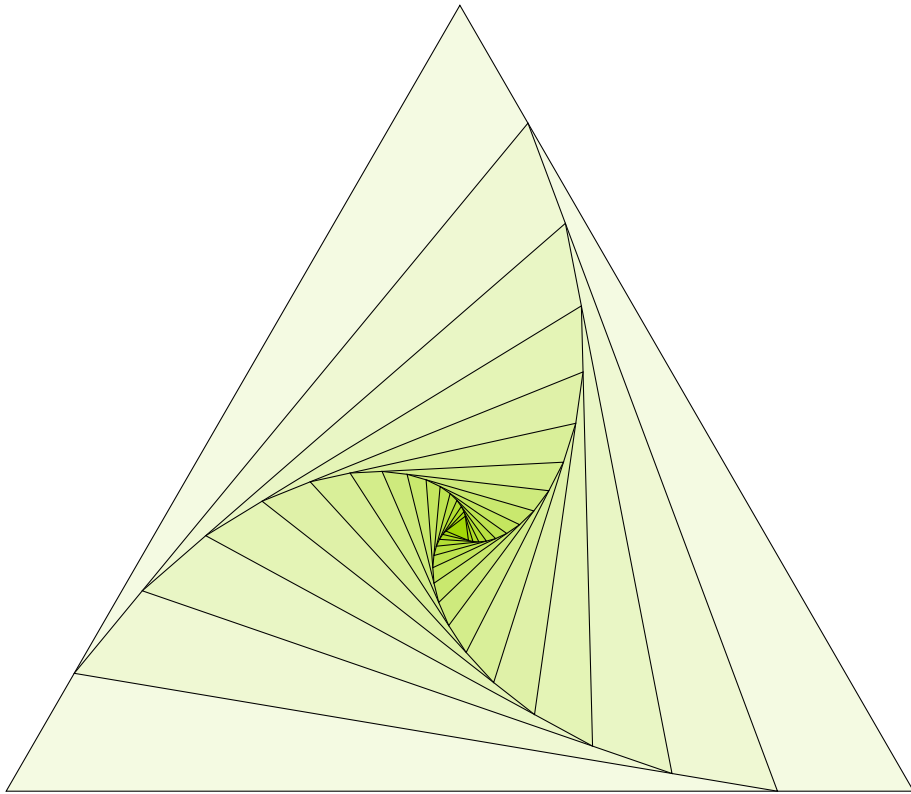


Analyse numérique élémentaire

Notes de cours

Sup Galilée, Ingénieurs MACS 1ère année



Francois Cuvelier
Université Paris XIII / Institut Galilée
L.A.G.A./Département de Mathématiques
<http://www.math.univ-paris13.fr/~cuvelier>

Table des matières

1	Représentation des nombres en machine, erreurs d'arrondis	1	2
1.1	Un exemple : calcul approché de π	1	3
1.2	Représentation scientifique des nombres dans différentes bases	2	4
1.2.1	Partie entière, mantisse et exposant	2	5
1.3	Nombres flottants : le système IEEE 754	5	6
1.3.1	Simple précision	6	7
1.3.2	Double précision	7	8
1.3.3	En MATLAB	7	9
1.4	Calculs sur les nombres flottants	8	10
1.4.1	Erreurs d'arrondi	8	11
1.4.2	Associativité	8	12
1.4.3	Monotonie	8	13
1.4.4	Erreurs d'annulation	8	14
1.5	Quelques catastrophes dues à l'arithmétique flottante	9	15
2	Langage algorithmique	13	16
2.1	Pseudo-langage algorithmique	13	17
2.1.1	Données et constantes	13	18
2.1.2	Variables	13	19
2.1.3	Opérateurs	14	20
2.1.4	Expressions	14	21
2.1.5	Instructions	15	22
2.1.6	Fonctions	16	23
2.2	Méthodologie d'élaboration d'un algorithme	18	24
2.2.1	Description du problème	18	25
2.2.2	Recherche d'une méthode de résolution	18	26
2.2.3	Réalisation d'un algorithme	19	27
2.2.4	Exercices	19	28
2.3	Principes de «bonne» programmation pour attaquer de «gros» problèmes	21	29
3	Résolution de systèmes non linéaires	23	30
3.1	Recherche des zéros d'une fonction	24	31
3.1.1	Méthode de dichotomie ou de bisection	24	32
3.2	Points fixes d'une fonction (dimension 1)	30	33

A.3	Receuil d'exercices	189	1
A.3.1	Algèbre linéaire	189	2
A.3.2	Normes	199	3
A.4	Listings	206	4
A.4.1	Codes sur la méthode de dichotomie/bissection	206	5

Chapitre 1

Représentation des nombres en machine, erreurs d'arrondis



Toute cette partie est le contenu quasi-intégral d'un document réalisé par C. Japhet

2

Ce chapitre est une introduction à la représentation des nombres en machine et aux erreurs d'arrondis, basé sur [5], [4].

3

4

1.1 Un exemple : calcul approché de π

5

Cet exemple est extrait de [5], [4]. Le nombre π est connu depuis l'antiquité, en tant que méthode de calcul du périmètre du cercle ou de l'aire du disque. Le problème de la quadrature du cercle étudié par les anciens Grecs consiste à construire un carré de même aire qu'un cercle donné à l'aide d'une règle et d'un compas. Ce problème resta insoluble jusqu'au 19^{ème} siècle, où la démonstration de la transcendance de π montra que le problème ne peut être résolu en utilisant une règle et un compas.

6

7

8

9

10

Nous savons aujourd'hui que l'aire d'un cercle de rayon r est $\mathcal{A} = \pi r^2$. Parmi les solutions proposées pour approcher $\mathcal{A}$, une méthode consiste à construire un polygone dont le nombre de côté augmenterait jusqu'à ce qu'il devienne équivalent au cercle circonscrit. C'est Archimède vers 250 avant J-C qui appliquera cette propriété au calcul des décimales du nombre π , en utilisant à la fois un polygone inscrit et circonscrit au cercle. Il utilise ainsi un algorithme pour le calcul et parvient à l'approximation de π dans l'intervalle $(3 + \frac{1}{7}, 3 + \frac{10}{71})$ en faisant tendre le nombre de côtés jusqu'à 96.

11

12

13

14

15

16

Regardons l'algorithme de calcul par les polygones inscrits. On considère un cercle de rayon $r = 1$ et on note $\mathcal{A}_n$ l'aire associée au polygone inscrit à n côtés. En notant $\alpha_n = \frac{2\pi}{n}$, $\mathcal{A}_n$ est égale à n fois l'aire du triangle ABC représenté sur la figure 1.1, c'est-à-dire

$$\mathcal{A}_n = n \cos \frac{\alpha_n}{2} \sin \frac{\alpha_n}{2},$$

que l'on peut réécrire

$$\mathcal{A}_n = \frac{n}{2} (2 \cos \frac{\alpha_n}{2} \sin \frac{\alpha_n}{2}) = \frac{n}{2} \sin \alpha_n = \frac{n}{2} \sin(\frac{2\pi}{n}).$$

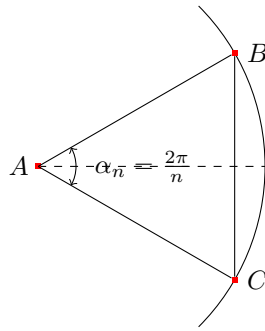


Figure 1.1: Quadrature du cercle

Comme on cherche à calculer π à l'aide de A_n , on ne peut pas utiliser l'expression ci-dessus pour calculer A_n , mais on peut exprimer A_{2n} en fonction de A_n en utilisant la relation

$$\sin \frac{\alpha_n}{2} = \sqrt{\frac{1 - \cos \alpha_n}{2}} = \sqrt{\frac{1 - \sqrt{1 - \sin^2 \alpha_n}}{2}}.$$

1 Ainsi, en prenant $n = 2^k$, on définit l'approximation de π par récurrence

$$x_k = A_{2^k} = \frac{2^k}{2} s_k, \quad \text{avec } s_k = \sin\left(\frac{2\pi}{2^k}\right) = \sqrt{\frac{1 - \sqrt{1 - s_{k-1}^2}}{2}}$$

2 En partant de $k = 2$ (i.e. $n = 4$ et $s = 1$) on obtient l'algorithme suivant:

Algorithme 1.1 Algorithme de calcul de π , version naïve

1: $s \leftarrow 1, n \leftarrow 4$	▷ Initialisations
2: Tantque $s > 1e - 10$ faire	▷ Arrêt si $s = \sin(\alpha)$ est petit
3: $s \leftarrow \text{sqrt}((1 - \text{sqrt}(1 - s * s))/2)$	▷ nouvelle valeur de $\sin(\alpha/2)$
4: $n \leftarrow 2 * n$	▷ nouvelle valeur de n
5: $A \leftarrow (n/2) * s$	▷ nouvelle valeur de l'aire du polygone
6: Fin Tantque	

3 On a $\lim_{k \rightarrow +\infty} x_k = \pi$. Ce n'est pourtant pas du tout ce que l'on va observer sur machine! Les
 4 résultats en Python (sous Sage) de la table 1.1 montre que l'algorithme commence par converger vers π
 5 puis pour $n > 65536$, l'erreur augmente et finalement on obtient $A_n = 0$!! "Although the theory and the
 6 program are correct, we obtain incorrect answers" ([5]).

7 Ceci résulte du codage des valeurs réelles sur un nombre fini de bits, ce que nous allons détailler dans
 8 ce chapitre.

1.2 Représentation scientifique des nombres dans différentes bases

11 Dans cette section nous introduisons les notions de mantisse, exposant, et la façon dont sont représentés
 12 les nombres sur une calculatrice ou un ordinateur.

1.2.1 Partie entière, mantisse et exposant

14 Exemple en base 10

15 La base 10 est la base naturelle avec laquelle on travaille et celle que l'on retrouve dans les calculatrices.

16 Un nombre à virgule, ou nombre décimal, a plusieurs écritures différentes en changeant simplement
 17 la position du point décimal et en rajoutant à la fin une puissance de 10 dans l'écriture de ce nombre. La
 18 partie à gauche du point décimal est la partie entière, celle à droite avant l'exposant s'appelle la mantisse.
 19 Par exemple le nombre $x = 1234.5678$ a plusieurs représentations :

$$x = 1234.5678 = 1234.5678 \cdot 10^0 = 1.2345678 \cdot 10^3 = 0.0012345678 \cdot 10^6, \quad (1.1)$$

n	A_n	$ A_n - \pi $	$\sin(\alpha_n)$
4	2.00000000000000	1.141593e+00	1.000000e+00
8	2.82842712474619	3.131655e-01	7.071068e-01
16	3.06146745892072	8.012519e-02	3.826834e-01
32	3.12144515225805	2.014750e-02	1.950903e-01
64	3.13654849054594	5.044163e-03	9.801714e-02
128	3.14033115695474	1.261497e-03	4.906767e-02
256	3.14127725093276	3.154027e-04	2.454123e-02
512	3.14151380114415	7.885245e-05	1.227154e-02
1024	3.14157294036788	1.971322e-05	6.135885e-03
2048	3.14158772527996	4.928310e-06	3.067957e-03
4096	3.14159142150464	1.232085e-06	1.533980e-03
8192	3.14159234561108	3.079787e-07	7.669903e-04
16384	3.14159257654500	7.704479e-08	3.834952e-04
32768	3.14159263346325	2.012654e-08	1.917476e-04
65536	3.14159265480759	1.217796e-09	9.587380e-05
131072	3.14159264532122	8.268578e-09	4.793690e-05
262144	3.14159260737572	4.621407e-08	2.396845e-05
524288	3.14159291093967	2.573499e-07	1.198423e-05
1048576	3.14159412519519	1.471605e-06	5.992115e-06
2097152	3.14159655370482	3.900115e-06	2.996060e-06
4194304	3.14159655370482	3.900115e-06	1.498030e-06
8388608	3.14167426502176	8.161143e-05	7.490335e-07
16777216	3.14182968188920	2.370283e-04	3.745353e-07
33554432	3.14245127249413	8.586189e-04	1.873047e-07
67108864	3.14245127249413	8.586189e-04	9.365235e-08
134217728	3.16227766016838	2.068501e-02	4.712161e-08
268435456	3.16227766016838	2.068501e-02	2.356080e-08
536870912	3.46410161513775	3.225090e-01	1.290478e-08
1073741824	4.00000000000000	8.584073e-01	7.450581e-09
2147483648	0.00000000000000	3.141593e+00	0.000000e+00

Table 1.1: Calcul de π avec l'algorithme naïf 1.1

avec

- *Partie entière* : 1234
- *Mantisse* : 0.5678 ou 1.2345678 ou 0.0012345678
- *Exposant* : 4 ou 6

Selon le décalage et l'exposant que l'on aura choisi, le couple mantisse-exposant va changer mais le nombre représenté est le même. Afin d'avoir une représentation unique, celle qui sera utilisée c'est la troisième dans (1.1), où la mantisse est 1.2345678 et l'exposant 3. C'est celle où le premier chiffre avant le point décimal dans la mantisse est non nul.

Exemple en base 2

C'est la base que les ordinateurs utilisent. Les chiffres utilisables en base 2 sont 0 et 1 que l'on appelle *bit* pour *binary digit*, les ordinateurs travaillent en binaire. Par exemple

$$39 = 32 + 4 + 2 + 1 = 2^5 + 2^2 + 2^1 + 2^0 = (100111)_2,$$
$$3.625 = 2^1 + 2^0 + 2^{-1} + 2^{-3} = (11.101)_2 = (1.1101)_2 2^1$$

Représentation d'un nombre en machine : nombres flottants

De façon générale tout nombre réel x sera représenté dans une base b ($b = 10$ pour une calculatrice $b = 2$ pour un ordinateur) par son signe (+ ou -), la mantisse m (appelée aussi significande), la base b et un

1 exposant e tel que le couple (m, e) caractérise le nombre. En faisant varier e , on fait « flotter » la virgule
 2 décimale. La limitation fondamentale est que la place mémoire d'un ordinateur est limitée, c'est-à-dire
 3 qu'il ne pourra stocker qu'un ensemble fini de nombres. Ainsi un nombre machine réel ou *nombre à*
 4 *virgule flottante* s'écrira :

$$\begin{aligned}\tilde{x} &= \pm m \cdot b^e \\ m &= D.D \cdots D \\ e &= D \cdots D\end{aligned}$$

5 où $D \in \{0, 1, \dots, b-1\}$ représente un chiffre. Des représentations approchées de π sont : $(0.031, 2)$, $(3.142, 0)$,
 6 $(0.003, 3)$ et on observe qu'elles ne donnent pas la même précision. Pour rendre la représentation unique
 7 et garder la meilleure précision, on utilisera une mantisse *normalisée* : le premier chiffre avant le point
 8 décimal dans la mantisse est non nul. Les nombres machine correspondants sont appelés *normalisés*.
 9 En base 2, le premier bit dans la mantisse sera donc toujours 1, et on n'écrit pas ce 1 ce qui permet
 10 d'économiser un bit. L'exposant est un nombre variant dans un intervalle fini de valeurs admissibles :
 11 $L \leq e \leq U$ (typiquement $L < 0$ et $U > 0$). Le système est donc caractérisé par quatre entiers :

- 12 • la base b ($b = 2$),
- 13 • le nombre de chiffres t dans la mantisse (en base b),
- 14 • l'exposant minimal L et maximal U .

15
 16 En mathématiques on effectue les calculs avec des nombres réels x provenant de l'intervalle continu
 17 $x \in [-\infty, \infty]$. A cause de la limitation ci-dessus, la plupart des réels seront approchés sur un ordinateur.
 18 Par exemple, $\frac{1}{3}$, $\sqrt{2}$, π possèdent une infinité de décimales et ne peuvent donc pas avoir de représentation
 19 exacte en machine. Le plus simple des calculs devient alors approché. L'expérience pratique montre que
 20 cette quantité limitée de nombres représentables est largement suffisante pour les calculs. Sur l'ordinateur,
 21 les nombres utilisés lors des calculs sont des nombres machine $\tilde{x}$ provenant d'un ensemble discret de
 22 nombres machine $\tilde{x} \in \{\tilde{x}_{min}, \dots, \tilde{x}_{max}\}$. Ainsi, chaque nombre réel x doit être transformé en un nombre
 23 machine $\tilde{x}$ afin de pouvoir être utilisé sur un ordinateur. Un exemple est donné sur la figure 1.1.

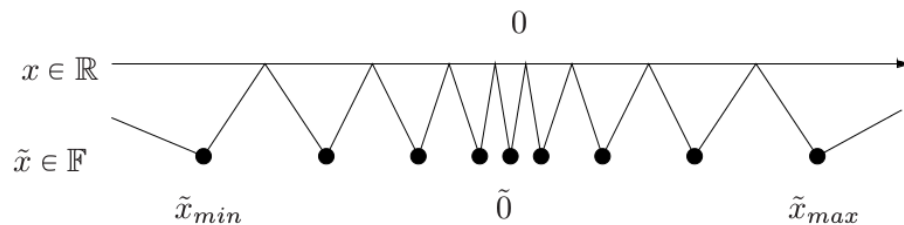


Figure 1.2: Représentation des nombres réels $\mathbb{R}$ par les nombres machine $\mathbb{F}$

24 Prenons un exemple, beaucoup trop simple pour être utilisé mais pour fixer les idées: $t = 3$, $L =$
 25 -1 , $U = 2$. Dans ce cas on a 3 chiffres significatifs et 33 nombres dans le système $\mathbb{F}$. Ils se répartissent
 26 avec 0 d'une part, 16 nombres négatifs que l'on ne représente pas ici, et 16 nombres positifs représentés
 sur la figure 1.3.

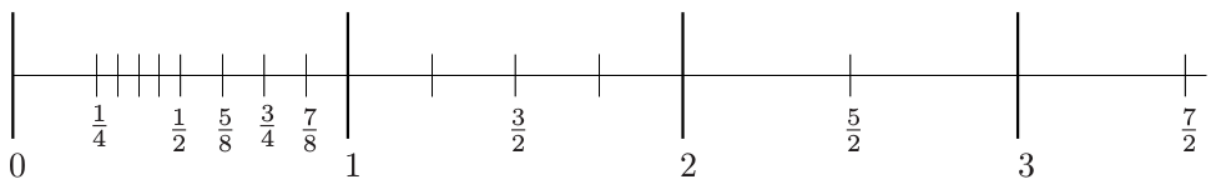


Figure 1.3: Nombres positifs de $\mathbb{F}$ dans le cas $t = 3$, $L = -1$, $U = 2$

Dans cet exemple, l'écriture en binaire des nombres entre $\frac{1}{2}$ et 1 est

$$\begin{aligned}\frac{1}{2} &= (0.100)_2, & \frac{3}{4} &= \frac{1}{2} + \frac{1}{4} = (0.110)_2, \\ \frac{5}{8} &= \frac{1}{2} + \frac{1}{8} = (0.101)_2, & \frac{7}{8} &= \frac{1}{2} + \frac{1}{4} + \frac{1}{8} = (0.111)_2.\end{aligned}$$

On obtient ensuite les autres nombres en multipliant par une puissance de 2. Le plus grand nombre représentable dans ce système est $\frac{7}{2}$ (en particulier 4 n'est pas représentable). On remarque que les nombres ne sont pas espacés régulièrement. Ils sont beaucoup plus resserrés du côté de 0 entre $\frac{1}{4}$ et $\frac{1}{2}$ que entre 1 et 2 et encore plus qu'entre 2 et 3. Plus précisément, chaque fois que l'on passe par une puissance de 2, l'espacement absolu est multiplié par 2, mais l'espacement relatif reste constant ce qui est une bonne chose pour un calcul d'ingénierie car on a besoin d'une précision absolue beaucoup plus grande pour des nombres petits (autour de un millièmme par exemple) que des nombres très grands (de l'ordre du million par exemple). Mais la précision ou l'erreur relative sera du même ordre. L'erreur absolue ou relative est définie à section 1.4.1.

Précision machine. elle est décrite par le nombre machine *eps*. *eps* est le plus petit nombre machine positif tel que $1 + \text{eps} > 1$ sur la machine. C'est la distance entre l'entier 1 et le nombre machine $\tilde{x} \in \mathbb{F}$ le plus proche, qui lui est supérieur. Dans l'exemple précédent $\text{eps} = 1/4$.

Sur une calculatrice

Le système utilisé est la base 10 ($b = 10$). Typiquement, il y a 10 chiffres pour la mantisse et 2 pour l'exposant ($L = -99$ et $U = 99$).

- Le plus grand nombre machine

$$\tilde{x}_{max} = 9.999999999 \times 10^{+99}$$

- Le plus petit nombre machine

$$\tilde{x}_{min} = -9.999999999 \times 10^{+99}$$

- Le plus petit nombre machine strictement positif

$$\tilde{x}_+ = 1.000000000 \times 10^{-99}$$

Notez qu'avec des nombres dénormalisés, ce nombre serait $0.000000001 \times 10^{-99}$, c'est-à-dire avec seulement un chiffre significatif!

Les différences de représentation des nombres flottants d'un ordinateur à un autre obligeaient à reprendre les programmes de calcul scientifique pour les porter d'une machine à une autre. Pour assurer la compatibilité entre les machines, depuis 1985 une norme a été proposée par l'IEEE (Institute of Electrical and Electronics Engineers), c'est la norme 754.

1.3 Nombres flottants : le système IEEE 754

Le système IEEE 754 est un standard pour la représentation des nombres à virgule flottante en binaire. Il définit les formats de représentation des nombres à virgule flottante (signe, mantisse, exposant, nombres dénormalisés) et valeurs spéciales (infinis et NaN). Le bit de poids fort est le bit de signe. Cela signifie que si ce bit est à 1, le nombre est négatif, et s'il est à 0, le nombre est positif. Les N_e bits suivants représentent l'exposant décalé, et les N_m bits suivants représentent la mantisse.

L'exposant est décalé de $2^{N_e-1} - 1$ (N_e représente le nombre de bits de l'exposant), afin de le stocker sous forme d'un nombre non signé.

1.3.1 Simple précision

- C'est le format 32 bits : 1 bit de signe, $N_e = 8$ bits d'exposant (-126 à 127), 23 bits de mantisse comme sur le tableau 1.2. L'exposant est décalé de $2^{N_e-1} - 1 = 2^7 - 1 = 127$.

signe s	exposant e	mantisse m
S	$EEEEEEEE$	$FFFFFFFFFFFFFFFFFFFFFFFFFFFF$
0	1 8	9 31

Table 1.2: Représentation en simple précision

$\tilde{x}$	exposant e	mantisse m
$\tilde{x} = 0$ (si $S = 0$) $\tilde{x} = -0$ (si $S = 1$)	$e = 0$	$m = 0$
Nombre <i>normalisé</i> $\tilde{x} = (-1)^S \times 2^{e-127} \times 1.m$	$0 < e < 255$	quelconque
Nombre <i>dénormalisé</i> $\tilde{x} = (-1)^S \times 2^{e-126} \times 0.m$	$e = 0$	$m \neq 0$
$\tilde{x} = \text{Inf}$ (si $S = 0$) $\tilde{x} = -\text{Inf}$ (si $S = 1$)	$e = 255$	$m = 0$
$\tilde{x} = \text{NaN}$ (<i>Not a Number</i>)	$e = 255$	$m \neq 0$

Table 1.3: Représentation en simple précision

- Le *plus petit nombre positif normalisé* différent de zéro, et le *plus grand nombre négatif normalisé* différent de zéro sont :
 $\pm 2^{-126} = \pm 1,175494351 \times 10^{-38}$
- Le *plus grand nombre positif fini*, et le *plus petit nombre négatif fini* sont : $\pm(2^{24} - 1) \times 2^{104} = \pm 3,4028235 \times 10^{38}$

$\tilde{x}$	signe	exposant e	mantisse m	valeur
zéro	0	0000 0000	000 0000 0000 0000 0000 0000	0,0
	1	0000 0000	000 0000 0000 0000 0000 0000	-0,0
1	0	0111 1111	000 0000 0000 0000 0000 0000	1,0
Plus grand nombre normalisé	0	1111 1110	111 1111 1111 1111 1111 1111	$3,4 \times 10^{38}$
Plus petit $\tilde{x} \geq 0$ normalisé	0	0000 0001	000 0000 0000 0000 0000 0000	2^{-126}
Plus petit $\tilde{x} \geq 0$ dénormalisé	0	0000 0000	000 0000 0000 0000 0000 0001	2^{-149}
Infini	0	1111 1111	000 0000 0000 0000 0000 0000	Inf
	1	1111 1111	000 0000 0000 0000 0000 0000	-Inf
NaN	0	1111 1111	010 0000 0000 0000 0000 0000	NaN
2	0	1000 0000	000 0000 0000 0000 0000 0000	2,0
exemple 1	0	1000 0001	101 0000 0000 0000 0000 0000	6,5
exemple 2	1	1000 0001	101 0000 0000 0000 0000 0000	-6,5
exemple 3	1	1000 0101	110 1101 0100 0000 0000 0000	-118,625

Table 1.4: Exemples en simple précision

- Détaillons l'exemple 3 : on code le nombre décimal -118,625 en utilisant le système IEEE 754.
- C'est un nombre négatif, le bit de signe est donc "1",
 - On écrit le nombre (sans le signe) en binaire. Nous obtenons 1110110,101,
 - On décale la virgule vers la gauche, en laissant seulement un 1 sur sa gauche (nombre flottant normalisé): $1110110,101 = 1,110110101 \times 2^6$. La mantisse est la partie à droite de la virgule, remplie de 0 vers la droite pour obtenir 23 bits. Cela donne 11011010100000000000000,
 - L'exposant est égal à 6, et nous devons le convertir en binaire et le décaler. Pour le format 32-bit IEEE 754, le décalage est 127. Donc $6 + 127 = 133$. En binaire, cela donne 10000101.

1.3.2 Double précision

C'est le format 64 bits : 1 bit de signe, 11 bits d'exposant (-1022 à 1023), 52 bits de mantisse. L'exposant est décalé de $2^{N_e-1} - 1 = 2^{10} - 1 = 1023$.

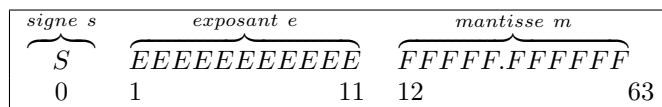


Table 1.5: Représentation en double précision

$\tilde{x}$	exposant e	mantisse m
$\tilde{x} = 0$ (si $S = 0$) $\tilde{x} = -0$ (si $S = 1$)	$e = 0$	$m = 0$
Nombre <i>normalisé</i> $\tilde{x} = (-1)^S \times 2^{e-1023} \times 1.m$	$0 < e < 2047$	quelconque
Nombre <i>dénormalisé</i> $\tilde{x} = (-1)^S \times 2^{e-1022} \times 0.m$	$e = 0$	$m \neq 0$
$\tilde{x} = \text{Inf}$ (si $S = 0$) $\tilde{x} = -\text{Inf}$ (si $S = 1$)	$e = 2047$	$m = 0$
$\tilde{x} = \text{NaN}$ (<i>Not a Number</i>)	$e = 2047$	$m \neq 0$

Table 1.6: Représentation en double précision

- La précision machine est $\text{eps} = 2^{-52}$.
- Le *plus grand nombre positif fini*, et le *plus petit nombre négatif fini* sont : $\tilde{x}_{\max}^{\pm} = \pm(2^{1024} - 2^{971}) = \pm 1,7976931348623157 \times 10^{308}$
- Overflow**: L'intervalle de calcul est $[\tilde{x}_{\max}^-, \tilde{x}_{\max}^+]$. Si un calcul produit un nombre x qui n'est pas dans cet intervalle, on dit qu'il y a *overflow*. La valeur de x est alors mise à $\pm \text{Inf}$. Cela peut arriver au milieu du calcul, même si le résultat final peut être représenté par un nombre machine.
- Le *plus petit nombre positif normalisé* différent de zéro, et le *plus grand nombre négatif normalisé* différent de zéro sont :
 $\tilde{x}_{\min}^{\pm} = \pm 2^{-1022} = \pm 2,2250738585072020 \times 10^{-308}$
- Le système IEEE permet les calculs avec des nombres dénormalisés dans l'intervalle $[\tilde{x}_{\min}^+ * \text{eps}, \tilde{x}_{\min}^+]$.
- Underflow**: Si un calcul produit un nombre positif x qui est plus petit que $\tilde{x}_{\min}^+ * \text{eps}$, on dit qu'il y a *underflow*. Néanmoins, le calcul ne s'arrête pas dans ce cas, il continue avec la valeur de x mise à zéro.

1.3.3 En MATLAB

En MATLAB, les calculs réels sont en double précision par défaut. La fonction `single` peut être utilisée pour convertir les nombres en simple précision. Pour voir la représentation des nombres réels en MATLAB, on peut les afficher au format hexadécimal avec la commande `format hex`.

Le système hexadécimal est celui en base 16 et utilise 16 symboles : 0 à 9 pour représenter les valeurs de 0 à 9, et A, B, C, D, E, F pour représenter les valeurs de 10 à 15. La lecture s'effectue de droite à gauche. La valeur vaut la somme des chiffres affectés de poids correspondant aux puissances successives du nombre 16. Par exemple, $5EB52_{16}$ vaut $2 * 16^0 + 5 * 16^1 + 11 * 16^2 + 14 * 16^3 + 5 * 16^4 = 387922$. Pour passer du binaire au format hexadécimal, c'est facile en regardant la chaîne binaire en groupe de 4 chiffres, et en représentant chaque groupe en un chiffre hexadécimal. Par exemple,

$$\begin{aligned}
 387922 &= 01011110101101010010_2 &= 0101 \quad 1110 \quad 1011 \quad 0101 \quad 0010_2 \\
 & &= 5 \quad E \quad B \quad 5 \quad 2_{16} \\
 & &= 5EB52_{16}
 \end{aligned}$$

La conversion de l'hexadécimal au binaire est le processus inverse.

1.4 Calculs sur les nombres flottants

1.4.1 Erreurs d'arrondi

Si $\tilde{x}$ et $\tilde{y}$ sont deux nombres machine, alors $z = \tilde{x} \times \tilde{y}$ ne correspondra pas en général à un nombre machine puisque le produit demande une quantité double de chiffres. Le résultat sera un nombre machine $\tilde{z}$ proche de z .

On définit l'*erreur absolue* entre un nombre réel x et le nombre machine correspondant $\tilde{x}$ par

$$r_a = |x - \tilde{x}|.$$

L'*erreur relative* entre ces nombres (si $x \neq 0$) est définie par

$$r = \frac{|x - \tilde{x}|}{|x|}.$$

Opérations machine: On désigne par *flop* (de l'anglais *floating operation*) une opération élémentaire à virgule flottante (addition, soustraction, multiplication ou division) de l'ordinateur. Sur les calculateurs actuels on peut s'attendre à la précision suivante, obtenue dans les opérations basiques:

$$\tilde{x} \tilde{\oplus} \tilde{y} = (\tilde{x} \oplus \tilde{y})(1 + r)$$

où $|r| < \epsilon_{ps}$, la précision machine, et $\oplus$ représente l'opération exacte, $\oplus \in \{+, -, *, /\}$ et $\tilde{\oplus}$ représente l'opération de l'ordinateur (*flop*).

1.4.2 Associativité

L'associativité des opérations élémentaires comme par exemple l'addition:

$$(x + y) + z = x + (y + z),$$

n'est plus valide en arithmétique finie. Par exemple, avec 6 chiffres de précision, si on prend les trois nombres

$$x = 1.23456e-3, \quad y = 1.00000e0, \quad z = -y,$$

on obtient $(x + y) + z = (0.00123 + 1.00000e0) - 1.00000e0 = 1.23000e-3$ alors que $x + (y + z) = x = 1.23456e-3$. Il est donc essentiel de considérer l'ordre des opérations et faire attention où l'on met les parenthèses.

1.4.3 Monotonie

Supposons que l'on a une fonction f strictement croissante sur un intervalle $[a, b]$. Peut-on assurer en arithmétique finie que

$$\tilde{x} < \tilde{y} \Rightarrow f(\tilde{x}) < f(\tilde{y})?$$

En général non. Dans la norme IEEE les *fonctions standard* sont implémentées de façon à respecter la monotonie (mais pas la stricte monotonie).

1.4.4 Erreurs d'annulation

Ce sont les erreurs dues à l'annulation numérique de chiffres significatifs, quand les nombres ne sont représentés qu'avec une quantité finie de chiffres, comme les nombres machine. Il est donc important en pratique d'être attentif aux signes dans les expressions, comme l'exemple suivant le montre :

Exemple : on cherche à évaluer sur l'ordinateur de façon précise, pour de petites valeurs de x , la fonction

$$f(x) = \frac{1}{1 - \sqrt{1 - x^2}}.$$

Pour $|x| < \sqrt{\epsilon_{ps}}$, on risque d'avoir le nombre $\sqrt{1 - x^2}$ remplacé par 1, par la machine, et donc lors du calcul de $f(x)$, on risque d'effectuer une division par 0. Par exemple, pour $x = \frac{\sqrt{\epsilon_{ps}}}{2}$, on obtient :

```
>> f=inline('1./(1-sqrt(1-x.^2))','x');
>> f(0.5*sqrt(eps))
```

```
ans =
```

```
Inf
```

On ne peut donc pas évaluer précisément $f(x)$ sur l'ordinateur dans ce cas. Maintenant, si on multiplie dans $f(x)$ le numérateur et le dénominateur par $1 + \sqrt{1 - x^2}$, on obtient

$$f(x) = \frac{1 + \sqrt{1 - x^2}}{x^2}$$

Cette fois, on peut évaluer $f(x)$ de façon précise :

```
>> f=inline('(1+sqrt(1-x.^2))./x.^2','x')
```

```
>> f(0.5*sqrt(eps))
```

```
ans =
```

```
3.6029e+16
```

C'est ce qui se passe dans la cas du calcul de π avec l'algorithme naïf 1.1. En utilisant la formule

$$\sin \frac{\alpha_n}{2} = \sqrt{\frac{1 - \cos \alpha_n}{2}} = \sqrt{\frac{1 - \sqrt{1 - \sin^2 \alpha_n}}{2}},$$

comme $\sin \alpha_n \rightarrow 0$, le numérateur à droite est de la forme

$$1 - \sqrt{1 - \varepsilon^2}, \quad \text{avec } \varepsilon = \sin \alpha_n \text{ petit,}$$

donc sujet aux erreurs d'annulation. Pour y palier, il faut reformuler les équations de façon à s'affranchir des erreurs d'annulations, par exemple en multipliant le numérateur et le dénominateur par $(1 + \sqrt{1 - \sin^2 \alpha_n})$.

$$\sin \frac{\alpha_n}{2} = \sqrt{\frac{1 - \sqrt{1 - \sin^2 \alpha_n}}{2}} = \sqrt{\frac{1 - (1 - \sin^2 \alpha_n)}{2(1 + \sqrt{1 - \sin^2 \alpha_n})}} = \frac{\sin \alpha_n}{\sqrt{2(1 + \sqrt{1 - \sin^2 \alpha_n})}}.$$

On peut alors écrire l'Algorithme 1.2 correspondant au calcul de π avec cette nouvelle formule.

Algorithme 1.2 Algorithme de calcul de π , version stable

1: $s \leftarrow 1, n \leftarrow 4,$	▷ Initialisations
2: Tantque $s > 1e - 10$ faire	▷ Arrêt si $s = \sin(\alpha)$ est petit
3: $s \leftarrow s/\text{sqrt}(2 * (1 - \text{sqrt}(1 - s * s)))$	▷ nouvelle valeur de $\sin(\alpha/2)$
4: $n \leftarrow 2 * n$	▷ nouvelle valeur de n
5: $A \leftarrow (n/2) * s$	▷ nouvelle valeur de l'aire du polygône
6: Fin Tantque	

1.5 Quelques catastrophes dues à l'arithmétique flottante

Il y a un petit nombre "connu" de catastrophes dans la vie réelle qui sont attribuables à une mauvaise gestion de l'arithmétique des ordinateurs (erreurs d'arrondis, d'annulation), voir [6]. Dans le premier exemple ci-dessous cela c'est payé en vies humaines.

n	A_n	$ A_n - \pi $	$\sin(\alpha_n)$
4	2.00000000000000	1.141593e+00	1.000000e+00
8	2.82842712474619	3.131655e-01	7.071068e-01
16	3.06146745892072	8.012519e-02	3.826834e-01
32	3.12144515225805	2.014750e-02	1.950903e-01
64	3.13654849054594	5.044163e-03	9.801714e-02
128	3.14033115695475	1.261497e-03	4.906767e-02
256	3.14127725093277	3.154027e-04	2.454123e-02
512	3.14151380114430	7.885245e-05	1.227154e-02
1024	3.14157294036709	1.971322e-05	6.135885e-03
2048	3.14158772527716	4.928313e-06	3.067957e-03
4096	3.14159142151120	1.232079e-06	1.533980e-03
8192	3.14159234557012	3.080197e-07	7.669903e-04
16384	3.14159257658487	7.700492e-08	3.834952e-04
32768	3.14159263433856	1.925123e-08	1.917476e-04
65536	3.14159264877699	4.812807e-09	9.587380e-05
131072	3.14159265238659	1.203202e-09	4.793690e-05
262144	3.14159265328899	3.008003e-10	2.396845e-05
524288	3.14159265351459	7.519985e-11	1.198422e-05
1048576	3.14159265357099	1.879963e-11	5.992112e-06
2097152	3.14159265358509	4.699352e-12	2.996056e-06
4194304	3.14159265358862	1.174172e-12	1.498028e-06
8388608	3.14159265358950	2.926548e-13	7.490141e-07
16777216	3.14159265358972	7.238654e-14	3.745070e-07
33554432	3.14159265358978	1.731948e-14	1.872535e-07
67108864	3.14159265358979	3.552714e-15	9.362676e-08
134217728	3.14159265358979	0.000000e+00	4.681338e-08
268435456	3.14159265358979	8.881784e-16	2.340669e-08
536870912	3.14159265358979	1.332268e-15	1.170334e-08
1073741824	3.14159265358979	1.332268e-15	5.851672e-09
2147483648	3.14159265358979	1.332268e-15	2.925836e-09
4294967296	3.14159265358979	1.332268e-15	1.462918e-09
8589934592	3.14159265358979	1.332268e-15	7.314590e-10
17179869184	3.14159265358979	1.332268e-15	3.657295e-10
34359738368	3.14159265358979	1.332268e-15	1.828648e-10
68719476736	3.14159265358979	1.332268e-15	9.143238e-11

Table 1.7: Calcul de π avec l'algorithme stable

1 Missile Patriot

En février 1991, pendant la Guerre du Golfe, une batterie américaine de missiles Patriot, à Dharan (Arabie Saoudite), a échoué dans l'interception d'un missile Scud irakien. Le Scud a frappé un baraquement de l'armée américaine et a tué 28 soldats. La commission d'enquête a conclu à un calcul incorrect du temps de parcours, dû à un problème d'arrondi. Les nombres étaient représentés en virgule fixe sur 24 bits, donc 24 chiffres binaires. Le temps était compté par l'horloge interne du système en 1/10 de seconde. Malheureusement, 1/10 n'a pas d'écriture finie dans le système binaire : $1/10 = 0,1$ (dans le système décimal) = 0,000110011001100110011... (dans le système binaire). L'ordinateur de bord arrondissait 1/10 à 24 chiffres, d'où une petite erreur dans le décompte du temps pour chaque 1/10 de seconde. Au moment de l'attaque, la batterie de missile Patriot était allumée depuis environ 100 heures, ce qui avait entraîné une accumulation des erreurs d'arrondi de 0,34 s. Pendant ce temps, un missile Scud parcourt environ 500 m, ce qui explique que le Patriot soit passé à côté de sa cible. Ce qu'il aurait fallu faire c'était redémarrer régulièrement le système de guidage du missile.

14 Explosion d'Ariane 5

Le 4 juin 1996, une fusée Ariane 5, à son premier lancement, a explosé 40 secondes après l'allumage. La fusée et son chargement avaient coûté 500 millions de dollars. La commission d'enquête a rendu son

rapport au bout de deux semaines. Il s'agissait d'une erreur de programmation dans le système inertielle 1
de référence. À un moment donné, un nombre codé en virgule flottante sur 64 bits (qui représentait la 2
vitesse horizontale de la fusée par rapport à la plate-forme de tir) était converti en un entier sur 16 bits. 3
Malheureusement, le nombre en question était plus grand que 32768 (overflow), le plus grand entier que 4
l'on peut coder sur 16 bits, et la conversion a été incorrecte. 5

Bourse de Vancouver 6

Un autre exemple où les erreurs de calcul ont conduit à une erreur notable est le cas de l'indice de la 7
Bourse de Vancouver. En 1982, elle a crée un nouvel indice avec une valeur nominale de 1000. Après 8
chaque transaction boursière, cet indice était recalculé et tronqué après le troisième chiffre décimal et, 9
au bout de 22 mois, la valeur obtenue était 524,881, alors que la valeur correcte était 1098.811. Cette 10
différence s'explique par le fait que toutes les erreurs d'arrondi étaient dans le même sens : l'opération 11
de troncature diminuait à chaque fois la valeur de l'indice. 12

Chapitre 2

Langage algorithmique

2.1 Pseudo-langage algorithmique

2

Pour uniformiser l'écriture des algorithmes nous employons un pseudo-langage contenant l'indispensable :

3

- variables,
- opérateurs (arithmétiques, relationnels, logiques),
- expressions,
- instructions (simples et composées),
- fonctions.

5

6

7

8

9

Ce pseudo-langage sera de fait très proche du langage de programmation de Matlab.

10

2.1.1 Données et constantes

11

Une donnée est une valeur introduite par l'utilisateur (par ex. une température, une vitesse, ...). Une constante est un symbole ou un identificateur non modifiable (par ex. π , la constante de gravitation,...).

12

13

2.1.2 Variables

14

Definition 2.1

Une variable est un objet dont la valeur est modifiable, qui possède un nom et un type (entier, caractère, réel, complexe, ...). Elle est rangée en mémoire à partir d'une certaine adresse.

15

2.1.3 Opérateurs

Opérateurs arithmétiques

Nom	Symbole	exemple
addition	+	$a + b$
soustraction	−	$a - b$
opposé	−	$-a$
produit	*	$a * b$
division	/	a/b
puissance a^b	^	a^b

Table 2.1: Opérateurs arithmétiques

Opérateurs relationnels

Nom	Symbole	exemple	Commentaires
identique	==	$a == b$	vrai si a et b ont même valeur, faux sinon.
différent	~=	$a ~= b$	faux si a et b ont même valeur, vrai sinon.
inférieur	<	$a < b$	vrai si a est plus petit que b , faux sinon.
supérieur	>	$a > b$	vrai si a est plus grand que b , faux sinon.
inférieur ou égal	<=	$a <= b$	vrai si a est plus petit ou égal à b , faux sinon.
supérieur ou égal	>=	$a >= b$	vrai si a est plus grand ou égal à b , faux sinon.

Table 2.2: Opérateurs relationnels

Opérateurs logiques

Nom	Symbole	exemple	Commentaires
négation	~	$\sim a$	vrai si a est faux (ou nul), faux sinon.
ou		$a b$	vrai si a ou b est vrai (non nul), faux sinon.
et	&	$a\&b$	vrai si a et b sont vrais (non nul), faux sinon.

Table 2.3: Opérateurs logiques

Opérateur d'affectation

Nom	Symbole	exemple	Commentaires
affectation	←	$a \leftarrow b$	On affecte à la variable a le contenu de b

Table 2.4: Opérateurs d'affectation

2.1.4 Expressions

♥ Definition 2.2

Une expression est un groupe d'opérandes (i.e. nombres, constantes, variables, ...) liées par certains opérateurs pour former un terme algébrique qui représente une valeur (i.e. un élément de donnée simple)

Exemple 2.3 • Voici un exemple classique d'expression numérique :

$$(b * b - 4 * a * c) / (2 * a).$$

On appelle **opérandes** les identifiants a , b et c , et les nombres 4 et 2. Les symboles $*$, $-$ et $/$ sont les **opérateurs**.

- Voici un exemple classique d'expression booléenne (logique) :

$$(x < 3.14)$$

On teste $(x < 3.14)$ avec x une variable numérique réelle. si $x < 3.14$ alors cette expression renverra la valeur *vraie* (i.e. 1), *faux* sinon (i.e. 0).

2.1.5 Instructions

♥ Definition 2.4

Une **instruction** est un ordre ou un groupe d'ordres qui déclenche l'exécution de certaines actions par l'ordinateur. Il y a deux types d'instructions : simple et structuré.
Les **instructions simples** sont essentiellement des ordres seuls et inconditionnels réalisant l'une des tâches suivantes :

1. affectation d'une valeur a une variable.
2. appel d'une fonction (procedure, subroutine, ... suivant les langages).

Les **instructions structurées** sont essentiellement :

1. les instructions composées, groupe de plusieurs instructions simples,
2. les instructions répétitives, permettant l'exécution répétée d'instructions simples, (i.e. boucles «pour», «tant que»)
3. les instructions conditionnelles, lesquelles ne sont exécutées que si une certaine condition est respectée (i.e. «si»)

Les exemples qui suivent sont écrits dans un pseudo langage algorithmique mais sont facilement transposable dans la plupart des langages de programmation.

Instructions simples

Voici un exemple de l'*instruction simple d'affectation* :

1: $a \leftarrow 3.14 * R$

On évalue l'expression $3.14 * R$ et affecte le résultat à la variable a .
Un autre exemple est donné par l'*instruction simple d'affichage* :

affiche('bonjour')

Affiche la chaine de caractères 'bonjour' à l'écran. Cette instruction fait appel à la fonction affiche.

Instructions composées

Instructions répétitives, boucle «pour»

Algorithme 2.1 Exemple de boucle «pour»

Données : n : un entier.

- 1: $S \leftarrow 0$
- 2: **Pour** $i \leftarrow 1$ à n **faire**
- 3: $S \leftarrow S + \cos(i^2)$
- 4: **Fin Pour**

1 Instruction répétitive, boucle «tant que»

Algorithme 2.2 Exemple de boucle «tant que»

```

1:  $i \leftarrow 0, x \leftarrow 1$ 
2: Tantque  $i < 1000$  faire
3:    $x \leftarrow x + i * i$ 
4:    $i \leftarrow i + 1$ 
5: Fin Tantque

```

2 Instruction répétitive, boucle «répéter ...jusqu'à»

Algorithme 2.3 Exemple de boucle «répéter ...jusqu'à»

```

1:  $i \leftarrow 0, x \leftarrow 1$ 
2: Répéter
3:    $x \leftarrow x + i * i$ 
4:    $i \leftarrow i + 1$ 
5: jusqu'à  $i \geq 1000$ 

```

3 Instructions conditionnelles «si»

Algorithme 2.4 Exemple d'instructions conditionnelle «si»

Données : *age* : un réel.

```

1: Si  $age \geq 18$  alors
2:   affiche('majeur')
3: Sinon Si  $age \geq 0$  alors
4:   affiche('mineur')
5: Sinon
6:   affiche('en devenir')
7: Fin Si

```

4 2.1.6 Fonctions

5 Les fonctions permettent

- 6 • d'automatiser certaines tâches répétitives au sein d'un même programme,
- 7 • d'ajouter à la clarté d'un programme,
- 8 • l'utilisation de portion de code dans un autre programme,
- 9 • ...

10 Fonctions prédéfinies

11 Pour faciliter leur usage, tous les langages de programmation possèdent des fonctions prédéfinies. On
 12 pourra donc supposer que dans notre langage algorithmique un grand nombre de fonctions soient prédéfinies
 13 : par exemple, les fonctions mathématiques sin, cos, exp, abs, ... (pour ne citer qu'elles)

14 Syntaxe

15 On utilise la syntaxe suivante pour la définition d'une fonction

16

Fonction $[args_1, \dots, args_n] \leftarrow \text{NOMFONCTION}(arge_1, \dots, arge_m)$
instructions
Fin Fonction

La fonction se nomme **NOMFONCTION** . Elle admet comme paramètres d'entrée (données) les m arguments $arge_1, \dots, arge_m$ et comme paramètres de sortie (résultats) les n arguments $args_1, \dots, args_n$. Ces derniers doivent être déterminés dans le corps de la fonction (partie instructions).
Dans le cas ou la fonction n'admet qu'un seul paramètre de sortie, l'écriture se simplifie :

Fonction $args \leftarrow \text{NOMFONCTION}(arge_1, \dots, arge_m)$
instructions
Fin Fonction

Ecrire ses propres fonctions

Pour écrire une fonction «propre», il faut tout d'abord déterminer exactement ce que devra faire cette fonction.

Puis, il faut pouvoir répondre à quelques questions :

- 1. Quelles sont les données (avec leurs limitations)?
- 2. Que doit-on calculer ?

Et, ensuite il faut la **commenter** : expliquer son usage, type des paramètres,

Exemple : résolution d'une équation du premier degré

Nous voulons écrire une fonction calculant la solution de l'équation

$$ax + b = 0,$$

où nous supposons que $a \in \mathbb{R}^*$ et $b \in \mathbb{R}$. La solution de ce problème est donc

$$x = -\frac{b}{a}.$$

Les données de cette fonction sont $a \in \mathbb{R}^*$ et $b \in \mathbb{R}$. Elle doit retourner $x = -\frac{b}{a}$ solution de $ax + b = 0$.

Algorithme 2.5 Exemple de fonction : Résolution de l'équation du premier degré $ax + b = 0$.

Données : a : nombre réel différent de 0
 b : nombre réel.

Résultat : x : un réel.

1: **Fonction** $x \leftarrow \text{REPD}(a, b)$
2: $x \leftarrow -b/a$
3: **Fin Fonction**

Remarque 2.5 Cette fonction est très simple, toutefois pour ne pas «alourdir» le code nous n'avons pas vérifié la validité des données fournies.

 **Exercice 2.1.1**
Ecrire un algorithme permettant de valider cette fonction.

1 Exemple : résolution d'une équation du second degré

2 Nous cherchons les solutions réelles de l'équation

$$ax^2 + bx + c = 0, \quad (2.1)$$

3 où nous supposons que $a \in \mathbb{R}^*$, $b \in \mathbb{R}$ et $c \in \mathbb{R}$ sont donnés.

4 Mathématiquement, l'étude des solutions réelles de cette équation nous amène à envisager trois cas
5 suivant les valeurs du discriminant $\Delta = b^2 - 4ac$

- 6 • si $\Delta < 0$ alors les deux solutions sont complexes,
- 7 • si $\Delta = 0$ alors la solution est $x = -\frac{b}{2a}$,
- 8 • si $\Delta > 0$ alors les deux solutions sont $x_1 = \frac{-b-\sqrt{\Delta}}{2*a}$ et $x_2 = \frac{-b+\sqrt{\Delta}}{2*a}$.



Exercice 2.1.2

1. Ecrire la fonction **discriminant** permettant de calculer le discriminant de l'équation (2.1).
2. Ecrire la fonction **RESD** permettant de résoudre l'équation (2.1) en utilisant la fonction **discriminant**.
3. Ecrire un programme permettant de valider ces deux fonctions.



Exercice 2.1.3

Même question que précédemment dans le cas complexe (solutions et coefficients).

2.2 Méthodologie d'élaboration d'un algorithme

2.2.1 Description du problème

- Spécification d'un ensemble de données
Origine : énoncé, hypothèses, sources externes, ...
- Spécification d'un ensemble de buts à atteindre
Origine : résultats, opérations à effectuer, ...
- Spécification des contraintes

2.2.2 Recherche d'une méthode de résolution

- Clarifier l'énoncé.
- Simplifier le problème.
- Ne pas chercher à le traiter directement dans sa globalité.
- S'assurer que le problème est soluble (sinon problème d'indécidabilité!)
- Recherche d'une stratégie de construction de l'algorithme
- Décomposer le problème en sous problèmes partiels plus simples : raffinement.
- Effectuer des raffinements successifs.
- Le niveau de raffinement le plus élémentaire est celui des instructions.

2.2.3

Réalisation d'un algorithme

1

Il doit être conçu indépendamment du langage de programmation et du système informatique (sauf cas très particulier)

23

- L'algorithme doit être exécuté en un nombre fini d'opérations.4
- L'algorithme doit être spécifié clairement, sans la moindre ambiguïté.5
- Le type de données doit être précisé.6
- L'algorithme doit fournir au moins un résultat.7
- L'algorithme doit être effectif : toutes les opérations doivent pouvoir être simulées par un homme en temps fini.89

Pour écrire un algorithme détaillé, il faut tout d'abord savoir répondre a quelques questions :

1011

- Que doit-il faire ? (i.e. Quel problème est-il censé résoudre?)12
- Quelles sont les données nécessaires à la résolution de ce problème?13
- Comment résoudre ce problème «à la main» (sur papier)?14

Si l'on ne sait pas répondre à l'une de ces questions, l'écriture de l'algorithme est fortement compromise.

15

2.2.4

Exercices

16



Exercice 2.2.1: Algorithme pour une somme

Ecrire un algorithme permettant de calculer

$$S(x) = \sum_{k=1}^n k \sin(2 * k * x)$$

Correction Exercice 2.2.1 L'énoncé de cet exercice est imprécis. On choisit alors $x \in \mathbb{R}$ et $n \in \mathbb{N}$ pour rendre possible le calcul. Le problème est donc de calculer

17

$$\sum_{k=1}^n k \sin(2kx).$$

Toutefois, on aurait pu choisir $x \in \mathbb{C}$ ou encore un tout autre problème :

Trouver $x \in \mathbb{R}$ tel que $S(x) = \sum_{k=1}^n k \sin(2kx)$

où $n \in \mathbb{N}$ et S , fonction de $\mathbb{R}$ à valeurs réelles, sont les données!

18

Algorithme 2.6 Calcul de $S = \sum_{k=1}^n k \sin(2kx)$

Données : x : nombre réel,
 n : nombre entier.

Résultat : S : un réel.

- 1: $S \leftarrow 0$
- 2: **Pour** $k \leftarrow 1$ à n **faire**
- 3: $S \leftarrow S + k * \sin(2 * k * x)$
- 4: **Fin Pour**

**Exercice 2.2.2: Algorithme pour un produit**

Ecrire un algorithme permettant de calculer

$$P(z) = \prod_{n=1}^k \sin(2 * k * z/n)^k$$

1

2 **Correction Exercice 2.2.2** L'énoncé de cet exercice est imprécis. On choisit alors $z \in \mathbb{R}$ et $k \in \mathbb{N}$ pour
3 rendre possible le calcul.

Algorithme 2.7 Calcul de $P = \prod_{n=1}^k \sin(2kz/n)^k$

Données : z : nombre réel,
 k : nombre entier.

Résultat : P : un réel.

1: $P \leftarrow 1$
2: **Pour** $n \leftarrow 1$ à k **faire**
3: $P \leftarrow P * \sin(2 * k * z/n)^k$
4: **Fin Pour**

4

◇

5

**Exercice 2.2.3: Série de Fourier**

Soit la série de Fourier

$$x(t) = \frac{4A}{\pi} \left\{ \cos \omega t - \frac{1}{3} \cos 3\omega t + \frac{1}{5} \cos 5\omega t - \frac{1}{7} \cos 7\omega t + \dots \right\}.$$

6

Ecrire la fonction SFT permettant de calculer $x_n(t)$.

Correction Exercice 2.2.3 Nous devons écrire la fonction permettant de calculer

$$x_n(t) = \frac{4A}{\pi} \sum_{k=1}^n (-1)^{k+1} \frac{1}{2k-1} \cos((2k-1)\omega t)$$

7 Les données de la fonction sont $A \in \mathbb{R}$, $\omega \in \mathbb{R}$, $n \in \mathbb{N}^*$ et $t \in \mathbb{R}$.

8 Grâce à ces renseignements nous pouvons déjà écrire l'entête de la fonction :

Algorithme 2.8 En-tête de la fonction SFT retournant valeur de la série de Fourier en t tronquée au n premiers termes de l'exercice 2.2.3.

Données : t : nombre réel,
 n : nombre entier strictement positif
 A, ω : deux nombres réels.

Résultat : x : un réel.

1: **Fonction** $x \leftarrow \text{SFT}(t, n, A, \omega)$
2: ...
3: **Fin Fonction**

9 Maintenant nous pouvons écrire progressivement l'algorithme pour aboutir au final à une version ne
10 contenant que des opérations élémentaires.

Algorithme 2.9 $\mathcal{R}_0$

1: $x \leftarrow \frac{4A}{\pi} \sum_{k=1}^n \left(\frac{(-1)^{k+1} \frac{1}{2k-1} \times}{\cos((2k-1)\omega t)} \right)$

Algorithme 2.9 $\mathcal{R}_1$

1: $S \leftarrow \sum_{k=1}^n (-1)^{k+1} \frac{1}{2k-1} \cos((2k-1)\omega t)$
 2: $x \leftarrow \frac{4A}{\pi} S$

Algorithme 2.9 $\mathcal{R}_1$

1: $S \leftarrow \sum_{k=1}^n (-1)^{k+1} \frac{1}{2k-1} \cos((2k-1)\omega t)$
 2: $x_n(t) \leftarrow \frac{4A}{\pi} S$

Algorithme 2.9 $\mathcal{R}_2$

1: $S \leftarrow 0$
 2: **Pour** $k = 1$ à n **faire**
 3: $S \leftarrow S + (-1)^{k+1} \frac{1}{2k-1} * \cos((2k-1)\omega t)$
 4: **Fin Pour**
 5: $x_n(t) \leftarrow \frac{4A}{\pi} S$

Finalement la fonction est

Algorithme 2.9 Fonction SFT retournant la valeur de la série de Fourier en t tronquée au n premiers termes de l'exercice ??.

Données : t : nombre réel,
 n : nombre entier strictement positif
 A, ω : deux nombres réels.

Résultat : x : un réel.

1: **Fonction** $x \leftarrow \text{SFT}(t, n, A, \omega)$
 2: $S \leftarrow 0$
 3: **Pour** $k = 1$ à n **faire**
 4: $S \leftarrow S + ((-1)^{(k+1)}) * \cos((2 * k - 1) * \omega * t) / (2 * k - 1)$
 5: **Fin Pour**
 6: $S \leftarrow 4 * A * S / \pi$
 7: **Fin Fonction**

**Exercice 2.2.4**

Reprendre les trois exercices précédents en utilisant les boucles «tant que».

2.3 Principes de «bonne» programmation pour attaquer de «gros» problèmes

Tous les exemples vus sont assez courts. Cependant, il peut arriver que l'on ait des programmes plus longs à écrire (milliers de lignes, voir des dizaines de milliers de lignes). Dans l'industrie, il arrive que des équipes produisent des codes de millions de lignes, dont certains mettent en jeux des vies humaines (contrôler un avion de ligne, une centrale nucléaire, ...). Le problème est évidemment d'écrire des programmes sûrs. Or, *un programme à 100% sûr, cela n'existe pas!* Cependant, plus un programme est simple, moins le risque d'erreur est grand : c'est sur cette remarque de bon sens que se basent les «bonnes» méthodes. Ainsi :

Tout problème compliqué doit être découpé en sous-problèmes plus simples

Il s'agit, lorsqu'on a un problème P à résoudre, de l'analyser et de le décomposer en un ensemble de problèmes $P_1, P_2, P_3, \dots$ plus simples. Puis, P_1 , est lui-même analysé et décomposé en $P_{11}, P_{12}, \dots$, et

- 1 P_2 en P_{21}, P_{22} , etc. On poursuit cette analyse jusqu'à ce qu'on n'ait plus que des problèmes élémentaires
- 2 à résoudre. Chacun de ces problèmes élémentaires est donc traité séparément dans un module, c'est à
- 3 dire un morceau de programme relativement indépendant du reste. Chaque module sera *testé et validé*
- 4 séparément dans la mesure du possible et naturellement *largement documenté*. Enfin ces modules
- 5 élémentaires sont assemblés en modules de plus en plus complexes, jusqu'à remonter au problème initiale.
- 6 A chaque niveau, il sera important de bien réaliser les phases de test, validation et documentation des
- 7 modules.

Par la suite, on s'évertue à écrire des algorithmes!
Ceux-ci ne seront pas optimisés^a!

^aaméliorés pour minimiser le nombre d'opérations élémentaires,
l'occupation mémoire, ..

8

Chapitre 3

Résolution de systèmes non linéaires

Un problème *simple* comme la recherche des zéros/racines d'un polynôme n'est pas ... *simple*. Depuis tout petit, on sait trouver les racines d'un polynôme de degré 2 : $a_2x^2 + a_1x + a_0 = 0$. Quid des racines de polynômes de degré plus élevé?



- **degré 2** : Babyloniens en 1600 avant J.-C.

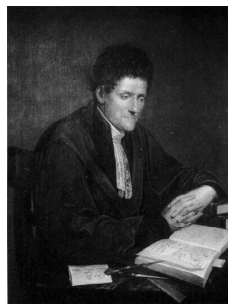
- **degré 3** : *Scipio del Ferro* (1465-1526, mathématicien italien) et *Niccolo Fontana* (1499-1557, mathématicien italien)

- **degré 4** : *Ludovico Ferrari* (1522-1565, mathématicien italien)

- **degré 5** : *Paolo Ruffini* (1765-1822, mathématicien italien) en 1799, *Niels Henrik Abel* (1802-1829, mathématicien norvégien) en 1824, montrent qu'il n'existe **pas de solution analytique**.



(a) *Niccolo Fontana* 1499-1557, (b) *Paolo Ruffini* 1765-1822, (c) *Niels Henrik Abel* 1802-1829, mathématicien italien



L'objectif de ce chapitre est de rechercher **numériquement** les *zéros/racines* d'une fonction ou d'un système d'équations lorsqu'ils existent :

- Soit $f : [a, b] \subset \mathbb{R} \longrightarrow \mathbb{R}$, trouver $x \in [a, b]$ tel que $f(x) = 0$, étudier en section 3.1

- Soit $f : \Omega \subset \mathbb{K}^n \longrightarrow \mathbb{K}^n$, trouver $x \in \Omega$ tel que $f(x) = 0$ ($\mathbb{K} = \mathbb{R}$ ou $\mathbb{K} = \mathbb{C}$). étudier en section 3.3

3.1 Recherche des zéros d'une fonction

Soit $f : [a, b] \subset \mathbb{R} \longrightarrow \mathbb{R}$ une fonction continue. On cherche à déterminer les zéros de f ou plus précisément l'ensemble des $x \in [a, b]$ tels que $f(x) = 0$.

Dans un premier temps, on étudiera la **méthode de dichotomie** qui est *assez naturelle*. Puis on étudiera plusieurs algorithmes liés à la méthode du point fixe.

3.1.1 Méthode de dichotomie ou de bisection

Principe et résultats

 **principe de la méthode de dichotomie** : Soit I un intervalle contenant un unique zéro de la fonction f , on le divise par son milieu en deux intervalles et on détermine lequel des deux contient le zéro. On itère ce processus sur le nouvel intervalle.

Plus précisément, on suppose que la fonction f vérifie $f(a)f(b) < 0$ et qu'il existe un unique $\xi \in]a, b[$ tel que $f(\xi) = 0$.

On note $I^{(0)} =]a, b[$ et $x_0 = (a+b)/2$. Si $f(x_0) = 0$ alors on a fini! Supposons $f(x_0) \neq 0$, alors ξ appartient à $]a, x_0[$ si $f(a)f(x_0) < 0$, sinon il appartient à $]x_0, b[$. On vient donc de déterminer une méthode permettant de diviser par 2 la longueur de l'intervalle de recherche de ξ . On peut bien évidemment itérer ce principe en définissant les trois suites $(a_k)_{k \in \mathbb{N}}$, $(b_k)_{k \in \mathbb{N}}$ et $(x_k)_{k \in \mathbb{N}}$ par

$$\bullet a_0 = a, b_0 = b \text{ et } x_0 = \frac{a+b}{2},$$

$$\bullet \forall k \in \mathbb{N}, x_k = (a_k + b_k)/2, \text{ et}$$

$$\begin{cases} a_{k+1} = b_{k+1} = x_k & \text{si } f(x_k) = 0, \\ a_{k+1} = x_k, b_{k+1} = b_k & \text{si } f(a_k)f(x_k) < 0, \\ a_{k+1} = a_k, b_{k+1} = x_k & \text{sinon (i.e. } f(a_k)f(x_k) < 0.) \end{cases}$$

En Figure 3.2, on représente les intervalles successifs obtenus par la méthode de dichotomie pour $a = 0.1$, $b = 2.4$ et $f : x \mapsto (x+2)(x+1)(x-1)$.

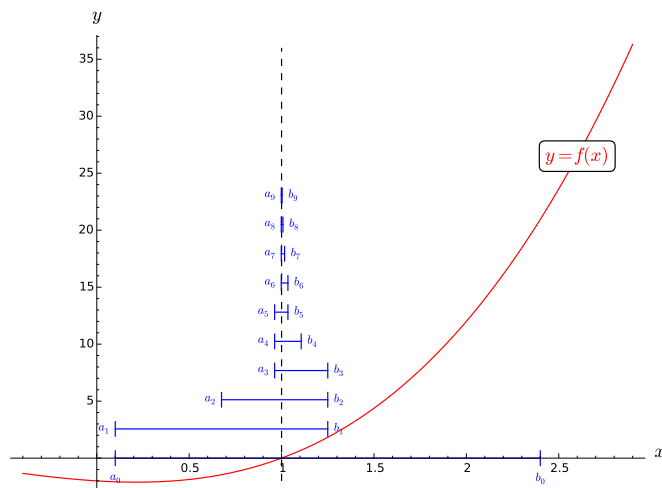


Figure 3.2: Méthode de dichotomie: $f(x) = (x+2)(x+1)(x-1)$,



Exercice 3.1.1

On suppose que la fonction f est continue sur $[a, b]$, vérifie $f(a)f(b) < 0$ et qu'il existe un unique $\xi \in]a, b[$ tel que $f(\xi) = 0$.

Q. 1 1. Montrer que les suites (a_k) et (b_k) convergent vers α .

2. En déduire que la suite (x_k) converge vers α .

Q. 2 1. Montrer que pour tout $k \in \mathbb{N}$, $|x_k - \alpha| \leq \frac{b-a}{2^{k+1}}$.

2. Soit $\epsilon > 0$. En déduire que si $k \geq \frac{\log(\frac{b-a}{\epsilon})}{\log(2)} - 1$ alors $|x_k - \alpha| \leq \epsilon$.

Correction Exercice 3.1.1

Q. 1 1. Supposons qu'il existe $k \in \mathbb{N}$ tel que $f(x_k) = 0$, (i.e. $x_k = \alpha$ car $x_k \in [a, b]$) alors par construction $a_{k+i} = b_{k+i} = x_{k+i} = \alpha$ pour tout $i \in \mathbb{N}^*$. Ceci assure la convergence des 3 suites vers α .

Supposons maintenant que $\forall k \in \mathbb{N}$, $f(x_k) \neq 0$. Par construction, nous avons $a_k \leq a_{k+1} \leq b$, $a \leq b_{k+1} \leq b_k$ et $a_k \leq b_k$. La suite (a_k) est convergente car elle est croissante et majorée. La suite (b_k) est décroissante et minorée : elle est donc convergente. De plus $0 \leq b_k - a_k \leq \frac{b_{k-1} - a_{k-1}}{2}$ et donc $0 \leq b_k - a_k \leq \frac{b-a}{2^{k+1}}$. On en déduit que les suites (a_k) et (b_k) ont même limite. Comme par construction, $\forall k \in \mathbb{N}$, $\alpha \in [a_k, b_k]$ ceci entraîne que α est la limite de ces 2 suites.

2. Par construction, $\forall k \in \mathbb{N}$, $a_k \leq x_k \leq b_k$ et comme les suites (a_k) et (b_k) convergent vers α , la suite (x_k) converge aussi vers α .

Q. 2 1. On a $\forall k \in \mathbb{N}$, $a_k \leq x_k \leq b_k$ et $a_k \leq \alpha \leq b_k$ d'où $|x_k - \alpha| \leq b_k - a_k$. Ce qui donne

$$|x_k - \alpha| \leq b_k - a_k \leq \frac{b_{k-1} - a_{k-1}}{2} \leq \frac{b-a}{2^{k+1}}.$$

2. Pour avoir $|x_k - \alpha| \leq \epsilon$, il suffit d'avoir $\frac{b-a}{2^{k+1}} \leq \epsilon$, et donc $k \geq \frac{\log(\frac{b-a}{\epsilon})}{\log(2)} - 1$.

◇

On peut alors écrire la proposition suivante :



Proposition 3.1

Soit $f : [a, b] \subset \mathbb{R} \rightarrow \mathbb{R}$ une fonction continue vérifiant $f(a)f(b) < 0$ et admettant $\alpha \in]a, b[$ comme **unique** solution de $f(x) = 0$. Alors la suite $(x_k)_{k \in \mathbb{N}}$ définie par la méthode de dichotomie converge vers α et

$$|x_k - \alpha| \leq \frac{b-a}{2^{k+1}}, \quad \forall k \in \mathbb{N}.$$

On a alors $\forall \epsilon > 0$, $\forall k \geq \frac{\log(\frac{b-a}{\epsilon})}{\log(2)} - 1$

$$|x_k - \alpha| \leq \epsilon.$$

A partir de ce résultat, on va proposer plusieurs variantes de l'algorithme de dichotomie.

On note tout d'abord que pour déterminer α il faut calculer la limite d'une suite et donc une infinité de termes! Numériquement et algorithmiquement, on se limite donc à déterminer une approximation de α . La proposition 3.1 permet d'obtenir une approximation α_ϵ de α en un nombre fini d'itérations avec une précision ϵ donnée: on choisit $\alpha_\epsilon = x_{k_{\min}}$ avec $k_{\min} = E(\frac{\log(\frac{b-a}{\epsilon})}{\log(2)})$ où $E(\cdot)$ est la fonction *partie entière*.

Algorithmique

Tout d'abord nous allons poser "correctement" le problème pour lever toute ambiguïté. En effet, si le problème est *rechercher une racine de f sur l'intervalle $[a, b]$ par la méthode de dichotomie*, dois-je uniquement rechercher une approximation d'une racine ou calculer la suite (x_k) ou l'ensemble des suites (x_k) , (a_k) et (b_k) ? C'est ce que nous appellerons le(s) **résultat(s)/objectif(s)** de l'algorithme. L'écriture

- 1 d'un algorithme est fortement corrélée aux objectifs souhaités.
 2 Sauf note contraire, on choisira dans la suite comme objectif de déterminer une approximation de la
 3 racine α de f :

4 **Résultat :** α_ϵ : un réel tel que $|\alpha_\epsilon - \alpha| \leq \epsilon$.

- 5 Ensuite, pour aboutir à cet objectif on détermine les données (avec hypothèses) nécessaires et suffi-
 6 santes :

7 **Données :** a, b : deux réels $a < b$,
 f : $f : [a, b] \subset \mathbb{R} \rightarrow \mathbb{R}$ vérifiant les hypothèses de la proposition 3.1,
 ϵ : un réel strictement positif.

On peut alors noter (formellement pour le moment) **DICHOTOMIE** la fonction permettant, à partir des données f, a, b , et ϵ , de trouver α_ϵ . Cette fonction algorithmique aura la syntaxe suivante :

$$\alpha_\epsilon \leftarrow \text{DICHOTOMIE}(f, a, b, \epsilon)$$

8 Nous allons maintenant proposer une manière d'aborder l'écriture de cette fonction dans un langage
 9 algorithmique présenté dans []. On doit s'affranchir de tout (ou presque) symbolismes mathématiques
 10 pour s'approcher au plus près des langages de programmation. Par exemple, la notation mathématique
 11 $(x_k)_{k=0}^{10}$ pour noter les 11 premiers itérés d'une suite n'est pas directement disponible dans les langages
 12 de programmations : on peut par exemple passer par un tableau $\mathbf{X}$ de dimension 11 (au moins) pour
 13 stocker les valeurs de la suite. Suivant les langages, l'accès aux composantes d'un tableau diffère :
 14 sous Matlab/Octave l'accès au 1er élément d'un tableau $\mathbf{X}$ se fait par $\mathbf{X}(1)$ et en C/C++/Python par
 15 $\mathbf{X}[0]$. Lors de l'écriture d'un algorithme, nous utiliserons l'opérateur $()$ ou (exclusif) $[]$ pour accéder
 16 aux différentes composantes d'un tableau : $()$ si le 1er élément est d'indice 1 ou $[]$ si le 1er élément est
 17 d'indice 0.

18 Pour écrire un algorithme non trivial, nous utiliserons une *technique de raffinements d'algorithme*
 19 (voir []) : par raffinements successifs, nous allons écrire des algorithmes **équivalents** pour aboutir au
 20 final à un algorithme n'utilisant que

- 21 • des instructions élémentaires (affectation, addition, somme, ...)
- 22 • des instructions composées (conditionnelle, boucles pour, tantque, ...)
- 23 • des fonctions usuelles, mathématiques (**cos** , **exp** , **E** partie entière, ...), entrées/sorties , ...

24 présentent dans tous les langages.

25 L'idée est donc de partir d'un algorithme formel facile à comprendre puis de le détailler de plus en plus
 26 au fur et à mesure des raffinements. Le passage entre deux raffinements successifs doit être simple.

27 Un petit rappel pour les algorithmes qui suivent : les données sont supposées ... données!

Algorithme 3.1 $\mathcal{R}_0$	Algorithme 3.1 $\mathcal{R}_1$
1: $k_{\min} \leftarrow \mathbf{E}(\frac{\log(\frac{b-a}{\epsilon})}{\log(2)}) \triangleright \mathbf{E}$, partie entière	1: $k_{\min} \leftarrow \mathbf{E}(\frac{\log(\frac{b-a}{\epsilon})}{\log(2)}) \triangleright \mathbf{E}$, partie entière
2: Calcul de la suite $(x_k)_{k=0}^{k_{\min}}$ par dichotomie	2: Initialisation de x_0
3: $\alpha_\epsilon \leftarrow x_{k_{\min}}$	3: Pour $k \leftarrow 0$ à $k_{\min} - 1$ faire
	4: Calcul de la suite (x_{k+1}) par dichotomie
	5: Fin Pour
	6: $\alpha_\epsilon \leftarrow x_{k_{\min}}$

28 Entre les raffinements $\mathcal{R}_0$ et $\mathcal{R}_1$, nous avons juste décrit le principe de calcul de toute suite récurrente
 29 d'ordre 1.

30 Pour faciliter la lecture, nous rappelons la méthode de dichotomie :

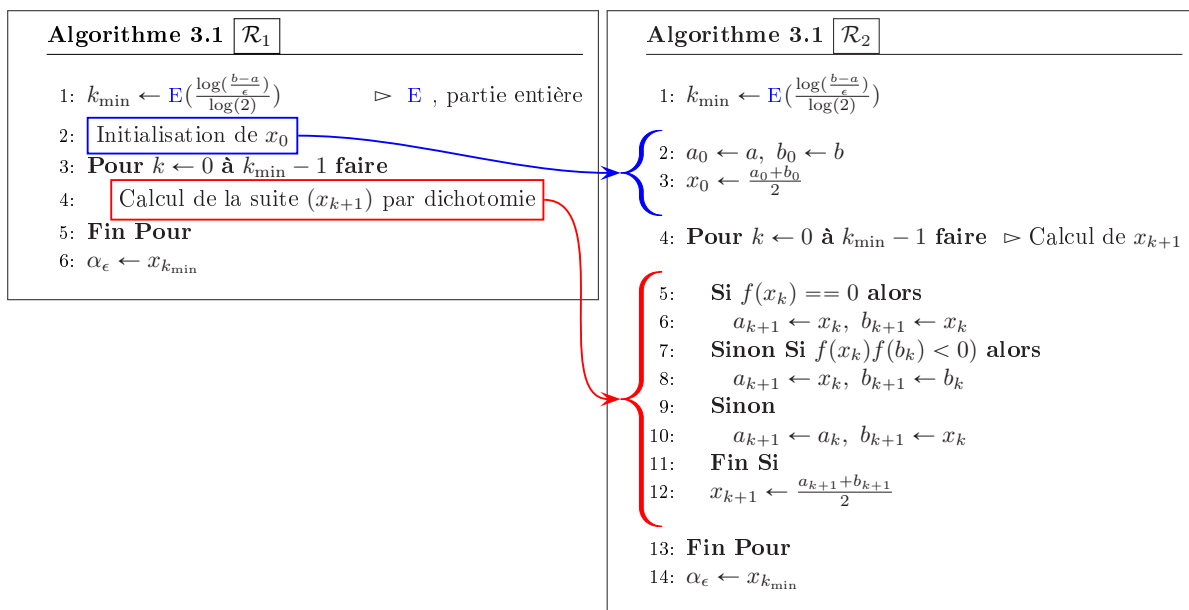
- 31 Pour faciliter la lecture, nous rappelons la méthode de dichotomie :
- 32 • $a_0 = a, b_0 = b$ et $x_0 = \frac{a+b}{2}$,
 - $\forall k \in \llbracket 0, k_{\min} - 1 \rrbracket$,

$$\begin{cases} a_{k+1} = b_{k+1} = x_k & \text{si } f(x_k) = 0, \\ a_{k+1} = x_k, b_{k+1} = b_k & \text{si } f(b_k)f(x_k) < 0, \\ a_{k+1} = a_k, b_{k+1} = x_k & \text{sinon (i.e. } f(a_k)f(x_k) < 0.) \end{cases}$$

et

$$x_{k+1} = \frac{a_{k+1} + b_{k+1}}{2}$$

Ecrit sous cette forme, le calcul de la suite (x_k) nécessite le calcul simultané des suites (a_k) et (b_k) .



La transcription de ce raffinement dans un langage de programmation n'est pas forcément triviale pour tout le monde. Quid de $a_k, \epsilon, \dots$ qui sont syntaxiquement incorrectes dans un langage de programmation? Le cas du ϵ est très simple : les lettres grecques étant prohibées, on la remplacera par exemple par **eps**. Pour les suites (a_k) , (b_k) et (x_k) , tronquées à $k_{\min}$, on pourra utiliser des tableaux (vecteurs) de $k_{\min} + 1$ réels notés **A**, **B** et **X** qui contiendront respectivement l'ensemble des valeurs a_k, b_k et x_k pour $0 \leq k \leq k_{\min}$. Plus précisément nous pourrions utiliser les opérateurs $()$ (indice des tableaux commence à 1) ou $[]$ (indice des tableaux commence à 0) pour accéder aux différents éléments des tableaux et nous aurons par convention $\mathbf{A}(k+1) = \mathbf{A}[k] = a_k, \forall k \in \llbracket 0, k_{\min} \rrbracket$. Par la suite nous utiliserons les opérateurs $()$ et nous aurons alors

$$\forall k \in \llbracket 0, k_{\min} \rrbracket, \mathbf{A}(k+1) = a_k, \mathbf{B}(k+1) = b_k \text{ et } \mathbf{X}(k+1) = x_k.$$

En utilisant ces changements de notations nous obtenons (enfin) l'Algorithme 3.1.

Algorithme 3.1 Méthode de dichotomie : version 1

Données : a, b : deux réels $a < b$,
 f : $f : [a, b] \subset \mathbb{R} \rightarrow \mathbb{R}$ vérifiant
les hypothèses de la proposition 3.1,
 eps : un réel strictement positif.

Résultat : x : un réel tel que $|x - \alpha| \leq \text{eps}$.

```

1: Fonction  $x \leftarrow \text{DICHOTOMIE1}(f, a, b, \text{eps})$ 
2:    $k_{\min} \leftarrow \lceil \log((b-a)/\text{eps}) / \log(2) \rceil$ 
3:    $A, B, X \in \mathbb{R}^{k_{\min}+1}$   $\triangleright A(k+1)$  contiendra  $a_k, \dots$ 
4:    $A(1) \leftarrow a, B(1) \leftarrow b, X(1) \leftarrow (a+b)/2$ 
5:   Pour  $k \leftarrow 1$  à  $k_{\min}$  faire
6:     Si  $f(X(k)) == 0$  alors
7:        $A(k+1) \leftarrow X(k), B(k+1) \leftarrow X(k)$ 
8:     Sinon Si  $f(B(k))f(X(k)) < 0$  alors
9:        $A(k+1) \leftarrow X(k), B(k+1) \leftarrow B(k)$ 
10:    Sinon
11:       $A(k+1) \leftarrow A(k), B(k+1) \leftarrow X(k)$ 
12:    Fin Si
13:     $X(k+1) \leftarrow (A(k+1) + B(k+1))/2$ 
14:  Fin Pour
15:   $x \leftarrow X(k_{\min} + 1)$ 
16: Fin Fonction

```

1 Des codes Matlab/Octave et C correspondant de cette fonction sont données en Annexe A.4.1, re-
2 spectivement en Listing A.1 et A.2. Les Listings A.3 et A.4 correspondent respectivement à un *script*
3 Matlab et un *main* C utilisant cette fonction.

4 On peut noter qu'une réécriture de la méthode de dichotomie permet d'éviter d'utiliser les deux suites
5 (a_k) et (b_k) . En effet, on a

6 • $A = a, B = b$ et $x_0 = \frac{A+B}{2}$,

• $\forall k \in \llbracket 0, k_{\min} - 1 \rrbracket$,

$$\begin{cases} A = B = x_k & \text{si } f(x_k) = 0, \\ A = x_k, B \text{ inchangé} & \text{si } f(B)f(x_k) < 0, \\ B = x_k, A \text{ inchangé} & \text{sinon (i.e. } f(A)f(x_k) < 0.) \end{cases}$$

et

$$x_{k+1} = \frac{A+B}{2}$$

7 On utilise, ces formules dans l'Algorithme 3.2.

Algorithme 3.2 Méthode de dichotomie : version 2

Données : a, b : deux réels $a < b$,
 f : $f : [a, b] \subset \mathbb{R} \longrightarrow \mathbb{R}$ vérifiant
les hypothèses de la proposition 3.1,
 eps : un réel strictement positif.

Résultat : x : un réel tel que $|x - \alpha| \leq \text{eps}$.

```

1: Fonction  $x \leftarrow \text{DICHOTOMIE2}(f, a, b, \text{eps})$ 
2:    $k_{\min} \leftarrow \lceil \log((b - a)/\text{eps}) / \log(2) \rceil$ 
3:    $X \in \mathbb{R}^{k_{\min}+1}$   $\triangleright X(k+1)$  contiendra  $x_k, \dots$ 
4:    $A \leftarrow a, B \leftarrow b, X(1) \leftarrow (A + B)/2$ 
5:   Pour  $k \leftarrow 1$  à  $k_{\min}$  faire
6:     Si  $f(X(k)) == 0$  alors
7:        $A \leftarrow X(k), B \leftarrow X(k)$ 
8:     Sinon Si  $f(B)f(X(k)) < 0$  alors
9:        $A \leftarrow X(k)$   $\triangleright B$  inchangé
10:    Sinon
11:       $B \leftarrow X(k)$   $\triangleright A$  inchangé
12:    Fin Si
13:     $X(k+1) \leftarrow (A + B)/2$ 
14:  Fin Pour
15:   $x \leftarrow X(k_{\min} + 1)$ 
16: Fin Fonction

```

Si notre objectif n'est que de calculer α_ϵ , on peut, sur le même principe, s'affranchir d'utiliser un tableau pour stocker tous les termes de la suite x_k : seul le dernier nous intéresse. On propose dans l'Algorithme 3.3, une version n'utilisant pas de tableaux.

Algorithme 3.3 Méthode de dichotomie : version 3

Données : a, b : deux réels $a < b$,
 f : $f : [a, b] \subset \mathbb{R} \longrightarrow \mathbb{R}$ vérifiant
les hypothèses de la proposition 3.1,
 eps : un réel strictement positif.

Résultat : x : un réel tel que $|x - \alpha| \leq \text{eps}$.

```

1: Fonction  $x \leftarrow \text{DICHOTOMIE2}(f, a, b, \text{eps})$ 
2:    $k_{\min} \leftarrow \lceil \log((b - a)/\text{eps}) / \log(2) \rceil$ 
3:    $A, B \in \mathbb{R}$ 
4:    $A \leftarrow a, B \leftarrow b, x \leftarrow (a + b)/2$ 
5:   Pour  $k \leftarrow 1$  à  $k_{\min}$  faire
6:     Si  $f(x) == 0$  alors
7:        $A \leftarrow x, B \leftarrow x$ 
8:     Sinon Si  $f(B)f(x) < 0$  alors
9:        $A \leftarrow x$   $\triangleright B$  inchangé
10:    Sinon
11:       $B \leftarrow x$   $\triangleright A$  inchangé
12:    Fin Si
13:     $x \leftarrow (A + B)/2$ 
14:  Fin Pour
15: Fin Fonction

```

Une autre écriture est possible sans utiliser le nombre $k_{\min}$ et donc en utilisant une boucle **Tantque** (pour les anglophobes!) ou **While** (pour les anglophiles!). Celle-ci est présentée dans l'Algorithme 3.4

Algorithme 3.4 Méthode de dichotomie : version 4

Données : a, b : deux réels $a < b$,
 f : $f : [a, b] \subset \mathbb{R} \longrightarrow \mathbb{R} \dots$,
 eps : un réel strictement positif.
Résultat : x : un réel tel que $|x - \alpha| \leq \text{eps}$.

```

1: Fonction  $x \leftarrow \text{DICHOTOMIE4}(f, a, b, \text{eps})$ 
2:    $A, B \in \mathbb{R}$ 
3:    $A \leftarrow a, B \leftarrow b, x \leftarrow (a + b)/2$ 
4:   Tantque  $|x - A| > \text{eps}$  faire
5:     Si  $f(x) == 0$  alors
6:        $A \leftarrow x, B \leftarrow x$ 
7:     Sinon Si  $f(B)f(x) < 0$  alors
8:        $A \leftarrow x$  ▷  $B$  inchangé
9:     Sinon
10:       $B \leftarrow x$  ▷  $A$  inchangé
11:   Fin Si
12:    $x \leftarrow (A + B)/2$ 
13: Fin Tantque
14: Fin Fonction

```

1 Que pensez-vous de l'algorithme suivant ?

Algorithme 3.5 Méthode de dichotomie : version 5

Données : a, b : deux réels $a < b$,
 f : $f : [a, b] \subset \mathbb{R} \longrightarrow \mathbb{R}$
Résultat : x : un réel tel que $f(x) = 0$.

```

1: Fonction  $x \leftarrow \text{DICHOTOMIE5}(f, a, b)$ 
2:    $A, B \in \mathbb{R}$ 
3:    $A \leftarrow a, B \leftarrow b, x \leftarrow (a + b)/2, xp \leftarrow a$ 
4:   Tantque  $x \sim xp$  faire
5:     Si  $f(B)f(x) < 0$  alors
6:        $A \leftarrow x$  ▷  $B$  inchangé
7:     Sinon
8:        $B \leftarrow x$  ▷  $A$  inchangé
9:     Fin Si
10:     $xp \leftarrow x$ 
11:     $x \leftarrow (A + B)/2$ 
12:  Fin Tantque
13: Fin Fonction

```

2 Des codes Matlab/Octave et C correspondant de cette fonction sont données en Annexe A.4.1, re-
 3 spectivement en Listing A.5 et A.6. Les Listings A.7 et A.8 correspondent respectivement à un *script*
 4 Matlab et un *main* C utilisant cette fonction.

3.2 Points fixes d'une fonction (dimension 1)

6 Soit $\Phi : [a, b] \subset \mathbb{R} \longrightarrow \mathbb{R}$ une fonction donnée. Rechercher un **point fixe** de Φ revient à

Trouver $\alpha \in [a, b]$ tel que

$$\alpha = \Phi(\alpha).$$

8 L'algorithme de la **méthode du point fixe** consiste en la construction, si elle existe, de la suite

$$x^{(k+1)} = \Phi(x^{(k)}), \quad \forall k \in \mathbb{N} \quad (3.1)$$

avec $x^{(0)} \in [a, b]$ donné.

♥ Definition 3.2

On dit qu'une suite $(x_k)_{k \in \mathbb{N}}$, obtenue par une méthode numérique, **converge vers α avec un ordre $p \geq 1$** si

$$\exists k_0 \in \mathbb{N}, \exists C > 0 \text{ tels que } \frac{|x_{k+1} - \alpha|}{|x_k - \alpha|^p} \leq C, \forall k \geq k_0. \quad (3.2)$$

où $C < 1$ si $p = 1$.

On a alors le théorème suivant

📖 Théorème 3.3: Théorème du point fixe dans $\mathbb{R}$

Soient $[a, b]$ un intervalle non vide de $\mathbb{R}$ et Φ une application continue de $[a, b]$ dans lui-même. Alors, il existe **au moins** un point $\alpha \in [a, b]$ vérifiant $\Phi(\alpha) = \alpha$. Le point α est appelé **point fixe de la fonction Φ** .

De plus, si Φ est contractante (lipschitzienne de rapport $L \in [0, 1[$), c'est à dire

$$\exists L < 1 \text{ t.q. } |\Phi(x) - \Phi(y)| \leq L|x - y| \forall (x, y) \in [a, b]^2, \quad (3.3)$$

alors Φ admet un **unique** point fixe $\alpha \in [a, b]$, la suite définie en (3.1) converge vers α avec un ordre 1 pour toute donnée initiale $x^{(0)}$ dans $[a, b]$, et l'on a les deux estimations suivantes :

$$|x_k - \alpha| \leq L^k |x_0 - \alpha|, \forall k \geq 0, \quad (3.4)$$

$$|x_k - \alpha| \leq \frac{L}{1 - L} |x_k - x_{k-1}|, \forall k \geq 0, \quad (3.5)$$

Proof. • On montre tout d'abord l'existence du point fixe. Pour cela, on note $f(x) = \Phi(x) - x$. f est donc une application continue de $[a, b]$ à valeurs réelles. On a $f(a) = \Phi(a) - a \geq 0$ et $f(b) = \Phi(b) - b \leq 0$ car $a \leq \Phi(x) \leq b$, pour tout $x \in [a, b]$. Si $f(a) = 0$ ou $f(b) = 0$, alors l'existence est immédiate. Sinon (i.e. $f(a) \neq 0$ et $f(b) \neq 0$), on a $f(a)f(b) < 0$ et par application directe du Théorème de Bolzano on obtient l'existence.

- On montre ensuite l'unicité sous l'hypothèse de contraction (3.3). On suppose qu'il existe α_1 et α_2 dans $[a, b]$ tels que $\Phi(\alpha_1) = \alpha_1$ et $\Phi(\alpha_2) = \alpha_2$. Dans ce cas on a

$$|\alpha_1 - \alpha_2| = |\Phi(\alpha_1) - \Phi(\alpha_2)| \leq L|\alpha_1 - \alpha_2|$$

Comme $L < 1$ on a nécessairement $\alpha_1 = \alpha_2$.

- Pour démontrer la convergence de la suite vers α , il suffit de démontrer qu'elle est de Cauchy. En effet, si elle est de Cauchy, elle converge vers une certaine limite β et Φ étant continue on a

$$\lim_{k \rightarrow +\infty} x_{k+1} = \beta = \lim_{k \rightarrow +\infty} \Phi(x_k) = \Phi(\beta).$$

Par unicité, du point fixe on obtient alors $\beta = \alpha$.

Il nous reste à montrer que la suite (x_k) est de Cauchy. Soit $k > 0$. On a

$$|x_{k+1} - x_k| = |\Phi(x_k) - \Phi(x_{k-1})| \leq L|x_k - x_{k-1}|,$$

et on obtient par récurrence

$$|x_{k+1} - x_k| \leq L^k |x_1 - x_0|$$

et

$$\forall l \geq 0, |x_{k+l} - x_{k+l-1}| \leq L^l |x_k - x_{k-1}|.$$

Soit $p > 2$. On en déduit par application répétée de l'inégalité triangulaire que

$$\begin{aligned}
 |x_{k+p} - x_k| &= |(x_{k+p} - x_{k+p-1}) + (x_{k+p-1} - x_{k+p-2}) + \dots + (x_{k+1} - x_k)| = \left| \sum_{l=0}^{p-1} x_{k+l+1} - x_{k+l} \right| \\
 &\leq \sum_{l=0}^{p-1} |x_{k+l+1} - x_{k+l}| \\
 &\leq \sum_{l=0}^{p-1} L^l |x_{k+1} - x_k| = \frac{1-L^p}{1-L} |x_{k+1} - x_k| \quad (\text{voir somme partielle d'une série géométrique}) \\
 &\leq \frac{1-L^p}{1-L} L^k |x_1 - x_0|.
 \end{aligned}$$

1 Comme $L^k \rightarrow 0$ quand $k \rightarrow +\infty$, on conclut que (x_k) est une suite de Cauchy.

On a $\forall k \geq 1$,

$$|x_{k+1} - \alpha| = |\Phi(x_k) - \Phi(\alpha)| \leq L|x_k - \alpha|$$

2 Comme $L < 1$, on a immédiatement l'ordre 1 (au moins).

- Pour démontrer la première estimation, on note que, $\forall k \geq 1$,

$$|x_k - \alpha| = |\Phi(x_{k-1}) - \Phi(\alpha)| \leq L|x_{k-1} - \alpha|$$

et donc par récurrence

$$|x_k - \alpha| \leq L^k |x_0 - \alpha|.$$

Pour la seconde, on note que, $\forall k \geq 1, \forall p \geq 1$,

$$|x_{k+p} - x_k| \leq \frac{1-L^p}{1-L} |x_{k+1} - x_k| \leq \frac{1-L^p}{1-L} L |x_k - x_{k-1}|$$

3 et en faisant tendre p vers l'infini on obtient le résultat souhaité.

4

□



Théorème 3.4: Convergence globale de la méthode du point fixe

Soit $\Phi \in \mathcal{C}^1([a, b])$ vérifiant $\Phi([a, b]) \subset [a, b]$ et

$$\exists L < 1 \text{ tel que } \forall x \in [a, b], |\Phi'(x)| \leq L, \quad (3.6)$$

Soit $x_0 \in [a, b]$ et $(x_k)_{k \in \mathbb{N}}$ la suite définie par $x_{k+1} = \Phi(x_k)$. On a alors

1. la fonction Φ admet un unique point fixe $\alpha \in [a, b]$,
2. $\forall k \in \mathbb{N}, x_k \in [a, b]$,
3. la suite (x_k) converge vers α avec un ordre 1 au moins.
- 4.

$$\lim_{k \rightarrow +\infty} \frac{x_{k+1} - \alpha}{x_k - \alpha} = \Phi'(\alpha). \quad (3.7)$$

5

Proof. Comme Φ est continue sur $[a, b]$ et vérifie $\Phi([a, b]) \subset [a, b]$, le théorème 3.3 (du point fixe) assure l'existence d'un point fixe. Pour l'unicité, on note $\alpha_1 \in [a, b]$ et $\alpha_2 \in [a, b]$, deux points fixes de Φ . On a donc $\Phi(\alpha_1) = \alpha_1$ et $\Phi(\alpha_2) = \alpha_2$. D'après le théorème des accroissements finis il existe $\xi \in]\min(\alpha_1, \alpha_2), \max(\alpha_1, \alpha_2)[\subset]a, b[$ tel que

$$\Phi(\alpha_1) - \Phi(\alpha_2) = (\alpha_1 - \alpha_2)\Phi'(\xi).$$

- 6 Comme $|\Phi'(\xi)| < 1$, on obtient $\alpha_1 - \alpha_2 = 0$, i.e. $\alpha_1 = \alpha_2$.

Pour le second point, on a immédiatement $\forall k \in \mathbb{N}$, $x_k \in [a, b]$ car $x_0 \in [a, b]$ et $\Phi([a, b]) \subset [a, b]$. Le troisième point découle du théorème des accroissements. En effet, pour tout $k > 0$, il existe $\xi_k \in]\min(\alpha, x_{k-1}), \max(\alpha, x_{k-1}) \subset]a, b[$ tel que

$$\Phi(x_{k-1}) - \Phi(\alpha) = (x_{k-1} - \alpha)\Phi'(\xi_k).$$

On obtient alors

$$|x_k - \alpha| = |\Phi(x_{k-1}) - \Phi(\alpha)| \leq L|x_{k-1} - \alpha|$$

et donc par récurrence

$$|x_k - \alpha| \leq L^k |x_0 - \alpha|.$$

Comme $L < 1$, on obtient

$$\lim_{k \rightarrow +\infty} |x_k - \alpha| = 0.$$

□ 1



Théorème 3.5: Convergence locale de la méthode du point fixe

Soit α un point fixe d'une fonction Φ de classe $\mathcal{C}^1$ au voisinage de α .

Si $|\Phi'(\alpha)| < 1$, alors il existe $\delta > 0$ pour lequel x_k converge vers α pour tout x_0 tel que $|x_0 - \alpha| < \delta$.

De plus, on a

$$\lim_{k \rightarrow +\infty} \frac{x_{k+1} - \alpha}{x_k - \alpha} = \Phi'(\alpha). \quad (3.8)$$

2

Proof. Comme Φ est continûment différentiable dans un voisinage de α avec $|\Phi'(\alpha)| < 1$, alors il existe $\delta > 0$ tel que $\forall x \in \mathcal{V} = [\alpha - \delta, \alpha + \delta]$, $|\Phi'(x)| < 1$.

Soient x et y appartenant à $\mathcal{V}$, d'après le théorème des accroissements finis il existe $\xi \in]\alpha - \delta, \alpha + \delta[$ tel que

$$f(x) - f(y) = (x - y)\Phi'(\xi).$$

Ceci entraîne que Φ est contractante sur $\mathcal{V}$. De plus $\Phi(\mathcal{V}) \subset \mathcal{V}$ car $\forall x \in \mathcal{V}$

$$|\Phi(x) - \Phi(\alpha)| = |x - \alpha||\Phi'(\xi)| \leq \delta|\Phi'(\xi)| \leq \delta$$

et donc $|\Phi(x) - \alpha| \leq \delta$ i.e. $\Phi(x) \in \mathcal{V}$. On peut donc appliquer le théorème du point fixe (Théorème 3.3) (avec $[a, b] = [\alpha - \delta, \alpha + \delta]$) ce qui donne la convergence de (x_k) vers α pour tout $x_0 \in \mathcal{V}$.

Pour démontrer (3.8), on utilise une nouvelle fois le théorème des accroissements finis : il existe $\xi_k \in]\min(x_k, \alpha), \max(x_k, \alpha)[$ tel que

$$\Phi(x_k) - \Phi(\alpha) = \Phi'(\xi_k)(x_k - \alpha).$$

Comme $\Phi(\alpha) = \alpha$ et $\Phi(x_k) = x_{k+1}$, on obtient

$$\frac{x_{k+1} - \alpha}{x_k - \alpha} = \Phi'(\xi_k).$$

Quand $k \rightarrow +\infty$, on a $x_k \rightarrow \alpha$ et donc $\xi_k \rightarrow \alpha$. Par continuité de la fonction Φ' on obtient (3.8). □ 3



Proposition 3.6

Soit $p \in \mathbb{N}^*$, et $\Phi \in \mathcal{C}^{p+1}(\mathcal{V})$ pour un certain voisinage $\mathcal{V}$ de α point fixe de Φ . Si $\Phi^{(i)}(\alpha) = 0$, pour $1 \leq i \leq p$ et si $\Phi^{(p+1)}(\alpha) \neq 0$, alors la méthode de point fixe associée à la fonction Φ est d'ordre $p + 1$ et

$$\lim_{k \rightarrow +\infty} \frac{x_{k+1} - \alpha}{(x_k - \alpha)^{p+1}} = \frac{\Phi^{(p+1)}(\alpha)}{(p+1)!}. \quad (3.9)$$

4

Proof. Les hypothèses du théorème 3.5 étant vérifiées ($\Phi'(\alpha) = 0$), la suite (x_k) converge vers α . D'après un développement de Taylor de Φ en α , il existe η_k entre x_k et α vérifiant

$$\Phi(x_k) = \sum_{i=0}^p \frac{(x_k - \alpha)^i}{i!} \Phi^{(i)}(\alpha) + \frac{(x_k - \alpha)^{p+1}}{(p+1)!} \Phi^{(p+1)}(\eta_k).$$

Comme $\alpha = \Phi(\alpha)$ et $\Phi^{(i)}(\alpha) = 0$, pour $1 \leq i \leq p$ on obtient

$$\begin{aligned} x_{k+1} &= \Phi(x_k) = \Phi(\alpha) + \sum_{i=1}^p \frac{(x_k - \alpha)^i}{i!} \Phi^{(i)}(\alpha) + \frac{(x_k - \alpha)^{p+1}}{(p+1)!} \Phi^{(p+1)}(\eta_k) \\ &= \alpha + \frac{(x_k - \alpha)^{p+1}}{(p+1)!} \Phi^{(p+1)}(\eta_k) \end{aligned}$$

De plus (η_k) converge vers α car η_k est entre x_k et α et donc comme $\Phi^{(p+1)}$ est continue au voisinage de α on obtient

$$\lim_{k \rightarrow +\infty} \frac{x_{k+1} - \alpha}{(x_k - \alpha)^{p+1}} = \lim_{k \rightarrow +\infty} \frac{\Phi^{(p+1)}(\eta_k)}{(p+1)!} = \frac{\Phi^{(p+1)}(\alpha)}{(p+1)!}.$$

□

3.2.1 Points fixes attractifs et répulsifs

Soit $\Phi : [a, b] \rightarrow [a, b]$ une application de classe $\mathcal{C}^1$ admettant un point fixe $\alpha \in [a, b]$.

Points fixes attractifs

On suppose $|\Phi'(\alpha)| < 1$. D'après le théorème 3.5, il existe $\delta > 0$ tel que

$$\forall x_0 \in [\alpha - h, \alpha + \delta], \quad \lim_{k \rightarrow +\infty} x_k = \alpha.$$

Dans ce cas on dit que α est un **point fixe attractif** pour Φ .



