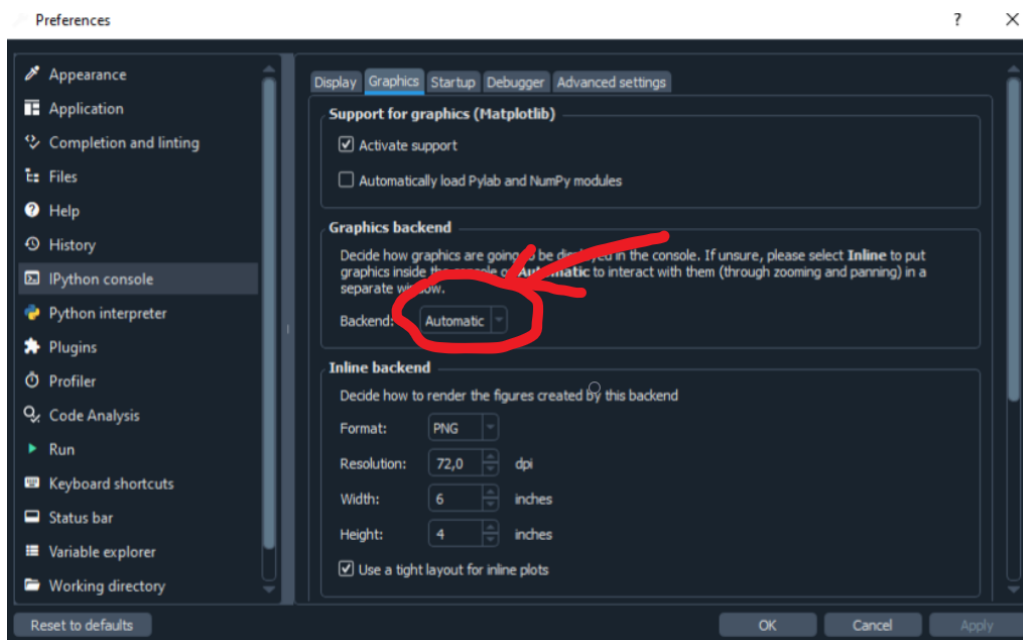


Figure 1: Sélectionner Graphics backend = Automatic dans les préférences de Spyder pour afficher les figures sur des fenêtres graphiques indépendantes. Les préférences de Spyder sont disponibles dans l'onglet Tools (ou Outils).

Le logiciel open-source **Spyder** est une IDE spécialement conçue pour la programmation scientifique en Python, qui inclut les bibliothèques les plus courantes du langage Python, telles que **Matplotlib** et **Numpy**. Son interface graphique et ses fonctionnalités sont très similaires à celles du logiciel **Matlab**, cf. figure 1. Le logiciel **Spyder** est disponible au téléchargement gratuit <https://www.spyder-ide.org/>. Il est aussi inclus dans la suite de logiciels **Anaconda** (<https://www.anaconda.com/download>).

¹Spyder: Scientific PYthon Developpement EnviRonnement, <https://www.spyder-ide.org/>.

²Lapageducoursest~: <https://arche.univ-lorraine.fr/course/view.php?id=6384>. Alternativement, vous pouvez chercher les mots clefs *botella analyse* sur la page d'accueil de Arche: <https://arche.univ-lorraine.fr/>.

Suivant les salles informatiques de la FST, **Spyder** est directement installé ou accessible à partir d'Anaconda.

1.2 Notation des travaux pratiques

La note de TP d'analyse numérique sera basée sur l'évaluation par l'encadrant de votre travail et assiduité lors des séances, et sur les compte-rendus individuels que vous devez rendre à la fin du semestre.

Vous devez rédiger un compte-rendu individuel pour chaque TP :

- Les compte-rendu devront répondre aux questions de l'énoncé, et comporter une dizaine de pages au maximum. Les réponses doivent être rédigées de façon précise et succincte (ne recopiez pas l'énoncé !), et seront illustrées par les programmes **Python** que vous écrirez et les figures que vous avez réalisées.
- Les compte-rendu doivent être rédigés avec le traitement de texte **Word** ou convertis en fichiers **Pdf**. Ne rendez pas de fichiers **Python**. N'oubliez pas de mettre votre nom et l'intitulé du TP sur les compte-rendu.

Les 3 compte-rendu de TP doivent être déposés sur la page ARCHE du cours, avant la date limite que vous fixera l'encadrant.

