



REPUBLIQUE ALGERIENNE DEMOCRATIQUE ET POPULAIRE
MINISTRE DE L'ENSEIGNEMENT SUPERIEUR ET DE LA
RECHERCHE SCIENTIFIQUE

Université Mohamed Boudiaf de M'sila

Faculté des Mathématiques et de l'Informatique
Département de Mathématiques



Mémoire de Master

Domaine: Mathématiques et Informatique

Filière: Mathématiques

Option: Analyse Mathématiques et numérique

Thème

Étude Comparative entre la Méthode de Collocation et la
Méthode de Galerkin pour la Résolution des Équations Intégrales

Présenté par:

SADI Omayma Nor El Yaqine

Soutenu publiquement le: 07/06/2026.

Devant le jury composé de:

SEGHIRI Fakhreddine	M.C.B	Université de M'sila	Président
LAKHAL Aissa	M.C.B	Université de M'sila	Encadreur
GAGUI Bachir	M.C.A	Université de M'sila	Examineur

Année universitaire 2025/2026

Remerciements

Avant tout, je remercie Dieu le Tout-Puissant de m'avoir donné la force, la patience et le courage pour accomplir ce travail.

Je tiens à remercier le Dr AISSA LAKHAL, directeur de ce mémoire, pour sa disponibilité et ses conseils judicieux tout au long de ce travail.

Je remercie également les membres du jury pour l'honneur qu'ils me font en acceptant d'évaluer ce travail.

Mes remerciements s'adressent aussi à tous mes enseignants qui ont contribué à ma formation universitaire et à l'enrichissement de mes connaissances.

Enfin, j'exprime ma profonde gratitude à mes parents, à ma famille ainsi qu'à mes amis pour leur soutien moral, leurs encouragements et leur confiance durant tout mon parcours.

Merci à toutes les personnes qui ont participé de près ou de loin à l'aboutissement de ce travail.

Dédicaces

À moi-même...

À celle qui a patienté, qui s'est fatiguée et qui a veillé jusqu'à arriver.

Merci de ne pas avoir abandonné malgré tout.

Ce succès est à toi, en premier.

À mon père...(Larbi)

cet homme exceptionnel qui a toujours été mon soutien, ma force et ma fierté.

À celui qui s'est fatigué pour mon bonheur,

qui m'a appris que la réussite se construit avec patience et courage.

Aucun mot ne pourra exprimer ma reconnaissance et l'amour immense que je te porte.

Que Dieu te protège et te garde toujours près de moi.

À ma mère ...(Badar Leila)

source infinie de tendresse, de sacrifices et de prières.

À celle qui a porté mes inquiétudes dans son cœur avant même que je ne les exprime,

et qui a illuminé mon chemin par son amour et ses encouragements.

Si j'ai atteint cette étape aujourd'hui, c'est grâce à toi après Dieu.

Tu es mon refuge, ma lumière et la plus belle bénédiction de ma vie.

À mes sœurs et mon frère ...(Imane , KHawla , Hadil , Aissa)

mes compagnons de vie, ma force dans les moments difficiles et ma joie dans les moments
heureux.

Merci pour votre présence, votre affection sincère et votre soutien permanent.

Votre amour m'a donné le courage de continuer et la volonté de ne jamais abandonner.

À toutes les personnes ...

qui m'ont soutenue de près ou de loin

par une parole encourageante, une prière sincère ou un simple geste rempli de bonté.

Merci d'avoir cru en moi et d'avoir partagé avec moi cette longue aventure.

Je vous dédie ce modeste travail avec tout mon amour, ma gratitude et mon profond
respect.

NOTATIONS

X	Espace de Banach
H	Espace de Hilbert
$\mathbb{R}$	Ensemble des nombres réels
$\ \cdot\ $	La norme
$\langle \cdot, \cdot \rangle$	Le produit scalaire
$C[a, b]$	Espace des fonctions continues Sur $[a, b]$
$L^2[a, b]$	Espace des fonctions de carré intégrable
$K(x, t)$	Noyau de l'équation intégrale
$U(x)$	La Solution exacte
$U_n(x)$	La Solution approchée
P_n	L'opérateur de projection
$R_n(x)$	Le résidu
x_i	Points de Collocation
$\varphi_i(x)$	Fonctions de base

Table des matières

Introduction	1
1 Notions préliminaires	2
1.1 Espaces Fonctionnels	2
1.1.1 Espaces continus $C[a,b]$	2
1.1.2 Espaces $L^2[a,b]$	3
1.1.3 Espace de Banach	3
1.1.4 Espace de Hilbert	5
1.2 Normes et produits scalaires	6
1.2.1 Norme L^2	6
1.2.2 Produit scalaire	7
1.3 Polynômes Orthogonaux	7
1.3.1 Polynôme de Tchebychev	7
1.3.2 Polynôme de Legendre	10
1.3.3 Polynôme de Laguerre	11
1.3.4 Polynômes de Jacobi	11
2 Équations intégrales et leur classification	14
2.1 Définition générale	14
2.2 Classification des équations intégrales	14
2.2.1 Équations intégrales de Fredholm	14
2.2.2 Équations intégrales de Volterra	15