Exercice 3.2.1

On suppose de plus que $\Phi'(\alpha) = 0$ et $\exists M \in \mathbb{R}^+$ tel que $|\Phi''(x)| \leq M$ pour tout $x \in [\alpha - h, \alpha + \delta]$.

Q. 1 1. Montrer que

$$\forall x_0 \in [\alpha - h, \alpha + \delta], \quad |x_k - \alpha| \leq \frac{2}{M} \left(\frac{1}{2} M |x_0 - \alpha| \right)^{2^k}$$

2. Quel est l'ordre de convergence dans ce cas.

Q. 2 A quelle condition a-t-on

$$|x_k - \alpha| \leq \frac{2}{M} 10^{2^k}.$$

Correction Exercice 3.2.1

Q. 1 1. D'après la formule de Taylor, on a $\exists \eta \in]\min(\alpha, x), \max(\alpha, x)[$ tel que

$$\begin{aligned} \Phi(x) &= \Phi(\alpha) + (x - \alpha)\Phi'(\alpha) + \frac{(x - \alpha)^2}{2!} \Phi''(\eta) \\ &= \alpha + \frac{1}{2} \Phi''(\eta)(x - \alpha)^2. \end{aligned}$$

On en déduit que $|\Phi(x) - \alpha| \leq \frac{M}{2} |x - \alpha|^2$ ce qui s'écrit encore $\frac{M}{2} |\Phi(x) - \alpha| \leq (\frac{M}{2} |x - \alpha|)^2$. Or si $x_0 \in [\alpha - h, \alpha + \delta]$, alors $x_{k-1} \in [\alpha - h, \alpha + \delta]$, $\forall k \in \mathbb{N}^*$. On a alors

$$\frac{M}{2} |\Phi(x_{k-1}) - \alpha| = \frac{M}{2} |x_k - \alpha| \leq \left(\frac{M}{2} |x_{k-1} - \alpha| \right)^2$$

et donc par récurrence

$$\frac{M}{2} |x_k - \alpha| \leq \left(\frac{M}{2} |x_0 - \alpha| \right)^{2^k}.$$

2. On a

$$|x_k - \alpha| \leq \frac{2}{M} \left(\frac{M}{2} |x_{k-1} - \alpha| \right)^2$$

c'est à dire

$$\frac{|x_k - \alpha|}{|x_{k-1} - \alpha|^2} \leq \frac{M}{2}$$

et donc la méthode est d'ordre 2.

Q. 2 En supposant de plus que $|x_0 - \alpha| \leq \frac{1}{5M}$, on obtient immédiatement le résultat.

◇

Points fixes répulsifs

Cas $|\Phi'(\alpha)| > 1$. Par définition de la dérivée en α on a

$$\lim_{\substack{x \rightarrow \alpha \\ x \neq \alpha}} \left| \frac{\Phi(x) - \Phi(\alpha)}{x - \alpha} \right| = |\Phi'(\alpha)| > 1.$$

Il existe alors $\delta > 0$ tel que

$$\forall x \in [\alpha - h, \alpha + \delta] \setminus \{\alpha\}, \quad |\Phi(x) - \alpha| > |x - \alpha|$$

et donc la suite (x_k) ne peut converger vers α et ceci même si x_0 est très proche de α . Dans ce cas on dit que α est un **point fixe répulsif** pour Φ .

Toutefois, on *peut rattrapper le coup*. En effet, comme Φ' est non nulle et de signe constant au voisinage de α : il existe $h > 0$ tel que la fonction Φ soit strictement monotone sur $I = [\alpha - h, \alpha + h]$. D'après le corollaire A.4 Φ est une bijection de I dans l'intervalle fermé $J = \Phi^{-1}(I)$. Comme $\alpha = \Phi(\alpha)$, on a $\alpha \in J$ et $\Phi^{-1}(\alpha) = \alpha$. Le problème de point fixe trouver $x \in I$ tel que $\Phi(x) = x$ revient alors à trouver $x \in J$ tel que $\Phi^{-1}(x) = x$. Or $\Phi'(\alpha) \neq 0$ et $\Phi(\alpha) = \alpha$, la proposition A.5 donne alors $(\Phi^{-1})'(\alpha) = 1/\Phi'(\alpha)$ et donc $|(\Phi^{-1})'(\alpha)| < 1$. Le point α est un **point fixe attractif** pour Φ^{-1} .

Points fixes indéterminés

Dans le cas $|\Phi'(\alpha)| = 1$, on ne peut conclure dans le cas général.

3.2.2 Interprétations graphiques de la méthode du point fixe

Exemple 1 : point fixe de la fonction $x \mapsto x^2$.

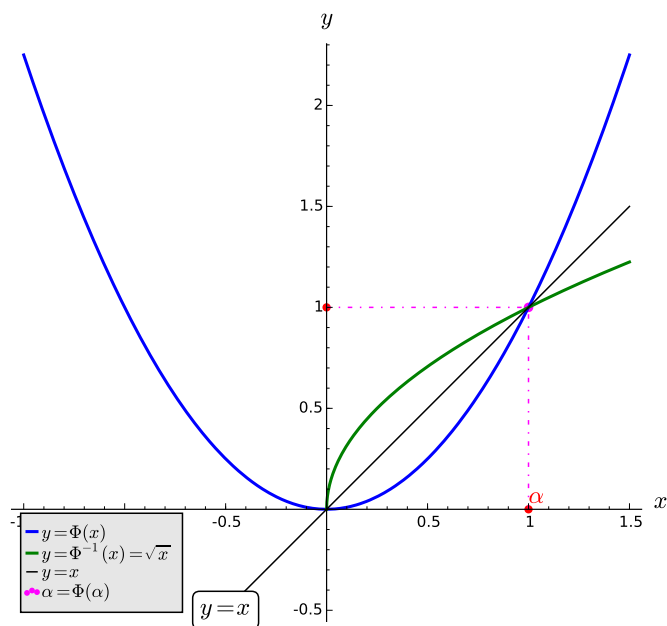
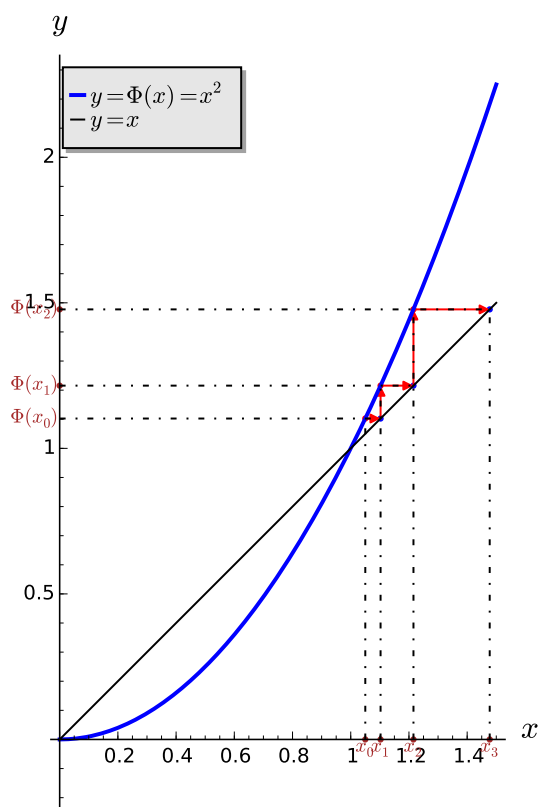
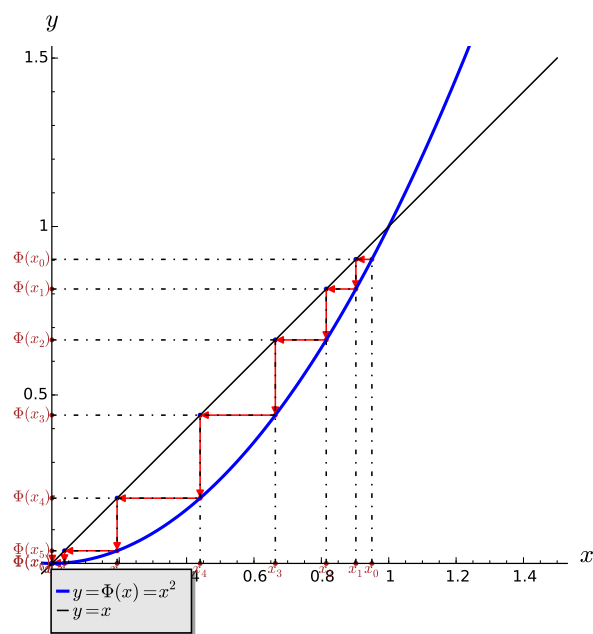
On s'intéresse ici au point fixe $\alpha = 1$ de la fonction $\Phi : x \mapsto x^2$.

Comme $\Phi'(\alpha) = 2$ le point fixe α est répulsif. On donne dans le tableau suivant les premiers termes de deux suites : l'une avec $x_0 = 1.05$ et l'autre avec $x_0 = 0.95$. Dans ce dernier cas, la suite va converger ... vers un autre point fixe.

k	x_k	x_k
0	1.05000	0.950000
1	1.10250	0.902500
2	1.21551	0.814506
3	1.47746	0.663420
⋮	⋮	⋮
8	265742.	1.98264×10^{-6}

Les premières itérations de ces deux suites sont représentées en Figure 3.4.

Sur l'intervalle $[0, +\infty[$, la fonction Φ admet comme fonction inverse $\Phi^{-1}(x) = \sqrt{x}$ et on a aussi $\Phi^{-1}(1) = 1$ (i.e. $\alpha = 1$ est un point fixe de Φ^{-1}). De plus $(\Phi^{-1})'(1) = 1/2 < 1$, ce qui permet d'affirmer que $\alpha = 1$ est un point fixe **attractif** de Φ^{-1} . On donne dans le tableau suivant les premiers termes de

Figure 3.3: fonction x^2 et son fonction inverse $\sqrt{x}$ sur $[0, +\infty[$ (a) $x_0 = 1.05$ (b) $x_0 = 0.950$ Figure 3.4: $\alpha = 1$, point fixe répulsif de $x \mapsto x^2$

deux suites : l'une avec $x_0 = 1.40$ et l'autre avec $x_0 = 0.200$.

k	x_k	x_k
0	1.40000	0.200000
1	1.18322	0.447214
2	1.08776	0.668740
3	1.04296	0.817765
$\vdots$	$\vdots$	$\vdots$
8	1.00132	0.993733

Les premières itérations de ces deux suites sont représentées en Figure 3.5.

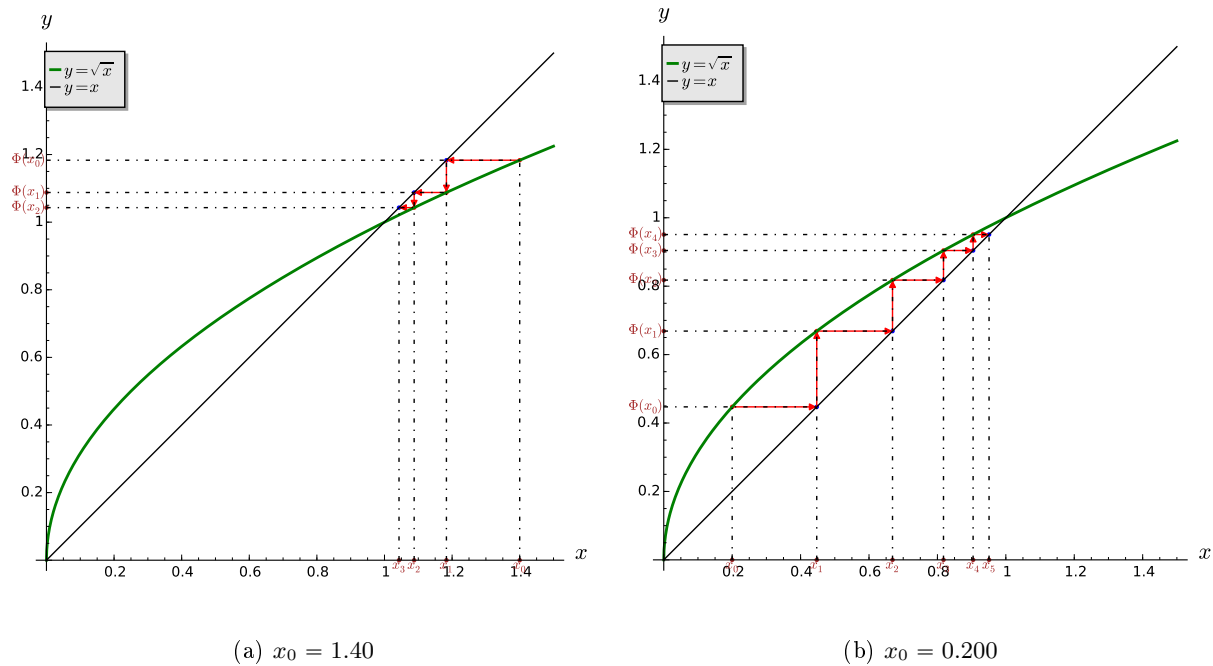


Figure 3.5: $\alpha = 1$, point fixe attractif de $x \mapsto \sqrt{x}$

Exemple 2 : point fixe de la fonction $x \mapsto x^2 - x + 1$.

On s'intéresse ici au point fixe $\alpha = 1$ de la fonction $f : x \mapsto x^2 - x + 1$ représenté en Figure 3.6. On note que $f'(\alpha) = 1$.

On donne dans le tableau suivant les premiers termes de deux suites : l'une avec $x_0 = 1.10$ diverge et l'autre avec $x_0 = 0.500$ converge vers α .

k	x_k	x_k
0	1.10000	0.500000
1	1.11000	0.750000
2	1.12210	0.812500
3	1.13701	0.847656
$\vdots$	$\vdots$	$\vdots$
8	1.32397	0.918223

Les premières itérations de ces deux suites sont représentées en Figure 3.7.

Exemple 3 : point fixe de fonctions affines particulières

On va étudier les points fixes des trois fonctions affines vérifiant $f_1(1) = f_2(1) = f_3(1) = 1$ et $f'_1(1) = -1$, $f'_2(1) = -3/4$ et $f'_3(1) = -4/3$. Ces trois fonctions sont données par $f_1(x) = -x + 2$, $f_2(x) = -\frac{3}{4}x + \frac{7}{4}$ et $f_3(x) = -\frac{4}{3}x + \frac{7}{3}$.

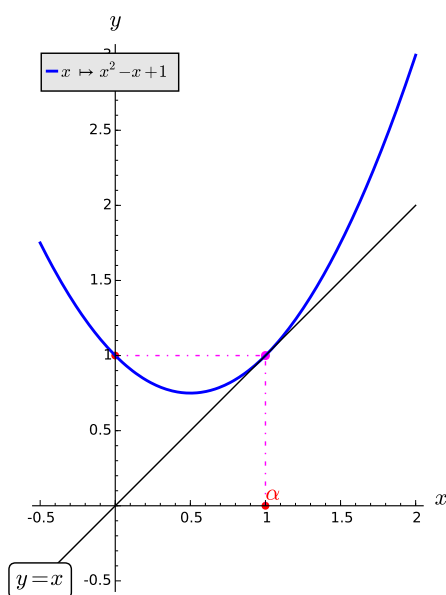
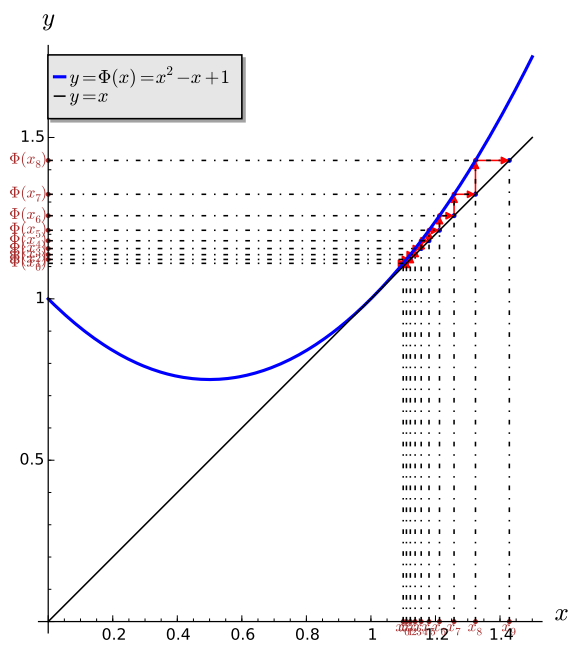
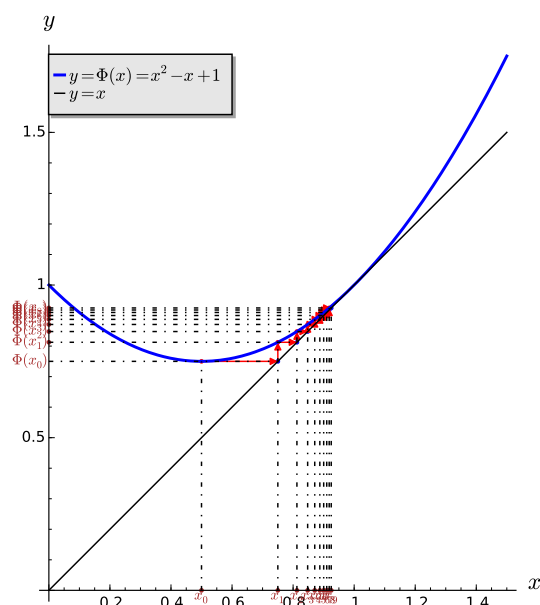


Figure 3.6: fonction $x^2 - x + 1$ et son point fixe $\alpha = 1$.



(a) $x_0 = 1.1$, α point répulsif



(b) $x_0 = 0.50$, α point attractif

Figure 3.7: $\alpha = 1$, point fixe attractif ou répulsif de $x \mapsto x^2 - x + 1$

Comme $\Phi'(\alpha) = 2$ le point fixe α est répulsif. On donne dans le tableau suivant les premiers termes de trois suites obtenues avec $x_0 = 1.5$ pour les suites associées à f_1 et f_2 , et avec $x_0 = 1.1$ pour celle associée à f_2 .

k	$x_k(f_1)$	$x_k(f_2)$	$x_k(f_3)$
0	1.50000	1.50000	1.10000
1	0.500000	0.625000	0.866667
2	1.50000	1.28125	1.17778
3	0.500000	0.789062	0.762963
$\vdots$	$\vdots$	$\vdots$	
19	0.500000	0.997886	-22.6503
20	1.50000	1.00159	32.5337

Les premières itérations des trois suites obtenues sont représentées en Figure 3.8.

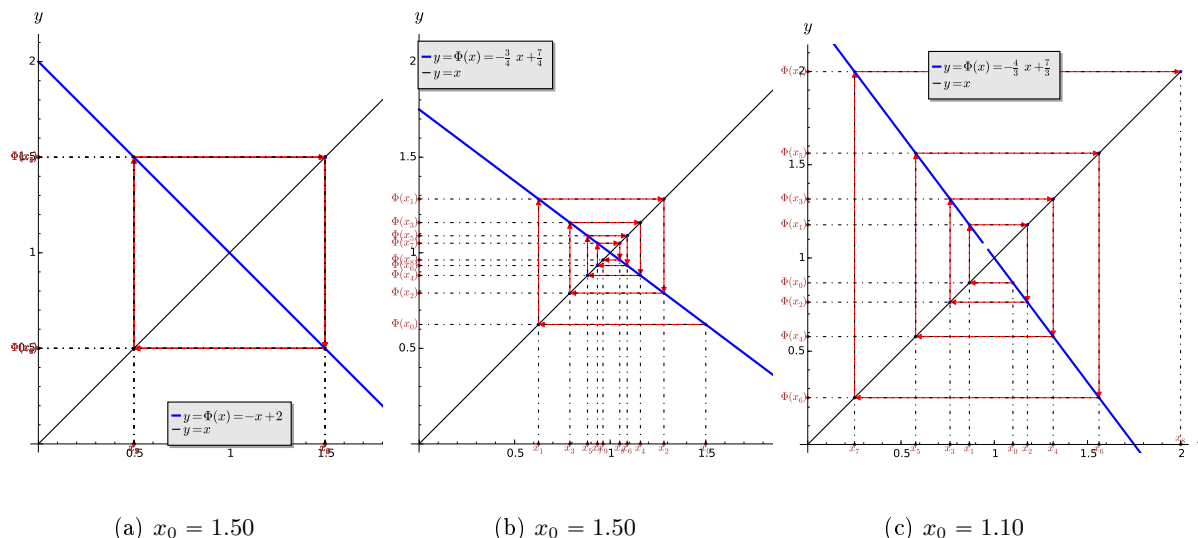


Figure 3.8: $\alpha = 1$, point fixe de fonctions affines particulières

3.2.3 Algorithme générique du point fixe

Le principe de base est très simple et donné par les Algorithmes 3.6 et 3.7, écrit respectivement avec une boucle **Tantque** et une boucle **Répéter**.

Algorithme 3.6 Méthode de point fixe : version **Tantque formel**

```

1:  $k \leftarrow 0$ 
2: Tantque non convergence faire
3:    $x_{k+1} \leftarrow \Phi(x_k)$ 
4:    $k \leftarrow k + 1$ 
5: Fin Tantque
6:  $\alpha_\epsilon \leftarrow x_k$  ▷ le dernier calculé.

```

Algorithme 3.7 Méthode de point fixe : version **Répéter formel**

```

1:  $k \leftarrow 0$ 
2: Répéter
3:    $k \leftarrow k + 1$ 
4:    $x_k \leftarrow \Phi(x_{k-1})$ 
5: jusqu'à convergence
6:  $\alpha_\epsilon \leftarrow x_k$  ▷ le dernier calculé.

```

Nous allons plus particulièrement étudier les critères d'arrêt des boucles **Tantque** et **Répéter** (négarion l'un de l'autre) qui sont bien trop flous : on s'arrête quand on converge!

Tout d'abord, on n'est pas sûr de converger : il faudra donc n'autoriser qu'un nombre maximum d'itérations $k_{\max}$. De plus, si l'on converge vers une certaine valeur α celle-ci n'étant pas connue à l'avance (sinon l'algo n'a aucun intérêt), on ne peut utiliser, comme condition du **Tantque**, $|x_k - \alpha| > \epsilon$ ($|x_k - \alpha| \leq \epsilon$ pour la condition du **Répéter**) où ϵ serait la précision souhaitée. L'idée serait de comparer x_{k+1} et x_k et donc de tester s'ils sont suffisamment proches : la condition serait alors $|\Phi(x_k) - x_k| = |x_{k+1} - x_k| > \text{tol}$, (boucle **Tantque**) ou $|\Phi(x_k) - x_k| \leq \text{tol}$, (boucle **Répéter**) avec tol la tolérance souhaitée. Il faut noter que dans ce cas la valeur tol doit être choisie correctement et *dépend* de l'ordre

de grandeur de α . Si α est de l'ordre de 10^8 , $\text{tol} = 1$ comme valeur est raisonnable. Toutefois on peut lui préférer une condition *relative* $\frac{|\Phi(x_k) - x_k|}{|x_k| + 1} > \text{tol}$ (boucle **Tantque**) ou $\frac{|\Phi(x_k) - x_k|}{|x_k| + 1} \leq \text{tol}$ (boucle **Répéter**) qui permet à tol de correspondre à une tolérance *relative* souhaitée que ce soit pour de grandes valeurs de α ou pour des petites. Les algorithmes du point fixe (recherche de α solution de $\Phi(x) = x$) avec notations mathématiques sont donnés par Algorithme 3.8 et 3.9, respectivement avec des boucles **Tantque** et **Répéter**.

Algorithme 3.8 Méthode de point fixe : version **Tantque** *formel* avec critères d'arrêt

```

1:  $k \leftarrow 0$ 
2:  $\text{err} \leftarrow |\Phi(x_0) - x_0|$   $\triangleright$  ou  $\frac{|\Phi(x_0) - x_0|}{|x_0| + 1}$ 
3: Tantque  $\text{err} > \epsilon$  et  $k \leq \text{kmax}$  faire
7 4:    $k \leftarrow k + 1$ 
5:    $x_k \leftarrow \Phi(x_{k-1})$ 
6:    $\text{err} \leftarrow |\Phi(x_k) - x_k|$   $\triangleright$  ou  $\frac{|\Phi(x_k) - x_k|}{|x_k| + 1}$ 
7: Fin Tantque
8: Si  $\text{err} \leq \text{tol}$  alors  $\triangleright$  Convergence
9:    $\alpha_{\text{tol}} \leftarrow x_k$   $\triangleright |\Phi(\alpha_{\text{tol}}) - \alpha_{\text{tol}}| \leq \text{tol}$ 
10: Fin Si
```

Algorithme 3.9 Méthode de point fixe : version **Répéter** *formel* avec critères d'arrêt

```

1:  $k \leftarrow 0$ 
2: Répéter
3:    $\text{err} \leftarrow |\Phi(x_k) - x_k|$   $\triangleright$  ou  $\frac{|\Phi(x_k) - x_k|}{|x_k| + 1}$ 
4:    $x_{k+1} \leftarrow \Phi(x_k)$ 
5:    $k \leftarrow k + 1$ 
6: jusqu'à  $\text{err} \leq \text{tol}$  ou  $k > \text{kmax}$ 
7: Si  $\text{err} \leq \text{tol}$  alors  $\triangleright$  Convergence
8:    $\alpha_{\text{tol}} \leftarrow x_k$   $\triangleright |\Phi(\alpha_{\text{tol}}) - \alpha_{\text{tol}}| \leq \text{tol}$ 
9: Fin Si
```

Des versions, sous forme de fonction, plus proches de la programmation sont proposées en Algorithme 3.10 et 3.11. Bien évidemment, suivant l'usage que l'on souhaite de ces fonctions, il est facile de les adapter pour retourner plus d'informations (nombre d'itération effectives, statut de convergence, ...)

Algorithme 3.10 Méthode de point fixe : version **Tantque** avec critères d'arrêt

Données :

Φ : $\Phi : \mathbb{R} \rightarrow \mathbb{R}$,
 x_0 : donnée initiale, $x_0 \in \mathbb{R}$,
 tol : la tolérance, $\text{tol} \in \mathbb{R}^+$,
 kmax : nombre maximum d'itérations, $\text{kmax} \in \mathbb{N}^*$

Résultat :

α_{tol} : un réel tel que $|\Phi(\alpha_{\text{tol}}) - \alpha_{\text{tol}}| \leq \text{tol}$
 (ou $\frac{|\Phi(\alpha_{\text{tol}}) - \alpha_{\text{tol}}|}{|\alpha_{\text{tol}}| + 1} \leq \text{tol}$)

```

11 1: Fonction  $\alpha_{\text{TOL}} \leftarrow \text{PTFIXE}(\Phi, x_0, \text{tol}, \text{kmax})$ 
2:    $k \leftarrow 0$ ,  $\alpha_{\text{tol}} \leftarrow \emptyset$ 
3:    $x \leftarrow x_0$ ,  $\text{fx} \leftarrow \Phi(x_0)$ ,
4:    $\text{err} \leftarrow |\text{fx} - x|$   $\triangleright$  ou  $\frac{|\text{fx} - x|}{|x| + 1}$ 
5:   Tantque  $\text{err} > \text{tol}$  et  $k \leq \text{kmax}$  faire
6:      $k \leftarrow k + 1$ 
7:      $x \leftarrow \text{fx}$ 
8:      $\text{fx} \leftarrow \Phi(x)$ 
9:      $\text{err} \leftarrow |\text{fx} - x|$   $\triangleright$  ou  $\frac{|\text{fx} - x|}{|x| + 1}$ 
10:   Fin Tantque
11:   Si  $\text{err} \leq \text{tol}$  alors  $\triangleright$  Convergence
12:      $\alpha_{\text{tol}} \leftarrow x$ 
13:   Fin Si
14: Fin Fonction
```

Algorithme 3.11 Méthode de point fixe : version **Répéter** avec critères d'arrêt

Données :

Φ : $\Phi : \mathbb{R} \rightarrow \mathbb{R}$,
 x_0 : donnée initiale, $x_0 \in \mathbb{R}$,
 tol : la tolérance, $\text{tol} \in \mathbb{R}^+$,
 kmax : nombre maximum d'itérations, $\text{kmax} \in \mathbb{N}^*$

Résultat :

α_{tol} : un réel tel que $|\Phi(\alpha_{\text{tol}}) - \alpha_{\text{tol}}| \leq \text{tol}$
 (ou $\frac{|\Phi(\alpha_{\text{tol}}) - \alpha_{\text{tol}}|}{|\alpha_{\text{tol}}| + 1} \leq \text{tol}$)

```

1: Fonction  $\alpha_{\text{TOL}} \leftarrow \text{PTFIXE}(\Phi, x_0, \text{tol}, \text{kmax})$ 
2:    $k \leftarrow 0$ ,  $\alpha_{\text{tol}} \leftarrow \emptyset$ 
3:    $x \leftarrow x_0$ 
4:   Répéter
5:      $x_p \leftarrow x$ 
6:      $x \leftarrow \Phi(x_p)$ 
7:      $\text{err} \leftarrow |x - x_p|$   $\triangleright$  ou  $\frac{|x - x_p|}{|x_p| + 1}$ 
8:      $k \leftarrow k + 1$ 
9:   jusqu'à  $\text{err} \leq \text{tol}$  ou  $k > \text{kmax}$ 
10:   Si  $\text{err} \leq \text{tol}$  alors  $\triangleright$  Convergence
11:      $\alpha_{\text{tol}} \leftarrow x$ 
12:   Fin Si
13: Fin Fonction
```

3.2.4 Méthodes de points fixes pour la recherche de racines

On peut noter que résoudre $f(x) = 0$ revient, par exemple, à résoudre $\Phi(x) := x + f(x) = x$. De manière plus générale, si $\mathcal{F}$ est une fonction continue vérifiant $\mathcal{F}(0) = 0$ alors $\Phi(x) = x + \mathcal{F}(f(x))$ est un autre choix possible.

Soit f une fonction de classe $\mathcal{C}^1$ dans un voisinage d'une de ses racines simple α .

Selectionner une version :

Version 1 En appliquant la formule de Taylor pour tout x dans ce voisinage, il existe $\xi \in]\min(x, \alpha), \max(x, \alpha)[$ tel que

$$f(\alpha) = 0 = f(x) + (\alpha - x)f'(\xi).$$

Ceci conduit à la méthode itérative suivante : x_0 étant donné dans le voisinage de α , on doit, $\forall k \in \mathbb{N}$, déterminer x_{k+1} vérifiant $f(x_k) + (x_{k+1} - x_k)q_k = 0$ sachant que q_k est une approximation de $f'(\xi)$.

Version 2 On cherche à déterminer une méthode itérative permettant de calculer x_{k+1} en fonction des valeurs précédentes en espérant que $|x_{k+1} - \alpha| \leq |x_k - \alpha|$. Pour cela, on écrit une formule de Taylor en x_k supposé proche de α et on note $h = \alpha - x_k$

$$f(\alpha) = 0 = f(x_k) + hf'(x_k) + o(h)$$

L'idéal serait de trouver la valeur exacte de h et de prendre $x_{k+1} = x_k + h$ mais ceci n'est pas possible dans un cadre général. C'est pourquoi on cherche une approximation $\tilde{h}$ de h . De plus, on ne connaît pas forcément explicitement la dérivée de f . On note donc q_k une approximation de $f'(x_k)$. On choisit alors $\tilde{h}$ comme solution de

$$f(x_k) + \tilde{h}q_k = 0$$

Si $q_k \neq 0$, on obtient la suite itérative

$$x_{k+1} = x_k - \frac{f(x_k)}{q_k}, \quad \forall k \in \mathbb{N} \quad (3.10)$$

La valeur x_{k+1} est de fait l'intersection de la droite passant par le point $((x_k), f(x_k))$ et de pente q_k avec l'axe des x .

Par la suite, on va écrire un algorithme du point fixe et étudier différentes méthodes :

- la **méthode de la corde** :

$$q_k = q = \frac{f(b) - f(a)}{b - a}$$

- la **méthode de la sécante** :

$$q_k = \frac{f(x_k) - f(x_{k-1})}{x_k - x_{k-1}}$$

où x_{-1} et x_0 sont données,

- la **méthode de Newton** : en supposant f' connu, on prend

$$q_k = f'(x_k).$$

La méthode de la corde

Soit f une fonction de classe $\mathcal{C}^1$ sur $[a, b]$ tel que $f(a) \neq f(b)$. Soit $x_0 \in [a, b]$ donné. La suite obtenue par la méthode de la corde est donnée par

$$x_{k+1} = x_k - \frac{b - a}{f(b) - f(a)} f(x_k), \quad \forall k \in \mathbb{N}. \quad (3.11)$$

On peut voir cette relation comme un cas particulier de la méthode du point fixe. En effet, en prenant $\Phi(x) = x - \frac{b-a}{f(b)-f(a)} f(x)$, la suite définie en (3.11) s'écrit $x_{k+1} = \Phi(x_k)$.

On a alors le résultat suivant



Proposition 3.7: convergence de la méthode de la corde

Soit $f \in \mathcal{C}^1([a, b])$ tel que $f(b) \neq f(a)$ et $\lambda = \frac{b-a}{f(b)-f(a)}$. On note $(x_k)_{k \in \mathbb{N}}$ la suite définie par $x_0 \in [a, b]$ et pour tout $k \geq 0$

$$x_{k+1} = x_k - \lambda f(x_k). \quad (3.12)$$

On suppose de plus que $\forall x \in [a, b]$

$$\min(\lambda(x-a), \lambda(x-b)) \leq f(x) \leq \max(\lambda(x-a), \lambda(x-b)) \quad (3.13)$$

$$\min(0, 2\lambda) < f'(x) < \max(0, 2\lambda) \quad (3.14)$$

alors la suite (x_k) converge vers l'unique racine $\alpha \in [a, b]$ de f .

Proof. voir correction Exercice 3.2.2 □



Exercice 3.2.2

Soit f une fonction de classe $\mathcal{C}^1$ sur $[a, b]$ vérifiant $f(a)f(b) < 0$. et $\lambda = \frac{b-a}{f(b)-f(a)}$. Soit $x_0 \in [a, b]$ donné. La suite obtenue par la méthode de la corde est donnée par

$$x_{k+1} = x_k - \frac{b-a}{f(b)-f(a)} f(x_k), \quad \forall k \in \mathbb{N}.$$

On note $\Phi(x) = x - \frac{b-a}{f(b)-f(a)} f(x)$.

Q. 1 Montrer que si pour tout $x \in [a, b]$ on a

$$\min(\lambda(x-a), \lambda(x-b)) \leq f(x) \leq \max(\lambda(x-a), \lambda(x-b)) \quad (3.15)$$

alors $\Phi([a, b]) \subset [a, b]$.

Q. 2 Montrer que si pour tout $x \in [a, b]$ on a

$$\min(0, 2\lambda) < f'(x) < \max(0, 2\lambda) \quad (3.16)$$

alors $|\Phi'(x)| < 1$.

Q. 3 En déduire que sous les deux conditions précédentes la méthode de la corde converge vers l'unique solution $\alpha \in [a, b]$ de $f(x) = 0$.

Correction Exercice 3.2.2

Q. 1 Si $\lambda > 0$, l'inéquation (3.15) devient

$$\begin{aligned} \lambda(x-b) \leq f(x) \leq \lambda(x-a) &\Leftrightarrow a \leq x - \frac{f(x)}{\lambda} \leq b \\ &\Leftrightarrow a \leq \Phi(x) \leq b. \end{aligned}$$

Si $\lambda < 0$, l'inéquation (3.15) devient

$$\begin{aligned} \lambda(x-a) \leq f(x) \leq \lambda(x-b) &\Leftrightarrow a \leq x - \frac{f(x)}{\lambda} \leq b \\ &\Leftrightarrow a \leq \Phi(x) \leq b. \end{aligned}$$

Q. 2 Si $\lambda > 0$, l'inéquation (3.16) devient

$$\begin{aligned} 0 < f'(x) < 2\lambda &\Leftrightarrow 0 < \frac{f'(x)}{\lambda} < 2 \\ &\Leftrightarrow -1 < 1 - \frac{f'(x)}{\lambda} < 1 \\ &\Leftrightarrow -1 < \Phi'(x) < 1. \end{aligned}$$

Si $\lambda < 0$, l'inéquation (3.16) devient

$$\begin{aligned} 2\lambda < f'(x) < 0 &\Leftrightarrow 0 < \frac{f'(x)}{\lambda} < 2 \\ &\Leftrightarrow -1 < 1 - \frac{f'(x)}{\lambda} < 1 \\ &\Leftrightarrow -1 < \Phi'(x) < 1. \end{aligned}$$

Q. 3 Sous les hypothèses (3.15) et (3.16) on a $\Phi([a, b]) \subset [a, b]$ et $\forall x \in [a, b], |\Phi'(x)| < 1$. Comme f est de classe $\mathcal{C}^1$ sur $[a, b]$, la fonction Φ l'est aussi. La suite (x_k) est définie par $x_{k+1} = \Phi(x_k)$. Ainsi les hypothèses du théorème 3.4 sont vérifiées ce qui assure l'unicité du point fixe ainsi que la convergence de la suite (x_k) vers ce point fixe.

◇



Proposition 3.8: ordre de convergence de la méthode de la corde

Soit $f \in \mathcal{C}^1([a, b])$ tel que $f(b) \neq f(a)$. Si la suite (x_k) définie par la méthode de la corde en (3.12) converge vers $\alpha \in]a, b[$ alors la convergence est au moins d'ordre 1.

De plus, si f est de classe $\mathcal{C}^2$ sur un certain voisinage $\mathcal{V}$ de α et si $f'(\alpha) = \frac{f(b)-f(a)}{b-a}$ alors la convergence est au moins d'ordre 2.

Proof. • **Order 1 :** On note $\lambda = \frac{b-a}{f(b)-f(a)}$. On a par définition $x_{k+1} = x_k - \lambda f(x_k)$ (ce qui suppose que $f(x_k)$ soit bien définie, i.e. $x_k \in [a, b]$). Comme $\lambda \neq 0$ et f continue, l'hypothèse (x_k) converge vers α entraîne que $f(\alpha) = 0$.

Pour définir l'ordre de convergence, on suppose de plus que $\forall k \in \mathbb{N}, x_k \neq \alpha$. On peut alors appliquer la formule de Taylor-Lagrange : il existe ξ_k compris entre x_k et α tel que

$$f(x_k) = \underbrace{f(\alpha)}_{=0} + (x_k - \alpha)f'(\xi_k) = (x_k - \alpha)f'(\xi_k).$$

On a alors en utilisant cette expression dans la définition de la suite x_k

$$x_{k+1} = x_k - \lambda(x_k - \alpha)f'(\xi_k).$$

En soustrayant α à cette équation on obtient

$$x_{k+1} - \alpha = x_k - \alpha - \lambda(x_k - \alpha)f'(\xi_k) = (x_k - \alpha)(1 - \lambda f'(\xi_k)).$$

Comme $x_k \neq \alpha$, on a alors

$$\frac{x_{k+1} - \alpha}{x_k - \alpha} = 1 - \lambda f'(\xi_k).$$

Or x_k converge vers α et ξ_k compris entre x_k et α , ce qui entraîne que ξ_k converge vers α . La fonction f' étant continue, on en déduit que $f'(\xi_k)$ converge vers $f'(\alpha)$. Ceci donne donc

$$\lim_{k \rightarrow +\infty} \frac{x_{k+1} - \alpha}{x_k - \alpha} = 1 - \lambda f'(\alpha).$$

La convergence est donc (au moins) d'ordre 1.

- **Order 2 :** La suite étant convergente, il existe $k_0 \in \mathbb{N}$ tel que $\forall k \geq k_0, x_k \in \mathcal{V}$. Soit $k \geq k_0$, comme $f \in \mathcal{C}^2(\mathcal{V})$, on peut appliquer la formule de Taylor-Lagrange : il existe $\eta_k \in \mathcal{V}$ compris entre x_k et α que

$$f(x_k) = \underbrace{f(\alpha)}_{=0} + (x_k - \alpha)f'(\alpha) + \frac{(x_k - \alpha)^2}{2!}f^{(2)}(\eta_k).$$

On a alors en utilisant cette expression dans la définition de la suite x_k

$$x_{k+1} = x_k - \lambda \left((x_k - \alpha)f'(\alpha) + \frac{(x_k - \alpha)^2}{2!}f^{(2)}(\eta_k) \right).$$

En soustrayant α à cette équation on obtient

$$x_{k+1} - \alpha = (x_k - \alpha) \underbrace{(1 - \lambda f'(\alpha))}_{=0 \text{ par hyp.}} + \lambda \frac{(x_k - \alpha)^2}{2!} f^{(2)}(\eta_k).$$

Comme $\eta_k \in \mathcal{V}$ converge vers α (car compris entre x_k et α) et $f^{(2)}$ continue sur $\mathcal{V}$, on en déduit

$$\lim_{k \rightarrow +\infty} \frac{x_{k+1} - \alpha}{(x_k - \alpha)^2} = \lambda \frac{f^{(2)}(\alpha)}{2}.$$

La convergence est donc (au moins) d'ordre 2.

□

On présente en Algorithme 3.12 l'implémentation usuelle de la méthode de la corde. Une autre version utilisant la fonction **PTFIXE** est donnée en Algorithme 3.13.

Algorithme 3.12 Méthode de la corde

Données :

f : $f : \mathbb{R} \rightarrow \mathbb{R}$,
 a, b : deux réels tels que $f(a) \neq f(b)$,
 x_0 : donnée initiale, $x_0 \in \mathbb{R}$,
 tol : la tolérance, $\text{tol} \in \mathbb{R}^+$,
 kmax : nombre maximum d'itérations, $\text{kmax} \in \mathbb{N}^*$

Résultat :

α_{tol} : un réel tel que $|f(\alpha_{\text{tol}})| \leq \text{tol}$

1: **Fonction** $\alpha_{\text{tol}} \leftarrow \text{CORDE}(f, a, b, x_0, \text{tol}, \text{kmax})$

2: $k \leftarrow 0, \alpha_{\text{tol}} \leftarrow \emptyset$

3: $q \leftarrow \frac{b-a}{f(b)-f(a)}$

4: $x \leftarrow x_0$,

5: $\text{err} \leftarrow \text{tol} + 1$

6: **Tantque** $\text{err} > \text{tol}$ et $k \leq \text{kmax}$ **faire**

7: $k \leftarrow k + 1$

8: $\text{xp} \leftarrow x$

9: $x \leftarrow \text{xp} - q * f(\text{xp})$

10: $\text{err} \leftarrow |x - \text{xp}|$

11: **Fin Tantque**

12: **Si** $\text{err} \leq \text{tol}$ **alors** $\triangleright$ Convergence

13: $\alpha_{\text{tol}} \leftarrow x$

14: **Fin Si**

15: **Fin Fonction**

Algorithme 3.13 Méthode de la corde

Données :

f : $f : \mathbb{R} \rightarrow \mathbb{R}$,

a, b : deux réels tels que $f(a) \neq f(b)$,

x_0 : donnée initiale, $x_0 \in \mathbb{R}$,

tol : la tolérance, $\text{tol} \in \mathbb{R}^+$,

kmax : nombre maximum d'itérations, $\text{kmax} \in \mathbb{N}^*$

Résultat :

α_{tol} : un réel tel que $|f(\alpha_{\text{tol}})| \leq \text{tol}$

1: **Fonction** $\alpha_{\text{tol}} \leftarrow \text{CORDE}(f, a, b, x_0, \text{tol}, \text{kmax})$

2: $q \leftarrow \frac{b-a}{f(b)-f(a)}$

3: $\Phi \leftarrow x \mapsto x - q * f(x)$

4: $\alpha_{\text{tol}} \leftarrow \text{PTFIXE}(\Phi, x_0, \text{tol}, \text{kmax})$

5: **Fin Fonction**

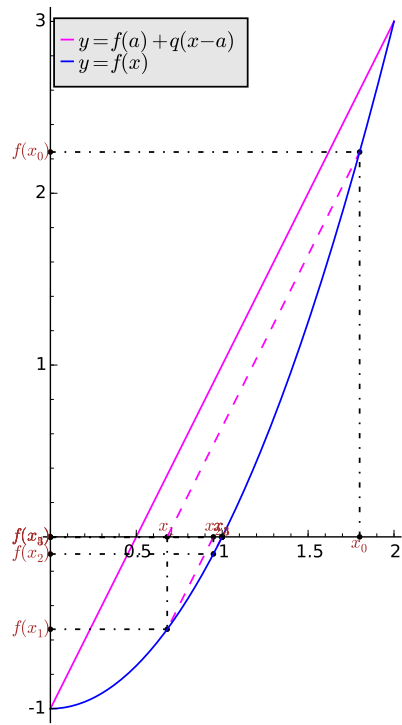
Pour illustrer ces résultats de convergence, on va rechercher la racine positive de $f(x) = x^2 - 1$ par la méthode de la corde avec comme données

• exemple 1 : $a = 0.000, b = 2.000, x_0 = 1.800$,

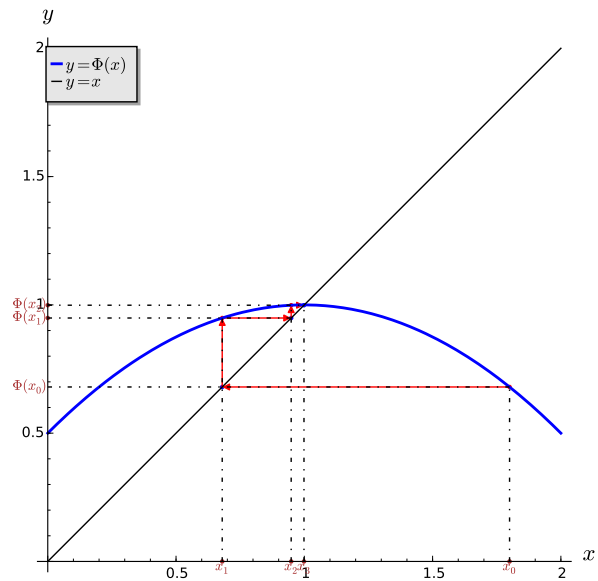
• exemple 2 : $a = 0.5000, b = 1.900, x_0 = 1.800$.

On représente en Figure 3.9 les itérations de la méthode de la corde pour l'exemple 1 à partir du graphe de la fonction f (figure de gauche) et à partir du graphe de la fonction Φ (figure de droite). Même chose pour l'exemple 2 avec la Figure 3.10.

Dans le tableau suivant on donne l'erreur commise par les suites dérivées des exemples 1 et 2 et ceci

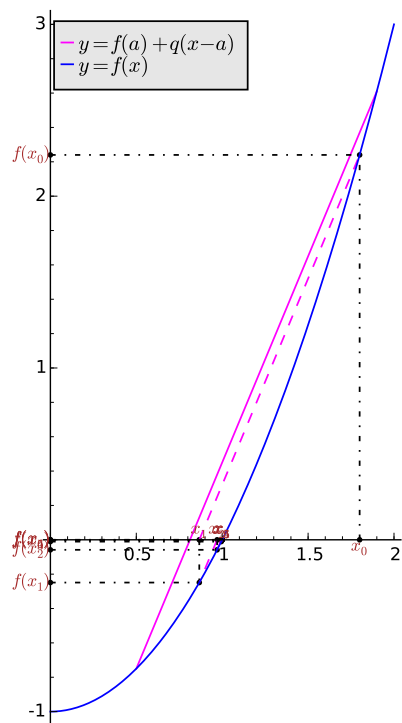


(a) représentation usuelle

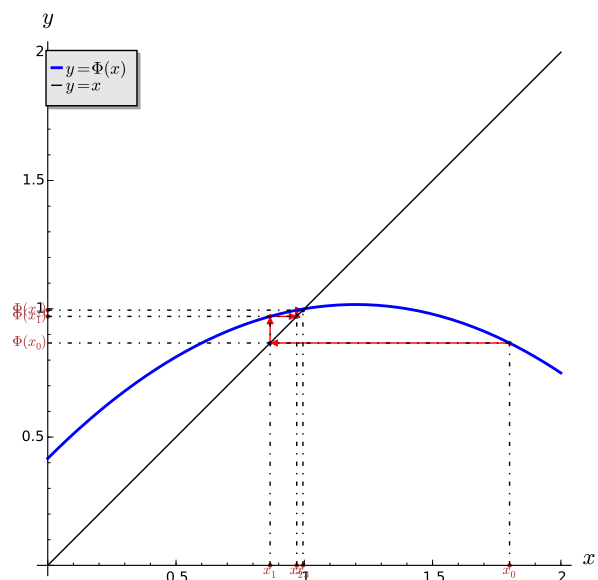


(b) Représentation point fixe

Figure 3.9: Exemple 1, méthode de la corde, $\alpha = 1$, racine de $f : x \mapsto x^2 - 1$ avec $a = 0.00$, $b = 2.00$, $x_0 = 1.80$,



(a) représentation usuelle



(b) Représentation point fixe

Figure 3.10: Exemple 2, méthode de la corde, $\alpha = 1$, racine de $f : x \mapsto x^2 - 1$ avec $a = 0.50$, $b = 1.90$, $x_0 = 1.80$,

pour quelques itérations.

	exemple 1	exemple 2
k	$ x_k - \alpha $	$ x_k - \alpha $
0	8.0000e-01	8.0000e-01
1	3.2000e-01	1.3333e-01
2	5.1200e-02	2.9630e-02
3	1.3107e-03	5.3041e-03
4	8.5899e-07	8.9573e-04
5	3.6893e-13	1.4962e-04
6	0.0000e+00	2.4947e-05
$\vdots$	$\vdots$	$\vdots$
15	0.0000e+00	2.4756e-12

On remarque qu'avec les données de l'exemple 1 la convergence est beaucoup plus rapide. Ceci est

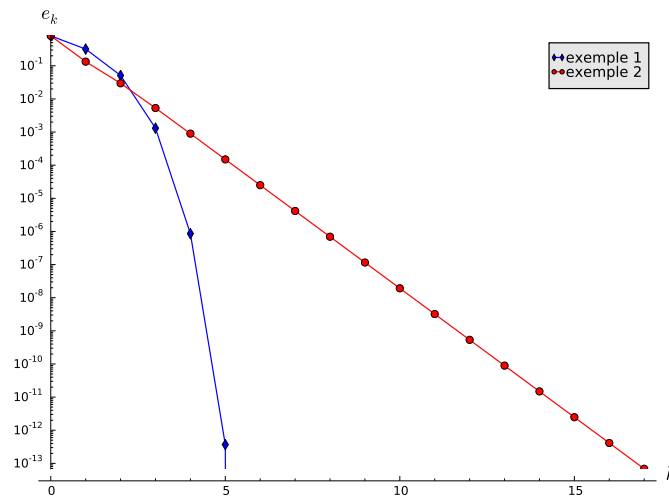


Figure 3.11: Erreurs en fonctions des itérations

illustré en Figure 3.11. En effet dans ce cas on a

$$\frac{f(b) - f(a)}{b - a} = 2 \quad \text{et} \quad f'(\alpha) = 2.$$

D'après la proposition 3.8, la convergence est d'ordre 2 pour l'exemple 1. Par contre, on a pour l'exemple 2

$$\frac{f(b) - f(a)}{b - a} = 2.400 \neq f'(\alpha) = 2$$

et donc la convergence est d'ordre 1. On retrouve ces résultats sur la Figure 3.12 où l'on a représenté en échelle logarithmique $e_{k+1} = |x_{k+1} - \alpha|$ en fonction de $e_k = |x_k - \alpha|$. En effet si une méthode est convergente d'ordre p exactement on aura pour k suffisamment grand $e_{k+1} \approx C e_k^p$.

Ajouter un autre exemple

La méthode de Newton

Soient f une fonction de classe $\mathcal{C}^1$ et x_0 un réel donné. La suite obtenue par la méthode de Newton est donnée par

$$x_{k+1} = x_k - \frac{f(x_k)}{f'(x_k)}, \quad \forall k \in \mathbb{N}. \quad (3.17)$$

Bien évidemment, il faudra s'assurer que $f'(x_k) \neq 0$. On peut voir cette relation comme un cas particulier de la méthode du point fixe. En effet, en prenant $\Phi(x) = x - \frac{f(x)}{f'(x)}$, la suite définie en (3.19) s'écrit $x_{k+1} = \Phi(x_k)$.

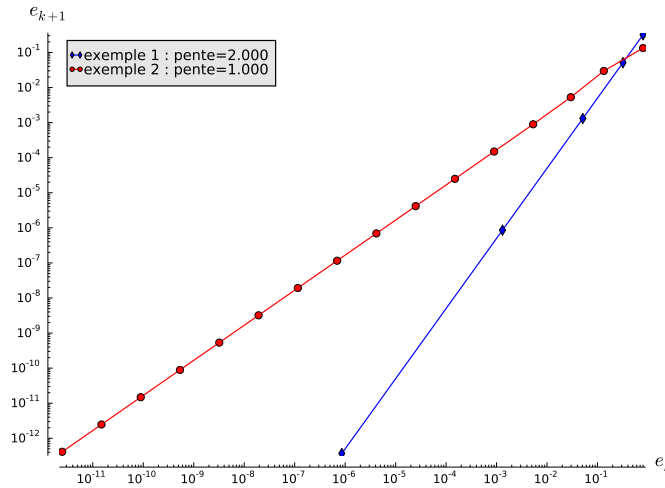


Figure 3.12: Représentation en échelle logarithmique de e_{k+1} en fonction de e_k . Les pentes sont calculées numériquement



Proposition 3.9: convergence de la méthode de Newton

Soit f une fonction de classe $\mathcal{C}^2$ sur un certain voisinage d'une racine simple α de f . Soit x_0 donné dans ce voisinage, la suite $(x_k)_{k \in \mathbb{N}}$ définie par la méthode de Newton

$$x_{k+1} = x_k - \frac{f(x_k)}{f'(x_k)}, \quad \forall k \in \mathbb{N}. \quad (3.18)$$

est localement convergente d'ordre 2.

Proof. Comme α est racine simple de f (i.e. $f(\alpha) = 0$ et $f'(\alpha) \neq 0$) et f' continue, il existe un voisinage $\mathcal{V}$ de α tel que pour tout $x \in \mathcal{V}$, $f'(x) \neq 0$. On peut alors définir la fonction Φ sur $\mathcal{V}$ par

$$\forall x \in \mathcal{V}, \quad \Phi(x) = x - \frac{f(x)}{f'(x)}.$$

On a alors $x_{k+1} = \Phi(x_k)$. La fonction Φ est de classe $\mathcal{C}^1$ sur $\mathcal{V}$ et

$$\Phi'(x) = 1 - \frac{(f'(x))^2 - f(x)f''(x)}{(f'(x))^2} = \frac{f(x)f''(x)}{(f'(x))^2}.$$

On a donc $\Phi'(\alpha) = 0$ car $f(\alpha) = 0$. D'après le théorème de convergence locale du point fixe (théorème 3.5), on en déduit que la suite (x_k) converge vers α (et que la convergence est au moins d'ordre 1. Pour démontrer qu'elle est d'ordre 2, on ne peut utiliser le théorème 3.6 car Φ n'est pas de classe $\mathcal{C}^2$. Toutefois comme f est de classe $\mathcal{C}^2$, on applique la formule de Taylor-Lagrange aux points x_k , α (en supposant $x_k \neq \alpha$) : il existe ξ_k compris entre x_k et α tel que

$$\underbrace{f(\alpha)}_{=0} = f(x_k) + (\alpha - x_k)f'(x_k) + \frac{(\alpha - x_k)^2}{2!}f^{(2)}(\xi_k).$$

Comme $f'(x_k) \neq 0$, l'équation précédente s'écrit aussi

$$\begin{aligned} \alpha &= x_k - \frac{f(x_k)}{f'(x_k)} + (\alpha - x_k)^2 \frac{f^{(2)}(\xi_k)}{2f'(x_k)} \\ &= x_{k+1} + (\alpha - x_k)^2 \frac{f^{(2)}(\xi_k)}{2f'(x_k)}. \end{aligned}$$

On obtient alors

$$\frac{\alpha - x_{k+1}}{(\alpha - x_k)^2} = \frac{f^{(2)}(\xi_k)}{2f'(x_k)}.$$

1 La fonction f est de classe $\mathcal{C}^2$ et la suite ξ_k converge vers α (car ξ_k compris entre x_k et α). Ceci
2 entraîne par passage à la limite que

$$\lim_{k \rightarrow +\infty} \frac{x_{k+1} - \alpha}{(x_k - \alpha)^2} = \frac{f^{(2)}(\alpha)}{2f'(\alpha)}.$$

3 La convergence est donc d'ordre 2. □

4 Soit f une fonction de classe $\mathcal{C}^2$ au voisinage de α , racine de f . On suppose que $f'(x) \neq 0$ pour tout
5 $x \in \mathcal{V}$ (i.e. α racine simple de f). Soit $x_0 \in \mathcal{V}$ donné. La suite obtenue par la méthode de Newton est
6 donnée par

$$x_{k+1} = x_k - \frac{f(x_k)}{f'(x_k)}, \quad \forall k \in \mathbb{N}. \quad (3.19)$$

7 On peut voir cette relation comme un cas particulier de la méthode du point fixe. En effet, en prenant
8 $\Phi(x) = x - \frac{f(x)}{f'(x)}$, la suite définie en (3.19) s'écrit $x_{k+1} = \Phi(x_k)$ et on note que $\Phi(\alpha) = \alpha$ (i.e. α point
9 fixe de Φ).

10 De plus f étant de classe $\mathcal{C}^3$ et sa dérivée non nulle sur $\mathcal{V}$, on obtient que Φ est de classe $\mathcal{C}^2$ sur $\mathcal{V}$, et
11 $\forall x \in \mathcal{V}$,

$$\Phi'(x) = 1 - \frac{(f'(x))^2 - f(x)f''(x)}{(f'(x))^2} = \frac{f(x)f''(x)}{(f'(x))^2}$$

12 et

$$\Phi''(x) = \frac{(f'(x)f''(x) + f(x)f^{(3)}(x))(f'(x))^2 - 2f(x)f'(x)(f''(x))^2}{(f'(x))^4}$$

13 On a alors

$$\Phi'(\alpha) = 0, \quad \Phi''(\alpha) = \frac{f''(\alpha)}{f'(\alpha)}.$$

14 D'après la proposition 3.6, si $f''(\alpha) \neq 0$ alors la **méthode de Newton est d'ordre 2**.

15 **Faire méthode de Newton modifiés dans le cas de racine de multiplicité $m > 1$???**



Exercice 3.2.3

En -1700 av. J.-C., les babyloniens ne connaissaient que les nombres rationnels (fractions) et ils utilisaient le système sexagésimal (base 60). Pour approcher la valeur $\sqrt{2}$, ils utilisaient comme approximation (voir tablette YBC 7289)

$$a = 1 + \frac{24}{60} + \frac{51}{60^2} + \frac{10}{60^3} = \frac{30547}{21600}$$

L'erreur commise est $|a - \sqrt{2}| \approx 5.994e - 7$.



Q. 1 Comment feriez-vous pour trouver à la main une méthode permettant de trouver des nombres rationnels approchant $\sqrt{2}$.

Q. 2 Généraliser la méthode pour trouver une approximation rationnelle de $\sqrt{a}$ où a est un réel positif.

Q. 3 Généraliser la méthode pour trouver une approximation rationnelle de $\sqrt[n]{a}$ où a est un réel positif et $n \in \mathbb{N}^*$.

Correction Exercice 3.2.3

Q. 1 Il suffit de voir que $\sqrt{2}$ est la racine positive de $f(x) = x^2 - 2$ et d'appliquer la méthode de Newton par exemple. La suite des itérés de Newton s'écrit alors

$$x_{k+1} = x_k - \frac{x_k^2 - 2}{2x_k} = \frac{x_k^2 + 2}{2x_k}$$

Avec $x_0 = 1$, on obtient

k	x_k	$ \sqrt{2} - x_k $
1	$\frac{3}{2}$	8.57864e-02
2	$\frac{17}{12}$	2.45310e-03
3	$\frac{577}{408}$	2.12390e-06

Avec $x_0 = \frac{3}{2}$, on obtient

k	x_k	$ \sqrt{2} - x_k $
1	$\frac{17}{12}$	2.45310e-03
2	$\frac{577}{408}$	2.12390e-06
3	$\frac{665857}{470832}$	1.59472e-12

Q. 2 Il suffit de voir que $\sqrt{a}$ est la racine positive de $f(x) = x^2 - a$ et d'appliquer la méthode de Newton par exemple. La suite des itérés de Newton s'écrit alors

$$x_{k+1} = x_k - \frac{x_k^2 - a}{2x_k} = \frac{x_k^2 + a}{2x_k}$$

Avec $a = 3$ et $x_0 = 1$, on obtient

k	x_k	$ \sqrt{3} - x_k $
1	2	2.67949e-01
2	$\frac{7}{4}$	1.79492e-02
3	$\frac{97}{56}$	9.20496e-05

Q. 3 Il suffit de voir que $\sqrt[n]{a}$ est la racine positive de $f(x) = x^n - a$ et d'appliquer la méthode de Newton par exemple. La suite des itérés de Newton s'écrit alors

$$x_{k+1} = x_k - \frac{x_k^n - a}{nx_k^{n-1}} = \frac{2x_k^n - x_k^n + a}{2x_k^{n-1}}$$

Avec $a = 3$, $n = 4$ et $x_0 = 1$, on obtient

k	x_k	$ \sqrt[4]{3} - x_k $
1	$\frac{3}{2}$	1.83926e-01
2	$\frac{97}{72}$	3.11482e-02
3	$\frac{115403137}{87616608}$	1.06368e-03
4	$\frac{236297297271008837816738085152257}{179546943199700984864483416264832}$	1.28780e-06

◇

On présente en Algorithme 3.14 l'implémentation standard de la méthode de Newton. Une autre version utilisant la fonction **PtFixe** est donnée en Algorithme 3.15.

Algorithme 3.14 Méthode de Newton**Données :**

f : $f : \mathbb{R} \rightarrow \mathbb{R}$,
 df : la dérivée de f ,
 x_0 : donnée initiale, $x_0 \in \mathbb{R}$,
 tol : la tolérance, $tol \in \mathbb{R}^+$,
 $kmax$: nombre maximum d'itérations, $kmax \in \mathbb{N}^*$

Résultat : α_{tol} : un réel tel que

```

1: Fonction  $\alpha_{TOL} \leftarrow \text{NEWTON}(f, df, x_0, tol, kmax)$ 
2:    $k \leftarrow 0, \alpha_{tol} \leftarrow \emptyset$ 
3:    $x \leftarrow x_0$ ,
4:    $err \leftarrow tol + 1$ 
5:   Tantque  $err > tol$  et  $k \leq kmax$  faire
6:      $k \leftarrow k + 1$ 
7:      $xp \leftarrow x$ 
8:      $x \leftarrow xp - f(xp)/df(xp)$   $\triangleright df(xp) \neq 0$ 
9:      $err \leftarrow |x - xp|$ 
10:  Fin Tantque
11:  Si  $err \leq tol$  alors  $\triangleright$  Convergence
12:     $\alpha_{tol} \leftarrow x$ 
13:  Fin Si
14: Fin Fonction

```

Algorithme 3.15 Méthode de Newton scalaire**Données :**

f : $f : \mathbb{R} \rightarrow \mathbb{R}$,
 df : la dérivée de f ,
 x_0 : donnée initiale, $x_0 \in \mathbb{R}$,
 tol : la tolérance, $tol \in \mathbb{R}^+$,
 $kmax$: nombre maximum d'itérations, $kmax \in \mathbb{N}^*$

Résultat : α_{tol} : un réel tel que

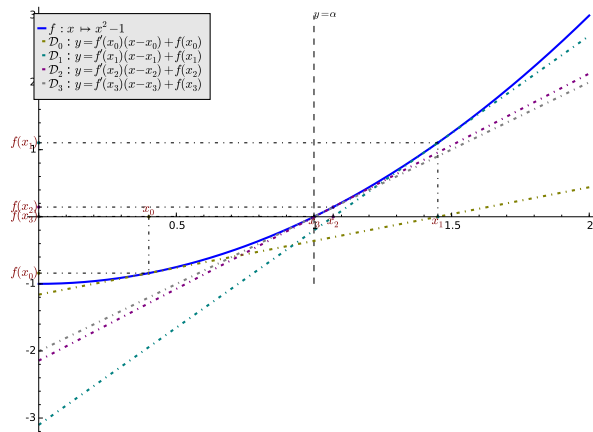
```

1: Fonction  $\alpha_{TOL} \leftarrow \text{NEWTON}(f, df, x_0, tol, kmax)$ 
2:    $\Phi \leftarrow x \mapsto x - f(x)/df(x)$ 
3:    $\alpha_{tol} \leftarrow \text{PTFIXE}(\Phi, x_0, tol, kmax)$ 
4: Fin Fonction

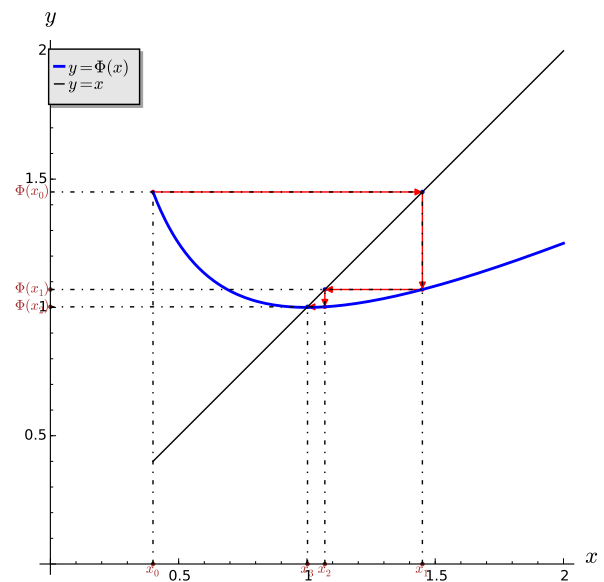
```

Comme premier exemple, on prend $f(x) = x^2 - 1$ avec $x_0 = 0.40$ et pour le second $f(x) = x^2 \cos(x)$ avec $x_0 = 2.00$.

On représente en Figures 3.13 and 3.14, respectivement pour les exemples 1 et 2, les itérations de la méthode de Newton à partir du graphe de la fonction f (figure de gauche) et à partir du graphe de la fonction Φ (figure de droite).



(a) représentation usuelle

(b) Représentation point fixe, $\Phi : x \mapsto x - \frac{x^2 - 1}{2x}$ Figure 3.13: Exemple 1, méthode de Newton, $\alpha = 1$, racine de $f : x \mapsto x^2 - 1$ avec $x_0 = 0.40$,

On illustre la convergence et l'ordre de convergence respectivement en Figures 3.15a et 3.15b. Pour cette dernière, on a représenté en échelle logarithmique $e_{k+1} = |x_{k+1} - \alpha|$ en fonction de $e_k = |x_k - \alpha|$. En effet si une méthode est convergente d'ordre p exactement on aura pour k suffisamment grand $e_{k+1} \approx C e_k^p$.

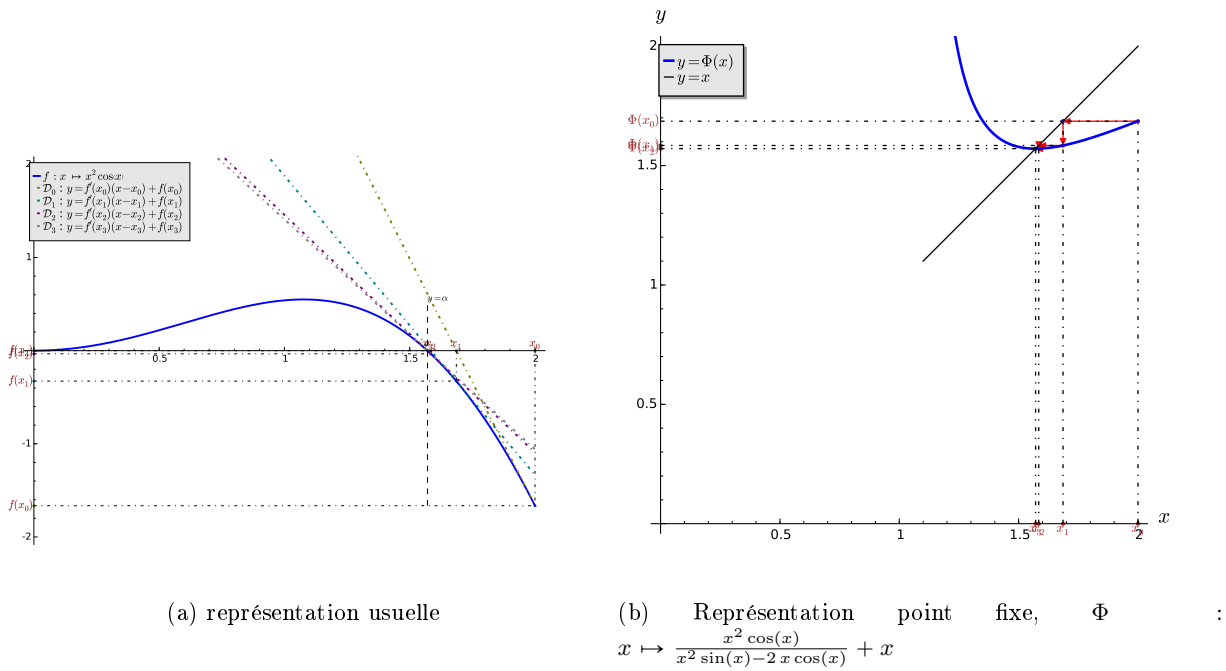


Figure 3.14: Exemple 2, méthode de Newton, $\alpha = \frac{1}{2} \pi$, racine de $f : x \mapsto x^2 \cos(x)$ avec $x_0 = 0.40$,

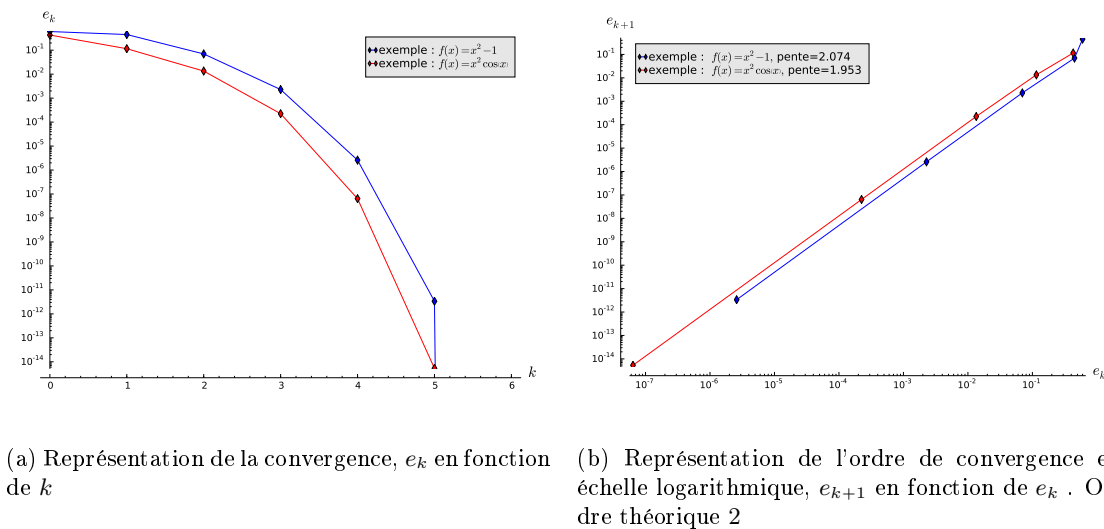


Figure 3.15: Méthode de Newton, convergence et ordre

3.2.5 La méthode de la sécante

Cette méthode est une alternative à la méthode de Newton lorsque l'on ne connaît pas la dérivée de la fonction f . A l'itération k , de la méthode de Newton, on approche $f'(x_k)$ par $\frac{f(x_k) - f(x_{k-1})}{x_k - x_{k-1}}$. En effet, d'après la formule de Taylor-Lagrange, il existe ξ_k compris entre x_{k-1} et x_k vérifiant

$$f(x_{k-1}) = f(x_k) + (x_{k-1} - x_k)f'(x_k) + \frac{(x_{k-1} - x_k)^2}{2!}f^{(2)}(\xi_k)$$

ce qui donne

$$f'(x_k) = \frac{f(x_k) - f(x_{k-1})}{x_k - x_{k-1}} + \frac{x_{k-1} - x_k}{2!}f^{(2)}(\xi_k).$$

- 2 Si la suite est convergente, on a $\frac{x_{k-1} - x_k}{2!}f^{(2)}(\xi_k) \rightarrow 0$, ce qui justifie l'approximation précédente. On a
 3 alors la méthode de la sécante donnée en (3.20). Cette méthode est une **méthode à deux pas** : le
 4 calcul de x_{k+1} nécessite de connaître x_k et x_{k-1} . Il faut donc deux valeurs d'initialisations x_{-1} et x_0
 5 pour définir la suite (x_k) .



Proposition 3.10: Convergence méthode de la sécante (Admis)

Soit f une fonction de classe C^2 sur un certain voisinage d'une racine simple α de f . Soient x_{-1} et x_0 donnés dans ce voisinage tels que $f(x_{-1}) \neq f(x_0)$, la suite $(x_k)_{k \in \mathbb{N}}$ définie par la méthode de la sécante

$$x_{k+1} = x_k - \frac{x_k - x_{k-1}}{f(x_k) - f(x_{k-1})}f(x_k), \quad \forall k \in \mathbb{N}. \quad (3.20)$$

est localement convergente d'ordre $\frac{1+\sqrt{5}}{2} \approx 1.618$.

Comme premier exemple, on recherche une racine de $x^2 - 1$ avec $x_{-1} = 0.000$ et $x_0 = 2.000$. Une représentation graphique des premières itérés de la suite est donnée en Figure 3.16. Sur cette figure, les droites $\mathcal{D}_k$ sont celles passant par les points $(x_{k-1}, f(x_{k-1}))$ et $(x_k, f(x_k))$. Pour deuxième exemple, on

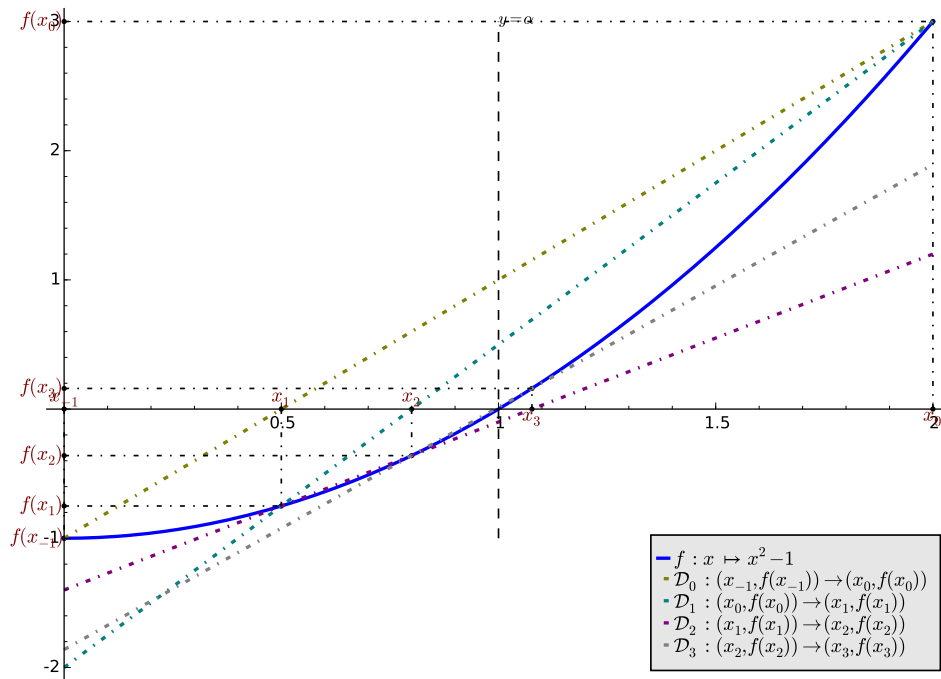
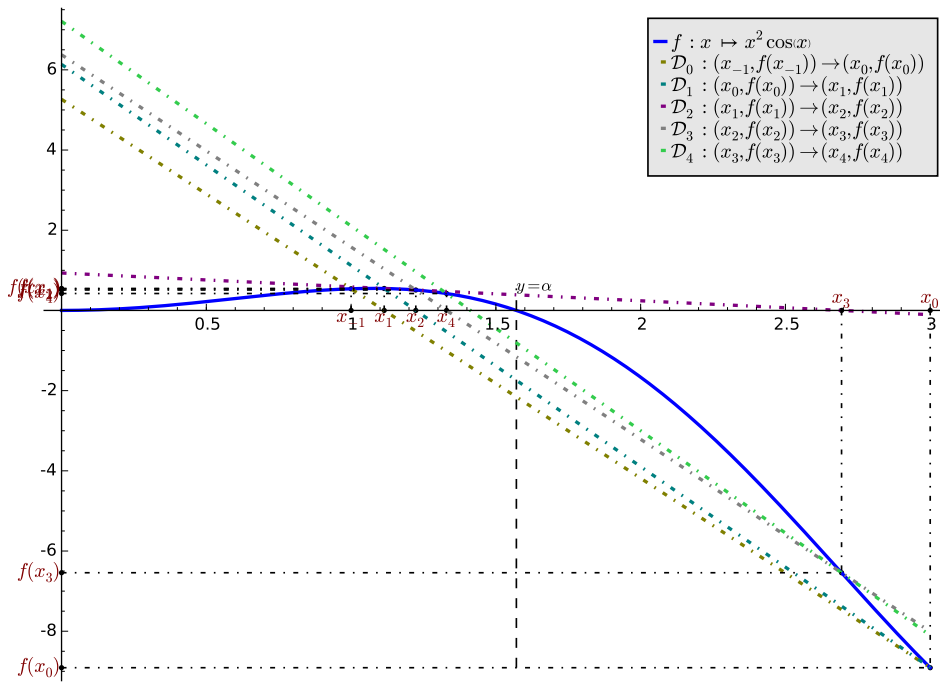


Figure 3.16: Méthode de la sécante pour $f(x) = x^2 - 1$, $x_{-1} = 0.000$ et $x_0 = 2.000$

recherche une racine de $x^2 \cos(x)$ avec $x_{-1} = 1.000$ et $x_0 = 3.000$. Une représentation graphique des premières itérés de la suite est donnée en Figure 3.17.

Figure 3.17: Méthode de la sécante pour $f(x) = x^2 \cos(x)$, $x_{-1} = 1.000$ et $x_0 = 3.000$

On illustre la convergence et l'ordre de convergence respectivement en Figures 3.18a et 3.18b. Pour cette dernière, on a représenté en échelle logarithmique $e_{k+1} = |x_{k+1} - \alpha|$ en fonction de $e_k = |x_k - \alpha|$. En effet si une méthode est convergente d'ordre p exactement on aura pour k suffisamment grand $e_{k+1} \approx C e_k^p$.

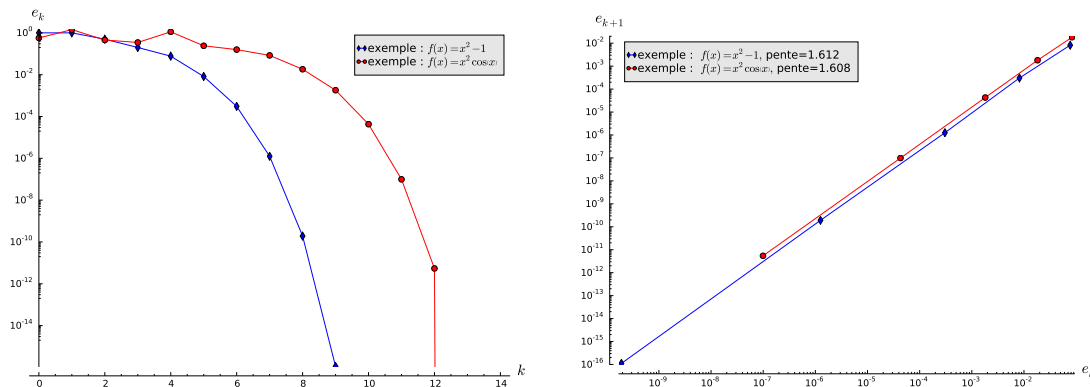
(a) Représentation de la convergence, e_k en fonction de k (b) Représentation de l'ordre de convergence en échelle logarithmique, e_{k+1} en fonction de e_k . Ordre théorique $\frac{1+\sqrt{5}}{2} \approx 1.618$

Figure 3.18: Méthode de la sécante, convergence et ordre

3.2.6 Méthode Regula-Falsi ou fausse position

Cette méthode diffère de la méthode de dichotomie par le choix de x_k à chaque itération. Pour la méthode de dichotomie on a pris $x_k = \frac{a_k + b_k}{2}$ point milieu du segment $[a_k, b_k]$. Pour la méthode Regula-Falsi, on prend pour x_k l'intersection de la droite passant par les points $(a_k, f(a_k))$ et $(b_k, f(b_k))$ avec l'axe des abscisses. Si $f(a_k)f(b_k) < 0$ cela nous assure que $x_k \in]a_k, b_k[$.

L'équation de la droite est donnée par

$$y = cx + d, \text{ avec } c = \frac{f(b_k) - f(a_k)}{b_k - a_k} \text{ et } d = \frac{b_k f(a_k) - a_k f(b_k)}{b_k - a_k}.$$

On a alors

$$x_k = -d/c = \frac{a_k f(b_k) - b_k f(a_k)}{f(b_k) - f(a_k)}.$$

1 En résumé, on définit les trois suites $(a_k)_{k \in \mathbb{N}}$, $(b_k)_{k \in \mathbb{N}}$ et $(x_k)_{k \in \mathbb{N}}$ par

$$2 \quad \bullet \quad a_0 = a, b_0 = b \text{ et } x_0 = \frac{a_0 f(b_0) - b_0 f(a_0)}{f(b_0) - f(a_0)},$$

$$\bullet \quad \forall k \in \mathbb{N},$$

$$\begin{cases} a_{k+1} = b_{k+1} = x_{k+1} = x_k, & \text{si } f(x_k) = 0, \\ a_{k+1} = x_k, b_{k+1} = b_k, & \text{si } f(b_k)f(x_k) < 0, \\ a_{k+1} = a_k, b_{k+1} = x_k, & \text{si } f(a_k)f(x_k) < 0, \\ x_{k+1} = \frac{a_{k+1}f(b_{k+1}) - b_{k+1}f(a_{k+1})}{f(b_{k+1}) - f(a_{k+1})}, & \text{si } f(x_k) \neq 0. \end{cases}$$



Exercice 3.2.4

On suppose que la fonction f est continue sur $[a, b]$, vérifie $f(a)f(b) < 0$ et qu'il existe un unique $\xi \in]a, b[$ tel que $f(\xi) = 0$.

Q. 1 Montrer que

$$a \leq \frac{af(b) - bf(a)}{f(b) - f(a)} \leq b.$$

Q. 2 Montrer que $a_0 \leq a_1 \leq \dots \leq a_k \leq x_k \leq b_k \leq \dots \leq b_1 \leq b_0$ pour tout $k \in \mathbb{N}$. et que si $f(x_k) \neq 0$ alors $f(a_k)f(b_k) < 0$.

Q. 3 En déduire la convergence de la suite (x_k) vers ξ .

Correction Exercice 3.2.4

Q. 1 On pose $x = \frac{af(b) - bf(a)}{f(b) - f(a)}$ qui est bien défini car $f(a) \neq f(b)$. En effet si $f(a) = f(b)$ alors $f(a)f(b) = f(a)^2 \geq 0$ qui est en contradiction avec l'hypothèse $f(a)f(b) < 0$. On a

$$x - a = \frac{af(b) - bf(a) - a(f(b) - f(a))}{f(b) - f(a)} = (b - a) \frac{f(a)}{f(a) - f(b)}$$

Comme $f(a)f(b) < 0$, $f(a) - f(b)$ est du même signe que $f(a)$ et alors $\frac{f(a)}{f(a) - f(b)} \geq 0$. De plus $b - a > 0$ et donc on a $x - a \geq 0$.

De la même manière, on a

$$x - b = \frac{af(b) - bf(a) - b(f(b) - f(a))}{f(b) - f(a)} = (a - b) \frac{f(b)}{f(b) - f(a)}$$

5 Comme $f(a)f(b) < 0$, $f(b) - f(a)$ est du même signe que $f(b)$ et alors $\frac{f(b)}{f(b) - f(a)} \geq 0$. De plus $a - b < 0$

6 et donc on a $x - b \leq 0$.

Q. 2 On va démontrer par récurrence la validité de la proposition $(\mathcal{P}_k)$ suivante $\forall k \in \mathbb{N}$:

$$(\mathcal{P}_k) \quad \begin{cases} \text{(i)} & f(a_k)f(b_k) < 0 \text{ si } f(x_k) \neq 0 \\ \text{(ii)} & a_0 \leq a_1 \leq \dots \leq a_k \leq x_k \leq b_k \leq \dots \leq b_1 \leq b_0, \end{cases}$$

7 **Initialisation** On vérifie la proposition $(\mathcal{P}_k)$ pour $k = 0$. On a $f(a)f(b) < 0$ donc $(\mathcal{P}_0) - (i)$ est vérifiée.

8 D'après **Q. 1**, on obtient $a_0 \leq x_0 \leq b_0$. La proposition $(\mathcal{P}_0)$ est donc vérifiée.

9 **Hérédité** Soit $k \geq 1$. On suppose la proposition $(\mathcal{P}_k)$ est vraie. Montrons que $(\mathcal{P}_{k+1})$ est vérifiée.

10 Si $f(x_k) = 0$ alors $a_{k+1} = b_{k+1} = x_{k+1} = x_k$ et la proposition $(\mathcal{P}_{k+1})$ est vérifiée.

11 On suppose maintenant que $f(x_k) \neq 0$.

12 De $(\mathcal{P}_k) - (i)$ on a $f(a_k) \neq 0$ et $f(b_k) \neq 0$. De plus $f(a_k) \neq f(b_k)$ car sinon $f(a_k)f(b_k) = f(a_k)^2 \geq 0$

13 ce qui est en contradiction avec $(\mathcal{P}_k) - (i)$. Par hypothèse $(\mathcal{P}_k)$, on a $a_k \leq x_k \leq b_k$ et $f(a_k)f(b_k) < 0$.

14 Par continuité de f , on a alors soit $f(b_k)f(x_k) < 0$ (et donc $f(a_k)f(x_k) > 0$) soit $f(b_k)f(x_k) > 0$

15 (et donc $f(a_k)f(x_k) < 0$).

- Si $f(b_k)f(x_k) < 0$ on a $a_{k+1} = x_k$ et $b_{k+1} = b_k$. Comme $f(a_k) \neq f(b_k)$, x_{k+1} est bien défini. D'après **Q. 1** en prenant $[a_{k+1}, b_{k+1}]$ comme intervalle, et sachant que $f(a_{k+1})f(b_{k+1}) = f(x_k)f(b_k) < 0$ on obtient

$$a_{k+1} \leq x_{k+1} \leq b_{k+1}.$$

De plus par hypothèse $a_k \leq x_k = a_{k+1}$ et donc $a_k \leq a_{k+1}$ et $b_{k+1} \leq b_k$. La proposition ($\mathcal{P}_{k+1}$) est donc vérifiée.

- Si $f(a_k)f(x_k) < 0$ on a $a_{k+1} = a_k$ et $b_{k+1} = x_k$. Comme $f(a_k) \neq f(b_k)$, x_{k+1} est bien défini. D'après **Q. 1** en prenant $[a_{k+1}, b_{k+1}]$ comme intervalle, et sachant que $f(a_{k+1})f(b_{k+1}) = f(a_k)f(x_k) < 0$ on obtient

$$a_{k+1} \leq x_{k+1} \leq b_{k+1}.$$

La proposition ($\mathcal{P}_{k+1}$) est donc vérifiée.

Q. 3 Supposons qu'il existe $s \in \mathbb{N}$ tel que $f(x_s) = 0$. Alors, pour tout $i \leq 1$, on a $a_{s+i} = b_{s+i} = x_{s+i} = x_s$. Les trois suites convergent donc vers x_s . D'après la question précédente, $x_s \in [a, b]$. Comme $f(a) \neq 0$ et $f(b) \neq 0$, on en déduit $x_s \in]a, b[$. Par hypothèse il existe un unique $\xi \in]a, b[$ tel que $f(\xi) = 0$, on a alors $x_s = \xi$.

Supposons que $\forall k \in \mathbb{N}$, $f(x_k) \neq 0$. D'après **Q. 2**, la suite (a_k) est croissante majorée par b et la suite (b_k) est décroissante minorée par a . Elles sont donc convergentes et l'on note respectivement l et L les limites de (a_k) et (b_k) . Comme $a \leq a_k \leq b_k \leq b$, on a $a \leq l \leq L \leq b$.

- Supposons $f(l) = f(L)$. On a $f(a_k)f(b_k) < 0$. Comme f est continue, à la limite on obtient $f(l)f(L) = f(l)^2 = f(L)^2 \leq 0$ et donc $f(l) = f(L) = 0$. On a nécessairement l et L dans $]a, b[$ car $f(a)f(b) < 0$ et donc $f(a)$ et $f(b)$ non nuls. Par unicité du zéro de f dans $]a, b[$ on obtient $l = L = \xi$. Comme $a_k \leq x_k \leq b_k$, on en déduit que la suite (x_k) converge aussi vers ξ .
- Supposons $f(l) \neq f(L)$. Par continuité de la fonction f la suite (x_k) converge alors vers $M = \frac{lf(L) - Lf(l)}{f(L) - f(l)}$. Comme $a_k \leq x_k \leq b_k$ on a aussi

$$l \leq M = \frac{lf(L) - Lf(l)}{f(L) - f(l)} \leq L. \quad (3.21)$$

De plus ayant $f(x_k) \neq 0 \forall k \in \mathbb{N}$, on a $f(a_k)f(b_k) < 0 \forall k \in \mathbb{N}$. En passant à la limite et par continuité de f on obtient $f(l)f(L) \leq 0$.

Montrons que $f(l) = 0$ ou $f(L) = 0$.

- Si $f(l) < f(L)$, alors de (3.21) on obtient

$$l(f(L) - f(l)) \leq lf(L) - Lf(l) \leq L(f(L) - f(l))$$

ce qui donne $lf(l) \geq Lf(l)$ et $lf(L) \leq Lf(L)$ i.e. $(l - L)f(l) \geq 0$ et $(L - l)f(L) \leq 0$. Comme $L - l \geq 0$, on en déduit $f(l) \leq 0$ et $f(L) \leq 0$. Or $f(l)f(L) \leq 0$, ce qui donne $f(l) = 0$ ou $f(L) = 0$.

- Si $f(l) > f(L)$, alors de (3.21) on obtient

$$l(f(L) - f(l)) \geq lf(L) - Lf(l) \geq L(f(L) - f(l))$$

ce qui donne $lf(l) \leq Lf(l)$ et $lf(L) \geq Lf(L)$ i.e. $(L - l)f(l) \geq 0$ et $(L - l)f(L) \leq 0$. Comme $L - l \geq 0$, on en déduit $f(l) \geq 0$ et $f(L) \geq 0$. Or $f(l)f(L) \leq 0$, ce qui donne $f(l) = 0$ ou $f(L) = 0$.

On a donc démontré que si $f(l) \neq f(L)$, alors $f(l) = 0$ ou $f(L) = 0$ et donc

- si $f(l) = 0$ alors $M = \frac{lf(L) - Lf(l)}{f(L) - f(l)} = l$ et donc $f(M) = 0$.
- si $f(L) = 0$ alors $M = \frac{lf(L) - Lf(l)}{f(L) - f(l)} = L$ et donc $f(M) = 0$.

Puisque l et L appartiennent à $]a, b[$, on a $M \in]a, b[$. Par unicité du zéro de f sur $]a, b[$ on en déduit que $M = \xi$.

◇

On vient de démontrer le théorème suivant

Théorème 3.11

Soit $f : [a, b] \subset \mathbb{R} \rightarrow \mathbb{R}$ une fonction continue vérifiant $f(a)f(b) < 0$ et admettant $\alpha \in]a, b[$ comme **unique** solution de $f(x) = 0$. Alors la suite $(x_k)_{k \in \mathbb{N}}$ définie par la **méthode Regula-Falsi** converge vers α

On a aussi la proposition suivante

Proposition 3.12: ordre de la méthode de Regula-Falsi

Soit $f : [a, b] \subset \mathbb{R} \rightarrow \mathbb{R}$ une fonction continue vérifiant $f(a)f(b) < 0$ et admettant $\alpha \in]a, b[$ comme **unique** solution de $f(x) = 0$. Si f est deux fois dérivable sur $[a, b]$ et si f'' est monotone sur $]a, b[$ alors il existe $C \in \mathbb{R}$ tel que

$$\lim_{k \rightarrow +\infty} \frac{|x_{k+1} - \alpha|}{|x_k - \alpha|} \leq C \quad (3.22)$$

La méthode de Regula-Falsi est alors à convergence linéaire (d'ordre 1).

3.3 Résolution de systèmes non linéaires

Nous allons (très rapidement) introduire les premières notions permettant la résolution de systèmes d'équations non linéaires. Par exemple, dans $\mathbb{R}^2$ nous allons regarder le problème suivant avec c une constante réelle

$$\begin{cases} f_1(x_1, x_2) = -x_1^3 + x_2 - \frac{1}{2} = 0 \\ f_2(x_1, x_2) = \frac{1}{25}(10x_2 + 1)^2 + c - x_1 = 0. \end{cases} \quad (3.23)$$

En Figures 3.19 et 3.21, on représente pour différentes valeurs de c les courbes $f_1(x_1, x_2) = 0$ et $f_2(x_1, x_2) = 0$: les intersections de ces deux courbes sont les solutions du problème (3.23). Comme on le voit graphiquement, il peut y avoir, suivant les valeurs de c , 0, 1, 2 ou 4 solutions.

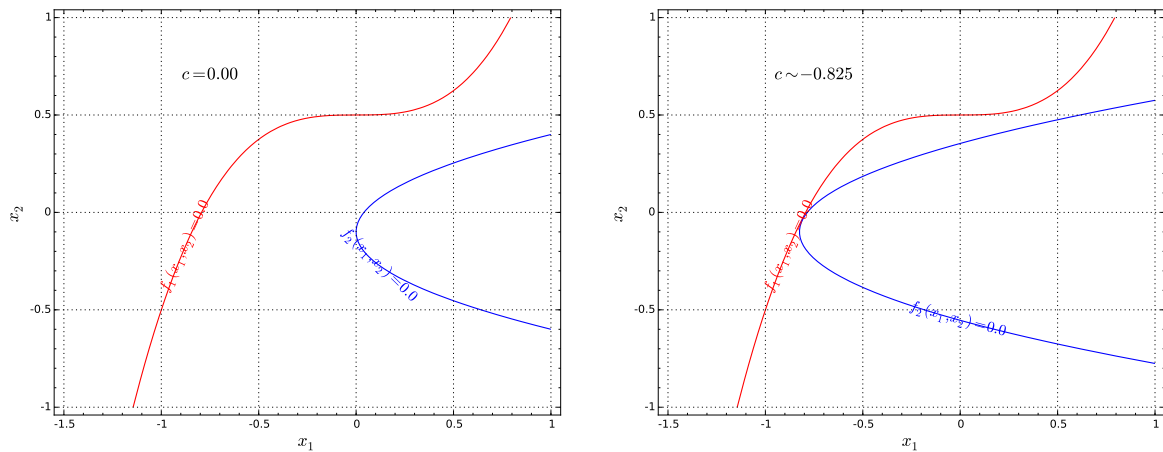


Figure 3.19: Résolution graphique de 3.23 avec $c = 0.00$ (gauche) et $c \sim -0.825$ (droite)

De manière plus générale, soient $U \subset \mathbb{R}^N$ un ouvert et $\mathbf{f}$ une application continue de U dans $\mathbb{R}^N$. Le problème que l'on souhaite résoudre est le suivant

Trouver $\alpha \in U \subset \mathbb{R}^N$ tel que

$$\mathbf{f}(\alpha) = 0 \iff \begin{cases} f_1(\alpha_1, \dots, \alpha_N) = 0 \\ f_2(\alpha_1, \dots, \alpha_N) = 0 \\ \vdots \\ f_N(\alpha_1, \dots, \alpha_N) = 0 \end{cases}$$

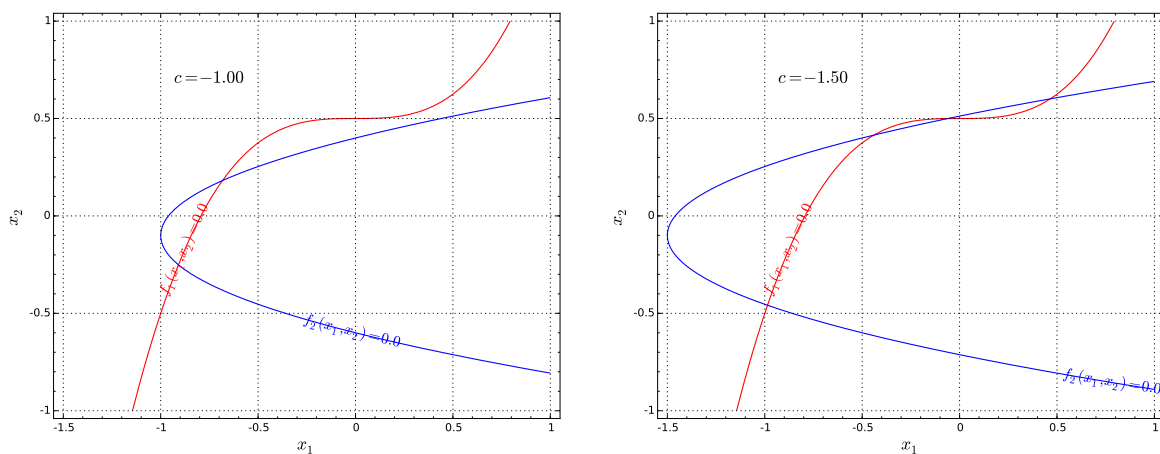


Figure 3.20: Résolution graphique de 3.23 avec $c = -1.00$ (gauche) et $c = -1.50$ (droite)

Comme dans le cas scalaire, pour résoudre numériquement ce genre de problème on utilise des suites itératives et plus particulièrement celles basées sur les méthodes de points fixes. En effet, nous allons voir que le théorème du point fixe se généralise très facilement (voir Théorème 3.13).

En définissant, par exemple, la fonction $\Phi \in C^0(U; \mathbb{R}^N)$ par $\Phi(\mathbf{x}) = \mathbf{x} + \mathbf{f}(\mathbf{x})$, on peut remarquer $\mathbf{f}(\mathbf{x}) = 0$ est équivalent à $\Phi(\mathbf{x}) = \mathbf{x}$. On peut donc se ramener à la recherche d'un point fixe (s'il existe) de la fonction Φ .

Trouver $\alpha \in U \subset \mathbb{R}^N$ tel que

$$\Phi(\alpha) = \alpha \iff \begin{cases} \Phi_1(\alpha_1, \dots, \alpha_N) = \alpha_1 \\ \Phi_2(\alpha_1, \dots, \alpha_N) = \alpha_2 \\ \vdots \\ \Phi_N(\alpha_1, \dots, \alpha_N) = \alpha_N \end{cases}$$

Les suites itératives sont donc de la forme

$$\mathbf{x}^{[k+1]} = \Phi(\mathbf{x}^{[k]})$$

où Φ est une fonction à déterminer et $\mathbf{x}^{[0]} \in \mathbb{R}^N$. Bien évidemment le choix d'une *bonne* fonction Φ est primordiale pour espérer avoir convergence.

Ce type de problème peut s'avérer délicat à traiter : comment choisir Φ ? $\mathbf{x}^{[0]}$? Si l'on converge vers quel point fixe?