2 Travaux pratiques d'analyse numérique avec Python

Vous disposez de 6 séances pour réaliser les 4 travaux pratiques ci-dessous. Le plus long TP, celui sur les EDO, peut nécessiter jusqu'à 2 séances.

2.1 Enoncés des TP

Travaux pratiques 0. (*Les aide-mémoire Python*)

Le répertoire `AIDE_MEMOIRE_PYTHON` contient des scripts qui résument l'essentiel des fonctionnalités Python utilisées pour réaliser les TP's d'analyse numérique, et donnent de nombreux exemples :

- Script `aide_memoire_graphisme.py` : pour réaliser des figures avec la bibliothèque `Matplotlib` <https://matplotlib.org/>. Les fonctions tracées sont importées à partir du fichier `fonctions_aide_memoire_g`
- Script `aide_memoire_python.py` : pour les bases du langage Python (boucles, tableaux `Numpy` <https://numpy.org/>, manipulation de fichiers, ...)

Ces scripts incluent des références Web pour approfondir les différentes notions abordées. les références les plus utilisées sont :

- <https://courspython.com/>
- <https://www.w3schools.com/python/default.asp>
- <https://www.w3schools.com/python/numpy/default.asp>

Travaux pratiques 1. (Intégration numérique par méthode des trapèzes)

Le but du TP est de calculer une approximation numérique I_h de l'intégrale (vue en TD du L3 MFE) :

$$I = \int_1^\pi \frac{\sin x}{2x^3} dx \simeq 0.198557 \dots, \quad (1)$$

à l'aide de la méthode des trapèzes, et de comparer la valeur numérique I_h avec la valeur exacte en vérifiant que l'erreur numérique :

$$E_h = |I - I_h|, \quad (2)$$

tends vers zéro lorsque $h \rightarrow 0$. Le répertoire de travail **Spyder** est **TP1_L3_QUADRATURE**, et le script principal est **main_TP_quadrature.py**.

- 1) Ecrivez la fonction **Python** :

```
def trapeze(a, b, f, n):  
    # ...  
    return Ih
```

qui calcule l'intégrale $I_h = \int_a^b f(x)dx$ par la formule du trapèze composite avec n sous-intervalles de taille uniforme h^3 .

- 2) Validez votre fonction **trapeze** en vérifiant qu'elle donne l'intégrale exacte d'une fonction linéaire (polynôme de degré 1).
3) Vous allez maintenant utiliser la fonction **trapeze** pour calculer l'intégrale I donnée par l'équation (1).

- a) A l'aide de la bibliothèque de calcul symbolique **Sympy**⁴, calculez une valeur précise (notée **Iex**) en tapant ces lignes dans votre programme **Python** :

```
import sympy as sp # Calcul symbolique avec sympy  
t = sympy.symbols('t')  
fex = sp.sin(t)/(2*t**3)  
Iex = sp.integrate(fex, (t, 1, sp.pi)).evalf()
```

- b) Vérifiez que $E_h \rightarrow 0$, lorsque $h \rightarrow 0$, c'est-à-dire lorsque n augmente. Pour des valeurs de h suffisamment petites ($n \geq 20$), calculez le facteur de réduction **FdR**⁵ de la méthode du trapèze. Comparez-le avec sa valeur théorique⁶.
c) Pour $n \geq 50$, représentez en échelle logarithmique l'erreur numérique E_n de la méthode du trapèze (voir les exemples données dans le script **aide_memoire_graphisme.py**) et montrez qu'elle est précise au second-ordre en traçant sur le même graphe la fonction $n \mapsto 1/n^2$ comme sur la figure 2. Pour cela, il est conseillé d'écrire dans un fichier **csv** plusieurs réalisations de la fonction **trapeze** sous la forme :

³On rappelle que, d'après le cours, cette formule s'écrit :

$$I_h = \sum_{k=0}^{n-1} h \left(\frac{1}{2} f(x_k) + \frac{1}{2} f(x_{k+1}) \right),$$

avec $h = (b - a)/n$ et $x_k = a + kh$.

⁴La page web de Sympy est <https://www.sympy.org/>

⁵La définition du TD est $\text{FdR} = E_h/E_{h/2}$ où E_h est donné par (2).

⁶D'après le TD, si h est suffisamment petit alors $\text{FdR} \simeq 2^q$ pour une méthode d'ordre de précision q .