2.2.3	Équations intégrales de Volterra-Fredholm	15
2.3	Existence et unicité de la solution des équations intégrales linéaires de Volterra-Fredholm	16
2.4	Importance des méthodes numériques	17
2.4.1	Méthode de collocation	17
2.4.2	Méthode de Galerkin	18
2.4.3	La méthode de Ritz-Galerkin	18
2.4.4	La méthode de Bubnov-Galerkin	18
2.4.5	La méthode de Petrov-Galerkin	18
3	Méthodes de Collocation et de Galerkin pour la résolution numérique	19
3.1	Principe des méthodes de projection	19
3.2	Méthode de collocation	20
3.2.1	Application de la Méthode de collocation	21
3.3	Méthode de Galerkin	22
3.4	Exemples Numériques	24
3.4.1	Exemple 1 : Solution polynomiale	25
3.4.2	Exemple 2 : Solution exponentielle	28
3.4.3	Exemple 3 : Solution trigonométrique	31
3.4.4	Comparaison globale	35
	Conclusion générale	37
	Bibliographie	39

Introduction

Les équations intégrales jouent un rôle fondamental dans plusieurs domaines des mathématiques appliquées, de la physique et de l'ingénierie. Elles apparaissent dans la modélisation de nombreux phénomènes scientifiques tels que les problèmes de transfert de chaleur, la mécanique des fluides, l'électromagnétisme et la théorie du potentiel. Cependant, la résolution analytique exacte de ces équations demeure souvent difficile, ce qui conduit au recours à des méthodes numériques d'approximation.

Parmi les méthodes numériques les plus utilisées pour résoudre les équations intégrales figurent la méthode de collocation et la méthode de Galerkin. Ces deux approches reposent sur l'approximation de la solution dans des espaces fonctionnels de dimension finie, mais elles diffèrent par leur manière de construire le système approché et par leurs propriétés numériques.

La méthode de collocation se caractérise par sa simplicité de mise en œuvre et son efficacité dans plusieurs applications pratiques, tandis que la méthode de Galerkin est reconnue pour sa stabilité et sa forte base théorique, notamment dans l'analyse fonctionnelle et les méthodes variationnelles.

L'objectif principal de ce mémoire est d'effectuer une étude comparative entre la méthode de collocation et la méthode de Galerkin pour la résolution des équations intégrales. Nous analyserons les fondements théoriques de chaque méthode, leurs avantages, leurs limites ainsi que leur efficacité à travers différents exemples et applications numériques.

Plan du mémoire :

Le premier chapitre : Notions préliminaires.

Le deuxième chapitre : Équations intégrales et leur classification.

Le troisième chapitre : Méthodes de Collocation et de Galerkin pour la résolution numérique.

On termine notre mémoire par une conclusion générale.

Chapitre 1

Notions préliminaires

1.1 Espaces Fonctionnels

1.1.1 Espaces continus $C[a,b]$

Soit X l'ensemble de toutes les fonctions à valeurs réelles $x, y, \dots$ qui sont des fonctions d'une variable réelle indépendante t et sont définies et continues sur un intervalle fermé donné $J = [a, b]$ en choisissant la métrique définie par

$$d(x, y) = \max_{t \in J} |x(t) - y(t)|.$$

où $\max$ désigne le maximum, L'ensemble $C[a, b]$ muni de cette métrique constitue un espace métrique. (la lettre C suggère "continue.") c'est un espace de fonctions car chaque point de $C[a, b]$ est une fonction.

Chaque élément de cet espace est une fonction réelle continue définie sur l'intervalle $[a, b]$. L'ensemble de toutes ces fonctions forme un espace vectoriel réel, muni des opérations algébriques définies de la manière usuelle :

$$\begin{aligned}(x + y)(t) &= x(t) + y(t) \\ (ax)(t) &= ax(t), (a \in \mathbb{R})\end{aligned}$$

En effet, si x et y sont des fonctions continues réelles définies sur $[a, b]$ et si a est un réel, alors $x + y$ et ax sont également des fonctions continues sur $[a, b]$.

1.1.2 Espaces $L^2[a, b]$

Définition 1.1.1 On note $L^2[a, b]$ l'ensemble des fonctions (réelles ou complexes) définies sur l'intervalle $[a, b]$ dont le carré de la valeur absolue est intégrable, c'est-à-dire :

$$L^2([a, b]) = \{f : [a, b] \longrightarrow \mathbb{R}, \int_a^b |f(x)|^2 dx < \infty\}.$$

Autrement dit, une fonction f appartient à $L^2[a, b]$ si :

- f est mesurable sur $[a, b]$.
- L'intégrale $\int_a^b |f(x)|^2$ est finie.

Cet espace est muni de la norme L^2 :

$$\|f\|_2 = \left(\int_a^b |f(x)|^2 dx \right)^{\frac{1}{2}}.$$

Avec cette norme, $L^2[a, b]$ est un espace vectoriel normé complet.

$$\langle f, g \rangle = \int_a^b f(x) \overline{g(x)} dx.$$

1.1.3 Espace de Banach

suite de Cauchy:

Soit (x_n) une suite d'éléments d'un espace normé $(E, \|\cdot\|)$; on dit que la suite (x_n) est de cauchy :

$$\forall \varepsilon > 0, \exists N_\varepsilon, \forall p, q \geq N_\varepsilon \text{ on a } \|x_p - x_q\| < \varepsilon.$$

Soit (x_n) une suite de Cauchy dans un espace normé $(E, \|\cdot\|)$ contient une sous-suite (x_{n_k}) convergente vers x alors la suite (x_n) est aussi convergente vers le même élément x .