3.3.1 Point fixe



Théorème 3.13

Soit $\mathcal{B}$ un espace de Banach et $U \subset \mathcal{B}$ un sous-ensemble fermé. On suppose que $\Phi : U \rightarrow U$ est une application strictement contractante, i.e.

$$\exists L \in]0, 1[, \quad \|\Phi(\mathbf{x}) - \Phi(\mathbf{y})\| \leq L \|\mathbf{x} - \mathbf{y}\|, \quad \forall (\mathbf{x}, \mathbf{y}) \in U \times U. \quad (3.24)$$

Alors

1. Φ admet un unique point fixe $\alpha \in U$ (i.e. unique solution de $\mathbf{x} = \Phi(\mathbf{x})$).
2. La suite des itérés $\mathbf{x}^{[k+1]} = \Phi(\mathbf{x}^{[k]})$ converge vers α pour toute valeur initiale $\mathbf{x}^{[0]} \in U$.
3. Pour tout $k \in \mathbb{N}$,

$$\|\alpha - \mathbf{x}^{[k]}\| \leq \frac{L^{k-l}}{1-L} \|\mathbf{x}^{[l+1]} - \mathbf{x}^{[l]}\|, \quad 0 \leq l \leq k \quad (3.25)$$

1

Proof. On démontre tout d'abord l'existence d'un point fixe. Pour cela on va démontrer que la suite $\mathbf{x}^{[k]}$ est de Cauchy dans U fermé d'un espace de Banach (donc elle converge dans U). Comme Φ est contractante, on a pour tout $k \in \mathbb{N}^*$

$$\|\mathbf{x}^{[k+1]} - \mathbf{x}^{[k]}\| = \|\Phi(\mathbf{x}^{[k]}) - \Phi(\mathbf{x}^{[k-1]})\| \leq L \|\mathbf{x}^{[k]} - \mathbf{x}^{[k-1]}\|$$

ce qui donne par récurrence pour tout $0 \leq j \leq k$

$$\|\mathbf{x}^{[k+1]} - \mathbf{x}^{[k]}\| \leq L^j \|\mathbf{x}^{[k+1-j]} - \mathbf{x}^{[k-j]}\| \quad (3.26)$$

ou encore pour tout $0 \leq l \leq k$, ($l = k - j$)

$$\|\mathbf{x}^{[k+1]} - \mathbf{x}^{[k]}\| \leq L^{k-l} \|\mathbf{x}^{[l+1]} - \mathbf{x}^{[l]}\| \quad (3.27)$$

On obtient aussi par récurrence

$$\forall l \geq 0, \quad \|\mathbf{x}^{[k+1+l]} - \mathbf{x}^{[k+l]}\| \leq L^l \|\mathbf{x}^{[k+1]} - \mathbf{x}^{[k]}\|. \quad (3.28)$$

Soit $p \geq 1$. On en déduit par application répétée de l'inégalité triangulaire que

$$\begin{aligned} \|\mathbf{x}^{[k+p]} - \mathbf{x}^{[k]}\| &= \|(\mathbf{x}^{[k+p]} - \mathbf{x}^{[k+p-1]}) + (\mathbf{x}^{[k+p-1]} - \mathbf{x}^{[k+p-2]}) + \dots + (\mathbf{x}^{[k+1]} - \mathbf{x}^{[k]})\| \\ &= \left\| \sum_{l=0}^{p-1} (\mathbf{x}^{[k+l+1]} - \mathbf{x}^{[k+l]}) \right\| \\ &\leq \sum_{l=0}^{p-1} \|\mathbf{x}^{[k+l+1]} - \mathbf{x}^{[k+l]}\| \\ &\stackrel{(3.28)}{\leq} \sum_{l=0}^{p-1} L^l \|\mathbf{x}^{[k+1]} - \mathbf{x}^{[k]}\| = \frac{1-L^p}{1-L} \|\mathbf{x}^{[k+1]} - \mathbf{x}^{[k]}\| \\ &\leq \frac{1-L^p}{1-L} L^k \|\mathbf{x}^{[1]} - \mathbf{x}^{[0]}\|. \quad (\text{en utilisant (3.27), avec } l=0) \end{aligned}$$

Comme $L^k \rightarrow 0$ quand $k \rightarrow +\infty$, on conclut que $(\mathbf{x}^{[k]})$ est une suite de Cauchy. De plus, par construction $\mathbf{x}^{[k]} \in U \subset \mathcal{B}$, pour tout $k \in \mathbb{N}$, et $\mathcal{B}$ étant un espace de Banach et U un fermé, la suite $(\mathbf{x}^{[k]})$ converge alors vers $\alpha \in U$. Comme Φ est contractante, elle est donc continue et en passant à la limite dans $\mathbf{x}^{[k+1]} = \Phi(\mathbf{x}^{[k]})$, on aboutit à $\alpha = \Phi(\alpha)$, i.e. α est un point fixe de Φ dans U .

L'unicité se déduit immédiatement par la contraction de la fonction Φ . En effet, soit α_1 et α_2 deux points fixes de Φ , alors

$$\|\alpha_1 - \alpha_2\| = \|\Phi(\alpha_1) - \Phi(\alpha_2)\| \leq L \|\alpha_1 - \alpha_2\|$$

2 Or $L < 1$, et donc nécessairement on a $\alpha_1 = \alpha_2$.

Il reste à démontrer l'inégalité (3.25). On a vu que pour $p \geq 1$

$$\|\mathbf{x}^{[k+p]} - \mathbf{x}^{[k]}\| \leq \frac{1 - L^p}{1 - L} \|\mathbf{x}^{[k+1]} - \mathbf{x}^{[k]}\|$$

L'application norme étant continue et $L < 1$, on obtient à la limite quand $p \rightarrow +\infty$

$$\|\boldsymbol{\alpha} - \mathbf{x}^{[k]}\| \leq \frac{1}{1 - L} \|\mathbf{x}^{[k+1]} - \mathbf{x}^{[k]}\|$$

On obtient l'inégalité en utilisant (3.27). □

3.3.2 Méthode de Newton

On commence par rappeler un résultat de calcul différentiel. Soit $\mathbf{f} : \mathbb{R}^N \rightarrow \mathbb{R}^N$ une fonction suffisamment régulière. On définit la **matrice Jacobienne de $\mathbf{f}$** , notée $\mathbb{J}_{\mathbf{f}}$, par

$$\mathbb{J}_{\mathbf{f}} = \begin{pmatrix} \frac{\partial f_1}{\partial x_1} & \frac{\partial f_1}{\partial x_2} & \cdots & \frac{\partial f_1}{\partial x_N} \\ \frac{\partial f_2}{\partial x_1} & \frac{\partial f_2}{\partial x_2} & \cdots & \frac{\partial f_2}{\partial x_N} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{\partial f_N}{\partial x_1} & \frac{\partial f_N}{\partial x_2} & \cdots & \frac{\partial f_N}{\partial x_N} \end{pmatrix}$$

On a alors $\forall \mathbf{h} \in \mathbb{R}^N$ à l'ordre 1

$$\mathbf{f}(\mathbf{x} + \mathbf{h}) \approx \mathbf{f}(\mathbf{x}) + \mathbb{J}_{\mathbf{f}}(\mathbf{x}) \cdot \mathbf{h}. \quad (3.29)$$

Nous voulons trouver $\boldsymbol{\alpha}$ tel que $\mathbf{f}(\boldsymbol{\alpha}) = 0$. Si $\mathbf{x}^{[k]}$ est proche de $\boldsymbol{\alpha}$, alors en utilisant (3.29) avec $\mathbf{x} = \mathbf{x}^{[k]}$ et $\boldsymbol{\alpha} = \mathbf{x}^{[k]} + \mathbf{h}$ (i.e. $\mathbf{h} = \boldsymbol{\alpha} - \mathbf{x}^{[k]}$) on obtient

$$\mathbf{f}(\boldsymbol{\alpha}) \approx \mathbf{f}(\mathbf{x}^{[k]}) + \mathbb{J}_{\mathbf{f}}(\mathbf{x}^{[k]}) \cdot \mathbf{h}$$

Au lieu de résoudre $\mathbf{f}(\mathbf{x}) = 0$, on résout le système linéarisé

$$\mathbf{f}(\mathbf{x}^{[k]}) + \mathbb{J}_{\mathbf{f}}(\mathbf{x}^{[k]}) \cdot \tilde{\mathbf{h}} = 0$$

c'est à dire le système linéaire

$$\mathbb{J}_{\mathbf{f}}(\mathbf{x}^{[k]}) \cdot \tilde{\mathbf{h}} = -\mathbf{f}(\mathbf{x}^{[k]}). \quad (3.30)$$

On espère alors que $\tilde{\mathbf{h}}$ est une bonne approximation de $\mathbf{h}$ au sens où $\mathbf{x}^{[k]} + \tilde{\mathbf{h}}$ est une meilleure approximation de $\boldsymbol{\alpha}$ que $\mathbf{x}^{[k]}$. On note alors $\mathbf{x}^{[k+1]} = \mathbf{x}^{[k]} + \tilde{\mathbf{h}}$. En posant $\Phi(\mathbf{x}) = \mathbf{x} - ((\mathbb{J}_{\mathbf{f}}(\mathbf{x}))^{-1} \mathbf{f}(\mathbf{x}))$ la méthode de Newton s'écrit alors

$$\mathbf{x}^{[k+1]} = \Phi(\mathbf{x}^{[k]}) = \mathbf{x}^{[k]} - ((\mathbb{J}_{\mathbf{f}}(\mathbf{x}^{[k]}))^{-1} \mathbf{f}(\mathbf{x}^{[k]})) \quad (3.31)$$

Cette formule est une généralisation de celle vue dans le cas scalaire (voir ??). Il faut noter qu'à chaque itération la matrice Jacobienne est modifiée et qu'il faut calculer le vecteur $-((\mathbb{J}_{\mathbf{f}}(\mathbf{x}^{[k]}))^{-1} \mathbf{f}(\mathbf{x}^{[k]}))$. Numériquement, on ne calcule que très rarement l'inverse d'une matrice car cela est très coûteux en temps mais on résout le système linéaire (3.30) ce qui est bien plus efficace.

On admet dans ce cours le théorème suivant



Théorème 3.14

Soit $\mathbf{f} \in \mathcal{C}^3(\mathbb{R}^N; \mathbb{R}^N)$. On suppose que la matrice Jacobienne appliquée en $\mathbf{x}$, $\mathbb{J}_{\mathbf{f}}(\mathbf{x})$ est inversible dans un voisinage de $\boldsymbol{\alpha}$, avec $\mathbf{f}(\boldsymbol{\alpha}) = 0$. Alors pour tout $\mathbf{x}^{[0]}$ suffisamment proche de $\boldsymbol{\alpha}$ la suite définie par

$$\mathbf{x}^{[k+1]} = \mathbf{x}^{[k]} - ((\mathbb{J}_{\mathbf{f}}(\mathbf{x}^{[k]}))^{-1} \mathbf{f}(\mathbf{x}^{[k]}))$$

converge vers $\boldsymbol{\alpha}$ et la convergence est d'ordre 2.

On donne ensuite l'algorithme 3.16 permettant de déterminer une approximation d'un point fixe d'une fonction $\mathbf{f}$. Dans cet algorithme on suppose donnée la fonction **SOLVE** permettant de résoudre un système linéaire.

Algorithme 3.16 Méthode de Newton**Données :**

$\mathbf{f}$: $\mathbf{f} : \mathbb{R}^N \rightarrow \mathbb{R}^N$,
 $\mathbf{Jf}$: la matrice Jacobienne de $\mathbf{f}$,
 $\mathbf{x0}$: donnée initiale, $\mathbf{x0} \in \mathbb{R}^N$,
 tol : la tolérance, $\text{tol} \in \mathbb{R}^+$,
 kmax : nombre maximum d'itérations, $\text{kmax} \in \mathbb{N}^*$

Résultat :

α_{tol} : un élément de $\mathbb{R}^N$ proche de α .

1: **Fonction** $\alpha_{\text{TOL}} \leftarrow \text{NEWTON}(\mathbf{f}, \mathbf{Jf}, \mathbf{x0}, \text{tol}, \text{kmax})$

2: $k \leftarrow 0, \alpha_{\text{tol}} \leftarrow \emptyset$

3: $\mathbf{x} \leftarrow \mathbf{x0}$,

4: $\text{err} \leftarrow \text{tol} + 1$

5: **Tantque** $\text{err} > \text{tol}$ et $k \leq \text{kmax}$ **faire**

6: $k \leftarrow k + 1$

7: $\mathbf{xp} \leftarrow \mathbf{x}$

8: $\mathbf{h} \leftarrow \text{SOLVE}(\mathbf{Jf}(\mathbf{xp}), -\mathbf{f}(\mathbf{xp}))$

$\triangleright \mathbf{x} \leftarrow \text{SOLVE}(\mathbf{A}, \mathbf{b})$: résoud le système linéaire $\mathbf{Ax} = \mathbf{b}$

9: $\mathbf{x} \leftarrow \mathbf{xp} + \mathbf{h}$

10: $\text{err} \leftarrow \text{NORM}(\mathbf{x} - \mathbf{xp})$

11: **Fin Tantque**

12: **Si** $\text{err} \leq \text{tol}$ **alors**

$\triangleright$ Convergence

13: $\alpha_{\text{tol}} \leftarrow \mathbf{x}$

14: **Fin Si**

15: **Fin Fonction**

- 1 **Remarque 3.15** Si l'on ne connaît pas explicitement la Jacobienne de $\mathbf{f}$, il est possible de calculer une
 2 approximation de celle-ci en utilisant des formules de dérivation numérique.

3.3.3 Exemples

Exemple modèle

Comme premier exemple, nous reprenons le système 3.23 avec $c = -1.5$

$$\begin{cases} f_1(x_1, x_2) = -x_1^3 + x_2 - \frac{1}{2} & = 0 \\ f_2(x_1, x_2) = \frac{1}{25}(10x_2 + 1)^2 - x_1 - \frac{3}{2} & = 0. \end{cases} \quad (3.32)$$

On représente en Figure 3.21 les itérées successives pour 4 suites avec une initialisation différentes. On remarque qu'il est très difficile, si l'on n'est pas suffisamment proche d'un point fixe, de prédire vers lequel on converge.

En Figure 3.22a, on représente les bassins d'attraction pour les itérées de Newton associés au système 3.32 : à chaque point initial $x_0 = (x, y)$ on associe le point fixe vers lequel la suite de Newton converge et chaque point fixe correspond une couleur. En Figure 3.22b, on représente le nombre d'itérations assurant la convergence des itérées de Newton : à chaque point initial $x_0 = (x, y)$ on associe le nombre d'itérations nécessaire à la convergence et une échelle de couleur permet de visualiser ces nombres.

Exemple complexe : $z^3 - 1 = 0$

On souhaite trouver les racines complexes de $z^3 - 1$. Pour cela on peut poser $z = x + iy$, et le système équivalent devient

$$\begin{cases} f_1(x, y) = x^3 - 3xy^2 - 1 & = 0 \\ f_2(x, y) = 3x^2y - y^3 & = 0. \end{cases} \quad (3.33)$$

Bien évidemment en restant dans le corps des complexes, l'algorithme de Newton est le même (encore faut-il que le langage de programmation utilisé le supporte).

On représente en Figure 3.23 les bassins d'attraction et le nombre d'itérations de convergence associé à des données initiales dans $[-1.5, 1.5] \times [-1.5, 1.5]$. On obtient alors une *fractale de Newton*. Pour illustrer ce caractère fractale des représentations, on donne en Figures 3.24 et 3.25 des zooms successifs sur les graphes.

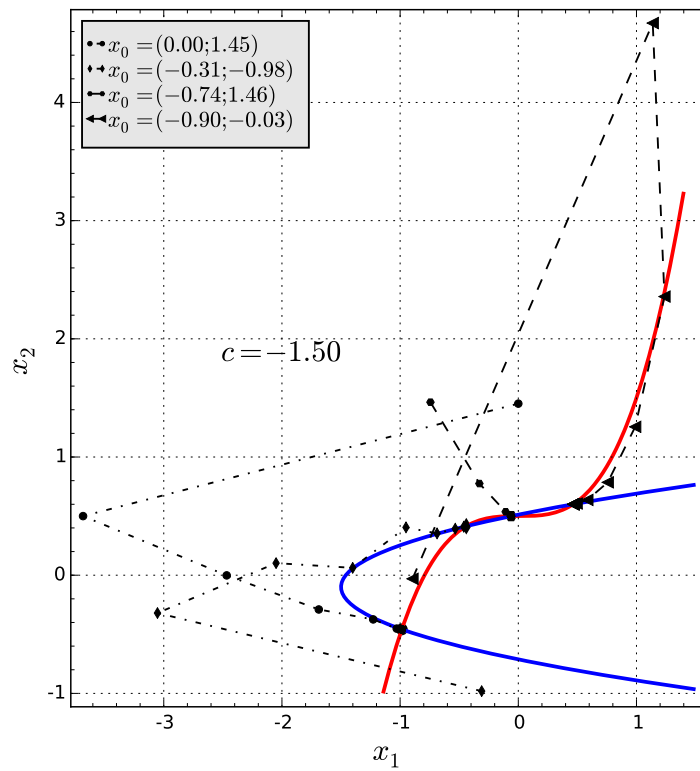


Figure 3.21: Représentation de 4 suites pour le système 3.32

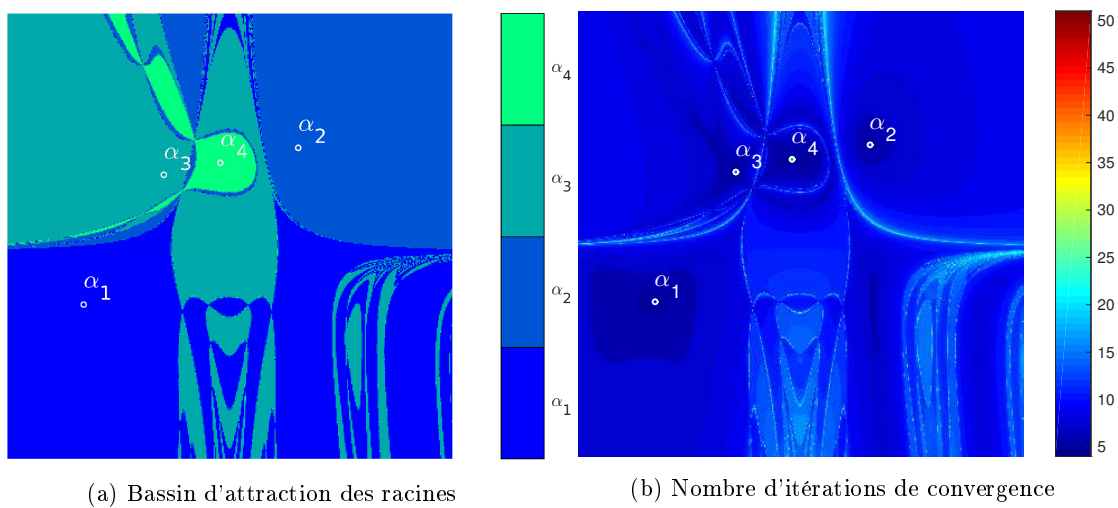


Figure 3.22: Méthode de Newton, système (3.32)

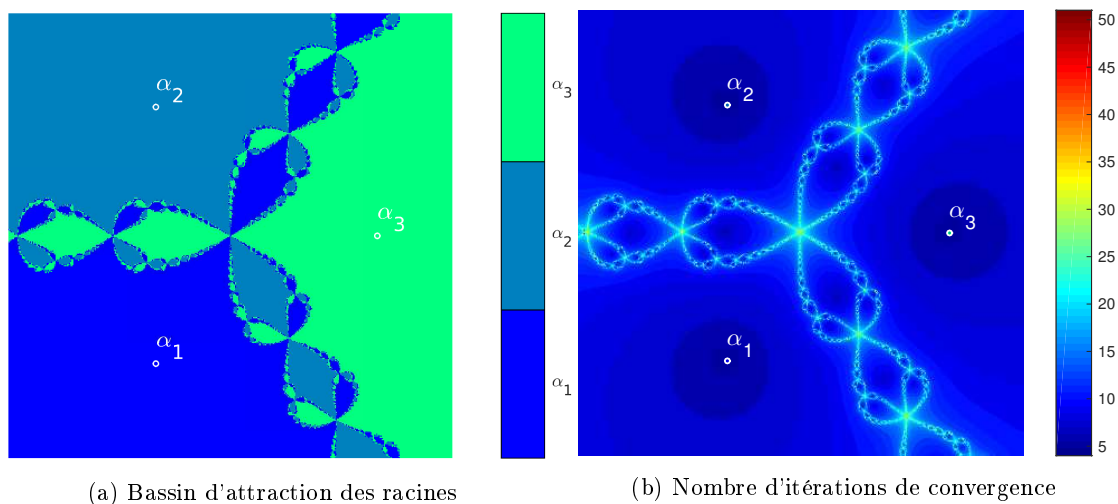


Figure 3.23: Méthode de Newton, système (3.33)

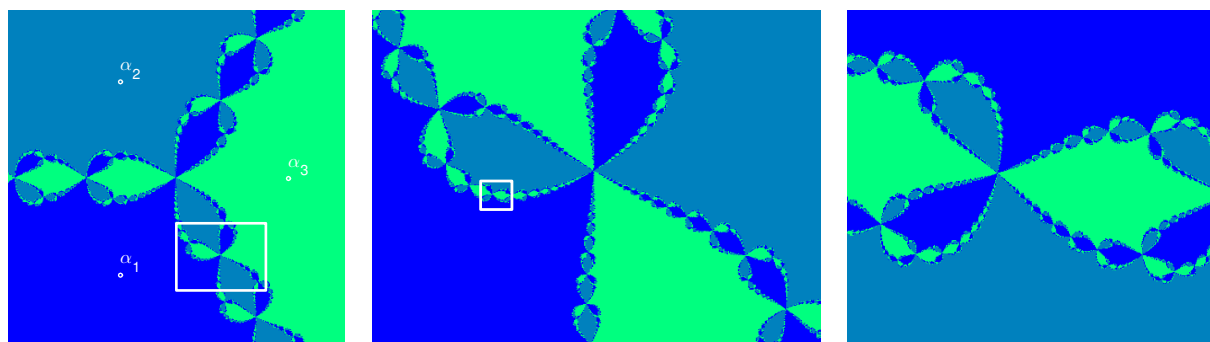


Figure 3.24: Méthode de Newton, système (3.33), zooms sur les bassins d'attraction

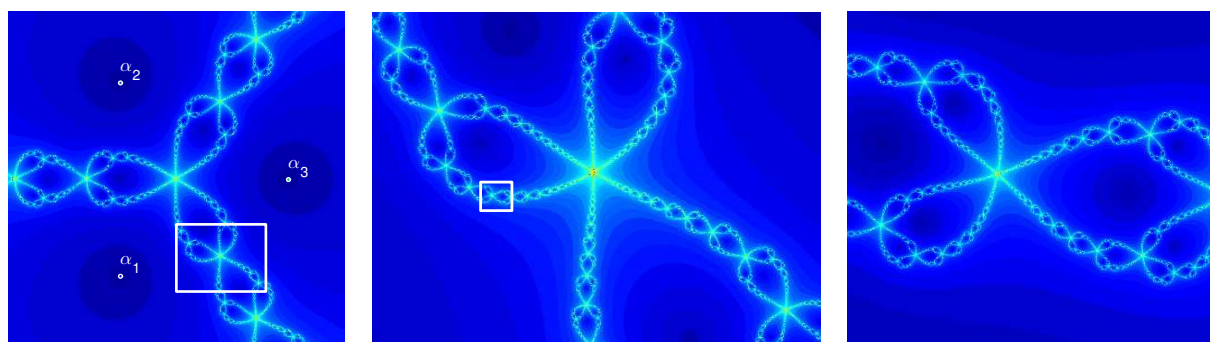
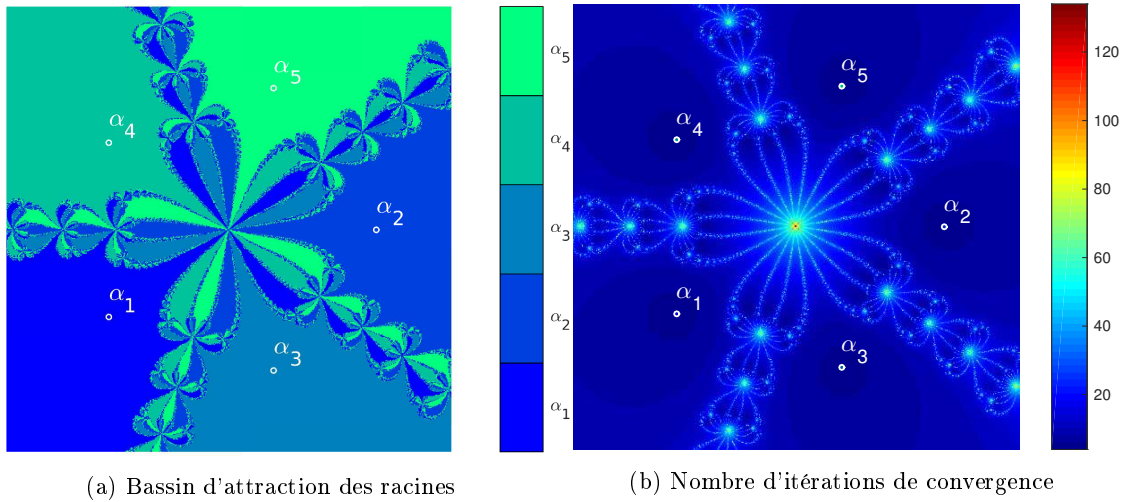


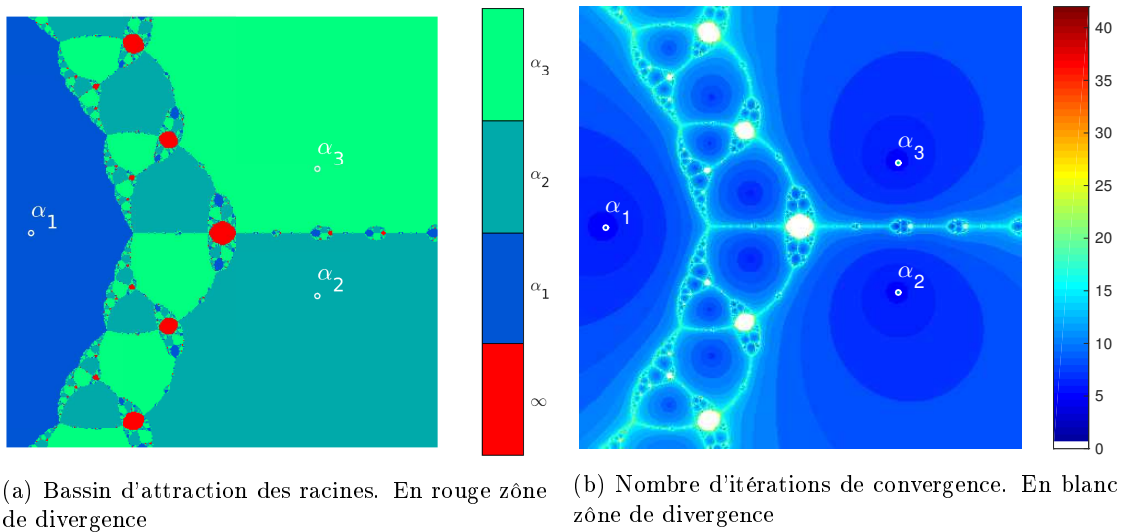
Figure 3.25: Méthode de Newton, système (3.33), zooms sur les nombres d'itérations

Exemple complexe : $z^5 - 1 = 0$

On représente en Figure 3.27 les bassins d'attraction et le nombre d'itérations de convergence associé à des données initiales dans $[-1.5, 1.5] \times [-1.5, 1.5]$.

Figure 3.26: Méthode de Newton pour $z^5 - 1 = 0$ **Exemple complexe :** $z^3 - 2z + 2 = 0$

On représente en Figure ?? les bassins d'attraction et le nombre d'itérations de convergence associé à des données initiales dans $[-2, 2] \times [-2, 2]$.

Figure 3.27: Méthode de Newton pour $z^3 - 2z + 2 = 0$

Chapitre 4

Résolution de systèmes linéaires

Dans cette partie nous allons considérer la résolution numérique d'un système linéaire $\mathbb{A}\mathbf{x} = \mathbf{b}$ dont la matrice $\mathbb{A}$ est inversible.

On pourrait penser que pour résoudre le système linéaire $\mathbf{Ax} = \mathbf{b}$, $\mathbb{A}$ inversible, le plus simple serait de calculer la matrice $\mathbb{A}^{-1}$, inverse de $\mathbb{A}$, puis d'effectuer un produit matrice-vecteur pour obtenir $\mathbf{x} = \mathbb{A}^{-1}\mathbf{b}$. Or pour calculer l'inverse d'une matrice d'ordre n on doit résoudre n systèmes linéaires d'ordre n ! En effet, déterminer l'inverse d'une matrice $\mathbb{A}$ revient à rechercher la matrice $\mathbb{X}$ solution de

$$\mathbb{A}\mathbb{X} = \mathbb{I} \iff \begin{pmatrix} A_{1,1} & A_{1,2} & \dots & A_{1,n} \\ A_{2,1} & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & A_{n-1,n} \\ A_{n,1} & \dots & A_{n,n-1} & A_{n,n} \end{pmatrix} \begin{pmatrix} X_{1,1} & X_{1,2} & \dots & X_{1,n} \\ X_{2,1} & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & X_{n-1,n} \\ X_{n,1} & \dots & X_{n,n-1} & X_{n,n} \end{pmatrix} = \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \dots & 0 & 1 \end{pmatrix}$$

Si on note $\mathbf{X}_{i,:}$ le i -ème vecteur colonne de la matrice $\mathbb{X}$ et $\mathbf{e}_i$ le i -ème de la base canonique de $\mathbb{R}^n$ alors le système précédant s'écrit

$$\begin{pmatrix} A_{1,1} & A_{1,2} & \dots & A_{1,n} \\ A_{2,1} & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & A_{n-1,n} \\ A_{n,1} & \dots & A_{n,n-1} & A_{n,n} \end{pmatrix} \left(\begin{array}{c|c|c|c} \mathbf{X}_1 & \dots & \mathbf{X}_n \end{array} \right) = \left(\begin{array}{c|c|c|c} \mathbf{e}_1 & \dots & \mathbf{e}_n \end{array} \right)$$

Ce dernier système est alors équivalent à résoudre les n systèmes linéaires

$$\mathbb{A}\mathbf{X}_j = \mathbf{e}_j, \quad \forall j \in \llbracket 1, n \rrbracket.$$



Pour résoudre un système linéaire, on ne calcule pas la matrice inverse associée.

Nous allons en section 4.2 étudier quelques **méthodes directes** pour la résolution d'un système linéaire basées sur la recherche d'une matrice $\mathbb{M}$ inversible telle que la matrice $\mathbb{M}\mathbb{A}$ soit triangulaire supérieure. Ceci conduit à la résolution du système linéaire équivalent

$$\mathbb{M}\mathbb{A}\mathbf{x} = \mathbb{M}\mathbf{b}.$$

par la *méthode de la remontée* décrite en section 4.2.1).

- 1 En section 4.5, nous nous intéresserons aux **méthodes itératives** pour la résolution d'un système
 2 linéaire qui peuvent s'écrire sous la forme

$$\mathbf{x}^{[k+1]} = \mathbb{B}\mathbf{x}^{[k]} + \mathbf{c}, \quad k \geq 0, \quad \mathbf{x}^{[0]} \text{ donné}$$

- 3 où la matrice $\mathbb{B}$ et le vecteur $\mathbf{c}$ sont construits à partir de la matrice $\mathbb{A}$ et du vecteur $\mathbf{b}$. On espère alors
 4 avoir $\lim_{k \rightarrow +\infty} \mathbf{x}^{[k]} = \mathbf{x}$.

4.1 Exercices et résultats préliminaires



Exercice 4.1.1: correction page 193

Soit $\mathbb{A} \in \mathcal{M}_{n,n}(\mathbb{C})$ une matrice et $(\lambda, \mathbf{u})$ un élément propre de $\mathbb{A}$ avec $\|\mathbf{u}\|_2 = 1$.

Q. 1 En s'aidant de la base canonique $\{\mathbf{e}_1, \dots, \mathbf{e}_n\}$, construire une base orthonormée $\{\mathbf{x}_1, \dots, \mathbf{x}_n\}$ telle que $\mathbf{x}_1 = \mathbf{u}$.

Notons $\mathbb{P}$ la matrice de changement de base canonique $\{\mathbf{e}_1, \dots, \mathbf{e}_n\}$ dans la base $\{\mathbf{x}_1, \dots, \mathbf{x}_n\}$:

$$\mathbb{P} = \begin{pmatrix} | & & | \\ \mathbf{x}_1 & \dots & \mathbf{x}_n \\ | & & | \end{pmatrix}$$

Soit $\mathbb{B}$ la matrice définie par $\mathbb{B} = \mathbb{P}^* \mathbb{A} \mathbb{P}$.

Q. 2 1. Exprimer les coefficients de la matrice $\mathbb{B}$ en fonction de la matrice $\mathbb{A}$ et des vecteurs $\mathbf{x}_i$, $i \in \llbracket 1, n \rrbracket$.

$$\mathbb{B} = \mathbb{P}^* \mathbb{A} \mathbb{P}.$$

2. En déduire que la première colonne de $\mathbb{B}$ est $(\lambda, 0, \dots, 0)^t$.

Q. 3 Montrer par récurrence sur l'ordre de la matrice que la matrice $\mathbb{A}$ s'écrit

$$\mathbb{A} = \mathbb{U} \mathbb{T} \mathbb{U}^*$$

où $\mathbb{U}$ est une matrice unitaire et $\mathbb{T}$ une matrice triangulaire supérieure.

Q. 4 En supposant $\mathbb{A}$ inversible et la décomposition $\mathbb{A} = \mathbb{U} \mathbb{T} \mathbb{U}^*$ connue, expliquer comment résoudre "simplement" le système linéaire $\mathbb{A} \mathbf{x} = \mathbf{b}$.

Correction Exercice 4.1.1

- 8 **Q. 1** La première chose à faire est de construire une base contenant $\mathbf{u}$ à partir de la base canon-
 9 ique $\{\mathbf{e}_1, \dots, \mathbf{e}_n\}$. Comme le vecteur propre $\mathbf{u}$ est non nul, il existe $j \in \llbracket 1, n \rrbracket$ tel que $\langle \mathbf{u}, \mathbf{e}_j \rangle \neq 0$. La
 10 famille $\{\mathbf{u}, \mathbf{e}_1, \dots, \mathbf{e}_{j-1}, \mathbf{e}_{j+1}, \dots, \mathbf{e}_n\}$ forme alors une base de $\mathbb{C}^n$ car $\mathbf{u}$ n'est pas combinaison linéaire des
 11 $\{\mathbf{e}_1, \dots, \mathbf{e}_{j-1}, \mathbf{e}_{j+1}, \dots, \mathbf{e}_n\}$.

- 12 On note $\{\mathbf{z}_1, \dots, \mathbf{z}_n\}$ la base dont le premier élément est $\mathbf{z}_1 = \mathbf{u}$:

$$\{\mathbf{z}_1, \dots, \mathbf{z}_n\} = \{\mathbf{u}, \mathbf{e}_1, \dots, \mathbf{e}_{j-1}, \mathbf{e}_{j+1}, \dots, \mathbf{e}_n\}.$$

On peut ensuite utiliser le **procédé de Gram-Schmidt**, rappelé en Proposition A.12, pour construire une base orthonormée à partir de cette base.

On calcule successivement les vecteurs $\mathbf{x}_i$ à partir de la base $\{\mathbf{z}_1, \dots, \mathbf{z}_n\}$ en construisant un vecteur $\mathbf{w}_i$ orthogonal aux vecteurs $\mathbf{x}_1, \dots, \mathbf{x}_{i-1}$.

$$\mathbf{w}_i = \mathbf{z}_i - \sum_{k=1}^{i-1} \langle \mathbf{x}_k, \mathbf{z}_i \rangle \mathbf{x}_k$$

puis on obtient le vecteur $\mathbf{x}_i$ en normalisant

$$\mathbf{x}_i = \frac{\mathbf{w}_i}{\|\mathbf{w}_i\|}$$

Q. 2 1. En conservant l'écriture colonne de la matrice $\mathbb{P}$ on obtient

$$\mathbb{B} = \begin{pmatrix} \mathbf{x}_1^* \\ \mathbf{x}_2^* \\ \vdots \\ \mathbf{x}_n^* \end{pmatrix} \mathbb{A} \begin{pmatrix} \mathbf{x}_1 & \mathbf{x}_2 & \dots & \mathbf{x}_n \end{pmatrix} = \begin{pmatrix} \mathbf{x}_1^* \\ \mathbf{x}_2^* \\ \vdots \\ \mathbf{x}_n^* \end{pmatrix} \begin{pmatrix} \mathbb{A}\mathbf{x}_1 & \mathbb{A}\mathbf{x}_2 & \dots & \mathbb{A}\mathbf{x}_n \end{pmatrix}$$

Ce qui donne

$$\mathbb{B} = \begin{pmatrix} \mathbf{x}_1^* \mathbb{A}\mathbf{x}_1 & \mathbf{x}_1^* \mathbb{A}\mathbf{x}_2 & \dots & \mathbf{x}_1^* \mathbb{A}\mathbf{x}_n \\ \mathbf{x}_2^* \mathbb{A}\mathbf{x}_1 & \mathbf{x}_2^* \mathbb{A}\mathbf{x}_2 & \dots & \mathbf{x}_2^* \mathbb{A}\mathbf{x}_n \\ \vdots & \vdots & \ddots & \vdots \\ \mathbf{x}_n^* \mathbb{A}\mathbf{x}_1 & \mathbf{x}_n^* \mathbb{A}\mathbf{x}_2 & \dots & \mathbf{x}_n^* \mathbb{A}\mathbf{x}_n \end{pmatrix}$$

On a donc

$$B_{i,j} = \mathbf{x}_i^* \mathbb{A}\mathbf{x}_j, \quad \forall (i,j) \in \llbracket 1, n \rrbracket^2$$

2. On a $\mathbb{A}\mathbf{u} = \lambda\mathbf{u}$, $\|\mathbf{u}\| = 1$, la base $\{\mathbf{x}_1, \dots, \mathbf{x}_n\}$ est orthonormée et $\mathbf{x}_1 = \mathbf{u}$. on obtient alors

$$\mathbb{B} = \begin{pmatrix} \lambda \mathbf{u}^* \mathbf{u} & \mathbf{u}^* \mathbb{A}\mathbf{x}_2 & \dots & \mathbf{u}^* \mathbb{A}\mathbf{x}_n \\ \lambda \mathbf{x}_2^* \mathbf{u} & \mathbf{x}_2^* \mathbb{A}\mathbf{x}_2 & \dots & \mathbf{x}_2^* \mathbb{A}\mathbf{x}_n \\ \vdots & \vdots & \ddots & \vdots \\ \lambda \mathbf{x}_n^* \mathbf{u} & \mathbf{x}_n^* \mathbb{A}\mathbf{x}_2 & \dots & \mathbf{x}_n^* \mathbb{A}\mathbf{x}_n \end{pmatrix} = \begin{pmatrix} \lambda & \mathbf{u}^* \mathbb{A}\mathbf{x}_2 & \dots & \mathbf{u}^* \mathbb{A}\mathbf{x}_n \\ 0 & \mathbf{x}_2^* \mathbb{A}\mathbf{x}_2 & \dots & \mathbf{x}_2^* \mathbb{A}\mathbf{x}_n \\ \vdots & \vdots & \ddots & \vdots \\ 0 & \mathbf{x}_n^* \mathbb{A}\mathbf{x}_2 & \dots & \mathbf{x}_n^* \mathbb{A}\mathbf{x}_n \end{pmatrix}$$

Q. 3 On veut démontrer, par récurrence faible, la proposition suivante pour $n \geq 2$

$(\mathcal{P}_n)$ $\forall \mathbb{A} \in \mathcal{M}_n(\mathbb{C})$, $\exists \mathbb{U} \in \mathcal{M}_n(\mathbb{C})$ unitaire, $\exists \mathbb{T} \in \mathcal{M}_n(\mathbb{C})$ triangulaire supérieure, telles que $\mathbb{A} = \mathbb{U}\mathbb{T}\mathbb{U}^*$.

Initialisation : Montrons que $(\mathcal{P}_2)$ est vérifié.

Soit $\mathbb{A}_2 \in \mathcal{M}_2(\mathbb{C})$. Elle admet au moins un élément propre $(\lambda, \mathbf{u})$ (voir Proposition A.37 par ex.) avec $\|\mathbf{u}\| = 1$. On peut donc appliquer le résultat de la question précédente : il existe une matrice unitaire $\mathbb{P}_2 \in \mathcal{M}_2(\mathbb{C})$ telle que la matrice $\mathbb{B}_2 = \mathbb{P}_2 \mathbb{A}_2 \mathbb{P}_2^*$ ait comme premier vecteur colonne $(\lambda, 0)^t$. La matrice $\mathbb{B}_2$ est donc triangulaire supérieure et comme $\mathbb{P}_2$ est unitaire on en déduit

$$\mathbb{A}_2 = \mathbb{P}_2^* \mathbb{B}_2 \mathbb{P}_2.$$

On pose $\mathbb{U}_2 = \mathbb{P}_2^*$ matrice unitaire et $\mathbb{T}_2 = \mathbb{B}_2$ matrice triangulaire supérieure pour conclure que la proposition $(\mathcal{P}_2)$ est vraie.

Hérédité : Supposons que $(\mathcal{P}_{n-1})$ soit vérifiée. Montrons que $(\mathcal{P}_n)$ est vraie.

Soit $\mathbb{A}_n \in \mathcal{M}_n(\mathbb{C})$. Elle admet au moins un élément propre $(\lambda, \mathbf{u})$ (voir Proposition A.37 par ex.) avec $\|\mathbf{u}\| = 1$. On peut donc appliquer le résultat de la question précédente : il existe une matrice unitaire $\mathbb{P}_n \in \mathcal{M}_n(\mathbb{C})$ telle que la matrice $\mathbb{B}_n = \mathbb{P}_n \mathbb{A}_n \mathbb{P}_n^*$ s'écrive

$$\mathbb{B}_n = \begin{pmatrix} \lambda & \mathbf{c}_{n-1}^* \\ 0 & \mathbb{A}_{n-1} \\ \vdots & \\ 0 & \end{pmatrix}$$

où $\mathbf{c}_{n-1} \in \mathcal{M}_{n-1}(\mathbb{C})$ et $\mathbb{A}_{n-1} \in \mathcal{M}_{n-1}(\mathbb{C})$. Par hypothèse de récurrence, $\exists \mathbb{U}_{n-1} \in \mathcal{M}_{n-1}(\mathbb{C})$ unitaire et $\mathbb{T}_{n-1} \in \mathcal{M}_{n-1}(\mathbb{C})$ triangulaire supérieure telles que

$$\mathbb{A}_{n-1} = \mathbb{U}_{n-1} \mathbb{T}_{n-1} \mathbb{U}_{n-1}^*$$

ou encore

$$\mathbb{T}_{n-1} = \mathbb{U}_{n-1}^* \mathbb{A}_{n-1} \mathbb{U}_{n-1}.$$

1 Soit $Q_n \in \mathcal{M}_n(\mathbb{C})$ la matrice définie par

$$Q_n = \left(\begin{array}{c|ccc} 1 & 0 & \dots & 0 \\ \hline 0 & & & \\ \vdots & & \mathbb{U}_{n-1} & \\ 0 & & & \end{array} \right).$$

2 La matrice Q_n est unitaire. En effet on a

$$Q_n Q_n^* = \left(\begin{array}{c|ccc} 1 & 0 & \dots & 0 \\ \hline 0 & & & \\ \vdots & & \mathbb{U}_{n-1} & \\ 0 & & & \end{array} \right) \left(\begin{array}{c|ccc} 1 & 0 & \dots & 0 \\ \hline 0 & & & \\ \vdots & & \mathbb{U}_{n-1}^* & \\ 0 & & & \end{array} \right) = \left(\begin{array}{c|ccc} 1 & 0 & \dots & 0 \\ \hline 0 & & & \\ \vdots & & \underbrace{\mathbb{U}_{n-1} \mathbb{U}_{n-1}^*}_{=\mathbb{I}_{n-1}} & \\ 0 & & & \end{array} \right) = \mathbb{I}_n.$$

On note T_n la matrice définie par $T_n = Q_n^* B_n Q_n$. On a alors

$$\begin{aligned} T_n &= \left(\begin{array}{c|ccc} 1 & 0 & \dots & 0 \\ \hline 0 & & & \\ \vdots & & \mathbb{U}_{n-1}^* & \\ 0 & & & \end{array} \right) \left(\begin{array}{c|ccc} \lambda & & \mathbf{c}_{n-1}^* & \\ \hline 0 & & & \\ \vdots & & \mathbb{A}_{n-1} & \\ 0 & & & \end{array} \right) \left(\begin{array}{c|ccc} 1 & 0 & \dots & 0 \\ \hline 0 & & & \\ \vdots & & \mathbb{U}_{n-1} & \\ 0 & & & \end{array} \right) \\ &= \left(\begin{array}{c|ccc} \lambda & & \mathbf{c}_{n-1}^* & \\ \hline 0 & & & \\ \vdots & & \mathbb{U}_{n-1}^* \mathbb{A}_{n-1} & \\ 0 & & & \end{array} \right) \left(\begin{array}{c|ccc} 1 & 0 & \dots & 0 \\ \hline 0 & & & \\ \vdots & & \mathbb{U}_{n-1} & \\ 0 & & & \end{array} \right) = \left(\begin{array}{c|ccc} \lambda & & \mathbf{c}_{n-1}^* \mathbb{U}_{n-1}^* & \\ \hline 0 & & & \\ \vdots & & \underbrace{\mathbb{U}_{n-1}^* \mathbb{A}_{n-1} \mathbb{U}_{n-1}}_{=T_{n-1}} & \\ 0 & & & \end{array} \right) \end{aligned}$$

3 La matrice T_n est donc triangulaire supérieure et on a par définition de B_n

$$T_n = Q_n^* P_n A_n P_n^* Q_n.$$

4 On note $U_n = P_n^* Q_n$. Cette matrice est unitaire car les matrices Q_n et P_n le sont. En effet, on a

$$U_n U_n^* = P_n^* Q_n (P_n^* Q_n)^* = P_n^* \underbrace{Q_n Q_n^*}_{=\mathbb{I}_n} P_n = P_n^* P_n = \mathbb{I}_n.$$

5 On a $T_n = U_n^* A_n U_n$ et en multipliant cette équation à gauche par U_n et à droite par U_n^* on obtient l'équation équivalente $A_n = U_n T_n U_n^*$. La propriété (P_n) est donc vérifiée. Ce qui achève la démonstration.

8 **Q. 4** Résoudre $Ax = b$ est équivalent à résoudre

$$UTU^*x = b. \quad (4.1)$$

Comme U est unitaire, on a $UU^* = \mathbb{I}$ et U^* inversible. Donc en multipliant (A.52) par U^* on obtient le système équivalent

$$\underbrace{U^*U}_{=\mathbb{I}} TU^*x = U^*b \iff TU^*x = U^*b.$$

9 On pose $y = U^*x$. Le système précédent se résout en deux étapes

1. on cherche y solution de $Ty = U^*b$. Comme U est unitaire on a $\det(U) \det(U^*) = \det(\mathbb{I}) = 1$ et donc

$$\begin{aligned} \det(A) &= \det(UTU^*) = \det(U) \det(T) \det(U^*) \\ &= \det(T) \end{aligned}$$

10 Or A inversible équivaut à $\det(A) \neq 0$ et donc la matrice T est inversible. La matrice T étant
11 triangulaire inférieure on peut résoudre facilement le système par la *méthode de remontée*.

12 2. une fois y déterminé, on résout $U^*x = y$. Comme U est unitaire, on obtient directement $x = Uy$.

13 ◇

14 On tire de cet exercice le théorème suivant

**Théorème 4.1:**

Soit $A \in \mathcal{M}_n(\mathbb{C})$. Il existe une matrice unitaire U et une matrice triangulaire supérieure T telles que

$$A = UTU^* \quad (4.2)$$

1

Proof. voir Exercice 4.1.1



2

**Théorème 4.2: Réduction de matrices**

1. Soit $A \in \mathcal{M}_n(\mathbb{C})$. Il existe une matrice **unitaire** U telle que $U^{-1}AU$ soit **triangulaire**.
2. Soit $A \in \mathcal{M}_n(\mathbb{C})$ une matrice **normale**. Il existe une matrice **unitaire** U telle que $U^{-1}AU$ soit **diagonale**.
3. Soit $A \in \mathcal{M}_n(\mathbb{R})$ une matrice **symétrique**. Il existe une matrice **orthogonale** P telle que $P^{-1}AP$ soit **diagonale**.

3

Proof. voir [1] Théorème 1.2-1 page 9



4

**Exercice 4.1.2: Matrice d'élimination**

5

Soit $\mathbf{v} \in \mathbb{C}^n$ avec $v_1 \neq 0$. On note $\mathbb{E}^{[\mathbf{v}]} \in \mathcal{M}_n(\mathbb{C})$ la matrice triangulaire inférieure à diagonale unité définie par

$$\mathbb{E}^{[\mathbf{v}]} = \left(\begin{array}{c|ccc} 1 & 0 & \dots & 0 \\ -v_2/v_1 & 1 & 0 & \dots \\ \vdots & 0 & \ddots & \vdots \\ \vdots & 0 & \ddots & 0 \\ -v_n/v_1 & 0 & \dots & 1 \end{array} \right) \quad (4.3)$$

Q. 1 1. Calculer le déterminant de $\mathbb{E}^{[\mathbf{v}]}$.

2. Déterminer l'inverse de $\mathbb{E}^{[\mathbf{v}]}$.

$\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$ avec $A_{1,1} \neq 0$. On note $\mathbf{A}_{:,j}$ le j ème vecteur colonne de $\mathbb{A}$ et $\mathbf{A}_{i,:}$ son i ème vecteur ligne. On pose $\mathbf{A}_1 = \mathbf{A}_{:,1}$.

Q. 2 1. Calculer $\tilde{\mathbb{A}} = \mathbb{E}^{[\mathbf{A}_1]}\mathbb{A}$ en fonction des vecteurs lignes de $\mathbb{A}$.

2. Montrer que la première colonne de $\tilde{\mathbb{A}}$ est le vecteur $(A_{1,1}, 0, \dots, 0)^t$ i.e.

$$\mathbb{E}^{[\mathbf{A}_1]}\mathbb{A}\mathbf{e}_1 = A_{1,1}\mathbf{e}_1 \quad (4.4)$$

où $\mathbf{e}_1$ est le premier vecteur de la base canonique de $\mathbb{C}^n$.

Soit $m \in \mathbb{N}^*$. On note $\mathbb{E}^{[m,\mathbf{v}]} \in \mathcal{M}_{m+n}(\mathbb{C})$ la matrice triangulaire inférieure à diagonale unité définie par

$$\mathbb{E}^{[m,\mathbf{v}]} = \left(\begin{array}{c|c} \mathbb{I}_m & \mathbf{0} \\ \hline \mathbf{0} & \mathbb{E}^{[\mathbf{v}]} \end{array} \right) \quad (4.5)$$

Q. 3 1. Calculer le déterminant de $\mathbb{E}^{[m,\mathbf{v}]}$.

2. Déterminer l'inverse de $\mathbb{E}^{[m,\mathbf{v}]}$ en fonction de l'inverse de $\mathbb{E}^{[\mathbf{v}]}$.

Soit $\mathbb{C}$ la matrice bloc définie par

$$\mathbb{C} = \left(\begin{array}{c|c} \mathbb{C}_{1,1} & \mathbb{C}_{1,2} \\ \hline \mathbf{0} & \mathbb{A} \end{array} \right)$$

où $\mathbb{C}_{1,1} \in \mathcal{M}_m(\mathbb{C})$ et $\mathbb{C}_{1,2} \in \mathcal{M}_{m,n}(\mathbb{C})$.

Q. 4 Déterminer la matrice produit $\mathbb{E}^{[m,\mathbf{A}_1]}\mathbb{B}$ en fonction des matrices $\mathbb{C}_{1,1}$, $\mathbb{C}_{1,2}$ et $\tilde{\mathbb{A}}$.

1

2 Correction Exercice 4.1.2

Q. 1 1. La matrice $\mathbb{E}^{[\mathbf{v}]}$ est triangulaire : son déterminant est donc le produit de ses éléments diagonaux (voir Proposition A.44, page 184). On a alors $\det(\mathbb{E}^{[\mathbf{v}]}) = 1$.

2. Pour calculer son inverse qui existe puisque $\det(\mathbb{E}^{[\mathbf{v}]}) \neq 0$, on écrit $\mathbb{E}^{[\mathbf{v}]}$ sous forme bloc :

$$\mathbb{E}^{[\mathbf{v}]} = \left(\begin{array}{c|ccc} 1 & 0 & \dots & 0 \\ \hline \mathbf{e} & \mathbb{I}_{n-1} & & \end{array} \right)$$

avec $\mathbf{e} = (-v_2/v_1, \dots, -v_n/v_1)^t \in \mathbb{C}^{n-1}$. On note $\mathbb{X} \in \mathcal{M}_n(\mathbb{C})$ son inverse qui s'écrit avec la même structure bloc

$$\mathbb{X} = \left(\begin{array}{c|c} a & \mathbf{b}^* \\ \hline \mathbf{c} & \mathbb{D} \end{array} \right)$$

avec $a \in \mathbb{K}$, $\mathbf{b} \in \mathbb{K}^{n-1}$, $\mathbf{c} \in \mathbb{K}^{n-1}$ et $\mathbb{D} \in \mathcal{M}_{n-1}(\mathbb{C})$.

La matrice $\mathbb{X}$ est donc solution de $\mathbb{E}^{[\mathbf{v}]} \mathbb{X} = \mathbb{I}$. Grace à l'écriture bloc des matrices on en déduit

rapidement la matrice $\mathbb{X}$. En effet, en utilisant les produits blocs des matrices, on obtient

$$\begin{aligned} \mathbb{E}^{[v]}\mathbb{X} &= \begin{pmatrix} 1 & \mathbf{0}_{n-1}^t \\ \mathbf{e} & \mathbb{I}_{n-1} \end{pmatrix} \begin{pmatrix} a & \mathbf{b}^* \\ \mathbf{c} & \mathbb{D} \end{pmatrix} = \begin{pmatrix} 1 \times a & 1 \times \mathbf{b}^* + \mathbf{0}_{n-1}^t \times \mathbb{D} \\ \mathbf{e} \times a + \mathbb{I}_{n-1} \times \mathbf{c} & \mathbf{e} \times \mathbf{b}^* + \mathbb{I}_{n-1} \times \mathbb{D} \end{pmatrix} \\ &= \begin{pmatrix} a & \mathbf{b}^* \\ a\mathbf{e} + \mathbf{c} & \mathbf{e}\mathbf{b}^* + \mathbb{D} \end{pmatrix} \end{aligned}$$

Comme $\mathbb{X}$ est l'inverse de $\mathbb{E}^{[v]}$, on a $\mathbb{E}^{[v]}\mathbb{X} = \mathbb{I}$ et donc en écriture bloc

$$\begin{pmatrix} a & \mathbf{b}^* \\ a\mathbf{e} + \mathbf{c} & \mathbf{e}\mathbf{b}^* + \mathbb{D} \end{pmatrix} = \begin{pmatrix} 1 & \mathbf{0}_{n-1}^t \\ \mathbf{0}_{n-1} & \mathbb{I}_{n-1} \end{pmatrix}.$$

Ceci revient à résoudre les 4 équations

$$a = 1, \quad \mathbf{b}^* = \mathbf{0}_{n-1}^t, \quad a\mathbf{e} + \mathbf{c} = \mathbf{0}_{n-1} \quad \text{et} \quad \mathbf{e}\mathbf{b}^* + \mathbb{D} = \mathbb{I}_{n-1}$$

qui donnent immédiatement $a = 1$, $\mathbf{b} = \mathbf{0}_{n-1}$, $\mathbf{c} = -\mathbf{e}$ et $\mathbb{D} = \mathbb{I}_{n-1}$. On obtient le résultat suivant

$$\begin{pmatrix} 1 & 0 & \dots & 0 \\ -\mathbf{e} & \mathbb{I}_{n-1} \end{pmatrix} \begin{pmatrix} 1 & 0 & \dots & 0 \\ \mathbf{e} & \mathbb{I}_{n-1} \end{pmatrix} = \mathbb{I}_n.$$



Il aurait été plus rapide d'utiliser la Proposition A.45, page 184.

Q. 2 1. Pour simplifier les notations, on note $\mathbb{E} = \mathbb{E}^{[A_1]}$. Par définition du produit de deux matrices on a

$$\tilde{A}_{i,j} = \sum_{k=1}^n E_{i,k} A_{k,j}, \quad \forall (i,j) \in \llbracket 1, n \rrbracket^2.$$

Quand $i = 1$, on a par construction $E_{1,k} = \delta_{1,k}$ et donc

$$\tilde{A}_{1,j} = A_{1,j}, \quad \forall j \in \llbracket 1, n \rrbracket \iff \tilde{\mathbf{A}}_{1,:} = \mathbf{A}_{1,:}. \quad (4.6)$$

Pour $i \geq 2$, on a $E_{i,1} = -\frac{v_i}{v_1}$ et $E_{i,k} = \delta_{i,k}$, $\forall k \in \llbracket 2, n \rrbracket$. On obtient alors pour tout $j \in \llbracket 1, n \rrbracket$

$$\tilde{A}_{i,j} = E_{i,1} A_{1,j} + \sum_{k=2}^n E_{i,k} A_{k,j} = -\frac{v_i}{v_1} A_{1,j} + \sum_{k=2}^n \delta_{i,k} A_{k,j} = -\frac{v_i}{v_1} A_{1,j} + A_{i,j}$$

ce qui donne pour tout $i \in \llbracket 2, n \rrbracket$

$$\tilde{A}_{i,j} = A_{i,j} - \frac{v_i}{v_1} A_{1,j}, \quad \forall j \in \llbracket 1, n \rrbracket \iff \tilde{\mathbf{A}}_{i,:} = -\frac{v_i}{v_1} \mathbf{A}_{1,:} + \mathbf{A}_{i,:} \quad (4.7)$$

En conclusion, la matrice $\tilde{\mathbf{A}}$ s'écrit

$$\tilde{\mathbf{A}} = \begin{pmatrix} \mathbf{A}_{1,:} \\ \mathbf{A}_{2,:} - (v_2/v_1)\mathbf{A}_{1,:} \\ \vdots \\ \mathbf{A}_{n,:} - (v_n/v_1)\mathbf{A}_{1,:} \end{pmatrix}$$

- 1 2. De (4.6), on tire $\tilde{A}_{1,1} = A_{1,1}$. A partir de (4.7) on obtient pour tout $i \in \llbracket 2, n \rrbracket$, $\tilde{A}_{i,1} = A_{i,1} - \frac{v_i}{v_1} A_{1,1}$.
 2 Par construction $v_j = A_{j,1}$ pour tout $j \in \llbracket 1, n \rrbracket$, ce qui donne $\tilde{A}_{i,1} = 0$.
 3 La première colonne de $\tilde{A}$ est $(1, 0, \dots, 0)^t$.

4 **Q. 3** 1. La matrice $\mathbb{E}^{[m, \mathbf{v}]}$ est triangulaire inférieure. Son déterminant est donc le produit de ses
 5 éléments diagonaux. Comme cette matrice est à diagonale unité (i.e. tous ses éléments diagonaux
 6 valent 1), on obtient $\det \mathbb{E}^{[m, \mathbf{v}]} = 1$.

7 Une autre manière de le démontrer. On peut voir que la matrice $\mathbb{E}^{[m, \mathbf{v}]}$ est bloc-diagonale. D'après
 8 la Proposition A.45, page 184, son déterminant est le produit des déterminant des blocs diagonaux
 9 : $\det \mathbb{E}^{[m, \mathbf{v}]} = \det \mathbb{I}_m \times \det \mathbb{E}^{[\mathbf{v}]} = 1$.

- 10 2. On note $\mathbb{X}$ l'inverse de la matrice $\mathbb{E}^{[m, \mathbf{v}]}$. Cette matrice s'écrit avec la même structure bloc

$$\mathbb{X} \left(\begin{array}{c|c} \mathbb{X}_{1,1} & \mathbb{X}_{1,2} \\ \hline \mathbb{X}_{2,1} & \mathbb{X}_{2,2} \end{array} \right) \text{ avec } \mathbb{X}_{1,1} \in \mathcal{M}_m(\mathbb{C}) \text{ et } \mathbb{X}_{2,2} \in \mathcal{M}_n(\mathbb{C})$$

11 On a donc $\mathbb{X} \mathbb{E}^{[m, \mathbf{v}]} = \mathbb{I}_{m+n}$ c'est à dire en écriture bloc

$$\left(\begin{array}{c|c} \mathbb{X}_{1,1} & \mathbb{X}_{1,2} \\ \hline \mathbb{X}_{2,1} & \mathbb{X}_{2,2} \end{array} \right) \left(\begin{array}{c|c} \mathbb{I}_m & 0 \\ \hline 0 & \mathbb{E}^{[\mathbf{v}]} \end{array} \right) = \left(\begin{array}{c|c} \mathbb{I}_m & 0 \\ \hline 0 & \mathbb{I}_n \end{array} \right) =$$

12 On doit donc résoudre les 4 équations suivantes :

$$\mathbb{X}_{1,1} \mathbb{I}_m = \mathbb{I}_m, \quad \mathbb{X}_{1,2} \mathbb{I}_n = 0, \quad \mathbb{X}_{2,1} \mathbb{I}_m = 0 \quad \text{et} \quad \mathbb{X}_{2,2} \mathbb{E}^{[\mathbf{v}]} = \mathbb{I}_n.$$

13 Comme la matrice $\mathbb{E}^{[\mathbf{v}]}$ est inversible, on obtient

$$\mathbb{X} = \left(\begin{array}{c|c} \mathbb{I}_m & 0 \\ \hline 0 & (\mathbb{E}^{[\mathbf{v}]})^{-1} \end{array} \right)$$

14  Plus rapidement, comme la matrice $\mathbb{E}^{[m, \mathbf{v}]}$ est bloc-diagonale, on en déduit (voir Proposition A.45,
 15 page 184) directement le résultat.

Q. 4 Le produit BC peut s'effectuer par bloc car les blocs sont de dimensions compatibles et on a

$$\begin{aligned} \text{BC} &= \left(\begin{array}{c|c} \mathbb{I}_m & 0_{m,n} \\ \hline 0_{n,m} & \mathbb{E} \end{array} \right) \left(\begin{array}{c|c} \mathbb{C}_{1,1} & \mathbb{C}_{1,2} \\ \hline 0_{n,m} & \mathbb{A} \end{array} \right) = \left(\begin{array}{c|c} \mathbb{I}_m \mathbb{C}_{1,1} + 0_{m,n} 0_{n,m} & \mathbb{I}_m \mathbb{C}_{1,2} + 0_{m,n} \mathbb{A} \\ \hline 0_{n,m} \mathbb{C}_{1,1} + \mathbb{E} 0_{n,m} & 0_{n,m} \mathbb{C}_{1,2} + \mathbb{E} \mathbb{A} \end{array} \right) \\ &= \left(\begin{array}{c|c} \mathbb{C}_{1,1} & \mathbb{C}_{1,2} \\ \hline 0_{n,m} & \mathbb{E} \mathbb{A} \end{array} \right) = \left(\begin{array}{c|c} \mathbb{C}_{1,1} & \mathbb{C}_{1,2} \\ \hline 0_{n,m} & \mathbb{A} \end{array} \right) \end{aligned}$$

◇

17 On tire de cet exercice le lemme suivant



Lemme 4.3

Soit $\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$ avec $A_{1,1} \neq 0$. Il existe une matrice $\mathbb{E} \in \mathcal{M}_n(\mathbb{C})$ triangulaire inférieure à diagonale
 unité telle que

$$\mathbb{E} \mathbb{A} \mathbf{e}_1 = A_{1,1} \mathbf{e}_1 \quad (4.8)$$

19 où $\mathbf{e}_1$ est le premier vecteur de la base canonique de $\mathbb{C}^n$.



Exercice 4.1.3: Matrice de permutation

Soit $(i, j) \in \llbracket 1, n \rrbracket^2$, on note $\mathbb{P}_n^{[i,j]} \in \mathcal{M}_n(\mathbb{R})$ la matrice identité dont on a permuté les lignes i et j .

Q. 1 Définir proprement cette matrice et la représenter.

Soient $\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$ et $\mathbb{B} \in \mathcal{M}_n(\mathbb{C})$. On note $\mathbf{A}_{r,:}$ le r ème vecteur ligne de $\mathbb{A}$ et $\mathbf{B}_{:,s}$ le s ème vecteur colonne de $\mathbb{B}$.

Q. 2 1. Déterminer $\mathbb{P}_n^{[i,j]} \mathbb{A}$ en fonction des vecteurs lignes de $\mathbb{A}$.

2. Déterminer $\mathbb{B} \mathbb{P}_n^{[i,j]}$ en fonction des vecteurs colonnes de $\mathbb{B}$.

Q. 3 1. Calculer le déterminant de $\mathbb{P}_n^{[i,j]}$.

2. Déterminer l'inverse de $\mathbb{P}_n^{[i,j]}$.

Correction Exercice 4.1.3 On note $\mathbb{P} = \mathbb{P}_n^{[i,j]}$.

Q. 1 On peut définir cette matrice par ligne, $\forall s \in \llbracket 1, n \rrbracket$,

$$\begin{cases} P_{r,s} &= \delta_{r,s}, \quad \forall r \in \llbracket 1, n \rrbracket \setminus \{i, j\}, \\ P_{i,s} &= \delta_{j,s}, \\ P_{j,s} &= \delta_{i,s}. \end{cases}$$

On peut noter que la matrice $\mathbb{P}$ est symétrique.

Q. 2 1. On note $\mathbb{D} = \mathbb{P} \mathbb{A}$. Par définition du produit matriciel on a

$$D_{r,s} = \sum_{k=1}^n P_{r,k} A_{k,s}.$$



Ne pas utiliser les indices i et j qui sont déjà fixés dans la définition de la matrice $\mathbb{P} = \mathbb{P}_n^{[i,j]}$.

On obtient, $\forall s \in \llbracket 1, n \rrbracket$,

$$\begin{cases} D_{r,s} &= \sum_{k=1}^n \delta_{r,k} A_{k,s} = A_{r,s}, \quad \forall r \in \llbracket 1, n \rrbracket \setminus \{i, j\}, \\ D_{i,s} &= \sum_{k=1}^n \delta_{j,k} A_{k,s} = A_{j,s}, \\ D_{j,s} &= \sum_{k=1}^n \delta_{i,k} A_{k,s} = A_{i,s}. \end{cases}$$

ce qui donne

$$\begin{cases} \mathbf{D}_{r,:} &= \mathbf{A}_{r,:}, \quad \forall r \in \llbracket 1, n \rrbracket \setminus \{i, j\}, \\ \mathbf{D}_{i,:} &= \mathbf{A}_{j,:}, \\ \mathbf{D}_{j,:} &= \mathbf{A}_{i,:}. \end{cases}$$

2. On note $\mathbb{E} = \mathbb{A} \mathbb{P}$. Par définition du produit matriciel et par symétrie de $\mathbb{P}$ on a

$$E_{r,s} = \sum_{k=1}^n A_{r,k} P_{k,s} = \sum_{k=1}^n A_{r,k} P_{s,k}.$$

On obtient en raisonnant par colonne, $\forall r \in \llbracket 1, n \rrbracket$,

$$\begin{cases} E_{r,s} &= \sum_{k=1}^n A_{r,k} \delta_{s,k} = A_{r,s}, \quad \forall s \in \llbracket 1, n \rrbracket \setminus \{i, j\}, \\ E_{r,i} &= \sum_{k=1}^n A_{r,k} \delta_{j,k} = A_{r,j}, \\ E_{r,j} &= \sum_{k=1}^n A_{r,k} \delta_{i,k} = A_{r,i}. \end{cases}$$

ce qui donne

$$\begin{cases} E_{:,s} = A_{:,s}, & \forall s \in \llbracket 1, n \rrbracket \setminus \{i, j\}, \\ E_{:,i} = A_{:,j}, \\ E_{:,j} = A_{:,i}. \end{cases}$$

Q. 3 1. $\det(\mathbb{P}) = -1$, si $i \neq j$ et $\det(\mathbb{P}) = 1$ sinon.

2. Immédiat par calcul direct on a $\mathbb{P}\mathbb{P} = \mathbb{I}$ et donc la matrice $\mathbb{P}$ est inversible et $\mathbb{P}^{-1} = \mathbb{P}$.

On tire de cet exercice le lemme suivant



Lemme 4.4

Soit $(i, j) \in \llbracket 1, n \rrbracket^2$. On note $\mathbb{P}_n^{[i,j]} \in \mathcal{M}_n(\mathbb{R})$ la matrice identité dont on a permuté les lignes i et j . Alors la matrice $\mathbb{P}_n^{[i,j]}$ est **symétrique et orthogonale**. Pour toute matrice $A \in \mathcal{M}_n(\mathbb{K})$,

1. la matrice $\mathbb{P}_n^{[i,j]}A$ est la matrice A dont on a permuté les **lignes** i et j ,
2. la matrice $A\mathbb{P}_n^{[i,j]}$ est la matrice A dont on a permuté les **colonnes** i et j ,

4.2 Méthodes directes

Pour résoudre le système linéaire $Ax = b$, nous allons le transformer en un système linéaire triangulaire supérieure équivalent

$$MAx = Mb$$

où M est une matrice inversible telle que MA soit triangulaire supérieure. Nous allons voir que ce nouveau système est très facile à résoudre par la *méthode de la remontée*.

Nous étudierons la *méthode de Gauss-Jordan* que nous réécrirons sous forme algébrique. Puis nous ferons le lien avec les méthodes utilisant la *factorisation LU* et la *factorisation de Cholesky*. Nous finirons par une méthode utilisant la factorisation QR d'une matrice, factorisation qui sera réutilisée pour le calcul de valeurs propres et vecteurs propres en section ??.

Nous allons tout d'abord regarder quelques cas particuliers : la matrice du système est diagonale, triangulaire inférieure ou triangulaire supérieure.

4.2.1 Matrices particulières

Matrices diagonales

Soit A une matrice de $\mathcal{M}_n(\mathbb{K})$ diagonale inversible et $b \in \mathbb{K}^n$. Dans ce cas les coefficients diagonaux de A sont tous non nuls et l'on a

$$x_i = b_i / A_{i,i}, \quad \forall i \in \llbracket 1, n \rrbracket. \quad (4.9)$$

On a immédiatement l'algorithme

Algorithme 4.1 Fonction **RSLMATDIAG** permettant de résoudre le système linéaire à matrice diagonale inversible

$$Ax = b.$$

Données : A : matrice diagonale de $\mathcal{M}_n(\mathbb{R})$ inversible.
 b : vecteur de $\mathbb{R}^n$.

Résultat : x : vecteur de $\mathbb{R}^n$.

```

1: Fonction  $x \leftarrow \text{RSLMATDIAG} (A, b)$ 
2:   Pour  $i \leftarrow 1$  à  $n$  faire
3:      $x(i) \leftarrow b(i) / A(i, i)$ 
4:   Fin Pour
5: Fin Fonction

```

Matrices triangulaires inférieures

Soit $\mathbb{A}$ une matrice de $\mathcal{M}_n(\mathbb{K})$ triangulaire inférieure inversible et $\mathbf{b} \in \mathbb{K}^n$. On veut résoudre le système linéaire

$$\mathbb{A}\mathbf{x} = \mathbf{b} \iff \begin{pmatrix} A_{1,1} & 0 & \dots & 0 \\ \vdots & \ddots & \ddots & \vdots \\ \vdots & & \ddots & 0 \\ A_{n,1} & \dots & \dots & A_{n,n} \end{pmatrix} \begin{pmatrix} x_1 \\ \vdots \\ \vdots \\ x_n \end{pmatrix} = \begin{pmatrix} b_1 \\ \vdots \\ \vdots \\ b_n \end{pmatrix}$$



Exercice 4.2.1

Soit $\mathbb{A} \in \mathcal{M}_n(\mathbb{K})$ une matrice **triangulaire**. Montrer que

$$\mathbb{A} \text{ inversible} \iff A_{i,i} \neq 0, \forall i \in \llbracket 1, n \rrbracket.$$

On remarque que l'on peut calculer successivement $x_1, x_1, \dots, x_n$, car il est possible de calculer x_i si on connaît $x_1, \dots, x_{i-1}$: c'est la **méthode de descente**. En effet, on a

$$(\mathbb{A}\mathbf{x})_i = b_i, \forall i \in \llbracket 1, n \rrbracket.$$

et donc, par définition d'un produit matrice-vecteur,

$$\sum_{j=1}^n A_{i,j}x_j = b_i, \forall i \in \llbracket 1, n \rrbracket.$$

Comme $\mathbb{A}$ est une matrice triangulaire inférieure, on a (voir Définition A.35) $A_{i,j} = 0$ si $j > i$. Ceci donne alors pour tout $i \in \llbracket 1, n \rrbracket$

$$\begin{aligned} b_i &= \sum_{j=1}^{i-1} A_{i,j}x_j + A_{i,i}x_i + \sum_{j=i+1}^n \underbrace{A_{i,j}}_{=0} x_j \\ &= \sum_{j=1}^{i-1} A_{i,j}x_j + A_{i,i}x_i \end{aligned}$$

De plus la matrice $\mathbb{A}$ étant triangulaire inversible ses éléments diagonaux sont tous non nuls. On obtient alors x_i en fonction des $x_{i+1}, \dots, x_n$:

$$x_i = \frac{1}{A_{i,i}} \left(b_i - \sum_{j=1}^{i-1} A_{i,j}x_j \right), \forall i \in \llbracket 1, n \rrbracket. \quad (4.10)$$

On écrit en détail les raffinements successifs permettant d'aboutir à l'Algorithme 4.2 final ne comportant que des opérations élémentaires (... à finaliser) de telle sorte que le passage entre deux raffinements successifs soit le plus compréhensible possible.

Algorithme 4.2 $\mathcal{R}_0$

1: Résoudre $\mathbb{A}\mathbf{x} = \mathbf{b}$ en calculant successivement $x_1, x_2, \dots, x_n$.

Algorithme 4.2 $\mathcal{R}_1$

1: **Pour** $i \leftarrow 1$ à n **faire**
 2: $x_i \leftarrow \frac{1}{A_{i,i}} \left(b_i - \sum_{j=1}^{i-1} A_{i,j}x_j \right)$
 3: **Fin Pour**

Dans le raffinement $\mathcal{R}_1$, la seule difficulté restante est le calcul de la somme $\sum_{j=1}^{i-1} A_{i,j}x_j$. En effet, l'opérateur mathématique $\sum$ n'est pas défini dans notre langage algorithmique : il va donc falloir détailler

1 un peu plus l'algorithme. Pour isoler le calcul de cette somme, on la note S . La ligne 2 peut donc s'écrire

$$x_i \leftarrow \frac{1}{A_{i,i}} (b_i - S).$$

2 Mais où calculer la valeur S ? 4 choix possible : avant la ligne 1, entre les lignes 1 et 2, entre les lignes 2
3 et 3 ou après la ligne 3.

4 On ne peut pas calculer S après utilisation de sa valeur dans le calcul de x_i ! ce qui élimine les 2 derniers
5 choix. Ensuite, on ne peut sortir le calcul de S de la boucle puisque la somme dépend de l'indice de boucle
6 i : ce qui élimine le premier choix. On doit donc calculer S dans la boucle et avant le calcul de x_i .

Algorithme 4.2 $\overline{\mathcal{R}_1}$

1: **Pour** $i \leftarrow 1$ à n **faire**

2: $x_i \leftarrow \frac{1}{A_{i,i}} \left(b_i - \sum_{j=1}^{i-1} A_{i,j} x_j \right)$

3: **Fin Pour**

Algorithme 4.2 $\overline{\mathcal{R}_2}$

1: **Pour** $i \leftarrow 1$ à n **faire**

2: $S \leftarrow \sum_{j=1}^{i-1} A_{i,j} x_j$

3: $x_i \leftarrow (b_i - S)/A_{i,i}$

4: **Fin Pour**

7
8
9 Maintenant que l'on a isolé la *difficulté*, il reste à détailler le calcul de S . Celui-ci se fait **intégralement**
10 en lieu et place de la ligne 2.

Algorithme 4.2 $\overline{\mathcal{R}_3}$

1: **Pour** $i \leftarrow 1$ à n **faire**

2: $S \leftarrow \sum_{j=1}^{i-1} A_{i,j} x_j$

3: $x_i \leftarrow (b_i - S)/A_{i,i}$

4: **Fin Pour**

Algorithme 4.2 $\overline{\mathcal{R}_4}$

1: **Pour** $i \leftarrow 1$ à n **faire**

2: $S \leftarrow 0$

3: **Pour** $j \leftarrow 1$ à $i-1$ **faire**

4: $S \leftarrow S + A(i,j) * x(j)$

5: **Fin Pour**

6: $x_i \leftarrow (b_i - S)/A_{i,i}$

7: **Fin Pour**

11
12
13 Insister sur $S \leftarrow 0$ à l'intérieur de la boucle en i ? (erreur trop courante des débutants)

14 On obtient alors l'algorithme final

Algorithme 4.2 Fonction **RSLTRIINF** permettant de résoudre le système linéaire triangulaire inférieur inversible

$$\mathbb{A}x = b.$$

Données : $\mathbb{A}$: matrice triangulaire de $\mathcal{M}_n(\mathbb{K})$ inférieure inversible.

b : vecteur de $\mathbb{K}^n$.

Résultat : x : vecteur de $\mathbb{K}^n$.

1: **Fonction** $x \leftarrow \text{RSLTRIINF}(\mathbb{A}, b)$

2: **Pour** $i \leftarrow 1$ à n **faire**

3: $S \leftarrow 0$

4: **Pour** $j \leftarrow 1$ à $i-1$ **faire**

5: $S \leftarrow S + A(i,j) * x(j)$

6: **Fin Pour**

7: $x(i) \leftarrow (b(i) - S)/A(i,i)$

8: **Fin Pour**

9: **Fin Fonction**

15 Matrices triangulaires supérieures

16 Soit $\mathbb{A}$ une matrice de $\mathcal{M}_n(\mathbb{R})$ triangulaire supérieure inversible et $b \in \mathbb{R}^n$. On veut résoudre le système
17 linéaire

$$\mathbb{A}\mathbf{x} = \mathbf{b} \iff \begin{pmatrix} A_{1,1} & \dots & \dots & A_{1,n} \\ 0 & \ddots & & \vdots \\ \vdots & \ddots & \ddots & \vdots \\ 0 & \dots & 0 & A_{n,n} \end{pmatrix} \begin{pmatrix} x_1 \\ \vdots \\ \vdots \\ x_n \end{pmatrix} = \begin{pmatrix} b_1 \\ \vdots \\ \vdots \\ b_n \end{pmatrix}$$

On remarque que l'on peut calculer successivement $x_n, x_{n-1}, \dots, x_1$, car il est possible de calculer x_i si on connaît $x_{i+1}, \dots, x_n$: c'est la **méthode de remontée**. En effet, on a

$$(\mathbb{A}\mathbf{x})_i = b_i, \forall i \in \llbracket 1, n \rrbracket.$$

et donc, par définition d'un produit matrice-vecteur,

$$\sum_{j=1}^n A_{i,j}x_j = b_i, \forall i \in \llbracket 1, n \rrbracket.$$

Comme $\mathbb{A}$ est une matrice triangulaire supérieure, on a (voir Définition A.35) $A_{i,j} = 0$ si $j < i$. Ceci donne alors pour tout $i \in \llbracket 1, n \rrbracket$

$$\begin{aligned} b_i &= \sum_{j=1}^{i-1} \underbrace{A_{i,j}}_{=0} x_j + A_{i,i}x_i + \sum_{j=i+1}^n A_{i,j}x_j \\ &= A_{i,i}x_i + \sum_{j=i+1}^n A_{i,j}x_j \end{aligned}$$

De plus la matrice $\mathbb{A}$ étant triangulaire inversible ses éléments diagonaux sont tous non nuls. On obtient donc x_i en fonction des $x_{i+1}, \dots, x_n$:

$$x_i = \frac{1}{A_{i,i}} \left(b_i - \sum_{j=i+1}^n A_{i,j}x_j \right), \forall i \in \llbracket 1, n \rrbracket. \quad (4.11)$$

Algorithme 4.3 $\mathcal{R}_0$

- 1: Résoudre $\mathbb{A}\mathbf{x} = \mathbf{b}$ en calculant successivement $x_n, x_{n-1}, \dots, x_1$.

Algorithme 4.3 $\mathcal{R}_1$

- 1: Pour $i \leftarrow n$ à 1 faire(pas de -1)
- 2: calculer x_i connaissant $x_{i+1}, \dots, x_n$ à l'aide de l'équation (??)
- 3: Fin Pour

Algorithme 4.3 $\mathcal{R}_1$

- 1: Pour $i \leftarrow n$ à 1 faire(pas de -1)
- 2: Calculer x_i connaissant $x_{i+1}, \dots, x_n$ à l'aide de l'équation (??)
- 3: Fin Pour

Algorithme 4.3 $\mathcal{R}_2$

- 1: Pour $i \leftarrow n$ à 1 faire(pas de -1)
- 2: $S \leftarrow \sum_{j=i+1}^n A_{i,j}x_j$
- 3: $x_i \leftarrow (b_i - S)/A_{i,i}$
- 4: Fin Pour

Algorithme 4.3 $\mathcal{R}_3$

- 1: Pour $i \leftarrow n$ à 1 faire(pas de -1)
- 2: $S \leftarrow \sum_{j=i+1}^n A_{i,j}x_j$
- 3: $x_i \leftarrow (b_i - S)/A_{i,i}$
- 4: Fin Pour

Algorithme 4.3 $\mathcal{R}_4$

- 1: Pour $i \leftarrow n$ à 1 faire(pas de -1)
- 2: $S \leftarrow 0$
- 3: Pour $j \leftarrow i+1$ à n faire
- 4: $S \leftarrow S + A_{i,j} * x(j)$
- 5: Fin Pour
- 6: $x_i \leftarrow (b_i - S)/A_{i,i}$
- 7: Fin Pour

1 On obtient alors l'algorithme final

Algorithme 4.3 Fonction **RSLTriSup** permettant de résoudre le système linéaire triangulaire supérieur inversible

$$\mathbb{A}\mathbf{x} = \mathbf{b}.$$

Données : $\mathbb{A}$: matrice triangulaire de $\mathcal{M}_n(\mathbb{R})$ supérieure inversible.

$\mathbf{b}$: vecteur de $\mathbb{R}^n$.

Résultat : $\mathbf{x}$: vecteur de $\mathbb{R}^n$.

```

1: Fonction  $\mathbf{x} \leftarrow \text{RSLTriSup}(\mathbb{A}, \mathbf{b})$ 
2:   Pour  $i \leftarrow n$  à 1 faire(pas de -1)
3:      $S \leftarrow 0$ 
4:     Pour  $j \leftarrow i + 1$  à  $n$  faire
5:        $S \leftarrow S + A(i, j) * x(j)$ 
6:     Fin Pour
7:      $x(i) \leftarrow (b(i) - S) / A(i, i)$ 
8:   Fin Pour
9: Fin Fonction
```

2 4.2.2 Méthode de Gauss-Jordan, écriture matricielle

3 Soient $\mathbb{A} \in \mathcal{M}_{\mathbb{K}}()$ une matrice inversible et $\mathbf{b} \in \mathbb{K}^n$.

4 On va tout d'abord rappeler (très) brièvement l'**algorithme d'élimination** ou **algorithme de Gauss-**
5 **Jordan** permettant de transformer le système linéaire $\mathbb{A}\mathbf{x} = \mathbf{b}$ en un système linéaire équivalent dont la
6 matrice est triangulaire supérieure. Ce dernier système se résout par la méthode de remontée.

7 Ensuite, on va réécrire cet algorithme sous forme algébrique pour obtenir le théorème ...



Cette méthode doit son nom aux mathématiciens Carl Friedrich Gauss (1777-1855, mathématicien, astronome et physicien allemand) et Wilhelm Jordan (1842-1899, mathématicien et géodésien allemand) mais elle est connue des Chinois depuis au moins le I^{er} siècle de notre ère. Elle est référencée dans l'important livre chinois *Jiuzhang suanshu* ou *Les Neuf Chapitres sur l'art mathématique*, dont elle constitue le huitième chapitre, sous le titre « Fang cheng » (la disposition rectangulaire). La méthode est présentée au moyen de dix-huit exercices. Dans son commentaire daté de 263, Liu Hui en attribue la paternité à Chang Ts'ang, chancelier de l'empereur de Chine au II^e siècle avant notre ère.

9 Algorithme de Gauss-Jordan usuel

10 Pour la résolution du système linéaire $\mathbb{A}\mathbf{x} = \mathbf{b}$ l'algorithme de Gauss-Jordan produit la forme échelon-
11 née (réduite) d'une matrice à l'aide d'opérations élémentaires sur les lignes du système. Trois types
12 d'opérations élémentaires sont utilisées:

- 13 • Permutation de deux lignes ;
- 14 • Multiplication d'une ligne par un scalaire non nul ;
- 15 • Ajout du multiple d'une ligne à une autre ligne.

16 A l'aide de ces opérations élémentaires cet algorithme permet donc de transformer le système linéaire
17 $\mathbb{A}\mathbf{x} = \mathbf{b}$ en le système équivalent $\mathbb{U}\mathbf{x} = \mathbf{f}$ où $\mathbb{U}$ est triangulaire supérieure. En fait, l'algorithme va
18 transformer la matrice $\mathbb{A}$ et le second membre $\mathbf{b}$ pour aboutir à un système dont la matrice est triangulaire
19 supérieure.

Algorithme 4.4 Algorithme de Gauss-Jordan formel pour la résolution de $\mathbb{A}\mathbf{x} = \mathbf{b}$

- 1: **Pour** $j \leftarrow 1$ à $n - 1$ **faire**
- 2: Rechercher l'indice k de la ligne du pivot (sur la colonne j , $k \in \llbracket j, n \rrbracket$)
- 3: Permuter les lignes j ($\mathcal{L}_j$) et k ($\mathcal{L}_k$) du système si besoin.
- 4: **Pour** $i \leftarrow j + 1$ à n **faire**
- 5: Eliminer en effectuant $\mathcal{L}_i \leftarrow \mathcal{L}_i - \frac{A_{i,j}}{A_{j,j}} \mathcal{L}_j$
- 6: **Fin Pour**
- 7: **Fin Pour**
- 8: Résoudre le système triangulaire supérieur par la méthode de la remontée.

On va maintenant voir comment écrire cet algorithme de manière plus détaillée. Pour conserver sa lisibilité, on choisit pour chaque ligne un peu délicate de créer et d'utiliser une fonction dédiée à cette tâche.

Algorithme 4.5 Algorithme de Gauss-Jordan avec fonctions pour la résolution de $\mathbb{A}\mathbf{x} = \mathbf{b}$

Données : $\mathbb{A}$: matrice de $\mathcal{M}_n(\mathbb{K})$ inversible.
 $\mathbf{b}$: vecteur de $\mathbb{K}^n$.

Résultat : $\mathbf{x}$: vecteur de $\mathbb{K}^n$.

- 1: **Fonction** $\mathbf{x} \leftarrow \text{RSLGAUSS}(\mathbb{A}, \mathbf{b})$
- 2: **Pour** $j \leftarrow 1$ à $n - 1$ **faire**
- 3: $k \leftarrow \text{CHERCHEINDPIVOT}(\mathbb{A}, j)$ ▷ **CHERCHEINDPIVOT** à écrire
- 4: $[\mathbb{A}, \mathbf{b}] \leftarrow \text{PERMLIGNESYS}(\mathbb{A}, \mathbf{b}, j, k)$ ▷ **PERMLIGNESYS** à écrire
- 5: **Pour** $i \leftarrow j + 1$ à n **faire**
- 6: $[\mathbb{A}, \mathbf{b}] \leftarrow \text{COMBLIGNESYS}(\mathbb{A}, \mathbf{b}, j, i, -A(i, j)/A(j, j))$ ▷ **COMBLIGNESYS** à écrire
- 7: **Fin Pour**
- 8: **Fin Pour**
- 9: $\mathbf{x} \leftarrow \text{RSLTRI}(\mathbb{A}, \mathbf{b})$ ▷ **RSLTRI** déjà écrite
- 10: **Fin Fonction**

Bien évidemment, il reste à décrire et écrire les différentes fonctions utilisées dans cette fonction :

Fonction $k \leftarrow \text{CHERCHEINDPIVOT}(\mathbb{A}, j)$: recherche $k \in \llbracket j, n \rrbracket$ tel que $\forall l \in \llbracket j, n \rrbracket$, $|A_{l,j}| \leq |A_{k,j}|$.

Fonction $[\mathbb{A}, \mathbf{b}] \leftarrow \text{PERMLIGNESYS}(\mathbb{A}, \mathbf{b}, i, k)$: permute les lignes i et k de la matrice $\mathbb{A}$ ainsi que celles du vecteur $\mathbf{b}$.

Fonction $[\mathbb{A}, \mathbf{b}] \leftarrow \text{COMBLIGNESYS}(\mathbb{A}, \mathbf{b}, j, i, \alpha)$: remplace la ligne i de la matrice $\mathbb{A}$ par la combinaison linéaire $\mathcal{L}_i \leftarrow \mathcal{L}_i + \alpha \mathcal{L}_j$. De même on remplace la ligne i de $\mathbf{b}$ par $b_i + \alpha b_j$.

Ces trois fonctions sont simples à écrire et sont données en Algorithmes 4.6, 4.7 et 4.8.

Algorithme 4.6 Recherche d'un pivot pour l'algorithme de Gauss-Jordan.

Données : $\mathbb{A}$: matrice de $\mathcal{M}_n(\mathbb{K})$.
 j : entier, $1 \leq j \leq n$.

Résultat : k : entier, indice ligne pivot

- 1: **Fonction** $k \leftarrow \text{CHERCHEINDPIVOT}(\mathbb{A}, j)$
- 2: $k \leftarrow j$, pivot $\leftarrow |A(j, j)|$
- 3: **Pour** $i \leftarrow j + 1$ à n **faire**
- 4: Si $|A(i, j)| > \text{pivot}$ alors
- 5: $k \leftarrow i$, pivot $\leftarrow |A(i, j)|$
- 6: **Fin Si**
- 7: **Fin Fonction**
- 8: **Fin Fonction**

Algorithme 4.7 Permute deux lignes d'une matrice et d'un vecteur.

Données : $\mathbb{A}$: matrice de $\mathcal{M}_n(\mathbb{K})$.
 $\mathbf{b}$: vecteur de $\mathbb{K}^n$.
 j, k : entiers, $1 \leq j, k \leq n$.

Résultat : $\mathbb{A}$ et $\mathbf{b}$ modifiés.

- 1: **Fonction** $[\mathbb{A}, \mathbf{b}] \leftarrow \text{PERMLIGNESYS}(\mathbb{A}, \mathbf{b}, j, k)$
- 2: **Pour** $l \leftarrow 1$ à n **faire**
- 3: $t \leftarrow A(j, l)$, $A(j, l) \leftarrow A(k, l)$, $A(k, l) \leftarrow t$
- 4: **Fin Pour**
- 5: $t \leftarrow b(j)$, $b(j) \leftarrow b(k)$, $b(k) \leftarrow t$
- 6: **Fin Fonction**

Algorithme 4.8 Combinaison linéaire $\mathcal{L}_i \leftarrow \mathcal{L}_i + \alpha \mathcal{L}_j$ appliqué à une matrice et à un vecteur.

Données : $\mathbb{A}$: matrice de $\mathcal{M}_n(\mathbb{K})$.
 $\mathbf{b}$: vecteur de $\mathbb{K}^n$.
 j, i : entiers, $1 \leq j, i \leq n$.
 α : scalaire de $\mathbb{K}$.

Résultat : $\mathbb{A}$ et $\mathbf{b}$ modifiés.

- 1: **Fonction** $[\mathbb{A}, \mathbf{b}] \leftarrow \text{COMBLIGNESYS}(\mathbb{A}, \mathbf{b}, j, i, \alpha)$
- 2: **Pour** $k \leftarrow 1$ à n **faire**
- 3: $A(i, k) \leftarrow A(i, k) + \alpha * A(j, k)$
- 4: **Fin Pour**
- 5: $b(i) \leftarrow b(i) + \alpha b(j)$
- 6: **Fin Fonction**

Ecriture algébrique

Sous forme d'exercice :

**Exercice 4.2.2**

Soit $\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$ inversible.

Q. 1 Montrer qu'il existe une matrice $\mathbb{G} \in \mathcal{M}_n(\mathbb{C})$ telle que $|\det(\mathbb{G})| = 1$ et $\mathbb{G}\mathbb{A}\mathbf{e}_1 = \alpha\mathbf{e}_1$ avec $\alpha \neq 0$ et $\mathbf{e}_1$ premier vecteur de la base canonique de $\mathbb{C}^n$.

Q. 2 1. Montrer par récurrence sur l'ordre des matrices que pour toute matrice $\mathbb{A}_n \in \mathcal{M}_n(\mathbb{C})$ inversible, il existe une matrice $\mathbb{S}_n \in \mathcal{M}_n(\mathbb{C})$ telle que $|\det \mathbb{S}_n| = 1$ et $\mathbb{S}_n\mathbb{A}_n = \mathbb{U}_n$ avec $\mathbb{U}_n$ matrice triangulaire supérieure inversible.

2. Soit $\mathbf{b} \in \mathbb{C}^n$. En supposant connue la décomposition précédente $\mathbb{S}_n\mathbb{A}_n = \mathbb{U}_n$, expliquer comment résoudre le système $\mathbb{A}_n\mathbf{x} = \mathbf{b}$.

Q. 3 Que peut-on dire si $\mathbb{A}$ est non inversible?

Indication : utiliser les résultats des exercices 4.1.2 et 4.1.3.

Correction Exercice 4.2.2

Q. 1 D'après le Lemme 4.3, si $A_{1,1} \neq 0$ le résultat est immédiat. Dans l'énoncé rien ne vient corroborer cette hypothèse. Toutefois, comme la matrice $\mathbb{A}$ est inversible, il existe au moins un $p \in \llbracket 1, n \rrbracket$ tel que $A_{p,1} \neq 0$. On peut même choisir le premier indice p tel que $|A_{p,1}| = \max_{i \in \llbracket 1, n \rrbracket} |A_{i,1}| > 0$ (pivot de l'algorithme de Gauss-Jordan). On note $\mathbb{P} = \mathbb{P}_n^{[1,p]}$ la matrice de permutation des lignes 1 et p (voir exercice 4.1.3, page 73). De plus on a

$$|\det \mathbb{P}| = 1 \quad \text{et} \quad \mathbb{P}^{-1} = \mathbb{P}.$$

Par construction $(\mathbb{P}\mathbb{A})_{1,1} = A_{p,1} \neq 0$, et on peut alors appliquer le Lemme 4.3 à la matrice $(\mathbb{P}\mathbb{A})$ pour obtenir l'existence d'une matrice $\mathbb{E} \in \mathcal{M}_n(\mathbb{C})$ vérifiant $\det \mathbb{E} = 1$ et telle que

$$\mathbb{E}(\mathbb{P}\mathbb{A})\mathbf{e}_1 = A_{p,1}\mathbf{e}_1.$$

En posant $\mathbb{G} = \mathbb{E}\mathbb{P}$ et $\alpha = A_{p,1}$, on obtient bien $\mathbb{G}\mathbb{A}\mathbf{e}_1 = \alpha\mathbf{e}_1$. De plus, on a

$$|\det \mathbb{G}| = |\det(\mathbb{E}\mathbb{P})| = |\det \mathbb{E} \times \det \mathbb{P}| = 1.$$

Remarque 4.5 La matrice $\mathbb{G}$ étant inversible, on a

$$\mathbb{A}\mathbf{x} = \mathbf{b} \iff \mathbb{G}\mathbb{A}\mathbf{x} = \mathbb{G}\mathbf{b}$$

ce qui correspond à la première *permutation/élimination* de l'algorithme de Gauss-Jordan.

Q. 2 1. On veut démontrer par récurrence la propriété $(\mathcal{P}_n)$,

$$(\mathcal{P}_n) \quad \begin{cases} \forall \mathbb{A}_n \in \mathcal{M}_n(\mathbb{C}), \text{ inversible } \exists \mathbb{S}_n \in \mathcal{M}_n(\mathbb{C}), |\det \mathbb{S}_n| = 1, \text{ tel que} \\ \text{la matrice } \mathbb{U}_n = \mathbb{S}_n\mathbb{A} \text{ soit une triangulaire supérieure inversible} \end{cases}$$

Initialisation : Pour $n = 2$. Soit $\mathbb{A}_2 \in \mathcal{M}_2(\mathbb{C})$ inversible. En utilisant la question précédente il existe $\mathbb{G}_2 \in \mathcal{M}_2(\mathbb{C})$ telle que $|\det \mathbb{G}_2| = 1$ et $\mathbb{G}_2\mathbb{A}_2\mathbf{e}_1 = \alpha\mathbf{e}_1$ avec $\alpha \neq 0$ et $\mathbf{e}_1$ premier vecteur de la base canonique de $\mathbb{C}^2$. On note $\mathbb{U}_2 = \mathbb{G}_2\mathbb{A}_2$. Cette matrice s'écrit donc sous la forme

$$\mathbb{U}_2 = \begin{pmatrix} \alpha & \bullet \\ 0 & \bullet \end{pmatrix}$$

et elle est triangulaire supérieure. Les matrices $\mathbb{G}_2$ et $\mathbb{A}_2$ étant inversible, leur produit $\mathbb{U}_2$ l'est aussi. La proposition $(\mathcal{P}_2)$ est donc vérifiée avec $\mathbb{S}_2 = \mathbb{G}_2$.

Hérédité : Soit $n \geq 3$. On suppose que $(\mathcal{P}_{n-1})$ est vraie. Montrons que $(\mathcal{P}_n)$ est vérifiée.

Soit $\mathbb{A}_n \in \mathcal{M}_n(\mathbb{C})$ inversible. En utilisant la question précédente il existe $\mathbb{G}_n \in \mathcal{M}_n(\mathbb{C})$ telle que $|\det \mathbb{G}_n| = 1$ et $\mathbb{G}_n\mathbb{A}_n\mathbf{e}_1 = \alpha_n\mathbf{e}_1$ avec $\alpha_n \neq 0$ et $\mathbf{e}_1$ premier vecteur de la base canonique de $\mathbb{C}^n$. On note $\mathbb{V}_n = \mathbb{G}_n\mathbb{A}_n$. Cette matrice s'écrit donc sous la forme

$$\mathbb{V}_n = \left(\begin{array}{c|ccc} \alpha_n & \bullet & \dots & \bullet \\ 0 & \bullet & \dots & \bullet \\ \vdots & \vdots & & \vdots \\ 0 & \bullet & \dots & \bullet \end{array} \right) \stackrel{\text{def}}{=} \left(\begin{array}{c|ccc} \alpha_n & \mathbf{c}_{n-1}^* & & \\ 0 & & & \\ \vdots & & \mathbb{B}_{n-1} & \\ 0 & & & \end{array} \right)$$

où $\mathbf{c}_{n-1} \in \mathbb{C}^{n-1}$ et $\mathbb{B}_{n-1} \in \mathcal{M}_{n-1}(\mathbb{C})$. Comme $\mathbb{G}_n$ et $\mathbb{A}_n$ sont inversibles, $\mathbb{V}_n$ l'est aussi. On en déduit donc que $\mathbb{B}_{n-1}$ est inversible car $0 \neq \det \mathbb{V}_n = \alpha_n \times \det \mathbb{B}_{n-1}$ et $\alpha_n \neq 0$.

On peut donc utiliser la propriété $(\mathcal{P}_{n-1})$ (hyp. de récurrence) sur la matrice $\mathbb{B}_{n-1}$: il existe donc $\mathbb{S}_{n-1} \in \mathcal{M}_{n-1}(\mathbb{C})$, avec $|\det \mathbb{S}_{n-1}| = 1$, tel que la matrice $\mathbb{U}_{n-1} = \mathbb{S}_{n-1}\mathbb{B}_{n-1}$ soit une triangulaire supérieure inversible.

Soit $\mathbb{Q}_n \in \mathcal{M}_n(\mathbb{C})$ la matrice définie par

$$\mathbb{Q}_n = \left(\begin{array}{c|ccc} 1 & 0 & \dots & 0 \\ \hline 0 & & & \\ \vdots & & \mathbb{S}_{n-1} & \\ 0 & & & \end{array} \right)$$

On a alors

$$\begin{aligned} \mathbb{Q}_n \mathbb{G}_n \mathbb{A}_n &= \mathbb{Q}_n \mathbb{V}_n = \left(\begin{array}{c|ccc} 1 & 0 & \dots & 0 \\ \hline 0 & & & \\ \vdots & & \mathbb{S}_{n-1} & \\ 0 & & & \end{array} \right) \left(\begin{array}{c|c} \alpha_n & \mathbf{c}_{n-1}^* \\ \hline 0 & \mathbb{B}_{n-1} \\ \vdots & \\ 0 & \end{array} \right) \\ &= \left(\begin{array}{c|c} \alpha_n & \mathbf{c}_{n-1}^* \\ \hline 0 & \mathbb{S}_{n-1}\mathbb{B}_{n-1} \\ \vdots & \\ 0 & \end{array} \right) = \left(\begin{array}{c|c} \alpha_n & \mathbf{c}_{n-1}^* \\ \hline 0 & \mathbb{U}_{n-1} \\ \vdots & \\ 0 & \end{array} \right) \stackrel{\text{def}}{=} \mathbb{U}_n \end{aligned}$$

La matrice $\mathbb{U}_n$ est triangulaire supérieure inversible car $\mathbb{U}_{n-1}$ l'est aussi et $\alpha_n \neq 0$.

On pose $\mathbb{S}_n = \mathbb{Q}_n \mathbb{G}_n$. On a donc $\mathbb{S}_n \mathbb{A}_n = \mathbb{U}_n$ et comme $|\det \mathbb{S}_n| = 1$ car $|\det \mathbb{G}_n| = 1$ et $\det \mathbb{G}_n = 1$, ceci prouve la véracité de la proposition $(\mathcal{P}_n)$.