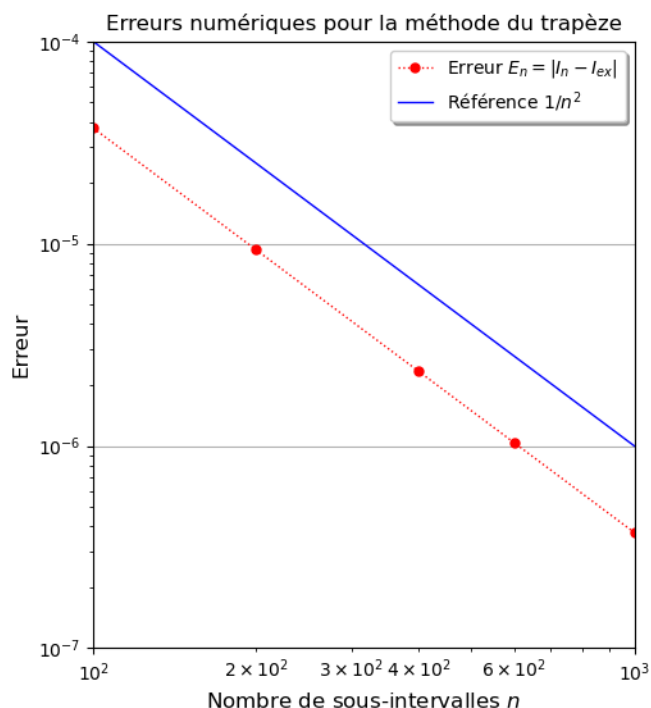


Figure 2: Résultats de la méthode du trapèze pour l'intégrale (1). Pour afficher les figures dans des fenêtres indépendantes de l'IDE, il faut sélectionner Graphics backend = Automatic dans les préférences de Spyder (cf. figure 1).

```
n1 E_n1
n2 E_n2
n3 E_n3
etc ...
```

puis lire les données du fichier dans un script Python, et tracer les données lues (un exemple est donné dans le script `aide_memoire_python.py`).

Travaux pratiques 2. (Résolution itérative d'équations non-linéaires)

Ce TP concerne la résolution de l'équation non-linéaire :

$$f(x) = 0 \quad \text{pour } x \in [a, b], \quad (1)$$

par les méthodes itératives de la bisection et de Newton qui ont été vues en cours et TD. Le répertoire de travail Python est 'TP2_L3_NEWTON_EQT_NON_LIN'.

- 1) Le script `main_TP_eqts_non_lin.py` applique l'algorithme de la bisection (donné dans le fichier `fonctions_TP_quadrature.py`) à la résolution du problème (1) pour la fonction

$$f(x) = \sin(2x) + x - 1,$$

que nous appellerons par la suite *problème test*. Le critère d'arrêt de l'algorithme est basé sur le nombre d'itérations maximal $k_{\max}$ et le résidu de l'équation $r_k = |f(x_k) - 0|$, et s'écrit :

$$k \leq k_{\max} \quad \text{et} \quad r_k < \epsilon, \quad k_{\max} \text{ et } \epsilon \text{ définis par l'utilisateur.}$$

Le code Python applique l'algorithme de la bisection pour l'intervalle de départ $I_0 = [-4, 4]$ et une tolérance $\epsilon = 10^{-10}$. Un exemple de tableau solution $\mathbf{xB}$ est donné sur la figure 3. Comprenez bien ce programme, car les algorithmes itératifs que vous programmerez par la suite suivront le même principe. Exportez les graphes dans votre compte-rendu *Word* et commentez-les⁷.

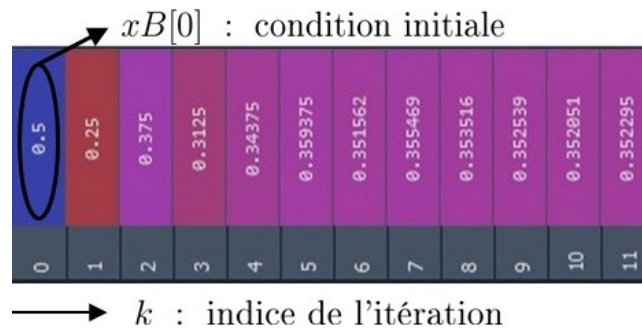


Figure 3: Tableau solution $\mathbf{xB}$ pour l'algorithme de la bisection, visible dans l'explorateur de variables de Spyder.