Soit (x_n) une suite de Cauchy alors il vient

$$\forall \varepsilon > 0, \exists N_\varepsilon, \forall p, q \geq N_\varepsilon \text{ on a } \|x_p - x_q\| < \varepsilon.$$

en particulier pour $n_k \geq N_\varepsilon$, on a

$$\forall p, n_k \geq N_\varepsilon, \|x_p - x_{n_k}\| < \varepsilon.$$

avec la convergence de la suite (x_{n_k}) vers x

$$n_k \geq N_\varepsilon, \|x_{n_k} - x\| < \varepsilon.$$

D'où la convergence de la suite (x_{n_k}) vers l'élément x

$$\begin{aligned} \forall p, n_k &\geq N_\varepsilon, \|x_p - x\| = \|x_p - x + x_{n_k} - x_{n_k}\| \\ &\leq \|x_p - x_{n_k}\| + \|x_{n_k} - x\| < \varepsilon \end{aligned}$$

Définition 1.1.2 *Un espace vectoriel normé $(E, \|\cdot\|)$ est dit complet, si toute suite de Cauchy (x_n) d'éléments de E est une suite convergente dans E .*

Autrement dit,

$$\forall \varepsilon > 0, \exists N_\varepsilon, \forall p, q \geq N_\varepsilon \text{ on a } \|x_p - x_q\| < \varepsilon.$$

Cela implique l'existence d'un élément $x \in E$ tel que

$$\lim_{n \rightarrow \infty} x_n = x$$

Espaces de Banach

Définition 1.1.3 *On appelle espace de Banach $(E; \|\cdot\|)$ tout espace vectoriel normé et complet pour la distance déduite de sa norme.*

1.1.4 Espace de Hilbert

Définition 1.1.4 *Un espace de Hilbert H est un espace préhilbertien complet pour la norme induite par son produit scalaire. En d'autres termes un espace de Hilbert est un espace de Banach dont la norme est induite par un produit scalaire.*

Généralement l'espace H est séparable et de dimension infinie, en d'autres termes, il existe un ensemble dénombrable partout dense dans H et pour tout $n \in \mathbb{N}$, il existe n vecteurs dans H linéairement indépendants

Exemple 1 L'espace $L^2(a,b)$

L'espace $L^2(I, \mathbb{R})$ des fonctions de carré intégrable défini par

$$L^2(I, \mathbb{R}) = \{f; \int_I |f(t)|^2 dt < \infty\}.$$

ou l'intégrale est prise au sens de Lebesgue, muni de la norme

$$\|f\|_2 = \left(\int_I |f(t)|^2 dt \right)^{\frac{1}{2}}.$$

induite par le produit scalaire définit par

$$\langle f, g \rangle = \int_I f(x)g(x)dx.$$

est un espace de Hilbert

1.2 Normes et produits scalaires

1.2.1 Norme L^2

Définition 1.2.1 La norme L^2 d'une fonction f est définie par :

$$\|f\|_{L^2([a,b])} = \left(\int_a^b |f(x)|^2 dx \right)^{\frac{1}{2}}.$$

Autrement dit :

1. On prend la valeur de la fonction $f(x)$
2. On calcule son module $|f(x)|$
3. On élève au carré $|f(x)|^2$
4. On intègre sur le domaine considéré
5. On prend la racine carrée

Interprétation intuitive La norme L^2 mesure la taille, la longueur ou l'énergie d'une fonction. Elle est l'analogue continu de la norme euclidienne d'un vecteur :

$$\|(x_1, x_2, \dots, x_n)\| = \sqrt{x_1^2 + x_2^2 + \dots + x_n^2}.$$

La somme devient une intégrale.

Quand une fonction appartient-elle à L^2 ?

$$f \in L^2([a,b]) \Leftrightarrow \int_a^b |f(x)|^2 dx < \infty.$$

Cela signifie que son énergie totale est finie.

- $f(x) = \frac{1}{1+x^2}$ sur $\mathbb{R}$ appartient à L^2
- $f(x) = 1$ sur $\mathbb{R}$ n'appartient pas à L^2

Propriétés de la norme L^2 : La norme L^2 vérifie les propriétés fondamentales d'une norme :

1. Positivité : $\|f_{L^2}\| \geq 0$
2. Séparation : $\|f_{L^2}\| = 0 \Leftrightarrow f = 0$ (presque partout)

3. Homogénéité : $\|\alpha f\|_{L^2} = |\alpha| \|f\|_{L^2}$
4. Inégalité triangulaire : $\|f + g\|_{L^2} \leq \|f\|_{L^2} + \|g\|_{L^2}$

1.2.2 Produit scalaire

Définition 1.2.2 Soit X un espace vectoriel réel ou complexe. Une application $(.,.) : X \times X \rightarrow \mathbb{R}$ (ou $\mathbb{C}$)

est appelée produit scalaire sur X si elle vérifie :

1. $\langle \varphi, \varphi \rangle \geq 0$ (positivité)
2. $\langle \varphi, \varphi \rangle = 0$ si et seulement si $\varphi = 0$ (définie positive)
3. $\langle \varphi, \psi \rangle = \overline{\langle \psi, \varphi \rangle}$ (symétrie)
4. $\langle \alpha \varphi + \beta \psi, \chi \rangle = \alpha \langle \varphi, \chi \rangle + \beta \langle \psi, \chi \rangle$ (linéarité)

Un produit scalaire définit une norme par

$$\|\varphi\| = \sqrt{\langle \varphi, \varphi \rangle}$$

Si l'espace X est complet pour cette norme, alors X est appelé espace de Hilbert ; sinon, c'est un préhilbertien.