2. Comme $\mathbb{S}_n$ est inversible, on a en multipliant à gauche le système par $\mathbb{S}_n$

$$\mathbb{A}_n \mathbf{x} = \mathbf{b} \iff \mathbb{S}_n \mathbb{A}_n \mathbf{x} = \mathbb{S}_n \mathbf{b} \iff \mathbb{U}_n \mathbf{x} = \mathbb{S}_n \mathbf{b}$$

Pour déterminer le vecteur $\mathbf{x}$, on peut alors résoudre le dernier système par l'algorithme de remontée.

Q. 3 (rapide) Si $\mathbb{A}$ est non inversible, alors dans la première question nous ne sommes pas assurés d'avoir $\alpha \neq 0$. Cependant l'existence de la matrice $\mathbb{G}$ reste avérée.

Pour la deuxième question, le seul changement vient du fait que la matrice $\mathbb{U}_n$ n'est plus inversible.

◇

On a donc démontré le théorème suivant



Théorème 4.6

Soit $\mathbb{A}$ une matrice carrée, inversible ou non. Il existe (au moins) une matrice inversible $\mathbb{G}$ telle que $\mathbb{G}\mathbb{A}$ soit triangulaire supérieure.

Proof. voir Exercice 4.2.2, page 80.

□

4.2.3 Factorisation LU

Avant de citer le théorème principal, on va faire un "petit" exercice...



Exercice 4.2.3:



Soit $\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$ une matrice dont les sous-matrices principales d'ordre i , notées Δ_i , $i \in \llbracket 1, n \rrbracket$ (voir Définition A.39, page 183) sont inversibles.

Montrer qu'il existe des matrices $\mathbb{E}^{[k]} \in \mathcal{M}_n(\mathbb{C})$, $k \in \llbracket 1, n-1 \rrbracket$, triangulaires inférieures à diagonale unité telles que la matrice $\mathbb{U}$ définie par

$$\mathbb{U} = \mathbb{E}^{[n-1]} \dots \mathbb{E}^{[1]} \mathbb{A}$$

soit triangulaire supérieure avec $U_{i,i} = \det \Delta_i / (U_{1,1} \times \dots \times U_{i-1,i-1})$, $\forall i \in \llbracket 1, n \rrbracket$.

Correction Exercice 4.2.3

On note $\mathbb{A}^{[0]} = \mathbb{A}$. On va démontrer par récurrence finie sur $k \in \llbracket 1, n-1 \rrbracket$, qu'il existe une matrice $\mathbb{E}^{[k]} \in \mathcal{M}_n(\mathbb{C})$, tel que la matrice $\mathbb{A}^{[k]}$ définie itérativement par

$$\mathbb{A}^{[k]} = \mathbb{E}^{[k]} \mathbb{A}^{[k-1]}$$

s'écrit sous la forme bloc

$$\mathbb{A}^{[k]} = \left(\begin{array}{cccc|cccc} \alpha_1 & \bullet & \dots & \bullet & \bullet & \dots & \bullet & \\ 0 & \ddots & \ddots & \vdots & \vdots & & \vdots & \\ \vdots & \ddots & \ddots & \bullet & \vdots & & \vdots & \\ 0 & \dots & 0 & \alpha_k & \bullet & \dots & \bullet & \\ \hline 0 & \dots & \dots & 0 & \bullet & \dots & \bullet & \\ \vdots & \ddots & \ddots & \vdots & \vdots & & \vdots & \\ \vdots & \ddots & \ddots & \vdots & \vdots & & \vdots & \\ 0 & \dots & \dots & 0 & \bullet & \dots & \bullet & \end{array} \right) \quad (4.12)$$

avec $\alpha_1 = A_{1,1}$ et $\forall i \in \llbracket 2, k \rrbracket$, $\alpha_i = \det \Delta_i / (\alpha_1 \times \dots \times \alpha_{i-1})$.

Initialisation ($k=1$): On a $A_{1,1} \neq 0$ car $\Delta_1 = A_{1,1}$ et $\det \Delta_1 \neq 0$. D'après le Lemme 4.3, il existe une matrice $\mathbb{E}^{[1]} \in \mathcal{M}_n(\mathbb{C})$, triangulaire inférieure à diagonale unité telle que $\mathbb{E}^{[1]} \mathbb{A} \mathbf{e}_1 = A_{1,1} \mathbf{e}_1$ où $\mathbf{e}_1$ est le premier vecteur de la base canonique de $\mathbb{C}^n$. On a alors

$$\mathbb{A}^{[1]} = \mathbb{E}^{[1]} \mathbb{A} = \left(\begin{array}{c|cccc} \alpha_1 & \bullet & \dots & \bullet & \\ \hline 0 & \bullet & \dots & \bullet & \\ \vdots & \vdots & & \vdots & \\ 0 & \bullet & \dots & \bullet & \end{array} \right)$$

avec $\alpha_1 = A_{1,1} = \det \Delta_1$.

Hérédité ($k < n-1$): Supposons construite la matrice $\mathbb{A}^{[k]}$. Il existe donc k matrices, $\mathbb{E}^{[1]}, \dots, \mathbb{E}^{[k]}$, triangulaires inférieures à diagonale unité telles que

$$\mathbb{A}^{[k]} = \mathbb{E}^{[k]} \dots \mathbb{E}^{[1]} \mathbb{A}.$$

• On va montrer que $\alpha_{k+1} \stackrel{\text{def}}{=} A_{k+1,k+1}^{[k]} \neq 0$. Pour cela, on réécrit la matrice $\mathbb{A}^{[k]}$ sous forme bloc, avec comme premier bloc diagonale le bloc de dimension $k+1$:

$$\mathbb{A}^{[k]} = \left(\begin{array}{cccc|cccc} \alpha_1 & \bullet & \dots & \bullet & \bullet & \dots & \bullet & \\ 0 & \ddots & \ddots & \vdots & \vdots & & \vdots & \\ \vdots & \ddots & \ddots & \bullet & \bullet & & \vdots & \\ 0 & \dots & 0 & \alpha_k & \bullet & \dots & \bullet & \\ \hline 0 & \dots & \dots & 0 & \alpha_{k+1} & \bullet & \dots & \bullet \\ \hline 0 & \dots & \dots & 0 & \bullet & \dots & \bullet & \\ \vdots & & & \vdots & \bullet & & \vdots & \\ 0 & \dots & \dots & 0 & \bullet & \dots & \bullet & \end{array} \right)$$

La matrice $\mathbb{G}^{[k]} \stackrel{\text{def}}{=} \mathbb{E}^{[k]} \dots \mathbb{E}^{[1]}$ est triangulaire inférieure à diagonale unité car produit de matrices triangulaires inférieures à diagonale unité (voir ...). Le produit de $\mathbb{G}^{[k]} \mathbb{A}$ s'écrit alors sous forme bloc

$$\mathbb{G}^{[k]} \mathbb{A} = \left(\begin{array}{c|c} \begin{matrix} \xrightarrow{k+1} \\ \begin{pmatrix} 1 & 0 & \dots & 0 & 0 & \dots & \dots & 0 \\ \bullet & \ddots & \ddots & \vdots & \vdots & & & \vdots \\ \vdots & \ddots & \ddots & 0 & \vdots & & & \vdots \\ \bullet & \dots & \bullet & 1 & 0 & \dots & \dots & 0 \\ \bullet & \dots & \dots & \bullet & 1 & 0 & \dots & 0 \\ \vdots & \ddots & \ddots & \vdots & \vdots & & & \vdots \\ \vdots & \ddots & \ddots & \vdots & \bullet & \ddots & \ddots & \vdots \\ \bullet & \dots & \dots & \bullet & \vdots & \ddots & \ddots & 0 \\ \bullet & \dots & \dots & \bullet & \bullet & \dots & \bullet & 1 \end{pmatrix} \end{matrix} & \begin{matrix} \xleftarrow{k+1} \\ \begin{pmatrix} \Delta_{k+1} & \vdots & \vdots & \vdots \\ \vdots & \ddots & \ddots & \vdots \\ \vdots & \vdots & \ddots & \vdots \\ \vdots & \vdots & \vdots & \ddots \end{pmatrix} \end{matrix} \end{array} \right)$$

Comme $\mathbb{A}^{[k]} = \mathbb{G}^{[k]} \mathbb{A}$, en utilisant les règles de multiplication par blocs des matrices on obtient

$$\begin{pmatrix} \alpha_1 & \bullet & \dots & \bullet & \bullet \\ 0 & \ddots & \ddots & \vdots & \bullet \\ \vdots & \ddots & \ddots & \bullet & \bullet \\ 0 & \dots & 0 & \alpha_k & \bullet \\ 0 & \dots & \dots & 0 & \alpha_{k+1} \end{pmatrix} = \begin{pmatrix} 1 & 0 & \dots & 0 \\ \bullet & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ \bullet & \dots & \bullet & 1 \end{pmatrix} \begin{pmatrix} \Delta_{k+1} \\ \vdots \\ \vdots \\ \vdots \end{pmatrix}$$

En prenant le déterminant de cette dernière équation, et en utilisant le fait que le déterminant de matrices triangulaires est le produit de ses coefficients diagonaux, on obtient

$$\prod_{i=1}^{k+1} \alpha_i = \det \Delta_{k+1}.$$

Or Δ_{k+1} inversible par hypothèse, ce qui entraîne $\det \Delta_{k+1} \neq 0$ et donc $\alpha_i \neq 0, \forall i \in \llbracket 1, k+1 \rrbracket$ et on a

$$\alpha_{k+1} = \frac{\det \Delta_{k+1}}{\prod_{i=1}^k \alpha_i} \neq 0.$$

- Montrons l'existence d'une matrice triangulaire inférieure à diagonale unité permettant d'éliminer les termes sous diagonaux de la colonne $k+1$ de $\mathbb{A}^{[k]}$.

Revenons à l'écriture bloc de premier bloc diagonal de dimension k . On a

$$\mathbb{A}^{[k]} = \left(\begin{array}{c|c} \begin{matrix} \begin{pmatrix} \alpha_1 & \bullet & \dots & \bullet \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & \bullet \\ 0 & \dots & 0 & \alpha_k \end{pmatrix} & \begin{pmatrix} \bullet & \dots & \dots & \bullet \\ \vdots & & & \vdots \\ \vdots & & & \vdots \\ \bullet & \dots & \dots & \bullet \end{pmatrix} \end{matrix} & \begin{matrix} \begin{pmatrix} \vdots & \vdots & \vdots & \vdots \\ \vdots & \ddots & \ddots & \vdots \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{k+1} & \bullet & \dots & \bullet \end{pmatrix} & \begin{pmatrix} \vdots & \vdots & \vdots & \vdots \\ \vdots & \ddots & \ddots & \vdots \\ \vdots & \vdots & \ddots & \vdots \\ \bullet & \dots & \bullet & \bullet \end{pmatrix} \end{matrix} \end{array} \right) \stackrel{\text{def}}{=} \begin{pmatrix} \mathbb{U}^{[k]} & \mathbb{F}^{[k]} \\ \mathbb{O} & \mathbb{V}^{[k]} \end{pmatrix}$$

Nous sommes exactement dans le cas de figure étudié dans l'exercice 4.1.2. En effet, avec les notations de cet exercice et si l'on pose $\mathbf{v} = \mathbb{V}_{:,1}^{[k]} = (A_{k+1,k+1}^{[k]}, \dots, A_{n,k+1}^{[k]})^t \in \mathbb{C}^{n-(k+1)}$ on a alors $v_1 = A_{k+1,k+1}^{[k]} = \alpha_{k+1} \neq 0$ et l'on peut définir la matrice $\mathbb{E}^{[k+1]}$, triangulaire inférieure à diagonale unité, par

$$\mathbb{E}^{[k+1]} = \mathbb{E}^{[k], \mathbf{v}} \stackrel{\text{def}}{=} \begin{pmatrix} \mathbb{I}_k & \mathbb{O} \\ \mathbb{O} & \mathbb{E}^{[\mathbf{v}]} \end{pmatrix} \in \mathcal{M}_n(\mathbb{C})$$

- 1 avec $\mathbb{E}^{[v]} \in \mathcal{M}_{n-k}(\mathbb{C})$ triangulaire inférieure à diagonale unité (définie dans l'exercice 4.1.2)
 2 telle que

$$\mathbb{E}^{[v]}\mathbb{V}^{[k]} = \begin{pmatrix} \alpha_{k+1} & \bullet & \cdots & \bullet \\ 0 & \vdots & & \vdots \\ \vdots & \vdots & & \vdots \\ 0 & \bullet & \cdots & \bullet \end{pmatrix}$$

On a alors

$$\begin{aligned} \mathbb{A}^{[k+1]} &\stackrel{\text{def}}{=} \mathbb{E}^{[k+1]}\mathbb{A}^{[k]} = \begin{pmatrix} \mathbb{I}_k & \mathbb{O} \\ \mathbb{O} & \mathbb{E}^{[v]} \end{pmatrix} \begin{pmatrix} \mathbb{U}^{[k]} & \mathbb{F}^{[k]} \\ \mathbb{O} & \mathbb{V}^{[k]} \end{pmatrix} \\ &= \begin{pmatrix} \mathbb{U}^{[k]} & \mathbb{F}^{[k]} \\ \mathbb{O} & \mathbb{E}^{[v]}\mathbb{V}^{[k]} \end{pmatrix} \end{aligned}$$

- 3 et donc $\mathbb{A}^{[k+1]}$ s'écrit bien sous la forme (4.12) au rang $k+1$.

- 4 **Final** ($k = n-1$): On a donc

$$\mathbb{U} = \mathbb{A}^{[n-1]} \stackrel{\text{def}}{=} \mathbb{E}^{[n-1]} \times \cdots \times \mathbb{E}^{[1]} \times \mathbb{A} = \begin{pmatrix} \alpha_1 & \bullet & \cdots & \cdots & \bullet \\ 0 & \ddots & & & \vdots \\ \vdots & \ddots & \ddots & & \vdots \\ 0 & \cdots & 0 & \alpha_{n-1} & \bullet \\ \hline 0 & \cdots & \cdots & 0 & \bullet \end{pmatrix} \quad (4.13)$$

- 5 où pour tout $k \in \llbracket 1, n-1 \rrbracket$ les matrices $\mathbb{E}^{[k]}$ sont triangulaires inférieures à diagonale unité.

- 6 Pour achever l'exercice, il reste à démontrer que

$$U_{n,n} = \det \Delta_n / (U_{1,1} \times \cdots \times U_{n-1,n-1}).$$

- 7 En effet, en prenant le déterminant dans (4.13) on obtient

$$\det \left(\mathbb{E}^{[n-1]} \times \cdots \times \mathbb{E}^{[1]} \times \mathbb{A} \right) = \det \begin{pmatrix} U_{1,1} & \bullet & \cdots & \cdots & \bullet \\ 0 & \ddots & & & \vdots \\ \vdots & \ddots & \ddots & & \vdots \\ 0 & \cdots & 0 & U_{n-1,n-1} & \bullet \\ \hline 0 & \cdots & \cdots & 0 & U_{n,n} \end{pmatrix}$$

Comme le déterminant d'un produit de matrices est égale au produit des déterminants des matrices on a

$$\begin{aligned} \det \left(\mathbb{E}^{[n-1]} \times \cdots \times \mathbb{E}^{[1]} \times \mathbb{A} \right) &= \det \mathbb{E}^{[n-1]} \times \cdots \times \det \mathbb{E}^{[1]} \times \det \mathbb{A} \\ &= \det \mathbb{A} \end{aligned}$$

- 8 car les matrices $\mathbb{E}^{[k]}$ sont triangulaires inférieures à diagonale unité et donc $\det \mathbb{E}^{[k]} = 1, \forall k \in \llbracket 1, n-1 \rrbracket$.
 9 De plus, le déterminant d'une matrice triangulaire supérieure est égale au produit de ses coefficients
 10 diagonaux et donc

$$\det \begin{pmatrix} U_{1,1} & \bullet & \cdots & \cdots & \bullet \\ 0 & \ddots & & & \vdots \\ \vdots & \ddots & \ddots & & \vdots \\ 0 & \cdots & 0 & U_{n-1,n-1} & \bullet \\ \hline 0 & \cdots & \cdots & 0 & U_{n,n} \end{pmatrix} = U_{n,n} \prod_{k=1}^{n-1} U_{k,k}.$$

- 11 On a alors

$$\det \mathbb{A} = \det \Delta_n = U_{n,n} \prod_{k=1}^{n-1} U_{k,k}, k \neq 0.$$

12

13

◇

**Théorème 4.7: Factorisation LU**

Soit $A \in \mathcal{M}_n(\mathbb{C})$ une matrice dont les sous-matrices principales sont inversibles alors il existe une unique matrice $L \in \mathcal{M}_n(\mathbb{C})$ triangulaire inférieure (*lower triangular* en anglais) à diagonale unité et une unique matrice $U \in \mathcal{M}_n(\mathbb{C})$ triangulaire supérieure (*upper triangular* en anglais) inversible telles que

$$A = LU.$$

1

Proof. En utilisant le résultat de l'exercice 4.2.3, il existe des matrices $E^{[k]} \in \mathcal{M}_n(\mathbb{C})$, $k \in \llbracket 1, n-1 \rrbracket$, triangulaires inférieures à diagonale unité telles que la matrice U définie par

2

3

$$U = E^{[n-1]} \dots E^{[1]} A$$

soit triangulaire supérieure avec $U_{i,i} \neq 0$, $\forall i \in \llbracket 1, n \rrbracket$. On pose $G = E^{[n-1]} \dots E^{[1]}$. La matrice G est donc aussi triangulaire inférieure à diagonale unité. Elle est donc inversible et son inverse est triangulaire inférieure à diagonale unité (voir Proposition A.37, page 182). En notant $L = G^{-1}$ on a $A = LU$.

4

5

6

On vient de démontrer l'existence d'une factorisation LU de la matrice A . Pour démontrer l'unicité, on va supposer qu'il existe deux factorisations LU de A i.e.

7

8

$$A = L_1 U_1 = L_2 U_2.$$

avec L_1, L_2 matrices triangulaires inférieures à diagonale unité et U_1, U_2 matrices triangulaires supérieures (inversibles). En multipliant l'équation $L_1 U_1 = L_2 U_2$ à gauche par L_1^{-1} et à droite par U_2^{-1} on obtient

9

10

$$U_1 U_2^{-1} = L_1^{-1} L_2. \quad (4.14)$$

La matrice $L_1^{-1} L_2$ est triangulaire inférieure à diagonale unité car produit de deux matrices triangulaires inférieures à diagonale unité. Elle est égale à la matrice $U_1 U_2^{-1}$ qui elle est triangulaire supérieure (car produit de deux matrices triangulaires supérieures). Donc $L_1^{-1} L_2$ est à la fois une matrice triangulaire supérieure et inférieure : elle est donc diagonale. Comme elle est à diagonale unité on en déduit que $L_1^{-1} L_2 = I$ et donc $L_1 = L_2$. De l'équation (4.14), on tire alors $U_1 = U_2$. $\square$

11

12

13

14

15

**Exercice 4.2.4**

Montrer que si A inversible admet une factorisation LU alors toutes ses sous-matrices principales sont inversibles.

16

Remarque 4.8 Si la matrice $A \in \mathcal{M}_n(\mathbb{C})$ est inversible mais que ses sous-matrices principales ne sont pas toutes inversibles, il est possible par des permutations préalables de lignes de la matrice de se ramener à une matrice telle que ses sous-matrices principales soient inversibles.

17

18

19

**Théorème 4.9: Factorisation LU avec permutations**

Soit $A \in \mathcal{M}_n(\mathbb{C})$ une matrice inversible. Il existe une matrice P , produit de matrices de permutation, une matrice $L \in \mathcal{M}_n(\mathbb{C})$ triangulaire inférieure à diagonale unité et une matrice $U \in \mathcal{M}_n(\mathbb{C})$ triangulaire supérieure telles que

$$PA = LU. \quad (4.15)$$

20

**Corollaire 4.10:**

Si $A \in \mathcal{M}_n(\mathbb{C})$ est une matrice hermitienne définie positive alors elle admet une unique factorisation LU.

21

Proof. (indication) Si la matrice A est hermitienne définie positive alors toutes ses sous-matrices principales sont définies positives et donc inversibles. $\square$

22

23

1 Résolution d'un système linéaire par factorisation LU

2 Soit $A \in \mathcal{M}_n(\mathbb{K})$ une matrice dont les sous-matrices principales sont inversibles et $b \in \mathbb{K}^n$. On veut résoudre le système linéaire $Ax = b$. Pour cela, grâce au théorème 4.7, on obtient :

Trouver $x \in \mathbb{K}^n$ tel que

$$Ax = b. \quad (4.16)$$

6 est équivalent à

Trouver $x \in \mathbb{K}^n$ solution de

$$Ux = y \quad (4.17)$$

avec $y \in \mathbb{K}^n$ solution de

$$Ly = b. \quad (4.18)$$

9 Ceci permet donc de découper le problème initial en trois sous-problèmes plus simples. De plus, ceux-ci peuvent se traiter de manière indépendante. On obtient alors l'algorithme suivant

Algorithme 4.9 Fonction RSLFactLU permettant de résoudre, par une factorisation LU, le système linéaire

$$Ax = b$$

où A une matrice de $\mathcal{M}_n(\mathbb{R})$ définie positive et $b \in \mathbb{R}^n$.

Données : A : matrice de $\mathcal{M}_n(\mathbb{R})$ dont les sous-matrices principales sont inversibles définie positive,
 b : vecteur de $\mathbb{R}^n$.

Résultat : x : vecteur de $\mathbb{R}^n$.

1: **Fonction** $x \leftarrow \text{RSLFactLU}(A, b)$

2: $[L, U] \leftarrow \text{FactLU}(A)$

▷ Factorisation LU

3: $y \leftarrow \text{RSLTriInf}(L, b)$

▷ Résolution du système $Ly = b$

4: $x \leftarrow \text{RSLTriSup}(U, y)$

▷ Résolution du système $Ux = y$

5: **Fin Fonction**

11 Il nous faut donc écrire la fonction **FactLU** (les deux fonctions **RSLTriInf** et **RSLTriSup** ayant déjà été écrites).

13 Détermination des matrices L et U

14 Soit $A \in \mathcal{M}_n(\mathbb{K})$ une matrice dont les sous-matrices principales sont inversibles. D'après le théorème 4.7, il existe une unique matrice $L \in \mathcal{M}_n(\mathbb{K})$ triangulaire inférieure avec $L_{i,i} = 1, \forall i \in \llbracket 1, n \rrbracket$, et une unique matrice $U \in \mathcal{M}_n(\mathbb{K})$ triangulaire supérieure telles que

$$A = LU \quad (4.19)$$

17 c'est à dire

$$\begin{pmatrix} A_{1,1} & \dots & A_{1,n} \\ \vdots & \ddots & \vdots \\ A_{n,1} & \dots & A_{n,n} \end{pmatrix} = \begin{pmatrix} 1 & 0 & \dots & 0 \\ L_{2,1} & \ddots & \ddots & \vdots \\ \vdots & & \ddots & 0 \\ L_{n,1} & \dots & L_{n,n-1} & 1 \end{pmatrix} \begin{pmatrix} U_{1,1} & \dots & \dots & U_{1,n} \\ 0 & \ddots & & \vdots \\ \vdots & \ddots & \ddots & \vdots \\ 0 & \dots & 0 & U_{n,n} \end{pmatrix}. \quad (4.20)$$

18 Pour déterminer les matrices L et U , on remarque que la 1ère ligne de L est déjà déterminée. On peut

alors l'utiliser pour calculer la première ligne de $\mathbb{U}$: $\forall j \in \llbracket 1, n \rrbracket$

$$\begin{aligned} A_{1,j} &= \sum_{k=1}^n L_{1,k} U_{k,j} \\ &= L_{1,1} U_{1,j} \text{ car } \mathbb{L} \text{ triangulaire inférieure} \\ &= U_{1,j} \end{aligned}$$

On a donc

$$U_{1,j} = A_{1,j}, \quad \forall j \in \llbracket 1, n \rrbracket. \quad (4.21)$$

La première colonne de $\mathbb{U}$ est aussi déterminée, on peut alors l'utiliser pour calculer la première colonne de $\mathbb{L}$: $\forall i \in \llbracket 2, n \rrbracket$

$$\begin{aligned} A_{i,1} &= \sum_{k=1}^n L_{i,k} U_{k,1} \\ &= L_{i,1} U_{1,1} \text{ car } \mathbb{U} \text{ triangulaire supérieure} \end{aligned}$$

On peut démontrer, de part les hypothèses sur la matrice $\mathbb{A}$, que $U_{1,1} \neq 0$ et alors

$$L_{1,j} = A_{i,j}/U_{1,1}, \quad \forall i \in \llbracket 2, n \rrbracket. \quad (4.22)$$

La première ligne de $\mathbb{U}$ et la première colonne de $\mathbb{L}$ sont donc déterminées par les formules (4.21) et (4.22).

Par récurrence, on suppose connues les $i-1$ premières lignes de $\mathbb{U}$ et les $i-1$ premières colonnes de $\mathbb{L}$. On va montrer que l'on peut expliciter la i -ème ligne de $\mathbb{U}$ et la i -ème colonne de $\mathbb{L}$. En effet, $\forall j \in \llbracket i, n \rrbracket$, on a

$$\begin{aligned} A_{i,j} &= \sum_{k=1}^n L_{i,k} U_{k,j} \\ &= \sum_{k=1}^{i-1} L_{i,k} U_{k,j} + L_{i,i} U_{i,j} + \sum_{k=i+1}^n L_{i,k} U_{k,j} \\ &= \sum_{k=1}^{i-1} L_{i,k} U_{k,j} + L_{i,i} U_{i,j} \text{ car } \mathbb{L} \text{ triangulaire inférieure} \end{aligned}$$

Dans l'expression $\sum_{k=1}^{i-1} L_{i,k} U_{k,j}$ tous les termes sont connus (hypothèse de récurrence) et $L_{i,i} = 1$. On en déduit,

$$\text{Pour } i \text{ allant de } 1 \text{ à } n : \quad U_{i,j} = \begin{cases} A_{i,j} - \sum_{k=1}^{i-1} L_{i,k} U_{k,j}, & \forall j \in \llbracket i, n \rrbracket. \\ 0, & \forall j \in \llbracket 1, i-1 \rrbracket. \end{cases} \quad (4.23)$$

Maintenant, on calcule, $\forall j \in \llbracket i+1, n \rrbracket$,

$$\begin{aligned} A_{j,i} &= \sum_{k=1}^n L_{j,k} U_{k,i} \\ &= \sum_{k=1}^{i-1} L_{j,k} U_{k,i} + L_{j,i} U_{i,i} + \sum_{k=i+1}^n L_{j,k} U_{k,i} \\ &= \sum_{k=1}^{i-1} L_{j,k} U_{k,i} + L_{j,i} U_{i,i} \text{ car } \mathbb{U} \text{ triangulaire supérieure} \end{aligned}$$

Dans l'expression $\sum_{k=1}^{i-1} L_{j,k} U_{k,i}$ tous les termes sont connus (hypothèse de récurrence). De plus $U_{i,i}$ est donné par (4.23) et on peut démontrer, de part les hypothèses sur la matrice $\mathbb{A}$, que $U_{i,i} \neq 0$. On a alors

$$\text{Pour } i \text{ allant de } 1 \text{ à } n : \quad L_{j,i} = \begin{cases} 0, & \forall j \in \llbracket 1, i-1 \rrbracket. \\ 1, & j = i \\ \frac{1}{U_{i,i}} \left(A_{j,i} - \sum_{k=1}^{i-1} L_{j,k} U_{k,i} \right), & \forall j \in \llbracket i+1, n \rrbracket, \end{cases} \quad (4.24)$$

- 1 On écrit en détail les raffinements successifs permettant d'aboutir à l'algorithme final de telle sorte
 2 que le passage entre deux raffinements successifs soit le plus compréhensible possible.
 3

Algorithme 4.10 $\mathcal{R}_0$

 1: Calculer les matrices $\mathbb{L}$ et $\mathbb{U}$
Algorithme 4.10 $\mathcal{R}_1$

 1: **Pour** $i \leftarrow 1$ à n faire
 2: Calculer la ligne i de $\mathbb{U}$.
 3: Calculer la colonne i de $\mathbb{L}$.
 4: **Fin Pour**
Algorithme 4.10 $\mathcal{R}_1$

 1: **Pour** $i \leftarrow 1$ à n faire
 2: Calculer la ligne i de $\mathbb{U}$.
 3: Calculer la colonne i de $\mathbb{L}$.
 4: **Fin Pour**
Algorithme 4.10 $\mathcal{R}_2$

 1: **Pour** $i \leftarrow 1$ à n faire
 2: **Pour** $j \leftarrow 1$ à $i - 1$ faire
 3: $U(i, j) \leftarrow 0$
 4: **Fin Pour**
 5: **Pour** $j \leftarrow i$ à n faire
 6: $U_{i,j} \leftarrow A_{i,j} - \sum_{k=1}^{i-1} L_{i,k} U_{k,j}$
 7: **Fin Pour**
 8: **Pour** $j \leftarrow 1$ à $i - 1$ faire
 9: $L_{j,i} \leftarrow 0$
 10: **Fin Pour**
 11: $L_{i,i} \leftarrow 1$
 12: **Pour** $j \leftarrow i + 1$ à n faire
 13: $L_{j,i} \leftarrow \frac{1}{U_{i,i}} \left(A_{j,i} - \sum_{k=1}^{i-1} L_{j,k} U_{k,i} \right)$
 14: **Fin Pour**
 15: **Fin Pour**
Algorithme 4.10 $\mathcal{R}_2$

 1: **Pour** $i \leftarrow 1$ à n faire
 2: **Pour** $j \leftarrow 1$ à $i - 1$ faire
 3: $U(i, j) \leftarrow 0$
 4: **Fin Pour**
 5: **Pour** $j \leftarrow i$ à n faire
 6: $U_{i,j} \leftarrow A_{i,j} - \sum_{k=1}^{i-1} L_{i,k} U_{k,j}$
 7: **Fin Pour**
 8: **Pour** $j \leftarrow 1$ à $i - 1$ faire
 9: $L_{j,i} \leftarrow 0$
 10: **Fin Pour**
 11: $L_{i,i} \leftarrow 1$
 12: **Pour** $j \leftarrow i + 1$ à n faire
 13: $L_{j,i} \leftarrow \frac{1}{U_{i,i}} \left(A_{j,i} - \sum_{k=1}^{i-1} L_{j,k} U_{k,i} \right)$
 14: **Fin Pour**
 15: **Fin Pour**
Algorithme 4.10 $\mathcal{R}_3$

 1: **Pour** $i \leftarrow 1$ à n faire
 2: **Pour** $j \leftarrow 1$ à $i - 1$ faire
 3: $U(i, j) \leftarrow 0$
 4: **Fin Pour**
 5: **Pour** $j \leftarrow i$ à n faire
 6: $S_1 \leftarrow \sum_{k=1}^{i-1} L_{i,k} U_{k,j}$
 7: $U_{i,j} \leftarrow A_{i,j} - S_1$
 8: **Fin Pour**
 9: **Pour** $j \leftarrow 1$ à $i - 1$ faire
 10: $L_{j,i} \leftarrow 0$
 11: **Fin Pour**
 12: $L_{i,i} \leftarrow 1$
 13: **Pour** $j \leftarrow i + 1$ à n faire
 14: $S_2 \leftarrow \sum_{k=1}^{i-1} L_{j,k} U_{k,i}$
 15: $L_{j,i} \leftarrow \frac{1}{U_{i,i}} (A_{j,i} - S_2)$
 16: **Fin Pour**
 17: **Fin Pour**

Algorithme 4.10 $\mathcal{R}_3$

```

1: Pour  $i \leftarrow 1$  à  $n$  faire
2:   Pour  $j \leftarrow 1$  à  $i - 1$  faire
3:      $U(i, j) \leftarrow 0$ 
4:   Fin Pour
5:   Pour  $j \leftarrow i$  à  $n$  faire
6:      $S_1 \leftarrow \sum_{k=1}^{i-1} L_{i,k} U_{k,j}$ 
7:      $U_{i,j} \leftarrow A_{i,j} - S_1$ 
8:   Fin Pour
9:   Pour  $j \leftarrow 1$  à  $i - 1$  faire
10:     $L_{j,i} \leftarrow 0$ 
11:  Fin Pour
12:   $L_{i,i} \leftarrow 1$ 
13:  Pour  $j \leftarrow i + 1$  à  $n$  faire
14:     $S_2 \leftarrow \sum_{k=1}^{i-1} L_{j,k} U_{k,i}$ 
15:     $L_{j,i} \leftarrow \frac{1}{U_{i,i}} (A_{j,i} - S_2)$ 
16:  Fin Pour
17: Fin Pour

```

Algorithme 4.10 $\mathcal{R}_4$

```

1: Pour  $i \leftarrow 1$  à  $n$  faire
2:   Pour  $j \leftarrow 1$  à  $i - 1$  faire
3:      $U(i, j) \leftarrow 0$ 
4:   Fin Pour
5:   Pour  $j \leftarrow i$  à  $n$  faire
6:      $S_1 \leftarrow 0$ 
7:     Pour  $k \leftarrow 1$  à  $i - 1$  faire
8:        $S_1 \leftarrow S_1 + L_{i,k} * U_{k,j}$ 
9:     Fin Pour
10:     $U_{i,j} \leftarrow A_{i,j} - S_1$ 
11:  Fin Pour
12:  Pour  $j \leftarrow 1$  à  $i - 1$  faire
13:     $L_{j,i} \leftarrow 0$ 
14:  Fin Pour
15:   $L_{i,i} \leftarrow 1$ 
16:  Pour  $j \leftarrow i + 1$  à  $n$  faire
17:     $S_2 \leftarrow 0$ 
18:    Pour  $k \leftarrow 1$  à  $i - 1$  faire
19:       $S_2 \leftarrow S_2 + L_{j,k} * U_{k,i}$ 
20:    Fin Pour
21:     $L_{j,i} \leftarrow \frac{1}{U_{i,i}} (A_{j,i} - S_2)$ 
22:  Fin Pour
23: Fin Pour

```

1
2

L'algorithme peut être amélioré, pour gagner en lisibilité... En effet, il est possible d'initialiser la matrice $\mathbb{U}$ par la matrice nulle et la matrice $\mathbb{L}$ par la matrice identité, ce qui permet alors de supprimer les boucles $U(i, j) \leftarrow 0$ et $L(j, i) \leftarrow 0$ ainsi que la commande $L(i, i) \leftarrow 1$. On obtient alors l'algorithme final

3
4
5
6

Algorithme 4.10 Fonction **FACTLU** permet de calculer les matrices $\mathbb{L}$ et $\mathbb{U}$ dites matrice de factorisation $\mathbb{LU}$ associée à la matrice $\mathbb{A}$, telle que

$$\mathbb{A} = \mathbb{L}\mathbb{U}$$

Données : $\mathbb{A}$: matrice de $\mathcal{M}_n(\mathbb{K})$ dont les sous-matrices principales sont inversibles.

Résultat : $\mathbb{L}$: matrice de $\mathcal{M}_n(\mathbb{K})$ triangulaire inférieure avec $L_{i,i} = 1, \forall i \in \llbracket 1, n \rrbracket$

$\mathbb{U}$: matrice de $\mathcal{M}_n(\mathbb{K})$ triangulaire supérieure.

```

1: Fonction  $[\mathbb{L}, \mathbb{U}] \leftarrow \mathbf{FACTLU}(\mathbb{A})$ 
2:    $\mathbb{U} \leftarrow \mathbb{O}_n$  ▷  $\mathbb{O}_n$  matrice nulle  $n \times n$ 
3:    $\mathbb{L} \leftarrow \mathbb{I}_n$  ▷  $\mathbb{I}_n$  matrice identité  $n \times n$ 
4:   Pour  $i \leftarrow 1$  à  $n$  faire ▷ Calcul de la ligne  $i$  de  $\mathbb{U}$ 
5:     Pour  $j \leftarrow i$  à  $n$  faire
6:        $S_1 \leftarrow 0$ 
7:       Pour  $k \leftarrow 1$  à  $i - 1$  faire
8:          $S_1 \leftarrow S_1 + L(i, k) * U(k, j)$ 
9:       Fin Pour
10:       $U(i, j) \leftarrow A(i, j) - S_1$ 
11:    Fin Pour
12:    Pour  $j \leftarrow i + 1$  à  $n$  faire ▷ Calcul de la colonne  $i$  de  $\mathbb{L}$ 
13:       $S_2 \leftarrow 0$ 
14:      Pour  $k \leftarrow 1$  à  $i - 1$  faire
15:         $S_2 \leftarrow S_2 + L(j, k) * U(k, i)$ 
16:      Fin Pour
17:       $L(j, i) \leftarrow (A_{j,i} - S_2) / U(i, i).$ 
18:    Fin Pour
19:  Fin Pour
20: Fin Fonction

```

Remarque 4.11 Pour optimiser en mémoire cette fonction, il est possible de stocker les matrices $\mathbb{L}$ et $\mathbb{U}$ dans une même matrice de $\mathcal{M}_n(\mathbb{R})$...

Pour faciliter la lecture d'un tel algorithme, il aurait pu être judicieux d'utiliser deux fonctions intermédiaires **FACTLULIGU** et **FACTLUCOLL** qui à l'étape i de l'algorithme calculent respectivement la ligne i de $\mathbb{U}$ et la colonne i de $\mathbb{L}$.

Algorithme 4.11 Fonction **FACTLULIGU** permet de calculer la ligne i de $\mathbb{U}$ à partir de (4.23)

Données : $\mathbb{A}$: matrice de $\mathcal{M}_n(\mathbb{K})$ dont les sous-matrices principales sont inversibles.
 $\mathbb{L}$: matrice de $\mathcal{M}_n(\mathbb{K})$ dont les $i - 1$ premières colonnes ont été calculées
 $\mathbb{U}$: matrice de $\mathcal{M}_n(\mathbb{K})$ dont les $i - 1$ premières lignes ont été calculées
Résultat : $\mathbb{U}$: matrice de $\mathcal{M}_n(\mathbb{K})$ dont les i premières lignes ont été calculées

```

1: Fonction  $\mathbb{U} \leftarrow \mathbf{FACTLULIGU}(\mathbb{A}, \mathbb{L}, \mathbb{U}, i)$ 
2:   Pour  $j \leftarrow 1$  à  $i - 1$  faire
3:      $U(i, j) \leftarrow 0$ 
4:   Fin Pour
5:   Pour  $j \leftarrow i$  à  $n$  faire
6:      $S_1 \leftarrow 0$ 
7:     Pour  $k \leftarrow 1$  à  $i - 1$  faire
8:        $S_1 \leftarrow S_1 + L(i, k) * U(k, j)$ 
9:     Fin Pour
10:     $U(i, j) \leftarrow A(i, j) - S_1$ 
11:   Fin Pour
12: Fin Fonction

```

Algorithme 4.12 Fonction **FACTLUCOLL** permet de calculer la colonne i de $\mathbb{L}$ à partir de (4.24)

Données : $\mathbb{A}$: matrice de $\mathcal{M}_n(\mathbb{K})$ dont les sous-matrices principales sont inversibles.
 $\mathbb{L}$: matrice de $\mathcal{M}_n(\mathbb{K})$ dont les $i - 1$ premières colonnes ont été calculées
 $\mathbb{U}$: matrice de $\mathcal{M}_n(\mathbb{K})$ dont les i premières lignes ont été calculées
Résultat : $\mathbb{L}$: matrice de $\mathcal{M}_n(\mathbb{K})$ dont les i premières colonnes ont été calculées

```

1: Fonction  $\mathbb{L} \leftarrow \mathbf{FACTLUCOLL}(\mathbb{A}, \mathbb{L}, \mathbb{U}, i)$ 
2:   Pour  $j \leftarrow 1$  à  $i - 1$  faire
3:      $L(j, i) \leftarrow 0$ 
4:   Fin Pour
5:    $L(i, i) \leftarrow 1$ 
6:   Pour  $j \leftarrow i + 1$  à  $n$  faire
7:      $S_2 \leftarrow 0$ 
8:     Pour  $k \leftarrow 1$  à  $i - 1$  faire
9:        $S_2 \leftarrow S_2 + L(j, k) * U(k, i)$ 
10:    Fin Pour
11:     $L(j, i) \leftarrow (A_{j,i} - S_2) / U(i, i).$ 
12:   Fin Pour
13: Fin Fonction

```

On a alors

Algorithme 4.13 Fonction **FACTLU** permet de calculer les matrices $\mathbb{L}$ et $\mathbb{U}$ dites matrice de factorisation $\mathbb{LU}$ associée à la matrice $\mathbb{A}$, telle que

$$\mathbb{A} = \mathbb{LU}$$

en utilisant des fonctions intermédiaires.

```

1: Fonction [ $\mathbb{L}, \mathbb{U}$ ]  $\leftarrow$  FACTLU(  $\mathbb{A}$  )
2:   Pour  $i \leftarrow 1$  à  $n$  faire
3:      $\mathbb{U} \leftarrow$  FACTLULIGU( $\mathbb{A}, \mathbb{L}, \mathbb{U}, i$ )
4:      $\mathbb{L} \leftarrow$  FACTLUCOLL( $\mathbb{A}, \mathbb{L}, \mathbb{U}, i$ )
5:   Fin Pour
6: Fin Fonction

```

4.2.4 Factorisation $\mathbb{LDL}^*$

Soit $\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$ une matrice hermitienne inversible admettant une factorisation $\mathbb{LU}$. On note $\mathbb{D}$ la matrice diagonale inversible définie par $\mathbb{D} = \text{diag } \mathbb{U}$ (i.e. $D_{i,i} = U_{i,i} \neq 0, \forall i \in \llbracket 1, n \rrbracket$) et $\mathbb{R} = \mathbb{D}^{-1}\mathbb{U}$. La matrice $\mathbb{R}$ est donc matrice triangulaire supérieure à diagonale unité car

$$R_{i,i} = (\mathbb{D}^{-1}\mathbb{U})_{i,i} = \sum_{k=1}^n (\mathbb{D}^{-1})_{i,k} U_{k,i} = (\mathbb{D}^{-1})_{i,i} U_{i,i} = \frac{1}{U_{i,i}} U_{i,i} = 1.$$

On a alors

$$\mathbb{A} = \mathbb{LU} = \mathbb{L}\mathbb{D}\mathbb{D}^{-1}\mathbb{U} = \mathbb{L}\mathbb{D}\mathbb{R}.$$

De plus comme $\mathbb{A}$ est hermitienne on a

$$\mathbb{A}^* = \mathbb{A} \iff \mathbb{R}^* \mathbb{D}^* \mathbb{L}^* = \mathbb{L}\mathbb{D}\mathbb{R}$$

Ce qui donne

$$\mathbb{A} = \mathbb{R}^* (\mathbb{D}^* \mathbb{L}^*) = \mathbb{L} (\mathbb{D}\mathbb{R})$$

et donc par unicité de la factorisation $\mathbb{LU}$ on a $\mathbb{R}^* = \mathbb{L}$ et $\mathbb{D}^* \mathbb{L}^* = \mathbb{D}\mathbb{R}$. On obtient alors

$$\mathbb{R}^* = \mathbb{L} \text{ et } \mathbb{D}^* = \mathbb{D}$$

et on a le théorème suivant



Théorème 4.12: Factorisation $\mathbb{LDL}^*$



Soit $\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$ une matrice hermitienne inversible admettant une factorisation $\mathbb{LU}$. Alors $\mathbb{A}$ s'écrit sous la forme

$$\mathbb{A} = \mathbb{LDL}^* \quad (4.25)$$

où $\mathbb{D} = \text{diag } \mathbb{U}$ est une matrice à coefficients réels.



Corollaire 4.13:



Une matrice $\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$ admet une factorisation $\mathbb{LDL}^*$ avec $\mathbb{L} \in \mathcal{M}_n(\mathbb{C})$ matrice triangulaire inférieure à diagonale unité et $\mathbb{D} \in \mathcal{M}_n(\mathbb{R})$ matrice diagonale à coefficients diagonaux strictement positifs **si et seulement si** la matrice $\mathbb{A}$ est hermitienne définie positive.

Proof. $\implies$ Soit $\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$ admettant une factorisation $\mathbb{LDL}^*$ avec $\mathbb{L} \in \mathcal{M}_n(\mathbb{C})$ matrice triangulaire inférieure à diagonale unité et $\mathbb{D} \in \mathcal{M}_n(\mathbb{R})$ matrice diagonale à coefficients diagonaux strictement positifs.

La matrice $\mathbb{A}$ est alors hermitienne car

$$\mathbb{A}^* = (\mathbb{LDL}^*)^* = \mathbb{L}^{**} \mathbb{D}^* \mathbb{L}^* = \mathbb{LDL}^*.$$

1 De plus $\forall \mathbf{x} \in \mathbb{C}^n \setminus \{0\}$ on a

$$\langle \mathbb{A} \mathbf{x}, \mathbf{x} \rangle = \langle \mathbb{L} \mathbb{D} \mathbb{L}^* \mathbf{x}, \mathbf{x} \rangle = \langle \mathbb{D} \mathbb{L}^* \mathbf{x}, \mathbb{L}^* \mathbf{x} \rangle$$

2 On pose $\mathbf{y} = \mathbb{L}^* \mathbf{x} \neq 0$ car $\mathbf{x} \neq 0$ et $\mathbb{L}^*$ inversible. On obtient alors

$$\langle \mathbb{A} \mathbf{x}, \mathbf{x} \rangle = \langle \mathbb{D} \mathbf{y}, \mathbf{y} \rangle = \sum_{i=1}^n D_{i,i} |y_i|^2 > 0$$

3 car $\mathbb{D}$ diagonale, $D_{i,i} > 0$, $\forall i \in \llbracket 1, n \rrbracket$ et $\mathbf{y} \neq 0$.

4 La matrice hermitienne $\mathbb{A}$ est donc bien définie positive.

5 $\Leftarrow$ Soit $\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$ une matrice hermitienne définie positive.

6 D'après le Corollaire 4.10, page 85, la matrice $\mathbb{A}$ admet une unique factorisation $\mathbb{L} \mathbb{U}$ et donc d'après
7 le Théorème 4.13, page 91, la matrice hermitienne $\mathbb{A}$ peut s'écrire sous la forme $\mathbb{A} = \mathbb{L} \mathbb{D} \mathbb{L}^*$ où $\mathbb{D}$ est
8 diagonale à coefficients réels et $\mathbb{L}$ triangulaire inférieure à diagonale unité. Il reste à démontrer que
9 $D_{i,i} > 0$, $\forall i \in \llbracket 1, n \rrbracket$.

10 Comme $\mathbb{A}$ est définie positive, on a $\forall \mathbf{x} \in \mathbb{C}^n \setminus \{0\}$, $\langle \mathbb{A} \mathbf{x}, \mathbf{x} \rangle > 0$. Or on a

$$\langle \mathbb{A} \mathbf{x}, \mathbf{x} \rangle = \langle \mathbb{L} \mathbb{D} \mathbb{L}^* \mathbf{x}, \mathbf{x} \rangle = \langle \mathbb{D} \mathbb{L}^* \mathbf{x}, \mathbb{L}^* \mathbf{x} \rangle$$

11 On note $\{\mathbf{e}_1, \dots, \mathbf{e}_n\}$, la base canonique de $\mathbb{C}^n$ et on rappelle que $\forall i \in \llbracket 1, n \rrbracket$, $\langle \mathbb{D} \mathbf{e}_i, \mathbf{e}_i \rangle = D_{i,i}$. Soit
12 $i \in \llbracket 1, n \rrbracket$. En choisissant $\mathbf{x} = (\mathbb{L}^*)^{-1} \mathbf{e}_i \neq 0$, on obtient alors

$$\langle \mathbb{D} \mathbb{L}^* \mathbf{x}, \mathbb{L}^* \mathbf{x} \rangle = \langle \mathbb{D} \mathbf{e}_i, \mathbf{e}_i \rangle = D_{i,i} > 0.$$

13

□

4.2.5 Factorisation de Cholesky

Definition 4.14

Une **factorisation régulière de Cholesky** d'une matrice $\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$ est une factorisation $\mathbb{A} = \mathbb{B} \mathbb{B}^*$ où $\mathbb{B}$ est une matrice triangulaire inférieure inversible.

Si les coefficients diagonaux de $\mathbb{B}$ sont positifs, on parle alors d'une **factorisation positive de Cholesky**.

15

Théorème 4.15: Factorisation de Cholesky



La matrice $\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$ admet une factorisation régulière de Cholesky **si et seulement si** la matrice $\mathbb{A}$ est hermitienne définie positive. Dans ce cas, elle admet une unique factorisation positive.

16

17 *Proof.* $\Rightarrow$ Soit $\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$ admettant une factorisation régulière de Cholesky $\mathbb{A} = \mathbb{B} \mathbb{B}^*$ avec $\mathbb{B}$ est une
18 matrice triangulaire inférieure inversible.

19 La matrice $\mathbb{A}$ est hermitienne car

$$\mathbb{A}^* = (\mathbb{B} \mathbb{B}^*)^* = (\mathbb{B}^*)^* \mathbb{B}^* = \mathbb{B} \mathbb{B}^* = \mathbb{A}.$$

20 Soit $\mathbf{x} \in \mathbb{C}^n \setminus \{0\}$, on a

$$\langle \mathbb{A} \mathbf{x}, \mathbf{x} \rangle = \langle \mathbb{B} \mathbb{B}^* \mathbf{x}, \mathbf{x} \rangle = \langle \mathbb{B}^* \mathbf{x}, \mathbb{B}^* \mathbf{x} \rangle = \|\mathbb{B}^* \mathbf{x}\|^2 > 0$$

21 car $\mathbb{B}^* \mathbf{x} \neq 0$ ($\mathbb{B}^*$ inversible et $\mathbf{x} \neq 0$). Donc la matrice $\mathbb{A}$ est bien hermitienne définie positive.

22 $\Leftarrow$ Soit $\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$ une matrice hermitienne définie positive.

23 D'après le Corollaire 4.13, page 91, il existe alors une matrice $\mathbb{L} \in \mathcal{M}_n(\mathbb{C})$ triangulaire inférieure à
24 diagonale unité et une matrice $\mathbb{D} \in \mathcal{M}_n(\mathbb{R})$ diagonale à coefficient strictement positifs telles que

$$\mathbb{A} = \mathbb{L} \mathbb{D} \mathbb{L}^*.$$

On note $\mathbb{H}\mathcal{M}_n(\mathbb{R})$ une matrice diagonale inversible vérifiant $\mathbb{H}^2 = \mathbb{D}$ (i.e. $H_{i,i} = \pm D_{i,i} \neq 0$, $\forall i \in \llbracket 1, n \rrbracket$). On a alors

$$\mathbb{A} = \mathbb{L}\mathbb{H}\mathbb{L}^* = (\mathbb{L}\mathbb{H})(\mathbb{L}\mathbb{H})^*$$

En posant $\mathbb{B} = \mathbb{L}\mathbb{H}$, la matrice $\mathbb{B}$ est bien triangulaire inférieure inversible car produit d'une matrice triangulaire inférieure inversible par une matrice diagonale inversible et on a $\mathbb{A} = \mathbb{B}\mathbb{B}^*$.

Montrons qu'une factorisation positive de Cholesky est unique.

Soient $\mathbb{B}_1$ et $\mathbb{B}_2$ deux factorisations positives de la matrice $\mathbb{A}$, on a donc

$$\mathbb{A} = \mathbb{B}_1\mathbb{B}_1^* = \mathbb{B}_2\mathbb{B}_2^*.$$

En multipliant à gauche par $\mathbb{B}_2^{-1}$ et à droite par $(\mathbb{B}_1^*)^{-1}$ cette équation on obtient

$$\mathbb{B}_2^{-1}\mathbb{B}_1 = \mathbb{B}_2^*(\mathbb{B}_1^*)^{-1} = \mathbb{B}_2^*(\mathbb{B}_1^{-1})^* = (\mathbb{B}_1^{-1}\mathbb{B}_2)^*$$

En notant $\mathbb{G} = \mathbb{B}_2^{-1}\mathbb{B}_1$, on tire de l'équation précédente

$$\mathbb{G} = (\mathbb{G}^{-1})^*. \quad (4.26)$$

On déduit de la (voir Proposition A.37, page 182), que l'inverse d'une matrice triangulaire inférieure à coefficients diagonaux réels strictement positifs est aussi une matrice triangulaire inférieure à coefficients diagonaux réels strictement positifs. De la (voir Proposition A.36, page 182), on obtient que le produit de matrices triangulaires inférieures à coefficients diagonaux réels strictement positifs reste triangulaire inférieure à coefficients diagonaux réels strictement positifs, on en déduit que les matrices $\mathbb{G} = \mathbb{B}_2^{-1}\mathbb{B}_1$ et $\mathbb{G}^{-1} = \mathbb{B}_1^{-1}\mathbb{B}_2$ sont triangulaires inférieures à coefficients diagonaux réels strictement positifs. Or l'équation (4.26) identifie la matrice triangulaire inférieure $\mathbb{G}$ à la matrice triangulaire supérieure $(\mathbb{G}^{-1})^*$: ce sont donc des matrices diagonales à coefficients diagonaux réels strictement positifs et on a alors $(\mathbb{G}^{-1})^* = \mathbb{G}^{-1}$. De l'équation (4.26), on obtient alors $\mathbb{G} = \mathbb{G}^{-1}$, c'est à dire $\mathbb{G} = \mathbb{I} = \mathbb{B}_2^{-1}\mathbb{B}_1$ et donc $\mathbb{B}_1 = \mathbb{B}_2$. $\square$

Résolution d'un système linéaire par la factorisation de Cholesky

Pour commencer, avec la factorisation de Cholesky point de salut sans matrice hermitienne définie positive!



N'utiliser la factorisation de Cholesky pour la résolution d'un système linéaire que si la matrice du système est hermitienne définie positive.

Soit $\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$ une matrice hermitienne définie positive et $\mathbf{b} \in \mathbb{C}^n$. Grâce au théorème 4.15, on obtient :

Trouver $\mathbf{x} \in \mathbb{C}^n$ tel que

$$\mathbb{A}\mathbf{x} = \mathbf{b}. \quad (4.27)$$

est équivalent à

Trouver $\mathbf{x} \in \mathbb{C}^n$ solution de

$$\mathbb{B}^*\mathbf{x} = \mathbf{y} \quad (4.28)$$

avec $\mathbb{B}$ la matrice de factorisation positive de Cholesky de la matrice $\mathbb{A}$ avec $\mathbf{y} \in \mathbb{C}^n$ solution de

$$\mathbb{B}\mathbf{y} = \mathbf{b}. \quad (4.29)$$

On est donc ramené à

Algorithme 4.14 Algorithme de base permettant de résoudre, par une factorisation de Cholesky positive, le système linéaire

$$\mathbb{A}\mathbf{x} = \mathbf{b}$$

où $\mathbb{A}$ une matrice de $\mathcal{M}_n(\mathbb{C})$ hermitienne définie positive et $\mathbf{b} \in \mathbb{C}^n$.

Données : $\mathbb{A}$: matrice de $\mathcal{M}_n(\mathbb{C})$ hermitienne définie positive,
 $\mathbf{b}$: vecteur de $\mathbb{C}^n$.

Résultat : $\mathbf{x}$: vecteur de $\mathbb{C}^n$.

- 1: Trouver la factorisation positive de Cholesky $\mathbb{B}$ de la matrice $\mathbb{A}$,
- 2: résoudre le système triangulaire inférieur $\mathbb{B}\mathbf{y} = \mathbf{b}$,
- 3: résoudre le système triangulaire supérieur $\mathbb{B}^*\mathbf{x} = \mathbf{y}$.

- 1 Ceci permet donc de découper le problème initial en trois sous-problèmes plus simples. De plus,
- 2 ceux-ci peuvent se traiter de manière indépendante.

Algorithme 4.15 Fonction **RSLCHOLESKY** permettant de résoudre, par une factorisation de Cholesky positive, le système linéaire

$$\mathbb{A}\mathbf{x} = \mathbf{b}$$

où $\mathbb{A}$ une matrice hermitienne de $\mathcal{M}_n(\mathbb{C})$ définie positive et $\mathbf{b} \in \mathbb{C}^n$.

Données : $\mathbb{A}$: matrice de $\mathcal{M}_n(\mathbb{C})$ symétrique définie positive,
 $\mathbf{b}$: vecteur de $\mathbb{C}^n$.

Résultat : $\mathbf{x}$: vecteur de $\mathbb{C}^n$.

- 1: **Fonction** $\mathbf{x} \leftarrow \text{RSLCHOLESKY}(\mathbb{A}, \mathbf{b})$
- 2: $\mathbb{B} \leftarrow \text{CHOLESKY}(\mathbb{A})$ ▷ Factorisation positive de Cholesky
- 3: $\mathbf{y} \leftarrow \text{RSLTRIINF}(\mathbb{B}, \mathbf{b})$ ▷ Résolution du système $\mathbb{B}\mathbf{y} = \mathbf{b}$
- 4: $\mathbb{U} \leftarrow \text{MATADJOINTE}(\mathbb{B})$ ▷ Calcul de la matrice adjointe de $\mathbb{B}$
- 5: $\mathbf{x} \leftarrow \text{RSLTRISUP}(\mathbb{U}, \mathbf{y})$ ▷ Résolution du système $\mathbb{B}^*\mathbf{x} = \mathbf{y}$
- 6: **Fin Fonction**

- 3 Il nous faut donc écrire la fonction **CHOLESKY** (les deux fonctions **RSLTRIINF** et **RSLTRISUP** ayant
- 4 déjà été écrites et la fonction **TRANSPOSE** étant simple à écrire).

5 Factorisation positive de Cholesky : écriture de l'algorithme

- 6 Soit $\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$ une matrice hermitienne définie positive. D'après le Théorème 4.15, il existe une unique
- 7 matrice $\mathbb{B} \in \mathcal{M}_n(\mathbb{C})$ triangulaire inférieure avec $B_{i,i} \in \mathbb{R}^{+*}$, $\forall i \in \llbracket 1, n \rrbracket$, telle que

$$\mathbb{A} = \mathbb{B}\mathbb{B}^* \quad (4.30)$$

- 8 c'est à dire

$$\begin{pmatrix} A_{1,1} & \dots & A_{1,n} \\ \vdots & \ddots & \vdots \\ A_{n,1} & \dots & A_{n,n} \end{pmatrix} = \begin{pmatrix} B_{1,1} & 0 & \dots & 0 \\ \vdots & \ddots & \ddots & \vdots \\ \vdots & & \ddots & 0 \\ B_{n,1} & \dots & \dots & B_{n,n} \end{pmatrix} \begin{pmatrix} \overline{B_{1,1}} & \dots & \dots & \overline{B_{n,1}} \\ 0 & \ddots & & \vdots \\ \vdots & \ddots & \ddots & \vdots \\ 0 & \dots & 0 & \overline{B_{n,n}} \end{pmatrix}. \quad (4.31)$$

- 9 Pour déterminer la matrice $\mathbb{B}$, on commence par calculer $B_{1,1}$ (la 1ère ligne de $\mathbb{B}$ est donc déterminée)
- 10 ce qui nous permet de calculer la 1ère colonne de $\mathbb{B}$.
- 11 Ensuite, on calcule $B_{2,2}$ (la 2ème ligne de $\mathbb{B}$ est donc déterminée) ce qui nous permet de calculer la 2ème
- 12 colonne de $\mathbb{B}$. Etc ...

- 13 On écrit en détail les raffinements successifs permettant d'aboutir à l'algorithme final de telle manière
- 14 que le passage entre deux raffinements successifs soit le plus simple possible.

15

Algorithme 4.16 $\mathcal{R}_0$ 1: Calculer la matrice $\mathbb{B}$ **Algorithme 4.16** $\mathcal{R}_1$

1: **Pour** $i \leftarrow 1$ à n **faire**
 2: Calculer $B_{i,i}$, connaissant les $i - 1$ premières colonnes de $\mathbb{B}$.
 3: Calculer la $i^{\text{ème}}$ colonne de $\mathbb{B}$.
 4: **Fin Pour**

Pour calculer $B_{i,i}$, connaissant les $i - 1$ premières colonnes de $\mathbb{B}$, on utilise (4.30) :

$$A_{i,i} = \sum_{j=1}^n B_{i,j}(\mathbb{B}^*)_{j,i} = \sum_{j=1}^n |B_{i,j}|^2,$$

et comme $\mathbb{B}$ est triangulaire inférieure on obtient

$$A_{i,i} = \sum_{j=1}^i |B_{i,j}|^2 = \sum_{j=1}^{i-1} |B_{i,j}|^2 + |B_{i,i}|^2.$$

On a donc

$$B_{i,i} = \left(A_{i,i} - \sum_{j=1}^{i-1} |B_{i,j}|^2 \right)^{1/2}, \quad \forall i \in \llbracket 1, n \rrbracket, \quad (4.32)$$

Comme les $i - 1$ premières colonnes de $\mathbb{B}$ ont déjà été calculées, $B_{i,i}$ est parfaitement déterminée par la formule précédente.

Remarque 4.16 Les hypothèses sur la matrice $\mathbb{A}$ permettent d'affirmer que, $\forall i \in \llbracket 1, n \rrbracket$, $A_{i,i} - \sum_{j=1}^{i-1} |B_{i,j}|^2 > 0$ et donc $B_{i,i} > 0$.

Pour calculer la $i^{\text{ème}}$ colonne de $\mathbb{B}$, il suffit de déterminer $B_{j,i}$, $\forall j \in \llbracket i + 1, n \rrbracket$. Pour cela on utilise (4.30)

$$A_{j,i} = \sum_{k=1}^n B_{j,k}(\mathbb{B}^*)_{k,i} = \sum_{k=1}^n B_{j,k} \overline{B_{i,k}}, \quad \forall j \in \llbracket i + 1, n \rrbracket$$

Comme $\mathbb{L}$ est triangulaire inférieure on obtient

$$A_{j,i} = \sum_{k=1}^i B_{j,k} \overline{B_{i,k}} = \sum_{k=1}^{i-1} B_{j,k} \overline{B_{i,k}} + B_{j,i} \overline{B_{i,i}}, \quad \forall j \in \llbracket i + 1, n \rrbracket$$

Or $B_{i,i} > 0$ vient d'être calculé et les $i - 1$ premières colonnes de $\mathbb{B}$ ont déjà été calculées, ce qui donne

$$B_{j,i} = \frac{1}{B_{i,i}} \left(A_{j,i} - \sum_{k=1}^{i-1} B_{j,k} \overline{B_{i,k}} \right), \quad \forall j \in \llbracket i + 1, n \rrbracket \quad (4.33)$$

$$B_{j,i} = 0, \quad \forall j \in \llbracket 1, i - 1 \rrbracket. \quad (4.34)$$

Avec (4.32) et (4.33), on a

Algorithme 4.16 $\mathcal{R}_1$

1: **Pour** $i \leftarrow 1$ à n **faire**
 2: Calculer $B_{i,i}$, connaissant les $i - 1$ premières colonnes de $\mathbb{B}$.
 3: Calculer la $i^{\text{ème}}$ colonne de $\mathbb{B}$.
 4: **Fin Pour**

Algorithme 4.16 $\mathcal{R}_2$

1: **Pour** $i \leftarrow 1$ à n **faire**
 2: $B_{i,i} \leftarrow \left(A_{i,i} - \sum_{j=1}^{i-1} |B_{i,j}|^2 \right)^{1/2}$
 3: **Pour** $j \leftarrow 1$ à $i - 1$ **faire**
 4: $B_{j,i} \leftarrow 0$
 5: **Fin Pour**
 6: **Pour** $j \leftarrow i + 1$ à n **faire**
 7: $B_{j,i} \leftarrow \frac{1}{B_{i,i}} \left(A_{j,i} - \sum_{k=1}^{i-1} B_{j,k} \overline{B_{i,k}} \right)$
 8: **Fin Pour**
 9: **Fin Pour**

1

Algorithme 4.16 $\overline{\mathcal{R}}_2$

```

1: Pour  $i \leftarrow 1$  à  $n$  faire
2:    $B_{i,i} \leftarrow \left( A_{i,i} - \sum_{j=1}^{i-1} |B_{i,j}|^2 \right)^{1/2}$ 
3:   Pour  $j \leftarrow 1$  à  $i-1$  faire
4:      $B_{j,i} \leftarrow 0$ 
5:   Fin Pour
6:   Pour  $j \leftarrow i+1$  à  $n$  faire
7:      $B_{j,i} \leftarrow \frac{1}{B_{i,i}} \left( A_{j,i} - \sum_{k=1}^{i-1} B_{j,k} \overline{B_{i,k}} \right)$ 
8:   Fin Pour
9: Fin Pour

```

Algorithme 4.16 $\overline{\mathcal{R}}_3$

```

1: Pour  $i \leftarrow 1$  à  $n$  faire
2:    $S_1 \leftarrow \sum_{j=1}^{i-1} |B_{i,j}|^2$ 
3:    $B_{i,i} \leftarrow (A_{i,i} - S_1)^{1/2}$ 
4:   Pour  $j \leftarrow 1$  à  $i-1$  faire
5:      $B_{j,i} \leftarrow 0$ 
6:   Fin Pour
7:   Pour  $j \leftarrow i+1$  à  $n$  faire
8:      $S_2 \leftarrow \sum_{k=1}^{i-1} B_{j,k} \overline{B_{i,k}}$ 
9:      $B_{j,i} \leftarrow \frac{1}{B_{i,i}} (A_{j,i} - S_2)$ 
10:  Fin Pour
11: Fin Pour

```

2

3

Algorithme 4.16 $\overline{\mathcal{R}}_3$

```

1: Pour  $i \leftarrow 1$  à  $n$  faire
2:    $S_1 \leftarrow \sum_{j=1}^{i-1} |B_{i,j}|^2$ 
3:    $B_{i,i} \leftarrow (A_{i,i} - S_1)^{1/2}$ 
4:   Pour  $j \leftarrow 1$  à  $i-1$  faire
5:      $B_{j,i} \leftarrow 0$ 
6:   Fin Pour
7:   Pour  $j \leftarrow i+1$  à  $n$  faire
8:      $S_2 \leftarrow \sum_{k=1}^{i-1} B_{j,k} \overline{B_{i,k}}$ 
9:      $B_{j,i} \leftarrow \frac{1}{B_{i,i}} (A_{j,i} - S_2)$ 
10:  Fin Pour
11: Fin Pour

```

Algorithme 4.16 $\overline{\mathcal{R}}_4$

```

1: Pour  $i \leftarrow 1$  à  $n$  faire
2:    $S_1 \leftarrow 0$ 
3:   Pour  $j \leftarrow 1$  à  $i-1$  faire
4:      $S_1 \leftarrow S_1 + |B_{i,j}|^2$ 
5:   Fin Pour
6:    $B_{i,i} \leftarrow (A_{i,i} - S_1)^{1/2}$ 
7:   Pour  $j \leftarrow 1$  à  $i-1$  faire
8:      $B_{j,i} \leftarrow 0$ 
9:   Fin Pour
10:  Pour  $j \leftarrow i+1$  à  $n$  faire
11:     $S_2 \leftarrow 0$ 
12:    Pour  $k \leftarrow 1$  à  $i-1$  faire
13:       $S_2 \leftarrow S_2 + B_{j,k} \overline{B_{i,k}}$ 
14:    Fin Pour
15:     $B_{j,i} \leftarrow \frac{1}{B_{i,i}} (A_{j,i} - S_2)$ 
16:  Fin Pour
17: Fin Pour

```

4

5

6 On obtient alors l'algorithme final

Algorithme 4.16 Fonction **CHOLESKY** permettant de calculer la matrice $\mathbb{B}$, dite matrice de factorisation positive de Cholesky associée à la matrice $\mathbb{A}$, telle que

$$\mathbb{A} = \mathbb{B}\mathbb{B}^*.$$

Données : $\mathbb{A}$: matrice de $\mathcal{M}_n(\mathbb{C})$ hermitienne définie positive.
Résultat : $\mathbb{B}$: matrice de $\mathcal{M}_n(\mathbb{C})$ triangulaire inférieure
 avec $\mathbb{B}(i, i) > 0, \forall i \in \llbracket 1, n \rrbracket$

```

1: Fonction  $\mathbb{B} \leftarrow \text{CHOLESKY}(\mathbb{A})$ 
2:   Pour  $i \leftarrow 1$  à  $n$  faire
3:      $S_1 \leftarrow 0$ 
4:     Pour  $j \leftarrow 1$  à  $i - 1$  faire
5:        $S_1 \leftarrow S_1 + |\mathbb{B}(i, j)|^2$ 
6:     Fin Pour
7:      $\mathbb{B}(i, i) \leftarrow \text{SQRT}(\mathbb{A}(i, i) - S_1)$ 
8:     Pour  $j \leftarrow 1$  à  $i - 1$  faire
9:        $\mathbb{B}(j, i) \leftarrow 0$ 
10:    Fin Pour
11:    Pour  $j \leftarrow i + 1$  à  $n$  faire
12:       $S_2 \leftarrow 0$ 
13:      Pour  $k \leftarrow 1$  à  $i - 1$  faire
14:         $S_2 \leftarrow S_2 + \mathbb{B}(j, k) * \overline{\mathbb{B}(i, k)}$ 
15:      Fin Pour
16:       $\mathbb{B}(j, i) \leftarrow (\mathbb{A}(j, i) - S_2) / \mathbb{B}(i, i).$ 
17:    Fin Pour
18:  Fin Pour
19: Fin Fonction

```

4.2.6 Factorisation QR

La tranformation de Householder

Definition 4.17: Matrice élémentaire de Householder

Soit $\mathbf{u} \in \mathbb{C}^n$ tel que $\|\mathbf{u}\|_2 = 1$. On appelle **matrice élémentaire de Householder** la matrice $\mathbb{H}(\mathbf{u}) \in \mathcal{M}_n(\mathbb{C})$ définie par

$$\mathbb{H}(\mathbf{u}) = \mathbb{I} - 2\mathbf{u}\mathbf{u}^*. \quad (4.35)$$

Propriété 4.18

Toute matrice élémentaire de Householder est hermitienne et unitaire.

Proof. Pour simplifier, on note $\mathbb{H} = \mathbb{H}(\mathbf{u})$. Cette matrice est hermitienne car

$$\mathbb{H}^* = (\mathbb{I} - 2\mathbf{u}\mathbf{u}^*)^* = \mathbb{I} - 2(\mathbf{u}\mathbf{u}^*)^* = \mathbb{I} - 2\mathbf{u}\mathbf{u}^* = \mathbb{H}.$$

Montrons qu'elle est unitaire (i.e. $\mathbb{H}^*\mathbb{H} = \mathbb{I}$). On a

$$\begin{aligned} \mathbb{H}^*\mathbb{H} &= \mathbb{H}\mathbb{H} = (\mathbb{I} - 2\mathbf{u}\mathbf{u}^*)(\mathbb{I} - 2\mathbf{u}\mathbf{u}^*) \\ &= \mathbb{I} - 4\mathbf{u}\mathbf{u}^* + 4\mathbf{u}\mathbf{u}^*\mathbf{u}\mathbf{u}^*. \end{aligned}$$

Or on a $\mathbf{u}^*\mathbf{u} = \|\mathbf{u}\|_2^2 = 1$ par hypothèse et donc

$$\mathbb{H}^*\mathbb{H} = \mathbb{I} - 4\mathbf{u}\mathbf{u}^* + 4\mathbf{u}(\mathbf{u}^*\mathbf{u})\mathbf{u}^* = \mathbb{I}.$$

□

**Propriété 4.19**

Soient $\mathbf{x} \in \mathbb{K}^n$ et $\mathbf{u} \in \mathbb{K}^n$, $\|\mathbf{u}\|_2 = 1$. On note $\mathbf{x}_{\parallel} = \text{proj}_{\mathbf{u}}(\mathbf{x}) \stackrel{\text{def}}{=} \langle \mathbf{u}, \mathbf{x} \rangle \mathbf{u}$ et $\mathbf{x}_{\perp} = \mathbf{x} - \mathbf{x}_{\parallel}$. On a alors

$$\mathbb{H}(\mathbf{u})(\mathbf{x}_{\perp} + \mathbf{x}_{\parallel}) = \mathbf{x}_{\perp} - \mathbf{x}_{\parallel}. \quad (4.36)$$

et

$$\mathbb{H}(\mathbf{u})\mathbf{x} = \mathbf{x}, \quad \text{si } \langle \mathbf{x}, \mathbf{u} \rangle = 0. \quad (4.37)$$

Proof. On note que par construction $\langle \mathbf{u}, \mathbf{x}_{\perp} \rangle = 0$. On a

$$\begin{aligned} \mathbb{H}(\mathbf{u})(\mathbf{x}_{\perp} + \mathbf{x}_{\parallel}) &= (\mathbb{I} - 2\mathbf{u}\mathbf{u}^*)(\mathbf{x}_{\perp} + \mathbf{x}_{\parallel}) = \mathbf{x}_{\perp} + \mathbf{x}_{\parallel} - 2\mathbf{u} \underbrace{\mathbf{u}^*\mathbf{x}_{\perp}}_{=0} - 2\mathbf{u}\mathbf{u}^*\mathbf{x}_{\parallel} \\ &= \mathbf{x}_{\perp} + \mathbf{x}_{\parallel} - 2\mathbf{u}\mathbf{u}^*\mathbf{u} \langle \mathbf{u}, \mathbf{x} \rangle = \mathbf{x}_{\perp} + \mathbf{x}_{\parallel} - 2\mathbf{u} \underbrace{\mathbf{u}^*\mathbf{u}}_{=1} \mathbf{u}^*\mathbf{x} \\ &= \mathbf{x}_{\perp} + \mathbf{x}_{\parallel} - 2\mathbf{u}\mathbf{u}^*\mathbf{x} = \mathbf{x}_{\perp} + \mathbf{x}_{\parallel} - 2\mathbf{x}_{\parallel} \\ &= \mathbf{x}_{\perp} - \mathbf{x}_{\parallel}. \end{aligned}$$

Si $\langle \mathbf{x}, \mathbf{u} \rangle = 0$ alors $\mathbf{x}_{\parallel} = 0$ et $\mathbf{x} = \mathbf{x}_{\perp}$. □

**Théorème 4.20**

Soient $\mathbf{a}, \mathbf{b}$ deux vecteurs non colinéaires de $\mathbb{C}^n$ avec $\|\mathbf{b}\|_2 = 1$. Soit $\alpha \in \mathbb{C}$ tel que $|\alpha| = \|\mathbf{a}\|_2$ et $\arg \alpha = -\arg \langle \mathbf{a}, \mathbf{b} \rangle \in [\pi]$. On a alors

$$\mathbb{H}\left(\frac{\mathbf{a} - \alpha\mathbf{b}}{\|\mathbf{a} - \alpha\mathbf{b}\|_2}\right)\mathbf{a} = \alpha\mathbf{b}. \quad (4.38)$$

Proof. (voir Exercice 4.2.5, page 98) □

**Exercice 4.2.5**

Soient $\mathbf{a}$ et $\mathbf{b}$ deux vecteurs non colinéaires de $\mathbb{C}^n$ avec $\|\mathbf{b}\|_2 = 1$. On va chercher $\alpha \in \mathbb{C}$ et $\mathbf{u} \in \mathbb{C}^n$ vérifiant

$$\mathbb{H}(\mathbf{u})\mathbf{a} = \alpha\mathbf{b}. \quad (4.39)$$

Q. 1 1. Montrer que si α vérifie (4.39) alors $|\alpha| = \|\mathbf{a}\|_2$.

2. Montrer que si $\arg \alpha = -\arg \langle \mathbf{a}, \mathbf{b} \rangle \in [\pi]$ alors $\alpha \langle \mathbf{a}, \mathbf{b} \rangle \in \mathbb{R}$.

Q. 2 Soient α et $\mathbf{u}$ vérifiant (4.39).

1. Montrer que

$$|\langle \mathbf{u}, \mathbf{a} \rangle|^2 = \frac{\langle \mathbf{a}, \mathbf{a} \rangle - \alpha \langle \mathbf{a}, \mathbf{b} \rangle}{2} \quad (4.40)$$

2. Montrer que si $\arg \alpha = -\arg \langle \mathbf{a}, \mathbf{b} \rangle \in [\pi]$ alors $\langle \mathbf{a}, \mathbf{a} \rangle - \alpha \langle \mathbf{a}, \mathbf{b} \rangle \in \mathbb{R}^{*+}$.

3. En déduire que

$$\mathbf{u} = \frac{1}{2\lambda}(\mathbf{a} - \alpha\mathbf{b}), \quad \text{avec } \lambda = \pm \left(\frac{\langle \mathbf{a}, \mathbf{a} \rangle - \alpha \langle \mathbf{a}, \mathbf{b} \rangle}{2} \right)^{1/2} \quad (4.41)$$

Correction Exercice 4.2.5 On pose $\mathbb{H} = \mathbb{H}(\mathbf{u})$ pour alléger les notations.

Q. 1 1. On a

$$\begin{aligned}\|a\|_2^2 &= \langle a, a \rangle = \langle H^* H a, a \rangle \quad \text{car } H \text{ unitaire} \\ &= \langle H a, H a \rangle \quad \text{par définition du produit scalaire} \\ &= \|H a\|_2^2 = \|\alpha b\|_2^2 = |\alpha|^2 \|b\|_2^2 = |\alpha|^2.\end{aligned}$$

2. On a par définition de l'argument $\alpha = |\alpha| e^{i \arg \alpha}$ et $\langle a, b \rangle = |\langle a, b \rangle| e^{i \arg(\langle a, b \rangle)}$ ce qui donne 1

$$\alpha \langle a, b \rangle = |\alpha| |\langle a, b \rangle| e^{i(\arg \alpha + \arg(\langle a, b \rangle))} \quad (4.42)$$

et donc $\alpha \langle a, b \rangle$ est réel si $\arg \alpha + \arg(\langle a, b \rangle) = 0 \pmod{\pi}$. 2

Q. 2 1. On a

$$\begin{aligned}H(u)a = \alpha b &\iff (I - 2uu^*)a = \alpha b \\ &\iff a - 2u(u^*a) = \alpha b\end{aligned}$$

et donc 3

$$a - 2\langle u, a \rangle u = \alpha b \quad (4.43)$$

En effectuant le produit scalaire avec a de cette dernière équation, on obtient 4

$$\langle a, a \rangle - 2\langle u, a \rangle \langle a, u \rangle = \alpha \langle a, b \rangle$$

ce qui prouve (4.40). 5

2. On a montré en Q.1 que $\alpha \langle a, b \rangle \in \mathbb{R}$ et donc $\langle a, a \rangle - \alpha \langle a, b \rangle \in \mathbb{R}$. Il reste donc à montrer que $\langle a, a \rangle - \alpha \langle a, b \rangle > 0$. 6

- Si $\arg \alpha = -\arg(\langle a, b \rangle) + \pi \pmod{2\pi}$, alors de (4.42) on obtient $\alpha \langle a, b \rangle \leq 0$ et donc $\langle a, a \rangle - \alpha \langle a, b \rangle \geq \|a\|_2^2 > 0$ car $a \neq 0$. 8
 - Si $\arg \alpha = -\arg(\langle a, b \rangle) \pmod{2\pi}$, alors de (4.42) on obtient $\alpha \langle a, b \rangle \geq 0$. 10
- Comme les vecteurs a et b ne sont pas colinéaires on a inégalité stricte dans Cauchy-Schwarz : 11

$$|\langle a, b \rangle| < \|a\|_2 \|b\|_2 = \|a\|_2.$$

On obtient donc 12

$$0 \leq \alpha \langle a, b \rangle \leq |\alpha| |\langle a, b \rangle| < |\alpha| \|a\|_2 = \|a\|_2^2$$

Attention, dans ce cas $\langle a, a \rangle - \alpha \langle a, b \rangle$ peut-être très petit. 13

3. De (4.43), on en déduit immédiatement (4.41). 14

Vérifions que $\|u\|_2 = 1$. On a 15

$$\|u\|_2^2 = \langle u, u \rangle = \frac{1}{4|\lambda|^2} \langle a - \alpha b, a - \alpha b \rangle$$

Or

$$\begin{aligned}\langle a - \alpha b, a - \alpha b \rangle &= \langle a, a \rangle - \bar{\alpha} \langle b, a \rangle - \alpha \langle a, b \rangle + |\alpha|^2 \langle b, b \rangle = \|a\|_2^2 - \bar{\alpha} \langle b, a \rangle - \alpha \langle a, b \rangle + |\alpha|^2 \\ &= 2\|a\|_2^2 - \bar{\alpha} \langle b, a \rangle - \alpha \langle a, b \rangle \\ &= 2\|a\|_2^2 - 2\alpha \langle a, b \rangle \quad \text{car } \alpha \langle a, b \rangle = \overline{(\alpha \langle a, b \rangle)} = \bar{\alpha} \langle b, a \rangle \in \mathbb{R}\end{aligned}$$

De plus

$$\begin{aligned}4|\lambda|^2 &= 2(\langle a, a \rangle - \alpha \langle b, a \rangle) \in \mathbb{R} \\ &= 2\|a\|_2^2 - 2\alpha \langle b, a \rangle\end{aligned}$$

◇ 16
17



Exercice 4.2.6

Soient $\mathbf{a}$ et $\mathbf{b}$ deux vecteurs non nuls et non colinéaires de $\mathbb{C}^n$ avec $\|\mathbf{b}\|_2 = 1$.

Q. 1 Ecrire la fonction algorithmique **HOUSEHOLDER** permettant de retourner une matrice de Householder $\mathbb{H}$ et $\alpha \in \mathbb{C}$ tels que $\mathbb{H}(\mathbf{u})\mathbf{a} = \alpha\mathbf{b}$. Le choix du α est fait par le paramètre δ (0 ou 1) de telle sorte que $\arg \alpha = -\arg(\langle \mathbf{a}, \mathbf{b} \rangle) + \delta\pi$ avec $|\alpha| = \|\mathbf{a}\|_2$.

Des fonctions comme **DOT**($\mathbf{a}, \mathbf{b}$) (produit scalaire de deux vecteurs), **NORM**($\mathbf{a}$) (norme d'un vecteur), **ARG**(z) (argument d'un nombre complexe), **MATPROD**($\mathbb{A}, \mathbb{B}$) (produit de deux matrices), **CTRANSPOSE**($\mathbb{A}$) (adjoint d'une matrice), ... pourront être utilisées

Q. 2 Proposer un programme permettant de tester cette fonction. On pourra utiliser la fonction **VECRAND**(n) retournant un vecteur aléatoire de $\mathbb{C}^n$, les parties réelles et imaginaires de chacune de ses composantes étant dans $]0, 1[$ (loi uniforme).

Q. 3 Proposer un programme permettant de vérifier que $\delta = 1$ est le "meilleur" choix.

1

2 **Correction Exercice 4.2.6** Soient $\mathbf{a}$ et $\mathbf{b}$ deux vecteurs non nuls et non colinéaires de $\mathbb{C}^n$.

3

4 **Q. 1** Les données du problème sont $\mathbf{a}$, $\mathbf{b}$ et δ . On veut calculer α et la matrice $\mathbb{H}(\mathbf{u})$.

Algorithme 4.17 Calcul du α et de la matrice de Householder $\mathbb{H}(\mathbf{u})$ telle que $\mathbb{H}(\mathbf{u})\mathbf{a} = \alpha\mathbf{b}$.

Données : $\mathbf{a}, \mathbf{b}$: deux vecteurs de $\mathbb{C}^n$ non nuls et non colinéaires.

δ : 0 ou 1, permet de déterminer α .

Résultat : $\mathbb{H}$: matrice de Householder dans $\mathcal{M}_n(\mathbb{C})$,

α : nombre complexe, donné par $-\arg(\langle \mathbf{a}, \mathbf{b} \rangle) + \delta\pi$.

1: **Fonction** $[\mathbb{H}, \alpha] \leftarrow \mathbf{HOUSEHOLDER}(\mathbf{a}, \mathbf{b}, \delta)$

2: $\mathbf{ab} \leftarrow \mathbf{DOT}(\mathbf{A}, \mathbf{B})$

$\triangleright$ **DOT** produit scalaire dans $\mathbb{C}$.

3: $\alpha \leftarrow \mathbf{NORM}(\mathbf{a}, 2) * \exp(i * (\delta * \pi - \mathbf{ARG}(\mathbf{ab})))$

4: $\mathbf{u} \leftarrow \mathbf{a} - \alpha * \mathbf{b}$

5: $\mathbf{u} \leftarrow \mathbf{u} / \mathbf{NORM}(\mathbf{u}, 2)$

6: $\mathbb{H} \leftarrow \mathbf{EYE}(n) - 2 * \mathbf{MATPROD}(\mathbf{u}, \mathbf{CTRANSPOSE}(\mathbf{u}))$

7: **Fin Fonction**

5 **Q. 2** 1: $n \leftarrow 100$

6 2: $\mathbf{a} \leftarrow \mathbf{VECRAND}(n)$

7 3: $\mathbf{b} \leftarrow \mathbf{VECRAND}(n)$

8 4: $\mathbf{b} \leftarrow \mathbf{NORM}(\mathbf{b}, 2)$

9 5: $[\mathbb{H}, \alpha] \leftarrow \mathbf{HOUSEHOLDER}(\mathbf{a}, \mathbf{b}, 0)$

10 6: $\mathbf{error} \leftarrow \mathbf{NORM}(\mathbb{H} * \mathbf{a} - \alpha * \mathbf{b}, 2)$

11 **Q. 3** 1: $n \leftarrow 100$

12 2: $\mathbf{a} \leftarrow \mathbf{VECRAND}(n)$

13 3: $\mathbf{b} \leftarrow \mathbf{a} + 1e - 6 * \mathbf{VECRAND}(n)$

14 4: $\mathbf{b} \leftarrow \mathbf{b} / \mathbf{NORM}(\mathbf{b}, 2)$

15 5: $[\mathbb{H}_1, \alpha_1] \leftarrow \mathbf{HOUSEHOLDER}(\mathbf{a}, \mathbf{b}, 1)$

16 6: $[\mathbb{H}_0, \alpha_0] \leftarrow \mathbf{HOUSEHOLDER}(\mathbf{a}, \mathbf{b}, 0)$

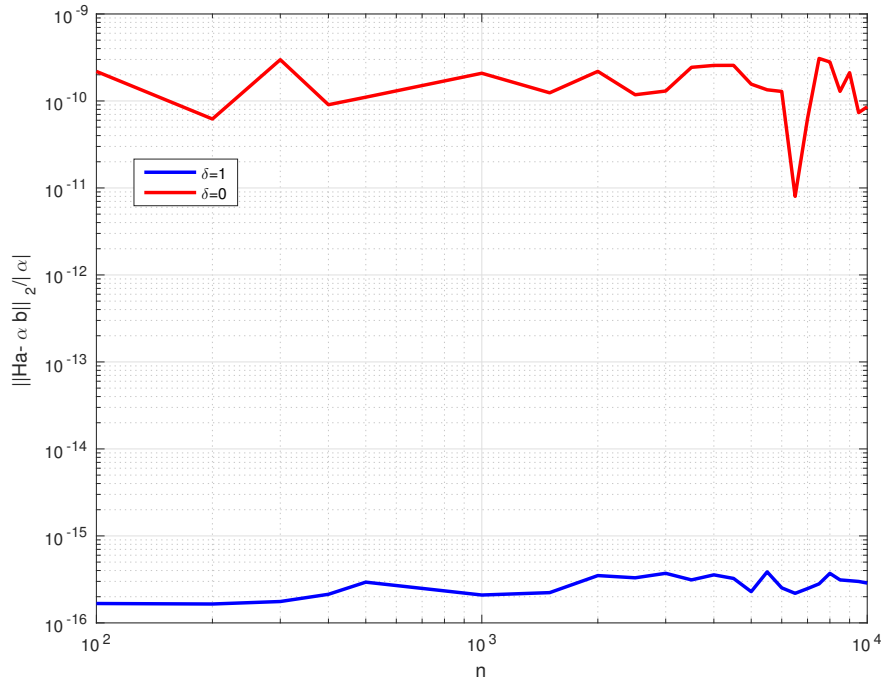
17 7: $\mathbf{error0} \leftarrow \mathbf{NORM}(\mathbb{H}_0 * \mathbf{a} - \alpha_0 * \mathbf{b}, 2) / \mathbf{ABS}(\alpha_0)$

18 8: $\mathbf{error1} \leftarrow \mathbf{NORM}(\mathbb{H}_1 * \mathbf{a} - \alpha_1 * \mathbf{b}, 2) / \mathbf{ABS}(\alpha_1)$

19

20

21 Si l'on souhaite calculer le produit matrice-vecteur $\mathbb{H}(\mathbf{u})\mathbf{x}$, il n'est pas nécessaire de calculer

Figure 4.1: Choix de α dans **HOUSEHOLDER** : erreur relative en norme L_2

explicitement la matrice $\mathbb{H}(\mathbf{u})$. En effet, on pose $\mathbf{v} = 2\lambda\mathbf{u} = \mathbf{a} - \alpha\mathbf{b}$ et $\beta = 2\lambda^2 = \|\mathbf{a}\|_2^2 - \alpha\langle\mathbf{a}, \mathbf{b}\rangle > 0$. On choisit α de manière à maximiser β pour éviter une division par un nombre trop petit. On prend donc

$$\alpha = \|\mathbf{a}\|_2 e^{i(\pi - \arg(\langle\mathbf{a}, \mathbf{b}\rangle))}$$

On obtient alors

$$\mathbb{H}(\mathbf{u})\mathbf{x} = \mathbf{x} - \frac{1}{\beta}\mathbf{v}\mathbf{v}^*\mathbf{x} = \mathbf{x} - \frac{\langle\mathbf{v}, \mathbf{x}\rangle}{\beta}\mathbf{v}. \quad (4.44)$$

1



Corollaire 4.21

Soit $\mathbf{a} \in \mathbb{C}^n$ avec $a_1 \neq 0$ et $\exists j \in \llbracket 2, n \rrbracket$ tel que $a_j \neq 0$. Soient $\theta = \arg a_1$ et

$$\mathbf{u}_{\pm} = \frac{\mathbf{a} \pm \|\mathbf{a}\|_2 e^{i\theta} \mathbf{e}_1}{\|\mathbf{a} \pm \|\mathbf{a}\|_2 e^{i\theta} \mathbf{e}_1\|}$$

Alors

$$\mathbb{H}(\mathbf{u}_{\pm})\mathbf{a} = \mp \|\mathbf{a}\|_2 e^{i\theta} \mathbf{e}_1 \quad (4.45)$$

où $\mathbf{e}_1$ désigne le premier vecteur de la base canonique de $\mathbb{C}^n$.

2

Proof. On va utiliser le Théorème 4.20.

On pose $\alpha = \pm \|\mathbf{a}\|_2 e^{i\theta}$ et $\mathbf{b} = \mathbf{e}_1$. Comme $\arg(\langle\mathbf{a}, \mathbf{b}\rangle) = \arg \overline{a_1} = -\arg a_1 = -\theta$, on a $\arg \alpha = \theta \pmod{\pi} = -\arg(\langle\mathbf{a}, \mathbf{b}\rangle) \pmod{\pi}$.

Les autres hypothèses du Théorème 4.20 sont vérifiées puisque le vecteur $\mathbf{a}$ n'est pas colinéaire à $\mathbf{e}_1$ et que $\|\mathbf{e}_1\|_2 = 1$. On a donc

$$\mathbb{H}\left(\frac{\mathbf{a} \mp \|\mathbf{a}\|_2 e^{i\theta} \mathbf{e}_1}{\|\mathbf{a} \mp \|\mathbf{a}\|_2 e^{i\theta} \mathbf{e}_1\|}\right)\mathbf{a} = \pm \|\mathbf{a}\|_2 e^{i\theta} \mathbf{e}_1.$$

□

8

Sur le même principe que l'écriture algébrique de la méthode de Gauss, nous allons transformer la matrice $\mathbb{A} \in \mathcal{M}_n(\mathbb{K})$ en une matrice triangulaire supérieure à l'aide de matrices de Householder. On a le théorème

9

10

11

**Théorème 4.22**

Soit $A \in \mathcal{M}_n(\mathbb{C})$ une matrice. Il existe une matrice unitaire $Q \in \mathcal{M}_n(\mathbb{C})$ produit d'au plus $n - 1$ matrices de Householder et une matrice triangulaire supérieure $R \in \mathcal{M}_n(\mathbb{C})$ telles que

$$A = QR. \quad (4.46)$$

Si A est réelle alors Q et R sont aussi réelles et l'on peut choisir Q de telle sorte que les coefficients diagonaux de R soient positifs. De plus, si A est inversible alors la factorisation est unique.

**Exercice 4.2.7**

Soit $B \in \mathcal{M}_{m+n}(\mathbb{K})$ la matrice bloc

$$B = \left(\begin{array}{c|c} B_{1,1} & B_{1,2} \\ \hline B_{2,1} & S \end{array} \right)$$

où $B_{1,1} \in \mathcal{M}_m(\mathbb{K})$ et $S \in \mathcal{M}_n(\mathbb{K})$. On note $s \in \mathbb{K}^n$ le premier vecteur colonne de S et on suppose que $s \neq 0$ et s non colinéaire à e_1^n premier vecteur de la base canonique de $\mathbb{K}^n$.

Q. 1 1. Montrer qu'il existe une matrice de Householder $H = H(\underline{u}) \in \mathcal{M}_n(\mathbb{K})$ et $\alpha \in \mathbb{K}^*$ tel que

$$HS = \left(\begin{array}{c|ccc} \pm\alpha & \bullet & \cdots & \bullet \\ \hline 0 & \bullet & \cdots & \bullet \\ \vdots & \vdots & & \vdots \\ 0 & \bullet & \cdots & \bullet \end{array} \right).$$

2. On note $u \in \mathbb{K}^{m+n}$, le vecteur défini par $u_i = 0, \forall i \in \llbracket 1, m \rrbracket$ et $u_{m+i} = \underline{u}_i, \forall i \in \llbracket 1, n \rrbracket$. Montrer que

$$H(u)B = \left(\begin{array}{c|c} B_{1,1} & B_{1,2} \\ \hline B_{2,1} & HS \end{array} \right).$$

Soient $k \in \llbracket 0, n - 1 \rrbracket$ et $A^{[k]} \in \mathcal{M}_n(\mathbb{K})$ la matrice bloc définie par

$$A^{[k]} = \left(\begin{array}{c|c} R^{[k]} & F^{[k]} \\ \hline 0 & A^{[k]} \end{array} \right)$$

où $R^{[k]}$ est une matrice triangulaire supérieure d'ordre k et $A^{[k]}$ une matrice d'ordre $n - k$.

Q. 2 1. Sous certaines hypothèses, montrer qu'il existe une matrice de Householder $H^{[k+1]}$ telle que $H^{[k+1]}A^{[k]} = A^{[k+1]}$.

2. Soit $A \in \mathcal{M}_n(\mathbb{K})$. Montrer qu'il existe une matrice unitaire $Q \in \mathcal{M}_n(\mathbb{K})$, produit d'au plus $n - 1$ matrices de Householder, et une matrice triangulaire supérieure R telles que $A = QR$.

3. Montrer que si A est réelle alors les coefficients diagonaux de R peuvent être choisis positifs.

4. Montrer que si A est réelle inversible alors la factorisation QR , avec R à coefficients diagonaux positifs, est unique.

Correction Exercice 4.2.7

Q. 1 1. D'après le (voir Corollaire 4.21, page 101) avec $a = s$, en posant $\alpha = \pm \|s\|_2 e^{i \arg s_1}$ et

$$\underline{u} = \frac{s - \alpha e_1^n}{\|s - \alpha e_1^n\|}$$

on obtient $H(\underline{u}) = \alpha e_1^n$.

On pose $\mathbb{H} = \mathbb{H}(\underline{\mathbf{u}})$. On a alors sous forme bloc

$$\mathbb{H}\mathbb{S} = \mathbb{H} \left(\begin{array}{c|ccc} & \bullet & \cdots & \bullet \\ \mathbf{s} & \vdots & & \vdots \\ & \bullet & \cdots & \bullet \end{array} \right) = \left(\begin{array}{c|ccc} \pm\alpha & \bullet & \cdots & \bullet \\ 0 & \bullet & \cdots & \bullet \\ \vdots & \vdots & & \vdots \\ 0 & \bullet & \cdots & \bullet \end{array} \right)$$

2. On a $\underline{\mathbf{u}} = \begin{pmatrix} \mathbf{0}_m \\ \underline{\mathbf{u}} \end{pmatrix}$ et

$$\begin{aligned} \mathbb{H}(\underline{\mathbf{u}}) &= \mathbb{I} - 2\underline{\mathbf{u}}\underline{\mathbf{u}}^* = \left(\begin{array}{c|c} \mathbb{I}_m & \mathbf{0}_{m,n} \\ \hline \mathbf{0}_{n,m} & \mathbb{I}_n \end{array} \right) - 2 \left(\begin{array}{c|c} \mathbf{0}_m \\ \underline{\mathbf{u}} \end{array} \right) \left(\begin{array}{c|c} \mathbf{0}_m^* & \underline{\mathbf{u}}^* \end{array} \right) \\ &= \left(\begin{array}{c|c} \mathbb{I}_m & \mathbf{0}_{m,n} \\ \hline \mathbf{0}_{n,m} & \mathbb{I}_n \end{array} \right) - 2 \left(\begin{array}{c|c} \mathbf{0}_m & \mathbf{0}_{m,n} \\ \hline \mathbf{0}_{n,m} & \underline{\mathbf{u}}\underline{\mathbf{u}}^* \end{array} \right) \\ &= \left(\begin{array}{c|c} \mathbb{I}_m & \mathbf{0}_{m,n} \\ \hline \mathbf{0}_{n,m} & \mathbb{I}_n - 2\underline{\mathbf{u}}\underline{\mathbf{u}}^* \end{array} \right) = \left(\begin{array}{c|c} \mathbb{I}_m & \mathbf{0}_{m,n} \\ \hline \mathbf{0}_{n,m} & \mathbb{H} \end{array} \right) \end{aligned}$$

Ce qui donne

$$\mathbb{H}(\underline{\mathbf{u}})\mathbb{B} = \left(\begin{array}{c|c} \mathbb{I}_m & \mathbf{0}_{m,n} \\ \hline \mathbf{0}_{n,m} & \mathbb{H} \end{array} \right) \left(\begin{array}{c|c} \mathbb{B}_{1,1} & \mathbb{B}_{1,2} \\ \hline \mathbb{B}_{2,1} & \mathbb{S} \end{array} \right) = \left(\begin{array}{c|c} \mathbb{B}_{1,1} & \mathbb{B}_{1,2} \\ \hline \mathbb{B}_{2,1} & \mathbb{H}\mathbb{S} \end{array} \right).$$

Q. 2 1. On note $\underline{\mathbf{s}} \in \mathbb{K}^{n-k}$ le premier vecteur colonne de $\underline{\mathbb{A}}^{[k]}$, et $\underline{\mathbf{u}} = \begin{pmatrix} \mathbf{0}_k \\ \underline{\mathbf{s}} \end{pmatrix}$. D'après la question précédente si $\underline{\mathbf{s}} \neq 0$ et $\underline{\mathbf{s}}$ non colinéaire à $\mathbf{e}_1^{n-k}$ premier vecteur de la base canonique de $\mathbb{K}^{n-k}$ alors il existe une matrice de Householder $\mathbb{H}^{[k+1]} = \mathbb{H}(\underline{\mathbf{u}})$ et $\alpha \in \mathbb{K}^*$ tels que

$$\underline{\mathbb{A}}^{[k+1]} \stackrel{\text{def}}{=} \mathbb{H}^{[k+1]}\underline{\mathbb{A}}^{[k]} = \left(\begin{array}{c|ccc} \mathbb{R}^{[k]} & & & \mathbb{F}^{[k]} \\ \hline & \left(\begin{array}{c|ccc} \pm\alpha & \bullet & \cdots & \bullet \\ 0 & \bullet & \cdots & \bullet \\ \vdots & \vdots & & \vdots \\ 0 & \bullet & \cdots & \bullet \end{array} \right) & & \end{array} \right) = \left(\begin{array}{c|c} \mathbb{R}^{[k+1]} & \mathbb{F}^{[k+1]} \\ \hline \mathbf{0} & \underline{\mathbb{A}}^{[k+1]} \end{array} \right)$$

On peut remarquer que si $\underline{\mathbf{s}} = 0$ ou $\underline{\mathbf{s}}$ colinéaire à $\mathbf{e}_1^{n-k}$ alors $\underline{\mathbb{A}}^{[k]}$ est déjà sous la forme $\underline{\mathbb{A}}^{[k+1]}$ et donc $\mathbb{H}^{[k]} = \mathbb{I}$.

2. il suffit d'appliquer itérativement le résultat précédent $n-1$ fois en posant $\underline{\mathbb{A}}^{[0]} = \underline{\mathbb{A}}$ et $\underline{\mathbb{A}}^{[k+1]} = \mathbb{H}^{[k+1]}\underline{\mathbb{A}}^{[k]}$ où $\mathbb{H}^{[k+1]}$ est soit une matrice de Householder soit la matrice identité. Par construction la matrice $\underline{\mathbb{A}}^{[n-1]}$ est triangulaire supérieure et l'on a

$$\underline{\mathbb{A}}^{[n-1]} = \mathbb{H}^{[n-1]} \times \cdots \times \mathbb{H}^{[1]} \underline{\mathbb{A}}$$

On pose $\mathbb{H} = \mathbb{H}^{[n-1]} \times \cdots \times \mathbb{H}^{[1]}$ et $\mathbb{R} = \underline{\mathbb{A}}^{[n-1]}$. La matrice $\mathbb{H}$ est unitaire car produit de matrices unitaires. On note $\mathbb{Q} = \mathbb{H}^*$ On a

$$\mathbb{Q} = \mathbb{H}^{[1]} \times \cdots \times \mathbb{H}^{[n-1]}$$

car les matrices de Householder et matrice identité sont unitaires et hermitiennes.