- 2) Ecrivez l'algorithme de Newton dans le même esprit que celui de la bisection. Plus précisément, créez la fonction Python :

```
def newton(fun, dfun, x0, eps, kmax):

    return x[0:k+1], res[0:k+1]
```

qui retourne le résidu r_k et la solution x_k à chaque itération k , avec pour variable d'entrée la fonction $f(x)$, sa dérivée $f'(x)$, la condition initiale x_0 , la tolérance ϵ et le nombre maximal d'itérations $k_{\max}$.

- a) Validez votre programme en l'appliquant au problème test à partir de $x_0 = 0$.
 - b) Que se passe-t-il si on initialise l'algorithme loin de la racine à trouver, par exemple $x_0 = 0.9$?
 - c) Comparez la vitesse de convergence de la méthode de Newton avec celle de l'algorithme de la bisection.
- 3) Dans cette question vous allez combiner les algorithmes de la bisection et de Newton pour résoudre le problème des tiges (*cf.* cours) avec $a_1 = 10$, $a_2 = 13$, $a_3 = 8$ et $a_4 = 10$. Pour une valeur de l'angle ω que vous choisirez, le but de la question est de déterminer la (ou les) racines de l'équation $f_\omega(\theta) = 0$ dans l'intervalle $I = [-\pi, \pi]$.

Afin de définir une condition initiale "suffisamment proche" pour que l'algorithme de Newton converge, vous allez appliquer la stratégie de l'exercice 2 du TD 2 :

- a) Pour chaque racine de l'intervalle I , déterminez graphiquement un encadrement initial I_0 .
- b) Faites quelques itérations avec l'algorithme de la bisection, pour obtenir une approximation $\tilde{x}$ avec la tolérance $\epsilon = 10^{-1}$.
- c) A partir de $\tilde{x}$, appliquez ensuite l'algorithme de Newton pour calculer la racine avec la tolérance de $\epsilon = 10^{-14}$.

⁷Il est intéressant d'observer que les graphes du résidu et de l'erreur ont le même comportement : ainsi, lorsqu'on ne connaît pas la solution exacte (ce qui est le cas dans les problèmes d'ingénierie), la valeur du résidu fournit un bon

Dans le compte-rendu, expliquez votre méthodologie (valeur de l'angle ω , encadrement initial, tolérance et nombre d'itérations pour chacun des algorithmes, résultat final) et montrez le graphe des résidus.

Travaux pratiques 3. (*Schémas numériques pour équations différentielles ordinaires*)

Ce TP concerne la résolution du problème à valeur initiale :

$$\begin{cases} \dot{y}(t) = f(t, y(t)), & t \in [t_0, t_f], \\ y(t_0) = y_0, \end{cases} \quad (1)$$

par les schémas numériques vus en cours et dans le TD 3 (schémas d'Euler progressif, Runge-Kutta, ...). Le répertoire de travail Python est `TP3_L3_SCHEMAS_EDO`.

- 1) Le script `main_TP_schemas_EDO.py` applique le schéma EP1 à la résolution de l'EDO du premier ordre (1) pour $t \in [0, 12]$, avec $f(t, y) = -y(1/10 - \sin(t))$ et la condition initiale $y(0) = 1$, dont la solution exacte est :

$$y(t) = \frac{e^{-1/10 t - \cos(t)}}{e^{-1}}.$$

Cette EDO sera appelée *exemple 1* dans la suite de l'énoncé.

- a) Tracez sur une même figure la solution numérique donnée par EP1 pour des valeurs de h de plus en plus petite : $h = 1; 0.5; 0.1$; etc Qu'observez-vous lorsque $h \rightarrow 0$?
- b) Programmez le schéma de Heun, dénommé `def Heun(f, t0, tf, y0, h):`, dans le même esprit que la fonction Python EP1. Comparez sur une même figure la solution donnée par les schémas EP1 et de Heun pour un même pas de temps, par exemple $h = 0.4$. Est-ce en accord avec les propriétés théoriques de ces schémas ?
- c) Déterminez analytiquement la condition de stabilité du schéma EP1⁸ pour l'exemple 1. Reportez sur une même figure les solutions données par EP1 pour des valeurs de h choisies au dessus et en dessous de la condition de stabilité⁹.