Exemple 1.2.1 $L^2[a, b]$

est un espace de Hilbert muni du produit scalaire

$$(\varphi, \psi) = \int_a^b \varphi(x) \overline{\psi(x)} dx.$$

1.3 Polynômes Orthogonaux

1.3.1 Polynôme de Tchebychev

En mathématiques, un polynôme de Tchebychev est un terme de l'une des deux suites de polynômes orthogonaux particulières reliées à la formule de Moivre. Les polynômes de Tchebychev sont nommés ainsi en l'honneur du mathématicien russe Pafnouti Lvovitch Tchebychev.

Définition 1.3.1 Soit $n \in \mathbb{N}$, on appelle polynôme de Tchebychev de degré n , l'application T_n définie de $[-1, 1]$ dans $\mathbb{R}$ par :

$$T_n(x) = \cos(n \arccos(x))$$

Il existe deux suites de polynômes de Tchebychev de degré n , l'une nommée polynômes de Tchebychev de première espèce et notée T_n et l'autre nommée polynômes de Tchebychev de seconde espèce et notée U_n .

Relation de récurrence entre les polynômes de Tchebychev de degré n La suite de polynômes de Tchebychev T_n peut être définie par la relation de récurrence. Ces deux suites de polynômes de Tchebychev peuvent être définies par la relation de récurrence

1. $T_0(x) = 1, T_1(x) = x$ et pour tout $n \in \mathbb{N}$

$$T_{n+2} = 2xT_{n+1} - T_n$$

2. $U_0(x) = 1, U_1(x) = 2x$ et pour tout $n \in \mathbb{N}$

$$U_{n+2} = 2xU_{n+1} - U_n$$

Une définition alternative de ces polynômes peut être donnée par les relations trigonométriques:

$$T_n(\cos \theta) = \cos(n\theta), \text{ soit: } T_n(x) = \cos(n \arccos x)$$

et

$$U_n(\theta) = \frac{\sin((n+1)\theta)}{\sin \theta}$$

ce qui revient, par exemple, à considérer $T_n(\cos \theta)$ comme le développement de $\cos n\theta$ sous forme de polynôme en $\cos \theta$:

Contrairement à d'autres familles de polynômes orthogonaux, tels ceux de Legendre, d'Hermite ou de Laguerre, les polynômes de Tchebychev n'ont pratiquement pas d'application directe en physique. En revanche, ils sont particulièrement utiles en analyse numérique pour l'interpolation polynomiale de fonctions. En premier lieu, en ce qui concerne le choix des

points d'interpolation, comme les zéros de $T_n(x)$ ou abscisses de Tchebychev, en vue de limiter le phénomène de Runge. Également, ils constituent une base alternative de polynômes par rapport à la base canonique X^n de $R[X]$ des polynômes de Lagrange, ce qui permet d'améliorer sensiblement la convergence.

Orthogonalité des polynômes de Tchebychev Chacune de T_n et U_n est une suite de polynômes orthogonaux par rapport à un produit scalaire de fonctions, associé à la fonction poids $w(x) = (1 - x^2)^{-\frac{1}{2}}$ sur l'intervalle $[-1, 1]$.

Les polynômes de Tchebychev de première espèce sont w -orthogonaux, c'est-à-dire orthogonaux relativement à la fonction poids définie par : $w(x) = (1 - x^2)^{-\frac{1}{2}}$.

En effet, pour $x = \cos(\theta)$, $dx = -\sin(\theta)d\theta$, on a alors

$$\begin{aligned} \int_{-1}^1 T_n(x) T_m(x) w(x) dx &= - \int_{\pi}^0 \frac{\cos(n\theta) \cos(m\theta) \sin(\theta)}{|\sin(\theta)|} d\theta \\ &= \int_0^{\pi} \cos(n\theta) \cos(m\theta) d\theta \\ &= \frac{1}{2} \int_0^{\pi} \cos((n+m)\theta) + \cos((n-m)\theta) d\theta \\ &= \begin{cases} \pi, & \text{si } n = m = 0 \\ \frac{\pi}{2}, & \text{si } n = m \neq 0 \\ 0, & \text{si } n \neq m \end{cases} . \end{aligned}$$

Racines des polynômes de Tchebychev de première espèce

Proposition 1.3.1 Soit $n \in \mathbb{N}^*$, Le polynôme de Tchebychev de première espèce T_n admet exactement n racines simples définies par :

$$x_k = \cos\left(\left(\frac{2k-1}{2n}\right)\pi\right), k = 1, \dots, n$$

Preuve. Soit $n \in \mathbb{N}^*$, T_n et $k = 1, \dots, n$

$$\begin{aligned}
 T_n(x_k) &= \cos(n \arccos(x_k)) \\
 &= \cos(n \arccos(\cos(\frac{2k-1}{2n}\pi))) \\
 &= \cos(\frac{2k-1}{2n}\pi) \\
 &= 0 \quad \blacksquare
 \end{aligned}$$