3. Si $\underline{\mathbb{A}}$ est réelle alors par construction $\mathbb{Q}$ et $\mathbb{R}$ sont réelles. Les coefficients diagonaux peuvent alors être choisis positifs lors de la construction de chaque matrice de Householder.

4. Pour montrer l'unicité d'une telle factorisation, on note $\mathbb{Q}_1, \mathbb{Q}_2$, deux matrices orthogonales et $\mathbb{R}_1, \mathbb{R}_2$, deux matrices triangulaires à coefficients diagonaux positifs telles que

$$\underline{\mathbb{A}} = \mathbb{Q}_1 \mathbb{R}_1 = \mathbb{Q}_2 \mathbb{R}_2.$$

Comme $\underline{\mathbb{A}}$ est inversible les coefficients diagonaux de $\mathbb{R}_1$ et $\mathbb{R}_2$ sont strictement positifs. On a alors

$$\mathbb{I} = \underline{\mathbb{A}} \underline{\mathbb{A}}^{-1} = \mathbb{Q}_1 \mathbb{R}_1 \mathbb{R}_2^{-1} \mathbb{Q}_2^{-1}$$

1 et donc

$$Q_1^{-1}Q_2 = R_1R_2^{-1} \stackrel{\text{def}}{=} T.$$

2 Comme Q_1 est orthogonale on a $T = Q_1^t Q_2$ et

$$T^t T = (Q_1^t Q_2)^t Q_1^t Q_2 = Q_2^t Q_1 Q_1^t Q_2 = \mathbb{I}.$$

3 De plus la matrice $T = R_1 R_2^{-1}$ est triangulaire supérieure à coefficients diagonaux strictement positifs
 4 puisque produit de triangulaire supérieure à coefficients diagonaux strictement positifs. D'après
 5 le théorème (factorisation de Cholesky) il existe une unique matrice L triangulaire inférieure à
 6 coefficients diagonaux strictement positifs telle que $LL^t = \mathbb{I}$ et Cette matrice L est la matrice
 7 identité. On en déduit que $T = L^t = \mathbb{I}$ et donc $Q_1 = Q_2$ et $R_1 = R_2$.

◇

Exercice 4.2.8: Algorithmique

Q. 1 Ecrire une fonction **FACTQR** permettant de calculer la factorisation QR d'une matrice $A \in \mathcal{M}_n(\mathbb{C})$.

On pourra utiliser la fonction **HOUSEHOLDER** (voir Exercice 4.2.6, page 100).

Q. 2 Ecrire un programme permettant de tester cette fonction.

11 Correction Exercice 4.2.8

12 **Q. 1** L'objectif est de déterminer les matrices Q , matrice unitaire, et R matrice triangulaire supérieure
 13 telle que $A = QR$.

Données : A : matrice de $\mathcal{M}_n(\mathbb{K})$.

14 **Résultat :** Q : matrice unitaire de $\mathcal{M}_n(\mathbb{K})$.

R : matrice triangulaire supérieure de $\mathcal{M}_n(\mathbb{K})$.

15 On rappelle la technique utilisée dans la correction de l'exercice 4.2.7 pour déterminer l'ensemble des
 16 matrices de Householder permettant de transformer la matrice A en une matrice triangulaire supérieure.
 17 On pose

$$A^{[0]} = A, \quad A^{[k+1]} = H^{[k+1]} A^{[k]}, \quad \forall k \in \llbracket 0, n-2 \rrbracket$$

18 où $H^{[k+1]}$ est soit une matrice de Householder soit la matrice identité. Plus précisément, on note $\underline{s} \in \mathbb{K}^{n-k}$
 19 le vecteur composé des $n-k$ dernières composantes de la $k+1$ -ème colonne de $A^{[k]}$ et $\underline{a} = \begin{pmatrix} 0_k \\ \underline{s} \end{pmatrix}$.

20 • Si $s_1 = 0$ ou $\underline{s}$ colinéaire à $\underline{e}_1^{n-k}$ premier vecteur de la base canonique de $\mathbb{K}^{n-k}$ alors

$$H^{[k+1]} = H(\underline{u}).$$

21 En notant $\underline{e}_{k+1}^n$ le $k+1$ -ème vecteur de la base canonique de $\mathbb{K}^n$, cette matrice peut-être calculée
 22 avec la fonction **HOUSEHOLDER** par

$$[H^{[k+1]}, \alpha] \leftarrow \text{HOUSEHOLDER}(\underline{a}, \underline{e}_{k+1}^n, 1)$$

23 • sinon $H^{[k+1]} = \mathbb{I}$.

24 On a vu que dans ce cas $A^{[n-1]}$ est triangulaire supérieure. On pose $H = H^{[n-1]} \times \dots \times H^{[1]}$ qui est une
 25 matrice unitaire. On a alors $R = A^{[n-1]} = HA$ et $Q = H^*$.

Algorithme 4.18 $\mathcal{R}_0$

1: Calculer Q et R

Algorithme 4.18 $\mathcal{R}_1$

1: $H \leftarrow H^{[n-1]} \times \dots \times H^{[1]}$
 2: $R \leftarrow H * A$
 3: $Q \leftarrow H^*$

Algorithme 4.18 $\mathcal{R}_1$

```

1:  $\mathbb{H} \leftarrow \mathbb{H}^{[n-1]} \times \dots \times \mathbb{H}^{[1]}$ 
2:  $\mathbb{R} \leftarrow \mathbb{H} * \mathbb{A}$ 
3:  $\mathbb{Q} \leftarrow \mathbb{H}^*$ 

```

Algorithme 4.18 $\mathcal{R}_2$

```

1:  $\mathbb{H} \leftarrow \mathbb{I}$ 
2:  $\mathbb{A}^{[0]} \leftarrow \mathbb{A}$ 
3: Pour  $k \leftarrow 0$  à  $n - 2$  faire
4:   Calculer  $\mathbb{H}^{[k+1]}$  à partir de  $\mathbb{A}^{[k]}$ 
5:    $\mathbb{A}^{[k+1]} \leftarrow \mathbb{H}^{[k+1]} * \mathbb{A}^{[k]}$ 
6:    $\mathbb{H} \leftarrow \mathbb{H}^{[k+1]} * \mathbb{H}$ 
7: Fin Pour
8:  $\mathbb{R} \leftarrow \mathbb{H} * \mathbb{A}$ 
9:  $\mathbb{Q} \leftarrow \mathbb{H}^*$ 

```

▷ ou $\mathbb{R} \leftarrow \mathbb{A}^{[n-1]}$

Algorithme 4.18 $\mathcal{R}_2$

```

1:  $\mathbb{H} \leftarrow \mathbb{I}$ 
2: Pour  $k \leftarrow 0$  à  $n - 2$  faire
3:   Calculer  $\mathbb{H}^{[k+1]}$  à partir de  $\mathbb{A}^{[k]}$ 
4:    $\mathbb{A}^{[k+1]} \leftarrow \mathbb{H}^{[k+1]} * \mathbb{A}^{[k]}$ 
5:    $\mathbb{H} \leftarrow \mathbb{H}^{[k+1]} * \mathbb{H}$ 
6: Fin Pour
7:  $\mathbb{R} \leftarrow \mathbb{A}^{[n-1]}$ 
8:  $\mathbb{Q} \leftarrow \mathbb{H}^*$ 

```

Algorithme 4.18 $\mathcal{R}_3$

```

1:  $\mathbb{H} \leftarrow \mathbb{I}$ 
2: Pour  $k \leftarrow 0$  à  $n - 2$  faire
3:    $\mathbf{a} \leftarrow [\mathbf{0}_k; \mathbb{A}^{[k]}(k + 1 : n, k + 1)]$ 
4:    $[\mathbb{H}^{[k+1]}, \alpha] \leftarrow \text{HOUSEHOLDER}(\mathbf{a}, \mathbf{e}_{k+1}^n, 1)$ 
5:    $\mathbb{A}^{[k+1]} \leftarrow \mathbb{H}^{[k+1]} * \mathbb{A}^{[k]}$ 
6:    $\mathbb{H} \leftarrow \mathbb{H}^{[k+1]} \mathbb{H}$ 
7: Fin Pour
8:  $\mathbb{R} \leftarrow \mathbb{A}^{[n-1]}$ 
9:  $\mathbb{Q} \leftarrow \mathbb{H}^*$ 

```

Ici, l'opérateur $[\bullet; \bullet]$ est l'opérateur de concaténation de deux vecteurs.

Algorithme 4.18 Fonction **FACTQR**

Données : $\mathbb{A}$: matrice de $\mathcal{M}_n(\mathbb{K})$.

Résultat : $\mathbb{Q}$: matrice unitaire de $\mathcal{M}_n(\mathbb{K})$.

$\mathbb{R}$: matrice triangulaire supérieure de $\mathcal{M}_n(\mathbb{K})$.

```

1: Fonction  $[\mathbb{Q}, \mathbb{R}] \leftarrow \text{FACTQR}(\mathbb{A})$ 
2:    $\mathbb{H} \leftarrow \mathbb{I}$ 
3:    $\mathbb{R} \leftarrow \mathbb{A}$ 
4:   Pour  $k \leftarrow 0$  à  $n - 2$  faire
5:      $\mathbf{a} \leftarrow [\mathbf{0}_k; \mathbb{R}(k + 1 : n, k + 1)]$ 
6:      $[\mathbb{S}, \alpha] \leftarrow \text{HOUSEHOLDER}(\mathbf{a}, \mathbf{e}_{k+1}^n, 1)$ 
7:      $\mathbb{R} \leftarrow \mathbb{S} * \mathbb{R}$ 
8:      $\mathbb{H} \leftarrow \mathbb{S} * \mathbb{H}$ 
9:   Fin Pour
10:   $\mathbb{Q} \leftarrow \mathbb{H}^*$ 
11: Fin Fonction

```

Q. 2

4.3 Normes vectorielles et normes matricielles

4.3.1 Normes vectorielles

♥ Définition 4.23

Une **norme** sur un espace vectoriel V est une application $\|\bullet\| : V \rightarrow \mathbb{R}^+$ qui vérifie les propriétés suivantes

- ◇ $\|\mathbf{v}\| = 0 \iff \mathbf{v} = 0$,
- ◇ $\|\alpha\mathbf{v}\| = |\alpha| \|\mathbf{v}\|$, $\forall \alpha \in \mathbb{K}$, $\forall \mathbf{v} \in V$,
- ◇ $\|\mathbf{u} + \mathbf{v}\| \leq \|\mathbf{u}\| + \|\mathbf{v}\|$, $\forall (\mathbf{u}, \mathbf{v}) \in V^2$ (inégalité triangulaire).

Une norme sur V est également appelée **norme vectorielle**. On appelle **espace vectoriel normé** un espace vectoriel muni d'une norme.

Les trois normes suivantes sont les plus couramment utilisées :

$$\begin{aligned}\|\mathbf{v}\|_1 &= \sum_{i=1}^n |v_i| \\ \|\mathbf{v}\|_2 &= \left(\sum_{i=1}^n |v_i|^2 \right)^{1/2} \\ \|\mathbf{v}\|_\infty &= \max_{i \in \llbracket 1, n \rrbracket} |v_i|.\end{aligned}$$

Proposition 4.24

Soit $\mathbf{v} \in \mathbb{K}^n$. Pour tout nombre réel $p \geq 1$, l'application $\|\bullet\|_p$ définie par

$$\|\mathbf{v}\|_p = \left(\sum_{i=1}^n |v_i|^p \right)^{1/p}$$

est une norme sur $\mathbb{K}^n$.

Lemme 4.25: Inégalité de Cauchy-Schwarz

$\forall \mathbf{x}, \mathbf{y} \in \mathbb{K}^n$

$$|\langle \mathbf{x}, \mathbf{y} \rangle| \leq \|\mathbf{x}\|_2 \|\mathbf{y}\|_2. \quad (4.47)$$

Cette inégalité s'appelle l'**inégalité de Cauchy-Schwarz**. On a égalité si et seulement si $\mathbf{x}$ et $\mathbf{y}$ sont colinéaires.

Proof. Exercice A.3.16, page 199

□

Lemme 4.26: Inégalité de Hölder

Pour $p > 1$ et $\frac{1}{p} + \frac{1}{q} = 1$, on a $\forall \mathbf{u}, \mathbf{v} \in \mathbb{K}^n$

$$\sum_{i=1}^n |u_i v_i| \leq \left(\sum_{i=1}^n |u_i|^p \right)^{1/p} \left(\sum_{i=1}^n |v_i|^q \right)^{1/q} = \|\mathbf{u}\|_p \|\mathbf{v}\|_q. \quad (4.48)$$

Cette inégalité s'appelle l'**inégalité de Hölder**.

Proof. Exercice A.3.18, page 203

□

♥ Definition 4.27

Deux **normes** $\|\bullet\|$ et $\|\bullet\|'$, définies sur un même espace vectoriel V , sont **équivalentes** s'il existe deux constantes C et C' telles que

$$\|\mathbf{v}\|' \leq C \|\mathbf{v}\| \quad \text{et} \quad \|\mathbf{v}\| \leq C' \|\mathbf{v}\|' \quad \text{pour tout } \mathbf{v} \in V. \quad (4.49)$$

📖 Proposition 4.28

Sur un espace vectoriel de dimension finie toutes les normes sont équivalentes.

4.3.2 Normes matricielles

♥ Definition 4.29

Une **norme matricielle** sur $\mathcal{M}_n(\mathbb{K})$ est une application $\|\bullet\| : \mathcal{M}_n(\mathbb{K}) \rightarrow \mathbb{R}^+$ vérifiant

1. $\|\mathbb{A}\| = 0 \iff \mathbb{A} = 0$,
2. $\|\alpha \mathbb{A}\| = |\alpha| \|\mathbb{A}\|, \forall \alpha \in \mathbb{K}, \forall \mathbb{A} \in \mathcal{M}_n(\mathbb{K})$,
3. $\|\mathbb{A} + \mathbb{B}\| \leq \|\mathbb{A}\| + \|\mathbb{B}\|, \forall (\mathbb{A}, \mathbb{B}) \in \mathcal{M}_n(\mathbb{K})^2$ (inégalité triangulaire)
4. $\|\mathbb{A}\mathbb{B}\| \leq \|\mathbb{A}\| \|\mathbb{B}\|, \forall (\mathbb{A}, \mathbb{B}) \in \mathcal{M}_n(\mathbb{K})^2$

📖 Proposition 4.30

Etant donné une norme vectorielle $\|\bullet\|$ sur $\mathbb{C}^n$, l'application $\|\bullet\|_s : \mathcal{M}_n(\mathbb{C}) \rightarrow \mathbb{R}^+$ définie par

$$\|\mathbb{A}\|_s = \sup_{\substack{\mathbf{v} \in \mathbb{C}^n \\ \mathbf{v} \neq 0}} \frac{\|\mathbb{A}\mathbf{v}\|}{\|\mathbf{v}\|} = \sup_{\substack{\mathbf{v} \in \mathbb{C}^n \\ \|\mathbf{v}\| \leq 1}} \|\mathbb{A}\mathbf{v}\| = \sup_{\substack{\mathbf{v} \in \mathbb{C}^n \\ \|\mathbf{v}\| = 1}} \|\mathbb{A}\mathbf{v}\|, \quad (4.50)$$

est une norme matricielle, appelée **norme matricielle subordonnée** (à la norme vectorielle donnée).

De plus

$$\|\mathbb{A}\mathbf{v}\| \leq \|\mathbb{A}\|_s \|\mathbf{v}\| \quad \forall \mathbf{v} \in \mathbb{C}^n \quad (4.51)$$

et la norme $\|\mathbb{A}\|$ peut se définir aussi par

$$\|\mathbb{A}\|_s = \inf \{ \alpha \in \mathbb{R} : \|\mathbb{A}\mathbf{v}\| \leq \alpha \|\mathbf{v}\|, \forall \mathbf{v} \in \mathbb{C}^n \}. \quad (4.52)$$

Il existe au moins un vecteur $\mathbf{u} \in \mathbb{C}^n$ tel que

$$\mathbf{u} \neq 0 \quad \text{et} \quad \|\mathbb{A}\mathbf{u}\| = \|\mathbb{A}\|_s \|\mathbf{u}\|. \quad (4.53)$$

Enfin une norme subordonnée vérifie toujours

$$\|\mathbb{I}\|_s = 1 \quad (4.54)$$

Proof. On note $\mathcal{B} = \{\mathbf{v} \in \mathbb{C}^n ; \|\mathbf{v}\| \leq 1\}$ la boule unité de $\mathbb{C}^n$ et $\mathcal{S} = \{\mathbf{v} \in \mathbb{C}^n ; \|\mathbf{v}\| = 1\}$ la sphère unité de $\mathbb{C}^n$. On note que les ensembles $\mathcal{B}$ et $\mathcal{S}$ sont des compacts car image réciproque de l'application continue $\mathbf{v} \mapsto \|\mathbf{v}\|$ par le fermé borné $[0, 1]$ (pour la boule) et le singleton $\{1\}$ (pour la sphère).

- Vérifions que les égalités suivantes sont vraies :

$$\sup_{\substack{\mathbf{v} \in \mathbb{C}^n \\ \mathbf{v} \neq 0}} \frac{\|\mathbb{A}\mathbf{v}\|}{\|\mathbf{v}\|} = \sup_{\mathbf{v} \in \mathcal{B}} \|\mathbb{A}\mathbf{v}\| = \sup_{\mathbf{v} \in \mathcal{S}} \|\mathbb{A}\mathbf{v}\|$$

1 On a

$$\sup_{\substack{\mathbf{v} \in \mathbb{C}^n \\ \mathbf{v} \neq 0}} \frac{\|\mathbb{A}\mathbf{v}\|}{\|\mathbf{v}\|} = \sup_{\substack{\mathbf{v} \in \mathbb{C}^n \\ \mathbf{v} \neq 0}} \left\| \frac{\mathbb{A}\mathbf{v}}{\|\mathbf{v}\|} \right\| = \sup_{\mathbf{v} \in \mathcal{S}} \|\mathbb{A}\mathbf{v}\|$$

2 Comme $\mathcal{S} \subset \mathcal{B}$ on a aussi

$$\sup_{\mathbf{v} \in \mathcal{B}} \|\mathbb{A}\mathbf{v}\| \geq \sup_{\mathbf{v} \in \mathcal{S}} \|\mathbb{A}\mathbf{v}\|.$$

3 De plus $\forall \mathbf{w} \in \mathcal{B}, \mathbf{u} = \frac{\mathbf{w}}{\|\mathbf{w}\|} \in \mathcal{S}$ et on a

$$\|\mathbb{A}\mathbf{w}\| = \|\mathbf{w}\| \|\mathbb{A}\mathbf{u}\| \leq \|\mathbb{A}\mathbf{u}\|$$

4 On en déduit

$$\sup_{\mathbf{w} \in \mathcal{B}} \|\mathbb{A}\mathbf{w}\| \leq \sup_{\mathbf{u} \in \mathcal{S}} \|\mathbb{A}\mathbf{u}\|.$$

5 et donc

$$\sup_{\mathbf{w} \in \mathcal{B}} \|\mathbb{A}\mathbf{w}\| = \sup_{\mathbf{u} \in \mathcal{S}} \|\mathbb{A}\mathbf{u}\|.$$

6 • Vérifions que l'application $\|\bullet\|_s$ est bien définie sur $\mathcal{M}_n(\mathbb{C})$ i.e. $\forall \mathbb{A} \in \mathcal{M}_n(\mathbb{C}), \|\mathbb{A}\|_s < +\infty$.

7 L'application $\mathbf{v} \mapsto \|\mathbb{A}\mathbf{v}\|$ est continue donc son sup sur la sphère unité qui est compacte est atteint.

8 • Montrons que $\forall \mathbf{v} \in \mathbb{C}^n, \|\mathbb{A}\mathbf{v}\| \leq \|\mathbb{A}\|_s \|\mathbf{v}\|$.

9 On a par définition du sup

$$\|\mathbb{A}\|_s = \sup_{\substack{\mathbf{u} \in \mathbb{C}^n \\ \mathbf{u} \neq 0}} \frac{\|\mathbb{A}\mathbf{u}\|}{\|\mathbf{u}\|} \geq \frac{\|\mathbb{A}\mathbf{v}\|}{\|\mathbf{v}\|}, \quad \forall \mathbf{v} \in \mathbb{C}^n \setminus \{0\}.$$

10 et donc

$$\|\mathbb{A}\mathbf{v}\| \leq \|\mathbb{A}\|_s \|\mathbf{v}\|, \quad \forall \mathbf{v} \in \mathbb{C}^n \setminus \{0\}.$$

• Montrons qu'il $\mathbf{u} \in \mathbb{C}^n$ tel que

$$\mathbf{u} \neq 0 \quad \text{et} \quad \|\mathbb{A}\mathbf{u}\| = \|\mathbb{A}\|_s \|\mathbf{u}\|.$$

11 On a

$$\|\mathbb{A}\|_s = \sup_{\mathbf{v} \in \mathcal{S}} \|\mathbb{A}\mathbf{v}\|$$

12 La sphère unité étant compacte et l'application $\mathbf{v} \mapsto \|\mathbb{A}\mathbf{v}\|$ étant continue, il existe $\mathbf{w} \in \mathcal{S}$ tel que

13 $\|\mathbb{A}\|_s = \|\mathbb{A}\mathbf{w}\|$. Soit $\lambda \in \mathbb{C}^*$ et $\mathbf{u} = \lambda \mathbf{w} \neq 0$. On a $\|\mathbf{u}\| = \lambda$ et

$$\|\mathbb{A}\|_s = \|\mathbb{A}\mathbf{w}\| = \left\| \mathbb{A} \frac{\mathbf{u}}{\|\mathbf{u}\|} \right\| = \frac{1}{\|\mathbf{u}\|} \|\mathbb{A}\mathbf{u}\|.$$

14 • On a immédiatement

$$\|\mathbb{I}\|_s = \sup_{\substack{\mathbf{v} \in \mathbb{C}^n \\ \mathbf{v} \neq 0}} \frac{\|\mathbf{v}\|}{\|\mathbf{v}\|} = \sup_{\substack{\mathbf{v} \in \mathbb{C}^n \\ \mathbf{v} \neq 0}} \frac{\|\mathbf{v}\|}{\|\mathbf{v}\|} = 1.$$

15 • Montrons que $\|\bullet\|_s$ est une norme matricielle.

16

1. $\|\mathbb{A}\|_s = 0 \iff \mathbb{A}_s = \mathbb{O}$?

$\Longleftarrow$ trivial.

$\Longrightarrow$ Soit $\mathbb{A} \in \mathcal{M}_n(\mathbb{K})$.

$$\begin{aligned} \|\mathbb{A}\|_s = 0 &= \sup_{\substack{\mathbf{v} \in \mathbb{K}^n \\ \mathbf{v} \neq 0}} \frac{\|\mathbb{A}\mathbf{v}\|}{\|\mathbf{v}\|} \implies \|\mathbb{A}\mathbf{v}\| = 0, \quad \forall \mathbf{v} \in \mathbb{K}^n \setminus \{0\} \\ &\implies \mathbb{A}\mathbf{v} = \mathbf{0}, \quad \forall \mathbf{v} \in \mathbb{K}^n \setminus \{0\} \end{aligned}$$

17 Soit $\{\mathbf{e}_1, \dots, \mathbf{e}_n\}$ la base canonique de $\mathbb{K}^n$. On a alors $\forall j \in \llbracket 1, n \rrbracket, \mathbb{A}\mathbf{e}_j = \mathbf{0}$ et on en déduit que

$$A_{i,j} = \langle \mathbf{e}_i, \mathbb{A}\mathbf{e}_j \rangle = 0, \quad \forall (i, j) \in \llbracket 1, n \rrbracket.$$

18 et donc $\mathbb{A} = \mathbb{O}$.

2. Montrons que $\|\alpha A\| = |\alpha| \|A\|$, $\forall \alpha \in \mathbb{K}$, $\forall A \in \mathcal{M}_n(\mathbb{K})$.

Soient $\alpha \in \mathbb{K}$ et $A \in \mathcal{M}_n(\mathbb{K})$. On a $\alpha A \in \mathcal{M}_n(\mathbb{K})$ (car $\mathcal{M}_n(\mathbb{K})$ est un espace vectoriel) et

$$\begin{aligned}\|\alpha A\|_s &= \sup_{\substack{\mathbf{v} \in \mathbb{K}^n \\ \mathbf{v} \neq 0}} \frac{\|\alpha A \mathbf{v}\|}{\|\mathbf{v}\|} = \sup_{\substack{\mathbf{v} \in \mathbb{K}^n \\ \mathbf{v} \neq 0}} \frac{|\alpha| \|A \mathbf{v}\|}{\|\mathbf{v}\|} \quad \text{car } \|\alpha \mathbf{u}\| = |\alpha| \|\mathbf{u}\| \\ &= |\alpha| \sup_{\substack{\mathbf{v} \in \mathbb{K}^n \\ \mathbf{v} \neq 0}} \frac{\|A \mathbf{v}\|}{\|\mathbf{v}\|} = |\alpha| \|A\|_s.\end{aligned}$$

3. Montrons que $\|A + B\|_s \leq \|A\|_s + \|B\|_s$, $\forall (A, B) \in \mathcal{M}_n(\mathbb{K})^2$

Soient A et B deux matrices de $\mathcal{M}_n(\mathbb{K})$. On a $A + B \in \mathcal{M}_n(\mathbb{K})$ car $\mathcal{M}_n(\mathbb{K})$ est un espace vectoriel et

$$\begin{aligned}\|A + B\|_s &= \sup_{\substack{\mathbf{v} \in \mathbb{K}^n \\ \mathbf{v} \neq 0}} \frac{\|(A + B)\mathbf{v}\|}{\|\mathbf{v}\|} = \sup_{\substack{\mathbf{v} \in \mathbb{K}^n \\ \mathbf{v} \neq 0}} \frac{\|A\mathbf{v} + B\mathbf{v}\|}{\|\mathbf{v}\|} \\ &\leq \sup_{\substack{\mathbf{v} \in \mathbb{K}^n \\ \mathbf{v} \neq 0}} \frac{\|A\mathbf{v}\| + \|B\mathbf{v}\|}{\|\mathbf{v}\|} \quad \text{par inégalité triangulaire dans } \mathbb{K}^n \\ &\leq \sup_{\substack{\mathbf{v} \in \mathbb{K}^n \\ \mathbf{v} \neq 0}} \frac{\|A\mathbf{v}\|}{\|\mathbf{v}\|} + \sup_{\substack{\mathbf{v} \in \mathbb{K}^n \\ \mathbf{v} \neq 0}} \frac{\|B\mathbf{v}\|}{\|\mathbf{v}\|} = \|A\|_s + \|B\|_s.\end{aligned}$$

4. Montrons que $\|AB\|_s \leq \|A\|_s \|B\|_s$, $\forall (A, B) \in \mathcal{M}_n(\mathbb{K})^2$.

Soient A et B deux matrices de $\mathcal{M}_n(\mathbb{K})$. On a $AB \in \mathcal{M}_n(\mathbb{K})$ par définition du produit matriciel et

$$\begin{aligned}\|AB\|_s &= \sup_{\substack{\mathbf{v} \in \mathbb{K}^n \\ \mathbf{v} \neq 0}} \frac{\|(AB)\mathbf{v}\|}{\|\mathbf{v}\|} = \sup_{\substack{\mathbf{v} \in \mathbb{K}^n \\ \mathbf{v} \neq 0}} \frac{\|A(B\mathbf{v})\|}{\|\mathbf{v}\|} \\ &\leq \sup_{\substack{\mathbf{v} \in \mathbb{K}^n \\ \mathbf{v} \neq 0}} \frac{\|A\|_s \|B\mathbf{v}\|}{\|\mathbf{v}\|} \quad \text{car } \|A\mathbf{u}\| \leq \|A\|_s \|\mathbf{u}\| \quad \forall \mathbf{u} \in \mathbb{C}^n \\ &\leq \|A\|_s \sup_{\substack{\mathbf{v} \in \mathbb{K}^n \\ \mathbf{v} \neq 0}} \frac{\|B\mathbf{v}\|}{\|\mathbf{v}\|} = \|A\|_s \|B\|_s.\end{aligned}$$

□

1



Théorème 4.31

Soit $A \in \mathcal{M}_n(\mathbb{K})$. On a

$$\|A\|_1 \stackrel{\text{def}}{=} \sup_{\substack{\mathbf{v} \in \mathbb{K}^n \\ \mathbf{v} \neq 0}} \frac{\|A\mathbf{v}\|_1}{\|\mathbf{v}\|_1} = \max_{j \in \llbracket 1, n \rrbracket} \sum_{i=1}^n |a_{ij}| \quad (4.55)$$

$$\|A\|_2 \stackrel{\text{def}}{=} \sup_{\substack{\mathbf{v} \in \mathbb{K}^n \\ \mathbf{v} \neq 0}} \frac{\|A\mathbf{v}\|_2}{\|\mathbf{v}\|_2} = \sqrt{\rho(A^*A)} = \sqrt{\rho(AA^*)} = \|A^*\|_2 \quad (4.56)$$

$$\|A\|_\infty \stackrel{\text{def}}{=} \sup_{\substack{\mathbf{v} \in \mathbb{K}^n \\ \mathbf{v} \neq 0}} \frac{\|A\mathbf{v}\|_\infty}{\|\mathbf{v}\|_\infty} = \max_{i \in \llbracket 1, n \rrbracket} \sum_{j=1}^n |a_{ij}| \quad (4.57)$$

La norme $\|\bullet\|_2$ est invariante par transformation unitaire :

$$UU^* = \mathbb{I} \implies \|A\|_2 = \|AU\|_2 = \|UA\|_2 = \|U^*AU\|_2. \quad (4.58)$$

2



Corollaire 4.32

3

1. Si une matrice $\mathbb{A}$ est hermitienne, ou symétrique (donc normale), on a $\|\mathbb{A}\|_2 = \rho(\mathbb{A})$.
2. Si une matrice $\mathbb{A}$ est unitaire, ou orthogonale (donc normale), on a $\|\mathbb{A}\|_2 = 1$.

**Théorème 4.33**

1. Soit $\mathbb{A}$ une matrice carrée quelconque et $\|\bullet\|$ une norme matricielle subordonnée ou non, quelconque. Alors

$$\rho(\mathbb{A}) \leq \|\mathbb{A}\|. \quad (4.59)$$

2. Etant donné une matrice $\mathbb{A}$ et un nombre $\varepsilon > 0$, il existe au moins une norme matricielle subordonnée telle que

$$\|\mathbb{A}\| \leq \rho(\mathbb{A}) + \varepsilon. \quad (4.60)$$

**Théorème 4.34**

L'application $\|\bullet\|_E : \mathcal{M}_n \rightarrow \mathbb{R}^+$ définie par

$$\|\mathbb{A}\|_E = \left(\sum_{(i,j) \in \llbracket 1,n \rrbracket^2} |a_{ij}|^2 \right)^{1/2} = \sqrt{\text{tr}(\mathbb{A}^* \mathbb{A})}, \quad (4.61)$$

pour toute matrice $\mathbb{A} = (a_{ij})$ d'ordre n , est une norme matricielle non subordonnée (pour $n \geq 2$), invariante par transformation unitaire et qui vérifie

$$\|\mathbb{A}\|_2 \leq \|\mathbb{A}\|_E \leq \sqrt{n} \|\mathbb{A}\|_2, \quad \forall \mathbb{A} \in \mathcal{M}_n. \quad (4.62)$$

De plus $\|\mathbb{I}\|_E = \sqrt{n}$.

**Théorème 4.35**

1. Soit $\|\bullet\|$ une norme matricielle subordonnée, et $\mathbb{B}$ une matrice vérifiant

$$\|\mathbb{B}\| < 1.$$

Alors la matrice $(\mathbb{I} + \mathbb{B})$ est inversible, et

$$\|(\mathbb{I} + \mathbb{B})^{-1}\| \leq \frac{1}{1 - \|\mathbb{B}\|}.$$

2. Si une matrice de la forme $(\mathbb{I} + \mathbb{B})$ est singulière, alors nécessairement

$$\|\mathbb{B}\| \geq 1$$

pour toute norme matricielle, subordonnée ou non.

4.3.3 Suites de vecteurs et de matrices

**Définition 4.36**

Soit V un espace vectoriel muni d'une norme $\|\bullet\|$, on dit qu'une suite $(\mathbf{v}_k)$ d'éléments de V **converge vers un élément** $\mathbf{v} \in V$, si

$$\lim_{k \rightarrow \infty} \|\mathbf{v}_k - \mathbf{v}\| = 0$$

et on écrit

$$\mathbf{v} = \lim_{k \rightarrow \infty} \mathbf{v}_k.$$



Théorème 4.37: admis

Soit $\mathbb{B}$ une matrice carrée. Les conditions suivantes sont équivalentes :

1. $\lim_{k \rightarrow \infty} \mathbb{B}^k = 0$,
2. $\lim_{k \rightarrow \infty} \mathbb{B}^k \mathbf{v} = 0$ pour tout vecteur $\mathbf{v}$,
3. $\rho(\mathbb{B}) < 1$,
4. $\|\mathbb{B}\| < 1$ pour au moins une norme matricielle subordonnée $\|\bullet\|$.



Théorème 4.38: admis

Soit $\mathbb{B}$ une matrice carrée, et $\|\bullet\|$ une norme matricielle quelconque. Alors

$$\lim_{k \rightarrow \infty} \|\mathbb{B}^k\|^{1/k} = \rho(\mathbb{B}).$$

4.4 Conditionnement d'un système linéaire

Pour la résolution numérique d'un système linéaire $\mathbb{A}\mathbf{x} = \mathbf{b}$, il est rare que les données $\mathbb{A}$ et $\mathbf{b}$ du problème ne soient pas entachées d'erreurs (aussi minimes soient-elles). La question qui se pose alors est de savoir si de petites perturbations sur les données ne peuvent pas entraîner des erreurs importantes sur le calcul de la solution.

Exemple de R.S. Wilson

Soient

$$\mathbb{A} = \begin{pmatrix} 10 & 7 & 8 & 7 \\ 7 & 5 & 6 & 5 \\ 8 & 6 & 10 & 9 \\ 7 & 5 & 9 & 10 \end{pmatrix}, \quad \Delta\mathbb{A} = \begin{pmatrix} 0 & 0 & \frac{1}{10} & \frac{1}{5} \\ \frac{2}{25} & \frac{1}{25} & 0 & 0 \\ 0 & -\frac{1}{50} & -\frac{11}{100} & 0 \\ -\frac{1}{100} & -\frac{1}{100} & 0 & -\frac{1}{50} \end{pmatrix}$$

et $\mathbf{b}^t = (32, 23, 33, 31)$, $(\Delta\mathbf{b})^t = (\frac{1}{100}, -\frac{1}{100}, \frac{1}{100}, -\frac{1}{100})$. Des calculs exacts donnent

$$\mathbb{A}\mathbf{x} = \mathbf{b} \iff \mathbf{x}^t = (1, 1, 1, 1)$$

$$\mathbb{A}\mathbf{u} = (\mathbf{b} + \Delta\mathbf{b}) \iff \mathbf{u}^t = \left(\frac{91}{50}, -\frac{9}{25}, \frac{27}{20}, \frac{79}{100}\right) \approx (1.8, -0.36, 1.3, 0.79)$$

$$(\mathbb{A} + \Delta\mathbb{A})\mathbf{v} = \mathbf{b} \iff \mathbf{v}^t = (-81, 137, -34, 22)$$

$$(\mathbb{A} + \Delta\mathbb{A})\mathbf{y} = (\mathbf{b} + \Delta\mathbf{b}) \iff \mathbf{y}^t = \left(-\frac{18283543}{461600}, \frac{31504261}{461600}, -\frac{3741501}{230800}, \frac{5235241}{461600}\right) \approx (-39.61, 68.25, -16.21, 11.34)$$

- 1 Il est clair sur cet exemple que de petites perturbations sur les données peuvent entraîner des erreurs
 2 importantes sur la solution exacte.
 3 On dit que le système linéaire précédent est **mal conditionné** ou qu'il a un **mauvais condition-**
 4 **nement** car il est sujet à de fortes variations de la solution pour de petites perturbations des données.
 5 A contrario, on dit qu'il est **bien conditionné** ou qu'il a un **bon conditionnement** si de petites
 6 perturbations des données n'entraînent qu'une variation *raisonnable* de la solution.
 7 Une nouvelle question : est-il possible de "mesurer" le **conditionnement** d'une matrice?

8 Définitions et résultats

Definition 4.39

Soit $\|\cdot\|$ une norme matricielle subordonnée, le conditionnement d'une matrice régulière A , associé à cette norme, est le nombre

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

Nous noterons $\text{cond}_p(A) = \|A\|_p \|A^{-1}\|_p$.

Proposition 4.40

Soit A une matrice régulière. On a les propriétés suivantes

1. $\forall \alpha \in \mathbb{K}^*, \text{cond}(\alpha A) = \text{cond}(A)$.
2. $\text{cond}_p(A) \geq 1, \forall p \in \llbracket 1, +\infty \rrbracket$.
3. $\text{cond}_2(A) = 1$ si et seulement si $A = \alpha Q$ avec $\alpha \in \mathbb{K}^*$ et Q matrice unitaire

11 *Proof.* Soit A une matrice régulière.

1. Soit $\alpha \in \mathbb{K}^*$, on a

$$\begin{aligned} \text{cond}(\alpha A) &\stackrel{\text{def}}{=} \|\alpha A\| \|(\alpha A)^{-1}\| = |\alpha| \|A\| \left\| \frac{1}{\alpha} A^{-1} \right\| \\ &= |\alpha| \|A\| \frac{1}{|\alpha|} \|A^{-1}\| = \|A\| \|A^{-1}\| \\ &= \text{cond}(A) \end{aligned}$$

2. On a $I = AA^{-1}$. Or pour toute norme subordonnée, on a $\|I\| = 1$ et donc $1 = \|AA^{-1}\| \leq \|A\| \|A^{-1}\| = \text{cond}(A)$.

3. **Admis.** Les valeurs singulières d'une matrice non pas été définies...

□

Théorème 4.41

Soit A une matrice inversible. Soient x et $x + \Delta x$ les solutions respectives de

$$Ax = b \quad \text{et} \quad A(x + \Delta x) = b + \Delta b.$$

Supposons $b \neq 0$, alors l'inégalité

$$\frac{\|\Delta x\|}{\|x\|} \leq \text{cond}(A) \frac{\|\Delta b\|}{\|b\|}$$

est satisfaite, et c'est la meilleure possible : pour une matrice A donnée, on peut trouver des vecteurs $b \neq 0$ et $\Delta b \neq 0$ tels qu'elle devienne une égalité.

Proof. On a

$$\mathbb{A}\mathbf{x} = \mathbf{b} \quad \text{et} \quad \mathbb{A}(\mathbf{x} + \Delta\mathbf{x}) = \mathbf{b} + \Delta\mathbf{b}$$

or $\mathbb{A}(\mathbf{x} + \Delta\mathbf{x}) = \mathbb{A}\mathbf{x} + \mathbb{A}\Delta\mathbf{x}$ et donc $\mathbb{A}\Delta\mathbf{x} = \Delta\mathbf{b}$ ou encore $\Delta\mathbf{x} = \mathbb{A}^{-1}\Delta\mathbf{b}$. Ceci donne

$$\|\Delta\mathbf{x}\| = \|\mathbb{A}^{-1}\Delta\mathbf{b}\| \leq \|\mathbb{A}^{-1}\| \|\Delta\mathbf{b}\|$$

et

$$\|\mathbf{b}\| = \|\mathbb{A}\mathbf{x}\| \leq \|\mathbb{A}\| \|\mathbf{x}\|.$$

On en déduit

$$\|\Delta\mathbf{x}\| \|\mathbf{b}\| \leq \|\mathbb{A}\| \|\mathbf{x}\| \|\mathbb{A}^{-1}\| \|\Delta\mathbf{b}\|.$$

Comme $\mathbf{b} \neq \mathbf{0}$, on a $\mathbf{x} = \mathbb{A}^{-1}\mathbf{b} \neq \mathbf{0}$ et donc, les normes étant positives,

$$\frac{\|\Delta\mathbf{x}\|}{\|\mathbf{x}\|} \leq \text{cond}(\mathbb{A}) \frac{\|\Delta\mathbf{b}\|}{\|\mathbf{b}\|}$$

D'après la Proposition 4.30, pour toute norme matricielle subordonnée il existe au moins un vecteur $\mathbf{u} \in \mathbb{K}^n \setminus \{0\}$ et un vecteur $\mathbf{v} \in \mathbb{K}^n \setminus \{0\}$ tel que

$$\|\mathbb{A}^{-1}\mathbf{u}\| = \|\mathbb{A}^{-1}\| \|\mathbf{u}\| \quad \text{et} \quad \|\mathbb{A}\mathbf{v}\| = \|\mathbb{A}\| \|\mathbf{v}\|.$$

En posant $\mathbf{b} = \mathbb{A}\mathbf{v}$ et $\Delta\mathbf{b} = \mathbf{u}$ on a bien égalité. □



Théorème 4.42

Soient $\mathbb{A}$ et $\mathbb{A} + \Delta\mathbb{A}$ deux matrices inversibles. Soient $\mathbf{x}$ et $\mathbf{x} + \Delta\mathbf{x}$ les solutions respectives de

$$\mathbb{A}\mathbf{x} = \mathbf{b} \quad \text{et} \quad (\mathbb{A} + \Delta\mathbb{A})(\mathbf{x} + \Delta\mathbf{x}) = \mathbf{b}.$$

Supposons $\mathbf{b} \neq \mathbf{0}$, alors on a

$$\frac{\|\Delta\mathbf{x}\|}{\|\mathbf{x} + \Delta\mathbf{x}\|} \leq \text{cond}(\mathbb{A}) \frac{\|\Delta\mathbb{A}\|}{\|\mathbb{A}\|}.$$

Proof. On a

$$\mathbb{A}\mathbf{x} = \mathbf{b} = (\mathbb{A} + \Delta\mathbb{A})(\mathbf{x} + \Delta\mathbf{x}) = \mathbb{A}\mathbf{x} + \mathbb{A}\Delta\mathbf{x} + \Delta\mathbb{A}(\mathbf{x} + \Delta\mathbf{x})$$

et donc

$$\mathbb{A}\Delta\mathbf{x} + \Delta\mathbb{A}(\mathbf{x} + \Delta\mathbf{x}) = \mathbf{0} \iff \Delta\mathbf{x} = -\mathbb{A}^{-1}\Delta\mathbb{A}(\mathbf{x} + \Delta\mathbf{x})$$

On en déduit alors

$$\|\Delta\mathbf{x}\| = \|\mathbb{A}^{-1}\Delta\mathbb{A}(\mathbf{x} + \Delta\mathbf{x})\| \leq \|\mathbb{A}^{-1}\| \|\Delta\mathbb{A}(\mathbf{x} + \Delta\mathbf{x})\| \leq \|\mathbb{A}^{-1}\| \|\Delta\mathbb{A}\| \|\mathbf{x} + \Delta\mathbf{x}\|$$

De plus, on a $\text{cond}(\mathbb{A}) \stackrel{\text{def}}{=} \|\mathbb{A}\| \|\mathbb{A}^{-1}\|$ et donc $\|\mathbb{A}^{-1}\| = \frac{\text{cond}(\mathbb{A})}{\|\mathbb{A}\|}$ ce qui donne

$$\|\Delta\mathbf{x}\| \leq \text{cond}(\mathbb{A}) \frac{\|\Delta\mathbb{A}\|}{\|\mathbb{A}\|} \|\mathbf{x} + \Delta\mathbf{x}\|.$$

Comme $\mathbf{b} \neq \mathbf{0}$, on a $\mathbf{x} + \Delta\mathbf{x} = (\mathbb{A} + \Delta\mathbb{A})^{-1}\mathbf{b} \neq \mathbf{0}$ et de l'inégalité précédente, on déduit alors

$$\frac{\|\Delta\mathbf{x}\|}{\|\mathbf{x} + \Delta\mathbf{x}\|} \leq \text{cond}(\mathbb{A}) \frac{\|\Delta\mathbb{A}\|}{\|\mathbb{A}\|}.$$

□

Remarque 4.43 Une matrice est donc **bien conditionnée** si son conditionnement est proche de 1.

4.5 Méthodes itératives

4.5.1 Principe

On souhaite résoudre le système linéaire $\mathbb{A}\mathbf{x} = \mathbf{b}$ par des **méthodes itératives**. Ces dernières consistent en la détermination d'une **matrice d'itération** $\mathbb{B}$ et d'un vecteur $\mathbf{c}$ tels que la suite de vecteurs $\mathbf{x}^{[k]}$ définie par

$$\mathbf{x}^{[k+1]} = \mathbb{B}\mathbf{x}^{[k]} + \mathbf{c}, \quad k \geq 0, \quad \mathbf{x}^{[0]} \text{ arbitraire}$$

soit telle que $\lim_{k \rightarrow \infty} \mathbf{x}^{[k]} = \tilde{\mathbf{x}}$ et $\tilde{\mathbf{x}}$ solution de $\mathbb{A}\mathbf{x} = \mathbf{b}$. Bien évidemment les matrices $\mathbb{B}$ et les vecteurs $\mathbf{c}$ dépendront de $\mathbb{A}$ et $\mathbf{b}$.

Nous allons étudier plusieurs méthodes itératives et pour chacune d'entre elles nous expliciterons la matrice d'itération $\mathbb{B}$ et le vecteur $\mathbf{c}$ associé.

4.5.2 Présentation des méthodes usuelles

Soit $\mathbb{A} \in \mathcal{M}_n(\mathbb{K})$ une matrice régulière, d'éléments diagonaux non-nuls, et $\mathbf{b} \in \mathbb{K}^n$. On note $\mathbb{D}$ la **matrice diagonale** telle que $\mathbb{D} = \text{diag}(\mathbb{A})$, $\mathbb{E}$ la **matrice triangulaire inférieure** à diagonale nulle définie par

$$\begin{cases} E_{ij} = 0, & i \leq j \\ E_{ij} = -A_{ij} & i > j \end{cases} \quad (4.63)$$

et $\mathbb{F}$ la **matrice triangulaire supérieure** à diagonale nulle définie par

$$\begin{cases} F_{ij} = 0, & i \geq j \\ F_{ij} = -A_{ij} & i < j \end{cases} \quad (4.64)$$

On a alors

$$\begin{aligned} \mathbb{A} &= \begin{pmatrix} A_{1,1} & 0 & \cdots & 0 \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \cdots & 0 & A_{n,n} \end{pmatrix} + \begin{pmatrix} 0 & \cdots & \cdots & 0 \\ A_{2,1} & \ddots & & \vdots \\ \vdots & \ddots & \ddots & \vdots \\ A_{n,1} & \cdots & A_{n,n-1} & 0 \end{pmatrix} + \begin{pmatrix} 0 & A_{1,2} & \cdots & A_{1,n} \\ \vdots & \ddots & \ddots & \vdots \\ \vdots & & \ddots & A_{n-1,n} \\ 0 & \cdots & \cdots & 0 \end{pmatrix} \\ &= \mathbb{D} - \mathbb{E} - \mathbb{F} = \begin{pmatrix} \ddots & & & -\mathbb{F} \\ & \mathbb{D} & & \\ -\mathbb{E} & & \ddots & \end{pmatrix} \end{aligned} \quad (4.65)$$

On rappelle que la i -ème équation du système linéaire $\mathbb{A}\mathbf{x} = \mathbf{b}$ s'écrit

$$b_i = \sum_{j=1}^n A_{i,j}x_j = \sum_{j=1}^{i-1} A_{i,j}x_j + A_{i,i}x_i + \sum_{j=i+1}^n A_{i,j}x_j. \quad (4.66)$$

Il faut aussi noter que la matrice diagonale $\mathbb{D}$ est inversible car les éléments diagonaux de $\mathbb{A}$ sont non nuls par hypothèse.

Méthode de Jacobi

Pour obtenir la méthode itérative de Jacobi, il suffit de *mettre*, dans la formule (4.5.2), l'itéré $k+1$ sur le terme diagonale et l'itéré k sur les autres termes pour avoir

$$b_i = \sum_{j=1}^{i-1} A_{i,j}x_j^{[k]} + A_{i,i}x_i^{[k+1]} + \sum_{j=i+1}^n A_{i,j}x_j^{[k]}.$$

ce qui donne

$$x_i^{[k+1]} = \frac{1}{A_{ii}} \left(b_i - \sum_{j=1, j \neq i}^n A_{ij}x_j^{[k]} \right) \quad \forall i \in \llbracket 1, n \rrbracket \quad (4.67)$$

Méthode de Gauss-Seidel

Pour obtenir la méthode itérative de Jacobi, il suffit de *mettre*, dans la formule (4.5.2), l'itéré $k + 1$ sur la partie *triangulaire inférieure* et l'itéré k sur les autres termes pour avoir

$$b_i = \sum_{j=1}^{i-1} A_{i,j} x_j^{[k+1]} + A_{i,i} x_i^{[k+1]} + \sum_{j=i+1}^n A_{i,j} x_j^{[k]}.$$

ce qui donne

$$x_i^{(k+1)} = \frac{1}{A_{ii}} \left(b_i - \sum_{j=1}^{i-1} A_{ij} x_j^{[k+1]} - \sum_{j=i+1}^n A_{ij} x_j^{[k]} \right) \quad \forall i \in \llbracket 1, n \rrbracket \quad (4.68)$$

Méthodes de relaxation

Ces méthodes sont basées sur un paramètre de relaxation $w \in \mathbb{R}^*$ et sont données par

$$x_i^{[k+1]} = w \hat{x}_i^{[k+1]} + (1 - w) x_i^{[k]}$$

où $\hat{x}_i^{[k+1]}$ est obtenu à partir de l'une des deux méthodes précédentes.

Avec la méthode de Jacobi

$$x_i^{(k+1)} = \frac{w}{A_{ii}} \left(b_i - \sum_{j=1, j \neq i}^n A_{ij} x_j^{[k]} \right) + (1 - w) x_i^{[k]} \quad \forall i \in \llbracket 1, n \rrbracket.$$

Avec la méthode de Gauss-Seidel

$$x_i^{[k+1]} = \frac{w}{A_{ii}} \left(b_i - \sum_{j=1}^{i-1} A_{ij} x_j^{[k+1]} - \sum_{j=i+1}^n A_{ij} x_j^{[k]} \right) + (1 - w) x_i^{[k]} \quad \forall i \in \llbracket 1, n \rrbracket$$

Cette dernière méthode de relaxation, utilisant la méthode de Gauss-Seidel, est appelée méthode S.O.R. (successive over relaxation)



Exercice 4.5.1

En écrivant $\mathbb{A}$ sous la forme $\mathbb{A} = \mathbb{D} - \mathbb{E} - \mathbb{F}$, montrer que les méthodes itératives de Jacobi, Gauss-Seidel et S.O.R. s'écrivent sous la forme $\mathbf{x}^{[k+1]} = \mathbb{B} \mathbf{x}^{[k]} + \mathbf{c}$, où l'on exprimera les matrices $\mathbb{B}$ et les vecteurs $\mathbf{c}$ en fonction de $\mathbb{D}$, $\mathbb{E}$, $\mathbb{F}$ et $\mathbf{b}$.

Correction Exercice 4.5.1 On rappelle que la i -ème équation du système linéaire $\mathbb{A} \mathbf{x} = \mathbf{b}$ s'écrit

$$b_i = \sum_{j=1}^n A_{i,j} x_j = \sum_{j=1}^{i-1} A_{i,j} x_j + A_{i,i} x_i + \sum_{j=i+1}^n A_{i,j} x_j.$$

et que ceci correspond à l'écriture matricielle

$$\mathbf{b} = \mathbb{D} \mathbf{x} - \mathbb{E} \mathbf{x} - \mathbb{F} \mathbf{x}$$

- Pour la **méthode de Jacobi** on a, $\forall i \in \llbracket 1, n \rrbracket$,

$$b_i = A_{i,i} x_i^{[k+1]} + \sum_{j=1}^{i-1} A_{i,j} x_j^{[k]} + \sum_{j=i+1}^n A_{i,j} x_j^{[k]}.$$

ce qui s'écrit matriciellement

$$\mathbf{b} = \mathbb{D} \mathbf{x}^{[k+1]} - \mathbb{E} \mathbf{x}^{[k]} - \mathbb{F} \mathbf{x}^{[k]}.$$

On a donc

$$\mathbb{D} \mathbf{x}^{[k+1]} = \mathbf{b} + (\mathbb{E} + \mathbb{F}) \mathbf{x}^{[k]}$$

Comme la matrice $\mathbb{D}$ est inversible, on obtient

$$\mathbf{x}^{[k+1]} = \mathbb{D}^{-1}(\mathbb{E} + \mathbb{F}) \mathbf{x}^{[k]} + \mathbb{D}^{-1} \mathbf{b}$$

La matrice d'itération de Jacobi est $\mathbb{B} = \mathbb{D}^{-1}(\mathbb{E} + \mathbb{F})$ et le vecteur $\mathbf{c} = \mathbb{D}^{-1} \mathbf{b}$.

- 1 • Pour la **méthode de Gauss-Seidel** on a, $\forall i \in \llbracket 1, n \rrbracket$,

$$b_i = A_{i,i}x_i^{[k+1]} + \sum_{j=1}^{i-1} A_{i,j}x_j^{[k+1]} + \sum_{j=i+1}^n A_{i,j}x_j^{[k]}.$$

- 2 ce qui s'écrit matriciellement

$$\mathbf{b} = \mathbb{D}\mathbf{x}^{[k+1]} - \mathbb{E}\mathbf{x}^{[k+1]} - \mathbb{F}\mathbf{x}^{[k]}.$$

- 3 On a donc

$$(\mathbb{D} - \mathbb{E})\mathbf{x}^{[k+1]} = \mathbf{b} + \mathbb{F}\mathbf{x}^{[k]}$$

- 4 Comme la matrice $\mathbb{D} - \mathbb{E}$ est inversible (matrice triangulaire inférieure d'éléments diagonaux non
5 nuls), on obtient

$$\mathbf{x}^{[k+1]} = (\mathbb{D} - \mathbb{E})^{-1}\mathbb{F}\mathbf{x}^{[k]} + (\mathbb{D} - \mathbb{E})^{-1}\mathbf{b}$$

- 6 La matrice d'itération de Gauss-Seidel est $\mathbb{B} = (\mathbb{D} - \mathbb{E})^{-1}\mathbb{F}$ et le vecteur $\mathbf{c} = (\mathbb{D} - \mathbb{E})^{-1}\mathbf{b}$.

- 7 • Pour la **méthode S.O.R.** on a, $\forall i \in \llbracket 1, n \rrbracket$,

$$x_i^{[k+1]} = \frac{w}{A_{ii}} \left(b_i - \sum_{j=1}^{i-1} A_{ij}x_j^{[k+1]} - \sum_{j=i+1}^n A_{ij}x_j^{[k]} \right) + (1-w)x_i^{[k]}$$

- 8 ce qui s'écrit aussi

$$\frac{A_{ii}}{w}x_i^{[k+1]} + \sum_{j=1}^{i-1} A_{ij}x_j^{[k+1]} = b_i - \sum_{j=i+1}^n A_{ij}x_j^{[k]} + \frac{1-w}{w}A_{ii}x_i^{[k]}$$

et matriciellement on obtient

$$\left(\frac{\mathbb{D}}{w} - \mathbb{E} \right) \mathbf{x}^{[k+1]} = \left(\frac{1-w}{w} \mathbb{D} + \mathbb{F} \right) \mathbf{x}^{[k]} + \mathbf{b}.$$

Comme la matrice $\left(\frac{\mathbb{D}}{w} - \mathbb{E} \right)$ est inversible (car triangulaire inférieure à éléments diagonaux non nuls), on a

$$\mathbf{x}^{[k+1]} = \left(\frac{\mathbb{D}}{w} - \mathbb{E} \right)^{-1} \left(\frac{1-w}{w} \mathbb{D} + \mathbb{F} \right) \mathbf{x}^{[k]} + \left(\frac{\mathbb{D}}{w} - \mathbb{E} \right)^{-1} \mathbf{b}$$

- 9 La matrice d'itération de S.O.R. est $\mathbb{B} = \left(\frac{\mathbb{D}}{w} - \mathbb{E} \right)^{-1} \left(\frac{1-w}{w} \mathbb{D} + \mathbb{F} \right)$ et le vecteur $\mathbf{c} = \left(\frac{\mathbb{D}}{w} - \mathbb{E} \right)^{-1} \mathbf{b}$.

10 $\diamond$

11

12



Proposition 4.44

Soit $\mathbb{A}$ une matrice régulière telle que tous ses éléments diagonaux soient non nuls. On note $\mathbb{D} = \text{diag}(\mathbb{A})$ et $\mathbb{E}, \mathbb{F}$, les matrices à diagonales nulles respectivement triangulaire inférieure et supérieure telles que $\mathbb{A} = \mathbb{D} - \mathbb{E} - \mathbb{F}$. On pose $\mathbb{L} = \mathbb{D}^{-1}\mathbb{E}$ et $\mathbb{U} = \mathbb{D}^{-1}\mathbb{F}$.

La matrice d'itération de la méthode de Jacobi, notée $\mathbb{J}$, est donnée par

$$\mathbb{J} = \mathbb{D}^{-1}(\mathbb{E} + \mathbb{F}) = \mathbb{L} + \mathbb{U}, \quad (4.69)$$

La matrice d'itération de la méthode S.O.R., notée $\mathcal{L}_w$, est donnée par

$$\mathcal{L}_w = \left(\frac{\mathbb{D}}{w} - \mathbb{E} \right)^{-1} \left(\frac{1-w}{w} \mathbb{D} + \mathbb{F} \right) = (\mathbb{I} - w\mathbb{L})^{-1} ((1-w)\mathbb{I} + w\mathbb{U}). \quad (4.70)$$

et elle vérifie

$$\rho(\mathcal{L}_w) \geq |w - 1|. \quad (4.71)$$

La matrice d'itération de Gauss-Seidel est $\mathcal{L}_1$ et elle correspond à

$$\mathcal{L}_1 = (\mathbb{D} - \mathbb{E})^{-1} \mathbb{F} = (\mathbb{I} - \mathbb{L})^{-1} \mathbb{U}. \quad (4.72)$$

Proof. Les résultats découlent de l'Exercice 4.5.1 pour l'écriture en fonction des matrices $\mathbb{D}$, $\mathbb{E}$ et $\mathbb{F}$. Pour l'écriture en fonction des matrices $\mathbb{L}$ et $\mathbb{U}$ seule l'équation (4.70) n'est pas forcément immédiate. On a vu que

$$\begin{aligned} \mathcal{L}_w &= \left(\frac{\mathbb{D}}{w} - \mathbb{E} \right)^{-1} \left(\frac{1-w}{w} \mathbb{D} + \mathbb{F} \right) = \left(\frac{\mathbb{D}}{w} [\mathbb{I} - w\mathbb{D}^{-1}\mathbb{E}] \right)^{-1} \left(\frac{1}{w} \mathbb{D} [(1-w)\mathbb{I} + w\mathbb{D}^{-1}\mathbb{F}] \right) \\ &= \left(\frac{\mathbb{D}}{w} [\mathbb{I} - w\mathbb{L}] \right)^{-1} \left(\frac{1}{w} \mathbb{D} [(1-w)\mathbb{I} + w\mathbb{U}] \right) \\ &= (\mathbb{I} - w\mathbb{L})^{-1} \left(\frac{\mathbb{D}}{w} \right)^{-1} \left(\frac{1}{w} \mathbb{D} [(1-w)\mathbb{I} + w\mathbb{U}] \right) \\ &= (\mathbb{I} - w\mathbb{L})^{-1} w\mathbb{D}^{-1} \left(\frac{1}{w} \mathbb{D} [(1-w)\mathbb{I} + w\mathbb{U}] \right) = (\mathbb{I} - w\mathbb{L})^{-1} ((1-w)\mathbb{I} + w\mathbb{U}) \end{aligned}$$

Il reste à démontrer l'inégalité (4.71). La matrice $\mathbb{L}$ est triangulaire inférieure à diagonale nulle car elle est le produit d'une matrice diagonale (et donc triangulaire inférieure) $\mathbb{D}^{-1}$ et d'une matrice triangulaire inférieure $\mathbb{E}$ à diagonale nulle. De même la matrice $\mathbb{U}$ est triangulaire supérieure à diagonale nulle.

On sait que le déterminant d'une matrice est égale aux produits de ses valeurs propres comptées avec leurs multiplicités. En notant n la dimension de la matrice $\mathcal{L}_w$, et en notant $\lambda_i(\mathcal{L}_w)$ ses n valeurs propres, on a donc

$$\det(\mathcal{L}_w) = \prod_{i=1}^n \lambda_i(\mathcal{L}_w).$$

Le rayon spectrale de $\mathcal{L}_w$, noté $\rho(\mathcal{L}_w)$, correspond au plus grand des modules des valeurs propres. On a alors

$$\rho(\mathcal{L}_w) = \max_{i \in [1, n]} |\lambda_i(\mathcal{L}_w)| \geq |\det(\mathcal{L}_w)|^{1/n}$$

De plus on a

$$\det(\mathcal{L}_w) = \det \left((\mathbb{I} - w\mathbb{L})^{-1} ((1-w)\mathbb{I} + w\mathbb{U}) \right) = \det \left((\mathbb{I} - w\mathbb{L})^{-1} \right) \det \left(((1-w)\mathbb{I} + w\mathbb{U}) \right)$$

La matrice $\mathbb{I} - w\mathbb{L}$ est triangulaire inférieure à diagonale unité donc son inverse aussi. On en déduit $\det \left((\mathbb{I} - w\mathbb{L})^{-1} \right) = 1$. La matrice $(1-w)\mathbb{I} + w\mathbb{U}$ est triangulaire supérieure avec tous ses éléments diagonaux valant $1-w$ et donc $\det \left(((1-w)\mathbb{I} + w\mathbb{U}) \right) = (1-w)^n$. On a alors $|\det(\mathcal{L}_w)| = |1-w|^n$ et

$$\rho(\mathcal{L}_w) \geq |\det(\mathcal{L}_w)|^{1/n} = |1-w|.$$

□

4.5.3 Etude de la convergence

**Théorème 4.45**

Soit $\mathbb{A}$ une matrice régulière décomposée sous la forme $\mathbb{A} = \mathbb{M} - \mathbb{N}$ avec $\mathbb{M}$ régulière. On pose

$$\mathbb{B} = \mathbb{M}^{-1}\mathbb{N} \quad \text{et} \quad \mathbf{c} = \mathbb{M}^{-1}\mathbf{b}.$$

Alors la suite définie par

$$\mathbf{x}^{[0]} \in \mathbb{K}^n \quad \text{et} \quad \mathbf{x}^{[k+1]} = \mathbb{B}\mathbf{x}^{[k]} + \mathbf{c}$$

converge vers $\bar{\mathbf{x}} = \mathbb{A}^{-1}\mathbf{b}$ quelque soit $\mathbf{x}^{[0]}$ si et seulement si $\rho(\mathbb{B}) < 1$.

Proof. Comme $\bar{\mathbf{x}} = \mathbb{A}^{-1}\mathbf{b}$ (sans présupposer la convergence) on a $\mathbb{M}\bar{\mathbf{x}} = \mathbb{N}\bar{\mathbf{x}} + \mathbf{b}$ et alors

$$\bar{\mathbf{x}} = \mathbb{M}^{-1}\mathbb{N}\bar{\mathbf{x}} + \mathbb{M}^{-1}\mathbf{b} = \mathbb{B}\bar{\mathbf{x}} + \mathbf{c}$$

On obtient donc

$$\bar{\mathbf{x}} - \mathbf{x}^{[k+1]} = \mathbb{B}(\bar{\mathbf{x}} - \mathbf{x}^{[k]})$$

Or la suite $\mathbf{x}^{[k]}$ converge vers $\bar{\mathbf{x}}$ si et seulement si la suite $\mathbf{e}^{[k]} \stackrel{\text{def}}{=} \bar{\mathbf{x}} - \mathbf{x}^{[k]}$ converge vers $\mathbf{0}$. On a

$$\mathbf{e}^{[k]} = \mathbb{B}^k \mathbf{e}^{[0]}, \quad \forall k \in \mathbb{N}.$$

D'après le Théorème 4.37, page 111, on a $\lim_{k \rightarrow +\infty} \mathbb{B}^k \mathbf{e}^{[0]} = \mathbf{0}$, $\forall \mathbf{e}^{[0]} \in \mathbb{K}^n$ si et seulement si $\rho(\mathbb{B}) < 1$. $\square$

**Corollaire 4.46**

Soit $\mathbb{A}$ une matrice vérifiant $A_{i,i} \neq 0 \, \forall i$. Une condition nécessaire de convergence pour la méthode S.O.R. est que $0 < w < 2$.

Proof. Soit $\mathbb{A}$ une matrice vérifiant $A_{i,i} \neq 0 \, \forall i$. Une condition nécessaire de convergence pour la méthode S.O.R. est que

$$0 < w < 2.$$

On a vu en (voir Proposition 4.44, page 117) que $\rho(\mathcal{L}_w) \geq |w - 1|$. Donc si $\rho(\mathcal{L}_w) \geq 1$, la non-convergence est certaine d'après le Théorème 4.37, page 111. Une condition nécessaire (mais non suffisante) de convergence est que $|w - 1| < 1$ i.e. $w \in]0, 2[$. $\square$

**Théorème 4.47**

Soit $\mathbb{A}$ une matrice à diagonale strictement dominante ou une matrice inversible à diagonale fortement dominante alors

- la méthode de Jacobi est convergente,
- si $w \in]0, 1]$ la méthode de Relaxation est convergente.

Proof. voir [7], vol.2, Théorème 19 et 20, pages 346 à 349. $\square$

**Théorème 4.48**

Soit $\mathbb{A}$ une matrice hermitienne inversible en décomposée en $\mathbb{A} = \mathbb{M} - \mathbb{N}$ où $\mathbb{M}$ est inversible. Soit $\mathbb{B} = \mathbb{I} - \mathbb{M}^{-1}\mathbb{A}$, la matrice de l'itération. Supposons que $\mathbb{M}^* + \mathbb{N}$ (qui est hermitienne) soit définie positive. Alors $\rho(\mathbb{B}) < 1$ si et seulement si $\mathbb{A}$ est définie positive.

Proof. (voir Exercice 4.5.2, page 119) □



Exercice 4.5.2

Soit $\mathbb{A} \in \mathcal{M}_{n,n}(\mathbb{C})$ une matrice hermitienne inversible décomposée en $\mathbb{A} = \mathbb{M} - \mathbb{N}$ où $\mathbb{M}$ est inversible. On note $\mathbb{B} = \mathbb{I} - \mathbb{M}^{-1}\mathbb{A}$.

Q. 1 Montrer que la matrice $\mathbb{M}^* + \mathbb{N}$ est hermitienne.

On suppose maintenant que $\mathbb{M}^* + \mathbb{N}$ est définie positive.

Q. 2 Soit $\mathbf{x}$ un vecteur quelconque de $\mathbb{C}^n$ et $\mathbf{y} = \mathbb{B}\mathbf{x}$.

1. Montrer que

$$\langle \mathbf{x}, \mathbb{A}\mathbf{x} \rangle - \langle \mathbf{y}, \mathbb{A}\mathbf{y} \rangle = \langle \mathbf{x}, \mathbb{A}\mathbb{M}^{-1}\mathbb{A}\mathbf{x} \rangle + \langle \mathbb{M}^{-1}\mathbb{A}\mathbf{x}, \mathbb{A}\mathbf{x} \rangle - \langle \mathbb{M}^{-1}\mathbb{A}\mathbf{x}, \mathbb{A}\mathbb{M}^{-1}\mathbb{A}\mathbf{x} \rangle \quad (4.73)$$

et

$$\mathbf{x} - \mathbf{y} = \mathbb{M}^{-1}\mathbb{A}\mathbf{x}. \quad (4.74)$$

2. En déduire que

$$\langle \mathbf{x}, \mathbb{A}\mathbf{x} \rangle - \langle \mathbf{y}, \mathbb{A}\mathbf{y} \rangle = \langle (\mathbf{x} - \mathbf{y}), (\mathbb{M}^* + \mathbb{N})(\mathbf{x} - \mathbf{y}) \rangle. \quad (4.75)$$

Q. 3 Montrer que si $\mathbb{A}$ est définie positive alors $\rho(\mathbb{B}) < 1$.

Q. 4 Démontrer par l'absurde que si $\rho(\mathbb{B}) < 1$ alors $\mathbb{A}$ est définie positive.



Théorème 4.49

Soit $\mathbb{A}$ une matrice hermitienne définie positive, alors la méthode de relaxation converge si et seulement si $w \in]0, 2[$.

Proof. voir [7], vol.2, Corollaire 24, page 351. □

4.5.4 Algorithmes

Nous allons écrire des algorithmes pour chacune des méthodes itératives proposées. L'objectif n'est pas d'écrire des algorithmes "optimisés" mais de voir la méthodologie permettant la construction d'algorithmes simples et fonctionnels.

Principe de base

Les méthodes itératives pour la résolution d'un système linéaire $\mathbb{A}\mathbf{x} = \mathbf{b}$ s'écrivent sous la forme

$$\mathbf{x}^{[0]} \in \mathbb{K}^n \text{ et } \mathbf{x}^{[k+1]} = \mathbb{B}\mathbf{x}^{[k]} + \mathbf{c}$$

La matrice d'itération $\mathbb{B}$ et le vecteur $\mathbf{c}$ sont construits de telle sorte que si la suite $(\mathbf{x}^{[k]})_{k \in \mathbb{N}}$ converge vers $\bar{\mathbf{x}}$ alors $\bar{\mathbf{x}}$ est aussi solution du système linéaire $\mathbb{A}\mathbf{x} = \mathbf{b}$.

Comme pour les méthodes de point fixe, La convergence de ces méthodes n'est pas assurées et si il y a convergence le nombre d'itération nécessaire n'est (à priori) pas connu. C'est pourquoi algorithmiquement on utilise une boucle **Tantque**. On définit alors un **nombre maximum d'itérations** au delà duquel les calculs itératifs sont stoppés et une valeur $\varepsilon > 0$ permettant d'arrêter les calculs lorsque $\mathbf{x}^{[k]}$ est suffisamment proche de $\bar{\mathbf{x}}$. Pour être plus précis, on note $\mathbf{r}^{[k]} = \mathbf{b} - \mathbb{A}\mathbf{x}^{[k]}$ le résidu. On a alors

$$\mathbf{r}^{[k]} = \mathbb{A}\bar{\mathbf{x}} - \mathbb{A}\mathbf{x}^{[k]} = \mathbb{A}\mathbf{e}^{[k]}$$

1 Si on prend comme critère d'arrêt

$$\frac{\|\mathbf{r}^{[k]}\|}{\|\mathbf{b}\|} \leq \varepsilon$$

alors

$$\|\mathbf{e}^{[k]}\| = \|\mathbb{A}^{-1}\mathbf{r}^{[k]}\| \leq \|\mathbb{A}^{-1}\| \|\mathbf{r}^{[k]}\| \leq \varepsilon \|\mathbb{A}^{-1}\| \|\mathbf{b}\| = \varepsilon \|\mathbb{A}^{-1}\| \|\mathbb{A}\bar{\mathbf{x}}\| \leq \varepsilon \|\mathbb{A}^{-1}\| \|\mathbb{A}\| \|\bar{\mathbf{x}}\| = \varepsilon \text{cond}(\mathbb{A}) \|\bar{\mathbf{x}}\|$$

2 et donc

$$\frac{\|\mathbf{e}^{[k]}\|}{\|\bar{\mathbf{x}}\|} \leq \varepsilon \text{cond}(\mathbb{A})$$

3 Pour éviter des soucis lorsque $\|\mathbf{b}\|$ est proche de zéro, on peut utiliser comme critère d'arrêt de con-
4 vergence

$$\frac{\|\mathbf{r}^{[k]}\|}{\|\mathbf{b}\| + 1} \leq \varepsilon.$$

5 Comme vu avec les méthodes itératives de type point fixe (voir ??), il n'est pas forcément utile de
6 stocker l'intégralité des termes de la suite $\mathbf{x}^{[k]}$ puisque seul le "dernier" nous intéresse. On a alors
7 L'Algorithme 4.19 correspondant à un algorithme itératif générique sans stockage des valeurs intermédi-
8 aires.

Algorithme 4.19 Méthode itérative pour la résolution d'un système linéaire $\mathbb{A}\mathbf{x} = \mathbf{b}$

Données :

$\mathbb{A}$: matrice de $\mathcal{M}_n(\mathbb{K})$,
 $\mathbf{b}$: vecteur de $\mathbb{K}^n$,
 $\mathbf{x}^0$: vecteur initial de $\mathbb{K}^n$,
 ε : la tolérance, $\varepsilon \in \mathbb{R}^+$,
 kmax : nombre maximum d'itérations, $\text{kmax} \in \mathbb{N}^*$

Résultat :

$\mathbf{x}^{\text{tol}}$: un vecteur de $\mathbb{K}^n$ si convergence, sinon $\emptyset$

9

```

1:  $k \leftarrow 0, \mathbf{x}^{\text{tol}} \leftarrow \emptyset$ 
2:  $\mathbf{x} \leftarrow \mathbf{x}^0, \mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x}$ ,
3:  $\text{tol} \leftarrow \varepsilon(\|\mathbf{b}\| + 1)$ 
4: Tantque  $\|\mathbf{r}\| > \text{tol}$  et  $k \leq \text{kmax}$  faire
5:    $k \leftarrow k + 1$ 
6:    $\mathbf{p} \leftarrow \mathbf{x}$  ▷  $\mathbf{p}$  contient le vecteur précédent
7:    $\mathbf{x} \leftarrow$  calcul de l'itérée suivante en fonction de  $\mathbf{p}, \mathbb{A}, \mathbf{b}, \dots$ 
8:    $\mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x}$ ,
9: Fin Tantque
10: Si  $\|\mathbf{r}\| \leq \text{tol}$  alors ▷ Convergence
11:    $\mathbf{x}^{\text{tol}} \leftarrow \mathbf{x}$ 
12: Fin Si
```

10 Méthode de Jacobi

11 Pour Jacobi, la suite des itérées est définie par

$$x_i^{[k+1]} = \frac{1}{A_{ii}} \left(b_i - \sum_{j=1, j \neq i}^n A_{ij} x_j^{[k]} \right), \quad \forall i \in \llbracket 1, n \rrbracket.$$

12 Cette formule donne explicitement les composantes du vecteur $\mathbf{x}^{[k+1]}$ en fonction de la matrice $\mathbb{A}$, et des
13 vecteurs $\mathbf{b}$ et $\mathbf{x}^{[k]}$. A partir de l'Algorithme 4.19, on va construire par raffinements successifs la **RSLJACOBI**
14 donnée dans l'Algorithme 4.20.

Algorithme 4.20 $\mathcal{R}_0$

```

1:  $k \leftarrow 0, \mathbf{x}^{\text{tol}} \leftarrow \emptyset$ 
2:  $\mathbf{x} \leftarrow \mathbf{x}^0, \mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x},$ 
3:  $\text{tol} \leftarrow \varepsilon(\|\mathbf{b}\| + 1)$ 
4: Tantque  $\|\mathbf{r}\| > \text{tol}$  et  $k \leq k_{\max}$  faire
5:    $k \leftarrow k + 1$ 
6:    $\mathbf{p} \leftarrow \mathbf{x}$ 
7:
8:    $\mathbf{x} \leftarrow \text{calcul par Jacobi}$ 
9:    $\mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x},$ 
10: Fin Tantque
11: Si  $\|\mathbf{r}\| \leq \text{tol}$  alors  $\triangleright$  Convergence
12:    $\mathbf{x}^{\text{tol}} \leftarrow \mathbf{x}$ 
13: Fin Si

```

Algorithme 4.20 $\mathcal{R}_1$

```

1:  $k \leftarrow 0, \mathbf{x}^{\text{tol}} \leftarrow \emptyset$ 
2:  $\mathbf{x} \leftarrow \mathbf{x}^0, \mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x},$ 
3:  $\text{tol} \leftarrow \varepsilon(\|\mathbf{b}\| + 1)$ 
4: Tantque  $\|\mathbf{r}\| > \text{tol}$  et  $k \leq k_{\max}$  faire
5:    $k \leftarrow k + 1$ 
6:    $\mathbf{p} \leftarrow \mathbf{x}$ 
7:   Pour  $i \leftarrow 1$  à  $n$  faire
8:      $x_i \leftarrow \frac{1}{A_{ii}} \left( b_i - \sum_{j=1, j \neq i}^n A_{ij} p_j \right)$ 
9:   Fin Pour
10:   $\mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x},$ 
11: Fin Tantque
12: Si  $\|\mathbf{r}\| \leq \text{tol}$  alors  $\triangleright$  Convergence
13:    $\mathbf{x}^{\text{tol}} \leftarrow \mathbf{x}$ 
14: Fin Si

```

1

Algorithme 4.20 $\mathcal{R}_1$

```

1:  $k \leftarrow 0, \mathbf{x}^{\text{tol}} \leftarrow \emptyset$ 
2:  $\mathbf{x} \leftarrow \mathbf{x}^0, \mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x},$ 
3:  $\text{tol} \leftarrow \varepsilon(\|\mathbf{b}\| + 1)$ 
4: Tantque  $\|\mathbf{r}\| > \text{tol}$  et  $k \leq k_{\max}$  faire
5:    $k \leftarrow k + 1$ 
6:    $\mathbf{p} \leftarrow \mathbf{x}$ 
7:   Pour  $i \leftarrow 1$  à  $n$  faire
8:
9:      $x_i \leftarrow \frac{1}{A_{ii}} \left( b_i - \sum_{j=1, j \neq i}^n A_{ij} p_j \right)$ 
10:  Fin Pour
11:   $\mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x},$ 
12: Fin Tantque
13: Si  $\|\mathbf{r}\| \leq \text{tol}$  alors
14:    $\mathbf{x}^{\text{tol}} \leftarrow \mathbf{x}$ 
15: Fin Si

```

Algorithme 4.20 $\mathcal{R}_2$

```

1:  $k \leftarrow 0, \mathbf{x}^{\text{tol}} \leftarrow \emptyset$ 
2:  $\mathbf{x} \leftarrow \mathbf{x}^0, \mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x},$ 
3:  $\text{tol} \leftarrow \varepsilon(\|\mathbf{b}\| + 1)$ 
4: Tantque  $\|\mathbf{r}\| > \text{tol}$  et  $k \leq k_{\max}$  faire
5:    $k \leftarrow k + 1$ 
6:    $\mathbf{p} \leftarrow \mathbf{x}$ 
7:   Pour  $i \leftarrow 1$  à  $n$  faire
8:      $S \leftarrow 0$ 
9:     Pour  $j \leftarrow 1$  à  $n$  ( $j \neq i$ ) faire
10:       $S \leftarrow S + A_{i,j} p_j$ 
11:     Fin Pour
12:      $x_i \leftarrow \frac{1}{A_{ii}} (b_i - S)$ 
13:   Fin Pour
14:   $\mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x},$ 
15: Fin Tantque
16: Si  $\|\mathbf{r}\| \leq \text{tol}$  alors
17:    $\mathbf{x}^{\text{tol}} \leftarrow \mathbf{x}$ 
18: Fin Si

```

2

On peut alors écrire la fonction **RSLJACOBI** :

3

Algorithme 4.20 Méthode itérative de Jacobi pour la résolution d'un système linéaire $\mathbb{A}\mathbf{x} = \mathbf{b}$

Données :

$\mathbb{A}$: matrice de $\mathcal{M}_n(\mathbb{K})$,
 $\mathbf{b}$: vecteur de $\mathbb{K}^n$,
 $\mathbf{x}^0$: vecteur initial de $\mathbb{K}^n$,
 ε : la tolérance, $\varepsilon \in \mathbb{R}^+$,
 kmax : nombre maximum d'itérations, $\text{kmax} \in \mathbb{N}^*$

Résultat :

$\mathbf{X}$: un vecteur de $\mathbb{K}^n$

```

1 2: Fonction  $\mathbf{X} \leftarrow \text{RSLJACOBI}(\mathbb{A}, \mathbf{b}, \mathbf{x}^0, \varepsilon, \text{kmax})$ 
3:    $k \leftarrow 0, \mathbf{X} \leftarrow \emptyset$ 
4:    $\mathbf{x} \leftarrow \mathbf{x}^0, \mathbf{r} \leftarrow \mathbf{b} - \mathbb{A} * \mathbf{x},$ 
5:   Tantque  $\|\mathbf{r}\| > \text{tol}$  et  $k \leq \text{kmax}$  faire
6:      $k \leftarrow k + 1$ 
7:      $\mathbf{p} \leftarrow \mathbf{x}$ 
8:     Pour  $i \leftarrow 1$  à  $n$  faire
9:        $\mathbf{S} \leftarrow 0$ 
10:      Pour  $j \leftarrow 1$  à  $n$  ( $j \neq i$ ) faire
11:         $\mathbf{S} \leftarrow \mathbf{S} + A(i, j) * p(j)$ 
12:      Fin Pour
13:       $x(i) \leftarrow (b(i) - \mathbf{S}) / A(i, i)$ 
14:    Fin Pour
15:     $\mathbf{r} \leftarrow \mathbf{b} - \mathbb{A} * \mathbf{x},$ 
16:  Fin Tantque
17:  Si  $\|\mathbf{r}\| \leq \text{tol}$  alors
18:     $\mathbf{X} \leftarrow \mathbf{x}$ 
19:  Fin Si
20: Fin Fonction

```

2 Méthode de Gauss-Seidel

3 Pour Gauss-Seidel, la suite des itérées est définie par

$$x_i^{(k+1)} = \frac{1}{A_{ii}} \left(b_i - \sum_{j=1}^{i-1} A_{ij} x_j^{[k+1]} - \sum_{j=i+1}^n A_{ij} x_j^{[k]} \right) \quad \forall i \in \llbracket 1, n \rrbracket$$

4 Cette formule donne explicitement la composante i du vecteur $\mathbf{x}^{[k+1]}$ en fonction de la matrice $\mathbb{A}$, et des
5 vecteurs $\mathbf{b}$ et $\mathbf{x}^{[k]}$, mais aussi des $i - 1$ premières composantes de $\mathbf{x}^{[k+1]}$. Contrairement à la méthode de
6 Jacobi, il est impératif de calculer successivement $x_1^{[k+1]}, x_2^{[k+1]}, \dots$. A partir de l'Algorithme 4.19, on va
7 construire par raffinements successifs la **RSLGAUSSSEIDEL** donnée dans l'Algorithme 4.21.

Algorithme 4.21 $\mathcal{R}_0$

```

1:  $k \leftarrow 0, \mathbf{x}^{\text{tol}} \leftarrow \emptyset$ 
2:  $\mathbf{x} \leftarrow \mathbf{x}^0, \mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x},$ 
3:  $\text{tol} \leftarrow \varepsilon(\|\mathbf{b}\| + 1)$ 
4: Tantque  $\|\mathbf{r}\| > \text{tol}$  et  $k \leq k_{\max}$  faire
5:    $k \leftarrow k + 1$ 
6:    $\mathbf{p} \leftarrow \mathbf{x}$ 
7:
8:    $\mathbf{x} \leftarrow$  calcul par Gauss-Seidel
9:    $\mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x},$ 
10: Fin Tantque
11: Si  $\|\mathbf{r}\| \leq \text{tol}$  alors
12:    $\mathbf{x}^{\text{tol}} \leftarrow \mathbf{x}$ 
13: Fin Si

```

Algorithme 4.21 $\mathcal{R}_1$

```

1:  $k \leftarrow 0, \mathbf{x}^{\text{tol}} \leftarrow \emptyset$ 
2:  $\mathbf{x} \leftarrow \mathbf{x}^0, \mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x},$ 
3:  $\text{tol} \leftarrow \varepsilon(\|\mathbf{b}\| + 1)$ 
4: Tantque  $\|\mathbf{r}\| > \text{tol}$  et  $k \leq k_{\max}$  faire
5:    $k \leftarrow k + 1$ 
6:    $\mathbf{p} \leftarrow \mathbf{x}$ 
7:   Pour  $i \leftarrow 1$  à  $n$  faire
8:      $x_i \leftarrow \frac{1}{A_{ii}} \left( b_i - \sum_{j=1}^{i-1} A_{ij}x_j - \sum_{j=i+1}^n A_{ij}p_j \right)$ 
9:   Fin Pour
10:   $\mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x},$ 
11: Fin Tantque
12: Si  $\|\mathbf{r}\| \leq \text{tol}$  alors
13:    $\mathbf{x}^{\text{tol}} \leftarrow \mathbf{x}$ 
14: Fin Si

```

1
2Algorithme 4.21 $\mathcal{R}_1$

```

1:  $k \leftarrow 0, \mathbf{x}^{\text{tol}} \leftarrow \emptyset$ 
2:  $\mathbf{x} \leftarrow \mathbf{x}^0, \mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x},$ 
3:  $\text{tol} \leftarrow \varepsilon(\|\mathbf{b}\| + 1)$ 
4: Tantque  $\|\mathbf{r}\| > \text{tol}$  et  $k \leq k_{\max}$  faire
5:    $k \leftarrow k + 1$ 
6:    $\mathbf{p} \leftarrow \mathbf{x}$ 
7:   Pour  $i \leftarrow 1$  à  $n$  faire
8:
9:      $x_i \leftarrow \frac{1}{A_{ii}} \left( b_i - \sum_{j=1}^{i-1} A_{ij}x_j - \sum_{j=i+1}^n A_{ij}p_j \right)$ 
10:  Fin Pour
11:   $\mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x},$ 
12: Fin Tantque
13: Si  $\|\mathbf{r}\| \leq \text{tol}$  alors
14:    $\mathbf{x}^{\text{tol}} \leftarrow \mathbf{x}$ 
15: Fin Si

```

Algorithme 4.21 $\mathcal{R}_2$

```

1:  $k \leftarrow 0, \mathbf{x}^{\text{tol}} \leftarrow \emptyset$ 
2:  $\mathbf{x} \leftarrow \mathbf{x}^0, \mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x},$ 
3:  $\text{tol} \leftarrow \varepsilon(\|\mathbf{b}\| + 1)$ 
4: Tantque  $\|\mathbf{r}\| > \text{tol}$  et  $k \leq k_{\max}$  faire
5:    $k \leftarrow k + 1$ 
6:    $\mathbf{p} \leftarrow \mathbf{x}$ 
7:   Pour  $i \leftarrow 1$  à  $n$  faire
8:      $S \leftarrow 0$ 
9:     Pour  $j \leftarrow 1$  à  $i-1$  faire
10:       $S \leftarrow S + A_{i,j}x_j$ 
11:     Fin Pour
12:     Pour  $j \leftarrow i+1$  à  $n$  faire
13:       $S \leftarrow S + A_{i,j}p_j$ 
14:     Fin Pour
15:      $x_i \leftarrow \frac{1}{A_{ii}} (b_i - S)$ 
16:   Fin Pour
17:   $\mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x},$ 
18: Fin Tantque
19: Si  $\|\mathbf{r}\| \leq \text{tol}$  alors
20:    $\mathbf{x}^{\text{tol}} \leftarrow \mathbf{x}$ 
21: Fin Si

```

3

On peut alors écrire la fonction **RSLGAUSSSEIDEL** :

4

Algorithme 4.21 Méthode itérative de Gauss-Seidel pour la résolution d'un système linéaire $\mathbb{A}\mathbf{x} = \mathbf{b}$

Données :

$\mathbb{A}$: matrice de $\mathcal{M}_n(\mathbb{K})$,
 $\mathbf{b}$: vecteur de $\mathbb{K}^n$,
 $\mathbf{x}^0$: vecteur initial de $\mathbb{K}^n$,
 ε : la tolérance, $\varepsilon \in \mathbb{R}^+$,
 kmax : nombre maximum d'itérations, $\text{kmax} \in \mathbb{N}^*$

Résultat :

$\mathbf{X}$: un vecteur de $\mathbb{K}^n$

```

1: Fonction  $\mathbf{X} \leftarrow \text{RSLGAUSSSEIDEL} (\mathbb{A}, \mathbf{b}, \mathbf{x}^0, \varepsilon, \text{kmax})$ 
2:    $k \leftarrow 0, \mathbf{X} \leftarrow \emptyset$ 
3:    $\mathbf{x} \leftarrow \mathbf{x}^0, \mathbf{r} \leftarrow \mathbf{b} - \mathbb{A} * \mathbf{x}$ ,
4:    $\text{tol} \leftarrow \varepsilon(\|\mathbf{b}\| + 1)$ 
5:   Tantque  $\|\mathbf{r}\| > \text{tol}$  et  $k \leq \text{kmax}$  faire
6:      $k \leftarrow k + 1$ 
7:      $\mathbf{p} \leftarrow \mathbf{x}$ 
8:     Pour  $i \leftarrow 1$  à  $n$  faire
9:        $\mathbf{S} \leftarrow 0$ 
10:      Pour  $j \leftarrow 1$  à  $i - 1$  faire
11:         $\mathbf{S} \leftarrow \mathbf{S} + A(i, j) * x(j)$ 
12:      Fin Pour
13:      Pour  $j \leftarrow i + 1$  à  $n$  faire
14:         $\mathbf{S} \leftarrow \mathbf{S} + A(i, j) * p(j)$ 
15:      Fin Pour
16:       $x(i) \leftarrow (b(i) - \mathbf{S}) / A(i, i)$ 
17:    Fin Pour
18:     $\mathbf{r} \leftarrow \mathbf{b} - \mathbb{A} * \mathbf{x}$ ,
19:  Fin Tantque
20:  Si  $\|\mathbf{r}\| \leq \text{tol}$  alors
21:     $\mathbf{X} \leftarrow \mathbf{x}$ 
22:  Fin Si
23: Fin Fonction

```

2 Jeux algorithmiques

3 On a bien sûr noté que les fonctions `RSLJACOBI` et `RSLGAUSSSEIDEL` ont la même ossature puisque
 4 toutes deux basées sur l'Algorithme 4.19 générique. En effet seule la transcription de la ligne 7 de
 5 l'Algorithme 4.19 diffère :

```

Fonction  $X \leftarrow \text{RSLJACOBI} (\mathbb{A}, \mathbf{b}, \mathbf{x}^0, \varepsilon, \text{kmax})$ 
 $k \leftarrow 0, X \leftarrow \emptyset$ 
 $\mathbf{x} \leftarrow \mathbf{x}^0, \mathbf{r} \leftarrow \mathbf{b} - \mathbb{A} * \mathbf{x},$ 
 $\text{tol} \leftarrow \varepsilon(\|\mathbf{b}\| + 1)$ 
Tantque  $\|\mathbf{r}\| > \text{tol}$  et  $k \leq \text{kmax}$  faire
   $k \leftarrow k + 1$ 
   $\mathbf{p} \leftarrow \mathbf{x}$ 
  Pour  $i \leftarrow 1$  à  $n$  faire
     $S \leftarrow 0$ 
    Pour  $j \leftarrow 1$  à  $n$  ( $j \neq i$ ) faire
       $S \leftarrow S + A(i, j) * p(j)$ 
    Fin Pour
     $x(i) \leftarrow (b(i) - S) / A(i, i)$ 
  Fin Pour
   $\mathbf{r} \leftarrow \mathbf{b} - \mathbb{A} * \mathbf{x},$ 
Fin Tantque
Si  $\|\mathbf{r}\| \leq \text{tol}$  alors
   $X \leftarrow \mathbf{x}$ 
Fin Si
Fin Fonction

```

```

Fonction  $X \leftarrow \text{RSLGAUSSSEIDEL} (\mathbb{A}, \mathbf{b}, \mathbf{x}^0, \varepsilon, \text{kmax})$ 
 $k \leftarrow 0, X \leftarrow \emptyset$ 
 $\mathbf{x} \leftarrow \mathbf{x}^0, \mathbf{r} \leftarrow \mathbf{b} - \mathbb{A} * \mathbf{x},$ 
 $\text{tol} \leftarrow \varepsilon(\|\mathbf{b}\| + 1)$ 
Tantque  $\|\mathbf{r}\| > \text{tol}$  et  $k \leq \text{kmax}$  faire
   $k \leftarrow k + 1$ 
   $\mathbf{p} \leftarrow \mathbf{x}$ 
  Pour  $i \leftarrow 1$  à  $n$  faire
     $S \leftarrow 0$ 
    Pour  $j \leftarrow 1$  à  $i - 1$  faire
       $S \leftarrow S + A(i, j) * x(j)$ 
    Fin Pour
    Pour  $j \leftarrow i + 1$  à  $n$  faire
       $S \leftarrow S + A(i, j) * p(j)$ 
    Fin Pour
     $x(i) \leftarrow (b(i) - S) / A(i, i)$ 
  Fin Pour
   $\mathbf{r} \leftarrow \mathbf{b} - \mathbb{A} * \mathbf{x},$ 
Fin Tantque
Si  $\|\mathbf{r}\| \leq \text{tol}$  alors
   $X \leftarrow \mathbf{x}$ 
Fin Si
Fin Fonction

```

On va écrire les fonctions algorithmiques **ITERJACOBI** et **ITERGAUSSSEIDEL** permettant le calcul d'une itérée respectivement pour les méthodes de Jacobi et Gauss-Seidel (voir Algorithmes 4.22 et 4.23).

Algorithme 4.22 Itération de Jacobi : calcul de $\mathbf{x}$ tel que

$$x_i = \frac{1}{A_{ii}} \left(b_i - \sum_{j=1, j \neq i}^n A_{ij} y_j \right), \quad \forall i \in \llbracket 1, n \rrbracket.$$

Données :

$\mathbf{A}$: matrice de $\mathcal{M}_n(\mathbb{K})$,
 $\mathbf{b}$: vecteur de $\mathbb{K}^n$,
 $\mathbf{y}$: vecteur de $\mathbb{K}^n$,

Résultat :

$\mathbf{x}$: un vecteur de $\mathbb{K}^n$

```

1: Fonction  $\mathbf{x} \leftarrow \text{ITERJACOBI} (\mathbf{A}, \mathbf{b}, \mathbf{y})$ 
2:   Pour  $i \leftarrow 1$  à  $n$  faire
3:      $S \leftarrow 0$ 
4:     Pour  $j \leftarrow 1$  à  $n$  ( $j \neq i$ ) faire
5:        $S \leftarrow S + A(i, j) * y(j)$ 
6:     Fin Pour
7:      $x(i) \leftarrow (b(i) - S) / A(i, i)$ 
8:   Fin Pour
9: Fin Fonction

```

Algorithme 4.23 Itération de Gauss-Seidel : calcul de $\mathbf{x}$ tel que

$$x_i = \frac{1}{A_{ii}} \left(b_i - \sum_{j=1}^{i-1} A_{i,j} x_j - \sum_{j=i+1}^n A_{i,j} y_j \right), \quad \forall i \in \llbracket 1, n \rrbracket.$$

Données :

$\mathbf{A}$: matrice de $\mathcal{M}_n(\mathbb{K})$,
 $\mathbf{b}$: vecteur de $\mathbb{K}^n$,
 $\mathbf{y}$: vecteur de $\mathbb{K}^n$,

Résultat :

$\mathbf{x}$: un vecteur de $\mathbb{K}^n$

```

1: Fonction  $\mathbf{x} \leftarrow \text{ITERGAUSSSEIDEL} (\mathbf{A}, \mathbf{b}, \mathbf{y})$ 
2:   Pour  $i \leftarrow 1$  à  $n$  faire
3:      $S \leftarrow 0$ 
4:     Pour  $j \leftarrow 1$  à  $i - 1$  faire
5:        $S \leftarrow S + A(i, j) * x(j)$ 
6:     Fin Pour
7:     Pour  $j \leftarrow i + 1$  à  $n$  faire
8:        $S \leftarrow S + A(i, j) * y(j)$ 
9:     Fin Pour
10:     $x(i) \leftarrow (b(i) - S) / A(i, i)$ 
11:   Fin Pour
12: Fin Fonction

```

En utilisant ces deux fonctions, les algorithmes de résolutions de systèmes linéaires par les méthodes de Jacobi et Gauss-Seidel peuvent se réécrire sous la forme suivante :

```

Fonction  $X \leftarrow \text{RSLJACOBIV2} (\mathbb{A}, \mathbf{b}, \mathbf{x}^0, \varepsilon, \text{kmax})$ 
 $k \leftarrow 0, X \leftarrow \emptyset$ 
 $\mathbf{x} \leftarrow \mathbf{x}^0, \mathbf{r} \leftarrow \mathbf{b} - \mathbb{A} * \mathbf{x},$ 
 $\text{tol} \leftarrow \varepsilon(\|\mathbf{b}\| + 1)$ 
Tantque  $\|\mathbf{r}\| > \text{tol}$  et  $k \leq \text{kmax}$  faire
   $k \leftarrow k + 1$ 
   $\mathbf{p} \leftarrow \mathbf{x}$ 
   $\mathbf{x} \leftarrow \text{ITERJACOBI}(\mathbb{A}, \mathbf{b}, \mathbf{p})$ 
   $\mathbf{r} \leftarrow \mathbf{b} - \mathbb{A} * \mathbf{x},$ 
Fin Tantque
Si  $\|\mathbf{r}\| \leq \text{tol}$  alors
   $X \leftarrow \mathbf{x}$ 
Fin Si
Fin Fonction

```

```

Fonction  $X \leftarrow \text{RSLGAUSSSEIDELV2} (\mathbb{A}, \mathbf{b}, \mathbf{x}^0, \varepsilon, \text{kmax})$ 
 $k \leftarrow 0, X \leftarrow \emptyset$ 
 $\mathbf{x} \leftarrow \mathbf{x}^0, \mathbf{r} \leftarrow \mathbf{b} - \mathbb{A} * \mathbf{x},$ 
 $\text{tol} \leftarrow \varepsilon(\|\mathbf{b}\| + 1)$ 
Tantque  $\|\mathbf{r}\| > \text{tol}$  et  $k \leq \text{kmax}$  faire
   $k \leftarrow k + 1$ 
   $\mathbf{p} \leftarrow \mathbf{x}$ 
   $\mathbf{x} \leftarrow \text{ITERGAUSSSEIDEL}(\mathbb{A}, \mathbf{b}, \mathbf{p})$ 
   $\mathbf{r} \leftarrow \mathbf{b} - \mathbb{A} * \mathbf{x},$ 
Fin Tantque
Si  $\|\mathbf{r}\| \leq \text{tol}$  alors
   $X \leftarrow \mathbf{x}$ 
Fin Si
Fin Fonction

```

- 1 En programmation, dès que l'on commence à faire des copier/coller¹ il faut se poser la question : est-il
 2 possible de faire sans?

La plupart du temps la réponse est oui, et cela permet souvent de simplifier, clarifier et raccourcir le code ce qui simplifie grandement sa maintenance. Dans notre cas, on écrit l'Algorithme générique 4.19 sous forme d'une fonction à laquelle on ajoute aux paramètres d'entrées une fonction formelle **ITERFONC** permettant le calcul d'une itérée (c'est donc une donnée du problème) :

$$\mathbf{x} \leftarrow \text{ITERFONC}(\mathbb{A}, \mathbf{b}, \mathbf{y}).$$

- 3 On a alors

Algorithme 4.24 Méthode itérative pour la résolution d'un système linéaire $\mathbb{A}\mathbf{x} = \mathbf{b}$

Données :

$\mathbb{A}$: matrice de $\mathcal{M}_n(\mathbb{K})$,
 $\mathbf{b}$: vecteur de $\mathbb{K}^n$,
ITERFONC : fonction de paramètres une matrice d'ordre n ,
 : et deux vecteurs de $\mathbb{K}^n$. retourne un vecteur de $\mathbb{K}^n$.
 $\mathbf{x}^0$: vecteur initial de $\mathbb{K}^n$,
 ε : la tolérance, $\varepsilon \in \mathbb{R}^+$,
 kmax : nombre maximum d'itérations, $\text{kmax} \in \mathbb{N}^*$

Résultat :

$\mathbf{x}^{\text{tol}}$: un vecteur de $\mathbb{K}^n$ si convergence, sinon $\emptyset$

- 4
- ```

1: Fonction $\mathbf{X} \leftarrow \text{RSLMETHITER}(\mathbb{A}, \mathbf{b}, \text{ITERFONC}, \mathbf{x}^0, \varepsilon, \text{kmax})$
2: $k \leftarrow 0, \mathbf{x}^{\text{tol}} \leftarrow \emptyset$
3: $\mathbf{x} \leftarrow \mathbf{x}^0, \mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x},$
4: $\text{tol} \leftarrow \varepsilon(\|\mathbf{b}\| + 1)$
5: Tantque $\|\mathbf{r}\| > \text{tol}$ et $k \leq \text{kmax}$ faire
6: $k \leftarrow k + 1$
7: $\mathbf{p} \leftarrow \mathbf{x}$
8: $\mathbf{x} \leftarrow \text{ITERFONC}(\mathbb{A}, \mathbf{b}, \mathbf{p})$
9: $\mathbf{r} \leftarrow \mathbf{b} - \mathbb{A}\mathbf{x},$
10: Fin Tantque
11: Si $\|\mathbf{r}\| \leq \text{tol}$ alors
12: $\mathbf{x}^{\text{tol}} \leftarrow \mathbf{x}$
13: Fin Si
14: Fin Fonction
```
- 

- 5 En utilisant cette fonction, les algorithmes de résolutions de systèmes linéaires par les méthodes de  
 6 Jacobi et Gauss-Seidel peuvent se réécrire sous la forme suivante :

- 7
- |                                                                                                                                                                                                                                                                                                                |                                                                                                                                                                                                                                                                                                                          |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <pre> <b>Fonction</b> <math>\mathbf{X} \leftarrow \text{RSLJACOBI}(\mathbb{A}, \mathbf{b}, \mathbf{x}^0, \varepsilon, \text{kmax})</math>   <math>\mathbf{X} \leftarrow \text{RSLMETHITER}(\mathbb{A}, \mathbf{b}, \text{ITERJACOBI}, \mathbf{x}^0, \varepsilon, \text{kmax})</math> <b>Fin Fonction</b></pre> | <pre> <b>Fonction</b> <math>\mathbf{X} \leftarrow \text{RSLGAUSSSEIDEL}(\mathbb{A}, \mathbf{b}, \mathbf{x}^0, \varepsilon, \text{kmax})</math>   <math>\mathbf{X} \leftarrow \text{RSLMETHITER}(\mathbb{A}, \mathbf{b}, \text{ITERGAUSSSEIDEL}, \mathbf{x}^0, \varepsilon, \text{kmax})</math> <b>Fin Fonction</b></pre> |
|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|

- 8 **Méthode S.O.R.**

- 9 Pour la méthode de relaxation utilisant Gauss-Seidel, avec  $w \in \mathbb{R}^*$ , la suite des itérées est définie par

$$x_i^{[k+1]} = \frac{w}{A_{ii}} \left( b_i - \sum_{j=1}^{i-1} A_{ij} x_j^{[k+1]} - \sum_{j=i+1}^n A_{ij} x_j^{[k]} \right) + (1-w)x_i^{[k]} \quad \forall i \in \llbracket 1, n \rrbracket$$

---

<sup>1</sup>Opération qui pour un bon programmeur est une chose très fatigante!