- 2) On considère maintenant l'EDO du second-ordre :

$$\begin{cases} \ddot{\theta}(t) = f(t, \theta(t)), & t \in [t_0, t_f], \\ \theta(t_0) = \theta_0, \quad \dot{\theta}(t_0) = \dot{\theta}_0. \end{cases} \quad (2)$$

avec $f(t, \theta) = -\omega^2 \theta$, où $\omega > 0$ est la pulsation propre. Les paramètres de l'EDO sont :

$$\omega = 4, \quad t_0 = 0, \quad t_f = 2\pi,$$

et les conditions initiales sont :

$$\theta_0 = 1, \quad \dot{\theta}_0 = 0.$$

renseignement sur la qualité de la solution numérique.

⁸Rappel de cours : le schéma EP1 est stable si

$$h < \frac{2}{\max_{t \in [t_0, t_f]} \{-\lambda\}}, \quad \text{avec } \lambda = \frac{\partial f}{\partial y}.$$

⁹Pour observer au mieux la stabilité du schéma au temps long, il est recommandé d'augmenter le temps d'intégration, par exemple `tf=50`, et de limiter l'axe des ordonnées, par exemple `plt.ylim(-2, 10)`.

Ce problème sera appelé dans la suite *exemple 2*, et sa solution exacte est :

$$\theta_{\text{exact}}(t) = \cos(\omega t). \quad (3)$$

Cette EDO est autonome (le second-membre f est indépendant du temps) et linéaire (sa solution analytique est une combinaison linéaire de sinus et cosinus), et elle représente la modélisation mathématique du pendule linéaire (tel que $\theta(t)$ représente l'angle du pendule et $\dot{\theta}(t)$ sa vitesse angulaire) ou de l'oscillateur harmonique sans frottement du TD. Cette EDO est de type conservatif car l'énergie mécanique est conservée au cours du mouvement (*cf.* TD).

- a) A partir de la fonction Python EP1, créez une version adaptée à la résolution d'EDO du second ordre (nommez-la par exemple `def EP1_2D(f, t0, tf, y0, h):`). Conseil : réécrivez l'EDO (2) en un système d'EDO du premier ordre (*cf.* cours et TD), de façon que le second membre s'écrive :

```
def fedo_2D(t, y):

    # ... Second-membre pour l'EDO de l'exemple 2
    w = 4
    F = np.zeros(2)
    F[0] = y[1]
    F[1] = -w**2*y[0]

    return F
```

Dans le script principal, le vecteur `y0= np.array([1, 0])` correspondra à la condition initiale $(\theta_0, \dot{\theta}_0) = (1, 0)$. Dans la fonction EP1_2D, le vecteur solution sera initialisé par le tableau `y = np.zeros((2,Nt + 1))`. La condition initiale sera `y[:, 0]= y0`, et la sous-matrice colonne `y[:, n]= [y[0, n], y[1, n]]` correspondra à la solution au temps $t_n = nh$: `y[0, n] =  $\theta(t_n)$` et `y[1, n] =  $\dot{\theta}(t_n)$` . Un exemple est donné sur la figure 4.

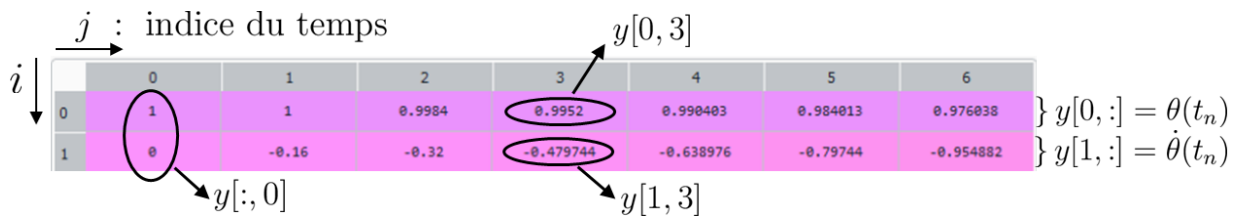


Figure 4: Exemple de tableau solution y pour la résolution d'EDO du second ordre.