T_n étant de degré n , les x_k sont exactement les zéros de T_n . Elles sont simples puisque pour tout $k \neq k'$, on a $x_k \neq x_{k'}$

Remarque 1.3.1 On peut généraliser cela à tout intervalle fermé $[a, b]$, d'où

$$x_k = \frac{a+b}{2} + \frac{b-a}{2} \cos\left(\left(\frac{2k-1}{2n}\right)\pi\right), k = 1, \dots, n$$

1.3.2 Polynôme de Legendre

En mathématiques et en physique théorique, les polynômes de Legendre constituent l'exemple le plus simple d'une suite de polynômes orthogonaux. Ce sont des solutions polynomiales $P_n(x)$ sur le segment $[-1, 1]$, de l'équation différentielle de Legendre :

$$\frac{d}{dx} \left[(1-x^2) \frac{d}{dx} P_n(x) \right] + n(n+1) P_n(x) = 0,$$

dans le cas particulier où le paramètre n est un entier naturel. De façon équivalente, les polynômes de Legendre sont les fonctions propres de l'endomorphisme de $\mathbb{R}[X]$ défini par :

$$P \rightarrow u(P) = \left[(1-x^2) \frac{dP}{dx} \right]$$

pour les valeurs propres $-n(n+1)$, $n \in \mathbb{N}$.

Ces polynômes orthogonaux ont de nombreuses applications tant en mathématiques, par exemple pour la décomposition d'une fonction en série de polynômes de Legendre, qu'en calcul numérique des intégrales notamment dans les méthodes de quadrature de Gauss, qu'en physique, où l'équation de Legendre apparaît naturellement lors de la résolution des équations de Laplace ou de Helmholtz en coordonnées sphériques.

1.3.3 Polynôme de Laguerre

En mathématiques, les polynômes de Laguerre, nommés d'après Edmond Laguerre, sont les solutions normalisées de l'équation de Laguerre

$$xy'' + (1 - x)y' + ny = 0$$

qui est une équation différentielle linéaire homogène d'ordre 2 et se réécrit sous la forme de Sturm-Liouville

$$-\frac{d}{dx} \left(x \exp(-x) \frac{dy}{dx} \right) - n \exp(-x) y = 0$$

Cette équation a des solutions non singulières seulement si n est un entier positif. Les solutions L_n forment une suite de polynômes orthogonaux dans $L^2(\mathbb{R}^+, \exp(-x))$, et la normalisation se fait en leur imposant d'être de norme 1, donc de former une base orthonormée. Ils forment même une base de Hilbert de $L^2(\mathbb{R}^+, \exp(-x) dx)$.

Cette suite de polynômes peut être définie par la formule de Rodrigues

$$L_n(x) = \frac{\exp(x)}{n!} \frac{d^n}{dx^n} (\exp(-x) x^n)$$

La suite des polynômes de Laguerre est une suite de Sheffer.

Les polynômes de Laguerre apparaissent en mécanique quantique dans la partie radiale de la solution de l'équation de Schrödinger pour un atome à un électron

Le coefficient dominant de L_n est $(-1)^n/n!$. Les physiciens utilisent souvent une définition des polynômes de Laguerre où ceux-ci sont multipliés par $(-1)^n n!$, obtenant ainsi des polynômes unitaires

1.3.4 Polynômes de Jacobi

Définition 1.3.2 *Les polynômes de Jacobi sont les polynômes orthogonaux du produit scalaire*

$$\langle P, Q \rangle = \int_{-1}^1 (1-x)^\alpha (1+x)^\beta P(x) Q(x) dx \quad \text{avec } \alpha > -1 \text{ et } \beta > -1.$$

Pour tout $n \in \mathbb{N}$, on note le polynôme de Jacobi d'indice n par $P_n^{(\alpha, \beta)}(x)$. Il sont des solutions de l'équation différentielle homogène de deuxième ordre

$$(1-x^2)z'' + [\beta - \alpha - (\alpha + \beta + 2)x]z' + n(n + \alpha + \beta + 1)z = 0,$$

avec n est le degré du polynôme.

Les polynômes de *Jacobi* est définis explicitement par trois formules :

Formule de Rodrigue

$$(1-x)^\alpha(1+x)^\beta P_n^{(\alpha, \beta)}(x) = \frac{(-1)^n}{2^n n!} \frac{d^n}{dx^n} [(1-x)^{n+\alpha}(1+x)^{n+\beta}],$$

tels que α, β sont deux paramètres vérifiant des conditions d'orthogonalité $\alpha > -1$ et $\beta > -1$.