**Algorithme 4.25** Itération S.O.R. : calcul de  $\mathbf{x}$  tel que

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

**Données :**

$\mathbf{A}$  : matrice de  $\mathcal{M}_n(\mathbb{K})$ ,  
 $\mathbf{b}$  : vecteur de  $\mathbb{K}^n$ ,  
 $\mathbf{y}$  : vecteur de  $\mathbb{K}^n$ ,  
 $w$  : réel non nul.

**Résultat :**

$\mathbf{x}$  : un vecteur de  $\mathbb{K}^n$

**Fonction**  $\mathbf{X} \leftarrow \text{RSLSORV3}(\mathbf{A}, \mathbf{b}, w, \mathbf{x}^0, \varepsilon, \text{kmax})$   
 $\text{ITERFUN} \leftarrow ((\mathbf{M}, \mathbf{r}, \mathbf{s}) \mapsto \text{ITERSOR}(\mathbf{M}, \mathbf{r}, \mathbf{s}, w))$   
 $\mathbf{X} \leftarrow \text{RSLMETHITER}(\mathbf{A}, \mathbf{b}, \text{ITERFUN}, \mathbf{x}^0, \varepsilon, \text{kmax})$   
**Fin Fonction**

```

1: Fonction $\mathbf{x} \leftarrow \text{ITERSOR}(\mathbf{A}, \mathbf{b}, \mathbf{y}, w)$
2: Pour $i \leftarrow 1$ à n faire
3: $S \leftarrow 0$
4: Pour $j \leftarrow 1$ à $i-1$ faire
5: $S \leftarrow S + A(i, j) * \mathbf{x}(j)$
6: Fin Pour
7: Pour $j \leftarrow i+1$ à n faire
8: $S \leftarrow S + A(i, j) * \mathbf{y}(j)$
9: Fin Pour
10: $\mathbf{x}(i) \leftarrow w * (\mathbf{b}(i) - S) / A(i, i) + (1-w) * \mathbf{x}(i)$
11: Fin Pour
12: Fin Fonction

```

1

### Une ch'tite remarque

2

Les méthodes itératives que l'on a étudiées sont basées sur l'expression

$$\mathbf{x}^{[k+1]} = \mathbb{B}\mathbf{x}^{[k]} + \mathbf{c}$$

Si l'on pose

$$\Phi : \mathbf{x} \mapsto \mathbb{B}\mathbf{x} + \mathbf{c}$$

alors

$$\mathbf{x}^{[k+1]} = \Phi(\mathbf{x}^{[k]}).$$

Cela vous rappelle-t'il quelque chose?

3

### 4.5.5 Exercices

4



#### Exercice 4.5.3

**Q. 1** Montrer que pour la matrice

$$\mathbb{A}_1 = \begin{pmatrix} 1 & 2 & -2 \\ 1 & 1 & 1 \\ 2 & 2 & 1 \end{pmatrix}$$

la méthode de Jacobi converge, tandis que la méthode de Gauss-Seidel diverge.

**Q. 2** Montrer que pour la matrice

$$\mathbb{A}_2 = \begin{pmatrix} 2 & -1 & 1 \\ 2 & 2 & 2 \\ -1 & -1 & 2 \end{pmatrix}$$

la méthode de Jacobi diverge, tandis que la méthode de Gauss-Seidel converge.

5



#### Exercice 4.5.4

6

Soit  $\mathbb{A} \in \mathcal{M}_n(\mathbb{R})$  une matrice symétrique définie positive décomposée (par points) sous la forme  $\mathbb{A} = \mathbb{D} - \mathbb{E} - \mathbb{F}$  où  $\mathbb{D} = \text{diag}(\mathbb{A})$ ,  $\mathbb{E}$  est triangulaire inférieure et d'éléments nuls sur la diagonale et  $\mathbb{F}$  est triangulaire supérieure et d'éléments nuls sur la diagonale.

On étudie une méthode itérative de résolution du système linéaire  $\mathbb{A}\mathbf{x} = \mathbf{b}$ . Soit  $\mathbf{x}_0 \in \mathbb{R}^n$ , on définit la suite  $(\mathbf{x}_k)_{k \in \mathbb{N}}$  par

$$(\mathbb{D} - \mathbb{E})\mathbf{x}_{k+1/2} = \mathbb{F}\mathbf{x}_k + \mathbf{b} \quad (4.76)$$

$$(\mathbb{D} - \mathbb{F})\mathbf{x}_{k+1} = \mathbb{E}\mathbf{x}_{k+1/2} + \mathbf{b} \quad (4.77)$$

**Q. 1** *Ecrire le vecteur  $\mathbf{x}_{k+1}$  sous la forme*

$$\mathbf{x}_{k+1} = \mathbb{B}\mathbf{x}_k + \mathbf{c} \quad (4.78)$$

*en explicitant la matrice  $\mathbb{B}$  et le vecteur  $\mathbf{c}$ .*

**Q. 2** 1. *Montrer que*

$$\mathbb{D}^{-1} = (\mathbb{D} - \mathbb{E})^{-1} - \mathbb{D}^{-1}\mathbb{E}(\mathbb{D} - \mathbb{E})^{-1}. \quad (4.79)$$

2. *Soit  $(\lambda, \mathbf{p})$  un élément propre de la matrice  $\mathbb{B}$ . Montrer que*

$$\lambda \mathbb{A}\mathbf{p} + (\lambda - 1)\mathbb{E}\mathbb{D}^{-1}\mathbb{F}\mathbf{p} = 0. \quad (4.80)$$

**Q. 3** *En déduire la convergence de cette méthode.*

**Q. 4** *Etendre ces résultats au cas d'une décomposition  $\mathbb{A} = \mathbb{D} - \mathbb{E} - \mathbb{F}$  par blocs.*

1

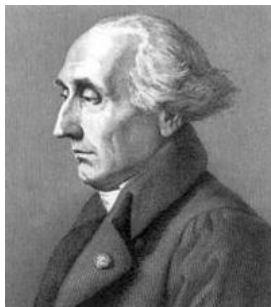
# Chapitre 5

## Interpolation

L'interpolation est un outil mathématique permettant de construire des fonctions à partir de la donnée d'un nombre fini de valeurs.

A développer...

Les protagonistes de cette histoire



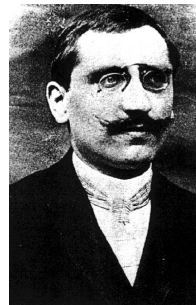
(a) *Joseph-Louis Lagrange* 1736-1813, mathématicien italien puis français



(b) *Pafnouti Lvovitch Tchebychev* 1821-1894, mathématicien russe



(c) *Charles Hermite* 1822-1901, mathématicien français



(d) *Henri-Léon Lebesgue* 1875-1941, mathématicien français

### 5.1 Polynôme d'interpolation de Lagrange

**Exercice 5.1.1**

Soient  $n \in \mathbb{N}^*$  et  $n + 1$  couples de  $\mathbb{R}^2$ ,  $(x_i, y_i)_{i \in \llbracket 0, n \rrbracket}$ , tels que les  $x_i$  sont distincts deux à deux. On note

**Q. 1** 1. Soit  $i \in \llbracket 0, n \rrbracket$ . Montrer qu'il existe un unique polynôme  $L_i$  de degré  $n$  vérifiant

$$L_i(x_j) = \delta_{ij}, \quad \forall j \in \llbracket 0, n \rrbracket. \quad (5.1)$$

2. Montrer que les  $(L_i)_{i \in \llbracket 0, n \rrbracket}$  forment une base de  $\mathbb{R}_n[X]$  (espace vectoriel des polynômes à coefficients réels de degré inférieur ou égal à  $n$ ).

On définit le polynôme  $P_n$  par

$$P_n(x) = \sum_{i=0}^n y_i L_i(x). \quad (5.2)$$

**Q. 2** Montrer que polynôme  $P_n$  est l'unique polynôme de degré au plus  $n$  vérifiant  $P_n(x_i) = y_i$ ,  $\forall i \in \llbracket 0, n \rrbracket$ .

1

**Correction Exercice**

2

**Q. 1** 1. De (5.1), on déduit que les  $n$  points distincts  $x_j$  pour  $j \in \llbracket 0, n \rrbracket \setminus \{i\}$  sont les  $n$  zéros du polynôme  $L_i$  de degré  $n$  : il s'écrit donc sous la forme

3

4

$$L_i(x) = C \prod_{\substack{j=0 \\ j \neq i}}^n (x - x_j) \quad \text{avec } C \in \mathbb{R}$$

5

Pour déterminer la constante  $C$ , on utilise (5.1) avec  $j = i$

$$L_i(x_i) = 1 = C \prod_{\substack{j=0 \\ j \neq i}}^n (x_i - x_j)$$

6

Les points  $x_i$  sont distincts deux à deux, on a  $\prod_{\substack{j=0 \\ j \neq i}}^n (x_i - x_j) \neq 0$  et donc

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

7

d'où

$$L_i(x) = \prod_{\substack{j=0 \\ j \neq i}}^n \frac{x - x_j}{x_i - x_j}, \quad \forall i \in \llbracket 0, n \rrbracket. \quad (5.3)$$

8

Il reste à démontrer l'unicité. On suppose qu'il existe  $L_i$  et  $U_i$  deux polynômes de  $\mathbb{R}_n[X]$  vérifiant (5.1). Alors  $Q_i = L_i - U_i$  est polynôme de degré  $n$  (au plus) admettant  $n + 1$  zéros distincts, c'est donc le polynôme nul et on a nécessairement  $L_i = U_i$ .

9

10

2. On sait que  $\dim \mathbb{R}_n[X] = n + 1$ . Pour que les  $\{L_i\}_{i \in \llbracket 0, n \rrbracket}$  forment une base de  $\mathbb{R}_n[X]$  il suffit de démontrer qu'ils sont linéairement indépendants.

11

12

Soit  $\lambda_0, \dots, \lambda_n$   $n + 1$  scalaires. Montrons pour cela que

13

$$\sum_{i=1}^n \lambda_i L_i = 0 \implies \lambda_i = 0, \quad \forall i \in \llbracket 0, n \rrbracket$$

14

Noter que la première égalité est dans l'espace vectoriel  $\mathbb{R}_n[X]$  et donc le 0 est pris au sens polynôme nul.

15

16

On a

$$\sum_{i=1}^n \lambda_i L_i = 0 \iff \sum_{i=1}^n \lambda_i L_i(x) = 0, \quad \forall x \in \mathbb{R}$$

17

Soit  $k \in \llbracket 0, n \rrbracket$ . En choisissant  $x = x_k$ , on a par (5.1)  $\sum_{i=1}^n \lambda_i L_i(x_k) = \lambda_k$  et donc

$$\sum_{i=1}^n \lambda_i L_i = 0 \implies \sum_{i=1}^n \lambda_i L_i(x_k) = 0, \quad \forall k \in \llbracket 0, n \rrbracket \iff \lambda_k = 0, \quad \forall k \in \llbracket 0, n \rrbracket.$$

Les  $\{L_i\}_{i \in \llbracket 0, n \rrbracket}$  sont donc linéairement indépendants.

**Q. 2** Par construction  $P_n \in \mathbb{R}_n[X]$  et

$$P_n(x_i) = y_i, \quad \forall i \in \llbracket 0, n \rrbracket \quad (5.4)$$

et c'est l'unique polynôme de degré au plus  $n$  vérifiant (??) (car la décomposition dans la base  $\{L_i\}_{i \in \llbracket 0, n \rrbracket}$  est unique).

**Q. 3** S'il existe  $i \in \llbracket 0, n \rrbracket$  tel que  $x = x_i$  alors l'équation (??) est immédiatement vérifiée.

Soit  $x \in [a, b]$  distinct de tous les  $x_i$ . Comme  $f \in \mathcal{C}^{n+1}([a, b]; \mathbb{R})$ ,  $P_n \in \mathbb{R}_n[X]$  et  $\pi_n \in \mathbb{R}_{n+1}[X]$ , on en déduit que la fonction  $F$  est dans  $\mathcal{C}^{n+1}([a, b]; \mathbb{R})$ . La fonction  $F$  admet aussi  $n+2$  zéros :  $x, x_0, \dots, x_n$ . On note  $\xi_{x,1}^{[0]}, \dots, \xi_{x,n+2}^{[0]}$  ces  $n+2$  zéros ordonnés  $\xi_{x,1}^{[0]} < \dots < \xi_{x,n+2}^{[0]}$ . La fonction  $F$  étant continue sur  $[a, b]$  et dérivable sur  $]a, b[$ , le théorème de Rolle dit qu'entre deux zéros consécutifs de  $F$ , il existe au moins un zéro de  $F' = F^{(1)}$ . Plus précisément on a

$$\forall i \in \llbracket 1, n+1 \rrbracket, \exists \xi_{x,i}^{[1]} \in ]\xi_{x,i}^{[0]}, \xi_{x,i+1}^{[0]}[ \text{ tels que } F^{(1)}(\xi_{x,i}^{[1]}) = 0$$

et on en déduit que la fonction  $F^{(1)}$  admet  $n+1$  zéros  $\xi_{x,1}^{[1]}, \dots, \xi_{x,n+1}^{[1]}$  et l'on a  $\xi_{x,1}^{[0]} < \xi_{x,1}^{[1]} < \dots < \xi_{x,n+1}^{[1]} < \xi_{x,n+2}^{[0]}$ . Il faut noter la dépendance en  $x$  des zéros de  $F'$  d'où la notation un peu "lourde".

Montrons par récurrence finie que  $(\mathcal{P}_k)$  est vraie pour tout  $k \in \llbracket 1, n+1 \rrbracket$

$$(\mathcal{P}_k) : \exists \xi_{x,i}^{[k]}, i \in \llbracket 1, n+2-k \rrbracket, \xi_{x,1}^{[0]} < \xi_{x,1}^{[k]} < \dots < \xi_{x,n+2-k}^{[k]} < \xi_{x,n+2}^{[0]} \text{ tels que } F^{(k)}(\xi_{x,i}^{[k]}) = 0$$

**Initialisation :** Pour  $k=1$ , la preuve a déjà été faite.

**Hérédité :** Soit  $1 < k-1 < n+1$ , on suppose  $(\mathcal{P}_{k-1})$  vérifiée. La fonction  $F^{(k-1)}$  étant continue sur  $[a, b]$  et dérivable sur  $]a, b[$ , le théorème de Rolle dit qu'entre deux zéros consécutifs de  $F^{(k-1)}$ , il existe au moins un zéro de  $F^{(k)}$ . Par hypothèse  $F^{(k-1)}$  admet  $n+2-(k-1)$  zéros vérifiant

$$\xi_{x,1}^{[0]} < \xi_{x,1}^{[k-1]} < \dots < \xi_{x,n+2-(k-1)}^{[k-1]} < \xi_{x,n+2}^{[0]}$$

La fonction  $F^{(k-1)}$  est continue sur  $[a, b]$  et dérivable sur  $]a, b[$  puisque  $F \in \mathcal{C}^{n+1}([a, b]; \mathbb{R})$ . Par application du théorème de Rolle, entre deux zéros de  $F^{(k-1)}$ , il existe au moins un zéro de  $F^{(k)}$ . Plus précisément pour tout  $i \in \llbracket 1, n+2-k \rrbracket$  on a

$$\exists \xi_{x,i}^{[k]} \in ]\xi_{x,i}^{[k-1]}, \xi_{x,i+1}^{[k-1]}[, \quad F^{(k)}(\xi_{x,i}^{[k]}) = 0$$

De plus, par construction,  $\xi_{x,1}^{[0]} < \xi_{x,1}^{[k]} < \dots < \xi_{x,n+2-k}^{[k]} < \xi_{x,n+2}^{[0]}$  et donc  $(\mathcal{P}_k)$  est vraie.

Avec  $k = n+1$  on obtient

$$(\mathcal{P}_{n+1}) : \exists \xi_{x,1}^{[n+1]} \in ]\xi_{x,1}^{[0]}, \xi_{x,n+2}^{[0]}[ \text{ tel que } F^{(n+1)}(\xi_{x,1}^{[n+1]}) = 0$$

et donc

$$0 = F^{(n+1)}(\xi_{x,1}^{[n+1]}) = f^{(n+1)}(\xi_{x,1}^{[n+1]}) - P_n^{(n+1)}(\xi_{x,1}^{[n+1]}) - \frac{f(x) - P_n(x)}{\pi_n(x)} \pi_n^{(n+1)}(\xi_{x,1}^{[n+1]})$$

Comme  $P_n \in \mathbb{R}_n[X]$ , on a  $P_n^{(n+1)} = 0$ . De plus  $\pi_n \in \mathbb{R}_{n+1}[X]$ , et comme  $\pi_n(x) = x^{n+1} + Q(x)$  avec  $Q \in \mathbb{R}_n[X]$  (i.e. son monôme de puissance  $n+1$  à pour coefficient 1) on obtient  $\pi_n^{(n+1)}(x) = (n+1)!$  On a alors

$$f^{(n+1)}(\xi_{x,1}^{[n+1]}) = \frac{f(x) - P_n(x)}{\pi_n(x)} (n+1)!$$

◇

### ♥ Definition 5.1

Soient  $n \in \mathbb{N}^*$  et  $(x_i, y_i)_{i \in \llbracket 0, n \rrbracket}$  avec  $(x_i, y_i) \in \mathbb{R}^2$  et les  $x_i$  distincts deux à deux. Le **polynôme d'interpolation de Lagrange** associé aux  $n + 1$  points  $(x_i, y_i)_{i \in \llbracket 0, n \rrbracket}$ , noté  $P_n$ , est donné par

$$P_n(x) = \sum_{i=0}^n y_i L_i(x), \quad \forall x \in \mathbb{R} \quad (5.5)$$

avec

$$L_i(x) = \prod_{\substack{j=0 \\ j \neq i}}^n \frac{x - x_j}{x_i - x_j}, \quad \forall i \in \llbracket 0, n \rrbracket, \quad \forall x \in \mathbb{R}. \quad (5.6)$$

### Théorème 5.2

Le **polynôme d'interpolation de Lagrange**,  $P_n$ , associé aux  $n + 1$  points  $(x_i, y_i)_{i \in \llbracket 0, n \rrbracket}$ , est l'unique polynôme de degré au plus  $n$ , vérifiant

$$P_n(x_i) = y_i, \quad \forall i \in \llbracket 0, n \rrbracket. \quad (5.7)$$

A titre d'exemple, on représente, En figure 5.2, le polynôme d'interpolation de Lagrange associé à 7 points donnés.

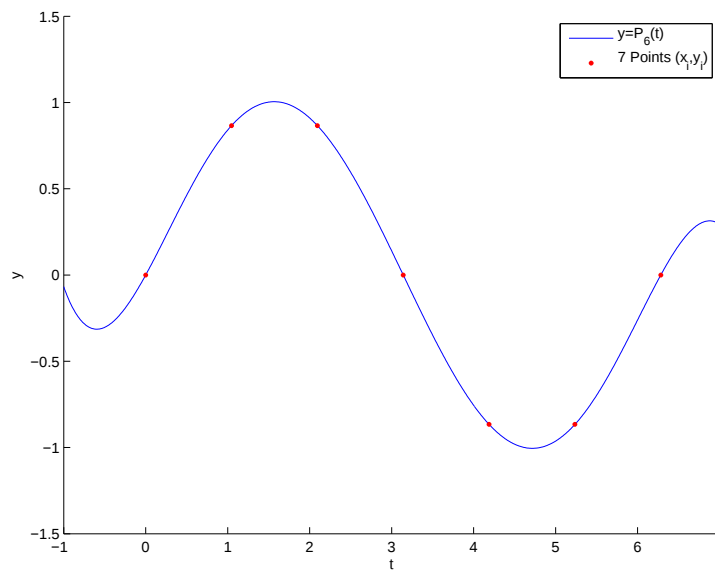


Figure 5.2: Polynôme d'interpolation de Lagrange avec 7 points donnés)

### Exercice 5.1.2

Ecrire la fonction **LAGRANGE** permettant de calculer  $P_n$  (polynôme d'interpolation de Lagrange associé aux  $n + 1$  points  $(x_i, y_i)_{i \in \llbracket 0, n \rrbracket}$ ) au point  $t \in \mathbb{R}$ .

#### Correction Exercice

**But :** Calculer le polynôme  $P_n(t)$  défini par (5.5)

**Données :**  $X$  : vecteur/tableau de  $\mathbb{R}^{n+1}$ ,  $X(i) = x_{i-1} \quad \forall i \in \llbracket 1, n+1 \rrbracket$  et  $X(i) \neq X(j)$  pour  $i \neq j$ ,

$Y$  : vecteur/tableau de  $\mathbb{R}^{n+1}$ ,  $Y(i) = y_{i-1} \quad \forall i \in \llbracket 1, n+1 \rrbracket$ ,  
 $t$  : un réel.

**Résultat :**  $y$  : le réel  $y = P_n(t)$ .

Algorithme 5.1  $\mathcal{R}_2$ 

1: Calcul de  $y = \mathcal{P}_n(t) = \sum_{i=1}^{n+1} Y(i) L_{i-1}(t)$

Algorithme 5.1  $\mathcal{R}_3$ 

1:  $y \leftarrow 0$   
 2: **Pour**  $i \leftarrow 1$  à  $n+1$  **faire**  
 3:    $y \leftarrow y + Y(i) * L_{i-1}(t)$   
 4: **Fin Pour**

Algorithme 5.1  $\mathcal{R}_3$ 

1:  $y \leftarrow 0$   
 2: **Pour**  $i \leftarrow 1$  à  $n+1$  **faire**  
 3:    $y \leftarrow y + Y(i) * L_{i-1}(t)$   
 4: **Fin Pour**

Algorithme 5.1  $\mathcal{R}_4$ 

1:  $y \leftarrow 0$   
 2: **Pour**  $i \leftarrow 1$  à  $n+1$  **faire**  
 3:    $L \leftarrow \prod_{\substack{j=1 \\ j \neq i}}^{n+1} \frac{t - X(j)}{X(i) - X(j)}$   
 4:    $y \leftarrow y + Y(i) * L$   
 5: **Fin Pour**

Algorithme 5.1  $\mathcal{R}_4$ 

1:  $y \leftarrow 0$   
 2: **Pour**  $i \leftarrow 1$  à  $n+1$  **faire**  
 3:    $L \leftarrow \prod_{\substack{j=1 \\ j \neq i}}^{n+1} \frac{t - X(j)}{X(i) - X(j)}$   
 4:    $y \leftarrow y + Y(i) * L$   
 5: **Fin Pour**

Algorithme 5.1  $\mathcal{R}_5$ 

1:  $y \leftarrow 0$   
 2: **Pour**  $i \leftarrow 1$  à  $n+1$  **faire**  
 3:    $L \leftarrow 1$   
 4:   **Pour**  $j \leftarrow 1$  à  $n+1$ , ( $j \sim i$ ) **faire**  
 5:      $L \leftarrow L * (t - X(j)) / (X(i) - X(j))$   
 6:   **Fin Pour**  
 7:    $y \leftarrow y + Y(i) * L$   
 8: **Fin Pour**

On obtient alors l'algorithme final

**Algorithme 5.1** Fonction **LAGRANGE** permettant de calculer le polynôme d'interpolation de Lagrange  $\mathcal{P}_n(x)$  défini par (5.5)

**Données :**  $\mathbf{X}$  : vecteur/tableau de  $\mathbb{R}^{n+1}$ ,  $X(i) = x_{i-1} \forall i \in \llbracket 1, n+1 \rrbracket$  et  $X(i) \neq X(j)$  pour  $i \neq j$ ,  
 $\mathbf{Y}$  : vecteur/tableau de  $\mathbb{R}^{n+1}$ ,  $Y(i) = y_{i-1} \forall i \in \llbracket 1, n+1 \rrbracket$ ,  
 $t$  : un réel.

**Résultat :**  $y$  : le réel  $y = \mathcal{P}_n(t)$ .

1: **Fonction**  $y \leftarrow \text{LAGRANGE} (t, \mathbf{X}, \mathbf{Y})$   
 2:    $y \leftarrow 0$   
 3:   **Pour**  $i \leftarrow 1$  à  $n+1$  **faire**  
 4:      $L \leftarrow 1$   
 5:     **Pour**  $j \leftarrow 1$  à  $n+1$ , ( $j \sim i$ ) **faire**  
 6:        $L \leftarrow L * (t - X(j)) / (X(i) - X(j))$   
 7:     **Fin Pour**  
 8:      $y \leftarrow y + Y(i) * L$   
 9:   **Fin Pour**  
 10:   **return**  $y$   
 11: **Fin Fonction**

◇

### 5.1.1 Erreur de l'interpolation

Soit une fonction  $f : [a, b] \rightarrow \mathbb{R}$ . On suppose que les  $y_i$  sont donnés par

$$y_i = f(x_i), \quad \forall i \in \llbracket 0, n \rrbracket. \quad (5.8)$$

On cherche à évaluer l'erreur  $E_n(t) = f(x) - \mathcal{P}_n(t)$ ,  $\forall t \in [a, b]$ .

On propose en figures 5.3 à 5.5 d'étudier deux exemples. Le premier (figure de gauche) correspond à l'interpolation de la fonction  $\sin(t)$  par le polynôme d'interpolation de Lagrange aux points équidistants  $t_i = a + ih$  avec  $a = 0$  et  $h = 2\pi/n$ . Le second (figure de droite) correspond à l'interpolation de la fonction  $\frac{1}{1+t^2}$  par le polynôme d'interpolation de Lagrange aux points  $x_i = a + ih$  avec  $a = -5$  et  $h = 10/n$ .

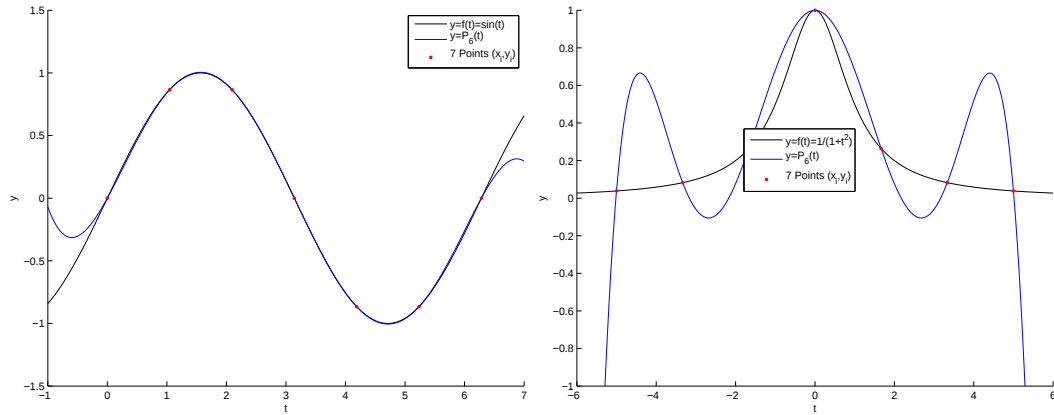


Figure 5.3: Polynômes d'interpolation de lagrange avec  $n = 6$  (7 points) uniformément répartis. À gauche pour la fonction  $f : x \rightarrow \sin(x)$  avec  $x_0 = 0$ ,  $x_6 = 2\pi$  et à droite pour la fonction  $f : x \rightarrow 1/(1 + 25x^2)$  avec  $x_0 = -5$ ,  $x_6 = 5$ .

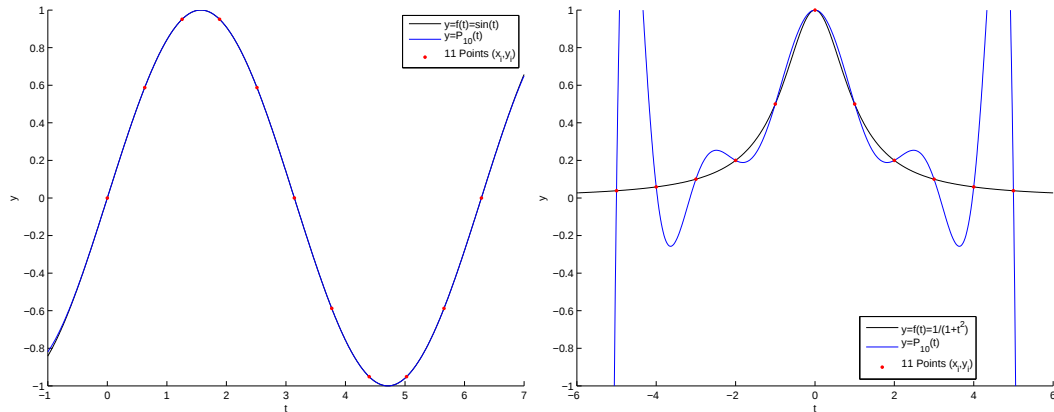


Figure 5.4: Polynômes d'interpolation de lagrange avec  $n = 10$  (11 points) uniformément répartis. À gauche pour la fonction  $f : x \rightarrow \sin(x)$  avec  $x_0 = 0$ ,  $x_{10} = 2\pi$  et à droite pour la fonction  $f : x \rightarrow 1/(1 + 25x^2)$  avec  $x_0 = -5$ ,  $x_{10} = 5$ .



#### Exercice 5.1.3

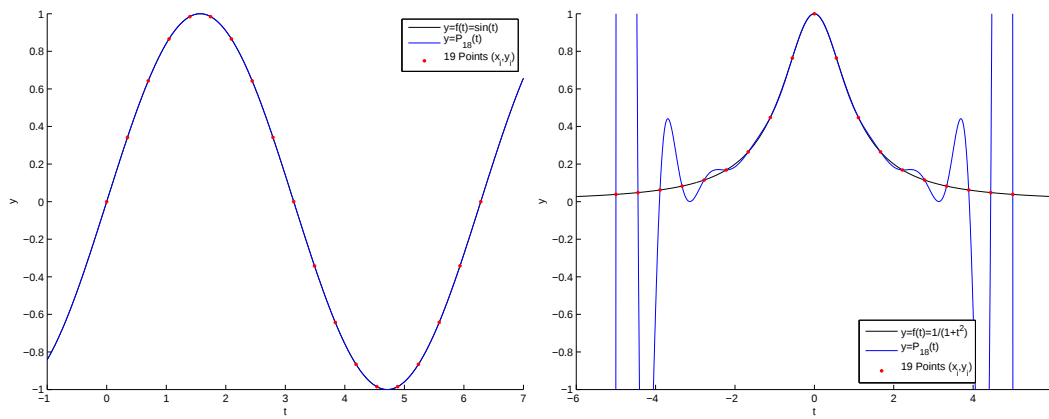


Figure 5.5: Polynômes d'interpolation de lagrange avec  $n = 18$  (19 points) uniformément répartis. A gauche pour la fonction  $f : x \rightarrow \sin(x)$  avec  $x_0 = 0$ ,  $x_{18} = 2\pi$  et à droite pour la fonction  $f : x \rightarrow 1/(1+x^2)$  avec  $x_0 = -5$ ,  $x_{18} = 5$ .

Soit  $f \in \mathcal{C}^{n+1}([a; b]; \mathbb{R})$ . Soient  $n \in \mathbb{N}^*$  et  $n + 1$  couples de  $\mathbb{R}^2$ ,  $(x_i, y_i)_{i \in \llbracket 0, n \rrbracket}$ , tels que les  $x_i$  sont distincts deux à deux et  $y_i = f(x_i)$ .

On note par  $P_n$  le polynôme d'interpolation de Lagrange associé aux points  $(x_i, y_i)_{i \in \llbracket 0, n \rrbracket}$  et  $\pi_n$  le polynôme de degré  $n + 1$  défini par

$$\pi_n(x) = \prod_{i=0}^n (x - x_i). \quad (5.9)$$

**Q. 1** Montrer que,  $\forall x \in [a; b]$ , il existe  $\xi_x$  appartenant au plus petit intervalle fermé contenant  $x, x_0, \dots, x_n$  tel que

$$f(x) - P_n(x) = \frac{\pi_n(x)}{(n+1)!} f^{(n+1)}(\xi_x). \quad (5.10)$$

**Indication :** Etudier les zéros de la fonction  $F(t) = f(t) - P_n(t) - \frac{f(x) - P_n(x)}{\pi_n(x)} \pi_n(t)$ .

### Correction Exercice

**Q. 1** S'il existe  $i \in \llbracket 0, n \rrbracket$  tel que  $x = x_i$  alors l'équation (5.10) est immédiatement vérifiée.

Soit  $x \in [a, b]$  distinct de tous les  $x_i$ . Comme  $f \in \mathcal{C}^{n+1}([a; b]; \mathbb{R})$ ,  $P_n \in \mathbb{R}_n[X]$  et  $\pi_n \in \mathbb{R}_{n+1}[X]$ , on en déduit que la fonction  $F$  est dans  $\mathcal{C}^{n+1}([a; b]; \mathbb{R})$ . La fonction  $F$  admet aussi  $n + 2$  zéros :  $x, x_0, \dots, x_n$ .

On note  $\xi_{x,1}^{[0]}, \dots, \xi_{x,n+2}^{[0]}$  ces  $n + 2$  zéros ordonnés  $\xi_{x,1}^{[0]} < \dots < \xi_{x,n+2}^{[0]}$ . La fonction  $F$  étant continue sur  $[a, b]$  et dérivable sur  $]a, b[$ , le théorème de Rolle dit qu'entre deux zéros consécutifs de  $F$ , il existe au moins un zéro de  $F' = F^{(1)}$ . Plus précisément on a

$$\forall i \in \llbracket 1, n + 1 \rrbracket, \exists \xi_{x,i}^{[1]} \in ]\xi_{x,i}^{[0]}, \xi_{x,i+1}^{[0]}[ \text{ tels que } F^{(1)}(\xi_{x,i}^{[1]}) = 0$$

et on en déduit que la fonction  $F^{(1)}$  admet  $n + 1$  zéros  $\xi_{x,1}^{[1]}, \dots, \xi_{x,n+1}^{[1]}$  et l'on a  $\xi_{x,1}^{[0]} < \xi_{x,1}^{[1]} < \dots < \xi_{x,n+1}^{[1]} < \xi_{x,n+2}^{[0]}$ . Il faut noter la dépendance en  $x$  des zéros de  $F'$  d'où la notation un peu "lourde".

Montrons par récurrence finie que  $(\mathcal{P}_k)$  est vraie pour tout  $k \in \llbracket 1, n + 1 \rrbracket$

$$(\mathcal{P}_k) : \exists \xi_{x,i}^{[k]}, i \in \llbracket 1, n + 2 - k \rrbracket, \xi_{x,1}^{[0]} < \xi_{x,1}^{[k]} < \dots < \xi_{x,n+2-k}^{[k]} < \xi_{x,n+2}^{[0]} \text{ tels que } F^{(k)}(\xi_{x,i}^{[k]}) = 0$$

**Initialisation :** Pour  $k = 1$ , la preuve a déjà été faite.

**Hérédité :** Soit  $1 < k - 1 < n + 1$ , on suppose  $(\mathcal{P}_{k-1})$  vérifiée. La fonction  $F^{(k-1)}$  étant continue sur  $[a, b]$  et dérivable sur  $]a, b[$ , le théorème de Rolle dit qu'entre deux zéros consécutifs de  $F^{(k-1)}$ , il existe au moins un zéro de  $F^{(k)}$ . Par hypothèse  $F^{(k-1)}$  admet  $n + 2 - (k - 1)$  zéros vérifiant

$$\xi_{x,1}^{[0]} < \xi_{x,1}^{[k-1]} < \dots < \xi_{x,n+2-(k-1)}^{[k-1]} < \xi_{x,n+2}^{[0]}$$

- 1 La fonction  $F^{(k-1)}$  est continue sur  $[a, b]$  et dérivable sur  $]a, b[$  puisque  $F \in \mathcal{C}^{n+1}([a, b]; \mathbb{R})$ . Par  
 2 application du théorème de Rolle, entre deux zéros de  $F^{(k-1)}$ , il existe au moins un zéro de  $F^{(k)}$ .  
 3 Plus précisément pour tout  $i \in \llbracket 1, n+2-k \rrbracket$  on a

$$\exists \xi_{x,i}^{[k]} \in ]\xi_{x,i}^{[k-1]}, \xi_{x,i+1}^{[k-1]}[, \quad F^{(k)}(\xi_{x,i}^{[k]}) = 0$$

- 4 De plus, par construction,  $\xi_{x,1}^{[0]} < \xi_{x,1}^{[k]} < \dots < \xi_{x,n+2-k}^{[k]} < \xi_{x,n+2}^{[0]}$  et donc  $(\mathcal{P}_k)$  est vraie.

- 5 Avec  $k = n+1$  on obtient

$$(\mathcal{P}_{n+1}) : \exists \xi_{x,1}^{[n+1]} \in ]\xi_{x,1}^{[0]}, \xi_{x,n+2}^{[0]}[ \text{ tel que } F^{(n+1)}(\xi_{x,1}^{[n+1]}) = 0$$

- 6 et donc

$$0 = F^{(n+1)}(\xi_{x,1}^{[n+1]}) = f^{(n+1)}(\xi_{x,1}^{[n+1]}) - P_n^{(n+1)}(\xi_{x,1}^{[n+1]}) - \frac{f(x) - P_n(x)}{\pi_n(x)} \pi_n^{(n+1)}(\xi_{x,1}^{[n+1]})$$

- 7 Comme  $P_n \in \mathbb{R}_n[X]$ , on a  $P_n^{(n+1)} = 0$ . De plus  $\pi_n \in \mathbb{R}_{n+1}[X]$ , et comme  $\pi_n(x) = x^{n+1} + Q(x)$  avec  
 8  $Q \in \mathbb{R}_n[X]$  (i.e. son monôme de puissance  $n+1$  à pour coefficient 1) on obtient  $\pi_n^{(n+1)}(x) = (n+1)!$  On  
 9 a alors

$$f^{(n+1)}(\xi_{x,1}^{[n+1]}) = \frac{f(x) - P_n(x)}{\pi_n(x)} (n+1)!$$

10

11

- 12 Le résultat suivant est du à Cauchy (1840)



### Théorème 5.3

Soient  $n \in \mathbb{N}^*$  et  $x_0, \dots, x_n$   $n+1$  points distincts de l'intervalle  $[a, b]$ . Soient  $f \in \mathcal{C}^{n+1}([a, b]; \mathbb{R})$  et  $\mathcal{P}_n$  le polynôme d'interpolation de Lagrange de degré  $n$  passant par  $(x_i, f(x_i))$ ,  $\forall i \in \llbracket 0, n \rrbracket$ . Alors,  $\forall x \in [a, b]$ ,  $\exists \xi_x \in (\min(x_i, x), \max(x_i, x))$ ,

$$f(x) - \mathcal{P}_n(x) = \frac{f^{(n+1)}(\xi_x)}{(n+1)!} \prod_{i=0}^n (x - x_i) \quad (5.11)$$

13

## 5.1.2 Points de Chebyshev

- 15 Pour minimiser l'erreur commise lors de l'interpolation d'une fonction  $f$  par un polynôme d'interpolation  
 16 de Lagrange, on peut, pour un  $n$  donné, "jouer" sur le choix des points  $x_i$  :  
 17 Trouver  $(\bar{x}_i)_{i=0}^n$ ,  $\bar{x}_i \in [a, b]$ , distincts deux à deux, tels que

$$\max_{t \in [a, b]} \prod_{i=0}^n |t - \bar{x}_i| \leq \max_{t \in [a, b]} \prod_{i=0}^n |t - x_i|, \quad \forall (x_i)_{i=0}^n, x_i \in [a, b], \text{ distincts 2 à 2} \quad (5.12)$$

- 18 On a alors le résultat suivant

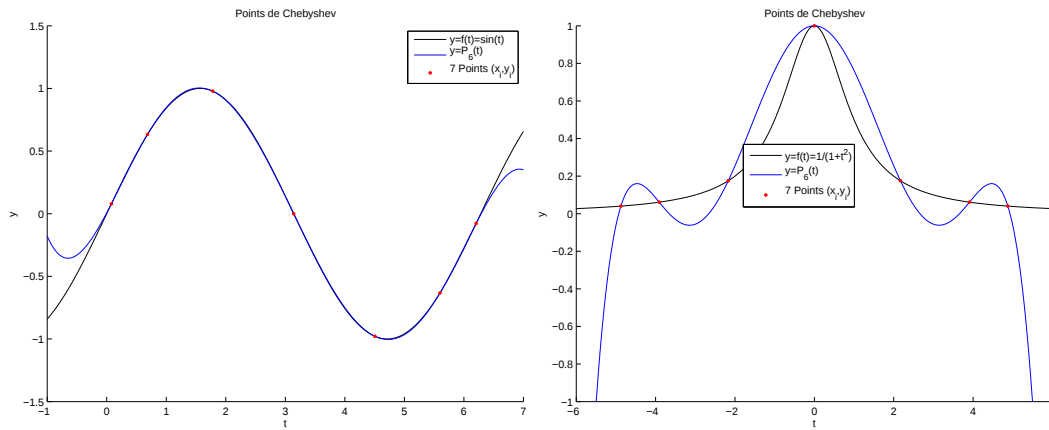
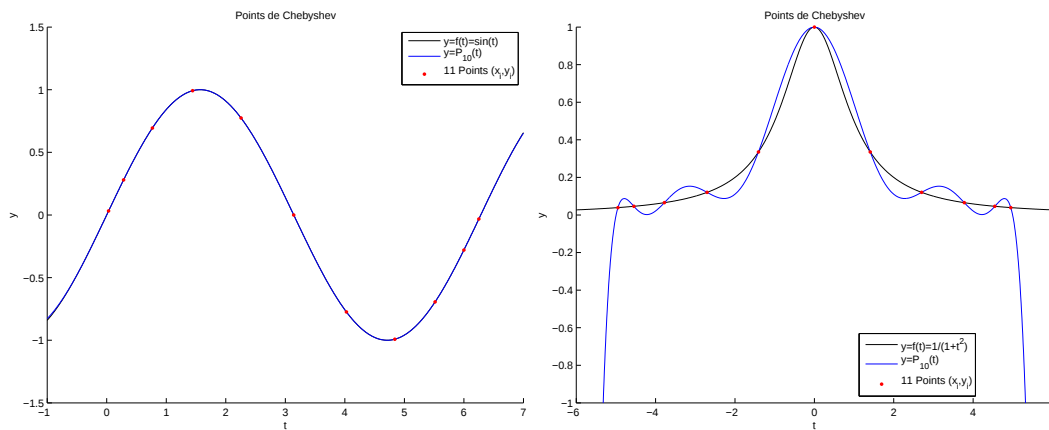
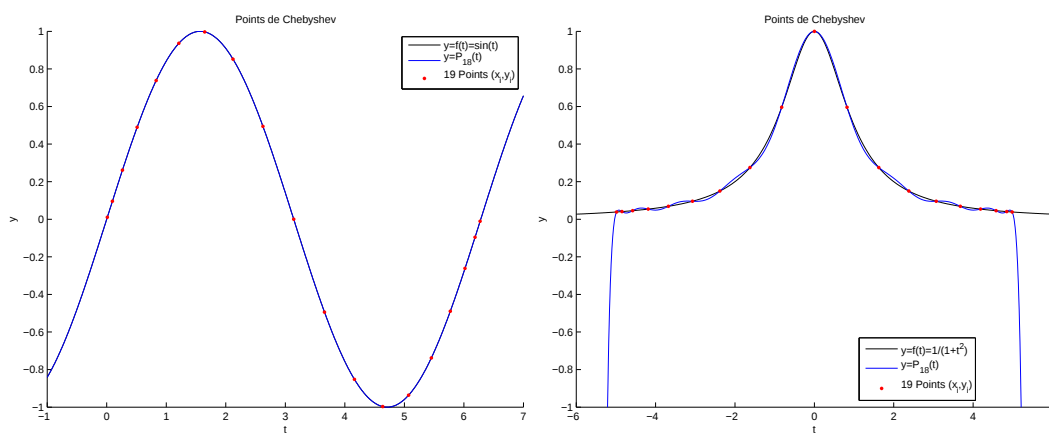


### Théorème 5.4

Les points réalisant (5.12) sont les points de Chebyshev donnés par

$$\bar{x}_i = \frac{a+b}{2} + \frac{b-a}{2} \cos\left(\frac{(2i+1)\pi}{2n+2}\right), \quad \forall i \in \llbracket 0, n \rrbracket. \quad (5.13)$$

19

Figure 5.6: Erreurs d'interpolation avec  $n = 6$ Figure 5.7: Erreurs d'interpolation avec  $n = 10$ Figure 5.8: Erreurs d'interpolation avec  $n = 18$

## 5.1.3 Stabilité

Inspiré du polycopié de G. Barles : ici

On suppose que l'on commet des erreurs lors du calcul des  $f(x_i)$  et l'on note  $f_i \approx f(x_i)$  les valeurs numériques obtenues. Dans ce cadre, on a deux polynômes d'interpolation de Lagrange l'un "exact" avec les couples de points  $(x_i, y_i) = (x_i, f(x_i))$  noté  $P_n$ , et l'autre "approché" avec les couples de points  $(x_i, f_i)$  noté  $\hat{P}_n$ . On a donc

$$P_n(x) = \sum_{i=0}^n f(x_i) L_i(x) \quad \text{et} \quad \hat{P}_n(x) = \sum_{i=0}^n f_i L_i(x)$$

L'erreur commise est alors majorée par :

$$\begin{aligned} |\hat{P}_n(x) - P_n(x)| &= \left| \sum_{i=0}^n (f_i - f(x_i)) L_i(x) \right| \\ &\leq \sum_{i=0}^n |f_i - f(x_i)| |L_i(x)| \\ &\leq \max_{i \in \llbracket 0, n \rrbracket} |f_i - f(x_i)| \sum_{i=0}^n |L_i(x)|. \end{aligned}$$

On note  $\Lambda_n = \max_{x \in [a, b]} \sum_{i=0}^n |L_i(x)|$ , dites *Constante de Lebesgue*. Cette constante est bien définie car l'application  $x \mapsto \sum_{i=0}^n |L_i(x)|$  est continue sur  $[a, b]$  intervalle fermé borné de  $\mathbb{R}$  et donc le maximum est bien atteint : il est donc fini.

On obtient alors

$$\|\hat{P}_n - P_n\|_{\infty} \leq \Lambda_n \max_{i \in \llbracket 0, n \rrbracket} |f_i - f(x_i)|. \quad (5.14)$$

**Proposition 5.5**

Soient  $n \in \mathbb{N}^*$  et  $x_0, \dots, x_n$  des points distincts de  $[a, b]$ . L'application  $\mathcal{L}_n : \mathcal{C}^0([a, b]; \mathbb{R}) \longrightarrow \mathbb{R}_n[X]$  qui à toute fonction  $f \in \mathcal{C}^0([a, b]; \mathbb{R})$  donne le polynôme d'interpolation de Lagrange  $P_n$  associés aux couples de  $(x_i, f(x_i))_{i \in \llbracket 0, n \rrbracket}$  est bien définie et linéaire. De plus on a

$$\|\mathcal{L}_n\| \stackrel{\text{def}}{=} \sup_{\substack{f \in \mathcal{C}^0([a, b]; \mathbb{R}) \\ f \neq 0}} \frac{\|\mathcal{L}_n(f)\|_{\infty}}{\|f\|_{\infty}} = \Lambda_n. \quad (5.15)$$

*Proof.* **Bien définie et linéaire : facile mais à écrire**

On montre tout d'abord que  $\|\mathcal{L}_n\| \leq \Lambda_n$ . En effet, on a pour tout  $x \in [a, b]$  et pour tout  $f \in \mathcal{C}^0([a, b]; \mathbb{R})$ ,

$$\begin{aligned} |\mathcal{L}_n(f)(x)| &= \left| \sum_{i=0}^n f(x_i) L_i(x) \right| \\ &\leq \sum_{i=0}^n |f(x_i) L_i(x)| \leq \|f\|_{\infty} \sum_{i=0}^n |L_i(x)| \\ &\leq \Lambda_n \|f\|_{\infty}. \end{aligned}$$

On obtient alors

$$\|\mathcal{L}_n(f)\|_{\infty} \leq \Lambda_n \|f\|_{\infty}.$$

et donc

$$\sup_{\substack{f \in \mathcal{C}^0([a, b]; \mathbb{R}) \\ f \neq 0}} \frac{\|\mathcal{L}_n(f)\|_{\infty}}{\|f\|_{\infty}} \leq \Lambda_n.$$

Pour obtenir l'égalité (5.15), il suffit donc, en reprenant les calculs précédents, de regarder si l'on peut trouver  $\bar{x} \in [a, b]$  et  $\bar{f} \in \mathcal{C}^0([a, b]; \mathbb{R})$  vérifiant

$$|\mathcal{L}_n(\bar{f})(\bar{x})| = \Lambda_n \|\bar{f}\|_{\infty}.$$

On détermine  $\bar{x}$  pour que la dernière majoration, correspondant à  $\sum_{i=0}^n |L_i(x)| \leq \Lambda_n$ , devienne une égalité. L'application  $x \mapsto \sum_{i=0}^n |L_i(x)|$  étant continue sur le fermé borné  $[a, b]$ , il existe alors  $\bar{x} \in [a, b]$  tel que

$$\Lambda_n = \max_{x \in [a, b]} \sum_{i=0}^n |L_i(x)| = \sum_{i=0}^n |L_i(\bar{x})|.$$

On va maintenant regarder s'il est possible de construire une fonction  $\bar{f} \in \mathcal{C}^0([a, b]; \mathbb{R})$  telle que

$$|\sum_{i=0}^n \bar{f}(x_i) L_i(\bar{x})| = \sum_{i=0}^n |\bar{f}(x_i) L_i(\bar{x})| = \|\bar{f}\|_\infty \sum_{i=0}^n |L_i(\bar{x})|.$$

Si c'est le cas, on aura bien le résultat souhaité.

Sans déroger à la généralité, on peut supposer les points  $x_i$  ordonnés :  $x_0 < x_1 < \dots < x_n$ . Pour avoir la première égalité, il faut que tous les  $\bar{f}(x_i) L_i(\bar{x})$  soient de même signe. On choisit le signe positif, et on impose par exemple les valeurs de  $\bar{f}(x_i)$ ,  $\forall i \in \llbracket 0, n \rrbracket$ ,

$$\begin{cases} \bar{f}(x_i) = 1, & \text{si } L_i(\bar{x}) \geq 0, \\ \bar{f}(x_i) = -1, & \text{si } L_i(\bar{x}) < 0. \end{cases}$$

Il existe une multitude de fonctions de  $\mathcal{C}^0([a, b]; \mathbb{R})$  mais il faut contrôler son maximum pour qu'il vaille 1 et avoir ainsi  $|\bar{f}(x_i)| = \|\bar{f}\|_\infty$ ,  $\forall i \in \llbracket 0, n \rrbracket$ . On peut alors choisir  $\bar{f}$  affine sur chacun des intervalles  $[x_k, x_{k+1}]$ ,  $\forall k \in \llbracket 0, n-1 \rrbracket$ , puisque l'on connaît les valeurs aux extrémités de chaque intervalle (+1 ou -1). En dehors de ces intervalles, on prend par exemple  $\bar{f}(x) = \bar{f}(x_0)$ ,  $\forall x \in [a, x_0]$ , et  $\bar{f}(x) = \bar{f}(x_n)$ ,  $\forall x \in [x_n, b]$ . Cette fonction est par construction continue sur  $[a, b]$  et vérifie

$$\begin{aligned} |\mathcal{L}_n(\bar{f})(\bar{x})| &= \left| \sum_{i=0}^n \bar{f}(x_i) L_i(\bar{x}) \right| \\ &= \sum_{i=0}^n |\bar{f}(x_i) L_i(\bar{x})| = \|\bar{f}\|_\infty \sum_{i=0}^n |L_i(\bar{x})| \\ &= \Lambda_n \|\bar{f}\|_\infty. \end{aligned}$$

Or on a

$$\|\mathcal{L}_n(\bar{f})\|_\infty = \sup_{x \in [a, b]} |\mathcal{L}_n(\bar{f})(x)| \geq |\mathcal{L}_n(\bar{f})(\bar{x})| = \Lambda_n \|\bar{f}\|_\infty.$$

ce qui donne

$$\sup_{\substack{f \in \mathcal{C}^0([a, b]; \mathbb{R}) \\ f \neq 0}} \frac{\|\mathcal{L}_n(f)\|_\infty}{\|f\|_\infty} \geq \Lambda_n.$$

□



### Théorème 5.6

Pour toute fonction  $f \in \mathcal{C}^0([a, b]; \mathbb{R})$ , on a

$$\|f - \mathcal{L}_n(f)\|_\infty \leq (1 + \Lambda_n) \inf_{Q \in \mathbb{R}_n[X]} \|f - Q\|_\infty \quad (5.16)$$

*Proof.* Soit  $Q \in \mathbb{R}_n[X]$ . Par unicité du théorème d'interpolation on a  $\mathcal{L}_n(Q) = Q$  et alors

$$\begin{aligned} \|f - \mathcal{L}_n(f)\|_\infty &= \|f - Q + \mathcal{L}_n(Q) - \mathcal{L}_n(f)\|_\infty \\ &\leq \|f - Q\|_\infty + \|\mathcal{L}_n(Q - f)\|_\infty \quad \text{par linéarité de } \mathcal{L}_n \\ &\leq \|f - Q\|_\infty + \Lambda_n \|f - Q\|_\infty \end{aligned}$$

d'où le résultat.

□

- 1 • Pour les points équidistants  $x_i = a + ih$ ,  $i \in \llbracket 0, n \rrbracket$  et  $h = (b - a)/n$ , on a la minoration suivante  
 2 (voir [3] p. 49)

$$\Lambda_n \geq \frac{2^n}{4n^2} \quad (5.17)$$

- 3 et le comportement asymptotique

$$\Lambda_n \approx \frac{2^{n+1}}{e \cdot n \ln(n)} \quad \text{quand } n \rightarrow +\infty \quad (5.18)$$

- 4 • Pour les points de Tchebychev, on a la majoration suivante : il existe  $C > 0$  tel que

$$\Lambda_n \leq C \ln(n) \quad (5.19)$$

- 5 et le comportement asymptotique

$$\Lambda_n \approx \frac{2}{\pi} \ln(n) \quad \text{quand } n \rightarrow +\infty \quad (5.20)$$



### Proposition 5.7: admis

- 6 Pour toute famille de points d'interpolation, il existe une fonction  $f \in \mathcal{C}^0([a, b]; \mathbb{R})$  telle que la suite des polynômes d'interpolation associés ne converge pas uniformément.



### Proposition 5.8: admis

- 7 Soit  $f$  une fonction lipschitzienne sur  $[a, b]$  à valeurs réelles, i.e. il existe une constante  $K \geq 0$  telle que  $\forall (x, y) \in [a, b]^2$ , on ait  $|f(x) - f(y)| \leq K|x - y|$ . Soient  $n \in \mathbb{N}^*$  et  $x_0, \dots, x_n$  les points de Tchebychev  $[a, b]$ . On note  $\mathcal{L}_n(f)$  le polynôme d'interpolation de Lagrange associés aux couples de  $(x_i, f(x_i))_{i \in \llbracket 0, n \rrbracket}$ .  
 8 Alors la suite  $(\mathcal{L}_n(f))_{n \geq 1}$  des polynômes d'interpolation converge uniformément vers  $f$  sur  $[a, b]$ .

- 9 Pour conclure, l'interpolation de Lagrange en des points équidistants n'est à utiliser qu'avec un nombre de points assez faible : des phénomènes d'instabilités pouvant apparaître.

## 5.2 Polynôme d'interpolation de Lagrange-Hermite



### Exercice 5.2.1

Soient  $(x_i, y_i, z_i)_{i \in \llbracket 0, n \rrbracket}$   $n + 1$  triplets de  $\mathbb{R}^3$ , où les  $x_i$  sont des points distincts deux à deux de l'intervalle  $[a, b]$ . Le polynôme d'interpolation de **Lagrange-Hermite**, noté  $H_n$ , associé aux  $n + 1$  triplets  $(x_i, y_i, z_i)_{i \in \llbracket 0, n \rrbracket}$ , est défini par

$$H_n(x_i) = y_i \text{ et } H'_n(x_i) = z_i, \quad \forall i \in \llbracket 0, n \rrbracket \quad (5.21)$$

**Q. 1** Quel est a priori le degré de  $H_n$  ?

On définit le polynôme  $P_n$  par

$$P_n(x) = \sum_{i=0}^n y_i A_i(x) + \sum_{i=0}^n z_i B_i(x) \quad (5.22)$$

avec, pour  $i \in \llbracket 0, n \rrbracket$ ,  $A_i$  et  $B_i$  polynômes de degré au plus  $2n + 1$  indépendants des valeurs  $y_i$  et  $z_i$ .

**Q. 2** 1. Déterminer des conditions suffisantes sur  $A_i$  et  $B_i$  pour que  $P_n \equiv H_n$ .

2. En déduire les expressions de  $A_i$  et  $B_i$  en fonction de  $L_i$  et de  $L'_i(x_i)$  où

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

**Q. 3** Démontrer qu'il existe un unique polynôme d'interpolation de Lagrange-Hermite de degré au plus  $2n + 1$  défini par (5.21).

### Correction Exercice

**Q. 1** On a  $2n + 2$  équations, donc à priori  $H_n$  est de degré  $2n + 1$ .

**Q. 2** 1. D'après (5.22) on a pour tout  $j \in \llbracket 0, n \rrbracket$

$$P_n(x_j) = \sum_{i=0}^n y_i A_i(x_j) + \sum_{i=0}^n z_i B_i(x_j)$$

Pour avoir  $P_n(x_j) = y_j$  il suffit d'avoir

$$A_i(x_j) = \delta_{i,j} \text{ et } B_i(x_j) = 0, \quad \forall i \in \llbracket 0, n \rrbracket. \quad (5.23)$$

De même, on a

$$P'_n(x_j) = \sum_{i=0}^n y_i A'_i(x_j) + \sum_{i=0}^n z_i B'_i(x_j)$$

et donc pour avoir  $P'_n(x_j) = z_j$  il suffit d'avoir

$$A'_i(x_j) = 0 \text{ et } B'_i(x_j) = \delta_{i,j}, \quad \forall i \in \llbracket 0, n \rrbracket. \quad (5.24)$$

2. Soit  $i \in \llbracket 0, n \rrbracket$ . On commence par déterminer le polynôme  $A_i \in \mathbb{R}_{2n+1}[X]$  vérifiant

$$A_i(x_j) = \delta_{i,j} \text{ et } A'_i(x_j) = 0, \quad \forall j \in \llbracket 0, n \rrbracket.$$

Les points  $(x_j)_{j \in \llbracket 0, n \rrbracket \setminus \{i\}}$  sont racines doubles de  $A_i$ . Le polynôme  $L_i \in \mathbb{R}_n[X]$  admet les mêmes racines (simples) que  $A_i$  et donc  $L_i^2 \in \mathbb{R}_{2n}[X]$  admet les mêmes racines doubles que  $A_i$ . On peut alors écrire

$$A_i(x) = \alpha_i(x) L_i^2(x) \text{ avec } \alpha_i(x) \in \mathbb{R}_1[X].$$

Il reste à déterminer le polynôme  $\alpha_i$ . Or on a

$$A_i(x_i) = 1 \text{ et } A'_i(x_i) = 0.$$

Comme  $L_i(x_i) = 1$ , on obtient

$$A_i(x_i) = \alpha_i(x_i) L_i^2(x_i) = \alpha_i(x_i) = 1$$

1 et

$$A'_i(x_i) = \alpha'_i(x_i)L_i^2(x_i) + 2\alpha_i(x_i)L'_i(x_i)L_i(x_i) = \alpha'_i(x_i) + 2\alpha_i(x_i)L'_i(x_i) = 0$$

2 c'est à dire

$$\alpha_i(x_i) = 1 \quad \text{et} \quad \alpha'_i(x_i) = -2L'_i(x_i).$$

3 Comme  $\alpha_i$  est un polynôme de degré 1 on en déduit

$$\alpha_i(x) = 1 - 2L'_i(x_i)(x - x_i)$$

4 et donc

$$A_i(x) = (1 - 2L'_i(x_i)(x - x_i))L_i^2(x). \quad (5.25)$$

5 On détermine ensuite le polynôme  $B_i \in \mathbb{R}_{2n+1}[X]$  vérifiant

$$B_i(x_j) = 0 \quad \text{et} \quad B'_i(x_j) = \delta_{i,j}, \quad \forall j \in \llbracket 0, n \rrbracket.$$

6 Les points  $(x_j)_{j \in \llbracket 0, n \rrbracket \setminus \{i\}}$  sont racines doubles de  $B_i$  et le point  $x_i$  est racine simple. Le polynôme  
7  $L_i^2 \in \mathbb{R}_{2n}[X]$  admet les mêmes racines doubles. On peut alors écrire

$$B_i(x) = C(x - x_i)L_i^2(x) \quad \text{avec} \quad C \in \mathbb{R}.$$

8 Il reste à déterminer la constante  $C$ . Or  $L_i(x_i) = 1$  et comme  $B'_i(x_i) = 1$  on obtient

$$B'_i(x_i) = CL_i^2(x_i) + 2C(x_i - x_i)L'_i(x_i)L_i(x_i) = C = 1$$

9 ce qui donne

$$B_i(x) = (x - x_i)L_i^2(x). \quad (5.26)$$

10 On vient de démontrer l'existence en construisant un polynôme de degré  $2n + 1$  vérifiant (5.21).

11 **Q. 3** Deux démonstrations pour l'unicité sont proposées (la deuxième donne aussi l'existence).

12 **dém. 1:** Soient  $P$  et  $Q$  deux polynômes de  $\mathbb{R}_{2n+1}[X]$  vérifiant (5.21). Le polynôme  $R = P - Q \in \mathbb{R}_{2n+1}[X]$   
13 admet alors  $n + 1$  racines doubles distinctes  $(x_0, \dots, x_n)$ . Or le seul polynôme de  $\mathbb{R}_{2n+1}[X]$  ayant  
14  $n + 1$  racines doubles est le polynôme nul et donc  $R = 0$ , i.e.  $P = Q$ .

15 **dém. 2:** Soit  $\Phi : \mathbb{R}_{2n+1}[X] \longrightarrow \mathbb{R}^{2n+2}$  définie par

$$\forall P \in \mathbb{R}_{2n+1}[X], \quad \Phi(P) = (P(x_0), \dots, P(x_n), P'(x_0), \dots, P'(x_n)).$$

16 L'existence et l'unicité du polynôme  $H_n$  est équivalente à la bijectivité de l'application  $\Phi$ . Or celle-ci  
17 est une application linéaire entre deux espaces de dimension  $2n + 2$ . Elle est donc bijective si et  
18 seulement si elle est injective (ou surjective). Pour vérifier l'injectivité de  $\Phi$  il est nécessaire et suffisant  
19 de vérifier que son noyau est réduit au polynôme nul.

20 Soit  $P \in \ker \Phi$ . On a alors  $\Phi(P) = \mathbf{0}_{2n+2}$  et donc  $(x_0, \dots, x_n)$  sont  $n + 1$  racines doubles distinctes  
21 de  $P$ . Or le seul polynôme de  $\mathbb{R}_{2n+1}[X]$  ayant  $n + 1$  racines doubles est le polynôme nul et donc  
22  $P = 0$ .

◇

25  **Definition 5.9**

Soient  $n \in \mathbb{N}^*$  et  $(x_i, y_i, z_i)_{i \in \llbracket 0, n \rrbracket}$   $n+1$  triplets de  $\mathbb{R}^3$ , où les  $x_i$  sont des points distincts deux à deux de l'intervalle  $[a, b]$ . Le **polynôme d'interpolation de Lagrange-Hermite**, noté  $H_n$ , associé aux  $n+1$  triplets  $(x_i, y_i, z_i)_{i \in \llbracket 0, n \rrbracket}$ , est défini par

$$H_n(x) = \sum_{i=0}^n y_i A_i(x) + \sum_{i=0}^n z_i B_i(x) \quad (5.27)$$

avec

$$A_i(x) = (1 - 2L'_i(x_i)(x - x_i))L_i^2(x) \quad \text{et} \quad B_i(x) = (x - x_i)L_i^2(x) \quad (5.28)$$

où

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

1

### Théorème 5.10

Le **polynôme d'interpolation de Lagrange-Hermite**,  $H_n$ , associé aux  $n+1$  triplets  $(x_i, y_i, z_i)_{i \in \llbracket 0, n \rrbracket}$ , est l'unique polynôme de degré au plus  $2n+1$ , vérifiant

$$H_n(x_i) = y_i \quad \text{et} \quad H'_n(x_i) = z_i, \quad \forall i \in \llbracket 0, n \rrbracket \quad (5.29)$$

2

### Exercice 5.2.2

Soit  $f \in \mathcal{C}^{2n+2}([a, b]; \mathbb{R})$ . On suppose de plus que,  $\forall i \in \llbracket 0, n \rrbracket$ ,  $x_i \in [a, b]$ ,  $y_i = f(x_i)$  et  $z_i = f'(x_i)$ . On note

$$\pi_n^2(x) = \prod_{i=0}^n (x - x_i)^2$$

et  $H_n$  le polynôme d'interpolation de Lagrange-Hermite associé aux triplets  $(x_i, f(x_i), f'(x_i))_{i \in \llbracket 0, n \rrbracket}$ .

**Q. 1** Montrer que

$$|f(x) - H_n(x)| \leq \frac{\|f^{(2n+2)}\|_\infty}{(2n+2)!} \pi_n^2(x). \quad (5.30)$$

**Indications :** Etudier les zéros de la fonction  $F(y) = f(y) - H_n(y) - \frac{f(x) - H_n(x)}{\pi_n^2(x)} \pi_n^2(y)$  et appliquer le théorème de Rolle.

3

### Correction Exercice

4

**Q. 1** Soit  $i \in \llbracket 1, n \rrbracket$ , on a  $f(x_i) - H_n(x_i) = 0$  et l'inégalité (5.30) est donc vérifiée pour  $x = x_i$ .

5

Soit  $x \in [a, b]$  tel que  $x \neq x_i, \forall i \in \llbracket 1, n \rrbracket$ . On a alors  $\pi_n^2(x) \neq 0$ . Comme  $f \in \mathcal{C}^{2n+2}([a, b]; \mathbb{R})$ ,  $H_n \in \mathbb{R}_{2n+1}[X]$  et  $\pi_n \in \mathbb{R}_{n+1}[X]$ , on en déduit que

6

$$F \in \mathcal{C}^{2n+2}([a, b]; \mathbb{R}).$$

7

On note que  $\pi_n^2$  admet  $(x_0, \dots, x_n)$  comme racines doubles distinctes. Par construction  $f - H_n$  admet les mêmes racines doubles. On en déduit alors que  $F$  admet aussi  $(x_0, \dots, x_n)$  comme racines doubles. De plus, on a  $F(x) = 0$  (i.e.  $x$  est racine simple) et donc

8

9

10

$F$  admet au moins  $2n+3$  racines (comptées avec leurs multiplicités).

11

Les points  $x, x_0, \dots, x_n$  étant distincts, la fonction  $F'$  admet par le théorème de Rolle  $n+1$  zéros distincts entre eux et distincts des points  $x, x_0, \dots, x_n$ . De plus les points  $x_0, \dots, x_n$  sont racines de  $F'$  puisque racines doubles de  $F$ . On en déduit alors que

12

13

14

$F'$  admet au moins  $2n+2$  racines distinctes deux à deux.

15

Par applications successives du théorème de Rolle, on abouti à :

16

1  $F^{(2n+2)}$  admet au moins une racine notée  $\xi_x \in ]a, b[$ .

2 On a alors

$$F^{(2n+2)}(\xi_x) = 0 = f^{(2n+2)}(\xi_x) - H_n^{(2n+2)}(\xi_x) - \frac{f(x) - H_n(x)}{\pi_n^2(x)} \frac{d^{2n+2}\pi_n^2}{dx^{2n+2}}(\xi_x)$$

3 Comme  $H_n \in \mathbb{R}_{2n+1}[X]$  on a  $H_n^{(2n+2)} \equiv 0$ . De plus comme  $\pi_n^2(x) = \prod_{i=0}^n (x - x_i)^2 \in \mathbb{R}_{2n+2}[X]$  sa dérivée

4 d'ordre  $2n + 2$  est constante et

$$\frac{d^{2n+2}\pi_n^2}{dx^{2n+2}} = (2n + 2)!$$

5 On en déduit alors

$$f^{(2n+2)}(\xi_x) = \frac{f(x) - H_n(x)}{\pi_n^2(x)} (2n + 2)!$$

6 On a donc montrer que  $\forall x \in [a, b] \exists \xi_x \in ]a, b[$  tels que

$$f(x) - H_n(x) = \frac{\pi_n^2(x)}{(2n + 2)!} f^{(2n+2)}(\xi_x).$$

7 Comme  $\pi_n^2(x) \geq 0$  on obtient bien (5.30).

8

9



### Théorème 5.11

Soient  $n \in \mathbb{N}^*$  et  $x_0, \dots, x_n, n + 1$  points distincts de l'intervalle  $[a, b]$ . Soient  $f \in \mathcal{C}^{2n+2}([a, b]; \mathbb{R})$  et  $H_n$  le polynôme d'interpolation de Lagrange-Hermite associé aux  $n + 1$  triplets  $(x_i, f(x_i), f'(x_i))_{i \in [0, n]}$ . On a alors  $\forall x \in [a, b], \exists \xi_x \in (\min(x_i, x), \max(x_i, x))$ , tels que

$$f(x) - H_n(x) = \frac{f^{(2n+2)}(\xi_x)}{(2n + 2)!} \prod_{i=0}^n (x - x_i)^2 \quad (5.31)$$

10

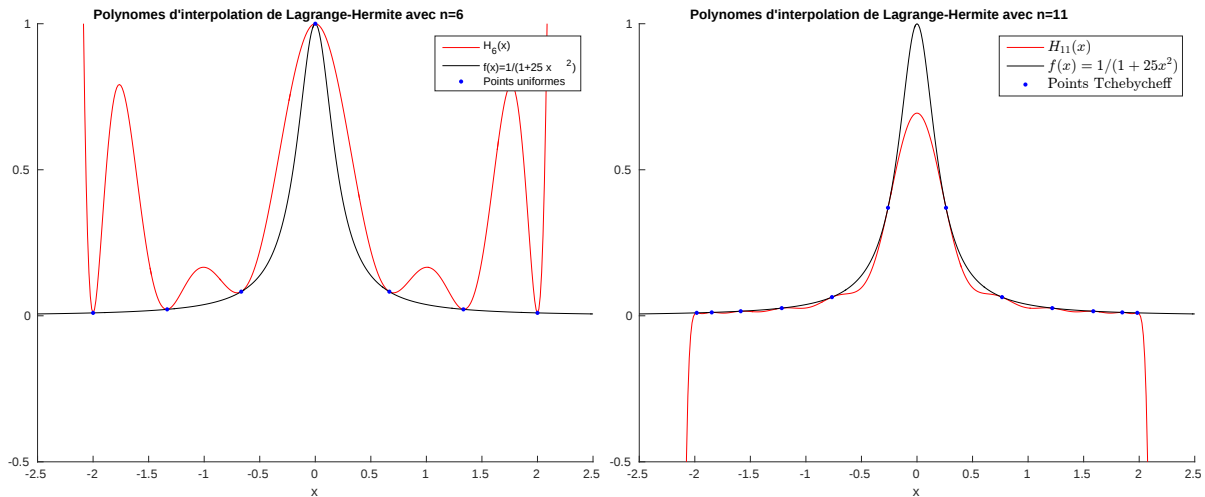


Figure 5.9: Polynôme d'interpolation de lagrange-Hermite avec  $n = 6$  (7 points) pour la fonction  $f : x \rightarrow 1/(1 + 25x^2)$ . A gauche avec des points uniformément répartis et à droite avec des points de Tchebychev

11



### Exercice 5.2.3

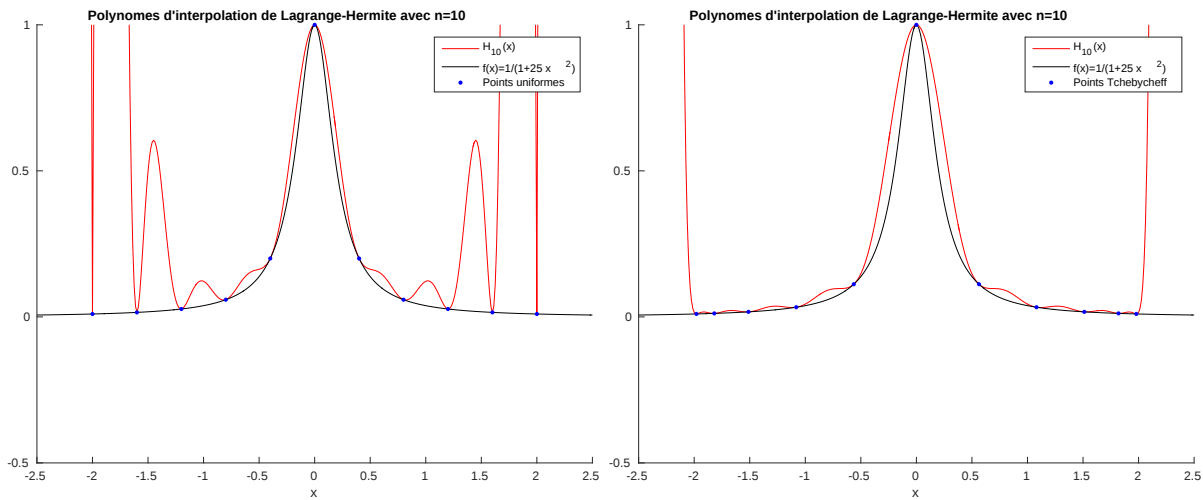


Figure 5.10: Polynôme d'interpolation de lagrange-Hermite avec  $n = 10$  (11 points) pour la fonction  $f : x \rightarrow 1/(1 + 25x^2)$ . A gauche avec des points uniformément répartis et à droite avec des points de Tchebychev

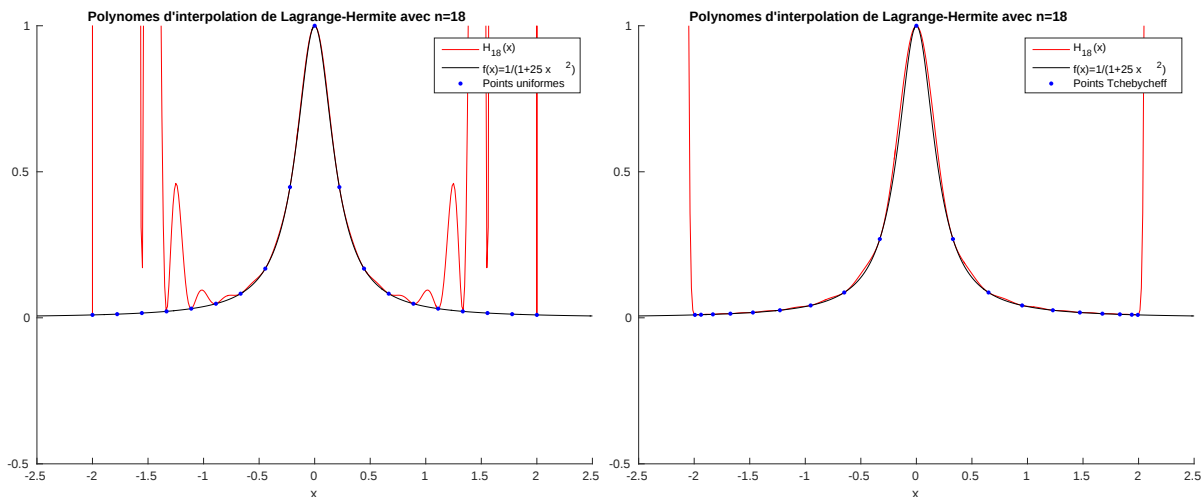


Figure 5.11: Polynôme d'interpolation de lagrange-Hermite avec  $n = 18$  (19 points) pour la fonction  $f : x \rightarrow 1/(1 + 25x^2)$ . A gauche avec des points uniformément répartis et à droite avec des points de Tchebychev

Ecrire une fonction algorithmique **HERMITE** permettant de calculer  $H_n$  (polynôme d'interpolation de Lagrange-Hermite associé aux  $n + 1$  triplets  $(x_i, y_i, z_i)_{i \in \llbracket 0, n \rrbracket}$ ) en  $t \in \mathbb{R}$ .

### Correction Exercice

**But :** Calculer le polynôme  $H_n(t)$  défini par (5.27)

**Données :**  $\mathbf{X}$  : vecteur/tableau de  $\mathbb{R}^{n+1}$ ,  $X(i) = x_{i-1} \forall i \in \llbracket 1, n+1 \rrbracket$  et  $X(i) \neq X(j)$  pour  $i \neq j$ ,  
 $\mathbf{Y}$  : vecteur/tableau de  $\mathbb{R}^{n+1}$ ,  $Y(i) = y_{i-1} \forall i \in \llbracket 1, n+1 \rrbracket$ ,  
 $\mathbf{Z}$  : vecteur/tableau de  $\mathbb{R}^{n+1}$ ,  $Z(i) = z_{i-1} \forall i \in \llbracket 1, n+1 \rrbracket$ ,  
 $t$  : un réel.

**Résultat :**  $\text{pH}$  : le réel  $\text{pH} = H_n(t)$ .

D'après la Définition 5.9, on a

$$H_n(t) = \sum_{i=0}^n y_i A_i(t) + \sum_{i=0}^n z_i B_i(t) = \sum_{i=0}^n (y_i A_i(t) + z_i B_i(t))$$

avec

$$A_i(t) = (1 - 2L'_i(x_i)(t - x_i))L_i^2(t) \quad \text{et} \quad B_i(t) = (t - x_i)L_i^2(t)$$

où

$$L_i(t) = \prod_{\substack{j=0 \\ j \neq i}}^n \frac{t - x_j}{x_i - x_j}.$$

Pour rendre effectif le calcul de  $H_n(t)$ , il reste à déterminer  $L'_i(x_i)$ . On a

$$L'_i(x) = \sum_{\substack{k=0 \\ k \neq i}}^n \frac{1}{x_i - x_k} \prod_{\substack{j=0 \\ j \neq i \\ j \neq k}}^n \frac{t - x_j}{x_i - x_j}$$

d'où

$$L'_i(x_i) = \sum_{\substack{k=0 \\ k \neq i}}^n \frac{1}{x_i - x_k}. \quad (5.32)$$

La fonction que l'on va écrire use (et certains diront abuse) de fonctions.

---

**Algorithme 5.2** Fonction **HERMITE** permettant de calculer le polynôme d'interpolation de Lagrange-Hermite  $H_n(t)$  défini par (5.27)

---

```

1: Fonction $\text{pH} \leftarrow \text{HERMITE} (X, Y, Z, t)$
2: $\text{pH} \leftarrow 0$
3: Pour $i \leftarrow 0$ à n faire
4: $\text{pH} \leftarrow \text{pH} + \text{POLYA}(i, X, t) * Y(i+1) + \text{POLYB}(i, X, t) * Z(i+1)$
5: Fin Pour
6: Fin Fonction
```

---

Les différentes fonctions utilisées pour la fonction **HERMITE** (directement ou indirectement) sont les suivantes :

**POLYA** : calcul du polynôme  $A_i$  en  $t$ , (données  $i, X, t$ )

**POLYB** : calcul du polynôme  $B_i$  en  $t$ , (données  $i, X, t$ )

**POLYL** : calcul du polynôme  $L_i$  en  $t$ , (données  $i, X, t$ )

**POLYLP** : calcul de  $L'_i(x_i)$ , (données  $i, X$ )

**Algorithme 5.3** Fonction **POLYA** permettant de calculer le polynôme  $A_i$  en  $t \in \mathbb{R}$  donné par  $A_i(t) = (1 - 2L'_i(x_i)(t - x_i))L_i^2(t)$

```
1: Fonction $y \leftarrow \text{POLYA} (i, X, t)$
2: $y \leftarrow (1 - 2 * \text{POLYLP}(i, X) * (t - X(i + 1))) * (\text{POLYL}(i, X, t))^2$
3: Fin Fonction
```

**Algorithme 5.5** Fonction **POLYL** permettant de calculer le polynôme  $L_i$  en  $t \in \mathbb{R}$  donné par  $L_i(t) = \prod_{j=0, j \neq i}^n \frac{t - x_j}{x_i - x_j}$

```
1: Fonction $y \leftarrow \text{POLYL} (i, X, t)$
2: $y \leftarrow 1$
3: Pour $j \leftarrow 0$ à n , ($j \sim i$) faire
4: $y \leftarrow y * (t - X(j + 1)) / (X(i + 1) - X(j + 1))$
5: Fin Pour
6: Fin Fonction
```

**Algorithme 5.4** Fonction **POLYB** permettant de calculer le polynôme  $B_i$  en  $t \in \mathbb{R}$  donné par  $B_i(t) = (t - x_i)L_i^2(t)$

```
1: Fonction $y \leftarrow \text{POLYB} (i, X, t)$
2: $y \leftarrow (t - X(i + 1)) * (\text{POLYL}(i, X, t))^2$
3: Fin Fonction
```

**Algorithme 5.6** Fonction **POLYLP** permettant de calculer  $L'_i(x_i) = \sum_{k=0, k \neq i}^n \frac{1}{x_i - x_k}$

```
1: Fonction $y \leftarrow \text{POLYLP} (i, X)$
2: $y \leftarrow 0$
3: Pour $k \leftarrow 0$ à n , ($k \sim i$) faire
4: $y \leftarrow y + 1 / (X(i + 1) - X(k + 1))$
5: Fin Pour
6: Fin Fonction
```

Bien évidemment une telle écriture est loin d'être optimale mais elle a l'avantage d'être facile à programmer et facile à lire car elle "colle" aux formules mathématiques.

On laisse le soin au lecteur d'écrire des fonctions plus performantes...

## 5.3 Exercices



### Exercice 5.3.1

**Q. 1** 1. Ecrire explicitement un polynôme  $P$  de degré 2 passant par les points  $A = (1; 2)$ ,  $B = (2; 6)$  et  $C = (3; 12)$ .

2. Démontrer que le polynôme  $P$  est l'unique polynôme de degré 2 passant par les points  $A$ ,  $B$  et  $C$ .

**Q. 2** Ecrire explicitement un polynôme  $Q$  de degré 3 tel que

$$\begin{aligned} Q(1) &= 4, & Q(2) &= 5, \\ Q'(1) &= 3, & Q'(2) &= 2. \end{aligned}$$



### Exercice 5.3.2

**Q. 1** Construire les polynômes  $h_{00}$ ,  $h_{10}$ ,  $h_{01}$  et  $h_{11}$  de degré 3 vérifiant

$$h_{00}(0) = 1, h'_{00}(0) = h_{00}(1) = h'_{00}(1) = 0, \quad (5.33)$$

$$h_{10}(1) = 1, h_{10}(0) = h'_{00}(0) = h'_{10}(1) = 0, \quad (5.34)$$

$$h'_{01}(0) = 1, h_{01}(0) = h_{01}(1) = h'_{01}(1) = 0, \quad (5.35)$$

$$h'_{11}(1) = 1, h_{11}(0) = h'_{11}(0) = h_{11}(1) = 0; \quad (5.36)$$

On pose

$$P(x) = \alpha h_{00}(x) + \beta h_{10}(x) + \gamma h_{01}(x) + \delta h_{11}(x). \quad (5.37)$$

**Q. 2** Quelles sont les particularités de  $P$ ?

Soient  $a$  et  $b$  deux réels,  $a < b$  et  $Q$  le polynôme de degré 3 vérifiant

$$Q(a) = u_a, \quad Q'(a) = v_a, \quad Q(b) = u_b \quad \text{et} \quad Q'(b) = v_b.$$

**Q. 3** Exprimer le polynôme  $Q$  avec les fonctions  $h_{00}$ ,  $h_{10}$ ,  $h_{01}$  et  $h_{11}$ .

**Exercice 5.3.3**

Soit  $t_0 < t_1$  deux nombres réels et soit  $\varepsilon$  un réel tel que  $0 < \varepsilon < \frac{t_1 - t_0}{2}$ .

**Q. 1** *Expliciter un polynôme  $P_\varepsilon$  de degré 3 tel que*

$$P_\varepsilon(t_0) = P_\varepsilon(t_0 + \varepsilon) = 1, \quad (5.38)$$

$$P_\varepsilon(t_1) = P_\varepsilon(t_1 + \varepsilon) = 0. \quad (5.39)$$

On note  $\Phi_0(t) = \lim_{\varepsilon \rightarrow 0} P_\varepsilon(t)$ .

**Q. 2** 1. *Montrer que  $\Phi_0(t_0) = 1$ ,  $\Phi'_0(t_0) = 0$ ,  $\Phi_0(t_1) = 0$  et  $\Phi'_0(t_1) = 0$  (i.e.  $\Phi_0$  est une fonction de base des polynômes de degré 3 pour l'interpolation de Hermite).*

2. *Peut-on obtenir toutes les fonctions de base de Hermite par des procédés analogues. Si oui, expliquer comment!*

1

**Exercice 5.3.4:**

Soient  $(x_i)_{i \in \mathbb{N}}$  une suite de points distincts de l'intervalle  $[a, b]$  et  $f$  une fonction définie sur  $[a, b]$  à valeurs réelles. On désigne par  $f[\ ]$  les **différences divisées** de la fonction  $f$  définie par

$$f[x_i] = f(x_i), \quad \forall i \in \mathbb{N}, \quad (\text{ordre } 0) \quad (5.40)$$

et

$$f[x_k, \dots, x_{k+r}] = \frac{f[x_{k+1}, \dots, x_{k+r}] - f[x_k, \dots, x_{k+r-1}]}{x_{k+r} - x_k}, \quad \forall k \in \mathbb{N}, \quad \forall r \in \mathbb{N}^*, \quad (\text{ordre } r) \quad (5.41)$$

**Q. 1** Montrer que

$$f[x_k, \dots, x_{k+r}] = \sum_{i=k}^{k+r} \frac{f(x_i)}{\prod_{\substack{j=k \\ j \neq i}}^{k+r} (x_i - x_j)}, \quad \forall k \in \mathbb{N}, \quad \forall r \in \mathbb{N}^*. \quad (5.42)$$

**Q. 2** Soit  $\sigma$  une permutation des entiers  $\{k, \dots, k+r\}$ . Montrer que

$$f[x_k, \dots, x_{k+r}] = f[x_{\sigma(1)}, \dots, x_{\sigma(r+1)}]. \quad (5.43)$$

On note  $Q_{k,r}$  le polynôme d'interpolation associé aux  $r+1$  couples  $(x_{k+i}, f(x_{k+i}))_{i \in [0, r]}$ .

**Q. 3** 1. Exprimer le polynôme  $Q_{k,1}$  en fonction de  $Q_{k,0}$ .

2. Exprimer le polynôme  $Q_{k,1}$  en fonction de  $Q_{k,0}$  et  $Q_{k+1,0}$ .

3. En déduire que

$$Q_{k,1}(x) = Q_{k,0}(x) + f[x_k, x_{k+1}](x - x_k). \quad (5.44)$$

**Q. 4** 1. Exprimer le polynôme  $Q_{k,2}$  en fonction de  $Q_{k,1}$ .

2. Exprimer le polynôme  $Q_{k,2}$  en fonction de  $Q_{k,1}$  et  $Q_{k+1,1}$ .

3. En déduire que

$$Q_{k,2}(x) = Q_{k,1}(x) + f[x_k, x_{k+1}, x_{k+2}](x - x_k)(x - x_{k+1}). \quad (5.45)$$

**Q. 5** 1. Montrer que

$$Q_{k,r}(x) = Q_{k,r-1}(x) + f[x_k, \dots, x_{k+r}] \prod_{j=0}^{r-1} (x - x_{k+j}) \quad (5.46)$$

**Indication :** Effectuer une démonstration par récurrence en écrivant le polynôme  $Q_{k,r}$  sous deux formes : l'une en fonction de  $Q_{k,r-1}$  et l'autre en fonction de  $Q_{k,r-1}$  et  $Q_{k+1,r-1}$ .

2. En déduire

$$Q_{k,r}(x) = f[x_k] + \sum_{i=1}^r f[x_k, \dots, x_{k+i}] \prod_{j=0}^{i-1} (x - x_{k+j}) \quad (5.47)$$

**Q. 6** On suppose que  $f \in C^r([a, b]; \mathbb{R})$ . Montrer qu'il existe  $\xi \in [\min_{i \in [0, r]} x_{k+i}, \max_{i \in [0, r]} x_{k+i}]$  tel que

$$f[x_k, \dots, x_{k+r}] = \frac{f^{(r)}(\xi)}{r!}. \quad (5.48)$$

**Q. 7** On suppose  $f \in C^{r+1}([a, b]; \mathbb{R})$ . Montrer que,  $\forall x \in [a, b]$ , il existe  $\xi_x$  appartenant au plus petit intervalle fermé contenant  $x, x_k, \dots, x_{k+r}$  tel que

$$f(x) - Q_{k,r}(x) = \frac{\prod_{j=0}^r (x - x_{k+j})}{(r+1)!} f^{(r+1)}(\xi_x). \quad (5.49)$$

**Exercice 5.3.5**

Soit  $\mathbf{X} = (x_i)_{i \in \llbracket 0, n \rrbracket}$   $n+1$  points deux à deux distincts de l'intervalle  $[a, b]$ . On note  $s$  le changement de variables  $s : t \longrightarrow a + (b-a)t$  de  $[0, 1]$  à valeurs dans  $[a, b]$ . Pour tout  $i \in \llbracket 0, n \rrbracket$ , on note  $t_i = s^{-1}(x_i) = \frac{x_i - a}{b - a}$  et  $\mathbf{T} = (t_i)_{i \in \llbracket 0, n \rrbracket}$ .

Les polynômes d'interpolation de Lagrange  $\mathcal{L}_n(f)$  et  $\mathcal{L}_n(g)$  associés respectivement aux points  $(x_i, f(x_i))_{i \in \llbracket 0, n \rrbracket}$  et  $(t_i, g(t_i))_{i \in \llbracket 0, n \rrbracket}$  sont définis par

$$\begin{aligned}\mathcal{L}_n(f)(x) &= \sum_{i=0}^n L_i^{\mathbf{X}}(x) f(x_i) \quad \text{avec} \quad L_i^{\mathbf{X}}(x) = \prod_{\substack{j=0 \\ j \neq i}}^n \frac{x - x_j}{x_i - x_j} \\ \mathcal{L}_n(g)(t) &= \sum_{i=0}^n L_i^{\mathbf{T}}(t) g(t_i) \quad \text{avec} \quad L_i^{\mathbf{T}}(t) = \prod_{\substack{j=0 \\ j \neq i}}^n \frac{t - t_j}{t_i - t_j}\end{aligned}$$

**Q. 1** Montrer que  $L_i^{\mathbf{X}} \circ s = L_i^{\mathbf{T}}$ .

**Q. 2** En déduire que  $\mathcal{L}_n(f \circ s) = \mathcal{L}_n(f) \circ s = \mathcal{L}_n(g)$ .

1

## 2 Correction Exercice

**Q. 1**

$$\begin{aligned}L_i^{\mathbf{X}} \circ s(t) &= \prod_{\substack{j=0 \\ j \neq i}}^n \frac{s(t) - x_j}{x_i - x_j} = \prod_{\substack{j=0 \\ j \neq i}}^n \frac{s(t) - s(t_j)}{s(t_i) - s(t_j)} \\ &= \prod_{\substack{j=0 \\ j \neq i}}^n \frac{a + t(b-a) - (a + t_j(b-a))}{a + t_i(b-a) - (a + t_j(b-a))} \\ &= \prod_{\substack{j=0 \\ j \neq i}}^n \frac{t - t_j}{t_i - t_j} = L_i^{\mathbf{T}}(t)\end{aligned}$$

**Q. 2** On a

$$\begin{aligned}\mathcal{L}_n(f \circ s)(t) &\stackrel{\text{def}}{=} \sum_{i=0}^n L_i^{\mathbf{T}}(t) f \circ s(t_i) \\ &= \sum_{i=0}^n L_i^{\mathbf{T}}(t) g(t_i) = \mathcal{L}_n(g)\end{aligned}$$

et

$$\begin{aligned}\mathcal{L}_n(f) \circ s(t) &\stackrel{\text{def}}{=} \sum_{i=0}^n L_i^{\mathbf{X}} \circ s(t) f(x_i) \\ &= \sum_{i=0}^n L_i^{\mathbf{T}}(t) f \circ s(t_i) = \mathcal{L}_n(g)\end{aligned}$$

5

6

◇

7



### Exercice 5.3.6: Spline cubique

Soient  $n \geq 3$  un entier et  $a = x_0 < x_1 < \dots < x_{n-1} < x_n = b$  une discrétisation régulière de l'intervalle  $[a, b]$ . On note  $h = x_{k+1} - x_k$ . Une fonction  $s$  définie sur  $[a, b]$  à valeurs réelles s'appelle **spline cubique** si elle est deux fois continûment différentiable et si, sur chaque intervalle  $[x_{k-1}, x_k]$ , elle est polynomiale de degré inférieur ou égal à 3.

Soit  $f \in \mathcal{C}^2([a, b]; \mathbb{R})$  et  $s$  une spline cubique vérifiant

$$s(x_i) = f(x_i) = f_i, \quad \forall i \in \llbracket 0, n \rrbracket. \quad (5.50)$$

**Q. 1** Montrer que si

$$s''(b)(f'(b) - s'(b)) = s''(a)(f'(a) - s'(a)) \quad (5.51)$$

alors

$$\int_a^b (s''(x))^2 dx \leq \int_a^b (f''(x))^2 dx. \quad (5.52)$$

**Indications :** Poser  $r = f - s$  et montrer par intégrations par parties que  $\int_a^b s''(x)r''(x)dx = 0$ .

Soient  $k \in \llbracket 1, n \rrbracket$  et  $S_k$  un polynôme de degré inférieur ou égal à 3 vérifiant

$$\begin{cases} S_k(x_{k-1}) &= f_{k-1} \end{cases} \quad (5.53a)$$

$$\begin{cases} S_k(x_k) &= f_k \end{cases} \quad (5.53b)$$

$$\begin{cases} S_k''(x_{k-1}) &= m_{k-1} \end{cases} \quad (5.53c)$$

$$\begin{cases} S_k''(x_k) &= m_k. \end{cases} \quad (5.53d)$$

**Q. 2** 1. Montrer l'existence et l'unicité du polynôme  $S_k$ .

2. Montrer que polynôme  $S_k$  peut s'écrire sous la forme

$$S_k(x) = a_k(x_k - x)^3 + b_k(x - x_{k-1})^3 + \alpha_k(x_k - x) + \beta_k(x - x_{k-1}) \quad (5.54)$$

en explicitant les coefficients  $(a_k, b_k, \alpha_k, \beta_k)$  en fonction de  $(f_{k-1}, f_k, m_{k-1}, m_k)$  et  $h$ .

On note  $g$  la fonction dont la restriction à chaque intervalle  $[x_{k-1}, x_k]$ ,  $k \in \llbracket 1, n \rrbracket$ , est  $S_k$ .

**Q. 3** 1. Vérifier que  $g$  est bien définie sur  $[a, b]$ .

2. Montrer que  $g$  est une spline cubique si et seulement si,  $\forall k \in \llbracket 1, n-1 \rrbracket$ ,

$$m_{k+1} + 4m_k + m_{k-1} = \frac{6}{h^2}(f_{k+1} - 2f_k + f_{k-1}). \quad (5.55)$$

**Q. 4** 1. Montrer qu'une condition nécessaire et suffisante pour que  $g$  soit une spline cubique et vérifie  $g''(a) = 0$ ,  $g''(b) = 0$ , est que le vecteur  $\mathbf{M} \in \mathbb{R}^{n+1} = (m_0, m_1, \dots, m_n)^t$  soit solution d'un système linéaire de la forme

$$\mathbb{A}\mathbf{M} = \mathbf{b} \quad (5.56)$$

que l'on précisera.

2. Montrer que la matrice  $\mathbb{A}$  est inversible.

1

## Correction Exercice

2



### Théorème 5.12: Intégration par parties

Soient  $u \in \mathcal{C}^1([a, b]; \mathbb{R})$  et  $v \in \mathcal{C}^1([a, b]; \mathbb{R})$  alors

$$\int_a^b u'(x)v(x)dx = [u(x)v(x)]_a^b - \int_a^b u(x)v'(x)dx.$$

3

**Q. 1** On pose  $r = f - s$ . On a alors  $r \in \mathcal{C}^2([a, b]; \mathbb{R})$  et,

4

$$\forall i \in \llbracket 0, n \rrbracket, \quad r(x_i) = 0. \quad (5.57)$$

1 De plus

$$\begin{aligned}\int_a^b (f''(x))^2 dx &= \int_a^b (s''(x) + r''(x))^2 dx \\ &= \int_a^b (s''(x))^2 dx + 2 \int_a^b s''(x)r''(x) dx + \int_a^b (r''(x))^2 dx\end{aligned}\quad (5.58)$$

Montrons que  $\int_a^b s''(x)r''(x) dx = 0$ .

On ne peut effectuer une intégration par partie pour  $\int_a^b s''(x)r''(x) dx$  car  $r''$  et  $s''$  ne sont pas dérivables. Par contre, on a

$$\int_a^b s''(x)r''(x) dx = \sum_{i=1}^n \int_{x_{i-1}}^{x_i} s''(x)r''(x) dx$$

2 et, sur chaque intervalle  $[x_{i-1}, x_i]$ ,  $s''$  est un polynôme de degré au plus 1. On a donc  $s'' \in \mathcal{C}^1([x_{i-1}, x_i]; \mathbb{R})$   
 3 et  $r' \in \mathcal{C}^1([x_{i-1}, x_i]; \mathbb{R})$  et il est alors possible de faire une intégration par partie avec  $u = r'$  et  $v = s''$   
 4 sur  $[x_{i-1}, x_i]$  :

$$\int_{x_{i-1}}^{x_i} s''(x)r''(x) dx = [s''(x)h'(x)]_{x_{i-1}}^{x_i} - \int_{x_{i-1}}^{x_i} s'''(x)r'(x) dx. \quad (5.59)$$

Or, sur  $[x_{i-1}, x_i]$ ,  $s \in \mathbb{R}^3[X]$  et donc  $s'''$  est constante, ce qui donne, en utilisant (5.57),

$$\int_{x_{i-1}}^{x_i} s'''(x)r'(x) dx = s'''(x_i) \int_{x_{i-1}}^{x_i} r'(x) dx = s'''(x_i)(r(x_i) - r(x_{i-1})) = 0.$$

De (5.59), On a alors,  $\forall i \in \llbracket 1, n \rrbracket$ ,

$$\int_{x_{i-1}}^{x_i} s'''(x)r'(x) dx = s''(x_i)h'(x_i) - s''(x_{i-1})h'(x_{i-1}).$$

En sommant, on abouti a

$$\int_a^b s''(x)r''(x) dx = s''(x_n)h'(x_n) - s''(x_0)h'(x_0) = s''(b)h'(b) - s''(a)h'(a).$$

Sous l'hypothèse (5.51) on a bien  $\int_a^b s''(x)r''(x) dx = 0$ . L'équation (5.58) devient alors

$$\int_a^b (f''(x))^2 dx = \int_a^b (s''(x))^2 dx + \int_a^b (r''(x))^2 dx.$$

D'où

$$\int_a^b (f''(x))^2 dx \geq \int_a^b (s''(x))^2 dx.$$

**Q. 2** 1. Soit  $\Phi_k : \mathbb{R}^3[X] \longrightarrow \mathbb{R}^4$  définie par

$$\Phi_k(P) = (P(x_{k-1}), P(x_k), P''(x_{k-1}), P''(x_k))^t.$$

L'existence et l'unicité du polynôme  $S_k$  est équivalente à la bijectivité de  $\Phi_k$ . Cette dernière étant une application entre deux espaces vectoriels de même dimension finie 4, elle est bijective si et seulement si elle est injective. Pour établir l'injectivité de  $\Phi_k$  il faut montrer que son noyau est réduit au polynôme nul.