- b) Utilisez le schéma EP1 pour résoudre l'exemple 2 et comparez la solution numérique $\theta(t_n) = y[0, n]$ et la solution analytique (3) sur une même figure¹⁰ : qu'observez-vous ? Est-ce en accord avec les propriétés de stabilité théorique de EP1 pour ce type d'EDO (*cf.* TD exercice 5) ? Que penser de la stabilité du schéma de Heun pour cet exemple ?
- c) Programmez le schéma RK4 (*cf.* TD). Énoncez la condition de stabilité pour ce schéma, et déterminez la valeur maximale de stabilité $h_{\max}$ du pas de temps. Utilisez RK4 avec un pas de temps h légèrement inférieur à $h_{\max}$, puis une valeur bien inférieure (*ie.*, $h_{\max}/10$). Comment qualifiez-vous le type d'erreur donné par RK4 par rapport à ce qui est représenté sur la figure 5 ?

¹⁰Il est à noter que la solution analytique est telle que $\theta(t) \in [-\theta_0, \theta_0]$ donc restreignez les ordonnées de la figure avec `plt.ylim(-5, 5)` par exemple.

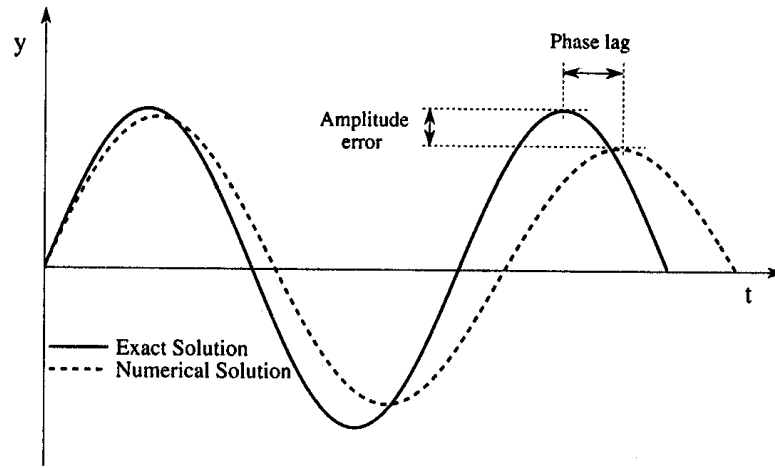


Figure 5: Décomposition de l'erreur numérique en une composante d'erreur d'amplitude (ou dissipation numérique) et une composante d'erreur de phase (ou dispersion numérique).

- d) Dans le plan $(\theta, \dot{\theta}) = [y[0, n], y[1, n]]$ appelé *espace des phases*, tracez les solutions numériques obtenues à partir de plusieurs conditions initiales, par exemple $y_0 = [1, 0]$, $y_0 = [2, 0]$ et $y_0 = [4, 0]$. Ce type de graphe est appelé *portrait de phases* de l'EDO, et on appelle *trajectoire* une solution partant d'une condition initiale donnée. Pourquoi les trajectoires de cette EDO correspondent-elles à des courbes fermées dans l'espace des phases ?
- 3) On considère maintenant l'EDO (2) avec $f(t, \theta) = -g \sin(\theta)/l$, qui correspond à l'équation du pendule non-linéaire de l'exercice 5 du TD. Les conditions initiales seront choisies telles que le pendule est lancé avec une vitesse angulaire $\dot{\theta}_0$ non-nulle à partir de la position d'équilibre stable $\theta_0 = 0$. On se placera dans les conditions de l'exercice 5, soit : $g = 9.81 \text{ m/s}^2$ et $l = 0.6 \text{ m}$.
 - a) Énoncez la condition de stabilité du schéma RK4 pour cette EDO.
 - b) Utilisez le schéma RK4 pour calculer plusieurs solutions numériques de vitesse angulaire initiale $\dot{\theta}_0$ choisie au dessous et au dessus de la valeur $\dot{\theta}_0^* = 2\sqrt{g/l}$, afin de construire un portrait de phases semblable à celui représenté sur la figure 6.

Travaux pratiques 4. (*Méthodes itératives pour systèmes d'équations linéaires*)

Ce TP concerne la résolution de systèmes linéaires de taille $n \times n$ de la forme $\mathcal{A}X = F$ avec les méthodes itératives vues en cours : algorithme de Jacobi, Gauss-Seidel et relaxation (ou *S.O.R.*). Ces algorithmes seront appliqués à la résolution du problème test :

$$\mathcal{K}X = F, \quad (1)$$

avec $\mathcal{K}$ la matrice tridiagonale $\mathcal{K} = \text{Tridiag}_n[-1, 2, -1]$ (vue en cours et dans l'exercice 4 du TD 4), et la solution exacte est choisie comme $X = (1, \dots, 1)$. Ce TP a pour but d'illustrer les propriétés de convergence des algorithmes qui ont été vues théoriquement dans l'exercice 4, et il est donc important que cet exercice ait été entièrement fait et compris. Le répertoire de travail Python du TP est TP4_L3_SYST_EQTS_LIN.