Formule analytique

$$P_n^{(\alpha, \beta)}(x) = \frac{1}{2^n} \sum_{\ell=0}^n \binom{n+\alpha}{n-\ell} \binom{n+\beta}{\ell} (x-1)^\ell (x+1)^{n-\ell}.$$

Relation de récurrence

Les polynômes de Jacobi vérifiant la relation de récurrence

$$\begin{aligned} P_0^{(\alpha, \beta)}(x) &= 1, \\ P_1^{(\alpha, \beta)}(x) &= \frac{\alpha + \beta + 2}{2}x + \frac{\alpha - \beta}{2}, \\ P_{n+1}^{(\alpha, \beta)}(x) &= \alpha_n \tau P_n^{(\alpha, \beta)}(x) + \beta_n P_n^{(\alpha, \beta)}(x) + \gamma_n P_{n-1}^{(\alpha, \beta)}(x), \end{aligned}$$

avec

$$\begin{aligned} \alpha_n &= \frac{(2n + \alpha + \beta + 1)(2n + \alpha + \beta + 2)}{2(n+1)(n + \alpha + \beta + 1)} \\ \beta_n &= \frac{(2n + \alpha + \beta + 2)(\alpha^2 - \beta^2)}{2(n+1)(n + \alpha + \beta + 1)(2n + \alpha + \beta)} \end{aligned}$$

$$\gamma_n = \frac{(n + \alpha)(n + \beta)(2n + \alpha + \beta + 2)}{(n + 1)(n + \alpha + \beta + 1)(2n + \alpha + \beta)^2}$$

Les polynômes orthogonaux standards sont des cas particuliers des polynômes de *Jacobi*.

Les cas les plus connus sont :

- Les polynômes de *Gegenbauer* (ou ultrasphérique), obtenus en posant $\alpha = \beta$.
- Les polynômes de *Legendre* L_n en posant $\alpha = \beta = 0$.
- Les différents types de polynômes de *Tchebyshev* T_n^i pour $i = 1, 2, 3$ et 4 qui sont obtenus en choisissant des valeurs spécifiques de α et β comme suit :

$$\left\{ \begin{array}{l} T_n^1(x) = P_n^{(-\frac{1}{2}, -\frac{1}{2})}(x) \\ T_n^2(x) = P_n^{(\frac{1}{2}, \frac{1}{2})}(x) \\ T_n^3(x) = P_n^{(-\frac{1}{2}, \frac{1}{2})}(x) \\ T_n^4(x) = P_n^{(\frac{1}{2}, -\frac{1}{2})}(x) \end{array} \right.$$

Chapitre 2

Équations intégrales et leur classification

2.1 Définition générale

Définition 2.1.1 Une équation fonctionnelle est connue sous le nom d'équation intégrale linéaire, avec une inconnue φ de la forme.

$$\varphi(x) = f(x) + \lambda \int_{\alpha(x)}^{\beta(x)} k(x,t) \varphi(t) dt$$

où k et f sont des fonctions donnés, $\alpha(x)$, $\beta(x)$ sont les bornes de l'intégration, λ est un paramètre numérique non nul, réel ou complexe.

2.2 Classification des équations intégrales

2.2.1 Équations intégrales de Fredholm

Définition 2.2.1 On appelle équation de Fredholm de première espèce une équation de la forme:

$$\int_a^b K(x,t) \varphi(t) dt = f(x).$$

Où φ est la fonction inconnue, f et k sont des fonctions connues, les bornes d'intégration sont constantes. C'est la caractéristique principale d'une équation de Fredholm.

Définition 2.2.2 On appelle équation intégrale de Fredholm de second espèce une équation de la forme :

$$\varphi(x) = f(x) + \lambda \int_a^b k(x,t) \varphi(t) dt.$$

Où $\varphi(x)$ est la fonction inconnue, $k(x,t)$ et $f(x)$ des fonctions données, λ est un paramètre réel ou complexe donné.

2.2.2 Équations intégrales de Volterra

les équations de Volterra sont des cas particuliers d'équations intégrales de Fredholm. il suffit de prendre le noyau k est tel que $k(x,t) = 0$ pour $x < t$.

Définition 2.2.3 On appelle équation intégrale linéaire de Volterra de première espèce une équation à une inconnue $\varphi(x)$ de la forme :

$$\int_a^x k(x,t) \varphi(t) dt = f(x)$$

Définition 2.2.4 On appelle équation intégrale de Volterra de seconde espèce une équation à inconnue $\varphi(x)$ de la forme :

$$\varphi(x) = f(x) + \lambda \int_a^x k(x,t) \varphi(t) dt$$

2.2.3 Équations intégrales de Volterra-Fredholm

Une équation intégrale de Volterra-Fredholm est une combinaison des intégrales de Volterra et Fredholm disjoints, apparaît dans une équation intégrale.

Définition 2.2.5 On appelle équation intégrale de Volterra-Fredholm une équation de la forme :

$$\varphi(x) = f(x) + \lambda_1 \int_a^x k_1(x,t) \varphi(t) dt + \lambda_2 \int_a^b k_2(x,t) \varphi(t) dt, \quad x \in [a,b]. \quad (2.2.1)$$

On appelle équation intégrale mixte une équation de la forme :

$$\varphi(x) = f(x) + \lambda \int_a^x \int_a^b k(s, t) \varphi(t) dt ds.$$

Où les fonctions k_1 , k_2 et f sont connues et $\varphi(x)$ la fonction inconnue.

Exemple 2.2.1

$$\varphi(x) = \exp(x) + x + 1 + \int_0^x (x-t) \varphi(t) dt + \int_0^1 \exp(x-t) \varphi(t) dt, \quad x \in [0,1]$$

$$\varphi(x) = \frac{17}{2}x^2 + 11x + \int_a^x \int_a^b (s-t) \varphi(t) dt ds$$

2.3 Existence et unicité de la solution des équations intégrales linéaires de Volterra-Fredholm

Dans cette section nous rappelons les théorèmes que nous allons utiliser pour obtenir des résultats d'existence et unicité de solution de l'équation (2.2.1).