Soit  $P \in \ker \Phi_k$  alors  $\Phi_k(P) = 0_{\mathbb{R}^4}$ . On en déduit que  $x_{k-1}$  et  $x_k$  sont racines de  $P$ , et  $P$  s'écrit alors sous la forme

$$P(x) = (x - x_{k-1})(x - x_k)Q(x)$$

avec  $Q(x) = \alpha x + \beta$  polynôme de degré 1.

On a

$$P'(x) = (x - x_{k-1})Q(x) + (x - x_k)Q(x) + \alpha(x - x_{k-1})(x - x_k)$$

et

$$P''(x) = 2Q(x) + \alpha(x - x_{k-1}) + \alpha(x - x_k).$$

Comme  $P''(x_{k-1}) = P''(x_k) = 0$ , on obtient

$$\begin{cases} Q(x_{k-1}) + \alpha(x_{k-1} - x_k) = 0, \\ Q(x_k) + \alpha(x_k - x_{k-1}) = 0, \end{cases} \iff \begin{cases} \alpha(x_{k-1} - h) + \beta = 0, \\ \alpha(x_k + h) + \beta = 0, \end{cases}$$

En soustrayant la première équation à la deuxième, on obtient  $3\alpha h = 0$ . Comme  $h \neq 0$ , on obtient  $\alpha = \beta = 0$ . D'où  $Q \equiv 0$  et donc  $P \equiv 0$ .

2. On a  $S''_k(x) = 6a_k(x_k - x) + 6b_k(x - x_{k-1})$ . Pour déterminer  $a_k$  et  $b_k$ , on utilise les équations (5.53c) et (5.53d) qui deviennent respectivement  $6ha_k = m_{k-1}$  et  $6hb_k = m_k$ . On obtient

$$a_k = \frac{m_{k-1}}{6h} \text{ et } b_k = \frac{m_k}{6h}.$$

Pour déterminer  $\alpha_k$  et  $\beta_k$  on utilise les équations (5.53c) et (5.53d) qui deviennent respectivement

$$a_k h^3 + \alpha_k h = f_{k-1} \text{ et } b_k h^3 + \beta_k h = f_k.$$

En remplaçant  $a_k$  et  $b_k$  par leurs valeurs, on obtient

$$\alpha_k = \frac{f_{k-1}}{h} - \frac{h}{6} m_{k-1} \text{ et } \beta_k = \frac{f_k}{h} - \frac{h}{6} m_k.$$

- Q. 3** 1. On a par définition  $\forall k \in \llbracket 1, n \rrbracket$ ,  $g(x) = S_k(x)$ ,  $\forall x \in [x_{k-1}, x_k]$ . Le problème de définition de  $g$  provient du fait que  $g$  est définie deux fois en  $x_k$ ,  $k \in \llbracket 1, n-1 \rrbracket$ . En effet, on a

$$g(x_k) = S_k(x_k) \text{ et } g(x_k) = S_{k+1}(x_k).$$

Or par construction des  $S_k$ , on a  $S_k(x_k) = S_{k+1}(x_k) = f_k$  et donc la fonction  $g$  est bien définie sur  $[a, b]$ .

2. Par construction, sur chaque intervalle  $[x_{k-1}, x_k]$ , la fonction  $g$  est polynomiale de degré inférieur ou égal à 3. Pour quelle soit un spline cubique, il reste à démontrer qu'elle est deux fois continûment différentiable sur  $[a, b]$ . Il suffit pour cela de vérifier qu'en chaque point  $x_k$ ,  $k \in \llbracket 1, n-1 \rrbracket$ , la fonction  $g$  est continue et admet des dérivées premières et secondes. La continuité est immédiate puisque  $g(x_k) = f_k$ . Pour les dérivées premières et secondes, il faut que leurs limites à gauche et à droite soient égales, c'est à dire

$$\forall k \in \llbracket 1, n-1 \rrbracket, S'_k(x_k) = S'_{k+1}(x_k) \text{ et } S''_k(x_k) = S''_{k+1}(x_k).$$

Par construction des  $S_k$ , la seconde équation est immédiate :  $S''_k(x_k) = S''_{k+1}(x_k) = m_k$ . On a  $S'_k(x) = -3a_k(x_k - x)^2 + 3b_k(x - x_{k-1})^2 - \alpha_k + \beta_k$  et donc

$$S'_k(x_k) = 3b_k h^2 - \alpha_k + \beta_k = \frac{h}{2} m_k + \frac{1}{h} (f_k - f_{k-1}) - \frac{h}{6} (m_k - m_{k-1})$$

De même, on obtient

$$S'_{k+1}(x_k) = -3a_{k+1} h^2 - \alpha_{k+1} + \beta_{k+1} = -\frac{h}{2} m_k + \frac{1}{h} (f_{k+1} - f_k) - \frac{h}{6} (m_{k+1} - m_k)$$

Donc  $g$  sera dérivable en  $x_k$  si  $S'_k(x_k) = S'_{k+1}(x_k)$ , c'est à dire si

$$\frac{h}{2} m_k + \frac{1}{h} (f_k - f_{k-1}) - \frac{h}{6} (m_k - m_{k-1}) = -\frac{h}{2} m_k + \frac{1}{h} (f_{k+1} - f_k) - \frac{h}{6} (m_{k+1} - m_k)$$

ce qui s'écrit encore,  $\forall k \in \llbracket 1, n-1 \rrbracket$ ,

$$m_{k+1} + 4m_k + m_{k-1} = \frac{6}{h^2} (f_{k+1} - 2f_k + f_{k-1})$$

- Q. 4** 1. La condition  $g''(a) = 0$  se traduit par  $S''_1(x_0) = 0$  or par (5.53c) avec  $k = 1$  on a  $S''_1(x_0) = m_0$  d'où  $m_0 = 0$ .  
La condition  $g''(b) = 0$  se traduit par  $S''_n(x_n) = 0$  or par (5.53d) avec  $k = n$  on a  $S''_n(x_n) = m_n$  d'où  $m_n = 0$ .

Pour déterminer les  $m_k$ ,  $k \in \llbracket 0, n \rrbracket$ , on a  $n + 1$  équations linéaires qui s'écrivent sous la forme matricielle  $\mathbb{A}\mathbf{M} = \mathbf{b}$  avec

$$\mathbb{A} = \begin{pmatrix} 1 & 0 & 0 & \dots & 0 \\ 1 & 4 & 1 & \ddots & \vdots \\ 0 & \ddots & \ddots & \ddots & 0 \\ \vdots & \ddots & 1 & 4 & 1 \\ 0 & \dots & 0 & 0 & 1 \end{pmatrix} \text{ et } \mathbf{b} = \frac{6}{h^2} \begin{pmatrix} 0 \\ f_0 - 2f_1 + f_2 \\ \vdots \\ f_{n-2} - 2f_{n-1} + f_n \\ 0 \end{pmatrix}$$

- 1      2. La matrice  $\mathbb{A}$  est à diagonale strictement dominante : elle est donc inversible.

2

◇

3

# Chapitre 6

## Intégration numérique

Soit  $f$  une fonction définie et intégrable sur un intervalle  $[a, b]$  donné. On propose de chercher des approximations de

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

dans le cas où l'on ne connaît pas de primitive de  $f$ .

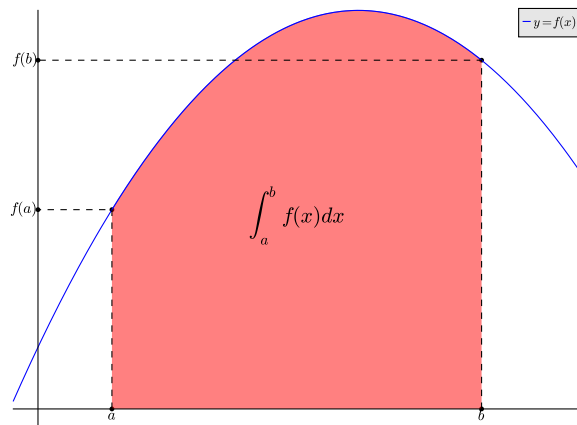


Figure 6.1: Représentation de  $\int_a^b f(x) dx$  (aire de la surface colorée)

### ♥ Definition 6.1

Soient  $f \in \mathcal{C}^0([a, b]; \mathbb{R})$  et  $\mathcal{Q}_n(f, a, b)$  la formule de quadrature donnée par :

$$\mathcal{Q}_n(f, a, b) \stackrel{\text{def}}{=} (b - a) \sum_{j=0}^n w_j f(x_j) \quad (6.1)$$

avec  $\forall j \in \llbracket 0, n \rrbracket$   $w_j \in \mathbb{R}$  et  $x_j \in [a, b]$ . L'erreur associée à cette formule de quadrature, notée  $\mathcal{E}_{a,b}(f)$ , est définie par

$$\mathcal{E}_{a,b}(f) = \int_a^b f(x) dx - \mathcal{Q}_n(f, a, b), \quad \forall f \in \mathcal{C}^0([a, b]; \mathbb{R}) \quad (6.2)$$

### ♥ Definition 6.2

On dit qu'une formule d'intégration (ou formule de quadrature) est d'ordre  $p$  ou a pour **degré d'exactitude**  $p$  si elle est exacte pour les polynômes de degré inférieur ou égal à  $p$ .

## 6.1 Méthodes de quadrature élémentaires

On suppose que les points  $x_j$  de la formule de quadrature (6.1) sont deux à deux distincts.

### 6.1.1 Méthodes simplistes

On peut approcher  $f$  par un polynôme constant. Les trois formules usuelles sont

**Méthode du rectangle à gauche :** En figure 6.2, on représente l'approximation de  $\int_a^b f(x) dx$  lorsque  $f$  est approché par le polynôme constant  $P(x) = f(a)$ .

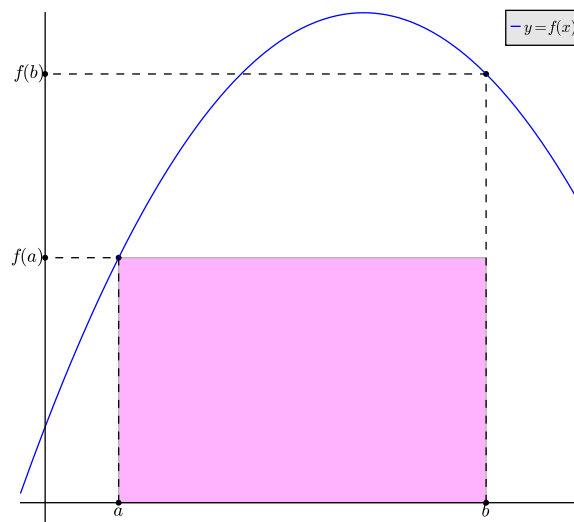


Figure 6.2: Formule du rectangle à gauche :  $\int_a^b f(x) dx \approx (b - a)f(a)$  (aire de la surface colorée)

On a alors

$$\int_a^b f(x) dx \approx \mathcal{Q}_0(f, a, b) = (b - a)f(a), \text{ formule du rectangle (à gauche)}$$

et son degré d'exactitude est 0.

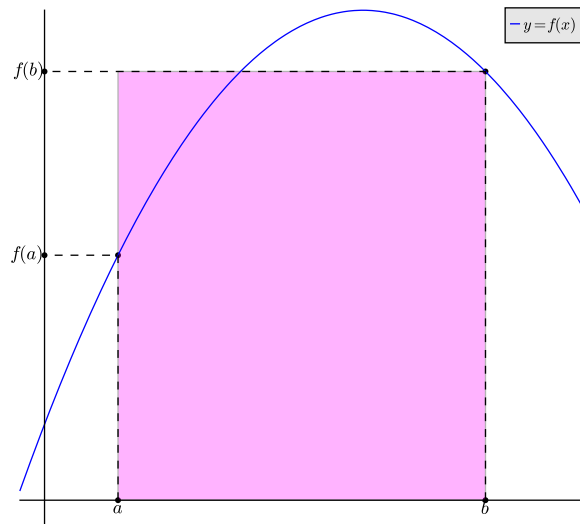


Figure 6.3: Formule du rectangle à droite :  $\int_a^b f(x)dx \approx (b-a)f(b)$  (aire de la surface colorée)

**Méthode du rectangle à droite :** En figure 6.3, on représente l'approximation de  $\int_a^b f(x)dx$  lorsque  $f$  est approché par le polynôme constant  $P(x) = f(b)$ . 1  
2

On a alors

$$\int_a^b f(x)dx \approx Q_0(f, a, b) = (b-a)f(b), \text{ formule du rectangle (à droite)}$$

et son degré d'exactitude est 0. 3

**Méthode du point milieu :** En figure 6.4, on représente l'approximation de  $\int_a^b f(x)dx$  lorsque  $f$  est approché par le polynôme constant  $P(x) = f((a+b)/2)$ . On a alors

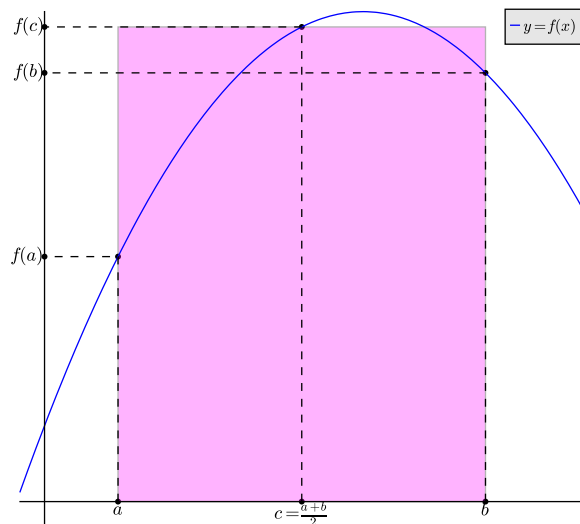


Figure 6.4: Formule du point milieu :  $\int_a^b f(x)dx \approx (b-a)f\left(\frac{a+b}{2}\right)$  (aire de la surface colorée)

$$\int_a^b f(x)dx \approx Q_0(f, a, b) = (b-a)f\left(\frac{a+b}{2}\right), \text{ formule du point milieu}$$

et son degré d'exactitude est 1. 4

- 1 La précision de ces formules est toute relative!  
 2 Nous allons maintenant voir comment *généraliser* en approchant la fonction  $f$  par des polynômes  
 3 d'interpolation de Lagrange de degré plus élevés.

### 6.1.2 Formules de quadrature élémentaires



#### Proposition 6.3

La formule de quadrature élémentaire (6.1) à  $n + 1$  points est d'ordre  $k$  si et seulement si

$$(b-a) \sum_{i=0}^n w_i x_i^r = \frac{b^{r+1} - a^{r+1}}{r+1}, \quad \forall r \in \llbracket 0, k \rrbracket. \quad (6.3)$$

- 6 *Proof.* •  $\Rightarrow$  Si la formule (6.1) est d'ordre  $k$ , elle est donc exacte pour tout polynôme de  $\mathbb{R}_k[X]$  et plus  
 7 particulièrement pour tous les monômes  $1, X, X^2, \dots, X^k$ . Soit  $r \in \llbracket 0, k \rrbracket$ . En prenant  $f(x) = x^r$ , la  
 8 formule (6.1) étant exacte par hypothèse, on obtient

$$\mathcal{Q}_n(x \mapsto x^r, a, b) = (b-a) \sum_{i=0}^n w_i x_i^r = \int_a^b x^r dx = \frac{b^{r+1} - a^{r+1}}{r+1}$$

- 9 •  $\Leftarrow$  On suppose que l'on a (6.3). Soit  $Q \in \mathbb{R}_k[X]$ . On va montrer que la formule de quadrature  
 10 (6.1) est alors exacte.  
 11 Le polynôme  $Q$  peut s'écrire comme combinaison linéaire des monômes de  $\{1, X, X^2, \dots, X^k\}$ , base  
 12 de  $\mathbb{R}_k[X]$  :

$$Q(x) = \sum_{j=0}^k \alpha_j x^j.$$

- 13 En prenant  $f = Q$ , la formule de quadrature (6.1) donne

$$\int_a^b Q(x) dx \approx (b-a) \sum_{i=0}^n w_i Q(x_i) = (b-a) \sum_{i=0}^n w_i \sum_{j=0}^k \alpha_j x_i^j$$

- 14 De plus, par linéarité de l'intégrale, on a

$$\int_a^b Q(x) dx = \sum_{j=0}^k \alpha_j \int_a^b x^j dx = \sum_{j=0}^k \alpha_j \frac{b^{j+1} - a^{j+1}}{j+1}$$

- 15 et en utilisant (6.3) on obtient

$$\int_a^b Q(x) dx = (b-a) \sum_{j=0}^k \alpha_j \sum_{i=0}^n w_i x_i^j$$

- 16 Ce qui donne

$$\int_a^b Q(x) dx = (b-a) \sum_{i=0}^n w_i Q(x_i).$$

- 17 La formule de quadrature est donc d'ordre  $k$ .  
 18 □



#### Proposition 6.4

Soient  $(x_i)_{i \in \llbracket 0, n \rrbracket}$  des points deux à deux distincts de l'intervalle  $[a, b]$  donnés. Il existe alors une  
 19 unique formule de quadrature élémentaire (6.1) à  $n + 1$  points d'ordre  $n$  au moins.

*Proof.* En fixant les points  $(x_i)_{i \in \llbracket 0, n \rrbracket}$  deux à deux distincts, pour obtenir explicitement la formule de quadrature de type (6.1) il faut déterminer les  $n + 1$  poids  $(w_i)_{i \in \llbracket 0, n \rrbracket}$ . Or, de (6.3), en prenant  $k = n$ , on obtient exactement  $n + 1$  équations linéaires en les  $(w_i)$  s'écrivant matriciellement sous la forme :

$$(b-a) \begin{pmatrix} 1 & 1 & \cdots & 1 \\ x_0 & x_1 & \cdots & x_n \\ \vdots & \vdots & & \vdots \\ x_0^n & x_1^n & \cdots & x_n^n \end{pmatrix} \begin{pmatrix} w_0 \\ w_1 \\ \vdots \\ w_n \end{pmatrix} = \begin{pmatrix} b-a \\ \frac{b^2-a^2}{2} \\ \vdots \\ \frac{b^{n+1}-a^{n+1}}{n+1} \end{pmatrix}$$

La matrice intervenant dans le système précédent s'appelle **la matrice de Vandermonde** et elle est inversible (car les  $(x_i)$  sont deux à deux distincts). Ceci établit donc l'existence de poids  $(w_i)_{i \in \llbracket 0, n \rrbracket}$  tels que la formule de quadrature élémentaire (6.1) soit d'ordre (au moins)  $n$ . Pour obtenir l'unicité, supposons qu'il existe  $(w_i)_{i \in \llbracket 0, n \rrbracket}$  et  $(\tilde{w}_i)_{i \in \llbracket 0, n \rrbracket}$  tels que pour tout  $P \in \mathbb{R}_n[X]$ , on ait

$$\int_a^b P(x) dx = (b-a) \sum_{i=0}^n w_i P(x_i) = (b-a) \sum_{i=0}^n \tilde{w}_i P(x_i).$$

On a alors  $\forall P \in \mathbb{R}_n[X]$ ,

$$\sum_{i=0}^n (w_i - \tilde{w}_i) P(x_i) = 0. \quad (6.4)$$

On rappelle que les fonctions de base de Lagrange associées aux  $(n+1)$  points  $(x_i)_{i \in \llbracket 0, n \rrbracket}$  définies en (5.6), notées  $L_i$ , sont dans  $\mathbb{R}_n[X]$  et vérifient

$$L_i(x_j) = \delta_{i,j}, \quad \forall j \in \llbracket 0, n \rrbracket$$

Soit  $j \in \llbracket 0, n \rrbracket$ . En choisissant  $P = L_j$  dans (6.4), on obtient

$$0 = \sum_{i=0}^n (w_i - \tilde{w}_i) L_j(x_i) = (w_j - \tilde{w}_j)$$

ce qui prouve l'unicité.  $\square$



### Proposition 6.5

Soit  $\mathcal{Q}_n(f, a, b)$  définie en (6.1), une formule de quadrature élémentaire à  $n + 1$  points (distincts deux à deux). On dit qu'elle est **symétrique** si

$$\forall i \in \llbracket 0, n \rrbracket, \quad \frac{x_i + x_{n-i}}{2} = \frac{a+b}{2} \quad \text{et} \quad w_i = w_{n-i}. \quad (6.5)$$

Dans ce cas si cette formule est exacte pour les polynômes de degré  $2m$  alors elle est nécessairement exacte pour les polynômes de degré  $2m + 1$ .

*Proof.* Soit  $P \in \mathbb{R}_{2m+1}[X]$ . Il peut alors s'écrire sous la forme

$$P(x) = C \left( x - \frac{a+b}{2} \right)^{2m+1} + R(x)$$

avec  $C$  une constante réelle et  $R \in \mathbb{R}_{2m}[X]$ . On a alors

$$\int_a^b P(x) dx = C \int_a^b \left( x - \frac{a+b}{2} \right)^{2m+1} dx + \int_a^b R(x) dx$$

et en appliquant la formule de quadrature au polynôme  $P$  on obtient

$$\sum_{i=0}^n w_i P(x_i) = C \sum_{i=0}^n w_i \left( x_i - \frac{a+b}{2} \right)^{2m+1} + \sum_{i=0}^n w_i R(x_i)$$

1 On veut donc démontrer que

$$\int_a^b P(x)dx = (b-a) \sum_{i=0}^n w_i P(x_i)$$

2 c'est à dire

$$C \int_a^b \left(x - \frac{a+b}{2}\right)^{2m+1} dx + \int_a^b R(x)dx = (b-a)C \sum_{i=0}^n w_i \left(x_i - \frac{a+b}{2}\right)^{2m+1} + (b-a) \sum_{i=0}^n w_i R(x_i)$$

3 Comme la formule de quadrature est supposée exacte pour les polynôme de degré  $2m$ , on a

$$\int_a^b R(x)dx = (b-a) \sum_{i=0}^n w_i R(x_i).$$

4 Il reste donc à démontrer que

$$\int_a^b \left(x - \frac{a+b}{2}\right)^{2m+1} dx = (b-a) \sum_{i=0}^n w_i \left(x_i - \frac{a+b}{2}\right)^{2m+1}.$$

5 Or en effectuant le changement de variable  $t \mapsto \frac{a+b}{2} + t \frac{b-a}{2}$  on obtient

$$\int_a^b \left(x - \frac{a+b}{2}\right)^{2m+1} dx = \frac{b-a}{2} \int_{-1}^1 t^{2m+1} dt = 0.$$

6 Des propriétés de symétrie de la formule, on déduit

$$x_i + x_{n-i} = a+b \Leftrightarrow x_i - \frac{a+b}{2} = -\left(x_{n-i} - \frac{a+b}{2}\right)$$

7 • Si  $n = 2k$ , ( $n$  paire), on a alors un nombre **impair** de points avec nécessairement  $x_k = x_{n-k} = \frac{a+b}{2}$   
8 et

$$\begin{aligned} \sum_{i=0}^n w_i \left(x_i - \frac{a+b}{2}\right)^{2m+1} &= \sum_{i=0}^{k-1} w_i \left(x_i - \frac{a+b}{2}\right)^{2m+1} + 0 \times w_k + \sum_{i=k+1}^{2k} w_i \left(x_i - \frac{a+b}{2}\right)^{2m+1} \\ &= \sum_{i=0}^{k-1} w_i \left(x_i - \frac{a+b}{2}\right)^{2m+1} - \sum_{i=k+1}^{2k} w_{n-i} \left(x_{n-i} - \frac{a+b}{2}\right)^{2m+1} \\ &= \sum_{i=0}^{k-1} w_i \left(x_i - \frac{a+b}{2}\right)^{2m+1} - \sum_{j=0}^{k-1} w_j \left(x_j - \frac{a+b}{2}\right)^{2m+1} \\ &= 0. \end{aligned}$$

9 • Si  $n = 2k-1$ , ( $n$  impaire), on a alors un nombre **pair** de points (avec  $x_i \neq \frac{a+b}{2}, \forall i \in \llbracket 0, n \rrbracket$ ) et

$$\begin{aligned} \sum_{i=0}^n w_i \left(x_i - \frac{a+b}{2}\right)^{2m+1} &= \sum_{i=0}^{k-1} w_i \left(x_i - \frac{a+b}{2}\right)^{2m+1} + \sum_{i=k}^{2k-1} w_i \left(x_i - \frac{a+b}{2}\right)^{2m+1} \\ &= \sum_{i=0}^{k-1} w_i \left(x_i - \frac{a+b}{2}\right)^{2m+1} - \sum_{i=k}^{2k-1} w_{n-i} \left(x_{n-i} - \frac{a+b}{2}\right)^{2m+1} \\ &= \sum_{i=0}^{k-1} w_i \left(x_i - \frac{a+b}{2}\right)^{2m+1} - \sum_{j=0}^{k-1} w_j \left(x_j - \frac{a+b}{2}\right)^{2m+1} \\ &= 0. \end{aligned}$$

10

□

11 **6.1.3 Liens avec le polynôme d'interpolation de Lagrange**

**Proposition 6.6**

Soient  $f \in \mathcal{C}^{n+1}([a, b]; \mathbb{R})$  et  $\mathcal{Q}_n(f, a, b)$  définie en (6.1), une formule de quadrature élémentaire à  $n + 1$  points  $(x_i)_{i \in \llbracket 0, n \rrbracket}$  (distincts deux à deux). Si, pour tout  $i \in \llbracket 0, n \rrbracket$ , les poids  $w_i$  sont donnés par

$$w_i = \frac{1}{b-a} \int_a^b \prod_{\substack{j=0 \\ j \neq i}}^n \frac{x - x_j}{x_i - x_j} dx = \int_0^1 \prod_{\substack{j=0 \\ j \neq i}}^n \frac{t - t_j}{t_i - t_j} dt, \quad \forall i \in \llbracket 0, n \rrbracket \quad (6.6)$$

avec  $t_i = (x_i - a)/(b - a)$  alors la formule de quadrature est d'ordre  $n$  au moins et l'on a

$$|\mathcal{E}_{a,b}(f)| \leq \frac{1}{(n+1)!} \|f^{(n+1)}\|_\infty \int_a^b \left| \prod_{i=0}^n (x - x_i) \right| dx \quad (6.7)$$

*Proof.* On note  $\mathcal{L}_n(f)$  le polynôme d'interpolation de Lagrange associés aux points  $(x_i, f(x_i))_{i \in \llbracket 0, n \rrbracket}$  :

$$\mathcal{L}_n(f)(x) = \sum_{i=0}^n L_i(x) f(x_i) \quad \text{avec} \quad L_i(x) = \prod_{\substack{j=0 \\ j \neq i}}^n \frac{x - x_j}{x_i - x_j}$$

On a alors

$$\int_a^b \mathcal{L}_n(f)(x) dx = \sum_{i=0}^n f(x_i) \int_a^b L_i(x) dx.$$

En prenant  $\forall i \in \llbracket 0, n \rrbracket$

$$w_i = \frac{1}{b-a} \int_a^b L_i(x) dx$$

on obtient la formule de quadrature

$$\int_a^b \mathcal{L}_n(f)(x) dx = \mathcal{Q}_n(f, a, b) = (b-a) \sum_{i=0}^n w_i f(x_i).$$

Soit  $P \in \mathbb{R}_n[X]$ . Par unicité du polynôme d'interpolation de Lagrange, on a  $P = \mathcal{L}_n(P)$  et donc

$$\int_a^b P(x) dx = \int_a^b \mathcal{L}_n(P)(x) dx = (b-a) \sum_{i=0}^n w_i P(x_i).$$

La formule de quadrature est donc d'ordre  $n$  au moins. De plus par le changement de variables  $s : t \rightarrow a + (b-a)t$  on obtient

$$\int_a^b L_i(x) dx = (b-a) \int_0^1 L_i \circ s(t) dt$$

et l'on a  $x_i = s(t_i) = a + (b-a)t_i$  où  $t_i = (x_i - a)/(b-a)$ . On en déduit

$$\begin{aligned} \int_0^1 L_i \circ s(t) dt &= \int_0^1 \prod_{\substack{j=0 \\ j \neq i}}^n \frac{s(t) - s(t_j)}{s(t_i) - s(t_j)} dt = \int_0^1 \prod_{\substack{j=0 \\ j \neq i}}^n \frac{(b-a)(t - t_j)}{(b-a)(t_i - t_j)} dt \\ &= \int_0^1 \prod_{\substack{j=0 \\ j \neq i}}^n \frac{t - t_j}{t_i - t_j} dt \end{aligned}$$

Ce qui achève la preuve de (6.6).

Comme  $f \in \mathcal{C}^{n+1}([a, b]; \mathbb{R})$ , pour démontrer l'inégalité (6.7) on peut appliquer le théorème 5.3 pour obtenir

$$\forall x \in [a, b], \exists \xi_x \in [a, b], \quad f(x) - \mathcal{L}_n(f)(x) = \frac{f^{(n+1)}(\xi_x)}{(n+1)!} \prod_{i=0}^n (x - x_i)$$

En intégrant cette équation sur l'intervalle  $[a, b]$ , on aboutit à

$$\int_a^b f(x)dx - \int_a^b \mathcal{L}_n(f)(x)dx = \frac{1}{(n+1)!} \int_a^b f^{(n+1)}(\xi_x) \prod_{i=0}^n (x - x_i)dx.$$

et donc

$$\left| \int_a^b f(x)dx - \int_a^b \mathcal{L}_n(f)(x)dx \right| \leq \frac{\|f^{(n+1)}\|_\infty}{(n+1)!} \int_a^b \left| \prod_{i=0}^n (x - x_i) \right| dx.$$

□

Dans la proposition précédente, le choix des points reste libre. Pour expliciter les poids  $w_i$ , donnés par (6.6), on peut choisir comme  $(n+1)$  points d'interpolations les points de la discrétisation régulière de l'intervalle  $[a, b]$  en  $n$  intervalles. Ceci est l'objet des Méthodes de Newton-Cotes.

Bien d'autres méthodes peuvent être obtenues (avec d'autres points), certaines permettant le calcul d'intégrales avec poids de la forme  $\int_a^b w(x)f(x)dx$  :

- méthode de Newton-Cotes ouvertes,
- méthode de Gauss-Legendre,
- méthode de Gauss-Jacobi,
- méthode de Gauss-Tchebychev,
- méthode de Gauss-Laguerre,
- méthode de Gauss-Hermitte,
- méthode de Gauss-Lobatto,
- méthode de Romberg...

#### 6.1.4 Formules élémentaires de Newton-Cotes

Soit  $(x_i)_{i \in \llbracket 0, n \rrbracket}$  une discrétisation régulière de l'intervalle  $[a, b]$ :  $x_i = a + ih$  avec  $h = (b - a)/n$ .



#### Proposition 6.7

Soient  $f \in \mathcal{C}^0([a, b]; \mathbb{R})$  et  $(x_i)_{i \in \llbracket 0, n \rrbracket}$  une discrétisation régulière de l'intervalle  $[a, b]$ :  $x_i = a + ih$  avec  $h = (b - a)/n$ .

Les formules de quadrature élémentaires de Newton-Cotes s'écrivent sous la forme

$$\int_a^b f(x) dx \approx (b - a) \sum_{i=0}^n w_i f(x_i)$$

| $n$ | ordre | $w_i$ (poids)       |                      |                      |                      |                      |                      |                      |                      |                     | nom        |
|-----|-------|---------------------|----------------------|----------------------|----------------------|----------------------|----------------------|----------------------|----------------------|---------------------|------------|
| 1   | 1     | $\frac{1}{2}$       | $\frac{1}{2}$        |                      |                      |                      |                      |                      |                      |                     | trapèze    |
| 2   | 3     | $\frac{1}{6}$       | $\frac{2}{3}$        | $\frac{1}{6}$        |                      |                      |                      |                      |                      |                     | Simpson    |
| 3   | 3     | $\frac{1}{8}$       | $\frac{3}{8}$        | $\frac{3}{8}$        | $\frac{1}{8}$        |                      |                      |                      |                      |                     | Newton     |
| 4   | 5     | $\frac{7}{90}$      | $\frac{16}{45}$      | $\frac{2}{15}$       | $\frac{16}{45}$      | $\frac{7}{90}$       |                      |                      |                      |                     | Villarcieu |
| 5   | 5     | $\frac{19}{288}$    | $\frac{25}{96}$      | $\frac{25}{144}$     | $\frac{25}{144}$     | $\frac{25}{96}$      | $\frac{19}{288}$     |                      |                      |                     | ?          |
| 6   | 7     | $\frac{41}{840}$    | $\frac{9}{35}$       | $\frac{9}{280}$      | $\frac{34}{105}$     | $\frac{9}{280}$      | $\frac{9}{35}$       | $\frac{41}{840}$     |                      |                     | Weddle     |
| 7   | 7     | $\frac{751}{17280}$ | $\frac{3577}{17280}$ | $\frac{49}{640}$     | $\frac{2989}{17280}$ | $\frac{2989}{17280}$ | $\frac{49}{640}$     | $\frac{3577}{17280}$ | $\frac{751}{17280}$  |                     | ?          |
| 8   | 9     | $\frac{989}{28350}$ | $\frac{2944}{14175}$ | $-\frac{464}{14175}$ | $\frac{5248}{14175}$ | $-\frac{454}{2835}$  | $\frac{5248}{14175}$ | $-\frac{464}{14175}$ | $\frac{2944}{14175}$ | $\frac{989}{28350}$ | ?          |

Table 6.1: Méthodes de Newton-Cotes

Par exemple, la formule de Simpson ( $n = 2$ ) est

$$\int_a^b f(x) dx \approx \frac{b-a}{6} \left( f(a) + 4f\left(\frac{a+b}{2}\right) + f(b) \right) \quad (6.8)$$

Pour un  $n$  donné, déterminer les coefficients  $w_i$  de la formule de Newton-Cotes est assez simple.

**Démonstration 1:** En effet, il suffit de remarquer que la formule doit être exacte si  $f$  est un polynôme de degré au plus  $n$ . Ensuite par linéarité de la formule, on est ramené à résoudre un système linéaire à  $n + 1$  inconnues (les  $(w_i)_{i \in \llbracket 0, n \rrbracket}$ ) en écrivant que pour chaque monôme

$$(b-a) \sum_{i=0}^n w_i f(x_i) = \int_a^b f(x) dx, \quad \forall f \in \{1, X, X^2, \dots, X^n\} \subset \mathbb{R}_n[X]. \quad (6.9)$$

Par exemple, pour  $n = 2$ , on a  $b = a + 2h$ . A partir de (6.9) on obtient les trois équations :

$$\begin{cases} w_0 + w_1 + w_2 &= 1, & (f(x) = 1) \\ aw_0 + (a+h)w_1 + (a+2h)w_2 &= \frac{1}{b-a} \int_a^b x dx = a+h, & (f(x) = x) \\ a^2w_0 + (a+h)^2w_1 + (a+2h)^2w_2 &= \frac{1}{b-a} \int_a^b x^2 dx = \frac{1}{3}(3a^2 + 6ah + 4h^2), & (f(x) = x^2). \end{cases}$$

On a alors

$$\begin{cases} w_0 = -w_1 - w_2 + 1, \\ a(-w_1 - w_2 + 1) + (a+h)w_1 + (a+2h)w_2 = a+h, \\ a^2(-w_1 - w_2 + 1) + (a+h)^2w_1 + (a+2h)^2w_2 = a^2 + 2ah + \frac{4}{3}h^2. \end{cases}$$

c'est à dire

$$\begin{cases} w_0 = -w_1 - w_2 + 1, \\ w_1 + 2w_2 = 1, \\ 2a(w_1 + 2w_2) + h(w_1 + 4w_2) = 2a + 4h/3, \end{cases}$$

Par substitution, de la deuxième équation dans la troisième, on obtient enfin

$$\begin{cases} w_0 = -w_1 - w_2 + 1, \\ w_1 + 2w_2 = 1, \\ w_1 + 4w_2 = 4/3, \end{cases}$$

ce qui donne  $w_0 = w_2 = \frac{1}{6}$  et  $w_1 = \frac{2}{3}$ .

1 **Démonstration 2:** En utilisant la Proposition 6.6, on a

$$w_i = \int_a^b L_i^{\mathbf{X}}(x) dx, \quad \forall i \in \llbracket 0, n \rrbracket$$

2 Or on a vu dans l'exercice ?? que

$$\int_a^b L_i^{\mathbf{X}}(x) dx = (b-a) \int_0^1 L_i^{\mathbf{X}} \circ s(t) dt = (b-a) \int_0^1 L_i^{\mathbf{T}}(t) dt$$

3 avec  $s : t \longrightarrow a + (b-a)t$  et pour tout  $i \in \llbracket 0, n \rrbracket$ ,  $t_i = s^{-1}(x_i) = \frac{x_i - a}{b-a} = \frac{i}{n}$ . On a donc dans le cadre  
4 des points équidistants

$$L_i^{\mathbf{T}}(t) = \prod_{\substack{j=0 \\ j \neq i}}^n \frac{t - t_j}{t_i - t_j} = \prod_{\substack{j=0 \\ j \neq i}}^n \frac{t - j/n}{(i - j)/n} = \prod_{\substack{j=0 \\ j \neq i}}^n \frac{nt - j}{i - j}$$

5 Par exemple, pour  $n = 2$ , on peut calculer  $(w_i)_{i \in \llbracket 0, 2 \rrbracket}$

6 On a

$$\begin{aligned} L_0(t) &= \frac{2t-1}{-1} \frac{2t-2}{-2} = (2t-1)(t-1) \\ L_1(t) &= \frac{2t}{1} \frac{2t-2}{-1} = -4(t-1)t \\ L_2(t) &= \frac{2t}{2} \frac{2t-1}{1} = (2t-1)t \end{aligned}$$

7 et

$$w_0 = \int_0^1 L_0(t) dt = \frac{1}{6}, \quad w_1 = \int_0^1 L_1(t) dt = \frac{2}{3}, \quad w_2 = \int_0^1 L_2(t) dt = \frac{1}{6}$$

8 **Remarque 6.8** Pour les méthode de Newton-Cotes, il ne faut pas trop "monter" en ordre car le phénomène  
9 de Runge (forte oscillation possible du polynôme d'interpolation sur les bords de l'intervalle) peut conduire  
10 à de très grande erreurs. Au delà de  $n = 7$ , des poids négatifs apparaissent dans les formules et les rendent  
11 beaucoup plus sensibles aux erreurs d'arrondis.

12 Sur un exemple très simple, il est possible d'illustrer ce phénomène. Soit  $f(x) = 3x + 2$ . On a démontré  
13 que toutes les formules de Newton-Cotes à  $n+1$  points,  $n \geq 1$ , sont exactes pour le calcul de  $\int_a^b f(x) dx$  car  
14  $f$  est un polynôme de degré 1. De plus les poids  $(w_i)_{i \in \llbracket 0, n \rrbracket}$  peuvent être calculés sous forme fractionnaire  
15 : il est immédiat à partir de la formule (6.6) que  $w_i \in \mathbb{Q}$ .

16 En choisissant, par exemple,  $a = 0$  et  $b = 1$  les  $n+1$  points  $x_i$  de la formule de Newton-Cotes sont aussi  
17 des nombres rationnels. La fonction  $f$  étant polynomiale, on obtient  $f(x_i) \in \mathbb{Q}$ . On en déduit que la  
18 formule de Newton-Cotes à  $n+1$  points donne dans ce cas un résultat (exacte) sous forme fractionnaire.  
19 Numériquement, les poids  $w_i \in \mathbb{Q}$  et les points  $x_i \in \mathbb{Q}$  vont être approchés en commettant de très petites  
20 erreurs. Sous Sage, logiciel gratuit de calcul formel (entre autres), on peut programmer à la fois le calcul  
21 des méthodes de Newton-Cotes exactes (en restant dans le corps  $\mathbb{Q}$ ) et le calcul "approché" correspondant  
22 aux méthodes de Newton-Cotes *approchées* (i.e. les  $w_i$  et les  $x_i$  sont approchés avant le calcul de la  
23 somme). On représente en Figure 6.5 les erreurs obtenues en fonction de  $n$  par ces deux approches. Ces  
24 courbes en échelle logarithmique ont été obtenues en ajoutant aux erreurs le nombre  $1e-16$  (pour éviter  
25 de prendre le log de zéro).

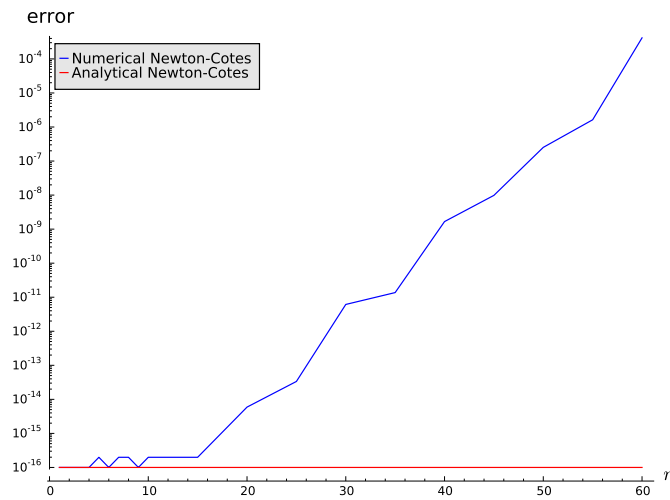


Figure 6.5: Instabilité des méthodes de Newton-Cotes élémentaires

Pour pallier ce problème, on étudie maintenant les méthodes de quadrature composées.

### 6.1.5 Méthodes de quadrature composées

Soit  $f$  une fonction définie et intégrable sur un intervalle  $[\alpha, \beta]$  donné. On propose de chercher une approximation de

$$I = \int_{\alpha}^{\beta} f(x) dx.$$

Ces méthodes consistent en l'utilisation de la relation de Chasles pour décomposer l'intégrale en une somme d'intégrales sur des domaines plus petits puis à approcher ces dernières par une formule de quadrature élémentaire à  $n + 1$  points.

#### ♥ Definition 6.9

Soit  $(\alpha_i)_{i \in \llbracket 0, k \rrbracket}$  une subdivision de l'intervalle  $[\alpha, \beta]$ :

$$\alpha = \alpha_0 < \alpha_1 < \dots < \alpha_k = \beta.$$

On a alors

$$\int_{\alpha}^{\beta} f(x) dx = \sum_{i=1}^k \int_{\alpha_{i-1}}^{\alpha_i} f(x) dx. \quad (6.10)$$

Soit  $\mathcal{Q}_n(g, a, b)$  la formule de quadrature élémentaire à  $n + 1$  points d'ordre  $p$  donnée par

$$\mathcal{Q}_n(g, a, b) \stackrel{\text{def}}{=} (b - a) \sum_{j=0}^n w_j g(x_j) \approx \int_a^b g(x) dx.$$

La **méthode de quadrature composée associée à  $\mathcal{Q}_n$** , notée  $\mathcal{Q}_{k,n}^{\text{comp}}$ , est donnée par

$$\mathcal{Q}_{k,n}^{\text{comp}}(f, \alpha, \beta) = \sum_{i=1}^k \mathcal{Q}_n(f, \alpha_{i-1}, \alpha_i) \approx \int_{\alpha}^{\beta} f(x) dx \quad (6.11)$$

La proposition suivante est immédiate



#### Proposition 6.10

Soit  $\mathcal{Q}_n$  une formule de quadrature élémentaire à  $n + 1$  points. Si  $\mathcal{Q}_n$  est d'ordre  $p$  alors la méthode de quadrature composée associée est aussi d'ordre  $p$  : elle est exacte pour tout polynôme de degré  $p$ .

Pour les trois formules composées qui suivent on choisit une discrétisation régulière. Soit  $(\alpha_i)_{i \in \llbracket 0, k \rrbracket}$  une discrétisation régulière de l'intervalle  $[\alpha; \beta]$  :  $\alpha_i = \alpha + ih$  avec  $h = (\beta - \alpha)/k$  le pas de la discrétisation.

#### Formule composite des points milieux

La formule élémentaire des points milieux est donnée par

$$\mathcal{Q}_0(g, a, b) \stackrel{\text{def}}{=} (b - a)g\left(\frac{a + b}{2}\right)$$

On note  $m_j = \frac{\alpha_{j-1} + \alpha_j}{2}$  le point milieu de l'intervalle  $[\alpha_{j-1}, \alpha_j]$ .

$$\int_{\alpha}^{\beta} f(x) dx = \sum_{j=1}^k \int_{\alpha_{j-1}}^{\alpha_j} f(x) dx \approx \sum_{j=1}^k \mathcal{Q}_0(f, \alpha_{j-1}, \alpha_j) = h \sum_{j=1}^k f(m_j) \quad (6.12)$$

On illustre graphiquement l'approximation d'une intégrale par cette formule en Figure 6.6.

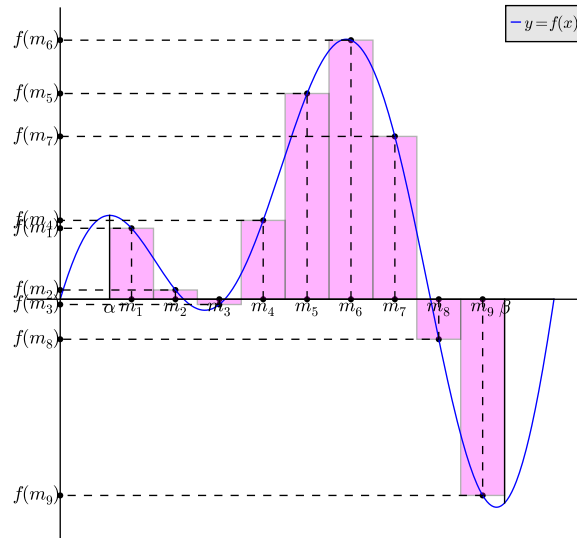


Figure 6.6: Formule composite des points milieux :  $\int_{\alpha}^{\beta} f(x) dx \approx h \sum_{j=1}^k f(m_j)$  (aire de la surface colorée)

#### Formule composite des trapèzes

La formule élémentaire des points trapèzes est donnée par

$$\mathcal{Q}_1(g, a, b) \stackrel{\text{def}}{=} \frac{b - a}{2}(g(a) + g(b))$$

On obtient alors

$$\begin{aligned} \int_a^b f(x) dx &= \sum_{j=1}^k \int_{\alpha_{j-1}}^{\alpha_j} f(x) dx \approx \sum_{j=1}^k \mathcal{Q}_1(f, \alpha_{j-1}, \alpha_j) = \frac{h}{2} \sum_{j=1}^k (f(\alpha_{j-1}) + f(\alpha_j)) \\ &\approx \frac{h}{2} \left( f(\alpha_0) + 2 \sum_{j=1}^{k-1} f(\alpha_j) + f(\alpha_k) \right) \end{aligned} \quad (6.13)$$

C'est une formule d'ordre 2 par rapport à  $h$ .

On illustre graphiquement l'approximation d'une intégrale par cette formule en Figure 6.7.

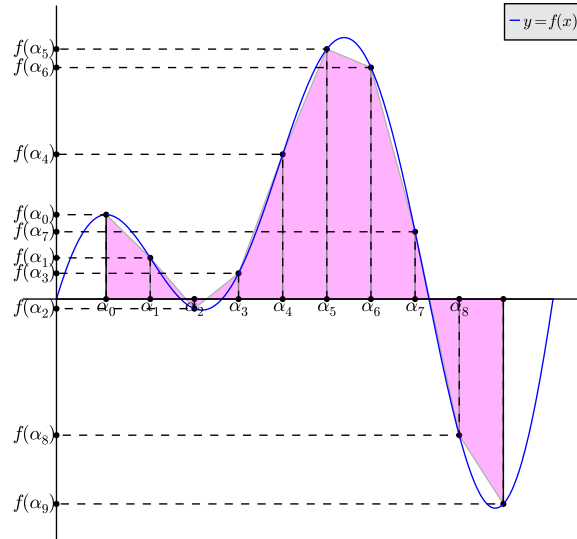


Figure 6.7: Formule composite des trapèzes :  $\int_a^b f(x)dx \approx \frac{h}{2} \left( f(\alpha_0) + 2 \sum_{j=1}^{k-1} f(\alpha_j) + f(\alpha_k) \right)$  (aire de la surface colorée)

### Formule composite de Simpson

La formule élémentaire de Simpson est donnée par

$$\mathcal{Q}_2(g, a, b) \stackrel{\text{def}}{=} \frac{b-a}{6} (g(a) + 4g(\frac{a+b}{2}) + g(b))$$

En notant  $m_j = \frac{\alpha_{j-1} + \alpha_j}{2}$  le point milieu de l'intervalle  $[\alpha_{j-1}, \alpha_j]$ , on obtient

$$\begin{aligned} \int_a^b f(x)dx &= \sum_{j=1}^k \int_{\alpha_{j-1}}^{\alpha_j} f(x)dx \approx \sum_{j=1}^k \mathcal{Q}_2(f, \alpha_{j-1}, \alpha_j) = \frac{h}{6} \sum_{j=1}^k (f(\alpha_{j-1}) + 4f(m_j) + f(\alpha_j)) \\ &\approx \frac{h}{6} \left( 4 \sum_{j=1}^k f(m_j) + f(\alpha_0) + 2 \sum_{j=1}^{k-1} f(\alpha_j) + f(\alpha_k) \right) \end{aligned} \quad (6.14)$$

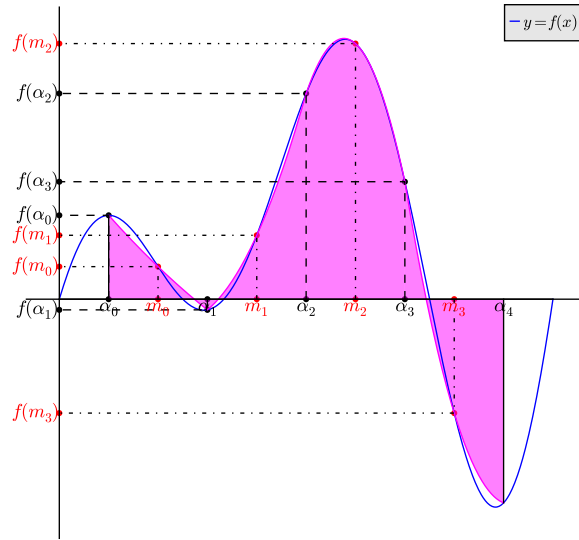


Figure 6.8: Formule composite de Simpson :  $\int_{\alpha}^{\beta} f(x)dx \approx \frac{h}{6} \left( 4 \sum_{j=1}^k f(m_j) + f(\alpha_0) + 2 \sum_{j=1}^{k-1} f(\alpha_j) + f(\alpha_k) \right)$   
(aire de la surface colorée)

## 6.2 Erreurs des méthodes de quadrature composées

- 1 Soit  $\mathcal{Q}_{k,n}^{\text{comp}}$  une méthode de quadrature composée associée à une méthode de quadrature élémentaire  $\mathcal{Q}_n$ .  
 2 On note  $\mathcal{E}_{\alpha,\beta}^{\text{comp}}(f)$  l'erreur de cette méthode de quadrature composée :

$$\mathcal{E}_{\alpha,\beta}^{\text{comp}}(f) = \int_{\alpha}^{\beta} f(x)dx - \mathcal{Q}_{k,n}^{\text{comp}}(f, \alpha, \beta). \quad (6.15)$$

- 3 On a alors

$$\mathcal{E}_{\alpha,\beta}^{\text{comp}}(f) = \sum_{j=1}^k \left( \int_{\alpha_{j-1}}^{\alpha_j} f(x)dx - \mathcal{Q}_n(f, \alpha_{j-1}, \alpha_j) \right) = \sum_{j=1}^k \mathcal{E}_{\alpha_{j-1}, \alpha_j}(f)$$

- 4 Dans le cadre des méthodes composées de Newton-Cotes, on peut démontrer (voir [2]), la majoration  
 5 suivante

$$\max_{x \in [a,b]} \left| \prod_{i=0}^n (x - x_i) \right| \leq C \frac{e^{-n}}{\sqrt{n} \log(n)} (b-a)^{n+1}. \quad (6.16)$$

- 6 où  $(x_i)_{i \in \llbracket 0, n \rrbracket}$  est la discrétisation régulière de  $[a, b]$  et  $C > 0$ . On suppose que  $f \in \mathcal{C}^{n+1}([a, b]; \mathbb{R})$ . En  
 7 notant  $h_j = \alpha_j - \alpha_{j-1}$  et en utilisant les majorations (6.7) et (6.16) on obtient

$$\begin{aligned} |\mathcal{E}_{\alpha,\beta}^{\text{comp}}(f)| &\leq \sum_{j=1}^k |\mathcal{E}_{\alpha_{j-1}, \alpha_j}(f)| \\ &\leq \sum_{j=1}^k \frac{K_n}{(n+1)!} \max_{x \in [\alpha_{j-1}, \alpha_j]} |f^{(n+1)}(x)| h_j^{n+2} \quad \text{avec } K_n = C \frac{e^{-n}}{\sqrt{n} \log(n)} \\ &\leq K_n \frac{h^{n+1}}{(n+1)!} \max_{x \in [\alpha, \beta]} |f^{(n+1)}(x)| \sum_{j=1}^k h_j \quad \text{avec } h = \max_{j \in \llbracket 1, k \rrbracket} h_j \\ &\leq K_n (\beta - \alpha) \frac{h^{n+1}}{(n+1)!} \|f^{(n+1)}\|_{\infty} \end{aligned}$$

- 8 Cette majoration facile à obtenir n'est pas optimale et n'est valable que pour les formules composées de  
 9 Newton-Cotes. A l'aide des **noyaux de Peano**, on peut démontrer le théorème suivant :


**Théorème 6.11: [2], page 43 (admis)**

Soient  $f \in \mathcal{C}^{p+1}([a, b]; \mathbb{R})$  et  $\mathcal{Q}_{k,n}^{\text{comp}}$  une méthode de quadrature composée associée à une méthode de quadrature élémentaire  $\mathcal{Q}_n$  d'ordre  $p \geq n$ . On a alors

$$|\mathcal{E}_{\alpha,\beta}^{\text{comp}}(f)| \leq C_p(\beta - \alpha)h^{p+1} \|f^{(p+1)}\|_{\infty} \quad (6.17)$$

avec  $h = \max_{j \in \llbracket 1, k \rrbracket} (\alpha_j - \alpha_{j-1})$  et  $C_p > 0$ .

On illustre ce théorème pour les méthodes de Newton-Cotes composées  $\mathcal{Q}_{k,n}^{\text{comp}}$ , pour  $n \in \llbracket 1, 8 \rrbracket$ , par les Figures 6.9 et 6.10.

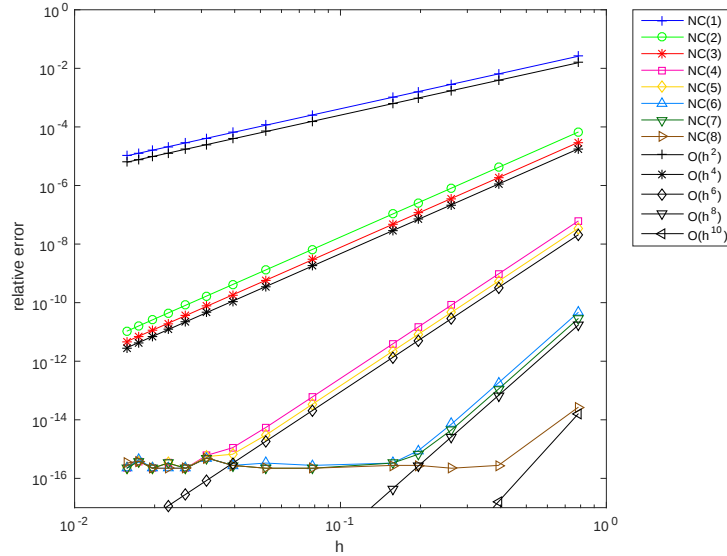


Figure 6.9: Erreur des méthodes de Newton-Cotes composées pour le calcul de  $\int_0^{\pi/2} \cos(x)dx$ , NC(n) correspondant à  $\mathcal{Q}_{k,n}^{\text{comp}}$  et  $h = \frac{\pi}{2k}$ .

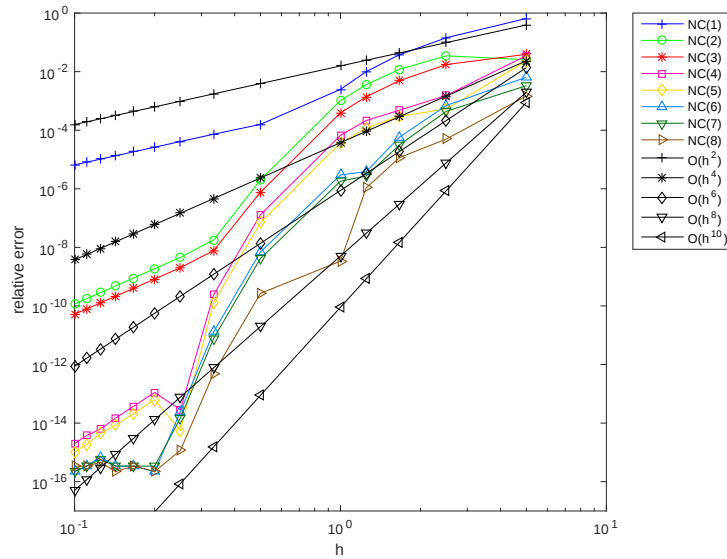


Figure 6.10: Erreur des méthodes de Newton-Cotes composées pour le calcul de  $\int_{-5}^5 \frac{1}{1+x^2}dx$ , NC(n) correspondant à  $\mathcal{Q}_{k,n}^{\text{comp}}$  et  $h = \frac{10}{k}$ .

## 6.3 Intégrales multiples

On veut approcher, en utilisant la formule de Simpson, l'intégrale

$$I = \int_a^b \int_c^d f(x, y) dy dx$$

Par utilisation de la formule de quadrature de Simpson (6.8) en  $y$  on a

$$g(x) = \int_c^d f(x, y) dy \approx \tilde{g}(x) = \frac{d-c}{6} \left( f(x, c) + 4f(x, \frac{c+d}{2}) + f(x, d) \right).$$

Une nouvelle utilisation de Simpson en  $x$  donne

$$\begin{aligned} I &= \int_a^b g(x) dx \approx \frac{b-a}{6} \left( g(a) + 4g(\frac{a+b}{2}) + g(b) \right) \\ &\approx \frac{b-a}{6} \left( \tilde{g}(a) + 4\tilde{g}(\frac{a+b}{2}) + \tilde{g}(b) \right) \end{aligned}$$

En posant  $\alpha = \frac{a+b}{2}$  et  $\beta = \frac{c+d}{2}$ , on obtient la formule de quadrature de Simpson "2D" :

$$I \approx \frac{b-a}{6} \frac{d-c}{6} \begin{pmatrix} f(a, c) + 4f(a, \beta) + f(a, d) \\ + 4(f(\alpha, c) + 4f(\alpha, \beta) + f(\alpha, d)) \\ + f(b, c) + 4f(b, \beta) + f(b, d) \end{pmatrix} \quad (6.18)$$

La méthodologie pour obtenir la formule composite de Simpson "2D" est la suivante

1. Discrétisation régulière de  $[a, b]$ :  $\forall k \in \llbracket 0, n \rrbracket$ ,  $x_k = a + kh_x$  avec  $h_x = \frac{b-a}{n}$ .
2. Discrétisation régulière de  $[c, d]$ :  $\forall l \in \llbracket 0, m \rrbracket$ ,  $y_l = c + lh_y$  avec  $h_y = \frac{d-c}{m}$ .
3. Relation de Chasles :

$$\int_a^b \int_c^d f(x, y) dy dx = \sum_{k=1}^n \sum_{l=1}^m \int_{x_{k-1}}^{x_k} \int_{y_{l-1}}^{y_l} f(x, y) dy dx.$$

4. Formule composite de Simpson "2D" :

$$\begin{aligned} &\int_a^b \int_c^d f(x, y) dy dx \\ &= \frac{h_x h_y}{36} \sum_{k=1}^n \sum_{l=1}^m \begin{pmatrix} f(x_{k-1}, y_{l-1}) + 4f(x_{k-1}, \beta_l) + f(x_{k-1}, y_l) \\ + 4(f(\alpha_k, y_{l-1}) + 4f(\alpha_k, \beta_l) + f(\alpha_k, y_l)) \\ + f(x_k, y_{l-1}) + 4f(x_k, \beta_l) + f(x_k, y_l) \end{pmatrix} \end{aligned}$$

avec  $\alpha_k = \frac{x_{k-1} + x_k}{2}$  et  $\beta_l = \frac{y_{l-1} + y_l}{2}$ .

# Chapitre 7

## Dérivation numérique



# Chapitre A

## Annexes

### A.1 Analyse : rappels



#### Théorème A.1: Théorème des valeurs intermédiaires

Soit  $f : [a, b] \subset \mathbb{R} \longrightarrow \mathbb{R}$  une application continue, alors

$$\forall y \in f([a, b]), \exists x \in [a, b] \text{ tel que } f(x) = y.$$



#### Corollaire A.2: Théorème de Bolzano

Soit  $f : [a, b] \subset \mathbb{R} \longrightarrow \mathbb{R}$  une application continue. Si  $f(a)$  et  $f(b)$  ne sont pas de même signe (i.e.  $f(a)f(b) < 0$ ) alors il existe au moins  $c \in ]a, b[$  tel que  $f(c) = 0$ .



#### Proposition A.3: Formule de Taylor-Lagrange

Soit  $n \in \mathbb{N}^*$  et  $f \in \mathcal{C}^n([a, b])$  dont la dérivée  $n$ -ième est dérivable. Alors pour tout  $x, y$  dans  $[a, b]$ ,  $x \neq y$ , il existe  $\xi \in ]x, y[$  tel que

$$f(x) = f(y) + \sum_{k=1}^n \frac{(x-y)^k}{k!} f^{(k)}(y) + \frac{(x-y)^{n+1}}{(n+1)!} f^{(n+1)}(\xi) \quad (\text{A.1})$$



#### Corollaire A.4: Théorème de la bijection

Si  $f$  est une fonction continue et strictement monotone sur un intervalle  $[a, b]$  et à valeurs réelles, alors elle constitue une bijection entre  $[a, b]$  et l'intervalle fermé dont les bornes sont  $f(a)$  et  $f(b)$ .

**1** *Proof.* Notons  $J = f^{-1}([a, b])$  cet intervalle fermé, c'est-à-dire l'ensemble des réels compris entre  $f(a)$  et  
**2**  $f(b)$ .

**3** • La monotonie de la fonction implique que l'image de l'intervalle  $[a, b]$  est contenue dans  $J$  :

**4** – si  $f$  est croissante, pour tout  $x \in [a, b]$  on a  $f(a) \leq f(x) \leq f(b)$

**5** – si  $f$  est décroissante, pour tout  $x \in [a, b]$  on a  $f(b) \leq f(x) \leq f(a)$ .

**6** • Le fait que cette monotonie soit stricte assure que deux réels distincts ne peuvent avoir la même  
**7** image, autrement dit la fonction est injective sur  $[a, b]$ .

**8** • Enfin, le théorème des valeurs intermédiaires (qui s'appuie sur l'hypothèse de continuité) garantit  
**9** que tout élément de  $J$  admet au moins un antécédent par  $f$ , c'est-à-dire que la fonction est surjective  
**10** dans  $J$ .

**11** □



### Proposition A.5

Soit  $f$  est une fonction bijective continue d'un intervalle ouvert  $I \subset \mathbb{R}$  sur un intervalle ouvert  $J \subset \mathbb{R}$ . Si  $f$  est dérivable en  $\alpha \in I$  et que  $f'(\alpha) \neq 0$  alors sa réciproque  $f^{-1}$  est dérivable en  $\beta = f(\alpha) \in J$  et

$$(f^{-1})'(\beta) = \frac{1}{f'(\alpha)} \quad \text{ou encore} \quad (f^{-1})'(\beta) = \frac{1}{f'(f^{-1}(\beta))}$$

**12**

*Proof.* On pose  $g = f^{-1}$  et on écrit son taux d'accroissement :

$$\frac{g(y) - g(\beta)}{y - \beta} = \frac{x - \alpha}{f(x) - f(\alpha)}$$

**13** avec  $y = f(x)$ . Cette fraction est l'inverse de  $\frac{f(x) - f(\alpha)}{x - \alpha}$  qui tend vers  $f'(\alpha) \neq 0$  quand  $x$  tend vers  $\alpha$ . □



### Théorème A.6: Théorème des accroissements finis

Soient  $a$  et  $b$  deux réels,  $a < b$  et  $f$  une fonction continue sur l'intervalle fermé  $[a, b]$ , dérivable sur l'intervalle ouvert  $]a, b[$ . Alors il existe  $\xi \in ]a, b[$  tel que

$$f'(\xi) = \frac{f(b) - f(a)}{b - a}.$$

**14**

## A.2 Algèbre linéaire

**15**



Toute cette partie peut être joyeusement omise par tout Homo sapiens *algebra linearis* compatible. Toutefois une lecture rapide permet de se rafraîchir la mémoire.

**16**

**17** Soit  $V$  un **espace vectoriel** de dimension finie  $n$ , sur le corps  $\mathbb{R}$  des nombres réels, ou sur le corps  $\mathbb{C}$   
**18** des nombres complexes. Notons plus généralement  $\mathbb{K}$  le corps  $\mathbb{R}$  ou  $\mathbb{C}$ .

## A.2.1 Vecteurs

1

Une **base** de  $V$  est un ensemble  $\{e_1, e_2, \dots, e_n\}$  de  $n$  **vecteurs linéairement indépendants**. Le vecteur  $= \sum_{i=1}^n v_i e_i$  sera représenté par le **vecteur colonne**

$$\mathbf{v} = \begin{pmatrix} v_1 \\ v_2 \\ \vdots \\ v_n \end{pmatrix}$$

et on désignera par  $\mathbf{v}^t$  et  $\mathbf{v}^*$  les **vecteurs lignes** suivants

$$\mathbf{v}^t = (v_1 \quad v_2 \quad \dots \quad v_n), \quad \mathbf{v}^* = (\overline{v_1} \quad \overline{v_2} \quad \dots \quad \overline{v_n})$$

où  $\overline{\alpha}$  est le nombre **complexe conjugué** du nombre  $\alpha$ .

2

### ♥ Definition A.7

- Le vecteur ligne  $\mathbf{v}^t$  est le **vecteur transposé** du vecteur colonne  $\mathbf{v}$ .
- Le vecteur ligne  $\mathbf{v}^*$  est le **vecteur adjoint** du vecteur colonne  $\mathbf{v}$ .

3

### ♥ Definition A.8

L'application  $\langle \bullet, \bullet \rangle : V \times V \rightarrow \mathbb{K}$  définie par

$$\langle \mathbf{u}, \mathbf{v} \rangle = \mathbf{u}^t \cdot \mathbf{v} = \mathbf{v}^t \cdot \mathbf{u} = \sum_{i=1}^n u_i v_i, \quad \text{si } \mathbb{K} = \mathbb{R} \quad (\text{A.2})$$

$$\langle \mathbf{u}, \mathbf{v} \rangle = \mathbf{u}^* \cdot \mathbf{v} = \overline{\mathbf{v}^* \cdot \mathbf{u}} = \overline{\langle \mathbf{v}, \mathbf{u} \rangle} = \sum_{i=1}^n \overline{u_i} v_i, \quad \text{si } \mathbb{K} = \mathbb{C} \quad (\text{A.3})$$

est appelée **produit scalaire** euclidien si  $\mathbb{K} = \mathbb{R}$ , hermitien<sup>a</sup> si  $\mathbb{K} = \mathbb{C}$ . Pour rappeler la dimension de l'espace, on écrit

$$\langle \mathbf{u}, \mathbf{v} \rangle = \langle \mathbf{u}, \mathbf{v} \rangle_n.$$

<sup>a</sup> La convention choisie pour le produit scalaire hermitien étant ici : linéarité à droite et semi-linéarité à gauche. Il est aussi possible de définir le produit scalaire hermitien par le complexe conjugué de A.3 :

$$\langle \mathbf{u}, \mathbf{v} \rangle = \mathbf{v}^* \cdot \mathbf{u} = \sum_{i=1}^n u_i \overline{v_i}.$$

Dans ce cas le produit scalaire est une forme sesquilinéaire à droite.

4

### ♥ Definition A.9

5

Soit  $V$  est un espace vectoriel muni d'un produit scalaire.

◇ Deux **vecteurs**  $\mathbf{u}$  et  $\mathbf{v}$  sont **orthogonaux** si  $\langle \mathbf{u}, \mathbf{v} \rangle = 0$ .

◇ Un **vecteur**  $\mathbf{v}$  est **orthogonal à une partie**  $U$  de  $V$  si

$$\forall \mathbf{u} \in U, \langle \mathbf{u}, \mathbf{v} \rangle = 0.$$

On note  $\mathbf{v} \perp U$ .

◇ Un ensemble de vecteurs  $\{\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_k\}$  de l'espace  $V$  est dit **orthonormal** si

$$\langle \mathbf{v}_i, \mathbf{v}_j \rangle = \delta_{ij}, \quad \forall (i, j) \in \llbracket 1, k \rrbracket^2$$

où  $\delta_{ij}$  est le **symbole de Kronecker** :  $\delta_{ij} = \begin{cases} 1 & \text{si } i = j, \\ 0 & \text{si } i \neq j. \end{cases}$

### ♥ Definition A.10

Le vecteur nul de  $\mathbb{K}^n$  est représenté par  $\mathbf{0}_n$  ou  $\mathbf{0}$  lorsqu'il n'y a pas d'ambiguïté.

### ♥ Definition A.11

Soit  $\mathbf{u} \in \mathbb{K}^n$  non nul. On définit l'**opérateur de projection** sur  $\mathbf{u}$  par

$$\text{proj}_{\mathbf{u}}(\mathbf{v}) = \frac{\langle \mathbf{u}, \mathbf{v} \rangle}{\langle \mathbf{u}, \mathbf{u} \rangle} \mathbf{u} = \frac{1}{\langle \mathbf{u}, \mathbf{u} \rangle} \mathbf{u} \mathbf{u}^* \mathbf{v}, \quad \forall \mathbf{v} \in \mathbb{K}^n. \quad (\text{A.4})$$

La matrice  $\mathbf{P}_{\mathbf{u}} = \mathbf{u} \mathbf{u}^*$  s'appelle la matrice de la projection orthogonale suivant le vecteur  $\mathbf{u}$ .

### 📖 Proposition A.12: Procédé de Gram-Schmidt

Soit  $\{\mathbf{v}_i\}_{i \in \llbracket 1, n \rrbracket}$  une base de  $\mathbb{K}^n$ . On construit successivement les vecteurs  $\mathbf{u}_i$

$$\mathbf{u}_i = \mathbf{v}_i - \sum_{k=1}^{i-1} \text{proj}_{\mathbf{u}_k}(\mathbf{v}_i) = \mathbf{v}_i - \sum_{k=1}^{i-1} \frac{\langle \mathbf{u}_k, \mathbf{v}_i \rangle}{\langle \mathbf{u}_k, \mathbf{u}_k \rangle} \mathbf{u}_k, \quad \forall i \in \llbracket 1, n \rrbracket.$$

Ils forment une **base orthogonale** de  $\mathbb{K}^n$  et  $\text{Vect}(\mathbf{u}_1, \dots, \mathbf{u}_i) = \text{Vect}(\mathbf{v}_1, \dots, \mathbf{v}_i)$ ,  $\forall i \in \llbracket 1, n \rrbracket$  (voir Exercice A.3.5, page 191).

Pour construire une **base orthonormale**  $\{\mathbf{z}_i\}_{i \in \llbracket 1, n \rrbracket}$ , il suffit de normaliser les vecteurs de la base orthogonale:

$$\mathbf{z}_i = \frac{\mathbf{u}_i}{\langle \mathbf{u}_i, \mathbf{u}_i \rangle^{1/2}}, \quad \forall i \in \llbracket 1, n \rrbracket.$$

## A.2.2 Matrices

### Généralités

Une matrice à  $m$  lignes et  $n$  colonnes est appelée *matrice de type*  $(m, n)$ , et on note  $\mathcal{M}_{m,n}(\mathbb{K})$ , ou simplement  $\mathcal{M}_{m,n}$ , l'espace vectoriel sur le corps  $\mathbb{K}$  formé par les matrices de type  $(m, n)$  à éléments dans  $\mathbb{K}$ .

Une matrice  $\mathbb{A} \in \mathcal{M}_{m,n}(\mathbb{K})$  d'éléments  $A_{ij} \in \mathbb{K}$  est notée

$$\mathbb{A} = (A_{ij})_{1 \leq i \leq m, 1 \leq j \leq n},$$

le premier indice  $i$  correspond aux lignes et le second  $j$  aux colonnes. On désigne par  $(\mathbb{A})_{ij}$  l'élément de la  $i^{\text{ème}}$  ligne et de la  $j^{\text{ème}}$  colonne. On peut aussi le noter  $A_{i,j}$ .

### ♥ Definition A.13

La matrice nulle de  $\mathcal{M}_{m,n}(\mathbb{K})$  est représentée par  $\mathbb{O}_{m,n}$  ou  $\mathbb{O}$  lorsqu'il n'y a pas d'ambiguïté. Si  $m = n$  on peut aussi noter  $\mathbb{O}_n$  cette matrice.

1

### ♥ Definition A.14

- ◇ Soit une matrice  $\mathbb{A} \in \mathcal{M}_{m,n}(\mathbb{C})$ , on note  $\mathbb{A}^* \in \mathcal{M}_{n,m}(\mathbb{C})$  la **matrice adjointe** de la matrice  $\mathbb{A}$ , définie de façon unique par

$$\langle \mathbb{A}\mathbf{u}, \mathbf{v} \rangle_m = \langle \mathbf{u}, \mathbb{A}^*\mathbf{v} \rangle_n, \quad \forall \mathbf{u} \in \mathbb{C}^n, \quad \forall \mathbf{v} \in \mathbb{C}^m$$

qui entraîne  $(\mathbb{A}^*)_{ij} = \overline{A_{ji}}$ .

- ◇ Soit une matrice  $\mathbb{A} \in \mathcal{M}_{m,n}(\mathbb{R})$ , on note  $\mathbb{A}^t \in \mathcal{M}_{n,m}(\mathbb{R})$  la **matrice transposée** de la matrice  $\mathbb{A}$ , définie de façon unique par

$$\langle \mathbb{A}\mathbf{u}, \mathbf{v} \rangle_m = \langle \mathbf{u}, \mathbb{A}^t\mathbf{v} \rangle_n, \quad \forall \mathbf{u} \in \mathbb{R}^n, \quad \forall \mathbf{v} \in \mathbb{R}^m$$

qui entraîne  $(\mathbb{A}^t)_{ij} = A_{ji}$ .

2

### ♥ Definition A.15

Si  $\mathbb{A} \in \mathcal{M}_{m,p}(\mathbb{K})$  et  $\mathbb{B} \in \mathcal{M}_{p,n}(\mathbb{K})$ , leur **produit**  $\mathbb{A}\mathbb{B} \in \mathcal{M}_{m,n}(\mathbb{K})$  est défini par

$$(\mathbb{A}\mathbb{B})_{ij} = \sum_{k=1}^p A_{ik}B_{kj}, \quad \forall i \in \llbracket 1, m \rrbracket, \quad \forall j \in \llbracket 1, n \rrbracket. \quad (\text{A.5})$$

3



### Exercice A.2.1: résultats à savoir

Soient  $\mathbb{A} \in \mathcal{M}_{m,p}(\mathbb{K})$  et  $\mathbb{B} \in \mathcal{M}_{p,n}(\mathbb{K})$ , montrer que

$$(\mathbb{A}\mathbb{B})^t = \mathbb{B}^t \mathbb{A}^t, \quad \text{si } \mathbb{K} = \mathbb{R}, \quad (\text{A.6})$$

$$(\mathbb{A}\mathbb{B})^* = \mathbb{B}^* \mathbb{A}^*, \quad \text{si } \mathbb{K} = \mathbb{C} \quad (\text{A.7})$$

4

Les matrices considérées jusqu'à la fin de ce paragraphe sont carrées.

5

### ♥ Definition A.16

Si  $\mathbb{A} \in \mathcal{M}_n$  alors les éléments  $A_{ii} = (\mathbb{A})_{ii}$  sont appelés **éléments diagonaux** et les éléments  $A_{ii} = (\mathbb{A})_{ij}, i \neq j$  sont appelés **éléments hors-diagonaux**.

6

### ♥ Definition A.17

On appelle **matrice identité** de  $\mathcal{M}_n$  la matrice dont les éléments diagonaux sont tous égaux à 1 et les éléments hors-diagonaux nuls. On la note  $\mathbb{I}$  ou encore  $\mathbb{I}_n$  et on a

$$(\mathbb{I})_{i,j} = \delta_{ij}, \quad \forall (i, j) \in \llbracket 1, n \rrbracket^2.$$

7

### ♥ Definition A.18

8

Une matrice  $A \in \mathcal{M}_n$  est **inversible** ou **régulière** s'il existe une **unique** matrice de  $\mathcal{M}_n$ , notée  $A^{-1}$  et appelée **matrice inverse** de la matrice  $A$ , telle que

$$AA^{-1} = A^{-1}A = \mathbb{I} \quad (\text{A.8})$$

Dans le cas contraire, on dit que la matrice  $A$  est **singulière** ou **non inversible**.



### Exercice A.2.2: résultats à savoir

Soient  $A$  et  $B$  deux matrices de  $\mathcal{M}_n$ . Montrer que

$$AB = \mathbb{I} \Rightarrow BA = \mathbb{I}.$$

Que peut-on en conclure?



### Exercice A.2.3: résultats à savoir

Soient  $A \in \mathcal{M}_n(\mathbb{K})$  et  $B \in \mathcal{M}_n(\mathbb{K})$  inversibles. Montrer que  $AB$  inversible et

$$(A^t)^{-1} = (A^{-1})^t, \text{ si } \mathbb{K} = \mathbb{R}, \quad (\text{A.9})$$

$$(A^*)^{-1} = (A^{-1})^*, \text{ si } \mathbb{K} = \mathbb{C}. \quad (\text{A.10})$$

$$(AB)^{-1} = B^{-1}A^{-1} \quad (\text{A.11})$$

$$(A^{-1})^{-1} = A \quad (\text{A.12})$$

### ♥ Definition A.19

Une matrice **carrée**  $A$  est :

- ◇ **symétrique** si  $A$  est réelle et  $A = A^t$ ,
- ◇ **hermitienne** si  $A = A^*$ ,
- ◇ **normale** si  $AA^* = A^*A$ ,
- ◇ **orthogonale** si  $A$  est réelle et  $AA^t = A^tA = \mathbb{I}$ ,
- ◇ **unitaire** si  $AA^* = A^*A = \mathbb{I}$ ,



### Proposition A.20

- une matrice symétrique ou hermitienne est nécessairement normale.
- une matrice orthogonale (resp. unitaire) est nécessairement normale et inversible d'inverse  $A^t$  (resp.  $A^*$ ).

### ♥ Definition A.21

Soit  $\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$  une matrice **hermitienne**.

◇ Elle est **définie positive** si

$$\langle \mathbb{A}u, u \rangle > 0, \quad \forall u \in \mathbb{C}^n \setminus \{0\} \quad (\text{A.13})$$

◇ Elle est **semi définie positive** si

$$\langle \mathbb{A}u, u \rangle \geq 0, \quad \forall u \in \mathbb{C}^n \setminus \{0\} \quad (\text{A.14})$$

1



#### Exercice A.2.4

Soit  $\mathbb{A} \in \mathcal{M}_n$ . Que peut-on dire de la matrice  $\mathbb{A}\mathbb{A}^*$ ? Et si la matrice  $\mathbb{A}$  est inversible? Proposer une technique permettant de générer une matrice semi-définie positive à partir d'une matrice aléatoire.

2



#### Definition A.22

Soit  $\mathbb{A} \in \mathcal{M}_n$ . La trace d'une matrice carrée  $\mathbb{A} = (a_{ij})$  est définie par

$$\text{tr}(\mathbb{A}) = \sum_{i=1}^n a_{ii}.$$

3



#### Definition A.23

Soit  $\mathcal{T}_n$  le **groupe des permutations** de l'ensemble  $\{1, 2, \dots, n\}$ . À tout élément  $\sigma \in \mathcal{T}_n$ , on associe la **matrice de permutation** de  $\mathbb{P}_\sigma \in \mathcal{M}_n$  est définie par

$$(\mathbb{P}_\sigma)_{i,j} = \delta_{i\sigma(j)}.$$

4



#### Exercice A.2.5: résultats à savoir

Montrer qu'une matrice de permutation est orthogonale.

5



#### Definition A.24

Soient  $\mathbb{A} = (A_{i,j})_{i,j=1}^n \in \mathcal{M}_n$  et  $\mathcal{T}_n$  le **groupe des permutations** de l'ensemble  $\{1, 2, \dots, n\}$ .. Le **déterminant** d'une matrice  $\mathbb{A}$  est défini par

$$\det(\mathbb{A}) = \sum_{\sigma \in \mathcal{T}_n} \varepsilon_\sigma \prod_{j=1}^n A_{\sigma(j),j}$$

où  $\varepsilon_\sigma$  désigne la signature de la permutation  $\sigma$ .

6



#### Proposition A.25: Méthode de Laplace ou des cofacteurs

7

Soit  $\mathbb{A} = (A_{i,j})_{i,j=1}^n \in \mathcal{M}_n$ . On note  $\mathbb{A}^{[i,j]} \in \mathcal{M}_{n-1}$  la matrice obtenue en supprimant la ligne  $i$  et la colonne  $j$  de  $\mathbb{A}$ . On a alors le **développement par rapport à la ligne  $i \in \llbracket 1, n \rrbracket$**

$$\det(\mathbb{A}) = \sum_{j=1}^n (-1)^{i+j} A_{i,j} \det(\mathbb{A}^{[i,j]}), \quad (\text{A.15})$$

et le **développement par rapport à la colonne  $j \in \llbracket 1, n \rrbracket$**

$$\det(\mathbb{A}) = \sum_{i=1}^n (-1)^{i+j} A_{i,j} \det(\mathbb{A}^{[i,j]}). \quad (\text{A.16})$$

Le terme  $(-1)^{i+j} \det(\mathbb{A}^{[i,j]})$  est appelé le **cofacteur** du terme  $A_{i,j}$ .

### ♥ Definition A.26

Soit  $\mathbb{A} \in \mathcal{M}_n(\mathbb{K})$ . On dit que  $\lambda \in \mathbb{C}$  est **valeur propre** de  $\mathbb{A}$  s'il existe  $\mathbf{u} \in \mathbb{C}^n$  **non nul** tel que

$$\mathbb{A}\mathbf{u} = \lambda\mathbf{u}. \quad (\text{A.17})$$

Le vecteur  $\mathbf{u}$  est appelé **vecteur propre** associé à la valeur propre  $\lambda$ .

Le couple  $(\lambda, \mathbf{u})$  est appelé **élément propre** de  $\mathbb{A}$ .

### ♥ Definition A.27

Soit  $\mathbb{A} \in \mathcal{M}_n(\mathbb{K})$ . Soit  $\lambda \in \mathbb{C}$  une valeur propre de  $\mathbb{A}$ . Le sous-espace

$$E_\lambda = \{\mathbf{u} \in \mathbb{C}^n : \mathbb{A}\mathbf{u} = \lambda\mathbf{u}\} = \ker(\mathbb{A} - \lambda\mathbb{I}) \quad (\text{A.18})$$

est appelé **sous-espace propre** associé à la valeur propre  $\lambda$ .

### ♥ Definition A.28

Soit  $\mathbb{A} \in \mathcal{M}_n(\mathbb{K})$ . Le polynôme de degré  $n$  défini par

$$\mathcal{P}_\mathbb{A}(\lambda) = \det(\mathbb{A} - \lambda\mathbb{I}) \quad (\text{A.19})$$

est appelé **polynôme caractéristique** de la matrice  $\mathbb{A}$ .

### 📖 Proposition A.29

Soit  $\mathbb{A} \in \mathcal{M}_n(\mathbb{K})$ .

- ◇ Les racines complexes du polynôme caractéristique  $\mathcal{P}_\mathbb{A}$  sont les valeurs propres de la matrice  $\mathbb{A}$ .
- ◇ Si la racine  $\lambda$  de  $\mathcal{P}_\mathbb{A}$  est de multiplicité  $k$ , on dit que la valeur propre  $\lambda$  est de multiplicité algébrique  $k$ .
- ◇ La matrice  $\mathbb{A}$  possède  $n$  valeurs propres distinctes ou non.

### ♥ Definition A.30

Soit  $A \in \mathcal{M}_n(\mathbb{K})$ . On note  $\lambda_i(A)$ ,  $i \in \llbracket 1, n \rrbracket$ , les  $n$  valeurs propres de  $A$ . Le **spectre** de la matrice  $A$  est le sous-ensemble

$$\text{Sp}(A) = \bigcup_{i=1}^n \{\lambda_i(A)\} \quad (\text{A.20})$$

du plan complexe.

1



### Proposition A.31

Soient  $A \in \mathcal{M}_n$  et  $B \in \mathcal{M}_n$ . On a les relations suivantes

$$\text{tr}(A) = \sum_{i=1}^n \lambda_i(A), \quad (\text{A.21})$$

$$\det(A) = \prod_{i=1}^n \lambda_i(A), \quad (\text{A.22})$$

$$\text{tr}(AB) = \text{tr}(BA), \quad (\text{A.23})$$

$$\text{tr}(A + B) = \text{tr} A + \text{tr} B, \quad (\text{A.24})$$

$$\det(AB) = \det(A) \det(B) = \det(BA), \quad (\text{A.25})$$

$$\det(A^*) = \overline{\det(A)}. \quad (\text{A.26})$$

2



### Definition A.32

Le **rayon spectral** d'une matrice  $A \in \mathcal{M}_n$  est le nombre  $\geq 0$  défini par

$$\rho(A) = \max \{|\lambda_i(A)|; i \in \llbracket 1, n \rrbracket\}$$

3



### Definition A.33

Soit  $X$  et  $Y$  deux espaces vectoriels

◇ On note  $\ker(A) = \{v \in X; Av = 0\}$  le **noyau** de l'application linéaire  $A : X \rightarrow Y$ .

◇ On note  $\text{im}(A) = \{Av \in Y; v \in X\}$  l'**image** de l'application linéaire  $A : X \rightarrow Y$ .

4



### Definition A.34

Le **rang** d'une matrice  $A$ , noté  $\text{rang}(A)$ , est défini comme le rang de l'application linéaire  $A$  qui lui est associé

$$\text{rang}(A) = \dim(\text{im}(A)).$$

5

## Matrices particulières

6



### Definition A.35

7

Une matrice carrée  $A \in \mathcal{M}_n$  est :

- ◇ **diagonale** si  $a_{ij} = 0$  pour  $i \neq j$ ,
- ◇ **triangulaire supérieure** si  $a_{ij} = 0$  pour  $i > j$ ,
- ◇ **triangulaire inférieure** si  $a_{ij} = 0$  pour  $i < j$ ,
- ◇ **triangulaire** si elle est triangulaire supérieure ou triangulaire inférieure
- ◇ **à diagonale dominante** si

$$|a_{ii}| \geq \sum_{j \neq i} |a_{ij}|, \quad \forall i \in \llbracket 1, n \rrbracket, \quad (\text{A.27})$$

- ◇ **à diagonale strictement dominante** si

$$|a_{ii}| > \sum_{j \neq i} |a_{ij}|, \quad \forall i \in \llbracket 1, n \rrbracket. \quad (\text{A.28})$$

### Proposition A.36

Soient  $A$  et  $B$  deux matrices de  $\mathcal{M}_n(\mathbb{K})$  triangulaires inférieures (resp. triangulaires supérieures). Alors la matrice  $AB$  est aussi triangulaire inférieure (resp. triangulaire supérieure). De plus on a

$$(AB)_{i,i} = A_{i,i}B_{i,i}, \quad \forall i \in \llbracket 1, n \rrbracket.$$

*Proof.* (voir Exercice A.3.2, page 190) □

### Proposition A.37

Soit  $A \in \mathcal{M}_n(\mathbb{K})$  une matrice triangulaire inférieure (resp. triangulaire supérieure).

1.  $A$  est inversible si et seulement si ses éléments diagonaux sont tous non nuls (i.e.  $A_{i,i} \neq 0$ ,  $\forall i \in \llbracket 1, n \rrbracket$ ).
2. Si  $A$  est inversible alors son inverse est triangulaire inférieure (resp. triangulaire supérieure) et

$$(A^{-1})_{i,i} = \frac{1}{(A)_{i,i}}$$

*Proof.* (voir Exercice A.3.10, page 196) □

### Definition A.38

On appelle **matrice bande** une matrice  $A$  telle que  $a_{ij} \neq 0$  pour  $|j - i| \leq c$ .  $c$  est la **demi largeur de bande**.

Lorsque  $c = 1$ , la matrice est dite **tridiagonale**. Lorsque  $c = 2$ , la matrice est dite **pentadiagonale**.

### Definition A.39

On appelle **sous-matrice** d'une matrice donnée, la matrice obtenue en supprimant certaines lignes et certaines colonnes. En particulier, si on supprime les  $(n - k)$  dernières lignes et colonnes d'une matrice carrée  $\mathbb{A}$  d'ordre  $n$ , on obtient la **sous matrice principale** d'ordre  $k$ .

1

### ♥ Définition A.40

On appelle **matrice bloc** une matrice  $\mathbb{A} \in \mathcal{M}_{N,M}$  écrite sous la forme

$$\mathbb{A} = \begin{pmatrix} \mathbb{A}_{1,1} & \cdots & \mathbb{A}_{1,q} \\ \vdots & & \vdots \\ \mathbb{A}_{p,1} & \cdots & \mathbb{A}_{p,q} \end{pmatrix}$$

où  $\forall i \in \llbracket 1, p \rrbracket, \forall j \in \llbracket 1, q \rrbracket, \mathbb{A}_{i,j}$  est une matrice de  $\mathcal{M}_{n_i, m_j}$ . On a  $N = \sum_{i=1}^p n_i$  et  $M = \sum_{j=1}^q m_j$ .

On dit que  $\mathbb{A}$  est une matrice **bloc-carrée** si  $p = q$  et si tous les blocs diagonaux sont des matrices carrées.

2

### 📖 Propriété A.41: Multiplication de matrices blocs

Soient  $\mathbb{A} \in \mathcal{M}_{N,M}$  et  $\mathbb{B} \in \mathcal{M}_{M,S}$ . Le produit  $\mathbb{P} = \mathbb{A}\mathbb{B} \in \mathcal{M}_{N,S}$  peut s'écrire sous forme bloc si les matrices  $\mathbb{A}$  et  $\mathbb{B}$  sont *compatibles par blocs* : il faut que le nombre de blocs colonne de  $\mathbb{A}$  soit égale au nombre de blocs ligne de  $\mathbb{B}$  avec correspondance des dimensions.

$$\mathbb{A} = \begin{pmatrix} \mathbb{A}_{1,1} & \cdots & \mathbb{A}_{1,q} \\ \vdots & & \vdots \\ \mathbb{A}_{p,1} & \cdots & \mathbb{A}_{p,q} \end{pmatrix} \text{ et } \mathbb{B} = \begin{pmatrix} \mathbb{B}_{1,1} & \cdots & \mathbb{B}_{1,r} \\ \vdots & & \vdots \\ \mathbb{B}_{q,1} & \cdots & \mathbb{B}_{q,r} \end{pmatrix}$$

avec  $\mathbb{A}_{i,k} \in \mathcal{M}_{n_i, m_k}$  et  $\mathbb{B}_{k,j} \in \mathcal{M}_{m_k, s_j}$  pour tout  $i \in \llbracket 1, p \rrbracket, k \in \llbracket 1, q \rrbracket$  et  $j \in \llbracket 1, r \rrbracket$ . La matrice produit  $\mathbb{P}$  s'écrit alors sous la forme bloc

$$\mathbb{P} = \begin{pmatrix} \mathbb{P}_{1,1} & \cdots & \mathbb{P}_{1,r} \\ \vdots & & \vdots \\ \mathbb{P}_{p,1} & \cdots & \mathbb{P}_{p,r} \end{pmatrix}$$

avec  $\forall i \in \llbracket 1, p \rrbracket, \forall j \in \llbracket 1, r \rrbracket \mathbb{P}_{i,j} \in \mathcal{M}_{n_i, s_j}$  et

$$\mathbb{P}_{i,j} = \sum_{k=1}^q \mathbb{A}_{i,k} \mathbb{B}_{k,j}.$$

3

### ♥ Définition A.42

On dit qu'une matrice bloc-carrée  $\mathbb{A}$  est **triangulaire inférieure** (resp. **supérieure**) **par blocs** si elle peut s'écrire sous la forme d'une matrice bloc avec les sous matrices  $\mathbb{A}_{i,j} = 0$  pour  $i < j$  (resp.  $i > j$ ). Elle s'écrit donc sous la forme

$$\mathbb{A} = \begin{pmatrix} \mathbb{A}_{1,1} & 0 & \cdots & 0 \\ \vdots & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ \mathbb{A}_{n,1} & \cdots & \cdots & \mathbb{A}_{n,n} \end{pmatrix} \quad (\text{resp. } \mathbb{A} = \begin{pmatrix} \mathbb{A}_{1,1} & \cdots & \cdots & \mathbb{A}_{n,1} \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & \vdots \\ 0 & \cdots & 0 & \mathbb{A}_{n,n} \end{pmatrix}).$$

4

### ♥ Definition A.43

On dit qu'une matrice bloc-carrée  $\mathbb{A}$  est **diagonale par blocs** ou **bloc-diagonale** si elle peut s'écrire sous la forme d'une matrice bloc avec les sous matrices  $\mathbb{A}_{i,j} = 0$  pour  $i \neq j$ . Elle s'écrit donc sous la forme

$$\mathbb{A} = \begin{pmatrix} \mathbb{A}_{1,1} & 0 & \cdots & 0 \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \cdots & 0 & \mathbb{A}_{n,n} \end{pmatrix}$$

1

### 📖 Proposition A.44

Soit  $\mathbb{A}$  une matrice bloc-carré décomposée en  $n \times n$  blocs. Si  $\mathbb{A}$  est **bloc-diagonale** ou **triangulaire par blocs** alors son déterminant est le produit des déterminant des blocs diagonaux :

$$\det \mathbb{A} = \prod_{i=1}^n \det \mathbb{A}_{i,i} \quad (\text{A.29})$$

2

### 📖 Proposition A.45

Soit  $\mathbb{A}$  une matrice bloc-carré **inversible** décomposée en  $n \times n$  blocs.

- Si  $\mathbb{A}$  est **bloc-diagonale** alors son inverse (décomposée en  $n \times n$  blocs) est aussi **bloc-diagonale**.
- Si  $\mathbb{A}$  est **triangulaire inférieure par blocs** (resp. supérieure) alors son inverse (décomposée en  $n \times n$  blocs) est aussi **triangulaire inférieure par blocs** (resp. supérieure).

Dans ces deux cas les blocs diagonaux de la matrice inverse sont les inverses des blocs diagonaux de  $\mathbb{A}$ . On a donc

$$\begin{aligned} \mathbb{A} &= \begin{pmatrix} \mathbb{A}_{1,1} & 0 & \cdots & 0 \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \cdots & 0 & \mathbb{A}_{n,n} \end{pmatrix} \quad \text{et} \quad \mathbb{A}^{-1} = \begin{pmatrix} \mathbb{A}_{1,1}^{-1} & 0 & \cdots & 0 \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \cdots & 0 & \mathbb{A}_{n,n}^{-1} \end{pmatrix} \\ \mathbb{A} &= \begin{pmatrix} \mathbb{A}_{1,1} & 0 & \cdots & 0 \\ \vdots & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ \mathbb{A}_{n,1} & \cdots & \cdots & \mathbb{A}_{n,n} \end{pmatrix} \quad \text{et} \quad \mathbb{A}^{-1} = \begin{pmatrix} \mathbb{A}_{1,1}^{-1} & 0 & \cdots & 0 \\ \bullet & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ \bullet & \cdots & \bullet & \mathbb{A}_{n,n}^{-1} \end{pmatrix} \\ \mathbb{A} &= \begin{pmatrix} \mathbb{A}_{1,1} & \cdots & \cdots & \mathbb{A}_{n,1} \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & \vdots \\ 0 & \cdots & 0 & \mathbb{A}_{n,n} \end{pmatrix} \quad \text{et} \quad \mathbb{A}^{-1} = \begin{pmatrix} \mathbb{A}_{1,1}^{-1} & \bullet & \cdots & \bullet \\ 0 & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & \bullet \\ 0 & \cdots & 0 & \mathbb{A}_{n,n}^{-1} \end{pmatrix} \end{aligned}$$

3

## A.2.3 Normes vectorielles et normes matricielles

4

### ♥ Définition A.46

Une **norme** sur un espace vectoriel  $V$  est une application  $\|\bullet\| : V \rightarrow \mathbb{R}^+$  qui vérifie les propriétés suivantes

- ◇  $\|\mathbf{v}\| = 0 \iff \mathbf{v} = 0$ ,
- ◇  $\|\alpha\mathbf{v}\| = |\alpha| \|\mathbf{v}\|$ ,  $\forall \alpha \in \mathbb{K}$ ,  $\forall \mathbf{v} \in V$ ,
- ◇  $\|\mathbf{u} + \mathbf{v}\| \leq \|\mathbf{u}\| + \|\mathbf{v}\|$ ,  $\forall (\mathbf{u}, \mathbf{v}) \in V^2$  (inégalité triangulaire).

Une norme sur  $V$  est également appelée **norme vectorielle**. On appelle **espace vectoriel normé** un espace vectoriel muni d'une norme.

1

Les trois normes suivantes sont les plus couramment utilisées :

$$\begin{aligned}\|\mathbf{v}\|_1 &= \sum_{i=1}^n |v_i| \\ \|\mathbf{v}\|_2 &= \left( \sum_{i=1}^n |v_i|^2 \right)^{1/2} \\ \|\mathbf{v}\|_\infty &= \max_i |v_i|.\end{aligned}$$

### 📖 Théorème A.47

Soit  $V$  un espace de dimension finie. Pour tout nombre réel  $p \geq 1$ , l'application  $\|\bullet\|_p$  définie par

$$\|\mathbf{v}\|_p = \left( \sum_{i=1}^n |v_i|^p \right)^{1/p}$$

est une norme.

2

**Claim A.48** Pour  $p > 1$  et  $\frac{1}{p} + \frac{1}{q} = 1$ , on a  $\forall \mathbf{u}, \mathbf{v} \in \mathbb{K}^n$

3

$$\sum_{i=1}^n |u_i v_i| \leq \left( \sum_{i=1}^n |u_i|^p \right)^{1/p} \left( \sum_{i=1}^n |v_i|^q \right)^{1/q} = \|\mathbf{u}\|_p \|\mathbf{v}\|_q. \quad (\text{A.30})$$

Cette inégalité s'appelle l'**inégalité de Hölder**.

4

### ♥ Définition A.49

Deux **normes**  $\|\bullet\|$  et  $\|\bullet\|'$ , définies sur un même espace vectoriel  $V$ , sont **équivalentes** s'il existe deux constantes  $C$  et  $C'$  telles que

$$\|\mathbf{v}\|' \leq C \|\mathbf{v}\| \quad \text{et} \quad \|\mathbf{v}\| \leq C' \|\mathbf{v}\|' \quad \text{pour tout } \mathbf{v} \in V. \quad (\text{A.31})$$

5

**Claim A.50** Sur un espace vectoriel de dimension finie toutes les normes sont équivalentes.

6

### ♥ Définition A.51

7

Une **norme matricielle** sur  $\mathcal{M}_n(\mathbb{K})$  est une application  $\|\bullet\| : \mathcal{M}_n(\mathbb{K}) \rightarrow \mathbb{R}^+$  vérifiant

1.  $\|\mathbb{A}\| = 0 \iff \mathbb{A} = 0$ ,
2.  $\|\alpha\mathbb{A}\| = |\alpha| \|\mathbb{A}\|, \forall \alpha \in \mathbb{K}, \forall \mathbb{A} \in \mathcal{M}_n(\mathbb{K})$ ,
3.  $\|\mathbb{A} + \mathbb{B}\| \leq \|\mathbb{A}\| + \|\mathbb{B}\|, \forall (\mathbb{A}, \mathbb{B}) \in \mathcal{M}_n(\mathbb{K})^2$  (inégalité triangulaire)
4.  $\|\mathbb{A}\mathbb{B}\| \leq \|\mathbb{A}\| \|\mathbb{B}\|, \forall (\mathbb{A}, \mathbb{B}) \in \mathcal{M}_n(\mathbb{K})^2$

**Claim A.52** Etant donné une norme vectorielle  $\|\bullet\|$  sur  $\mathbb{K}^n$ , l'application  $\|\bullet\|_s : \mathcal{M}_n(\mathbb{K}) \rightarrow \mathbb{R}^+$  définie par

$$\|\mathbb{A}\|_s = \sup_{\substack{\mathbf{v} \in \mathbb{K}^n \\ \mathbf{v} \neq 0}} \frac{\|\mathbb{A}\mathbf{v}\|}{\|\mathbf{v}\|} = \sup_{\substack{\mathbf{v} \in \mathbb{K}^n \\ \|\mathbf{v}\| \leq 1}} \|\mathbb{A}\mathbf{v}\| = \sup_{\substack{\mathbf{v} \in \mathbb{K}^n \\ \|\mathbf{v}\| = 1}} \|\mathbb{A}\mathbf{v}\|, \quad (\text{A.32})$$

est une norme matricielle, appelée **norme matricielle subordonnée** (à la norme vectorielle donnée).