- 1) Le script principal `main_TP_systemes_lineaires.py` applique la méthode de Jacobi à la résolution de (1). Comprenez et exécutez ce programme pour un système linéaire de petite taille, tel que $n = 2$. La suite des questions a pour but de comparer la performance des divers algorithmes pour le cas $n = 2$. Tous les tests seront effectués pour la même condition initiale $X_0 = 0$ et la tolérance $\epsilon = 10^{-10}$.

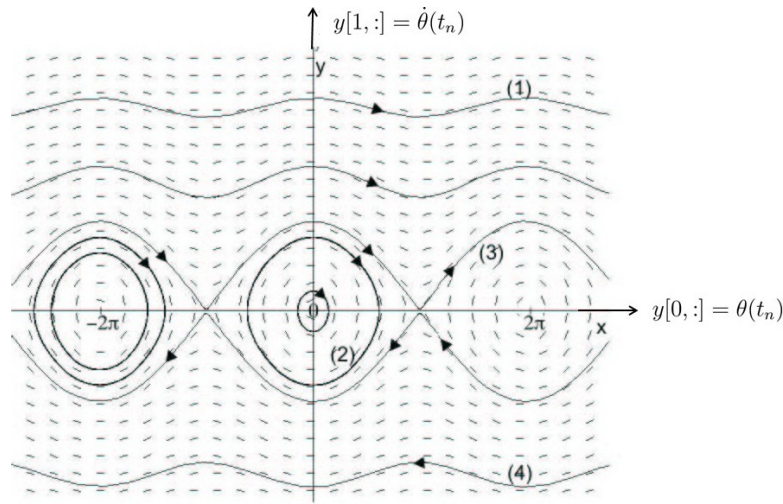


Figure 6: Portrait de phases du pendule non-linéaire pour diverses trajectoires partant de la condition initiale $\theta_0 = 0$ et $\dot{\theta}_0$. La trajectoire (3) de vitesse initiale $\dot{\theta}_0^* = 2\sqrt{g/l}$ correspond à des solutions joignant deux positions d'équilibre instable $\theta(t) = \pi \bmod(\pi)$, la trajectoire (2) de vitesse initiale faible (telle que $|\dot{\theta}_0| < |\dot{\theta}_0^*|$) correspond à des oscillations du pendule autour de la position d'équilibre stable, les trajectoires (1) et (4) correspondent au cas où le pendule est lancé avec une vitesse initiale suffisamment grande ($|\dot{\theta}_0| > |\dot{\theta}_0^*|$) pour lui permettre de faire des tours sur lui-même.

- Programmez la méthode *S.O.R.* pour un paramètre de relaxation ω quelconque. Conseil : utilisez au maximum la manière utilisée pour programmer la méthode de Jacobi.
- Programmez la fonction `def w_opt(n)` : qui donne la valeur du paramètre de relaxation optimal pour le problème test (1).
- Utilisez ces fonctions pour résoudre le problème test avec les méthodes de Gauss-Seidel et *S.O.R.* avec paramètre de relaxation optimal ω_{opt} .
- Comparez sur le même graphe les résultats donnés par les méthodes de Jacobi, G-S et S.O.R. (résidu, erreur et facteur de réduction). Lorsque vous aurez confirmé les résultats théoriques du cours sur la valeur asymptotique du facteur de réduction, les méthodes que vous avez programmées seront considérées comme validées, et vous pourrez passer aux questions suivantes.

2) On considère dans cette question le problème test (1) avec $n = 31$.

- Utilisez les trois méthodes itératives pour résoudre ce problème avec la tolérance de $\epsilon = 10^{-8}$.
- Comparez les divers résultats des méthodes sur les mêmes graphes. Pourquoi observe-t-on que l'erreur relative est-elle bien supérieure à la norme du résidu ?
- Pour chaque méthode, estimez le nombre minimal d'itérations pour que l'erreur relative soit inférieure à $\epsilon = 10^{-6}$. Est-ce en accord avec les résultats théoriques vus en TD ?