Définition 2.3.1 Soit X un espace normé et $T : X \rightarrow X$ un opérateur T est dit un opérateur de Picard s'il existe $\varphi_0 \in X$ unique tel que

$$T(\varphi) = \varphi_0, \text{ et } \lim_{n \rightarrow \infty} T^n(\varphi_0) = \varphi_0, \text{ pour tout } \varphi \text{ de } X$$

Théorème 2.3.1 (principe de contraction) Soit X un espace normé, si $T : X \rightarrow X$ un opérateur de contraction admet un point fixe unique φ , alors T est un opérateur de Picard

$$\|\varphi_0 - T^n(\varphi_0)\| \leq \frac{a^n}{1-a} \|\varphi - T(\varphi)\| \text{ pour tout } n \in \mathbb{N}.$$

Théorème d'existence et d'unicité

On considère l'équation linéaire de Volterra-Fredholm

$$\varphi(x) = f(x) + \lambda_1 \int_a^x k_1(x, t) \varphi(t) dt + \lambda_2 \int_a^b k_2(x, t) \varphi(t) dt, \quad x \in [a, b]. \quad (2.3.1)$$

Où

- 1) $f \in C[a, b], k_1(x, t) \in C(D_1)$, avec $D_1 = \{(x, t) \in \mathbb{R}^2 \text{ tel que } a \leq t \leq x \leq b\}$
- 2) $\varphi \in C[a, b], k_2(x, t) \in C(D_2)$, avec $D_2 = [a, b] \times [a, b]$
- 3) $M_1 = \max_{(x, t) \in D_1} |k_1(x, t)|$, et $M_2 = \max_{(x, t) \in D_2} |k_2(x, t)|$

Théorème 2.3.2 Dans les condition de continuité ci-dessus, supposons qu'il existe une constante $c > 0$ tel que:

$$\frac{1}{c} [M_1 + M_2 \exp(c(b-a))] < 1.$$

Alors l'équation (2.2) a une solution unique $\varphi \in C[a, b]$, et cette solution peut être obtenue par la méthode d'approximation successive, à partir de $\varphi \in C[a, b]$

Preuve. voir[7]. ■

2.4 Importance des méthodes numériques

Les méthodes numériques sont importantes car elles permettent de résoudre des problèmes mathématiques dont les solutions analytiques sont difficiles ou impossibles à obtenir. Elles jouent un rôle essentiel dans la simulation, la modélisation et l'optimisation dans plusieurs domaines tels que les sciences, l'ingénierie, l'économie et la physique.

Il existe diverses méthodes de projection. Les plus employées sont :

2.4.1 Méthode de collocation

En mathématiques, une méthode de collocation est une méthode de résolution numérique d'équations différentielles ordinaires, d'équations aux dérivées partielles et d'équations intégrales. L'idée est de choisir un espace de dimension finie de solutions candidates (généralement des polynômes jusqu'à un certain degré) et un certain nombre de points dans le domaine (appelés points de collocation), et de sélectionner la solution qui satisfait l'équation donnée aux points de collocation

2.4.2 Méthode de Galerkin

Est une technique de projection dans laquelle l'espace d'approximation est choisi de manière à ce que l'erreur soit orthogonal au sous-espace des fonctions de base.

2.4.3 La méthode de Ritz-Galerkin

Est une méthode utilisée pour résoudre des équations intégrales (ou des problèmes aux dérivées partielles) en approximant la solution exacte par une combinaison linéaire de fonctions de test choisies. Pour une équation intégrale, l'idée générale est de transformer le problème en un système d'équations algébriques de taille finie, que l'on peut résoudre numériquement.

2.4.4 La méthode de Bubnov-Galerkin

Est une technique d'approximation utilisée principalement pour résoudre des équations aux dérivées partielles (EDP), mais elle peut également être adaptée pour résoudre des équations intégrales. Cette méthode fait partie des méthodes des éléments finis, mais elle peut aussi être appliquée à des problèmes d'équations intégrales en utilisant des espaces fonctionnels appropriés.

2.4.5 La méthode de Petrov-Galerkin

Est une méthode numérique utilisée pour résoudre des équations intégrales, en particulier les équations intégrales de type Fredholm ou Volterra. Elle est une généralisation de la méthode classique de Galerkin, et elle est souvent utilisée dans des situations où l'on veut obtenir des approximations plus flexibles ou plus précises que celles fournies par la méthode de Galerkin standard.

Chapitre 3

Méthodes de Collocation et de Galerkin pour la résolution numérique

3.1 Principe des méthodes de projection

Dans toutes les méthodes de projection, on étudie la résolution de l'équation intégrale de Fredholm de deuxième espèce de la forme

$$\lambda \varphi(s) - \int_D k(s, t) \varphi(t) dt = f(s), \quad (s \in D) \quad (3.1.1)$$