De plus

$$\|\mathbb{A}\mathbf{v}\| \leq \|\mathbb{A}\|_s \|\mathbf{v}\| \quad \forall \mathbf{v} \in \mathbb{K}^n \quad (\text{A.33})$$

et la norme  $\|\mathbb{A}\|$  peut se définir aussi par

$$\|\mathbb{A}\|_s = \inf \{ \alpha \in \mathbb{R} : \|\mathbb{A}\mathbf{v}\| \leq \alpha \|\mathbf{v}\|, \forall \mathbf{v} \in \mathbb{K}^n \}. \quad (\text{A.34})$$

Il existe au moins un vecteur  $\mathbf{u} \in \mathbb{K}^n$  tel que

$$\mathbf{u} \neq 0 \quad \text{et} \quad \|\mathbb{A}\mathbf{u}\| = \|\mathbb{A}\|_s \|\mathbf{u}\|. \quad (\text{A.35})$$

Enfin une norme subordonnée vérifie toujours

$$\|\mathbb{I}\|_s = 1 \quad (\text{A.36})$$



### Théorème A.53

Soit  $\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$ . On a

$$\|\mathbb{A}\|_1 \stackrel{\text{déf.}}{=} \sup_{\substack{\mathbf{v} \in \mathbb{C}^n \\ \mathbf{v} \neq 0}} \frac{\|\mathbb{A}\mathbf{v}\|_1}{\|\mathbf{v}\|_1} = \max_{j \in \llbracket 1, n \rrbracket} \sum_{i=1}^n |a_{ij}| \quad (\text{A.37})$$

$$\|\mathbb{A}\|_2 \stackrel{\text{déf.}}{=} \sup_{\substack{\mathbf{v} \in \mathbb{C}^n \\ \mathbf{v} \neq 0}} \frac{\|\mathbb{A}\mathbf{v}\|_2}{\|\mathbf{v}\|_2} = \sqrt{\rho(\mathbb{A}^* \mathbb{A})} = \sqrt{\rho(\mathbb{A} \mathbb{A}^*)} = \|\mathbb{A}^*\|_2 \quad (\text{A.38})$$

$$\|\mathbb{A}\|_\infty \stackrel{\text{déf.}}{=} \sup_{\substack{\mathbf{v} \in \mathbb{C}^n \\ \mathbf{v} \neq 0}} \frac{\|\mathbb{A}\mathbf{v}\|_\infty}{\|\mathbf{v}\|_\infty} = \max_{i \in \llbracket 1, n \rrbracket} \sum_{j=1}^n |a_{ij}| \quad (\text{A.39})$$

La norme  $\|\bullet\|_2$  est invariante par transformation unitaire :

$$\mathbb{U}\mathbb{U}^* = \mathbb{I} \implies \|\mathbb{A}\|_2 = \|\mathbb{A}\mathbb{U}\|_2 = \|\mathbb{U}\mathbb{A}\|_2 = \|\mathbb{U}^* \mathbb{A} \mathbb{U}\|_2. \quad (\text{A.40})$$

Par ailleurs, si la matrice  $\mathbb{A}$  est normale :

$$\mathbb{A}\mathbb{A}^* = \mathbb{A}^* \mathbb{A} \implies \|\mathbb{A}\|_2 = \rho(\mathbb{A}). \quad (\text{A.41})$$

**Remarque A.54** 1. Si une matrice  $\mathbb{A}$  est hermitienne, ou symétrique (donc normale), on a  $\|\mathbb{A}\|_2 = \rho(\mathbb{A})$ .

2. Si une matrice  $\mathbb{A}$  est unitaire, ou orthogonale (donc normale), on a  $\|\mathbb{A}\|_2 = 1$ .

**Théorème A.55**

1. Soit  $\mathbb{A}$  une matrice carrée quelconque et  $\|\bullet\|$  une norme matricielle subordonnée ou non, quelconque. Alors

$$\rho(\mathbb{A}) \leq \|\mathbb{A}\|. \quad (\text{A.42})$$

2. Etant donné une matrice  $\mathbb{A}$  et un nombre  $\varepsilon > 0$ , il existe au moins une norme matricielle subordonnée telle que

$$\|\mathbb{A}\| \leq \rho(\mathbb{A}) + \varepsilon. \quad (\text{A.43})$$

1

**Théorème A.56**

L'application  $\|\bullet\|_E : \mathcal{M}_n \rightarrow \mathbb{R}^+$  définie par

$$\|\mathbb{A}\|_E = \left( \sum_{(i,j) \in \llbracket 1, n \rrbracket^2} |a_{ij}|^2 \right)^{1/2} = \sqrt{\text{tr}(\mathbb{A}^* \mathbb{A})}, \quad (\text{A.44})$$

pour toute matrice  $\mathbb{A} = (a_{ij})$  d'ordre  $n$ , est une norme matricielle non subordonnée (pour  $n \geq 2$ ), invariante par transformation unitaire et qui vérifie

$$\|\mathbb{A}\|_2 \leq \|\mathbb{A}\|_E \leq \sqrt{n} \|\mathbb{A}\|_2, \quad \forall \mathbb{A} \in \mathcal{M}_n. \quad (\text{A.45})$$

De plus  $\|\mathbb{I}\|_E = \sqrt{n}$ .

2

**Théorème A.57**

1. Soit  $\|\bullet\|$  une norme matricielle subordonnée, et  $\mathbb{B}$  une matrice vérifiant

$$\|\mathbb{B}\| < 1.$$

Alors la matrice  $(\mathbb{I} + \mathbb{B})$  est inversible, et

$$\|(\mathbb{I} + \mathbb{B})^{-1}\| \leq \frac{1}{1 - \|\mathbb{B}\|}.$$

2. Si une matrice de la forme  $(\mathbb{I} + \mathbb{B})$  est singulière, alors nécessairement

$$\|\mathbb{B}\| \geq 1$$

pour toute norme matricielle, subordonnée ou non.

3

**A.2.4 Réduction des matrices**

4

**Definition A.58**

5

Soit  $A : V \rightarrow V$  une application linéaire, représenté par une matrice carrée  $\mathbb{A} \in \mathcal{M}_n$  relativement à une base  $\{\mathbf{e}_i\}_{i \in \llbracket 1, n \rrbracket}$ . Relativement à une autre base  $\{\mathbf{f}_i\}_{i \in \llbracket 1, n \rrbracket}$ , la même application est représentée par la matrice

$$\mathbb{B} = \mathbb{P}^{-1} \mathbb{A} \mathbb{P} \quad (\text{A.46})$$

où  $\mathbb{P}$  est la matrice inversible dont le  $j$ -ème vecteur colonne est formé des composantes du vecteur  $\mathbf{f}_j$  dans la base  $\{\mathbf{e}_i\}_{i \in \llbracket 1, n \rrbracket}$  :

$$\mathbb{P} = \begin{pmatrix} \langle \mathbf{e}_1, \mathbf{f}_1 \rangle & \langle \mathbf{e}_1, \mathbf{f}_2 \rangle & \cdots & \langle \mathbf{e}_1, \mathbf{f}_n \rangle \\ \langle \mathbf{e}_2, \mathbf{f}_1 \rangle & \langle \mathbf{e}_2, \mathbf{f}_2 \rangle & \ddots & \vdots \\ \vdots & \ddots & \ddots & \langle \mathbf{e}_{n-1}, \mathbf{f}_n \rangle \\ \langle \mathbf{e}_n, \mathbf{f}_1 \rangle & \cdots & \langle \mathbf{e}_n, \mathbf{f}_{n-1} \rangle & \langle \mathbf{e}_n, \mathbf{f}_n \rangle \end{pmatrix} \quad (\text{A.47})$$

La matrice  $\mathbb{P}$  est appelée **matrice de passage de la base**  $\{\mathbf{e}_i\}_{i \in \llbracket 1, n \rrbracket}$  **dans le base**  $\{\mathbf{f}_i\}_{i \in \llbracket 1, n \rrbracket}$ .

### ♥ Definition A.59

On dit que la matrice carrée  $\mathbb{A}$  est diagonalisable s'il existe une matrice inversible  $\mathbb{P}$  telle que la matrice  $\mathbb{P}^{-1} \mathbb{A} \mathbb{P}$  soit diagonale.

**Claim A.60** On notera que, dans le cas où  $\mathbb{A} \in \mathcal{M}_n$  est diagonalisable, les éléments diagonaux de la matrice  $\mathbb{P}^{-1} \mathbb{A} \mathbb{P}$  sont les valeurs propres  $\lambda_1, \lambda_2, \dots, \lambda_n$  de la matrice  $\mathbb{A}$ , et que le  $j$ -ème vecteur colonne  $\mathbf{p}_j$  de la matrice  $\mathbb{P}$  est formé des composantes, dans la même base que  $\mathbb{A}$ , d'un vecteur propre associé à la valeur propre  $\lambda_j$ . On a

$$\mathbb{P}^{-1} \mathbb{A} \mathbb{P} = \text{diag}(\lambda_1, \dots, \lambda_n) \iff \mathbb{A} \mathbf{p}_j = \lambda_j \mathbf{p}_j, \quad \forall j \in \llbracket 1, n \rrbracket. \quad (\text{A.48})$$

C'est à dire qu'une matrice est diagonalisable si, et seulement si, il existe une base de vecteurs propres.

### 📖 Théorème A.61

1. Etant donnée une matrice **carrée**  $\mathbb{A}$ , il existe une matrice **unitaire**  $\mathbb{U}$  telle que la matrice  $\mathbb{U}^{-1} \mathbb{A} \mathbb{U}$  soit **triangulaire**.
2. Etant donnée une matrice **normale**  $\mathbb{A}$ , il existe une matrice **unitaire**  $\mathbb{U}$  telle que la matrice  $\mathbb{U}^{-1} \mathbb{A} \mathbb{U}$  soit **diagonale**.
3. Etant donnée une matrice **symétrique**  $\mathbb{A}$ , il existe une matrice **orthogonale**  $\mathbb{O}$  telle que la matrice  $\mathbb{O}^{-1} \mathbb{A} \mathbb{O}$  soit **diagonale**.

## A.2.5 Suites de vecteurs et de matrices

### ♥ Definition A.62

Soit  $V$  un espace vectoriel muni d'une norme  $\|\bullet\|$ , on dit qu'une suite  $(\mathbf{v}_k)$  d'éléments de  $V$  **converge vers un élément**  $\mathbf{v} \in V$ , si

$$\lim_{k \rightarrow \infty} \|\mathbf{v}_k - \mathbf{v}\| = 0$$

et on écrit

$$\mathbf{v} = \lim_{k \rightarrow \infty} \mathbf{v}_k.$$

**Théorème A.63**

Soit  $\mathbb{B}$  une matrice carrée. Les conditions suivantes sont équivalentes :

1.  $\lim_{k \rightarrow \infty} \mathbb{B}^k = 0$ ,
2.  $\lim_{k \rightarrow \infty} \mathbb{B}^k \mathbf{v} = 0$  pour tout vecteur  $\mathbf{v}$ ,
3.  $\rho(\mathbb{B}) < 1$ ,
4.  $\|\mathbb{B}\| < 1$  pour au moins une norme matricielle subordonnée  $\|\bullet\|$ .

1

**Théorème A.64**

Soit  $\mathbb{B}$  une matrice carrée, et  $\|\bullet\|$  une norme matricielle quelconque. Alors

$$\lim_{k \rightarrow \infty} \|\mathbb{B}^k\|^{1/k} = \rho(\mathbb{B}).$$

2

## A.3 Receuil d'exercices

3

### A.3.1 Algèbre linéaire

4

Sur les matrices

5

**Exercice A.3.1**

Soit  $\mathbb{A} \in \mathcal{M}_{m,n}(\mathbb{R})$  et  $\mathbb{B} \in \mathcal{M}_{n,m}(\mathbb{R})$  telles que

$$\langle \mathbb{A}\mathbf{u}, \mathbf{v} \rangle_m = \langle \mathbf{u}, \mathbb{B}\mathbf{v} \rangle_n, \quad \forall \mathbf{u} \in \mathbb{R}^n, \quad \forall \mathbf{v} \in \mathbb{R}^m.$$

Exprimer les éléments de la matrice  $\mathbb{B}$  en fonction de ceux de la matrice  $\mathbb{A}$ .

6

**Exercice A.3.2**

7

Soient  $\mathbb{A}$  et  $\mathbb{B}$  deux matrices triangulaires supérieures de  $\mathcal{M}_n$ . Soient  $\mathbb{E}$  et  $\mathbb{F}$  deux matrices triangulaires inférieures de  $\mathcal{M}_n$ .

**Q. 1** 1. Que peut-on dire des matrices  $\mathbb{A}^*$  et  $(\mathbb{A}^*)^*$  ?

2. Montrer que  $\mathbb{C} = \mathbb{A}\mathbb{B}$  est triangulaire supérieure et que  $C_{i,i} = A_{i,i}B_{i,i}$ ,  $\forall i \in \llbracket 1, n \rrbracket$ .

3. Montrer que  $\mathbb{G} = \mathbb{E}\mathbb{F}$  est triangulaire inférieure et que  $G_{i,i} = E_{i,i}F_{i,i}$ ,  $\forall i \in \llbracket 1, n \rrbracket$ .

4. Que peut-on dire des matrices  $\mathbb{A}\mathbb{E}$  et  $\mathbb{E}\mathbb{A}$  ?

**Q. 2** 1. Calculer  $\det(\mathbb{A})$ .

2. Déterminer les valeurs propres de  $\mathbb{A}$ .

3. Que peut-on dire si les éléments diagonaux de  $\mathbb{A}$  sont tous distincts ?

**Q. 3** Soit  $\mathbb{D}$  la matrice définie par

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

1. La matrice  $\mathbb{D}$  est-elle inversible ? Si oui calculer son inverse.

2. Pour chacune des valeurs propres, déterminer l'espace propre associé.

3. La matrice  $\mathbb{D}$  est-elle diagonalisable ? Justifier.

1



### Exercice A.3.3

**Q. 1** Soit  $\mathbb{T} \in \mathcal{M}_{n,n}(\mathbb{C})$  une matrice triangulaire supérieure. Montrer que si  $\mathbb{T}$  est une matrice normale alors elle est diagonale.

**Q. 2** Montrer que  $\mathbb{A} \in \mathcal{M}_{n,n}(\mathbb{C})$  est une matrice normale si et seulement si il existe  $\mathbb{U} \in \mathcal{M}_{n,n}(\mathbb{C})$  unitaire et  $\mathbb{D} \in \mathcal{M}_{n,n}(\mathbb{C})$  diagonale telle que  $\mathbb{A} = \mathbb{U}\mathbb{D}\mathbb{U}^*$ .

**Q. 3** En déduire qu'une matrice normale est diagonalisable et que ses vecteurs propres sont orthogonaux.

2



### Exercice A.3.4

Soit  $\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$  une matrice hermitienne

**Q. 1** Montrer que

$$\langle \mathbb{A}\mathbf{u}, \mathbf{u} \rangle \in \mathbb{R}, \quad \forall \mathbf{u} \in \mathbb{C}^n. \quad (\text{A.49})$$

On suppose de plus que la matrice  $\mathbb{A}$  est définie positive.

**Q. 2** 1. Montrer que les éléments diagonaux de  $\mathbb{A}$  sont strictement positifs.

2. Montrer que les sous matrices principales de  $\mathbb{A}$  sont elles aussi hermitiennes et définies positives.

3



### Exercice A.3.5: Procédé de Gram-Schmidt

4

Soit  $\{\mathbf{v}_i\}_{i \in \llbracket 1, n \rrbracket}$  une base de  $\mathbb{K}^n$ . On construit successivement les vecteurs  $\mathbf{u}_i$

$$\mathbf{u}_i = \mathbf{v}_i - \sum_{k=1}^{i-1} \frac{\langle \mathbf{u}_k, \mathbf{v}_i \rangle}{\langle \mathbf{u}_k, \mathbf{u}_k \rangle} \mathbf{u}_k, \quad \forall i \in \llbracket 1, n \rrbracket.$$

Montrer qu'ils forment une **base orthogonale** de  $\mathbb{K}^n$  et que  $\text{Vect}(\mathbf{u}_1, \dots, \mathbf{u}_i) = \text{Vect}(\mathbf{v}_1, \dots, \mathbf{v}_i)$ ,  $\forall i \in \llbracket 1, n \rrbracket$ .

**Correction Exercice A.3.5** Montrons par récurrence sur  $i$  que

$$(\mathcal{H})_i : \text{Vect}(\mathbf{u}_1, \dots, \mathbf{u}_i) \text{ est une famille orthogonale et } \text{Vect}(\mathbf{u}_1, \dots, \mathbf{u}_i) = \text{Vect}(\mathbf{v}_1, \dots, \mathbf{v}_i)$$

**Initialisation :** Pour  $i = 1$ , on a  $\mathbf{u}_1 = \mathbf{v}_1$  et  $(\mathcal{H})_1$  est vérifiée.

**Hérédité :** Soit  $i < n$ . Supposons  $(\mathcal{H})_i$  vérifiée. Montrons alors que  $(\mathcal{H})_{i+1}$  est vraie.

On a

$$\mathbf{u}_{i+1} = \mathbf{v}_{i+1} - \sum_{k=1}^i \frac{\langle \mathbf{u}_k, \mathbf{v}_{i+1} \rangle}{\langle \mathbf{u}_k, \mathbf{u}_k \rangle} \mathbf{u}_k. \quad (\text{A.50})$$

- En effectuant le produit scalaire de (A.50) par  $\mathbf{u}_j$  avec  $j \in \llbracket 1, i \rrbracket$  on obtient

$$\langle \mathbf{u}_j, \mathbf{u}_{i+1} \rangle = \langle \mathbf{u}_j, \mathbf{v}_{i+1} \rangle - \sum_{k=1}^i \frac{\langle \mathbf{u}_k, \mathbf{v}_{i+1} \rangle}{\langle \mathbf{u}_k, \mathbf{u}_k \rangle} \langle \mathbf{u}_j, \mathbf{u}_k \rangle.$$

Par hypothèse de récurrence, la famille  $\text{Vect}(\mathbf{u}_1, \dots, \mathbf{u}_i)$  est orthogonale, c'est à dire  $\forall (r, s) \in \llbracket 1, i \rrbracket^2$ ,  $\langle \mathbf{u}_r, \mathbf{u}_s \rangle = 0$  si  $r \neq s$  et  $\mathbf{u}_r \neq 0$ . On obtient donc

$$\langle \mathbf{u}_j, \mathbf{u}_{i+1} \rangle = \langle \mathbf{u}_j, \mathbf{v}_{i+1} \rangle - \frac{\langle \mathbf{u}_j, \mathbf{v}_{i+1} \rangle}{\langle \mathbf{u}_j, \mathbf{u}_j \rangle} \langle \mathbf{u}_j, \mathbf{u}_j \rangle = 0, \quad \forall j \in \llbracket 1, i \rrbracket.$$

- On montre maintenant par l'absurde que  $\mathbf{u}_{i+1} \neq 0$ .  
Supposons  $\mathbf{u}_{i+1} = 0$ . Alors de (A.50), on obtient

$$\mathbf{v}_{i+1} = \sum_{k=1}^i \frac{\langle \mathbf{u}_k, \mathbf{v}_{i+1} \rangle}{\langle \mathbf{u}_k, \mathbf{u}_k \rangle} \mathbf{u}_k$$

et donc  $\mathbf{v}_{i+1} \in \text{Vect}(\mathbf{u}_1, \dots, \mathbf{u}_i) \stackrel{(\mathcal{H})_i}{=} \text{Vect}(\mathbf{v}_1, \dots, \mathbf{v}_i)$ . Ceci entre en contradiction avec  $\text{Vect}(\mathbf{v}_1, \dots, \mathbf{v}_n)$  base de  $\mathbb{K}^n$ .

- On déduit de (A.50) que  $\mathbf{u}_{i+1} \in \text{Vect}(\mathbf{u}_1, \dots, \mathbf{u}_i, \mathbf{v}_{i+1})$ . Par hypothèse de récurrence,  $\text{Vect}(\mathbf{u}_1, \dots, \mathbf{u}_i) = \text{Vect}(\mathbf{v}_1, \dots, \mathbf{v}_i)$ , ce qui donne  $\mathbf{u}_{i+1} \in \text{Vect}(\mathbf{v}_1, \dots, \mathbf{v}_{i+1})$  et donc

$$\text{Vect}(\mathbf{u}_1, \dots, \mathbf{u}_{i+1}) = \text{Vect}(\mathbf{v}_1, \dots, \mathbf{v}_{i+1}).$$

◇



### Exercice A.3.6: factorisation $\mathbb{QR}$

Soit  $A \in \mathcal{M}_n(\mathbb{C})$  une matrice inversible. Pour tout  $i \in \llbracket 1, n \rrbracket$ , on note  $\mathbf{a}_i = A_{:,i}$  ses  $n$  vecteurs colonnes. En utilisant le procédé de Gram-Schmidt sur la base  $\{\mathbf{a}_1, \dots, \mathbf{a}_n\}$  montrer qu'il existe une matrice  $Q$  unitaire et une matrice triangulaire supérieure  $R$  à coefficients diagonaux strictement positifs tel que  $A = QR$ .

**Correction Exercice A.3.6** On utilise le procédé d'orthonormalisation de Gram-Schmidt (voir Proposition A.12, page 176) pour obtenir la base orthogonale  $\{\mathbf{u}_1, \dots, \mathbf{u}_n\}$  en calculant successivement

$$\mathbf{u}_i = \mathbf{a}_i - \sum_{k=1}^{i-1} \frac{\langle \mathbf{u}_k, \mathbf{a}_i \rangle}{\langle \mathbf{u}_k, \mathbf{u}_k \rangle} \mathbf{u}_k, \quad \forall i \in \llbracket 1, n \rrbracket. \quad (\text{A.51})$$

- 1 De plus on a  $\text{Vect}(\mathbf{u}_1, \dots, \mathbf{u}_i) = \text{Vect}(\mathbf{a}_1, \dots, \mathbf{a}_i), \forall i \in \llbracket 1, n \rrbracket$  On normalise la base orthogonale  $\{\mathbf{u}_1, \dots, \mathbf{u}_n\}$   
 2 pour obtenir la base orthonormée  $\mathcal{Q} = \{\mathbf{q}_1, \dots, \mathbf{q}_n\}$  :

$$\mathbf{q}_i = \frac{\mathbf{u}_i}{\|\mathbf{u}_i\|_2}, \forall i \in \llbracket 1, n \rrbracket$$

- 3 et l'on a aussi  $\text{Vect}(\mathbf{q}_1, \dots, \mathbf{q}_i) = \text{Vect}(\mathbf{a}_1, \dots, \mathbf{a}_i), \forall i \in \llbracket 1, n \rrbracket$ .  
 4 On note  $\mathbb{Q} \in \mathcal{M}_n(\mathbb{K})$  la matrice définie par

$$\mathbb{Q} = \left( \begin{array}{c|c|c} \mathbf{q}_1 & \cdots & \mathbf{q}_n \end{array} \right)$$

- 5 Cette matrice est clairement unitaire puisque la base  $\{\mathbf{q}_1, \dots, \mathbf{q}_n\}$  est orthonormée.  
 6 Montrons que  $\mathbb{Q}^* \mathbb{A}$  est triangulaire supérieure. On a

$$\mathbb{Q}^* \mathbb{A} = \begin{pmatrix} \mathbf{q}_1^* \\ \vdots \\ \mathbf{q}_n^* \end{pmatrix} \begin{pmatrix} \mathbf{a}_1 & \cdots & \mathbf{a}_n \end{pmatrix} = \begin{pmatrix} \langle \mathbf{q}_1, \mathbf{a}_1 \rangle & \cdots & \langle \mathbf{q}_1, \mathbf{a}_n \rangle \\ \vdots & \ddots & \vdots \\ \langle \mathbf{q}_n, \mathbf{a}_1 \rangle & \cdots & \langle \mathbf{q}_n, \mathbf{a}_n \rangle \end{pmatrix}$$

- 7 c'est à dire  $(\mathbb{Q}^* \mathbb{A})_{i,j} = \langle \mathbf{q}_i, \mathbf{a}_j \rangle, \forall (i, j) \in \llbracket 1, n \rrbracket$ . Par définition cette matrice est triangulaire supérieure si  
 8  $(\mathbb{Q}^* \mathbb{A})_{i,j} = 0$  pour tout  $i > j$ . Soit  $i \in \llbracket 1, n-1 \rrbracket$ . La base  $\mathcal{Q}$  étant orthonormée, on a  $\mathbf{q}_i \perp \text{Vect}(\mathbf{q}_1, \dots, \mathbf{q}_{i-1})$ .  
 9 Comme  $\text{Vect}(\mathbf{q}_1, \dots, \mathbf{q}_{i-1}) = \text{Vect}(\mathbf{a}_1, \dots, \mathbf{a}_{i-1})$ , on en déduit que

$$\langle \mathbf{q}_i, \mathbf{a}_j \rangle = 0, \forall j \in \llbracket 1, i-1 \rrbracket.$$

- 10 La matrice  $\mathbb{Q}^* \mathbb{A}$  est donc triangulaire supérieure.  
 11 De plus, on a

$$(\mathbb{Q}^* \mathbb{A})_{i,i} = \langle \mathbf{q}_i, \mathbf{a}_i \rangle = \frac{\langle \mathbf{u}_i, \mathbf{a}_i \rangle}{\|\mathbf{u}_i\|_2}$$

En prenant le produit scalaire de (A.51) avec  $\mathbf{u}_i$  on obtient

$$\begin{aligned} \langle \mathbf{u}_i, \mathbf{u}_i \rangle &= \langle \mathbf{u}_i, \mathbf{a}_i \rangle - \sum_{k=1}^{i-1} \frac{\langle \mathbf{u}_k, \mathbf{a}_i \rangle}{\langle \mathbf{u}_k, \mathbf{u}_k \rangle} \langle \mathbf{u}_i, \mathbf{u}_k \rangle \\ &= \langle \mathbf{u}_i, \mathbf{a}_i \rangle \quad \text{car } \langle \mathbf{u}_i, \mathbf{u}_k \rangle = 0, \forall k \neq i \text{ (base orthogonale)} \end{aligned}$$

- 12 Comme  $\mathbf{u}_i \neq 0$ , on obtient  $(\mathbb{Q}^* \mathbb{A})_{i,i} > 0$ .  
 13 On note  $\mathbb{R} = \mathbb{Q}^* \mathbb{A}$  cette matrice triangulaire supérieure avec  $R_{i,i} > 0, \forall i \in \llbracket 1, n \rrbracket$ . La matrice  $\mathbb{Q}$  étant  
 14 unitaire (i.e.  $\mathbb{Q} \mathbb{Q}^* = \mathbb{I}$ ) alors  $\mathbb{A} = \mathbb{Q} \mathbb{R}$ .  $\diamond$

15

### 16 Exercice: 4.1.1

Soit  $\mathbb{A} \in \mathcal{M}_{n,n}(\mathbb{C})$  une matrice et  $(\lambda, \mathbf{u})$  un élément propre de  $\mathbb{A}$  avec  $\|\mathbf{u}\|_2 = 1$ .

**Q. 1** En s'aidant de la base canonique  $\{\mathbf{e}_1, \dots, \mathbf{e}_n\}$ , construire une base orthonormée  $\{\mathbf{x}_1, \dots, \mathbf{x}_n\}$  telle que  $\mathbf{x}_1 = \mathbf{u}$ .

Notons  $\mathbb{P}$  la matrice de changement de base canonique  $\{\mathbf{e}_1, \dots, \mathbf{e}_n\}$  dans la base  $\{\mathbf{x}_1, \dots, \mathbf{x}_n\}$  :

$$\mathbb{P} = \left( \begin{array}{c|c|c} \mathbf{x}_1 & \dots & \mathbf{x}_n \end{array} \right)$$

Soit  $\mathbb{B}$  la matrice définie par  $\mathbb{B} = \mathbb{P}^* \mathbb{A} \mathbb{P}$ .

**Q. 2** 1. Exprimer les coefficients de la matrice  $\mathbb{B}$  en fonction de la matrice  $\mathbb{A}$  et des vecteurs  $\mathbf{x}_i$ ,  $i \in \llbracket 1, n \rrbracket$ .

$$\mathbb{B} = \mathbb{P}^* \mathbb{A} \mathbb{P}.$$

2. En déduire que la première colonne de  $\mathbb{B}$  est  $(\lambda, 0, \dots, 0)^t$ .

**Q. 3** Montrer par récurrence sur l'ordre de la matrice que la matrice  $\mathbb{A}$  s'écrit

$$\mathbb{A} = \mathbb{U} \mathbb{T} \mathbb{U}^*$$

où  $\mathbb{U}$  est une matrice unitaire et  $\mathbb{T}$  une matrice triangulaire supérieure.

**Q. 4** En supposant  $\mathbb{A}$  inversible et la décomposition  $\mathbb{A} = \mathbb{U} \mathbb{T} \mathbb{U}^*$  connue, expliquer comment résoudre "simplement" le système linéaire  $\mathbb{A} \mathbf{x} = \mathbf{b}$ .

### Correction Exercice 4.1.1

**Q. 1** La première chose à faire est de construire une base contenant  $\mathbf{u}$  à partir de la base canonique  $\{\mathbf{e}_1, \dots, \mathbf{e}_n\}$ . Comme le vecteur propre  $\mathbf{u}$  est non nul, il existe  $j \in \llbracket 1, n \rrbracket$  tel que  $\langle \mathbf{u}, \mathbf{e}_j \rangle \neq 0$ . La famille  $\{\mathbf{u}, \mathbf{e}_1, \dots, \mathbf{e}_{j-1}, \mathbf{e}_{j+1}, \dots, \mathbf{e}_n\}$  forme alors une base de  $\mathbb{C}^n$  car  $\mathbf{u}$  n'est pas combinaison linéaire des  $\{\mathbf{e}_1, \dots, \mathbf{e}_{j-1}, \mathbf{e}_{j+1}, \dots, \mathbf{e}_n\}$ .

On note  $\{\mathbf{z}_1, \dots, \mathbf{z}_n\}$  la base dont le premier élément est  $\mathbf{z}_1 = \mathbf{u}$  :

$$\{\mathbf{z}_1, \dots, \mathbf{z}_n\} = \{\mathbf{u}, \mathbf{e}_1, \dots, \mathbf{e}_{j-1}, \mathbf{e}_{j+1}, \dots, \mathbf{e}_n\}.$$

On peut ensuite utiliser le **procédé de Gram-Schmidt**, rappelé en Proposition A.12, pour construire une base orthonormée à partir de cette base.

On calcule successivement les vecteurs  $\mathbf{x}_i$  à partir de la base  $\{\mathbf{z}_1, \dots, \mathbf{z}_n\}$  en construisant un vecteur  $\mathbf{w}_i$  orthogonal aux vecteurs  $\mathbf{x}_1, \dots, \mathbf{x}_{i-1}$ .

$$\mathbf{w}_i = \mathbf{z}_i - \sum_{k=1}^{i-1} \langle \mathbf{x}_k, \mathbf{z}_i \rangle \mathbf{x}_k$$

puis on obtient le vecteur  $\mathbf{x}_i$  en normalisant

$$\mathbf{x}_i = \frac{\mathbf{w}_i}{\|\mathbf{w}_i\|}$$

**Q. 2** 1. En conservant l'écriture colonne de la matrice  $\mathbb{P}$  on obtient

$$\mathbb{B} = \begin{pmatrix} \mathbf{x}_1^* \\ \mathbf{x}_2^* \\ \vdots \\ \mathbf{x}_n^* \end{pmatrix} \mathbb{A} \begin{pmatrix} \mathbf{x}_1 & \mathbf{x}_2 & \dots & \mathbf{x}_n \end{pmatrix} = \begin{pmatrix} \mathbf{x}_1^* \\ \mathbf{x}_2^* \\ \vdots \\ \mathbf{x}_n^* \end{pmatrix} \begin{pmatrix} \mathbb{A} \mathbf{x}_1 & \mathbb{A} \mathbf{x}_2 & \dots & \mathbb{A} \mathbf{x}_n \end{pmatrix}$$

Ce qui donne

$$\mathbb{B} = \begin{pmatrix} \mathbf{x}_1^* \mathbb{A} \mathbf{x}_1 & \mathbf{x}_1^* \mathbb{A} \mathbf{x}_2 & \dots & \mathbf{x}_1^* \mathbb{A} \mathbf{x}_n \\ \mathbf{x}_2^* \mathbb{A} \mathbf{x}_1 & \mathbf{x}_2^* \mathbb{A} \mathbf{x}_2 & \dots & \mathbf{x}_2^* \mathbb{A} \mathbf{x}_n \\ \vdots & \vdots & \ddots & \vdots \\ \mathbf{x}_n^* \mathbb{A} \mathbf{x}_1 & \mathbf{x}_n^* \mathbb{A} \mathbf{x}_2 & \dots & \mathbf{x}_n^* \mathbb{A} \mathbf{x}_n \end{pmatrix}$$

1 On a donc

$$B_{i,j} = \mathbf{x}_i^* \mathbb{A} \mathbf{x}_j, \quad \forall (i, j) \in \llbracket 1, n \rrbracket^2$$

2 .

2. On a  $\mathbb{A} \mathbf{u} = \lambda \mathbf{u}$ ,  $\|\mathbf{u}\| = 1$ , la base  $\{\mathbf{x}_1, \dots, \mathbf{x}_n\}$  est orthonormée et  $\mathbf{x}_1 = \mathbf{u}$ . on obtient alors

$$\mathbb{B} = \left( \begin{array}{c|ccc} \lambda \mathbf{u}^* \mathbf{u} & \mathbf{u}^* \mathbb{A} \mathbf{x}_2 & \dots & \mathbf{u}^* \mathbb{A} \mathbf{x}_n \\ \hline \lambda \mathbf{x}_2^* \mathbf{u} & \mathbf{x}_2^* \mathbb{A} \mathbf{x}_2 & \dots & \mathbf{x}_2^* \mathbb{A} \mathbf{x}_n \\ \vdots & \vdots & \ddots & \vdots \\ \hline \lambda \mathbf{x}_n^* \mathbf{u} & \mathbf{x}_n^* \mathbb{A} \mathbf{x}_2 & \dots & \mathbf{x}_n^* \mathbb{A} \mathbf{x}_n \end{array} \right) = \left( \begin{array}{c|ccc} \lambda & \mathbf{u}^* \mathbb{A} \mathbf{x}_2 & \dots & \mathbf{u}^* \mathbb{A} \mathbf{x}_n \\ \hline 0 & \mathbf{x}_2^* \mathbb{A} \mathbf{x}_2 & \dots & \mathbf{x}_2^* \mathbb{A} \mathbf{x}_n \\ \vdots & \vdots & \ddots & \vdots \\ \hline 0 & \mathbf{x}_n^* \mathbb{A} \mathbf{x}_2 & \dots & \mathbf{x}_n^* \mathbb{A} \mathbf{x}_n \end{array} \right)$$

3 **Q. 3** On veut démontrer, par récurrence faible, la proposition suivante pour  $n \geq 2$

$(\mathcal{P}_n) \quad \forall \mathbb{A} \in \mathcal{M}_n(\mathbb{C}), \exists \mathbb{U} \in \mathcal{M}_n(\mathbb{C})$  unitaire,  $\exists \mathbb{T} \in \mathcal{M}_n(\mathbb{C})$  triangulaire supérieure, telles que  $\mathbb{A} = \mathbb{U} \mathbb{T} \mathbb{U}^*$ .

4 **Initialisation :** Montrons que  $(\mathcal{P}_2)$  est vérifié.

5 Soit  $\mathbb{A}_2 \in \mathcal{M}_2(\mathbb{C})$ . Elle admet au moins un élément propre  $(\lambda, \mathbf{u})$  (voir Proposition A.37 par ex.)  
 6 avec  $\|\mathbf{u}\| = 1$ . On peut donc appliquer le résultat de la question précédente : il existe une matrice  
 7 unitaire  $\mathbb{P}_2 \in \mathcal{M}_2(\mathbb{C})$  telle que la matrice  $\mathbb{B}_2 = \mathbb{P}_2 \mathbb{A}_2 \mathbb{P}_2^*$  ait comme premier vecteur colonne  $(\lambda, 0)^t$ .  
 8 La matrice  $\mathbb{B}_2$  est donc triangulaire supérieure et comme  $\mathbb{P}_2$  est unitaire on en déduit

$$\mathbb{A}_2 = \mathbb{P}_2^* \mathbb{B}_2 \mathbb{P}_2.$$

9 On pose  $\mathbb{U}_2 = \mathbb{P}_2^*$  matrice unitaire et  $\mathbb{T}_2 = \mathbb{B}_2$  matrice triangulaire supérieure pour conclure que la  
 10 proposition  $(\mathcal{P}_2)$  est vraie.

11 **Hérédité :** Supposons que  $(\mathcal{P}_{n-1})$  soit vérifiée. Montrons que  $(\mathcal{P}_n)$  est vraie.

12 Soit  $\mathbb{A}_n \in \mathcal{M}_n(\mathbb{C})$ . Elle admet au moins un élément propre  $(\lambda, \mathbf{u})$  (voir Proposition A.37 par ex.)  
 13 avec  $\|\mathbf{u}\| = 1$ . On peut donc appliquer le résultat de la question précédente : il existe une matrice  
 14 unitaire  $\mathbb{P}_n \in \mathcal{M}_n(\mathbb{C})$  telle que la matrice  $\mathbb{B}_n = \mathbb{P}_n \mathbb{A}_n \mathbb{P}_n^*$  s'écrive

$$\mathbb{B}_n = \left( \begin{array}{c|ccc} \lambda & & & \mathbf{c}_{n-1}^* \\ \hline 0 & & & \\ \vdots & & \mathbb{A}_{n-1} & \\ \hline 0 & & & \end{array} \right)$$

15 où  $\mathbf{c}_{n-1} \in \mathcal{M}_{n-1}(\mathbb{C})$  et  $\mathbb{A}_{n-1} \in \mathcal{M}_{n-1}(\mathbb{C})$ . Par hypothèse de récurrence,  $\exists \mathbb{U}_{n-1} \in \mathcal{M}_{n-1}(\mathbb{C})$  unitaire  
 16 et  $\mathbb{T}_{n-1} \in \mathcal{M}_{n-1}(\mathbb{C})$  triangulaire supérieure telles que

$$\mathbb{A}_{n-1} = \mathbb{U}_{n-1} \mathbb{T}_{n-1} \mathbb{U}_{n-1}^*$$

17 ou encore

$$\mathbb{T}_{n-1} = \mathbb{U}_{n-1}^* \mathbb{A}_{n-1} \mathbb{U}_{n-1}.$$

18 Soit  $\mathbb{Q}_n \in \mathcal{M}_n(\mathbb{C})$  la matrice définie par

$$\mathbb{Q}_n = \left( \begin{array}{c|ccc} 1 & 0 & \dots & 0 \\ \hline 0 & & & \\ \vdots & & \mathbb{U}_{n-1} & \\ \hline 0 & & & \end{array} \right).$$

19 La matrice  $\mathbb{Q}_n$  est unitaire. En effet on a

$$\mathbb{Q}_n \mathbb{Q}_n^* = \left( \begin{array}{c|ccc} 1 & 0 & \dots & 0 \\ \hline 0 & & & \\ \vdots & & \mathbb{U}_{n-1} & \\ \hline 0 & & & \end{array} \right) \left( \begin{array}{c|ccc} 1 & 0 & \dots & 0 \\ \hline 0 & & & \\ \vdots & & \mathbb{U}_{n-1}^* & \\ \hline 0 & & & \end{array} \right) = \left( \begin{array}{c|ccc} 1 & 0 & \dots & 0 \\ \hline 0 & & \underbrace{\mathbb{U}_{n-1} \mathbb{U}_{n-1}^*}_{=\mathbb{I}_{n-1}} & \\ \vdots & & & \\ \hline 0 & & & \end{array} \right) = \mathbb{I}_n.$$

On note  $\mathbb{T}_n$  la matrice définie par  $\mathbb{T}_n = \mathbb{Q}_n^* \mathbb{B}_n \mathbb{Q}_n$ . On a alors

$$\begin{aligned} \mathbb{T}_n &= \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & & & \\ \vdots & & \mathbb{U}_{n-1}^* & \\ 0 & & & \end{pmatrix} \begin{pmatrix} \lambda & \mathbf{c}_{n-1}^* \\ 0 & \\ \vdots & \mathbb{A}_{n-1} \\ 0 & \end{pmatrix} \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & & & \\ \vdots & & \mathbb{U}_{n-1} & \\ 0 & & & \end{pmatrix} \\ &= \begin{pmatrix} \lambda & \mathbf{c}_{n-1}^* \\ 0 & \\ \vdots & \mathbb{U}_{n-1}^* \mathbb{A}_{n-1} \\ 0 & \end{pmatrix} \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & & & \\ \vdots & & \mathbb{U}_{n-1} & \\ 0 & & & \end{pmatrix} = \begin{pmatrix} \lambda & \mathbf{c}_{n-1}^* \mathbb{U}_{n-1}^* \\ 0 & \\ \vdots & \underbrace{\mathbb{U}_{n-1}^* \mathbb{A}_{n-1} \mathbb{U}_{n-1}}_{=\mathbb{T}_{n-1}} \\ 0 & \end{pmatrix} \end{aligned}$$

La matrice  $\mathbb{T}_n$  est donc triangulaire supérieure et on a par définition de  $\mathbb{B}_n$

$$\mathbb{T}_n = \mathbb{Q}_n^* \mathbb{P}_n \mathbb{A}_n \mathbb{P}_n^* \mathbb{Q}_n.$$

On note  $\mathbb{U}_n = \mathbb{P}_n^* \mathbb{Q}_n$ . Cette matrice est unitaire car les matrices  $\mathbb{Q}_n$  et  $\mathbb{P}_n$  le sont. En effet, on a

$$\mathbb{U}_n \mathbb{U}_n^* = \mathbb{P}_n^* \mathbb{Q}_n (\mathbb{P}_n^* \mathbb{Q}_n)^* = \mathbb{P}_n^* \underbrace{\mathbb{Q}_n \mathbb{Q}_n^*}_{=\mathbb{I}_n} \mathbb{P}_n = \mathbb{P}_n^* \mathbb{P}_n = \mathbb{I}_n.$$

On a  $\mathbb{T}_n = \mathbb{U}_n^* \mathbb{A}_n \mathbb{U}_n$  et en multipliant cette équation à gauche par  $\mathbb{U}_n$  et à droite par  $\mathbb{U}_n^*$  on obtient l'équation équivalente  $\mathbb{A}_n = \mathbb{U}_n \mathbb{T}_n \mathbb{U}_n^*$ . La propriété  $(\mathcal{P}_n)$  est donc vérifiée. Ce qui achève la démonstration.

**Q. 4** Résoudre  $\mathbb{A}\mathbf{x} = \mathbf{b}$  est équivalent à résoudre

$$\mathbb{U}\mathbb{T}\mathbb{U}^*\mathbf{x} = \mathbf{b}. \quad (\text{A.52})$$

Comme  $\mathbb{U}$  est unitaire, on a  $\mathbb{U}\mathbb{U}^* = \mathbb{I}$  et  $\mathbb{U}^*$  inversible. Donc en multipliant (A.52) par  $\mathbb{U}^*$  on obtient le système équivalent

$$\underbrace{\mathbb{U}^* \mathbb{U}}_{=\mathbb{I}} \mathbb{T} \mathbb{U}^* \mathbf{x} = \mathbb{U}^* \mathbf{b} \iff \mathbb{T} \mathbb{U}^* \mathbf{x} = \mathbb{U}^* \mathbf{b}.$$

On pose  $\mathbf{y} = \mathbb{U}^* \mathbf{x}$ . Le système précédant se résout en deux étapes

1. on cherche  $\mathbf{y}$  solution de  $\mathbb{T}\mathbf{y} = \mathbb{U}^* \mathbf{b}$ . Comme  $\mathbb{U}$  est unitaire on a  $\det(\mathbb{U}) \det(\mathbb{U}^*) = \det(\mathbb{I}) = 1$  et donc

$$\begin{aligned} \det(\mathbb{A}) &= \det(\mathbb{U}\mathbb{T}\mathbb{U}^*) = \det(\mathbb{U}) \det(\mathbb{T}) \det(\mathbb{U}^*) \\ &= \det(\mathbb{T}) \end{aligned}$$

Or  $\mathbb{A}$  inversible équivaut à  $\det(\mathbb{A}) \neq 0$  et donc la matrice  $\mathbb{T}$  est inversible. La matrice  $\mathbb{T}$  étant triangulaire inférieure on peut résoudre facilement le système par la *méthode de remontée*.

2. une fois  $\mathbf{y}$  déterminé, on résout  $\mathbb{U}^* \mathbf{x} = \mathbf{y}$ . Comme  $\mathbb{U}$  est unitaire, on obtient directement  $\mathbf{x} = \mathbb{U}\mathbf{y}$ .

◇

**Inverse d'une matrice**



### Exercice A.3.7

Soit  $A \in \mathcal{M}_3(\mathbb{R})$  définie par

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

**Q. 1** Calculer le déterminant de la matrice  $A$ . Que peut-on en conclure?

**Q. 2** Calculer si possible l'inverse de la matrice  $A$  en utilisant la technique de la matrice augmentée.



### Exercice A.3.8

Soient  $A$  et  $B$ , deux matrices de  $\mathcal{M}_n(\mathbb{K})$ .

**Q. 1** Montrer que

$$AB = I \Rightarrow BA = I \quad (\text{A.53})$$

Conclure.



### Exercice A.3.9

**Q. 1** Soit  $A$  une matrice inversible et symétrique, montrer que  $A^{-1}$  est symétrique.

**Q. 2** Soit  $A$  une matrice carrée telle que  $I - A$  est inversible. Montrer que

$$A(I - A)^{-1} = (I - A)^{-1}A.$$

**Q. 3** Soient  $A, B$  des matrices carrées inversibles de même dimension telle que  $A + B$  soit inversible. Montrer que

$$A(A + B)^{-1}B = B(A + B)^{-1}A = (A^{-1} + B^{-1})^{-1}$$



### Exercice A.3.10

Soit  $L \in \mathcal{M}_n(\mathbb{C})$  une matrice triangulaire inférieure.

**Q. 1** A quelle(s) condition(s) la matrice  $L$  est-elle inversible?

On suppose  $L$  inversible et on note  $X = L^{-1}$ .

**Q. 2** Montrer que  $X$  est une matrice triangulaire inférieure avec

$$X_{i,i} = \frac{1}{L_{i,i}}, \quad \forall i \in \llbracket 1, n \rrbracket.$$

### Correction Exercice A.3.10

**Q. 1** La matrice  $L$  est inversible si et seulement si son déterminant est non nul. Or le déterminant d'une matrice triangulaire est égal au produit de ses éléments diagonaux. Pour avoir  $L$ , matrice triangulaire, inversible, il est nécessaire et suffisant d'avoir

$$L_{ii} \neq 0, \quad \forall i \in \llbracket 1, n \rrbracket.$$

**Q. 2** La matrice  $X$  étant la matrice inverse de  $L$ , on a

$$LX = I \quad (\text{A.54})$$

On note  $X^{[j]} = X_{:,j}$  le  $j$ -ème vecteur colonne de la matrice  $X$  et  $e^{[j]}$  le  $j$ -ème vecteur de la base canonique de  $\mathbb{C}^n$  ( $e_i^{[j]} = \delta_{i,j}$ ).

L'équation (A.54) peut donc se réécrire

$$\mathbb{L} \left( \begin{array}{c|c|c} \mathbf{X}^{[1]} & \cdots & \mathbf{X}^{[n]} \end{array} \right) = \left( \begin{array}{c|c|c} \mathbf{e}^{[1]} & \cdots & \mathbf{e}^{[n]} \end{array} \right)$$

ou encore, déterminer la matrice  $\mathbb{X}$  inverse de  $\mathbb{L}$  revient à résoudre les  $n$  systèmes linéaires suivants:

$$\mathbb{L}\mathbf{X}^{[j]} = \mathbf{e}^{[j]}, \quad \forall j \in \llbracket 1, n \rrbracket. \quad (\text{A.55})$$

- Pour montrer que  $\mathbb{X}$  est triangulaire inférieure il suffit de vérifier que pour tout  $j \in \llbracket 2, n \rrbracket$

$$X_i^{[j]} = 0, \quad \forall i \in \llbracket 1, j-1 \rrbracket.$$

Soit  $j \in \llbracket 2, n \rrbracket$ . On décompose la matrice  $\mathbb{L}$  en la matrice bloc carré 2 par 2 ou le premier bloc diagonal, noté  $\mathbb{L}_{j-1}$ , est une matrice triangulaire inférieure inversible de dimension  $j-1$ . Le système (A.55) s'écrit alors

$$\mathbb{L}\mathbf{X}^{[j]} = \mathbf{e}^{[j]} \iff \left( \begin{array}{c|cccc} \xrightarrow{j-1} & L_{1,1} & 0 & \cdots & 0 \\ \bullet & \ddots & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & \ddots & \vdots \\ \bullet & \cdots & \bullet & L_{j-1,j-1} & 0 \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ \bullet & \cdots & \cdots & \cdots & L_{j,j} \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ \bullet & \cdots & \cdots & \cdots & L_{n,n} \end{array} \right) \left( \begin{array}{c} X_1^{[j]} \\ \vdots \\ X_{j-1}^{[j]} \\ X_j^{[j]} \\ \vdots \\ X_n^{[j]} \end{array} \right) = \left( \begin{array}{c} 0 \\ \vdots \\ 0 \\ 1 \\ 0 \\ \vdots \\ 0 \end{array} \right) \quad \xrightarrow{j-1}$$

On en déduit donc que nécessairement

$$\left( \begin{array}{c|cccc} L_{1,1} & 0 & \cdots & 0 \\ \bullet & \ddots & \ddots & \vdots \\ \vdots & \ddots & \ddots & \vdots \\ \bullet & \cdots & \bullet & L_{j-1,j-1} \end{array} \right) \left( \begin{array}{c} X_1^{[j]} \\ \vdots \\ X_{j-1}^{[j]} \end{array} \right) = \left( \begin{array}{c} 0 \\ \vdots \\ 0 \end{array} \right)$$

Comme la matrice de ce système est inversible, on a bien  $X_i^{[j]} = 0, \forall i \in \llbracket 1, j-1 \rrbracket$  et donc la matrice  $\mathbb{X}$  est triangulaire inférieure.

- Par définition du produit matricielle, de l'équation (A.54) on tire pour tout  $j \in \llbracket 1, n \rrbracket$

$$(\mathbb{L}\mathbb{X})_{j,j} = (\mathbb{I})_{j,j} \iff \sum_{k=1}^n L_{j,k} X_{k,j} = 1 \iff \sum_{k=1}^{j-1} L_{j,k} X_{k,j} + L_{j,j} X_{j,j} + \sum_{k=j+1}^n L_{j,k} X_{k,j} = 1$$

Or les matrices  $\mathbb{L}$  et  $\mathbb{X}$  sont triangulaires inférieures et donc  $X_{k,j} = 0, \forall k \in \llbracket 1, j-1 \rrbracket$ , et  $L_{j,k} = 0, \forall k \in \llbracket j+1, n \rrbracket$ . On obtient alors  $L_{j,j} X_{j,j} = 1$  et comme  $L_{j,j} \neq 0$  on a bien  $X_{j,j} = 1/L_{j,j}$ .

◇

### Exercice A.3.11

Soit  $A \in \mathcal{M}_{n,n}(\mathbb{K})$  et  $U, B, V$  trois matrices rectangulaires.

**Q. 1** Sous quelles hypothèses peut-on définir la matrice  $G$  suivante

$$G = A^{-1} - A^{-1}U(I + BVA^{-1}U)^{-1}BVA^{-1} \quad (\text{A.56})$$

**Q. 2** Montrer que  $(A + UBV)G = I$ . Conclure.

**Q. 3** Soit  $\beta \in \mathbb{R}$  et  $u, v \in \mathbb{C}^n$ . Calculer  $(A + \beta uv^t)^{-1}$ .



### Exercice A.3.12

Etant donnée une matrice  $D \in \mathcal{M}_{n,n}(\mathbb{C})$ , on pose

$$D = A + iB \text{ avec } A, B \in \mathcal{M}_{n,n}(\mathbb{R})$$

Sous certaines hypothèses à préciser, établir la relation

$$D^{-1} = (A + BA^{-1}B)^{-1} - iA^{-1}B(A + BA^{-1}B)^{-1}.$$

## 3 Matrices blocs



### Exercice A.3.13

On considère les matrices blocs suivantes

$$A = \left( \begin{array}{cc|cc} 1 & 2 & 1 & 0 \\ 3 & 4 & 0 & 1 \\ \hline 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \end{array} \right) = \left( \begin{array}{c|c} C & I \\ \hline I & 0 \end{array} \right) \text{ et } B = \left( \begin{array}{cc|cc} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ \hline 1 & 2 & 1 & 2 \\ 3 & 4 & 3 & 4 \end{array} \right) = \left( \begin{array}{c|c} I & 0 \\ \hline C & C \end{array} \right)$$

avec par identification

$$I = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \text{ et } C = \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix}$$

**Q. 1** Calculer les matrices  $AB$  et  $BA$  en utilisant l'écriture bloc.

**Q. 2** Exprimer les matrices  $A(A+B)$  et  $(2B-A)(B+A)$  en fonction des matrices  $C$  et  $I$ .



### Exercice A.3.14

Soient  $A \in \mathcal{M}_{n,k}(\mathbb{K})$  et  $B \in \mathcal{M}_{k,n}(\mathbb{K})$ . On note  $L$  la matrice

$$L = \left( \begin{array}{c|c} I - BA & B \\ \hline 2A - ABA & AB - I \end{array} \right).$$

**Q. 1** Montrer que la matrice  $L$  est bien définie et spécifier les dimensions des blocs.

**Q. 2** Calculer  $L^2$ . Que peut-on en conclure?



### Exercice A.3.15: résultats à savoir



Soient  $A \in \mathcal{M}_m(\mathbb{C})$ ,  $B \in \mathcal{M}_n(\mathbb{C})$  et  $D \in \mathcal{M}_{m,n}(\mathbb{C})$ .

**Q. 1** Calculer, en fonction des déterminant de  $A$  et  $B$ , le déterminant des matrices

$$E = \left( \begin{array}{c|c} A & 0 \\ \hline 0 & I_n \end{array} \right), \quad F = \left( \begin{array}{c|c} I_m & 0 \\ \hline 0 & B \end{array} \right), \quad \text{et } G = \left( \begin{array}{c|c} A & 0 \\ \hline 0 & B \end{array} \right).$$

**Q. 2** Soit  $H = \left( \begin{array}{c|c} A & D \\ \hline 0 & B \end{array} \right)$ . En utilisant les factorisations QR des matrices  $A$  et  $B$ , montrer que

$$\det(H) = \det(A) \det(B). \quad (\text{A.57})$$

**Q. 3** En déduire qu'une matrice triangulaire supérieure par blocs est inversible si et seulement si ses matrices blocs diagonales sont inversibles.

**Q. 4** En déduire qu'une matrice triangulaire inférieure par blocs est inversible si et seulement si ses matrices blocs diagonales sont inversibles.

### A.3.2 Normes

#### Normes vectorielles



#### Exercice A.3.16

Soient  $x$  et  $y$  deux vecteurs de  $\mathbb{C}^n$ .

**Q. 1** Trouver  $\alpha \in \mathbb{C}$  tel que  $\langle \alpha x - y, x \rangle = 0$ .

**Q. 2** En calculant  $\|\alpha x - y\|_2^2$ , montrer que

$$|\langle x, y \rangle| \leq \|x\|_2 \|y\|_2. \quad (\text{A.58})$$

**Q. 3** Soit  $x \neq 0$ . Montrer alors que l'inégalité (A.58) est une égalité si et seulement si  $y = \alpha x$ .

#### Correction Exercice A.3.16

**Q. 1** • Si  $x = 0$ , alors  $\alpha$  quelconque.

• Si  $x \neq 0$ , alors

$$\langle \alpha x - y, x \rangle = 0 \iff \alpha \langle x, x \rangle - \langle y, x \rangle = 0$$

Or  $x \neq 0$ , ce qui donne

$$\alpha = \frac{\langle y, x \rangle}{\langle x, x \rangle}. \quad (\text{A.59})$$

**Q. 2** On a

$$\begin{aligned} \|\alpha x - y\|_2^2 &= \langle \alpha x - y, \alpha x - y \rangle \\ &= \bar{\alpha} \langle \alpha x - y, x \rangle - \langle \alpha x - y, y \rangle \\ &= -\langle \alpha x - y, y \rangle, \quad \text{car } \langle \alpha x - y, x \rangle = 0 \\ &= -\alpha \langle x, y \rangle + \langle y, y \rangle \end{aligned}$$

En utilisant (A.59), on obtient alors

$$\begin{aligned} \|\alpha x - y\|_2^2 &= -\frac{\langle y, x \rangle}{\langle x, x \rangle} \langle x, y \rangle + \langle y, y \rangle \\ &= \frac{-\langle y, x \rangle \langle x, y \rangle + \langle y, y \rangle \langle x, x \rangle}{\langle x, x \rangle} \end{aligned}$$

1 Comme  $\langle \mathbf{y}, \mathbf{x} \rangle = \overline{\langle \mathbf{x}, \mathbf{y} \rangle}$ , on a  $\langle \mathbf{y}, \mathbf{x} \rangle \langle \mathbf{x}, \mathbf{y} \rangle = |\langle \mathbf{x}, \mathbf{y} \rangle|^2$  et donc

$$\begin{aligned} \|\alpha \mathbf{x} - \mathbf{y}\|_2^2 &= \frac{1}{\langle \mathbf{x}, \mathbf{x} \rangle} \left( -|\langle \mathbf{x}, \mathbf{y} \rangle|^2 + \|\mathbf{x}\|_2^2 \|\mathbf{y}\|_2^2 \right) \\ &\geq 0. \end{aligned} \quad (\text{A.60})$$

On a alors

$$|\langle \mathbf{x}, \mathbf{y} \rangle|^2 \leq \|\mathbf{x}\|_2^2 \|\mathbf{y}\|_2^2$$

2 La fonction  $x \mapsto \sqrt{x}$  étant croissante sur  $[0; +\infty[$ , on obtient (A.58).

**Q. 3** Soit  $\mathbf{x} \neq 0$ . On veut montrer que

$$|\langle \mathbf{x}, \mathbf{y} \rangle| = \|\mathbf{x}\|_2 \|\mathbf{y}\|_2 \iff \mathbf{y} = \alpha \mathbf{x}$$

$\Leftarrow$  On suppose  $\mathbf{y} = \alpha \mathbf{x}$ . On a alors

$$\langle \mathbf{x}, \mathbf{y} \rangle = \bar{\alpha} \langle \mathbf{x}, \mathbf{x} \rangle = \bar{\alpha} \|\mathbf{x}\|_2^2 \implies |\langle \mathbf{x}, \mathbf{y} \rangle| = |\alpha| \|\mathbf{x}\|_2^2.$$

Comme  $\|\mathbf{y}\|_2 = |\alpha| \|\mathbf{x}\|_2$ , on a aussi

$$\|\mathbf{x}\|_2 \|\mathbf{y}\|_2 = |\alpha| \|\mathbf{x}\|_2^2.$$

On en déduit alors

$$|\langle \mathbf{x}, \mathbf{y} \rangle| = \|\mathbf{x}\|_2 \|\mathbf{y}\|_2.$$

$\Rightarrow$  On suppose  $|\langle \mathbf{x}, \mathbf{y} \rangle| = \|\mathbf{x}\|_2 \|\mathbf{y}\|_2$ . Avec cette hypothèse, l'équation (A.60) devient

$$\|\alpha \mathbf{x} - \mathbf{y}\|_2^2 = 0$$

3 et donc  $\alpha \mathbf{x} - \mathbf{y} = 0$ , c'est à dire  $\mathbf{y} = \alpha \mathbf{x}$ .

4  $\diamond$

5



### Exercice A.3.17

Soient  $\mathbf{x}$  et  $\mathbf{y}$  deux vecteurs de  $\mathbb{C}^n$ .

**Q. 1** Démontrer l'inégalité triangulaire

$$\|\mathbf{x} + \mathbf{y}\|_2 \leq \|\mathbf{x}\|_2 + \|\mathbf{y}\|_2. \quad (\text{A.61})$$

**Q. 2** Si  $\mathbf{x}$  et  $\mathbf{y}$  sont non nuls, prouver que l'inégalité (A.61) est une égalité si et seulement si  $\mathbf{y} = \alpha \mathbf{x}$  avec  $\alpha$  un réel strictement positif.

**Q. 3** Dédurre de (A.61) l'inégalité suivante :

$$|\|\mathbf{x}\|_2 - \|\mathbf{y}\|_2| \leq \|\mathbf{x} - \mathbf{y}\|_2. \quad (\text{A.62})$$

**Q. 4** Soient  $\mathbf{x}_1, \dots, \mathbf{x}_p$ ,  $p$  vecteurs de  $\mathbb{C}^n$ . Montrer que

$$\left\| \sum_{i=1}^p \mathbf{x}_i \right\|_2 \leq \sum_{i=1}^p \|\mathbf{x}_i\|_2. \quad (\text{A.63})$$

6

7 **Correction Exercice A.3.17**

8 **Q. 1** On rappelle que  $\|\mathbf{x}\|_2^2 = \langle \mathbf{x}, \mathbf{x} \rangle$ . On obtient, en utilisant les propriétés du produit scalaire,

$$\begin{aligned} \|\mathbf{x} + \mathbf{y}\|_2^2 &= \langle \mathbf{x} + \mathbf{y}, \mathbf{x} + \mathbf{y} \rangle \\ &= \langle \mathbf{x}, \mathbf{x} \rangle + \langle \mathbf{x}, \mathbf{y} \rangle + \langle \mathbf{y}, \mathbf{x} \rangle + \langle \mathbf{y}, \mathbf{y} \rangle \\ &= \|\mathbf{x}\|_2^2 + \|\mathbf{y}\|_2^2 + \langle \mathbf{x}, \mathbf{y} \rangle + \langle \mathbf{y}, \mathbf{x} \rangle. \end{aligned}$$

Pour tout nombre complexe  $z$ , on a  $z + \bar{z} = 2\operatorname{Re}(z)$ , et  $|z| \geq \operatorname{Re}(z)$ .<sup>1</sup> Comme  $\langle \mathbf{y}, \mathbf{x} \rangle = \overline{\langle \mathbf{x}, \mathbf{y} \rangle}$ , on a

$$\begin{aligned} \|\mathbf{x} + \mathbf{y}\|_2^2 &= \|\mathbf{x}\|_2^2 + \|\mathbf{y}\|_2^2 + 2\operatorname{Re}(\langle \mathbf{x}, \mathbf{y} \rangle) \\ &\leq \|\mathbf{x}\|_2^2 + \|\mathbf{y}\|_2^2 + 2|\langle \mathbf{x}, \mathbf{y} \rangle|. \end{aligned} \quad (\text{A.64})$$

L'inégalité de Cauchy-Schwarz donne

$$|\langle \mathbf{x}, \mathbf{y} \rangle| \leq \|\mathbf{x}\|_2 \|\mathbf{y}\|_2$$

et donc

$$\begin{aligned} \|\mathbf{x} + \mathbf{y}\|_2^2 &\leq \|\mathbf{x}\|_2^2 + \|\mathbf{y}\|_2^2 + 2\|\mathbf{x}\|_2 \|\mathbf{y}\|_2 \\ &= (\|\mathbf{x}\|_2 + \|\mathbf{y}\|_2)^2. \end{aligned}$$

La fonction  $x \mapsto \sqrt{x}$  étant croissante sur  $[0; +\infty[$ , on obtient

$$\|\mathbf{x} + \mathbf{y}\|_2 \leq \|\mathbf{x}\|_2 + \|\mathbf{y}\|_2.$$

**Q. 2** Soient  $\mathbf{x}$  et  $\mathbf{y}$  deux vecteurs de  $\mathbb{C}^n$  non nuls. On veut démontrer que

$$\|\mathbf{x} + \mathbf{y}\|_2 = \|\mathbf{x}\|_2 + \|\mathbf{y}\|_2 \iff \mathbf{y} = \alpha \mathbf{x}, \alpha > 0.$$

$\Leftarrow$  On suppose  $\mathbf{y} = \alpha \mathbf{x}$  avec  $\alpha > 0$ .  
On a alors

$$\begin{aligned} \|\mathbf{x} + \mathbf{y}\|_2 &= \|\mathbf{x} + \alpha \mathbf{x}\|_2 \\ &= \|(1 + \alpha)\mathbf{x}\|_2 \\ &= |1 + \alpha| \|\mathbf{x}\|_2. \end{aligned}$$

Comme  $\alpha > 0$ , on a  $|1 + \alpha| = 1 + \alpha$  et donc

$$\|\mathbf{x} + \mathbf{y}\|_2 = \|\mathbf{x}\|_2 + \alpha \|\mathbf{x}\|_2. \quad (\text{A.65})$$

De plus, on a

$$\begin{aligned} \|\mathbf{y}\|_2 &= \|\alpha \mathbf{x}\|_2 \\ &= |\alpha| \|\mathbf{x}\|_2 \\ &= \alpha \|\mathbf{x}\|_2, \text{ car } \alpha > 0. \end{aligned}$$

D'après (A.65), on en déduit

$$\|\mathbf{x} + \mathbf{y}\|_2 = \|\mathbf{x}\|_2 + \|\mathbf{y}\|_2.$$

$\Rightarrow$  On suppose que

$$\|\mathbf{x} + \mathbf{y}\|_2 = \|\mathbf{x}\|_2 + \|\mathbf{y}\|_2. \quad (\text{A.66})$$

Ce qui donne

$$\|\mathbf{x} + \mathbf{y}\|_2^2 = \|\mathbf{x}\|_2^2 + \|\mathbf{y}\|_2^2 + 2\|\mathbf{x}\|_2 \|\mathbf{y}\|_2 \quad (\text{A.67})$$

On déduit alors de l'égalité (A.64) que

$$\operatorname{Re}(\langle \mathbf{x}, \mathbf{y} \rangle) = \|\mathbf{x}\|_2 \|\mathbf{y}\|_2. \quad (\text{A.68})$$

Or, on a,  $\forall z \in \mathbb{C}$ ,  $\operatorname{Re}(z) \leq |z|$  et donc, en utilisant l'inégalité de Cauchy-Schwarz

$$\|\mathbf{x}\|_2 \|\mathbf{y}\|_2 = \operatorname{Re}(\langle \mathbf{x}, \mathbf{y} \rangle) \leq |\langle \mathbf{x}, \mathbf{y} \rangle| \leq \|\mathbf{x}\|_2 \|\mathbf{y}\|_2$$

ce qui impose

$$|\langle \mathbf{x}, \mathbf{y} \rangle| = \|\mathbf{x}\|_2 \|\mathbf{y}\|_2.$$

D'après l'exercice précédant, cette égalité est vérifiée si et seulement si  $\mathbf{y} = \alpha \mathbf{x}$  avec  $\alpha = \frac{\langle \mathbf{y}, \mathbf{x} \rangle}{\langle \mathbf{x}, \mathbf{x} \rangle}$ .

Il nous reste à vérifier que  $\alpha > 0$ .

L'hypothèse (A.66) avec  $\mathbf{y} = \alpha \mathbf{x}$  devient

$$\|\mathbf{x} + \mathbf{y}\|_2 = |1 + \alpha| \|\mathbf{x}\|_2 = \|\mathbf{x}\|_2 + \|\mathbf{y}\|_2 = 1 + |\alpha| \|\mathbf{x}\|_2. \quad (\text{A.69})$$

<sup>1</sup> En effet,  $z = a + ib$  et  $\bar{z} = a - ib$  d'où  $z + \bar{z} = 2a$ . De plus,  $|z|^2 = a^2 + b^2 \geq a^2$  ce qui donne  $|z| \geq |a| \geq a$ .

1 On a donc

$$\begin{aligned}
 |1 + \alpha| = 1 + |\alpha| &\iff |1 + \alpha|^2 = (1 + |\alpha|)^2 \\
 &\iff (1 + \alpha)(1 + \bar{\alpha}) = 1 + 2|\alpha| + |\alpha|^2 \\
 &\iff (1 + \alpha)(1 + \bar{\alpha}) = 1 + 2|\alpha| + |\alpha|^2 \\
 &\iff 1 + \alpha\bar{\alpha} + \alpha + \bar{\alpha} = 1 + 2|\alpha| + |\alpha|^2 \\
 &\iff \operatorname{Re}(\alpha) = |\alpha|
 \end{aligned}$$

2 Donc  $\alpha$  est un réel et  $\alpha = \operatorname{Re}(\alpha) = |\alpha| \geq 0$ .

3 Comme  $\mathbf{y} = \alpha \mathbf{x}$  et  $\mathbf{y} \neq 0$ , on a  $\alpha \neq 0$  et donc  $\alpha > 0$ .

**Q. 3** On a  $\mathbf{x} = (\mathbf{x} - \mathbf{y}) + \mathbf{y}$ , et par application de l'inégalité triangulaire (A.61) on obtient

$$\|\mathbf{x}\|_2 = \|(\mathbf{x} - \mathbf{y}) + \mathbf{y}\|_2 \leq \|\mathbf{x} - \mathbf{y}\|_2 + \|\mathbf{y}\|_2 \implies \|\mathbf{x}\|_2 - \|\mathbf{y}\|_2 \leq \|\mathbf{x} - \mathbf{y}\|_2.$$

De même, avec  $\mathbf{y} = (\mathbf{y} - \mathbf{x}) + \mathbf{x}$ , et par application de l'inégalité triangulaire (A.61) on a

$$\|\mathbf{y}\|_2 = \|(\mathbf{y} - \mathbf{x}) + \mathbf{x}\|_2 \leq \|\mathbf{y} - \mathbf{x}\|_2 + \|\mathbf{x}\|_2 \implies \|\mathbf{y}\|_2 - \|\mathbf{x}\|_2 \leq \|\mathbf{x} - \mathbf{y}\|_2.$$

En combinant ces deux inégalités, on obtient alors

$$|\|\mathbf{x}\|_2 - \|\mathbf{y}\|_2| \leq \|\mathbf{x} - \mathbf{y}\|_2.$$

4 **Q. 4** On effectue une démonstration par récurrence.

5 Soit  $n \in \mathbb{N}$ ,  $n > 1$ . On définit la propriété  $\mathcal{P}(n)$  par

$$\mathcal{P}(n) : \left\| \sum_{i=1}^n \mathbf{x}_i \right\|_2 \leq \sum_{i=1}^n \|\mathbf{x}_i\|_2. \quad (\text{A.70})$$

6 • On a démontré, en Q.1, que la propriété  $\mathcal{P}(2)$  est vraie.

7 • Soit  $n > 2$ , on suppose que  $\mathcal{P}(n)$  est vérifiée (hypothèse de récurrence). On veut alors montrer que  
8  $\mathcal{P}(n+1)$  est vraie.

9 On a

$$\begin{aligned}
 \left\| \sum_{i=1}^{n+1} \mathbf{x}_i \right\|_2 &= \left\| \sum_{i=1}^n \mathbf{x}_i + \mathbf{x}_{n+1} \right\|_2 \\
 &\leq \left\| \sum_{i=1}^n \mathbf{x}_i \right\|_2 + \|\mathbf{x}_{n+1}\|_2, \text{ d'après (A.61)} \\
 &\leq \sum_{i=1}^n \|\mathbf{x}_i\|_2 + \|\mathbf{x}_{n+1}\|_2, \text{ car } \mathcal{P}(n) \text{ est vérifiée} \\
 &= \sum_{i=1}^{n+1} \|\mathbf{x}_i\|_2.
 \end{aligned}$$

10 La propriété  $\mathcal{P}(n+1)$  est donc vérifiée.

11 On a donc démontré par récurrence que la propriété  $\mathcal{P}(n)$  est vraie pour tout  $n \geq 2$ .

12

13

14



### Exercice A.3.18

**Q. 1** Soit la fonction  $f(t) = (1 - \lambda) + \lambda t - t^\lambda$  avec  $0 < \lambda < 1$ . Montrer que pour tous  $\alpha \geq 0$  et  $\beta \geq 0$  on a

$$\alpha^\lambda \beta^{1-\lambda} \leq \lambda \alpha + (1 - \lambda) \beta. \quad (\text{A.71})$$

Soient  $\mathbf{x}$  et  $\mathbf{y}$  deux vecteurs non nuls de  $\mathbb{C}^n$ . Soient  $p > 1$  et  $q > 1$  vérifiant  $\frac{1}{p} + \frac{1}{q} = 1$ .

**Q. 2** On pose  $\mathbf{u} = \frac{\mathbf{x}}{\|\mathbf{x}\|_p}$  et  $\mathbf{v} = \frac{\mathbf{y}}{\|\mathbf{y}\|_q}$ . En utilisant l'inégalité (A.71), montrer que l'on a l'inégalité de Hölder

$$\sum_{i=1}^n |u_i v_i| \leq \frac{1}{p} \sum_{i=1}^n |u_i|^p + \frac{1}{q} \sum_{i=1}^n |v_i|^q = 1. \quad (\text{A.72})$$

**Q. 3** En déduire l'inégalité de suivante

$$|\langle \mathbf{x}, \mathbf{y} \rangle| \leq \|\mathbf{x}\|_p \|\mathbf{y}\|_q. \quad (\text{A.73})$$

Quel est le lien entre l'inégalité de Hölder et l'inégalité de Cauchy-Schwarz?

### Correction Exercice A.3.18

**Q. 1** L'inégalité (A.71) est vérifiée si  $\alpha = 0$  ou  $\beta = 0$ . Il nous reste donc à la vérifier pour  $\alpha > 0$  et  $\beta > 0$ . Dans ce cas (A.71) s'écrit

$$\left(\frac{\alpha}{\beta}\right)^\lambda \leq \lambda \frac{\alpha}{\beta} + (1 - \lambda)$$

c'est à dire

$$f\left(\frac{\alpha}{\beta}\right) \geq 0.$$

Montrons que  $f(t) \geq 0, \forall t \in ]0, +\infty[$ .

On a  $f'(t) = \lambda(1 - t^{\lambda-1})$  et

$$\begin{aligned} f'(t) = 0 &\Leftrightarrow 1 - t^{\lambda-1} = 0 \\ &\Leftrightarrow t = 1 \text{ car } -1 < \lambda - 1 < 0 \end{aligned}$$

De plus, on a  $t^{\lambda-1} = e^{(\lambda-1)\ln(t)}$ .

- Etudions la fonction sur  $]0, 1[$ . On a pour  $t \in ]0, 1[$ ,  $\ln(t) < 0$  et donc  $(\lambda - 1)\ln(t) > 0$ . Comme la fonction exp est croissante, on en déduit  $\exp((\lambda - 1)\ln(t)) > 1$  et alors  $f'(t) < 0$ .

- Etudions la fonction sur  $]1, +\infty[$ . On a pour  $t \in ]1, +\infty[$ ,  $\ln(t) > 0$  et donc  $(\lambda - 1)\ln(t) < 0$ . Comme la fonction exp est croissante, on en déduit  $0 < \exp((\lambda - 1)\ln(t)) < 1$  et alors  $f'(t) > 0$ .

Le minimum de  $f$  est donc atteint en  $t = 1$  et on a

$$\forall t \in ]0, +\infty[, f(t) \geq f(1) = 0.$$

L'inégalité (A.71) est donc vérifiée  $\forall \alpha \geq 0, \forall \beta \geq 0$  et  $\forall \lambda \in ]0, 1[$ .

**Q. 2** On pose  $\lambda = \frac{1}{p} \in ]0, 1[$ . on a alors  $1 - \lambda = \frac{1}{q}$ . On pose

$$\alpha = |u_i|^p \geq 0, \quad \beta = |v_i|^q \geq 0.$$

En utilisant (A.71), on obtient directement

$$|u_i| |v_i| \leq \frac{1}{p} |u_i|^p + \frac{1}{q} |v_i|^q, \quad \forall i \in \llbracket 1, n \rrbracket.$$

En sommant sur  $i$  on obtient:

$$\sum_{i=1}^n |u_i v_i| \leq \frac{1}{p} \sum_{i=1}^n |u_i|^p + \frac{1}{q} \sum_{i=1}^n |v_i|^q = \frac{1}{p} \|\mathbf{u}\|_p^p + \frac{1}{q} \|\mathbf{v}\|_q^q$$

Comme par construction  $\|\mathbf{u}\|_p = \|\mathbf{v}\|_q = 1$ , on obtient

$$\sum_{i=1}^n |u_i v_i| \leq \frac{1}{p} + \frac{1}{q} = 1.$$

**Q. 3** Par construction, on a

$$\sum_{i=1}^n |u_i v_i| = \frac{1}{\|\mathbf{x}\|_p \|\mathbf{y}\|_q} \sum_{i=1}^n |x_i y_i|$$

et donc en utilisant l'inégalité (A.73) on obtient

$$\sum_{i=1}^n |x_i y_i| \leq \|\mathbf{x}\|_p \|\mathbf{y}\|_q.$$

De plus

$$|\langle \mathbf{x}, \mathbf{y} \rangle| = \left| \sum_{i=1}^n \bar{x}_i y_i \right| \leq \sum_{i=1}^n |\bar{x}_i y_i| = \sum_{i=1}^n |x_i y_i| \leq \|\mathbf{x}\|_p \|\mathbf{y}\|_q.$$

1 Pour  $p = q = 2$ , l'inégalité de Hölder entraîne l'inégalité de Cauchy-Schwarz.

2

3

### Exercice A.3.19

Soit  $p > 1$  et  $q$  le nombre tel que  $\frac{1}{q} = 1 - \frac{1}{p}$ .

**Q. 1** Vérifier que  $\forall (\alpha, \beta) \in \mathbb{C}^2$  on a

$$|\alpha + \beta|^p \leq |\alpha| |\alpha + \beta|^{p/q} + |\beta| |\alpha + \beta|^{p/q}. \quad (\text{A.74})$$

**Q. 2** En utilisant l'inégalité de Hölder et (A.74), démontrer l'inégalité de Minkowski :

$$\|\mathbf{x} + \mathbf{y}\|_p \leq \|\mathbf{x}\|_p + \|\mathbf{y}\|_p, \quad \forall \mathbf{x} \in \mathbb{C}^n, \quad \forall \mathbf{y} \in \mathbb{C}^n, \quad p \geq 1. \quad (\text{A.75})$$

4

## 5 Normes matricielles

### Exercice A.3.20: Norme de Frobenius

Soit  $\mathbb{A} \in \mathcal{M}_n(\mathbb{C})$ . On définit l'application  $\|\bullet\|_F$  par

$$\|\mathbb{A}\|_F^2 = \sum_{i=1}^n \sum_{j=1}^n |a_{ij}|^2. \quad (\text{A.76})$$

**Q. 1** On note, respectivement,  $\mathbf{A}_{i,:}$ ,  $\forall i \in \llbracket 1, n \rrbracket$  et  $\mathbf{A}_{:,j}$ ,  $\forall j \in \llbracket 1, n \rrbracket$  les vecteurs lignes et colonnes de  $\mathbb{A}$ . Montrer que

$$\|\mathbb{A}\|_F^2 = \sum_{i=1}^n \|\mathbf{A}_{i,:}\|_2^2 = \sum_{j=1}^n \|\mathbf{A}_{:,j}\|_2^2 = \text{tr } \mathbb{A}^* \mathbb{A}. \quad (\text{A.77})$$

**Q. 2** Montrer que

$$\|\mathbb{A} \mathbf{x}\|_2 \leq \|\mathbb{A}\|_F \|\mathbf{x}\|_2, \quad \forall \mathbf{x} \in \mathbb{C}^n. \quad (\text{A.78})$$

**Q. 3** Montrer que cette application est une norme matricielle (nommée norme de Frobenius).

**Q. 4** Calculer  $\|\mathbb{A}^*\|_F$  et  $\|\mathbb{I}_n\|_F$  où  $\mathbb{I}_n$  est la matrice identité de  $\mathcal{M}_n(\mathbb{C})$ .

6

### Exercice A.3.21

7

Soit  $A \in \mathcal{M}_{m,n}(\mathbb{C})$ . Montrer les propriétés suivantes

1.  $\|A\|_2 = \max_{\substack{\|x\|_2=1 \\ \|y\|_2=1}} |\langle Ax, y \rangle|.$
2.  $\|A\|_2 = \|A^*\|_2.$
3.  $\|A^*A\|_2 = \|A\|_2^2.$
4.  $\|U^*AV\|_2 = \|A\|_2$  quand  $UU^* = I$  et  $VV^* = I$

1



### Exercice A.3.22

Soient  $A \in \mathcal{M}_{m,n}(\mathbb{C})$ ,  $B \in \mathcal{M}_{n,l}(\mathbb{C})$ ,  $x \in \mathbb{C}^n$  et  $p \geq 1$ . Montrer que

$$\|Ax\|_p \leq \|A\|_p \|x\|_p \quad (\text{A.79})$$

$$\|A\|_p = \max_{\|x\|_p=1} \|Ax\|_p \quad (\text{A.80})$$

$$\|A\|_p = \max_{\|x\|_p \leq 1} \|Ax\|_p \quad (\text{A.81})$$

$$\|AB\|_p \leq \|A\|_p \|B\|_p \quad (\text{A.82})$$

2



### Exercice A.3.23

Soit  $A \in \mathcal{M}_n(\mathbb{C})$ . Montrer que l'on a

$$\|A\|_1 = \max_{j \in \llbracket 1, n \rrbracket} \sum_{i=1}^n |a_{ij}|, \quad (\text{A.83})$$

$$\|A\|_2 = \rho(AA^*)^{1/2}, \quad (\text{A.84})$$

$$\|A\|_\infty = \max_{i \in \llbracket 1, n \rrbracket} \sum_{j=1}^n |a_{ij}|. \quad (\text{A.85})$$

3

## A.4 Listings

### A.4.1 Codes sur la méthode de dichotomie/bissection

Transcription de l'Algorithme 3.1 (page 28)

Listing A.1: fonction dichotomie1 (Matlab)

```

1 function x=dichotomie1(f,a,b,epsilon)
2 kmin=floor(log((b-a)/epsilon)/log(2));
3 X=zeros(kmin+1,1);
4 A=zeros(kmin+1,1); B=zeros(kmin+1,1);
5 A(1)=a; B(1)=b; X(1)=(a+b)/2;
6 for k=1:kmin
7 if f(X(k))==0
8 A(k+1)=X(k); B(k+1)=X(k);
9 elseif f(B(k))*f(X(k))<0
10 A(k+1)=X(k); B(k+1)=B(k);
11 else
12 A(k+1)=A(k); B(k+1)=X(k);
13 end
14 X(k+1)=(A(k+1)+B(k+1))/2;
15 end
16 x=X(kmin+1);
17 end

```

Listing A.2: fonction dichotomie1 (C)

```

double dichotomie1(double (*f)(double),
double a, double b, double eps){
double *A,*B,*X,x;
int kmin,k;
size_t Size;
assert(f(a)*f(b)<0);
kmin=(int) floor(log((b-a)/eps)/log(2.));
Size=(kmin+1)*sizeof(double);
assert(X=(double*) malloc(Size));
assert(A=(double*) malloc(Size));
assert(B=(double*) malloc(Size));
// ou assert(A && B && X);
A[0]=a; B[0]=b; X[0]=(a+b)/2.;
for (k=0;k<kmin;k++){
if (f(X[k])==0){
A[k+1]=X[k]; B[k+1]=X[k];
} else if (f(B[k])*f(X[k])<0){
A[k+1]=X[k]; B[k+1]=B[k];
} else {
A[k+1]=A[k]; B[k+1]=X[k];
}
X[k+1]=(A[k+1]+B[k+1])/2;
}
x=X[kmin];
free(X); free(A); free(B);
return x;
}

```

Listing A.4: main dichotomie1 (C)

```

#include <stdio.h>
#include <stdlib.h>
#include <math.h>
#include <assert.h>

double dichotomie1(
double (*f)(double),
double a, double b, double eps
);
double g1(double x){
return (x+2)*(x+2)*(x-M_PI);
}

int main(){
double x;
x=dichotomie1(g1,-1,2*M_PI,1e-8);
printf("x=%.16lf, error=%.6e\n",
x, fabs(x-M_PI));
x=dichotomie1(cos,-1,M_PI,1e-8);
printf("x=%.16lf, error=%.6e\n",
x, fabs(x-M_PI_2));
return 1;
}
// Definition de dichotomie1 ensuite ...

```

Listing A.3: script dichotomie1 (Matlab)

```

1 clear all
2 close all
3 f=@(x) (x+2)*(x+2)*(x-pi);
4 x=dichotomie1(f,-1,2*pi,1e-8);
5 fprintf('x=%.16f\n',x)
6 x=dichotomie1(@cos,2,pi,1e-8);
7 fprintf('x=%.16f\n',x)

```

## Transcription de l'Algorithme 3.5 (page 30)

1

Listing A.5: fonction dichotomie5 (Matlab)

```

1 function x=dichotomie5(f,a,b)
2 assert(f(a)*f(b)<0, ...
3 'test_f(a)*f(b)<0_failed');
4 A=a;B=b;x=(A+B)/2;xp=A;
5 while x~=xp
6 if f(B)*f(x)<0
7 A=x;
8 else
9 B=x;
10 end
11 xp=x;
12 x=(A+B)/2;
13 end
14 end

```

Listing A.7: script dichotomie5 (Matlab)

```

1 clear all
2 close all
3 f=@(x) (x+2)*(x+2)*(x-pi);
4 x=dichotomie5(f,-1,2*pi);
5 fprintf('x=%.16f\n',x)
6 x=dichotomie5(@cos,2,pi);
7 fprintf('x=%.16f\n',x)

```

Listing A.6: fonction dichotomie5 (C)

```

double dichotomie5(double (*f)(double),
 double a, double b){
1
2
3 double A,B,x,xp;
4 assert(f(a)*f(b)<0);
5 A=a;B=b;x=(A+B)/2.;xp=A;
6 while (x!=xp){
7 if (f(B)*f(x)<0)
8 A=x;
9 else
10 B=x;
11 xp=x;
12 x=(A+B)/2;
13 }
14 return x;
15 }

```

Listing A.8: main dichotomie5 (C)

```

#include <stdio.h>
#include <math.h>
#include <assert.h>
1
2
3
4
5 double dichotomie5(double (*f)(double),
6 double a, double b);
7
8 double g1(double x){
9 return (x+2)*(x+2)*(x-M_PI);
10 }
11
12 int main(){
13 double x;
14 x=dichotomie5(g1,-1,2*M_PI);
15 printf("x=%.16f\n",x);
16 x=dichotomie5(cos,-1,M_PI);
17 printf("x=%.16f\n",x);

```



## Liste des algorithmes

|      |                                                                                                                                                    |    |    |
|------|----------------------------------------------------------------------------------------------------------------------------------------------------|----|----|
| 1.1  | Algorithme de calcul de $\pi$ , version naïve . . . . .                                                                                            | 2  | 2  |
| 1.2  | Algorithme de calcul de $\pi$ , version stable . . . . .                                                                                           | 9  | 3  |
| 2.1  | Exemple de boucle «pour» . . . . .                                                                                                                 | 15 | 4  |
| 2.2  | Exemple de boucle «tant que» . . . . .                                                                                                             | 16 | 5  |
| 2.3  | Exemple de boucle «répéter ...jusqu'à» . . . . .                                                                                                   | 16 | 6  |
| 2.4  | Exemple d'instructions conditionnelle «si» . . . . .                                                                                               | 16 | 7  |
| 2.5  | Exemple de fonction : Résolution de l'équation du premier degré $ax + b = 0$ . . . . .                                                             | 17 | 8  |
| 2.6  | Calcul de $S = \sum_{k=1}^n k \sin(2kx)$ . . . . .                                                                                                 | 19 | 9  |
| 2.7  | Calcul de $P = \prod_{n=1}^k \sin(2kz/n)^k$ . . . . .                                                                                              | 20 | 10 |
| 2.8  | En-tête de la fonction SFT retournant valeur de la série de Fourier en $t$ tronquée au $n$ premiers termes de l'exercice 2.2.3. . . . .            | 20 | 11 |
| 2.9  | Fonction SFT retournant la valeur de la série de Fourier en $t$ tronquée au $n$ premiers termes de l'exercice ?? . . . . .                         | 21 | 12 |
| 3.1  | Méthode de dichotomie : version 1 . . . . .                                                                                                        | 28 | 13 |
| 3.2  | Méthode de dichotomie : version 2 . . . . .                                                                                                        | 29 | 14 |
| 3.3  | Méthode de dichotomie : version 3 . . . . .                                                                                                        | 29 | 15 |
| 3.4  | Méthode de dichotomie : version 4 . . . . .                                                                                                        | 30 | 16 |
| 3.5  | Méthode de dichotomie : version 5 . . . . .                                                                                                        | 30 | 17 |
| 3.6  | Méthode de point fixe : version <b>Tantque</b> <i>formel</i> . . . . .                                                                             | 39 | 18 |
| 3.7  | Méthode de point fixe : version <b>Répéter</b> <i>formel</i> . . . . .                                                                             | 39 | 19 |
| 3.8  | Méthode de point fixe : version <b>Tantque</b> <i>formel</i> avec critères d'arrêt . . . . .                                                       | 40 | 20 |
| 3.9  | Méthode de point fixe : version <b>Répéter</b> <i>formel</i> avec critères d'arrêt . . . . .                                                       | 40 | 21 |
| 3.10 | Méthode de point fixe : version <b>Tantque</b> avec critères d'arrêt . . . . .                                                                     | 40 | 22 |
| 3.11 | Méthode de point fixe : version <b>Répéter</b> avec critères d'arrêt . . . . .                                                                     | 40 | 23 |
| 3.12 | Méthode de la corde . . . . .                                                                                                                      | 44 | 24 |
| 3.13 | Méthode de la corde . . . . .                                                                                                                      | 44 | 25 |
| 3.14 | Méthode de Newton . . . . .                                                                                                                        | 50 | 26 |
| 3.15 | Méthode de Newton scalaire . . . . .                                                                                                               | 50 | 27 |
| 3.16 | Méthode de Newton . . . . .                                                                                                                        | 60 | 28 |
| 4.1  | Fonction <b>RSLMATDIAG</b> permettant de résoudre le système linéaire à matrice diagonale inversible $\mathbb{A}\mathbf{x} = \mathbf{b}$ . . . . . | 74 | 29 |
|      | . . . . .                                                                                                                                          |    | 30 |
|      | . . . . .                                                                                                                                          |    | 31 |
|      | . . . . .                                                                                                                                          |    | 32 |
|      | . . . . .                                                                                                                                          |    | 33 |

|    |      |                                                                                                                                                                                          |                                         |     |
|----|------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------|-----|
| 1  | 4.2  | Fonction <b>RSLTRIINF</b> permettant de résoudre le système linéaire triangulaire inférieur inversible                                                                                   | $\mathbb{A}\mathbf{x} = \mathbf{b}$ .   |     |
| 2  |      |                                                                                                                                                                                          |                                         |     |
| 3  |      |                                                                                                                                                                                          |                                         | 76  |
| 4  | 4.3  | Fonction <b>RSLTRISUP</b> permettant de résoudre le système linéaire triangulaire supérieur inversible                                                                                   | $\mathbb{A}\mathbf{x} = \mathbf{b}$ .   |     |
| 5  |      |                                                                                                                                                                                          |                                         |     |
| 6  |      |                                                                                                                                                                                          |                                         | 78  |
| 7  | 4.4  | Algorithme de Gauss-Jordan formel pour la résolution de $\mathbb{A}\mathbf{x} = \mathbf{b}$                                                                                              |                                         | 79  |
| 8  | 4.5  | Algorithme de Gauss-Jordan avec fonctions pour la résolution de $\mathbb{A}\mathbf{x} = \mathbf{b}$                                                                                      |                                         | 79  |
| 9  | 4.6  | Recherche d'un pivot pour l'algorithme de Gauss-Jordan.                                                                                                                                  |                                         | 79  |
| 10 | 4.7  | Permutte deux lignes d'une matrice et d'un vecteur.                                                                                                                                      |                                         | 79  |
| 11 | 4.8  | Combinaison linéaire $\mathcal{L}_i \leftarrow \mathcal{L}_i + \alpha \mathcal{L}_j$ appliqué à une matrice et à un vecteur.                                                             |                                         | 79  |
| 12 | 4.9  | Fonction <b>RSLFactLU</b> permettant de résoudre, par une factorisation $\mathbb{L}\mathbb{U}$ , le système linéaire                                                                     | $\mathbb{A}\mathbf{x} = \mathbf{b}$     |     |
| 13 |      | où $\mathbb{A}$ une matrice de $\mathcal{M}_n(\mathbb{R})$ définie positive et $\mathbf{b} \in \mathbb{R}^n$ .                                                                           |                                         | 86  |
| 14 | 4.10 | Fonction <b>FACTLU</b> permet de calculer les matrices $\mathbb{L}$ et $\mathbb{U}$ dites matrice de factorisation $\mathbb{L}\mathbb{U}$ associée à la matrice $\mathbb{A}$ , telle que | $\mathbb{A} = \mathbb{L}\mathbb{U}$     |     |
| 15 |      |                                                                                                                                                                                          |                                         |     |
| 16 | 4.11 | Fonction <b>FACTLUliGU</b> permet de calculer la ligne $i$ de $\mathbb{U}$ à partir de (4.23)                                                                                            |                                         | 90  |
| 17 | 4.12 | Fonction <b>FACTLUcoll</b> permet de calculer la colonne $i$ de $\mathbb{L}$ à partir de (4.24)                                                                                          |                                         | 90  |
| 18 | 4.13 | Fonction <b>FACTLU</b> permet de calculer les matrices $\mathbb{L}$ et $\mathbb{U}$ dites matrice de factorisation $\mathbb{L}\mathbb{U}$ associée à la matrice $\mathbb{A}$ , telle que | $\mathbb{A} = \mathbb{L}\mathbb{U}$     |     |
| 19 |      | en utilisant des fonctions intermédiaires.                                                                                                                                               |                                         | 91  |
| 20 | 4.14 | Algorithme de base permettant de résoudre, par une factorisation de Cholesky positive, le système linéaire                                                                               | $\mathbb{A}\mathbf{x} = \mathbf{b}$     |     |
| 21 |      | où $\mathbb{A}$ une matrice de $\mathcal{M}_n(\mathbb{C})$ hermitienne définie positive et $\mathbf{b} \in \mathbb{C}^n$ .                                                               |                                         | 94  |
| 22 | 4.15 | Fonction <b>RSLCHOLESKY</b> permettant de résoudre, par une factorisation de Cholesky positive, le système linéaire                                                                      | $\mathbb{A}\mathbf{x} = \mathbf{b}$     |     |
| 23 |      | où $\mathbb{A}$ une matrice hermitienne de $\mathcal{M}_n(\mathbb{C})$ définie positive et $\mathbf{b} \in \mathbb{C}^n$ .                                                               |                                         | 94  |
| 24 | 4.16 | Fonction <b>CHOLESKY</b> permettant de calculer la matrice $\mathbb{B}$ , dites matrice de factorisation positive de Cholesky associée à la matrice $\mathbb{A}$ , telle que             | $\mathbb{A} = \mathbb{B}\mathbb{B}^*$ . |     |
| 25 | 4.17 | Calcul du $\alpha$ et de la matrice de Householder $\mathbb{H}(\mathbf{u})$ telle que $\mathbb{H}(\mathbf{u})\mathbf{a} = \alpha\mathbf{b}$ .                                            |                                         | 100 |
| 26 | 4.18 | Fonction <b>FACTQR</b>                                                                                                                                                                   |                                         | 105 |
| 27 | 4.19 | Méthode itérative pour la résolution d'un système linéaire $\mathbb{A}\mathbf{x} = \mathbf{b}$                                                                                           |                                         | 120 |
| 28 | 4.20 | Méthode itérative de Jacobi pour la résolution d'un système linéaire $\mathbb{A}\mathbf{x} = \mathbf{b}$                                                                                 |                                         | 122 |
| 29 | 4.21 | Méthode itérative de Gauss-Seidel pour la résolution d'un système linéaire $\mathbb{A}\mathbf{x} = \mathbf{b}$                                                                           |                                         | 124 |
| 30 | 4.22 | Itération de Jacobi : calcul de $\mathbf{x}$ tel que                                                                                                                                     |                                         |     |
| 31 |      | $x_i = \frac{1}{A_{ii}} \left( b_i - \sum_{j=1, j \neq i}^n A_{ij} y_j \right), \quad \forall i \in \llbracket 1, n \rrbracket.$                                                         |                                         | 125 |
| 32 | 4.23 | Itération de Gauss-Seidel : calcul de $\mathbf{x}$ tel que                                                                                                                               |                                         |     |
| 33 |      | $x_i = \frac{1}{A_{ii}} \left( b_i - \sum_{j=1}^{i-1} A_{i,j} x_j - \sum_{j=i+1}^n A_{i,j} y_j \right), \quad \forall i \in \llbracket 1, n \rrbracket.$                                 |                                         | 125 |
| 34 | 4.24 | Méthode itérative pour la résolution d'un système linéaire $\mathbb{A}\mathbf{x} = \mathbf{b}$                                                                                           |                                         | 126 |
| 35 | 4.25 | Itération S.O.R. : calcul de $\mathbf{x}$ tel que                                                                                                                                        |                                         |     |
| 36 |      | $x_i = \frac{w}{A_{ii}} \left( b_i - \sum_{j=1}^{i-1} A_{ij} x_j - \sum_{j=i+1}^n A_{ij} y_j \right) + (1-w)x_i$                                                                         |                                         | 127 |
| 37 | 5.1  | Fonction <b>LAGRANGE</b> permettant de calculer le polynôme d'interpolation de Lagrange $\mathcal{P}_n(x)$ définit par (5.5)                                                             |                                         | 133 |
| 38 | 5.2  | Fonction <b>HERMITE</b> permettant de calculer le polynôme d'interpolation de Lagrange-Hermite $H_n(t)$ définit par (5.27)                                                               |                                         | 146 |
| 39 | 5.3  | Fonction <b>POLYA</b> permettant de calculer le polynôme $A_i$ en $t \in \mathbb{R}$ donné par $A_i(t) = (1 - 2L'_i(x_i)(t - x_i))L_i^2(t)$                                              |                                         | 147 |
| 40 | 5.4  | Fonction <b>POLYB</b> permettant de calculer le polynôme $B_i$ en $t \in \mathbb{R}$ donné par $B_i(t) = (t - x_i)L_i^2(t)$                                                              |                                         | 147 |
| 41 | 5.5  | Fonction <b>POLYL</b> permettant de calculer le polynôme $L_i$ en $t \in \mathbb{R}$ donné par $L_i(t) = \prod_{j=0, j \neq i}^n \frac{t - x_j}{x_i - x_j}$                              |                                         | 147 |
| 42 |      |                                                                                                                                                                                          |                                         |     |
| 43 |      |                                                                                                                                                                                          |                                         |     |
| 44 |      |                                                                                                                                                                                          |                                         |     |
| 45 |      |                                                                                                                                                                                          |                                         |     |
| 46 |      |                                                                                                                                                                                          |                                         |     |
| 47 |      |                                                                                                                                                                                          |                                         |     |
| 48 |      |                                                                                                                                                                                          |                                         |     |
| 49 |      |                                                                                                                                                                                          |                                         |     |
| 50 |      |                                                                                                                                                                                          |                                         |     |
| 51 |      |                                                                                                                                                                                          |                                         |     |

5.6 Fonction **POLYLP** permettant de calculer  $L'_i(x_i) = \sum_{k=0, k \neq i}^n \frac{1}{x_i - x_k}$  . . . . . 147 1



## Bibliography

- [1] P.G. Ciarlet. *Introduction à l'analyse numérique matricielle et à l'optimisation*. DUNOD, 2006. 2
- [2] M. Crouzeix and A.L. Mignot. *Analyse numérique des équations différentielles*. Mathématiques appliquées pour la maîtrise. Masson, 1992. 3  
4
- [3] J.-P. Demailly. *Analyse numérique et équations différentielles*. Grenoble Sciences. EDP Sciences, 2006. 5  
6
- [4] J.P. Demailly. *Analyse Numérique et Equations Différentielles*. PUG, 1994. 7
- [5] W. Gander, M.J. Gander, and F. Kwok. *Scientific computing : an introduction using Maple and MATLAB*. Springer, Cham, 2014. 8  
9
- [6] T. Huckle. Collection of software bugs [http://www.zenger.informatik.tu-](http://www.zenger.informatik.tu-muenchen.de/persons/huckle/bugse.html) 10  
[muenchen.de/persons/huckle/bugse.html](http://www.zenger.informatik.tu-muenchen.de/persons/huckle/bugse.html). 11
- [7] P. Lascaux and R. Théodor. *Analyse numérique matricielle appliquée à l'art de l'ingénieur*. Number 12  
vol. 1 et 2 in *Analyse numérique matricielle appliquée à l'art de l'ingénieur*. Dunod, 2004. 13