L'ensemble D est fermé et borné d'un espace complet de fonctions V , tel que $V = C(D)$ ou $V = L^2(D)$ presque partout. On choisit une suite finie d'approximation de sous espace V_n , tel que $V_n \subset V, n \geq 1$, et V_n de dimension k_n . Soit la base $(\varphi_1, \varphi_2, \dots, \varphi_{k_n})$ de V_n , ($k = k_n$). Alors le principe de méthode de projection consiste à trouver une suite de fonctions $\varphi_n \in V_n$, tel que

$$\varphi_n(s) = \sum_{j=1}^{k_n} c_j \varphi_j(s) \quad s \in D$$

On introduit le résidu $r_n(s)$ pour approcher la solution de (3.1)

$$\begin{aligned} r_n(s) &= \lambda \varphi_n(s) - \int_D k(s; t) \varphi_n(t) dt - f(s) \\ &= \sum_{j=1}^{k_n} c_j \left(\lambda \varphi_j(s) - \int_D k(s; t) \varphi_j(t) dt \right) - f(s) \end{aligned}$$

On note cette approximation par $\varphi \approx \varphi_n$ L'équation (3.1) qui peut être écrite sous la forme de notation d'opérateur, tel que

$$(\lambda - k) \varphi = f$$

et le résidu r_n peut être écrite sous forme

$$r_n = (\lambda - K) \varphi_n - f$$

Les coefficients $\{c_1, c_2, \dots, c_n\}$ doivent être choisis de telle sorte que

$$r_n(s) \rightarrow 0$$

3.2 Méthode de collocation

La méthode de collocation est une méthode numérique utilisée pour résoudre des équations intégrales. Elle repose sur la réduction du problème d'intégrale à un système d'équations algébriques en choisissant judicieusement un ensemble de points (appelés points de collocation) où la solution approximative est forcée d'être exacte.

Pour résoudre l'équation de Fredholm de deuxième espèce

$$\varphi(s) - \int k(s, t) \varphi(t) dt = f(s) \quad (3.2.1)$$

on choisit un ensemble de noeuds distincts $\{s_0, s_1, \dots, s_n\} \subset D$ tels que :

$$r_n(s_i) = \sum_{j=0}^n c_j \left(\varphi_j(s_i) - \int_D k(s_i, t) \varphi_j(t) dt \right) - f(s_i) = 0 \quad i = 0, 1, 2, \dots, n$$

ce qui permet de déterminer les coefficients $\{c_0, c_1, \dots, c_n\}$ comme solution du système linéaire suivant :

$$\sum_{j=0}^n c_j \left(\varphi_j(s_i) - \int_D k(s_i, t) \varphi_j(t) dt \right) = f(s_i), \quad i = 0, 1, 2, \dots, n$$

D'où

$$\sum_{j=0}^n c_j \varphi_j(s_i) = \sum_{j=0}^n c_j \int_D k(s_i, t) \varphi_j(t) dt + f(s_i) \quad i = 1, 2, \dots, n.$$

Comme

$$P_n \varphi(s) = \sum_{j=0}^n c_j \varphi_j(s),$$

on peut aussi écrire :

$$P_n \varphi(s_i) = \sum_{j=0}^n \left(\int_D c_j k(s_i, t) \varphi_j(t) dt + f(s_i) \right) = \sum_{j=0}^n c_j \varphi_j(s_i)$$

avec $s \in D$, et $P_n : V = C(D) \rightarrow V_n$, est un opérateur de projection. Pour déterminer les coefficients c_j , on résout le système

$$\sum_{j=0}^n c_j \left(\varphi_j(s_i) - \int_D k(s_i, t) \varphi_j(t) dt \right) = f(s_i), \quad i = 0, 1, 2, \dots, n \quad (3.2.2)$$

Ce système admet une solution unique si et seulement si

$$\det [\psi_j(s_i)] \neq 0$$

où $[\psi_j(s_i)]$ est la matrice induite du système (3.3)

Il faut noter que cette condition implique que le système des fonctions $\{\varphi_0, \varphi_1, \dots, \varphi_n\}$ est linéairement indépendant.

3.2.1 Application de la Méthode de collocation

Si l'on considère le système $\{1, s, \dots, s^n\}$ des monômes linéairement indépendants, alors on obtient le déterminant de Vandermonde, pour tout $i \in \{0, 1, 2, \dots, n\}$, on construit une fonction $l_i \in V_n$, telle que :

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

D'où

$$P_n \varphi(s) = \sum_{j=0}^n \varphi(s_j) l_j(s), \quad s \in D$$

Le système $\{l_1, l_2, \dots, l_k\}$ est appelé la base des fonctions de Lagrange vérifiant :

$$\|P_n\| = \max_{s \in D} \sum_{j=0}^n |l_j(s)|.$$

Si on prend $V_n = \text{span}\{1, s, \dots, s^n\}$ alors la base des fonctions de Lagrange est sont données par :

$$l_i(s) = \prod_{j=0, j \neq i}^n \frac{(s - s_j)}{(s_i - s_j)}, \quad (i = 0, 1, \dots, n)$$

3.3 Méthode de Galerkin

Soit $V = L^2(D)$ ou un autre espace de fonctions de Hilbert, et donne $\langle \cdot \rangle$ muni d'un produit scalaire pour. Exiger le résiduel r_n pour satisfaire

$$\langle r_n, \varphi \rangle = 0 \tag{3.3.1}$$

Le côté gauche est le coefficient de Fourier de r_n associé à u_i . Si $\{\varphi_1, \varphi_2, \dots, \varphi_{k_n}\}$ se compose des membres principaux d'une famille orthonormée $s = (\varphi_i)_{i \geq 1}$ qui s'étend sur V , alors (3.3.1) nécessite les termes avancés à être zéro dans l'expansion de Fourier de r_n par rapport à s .

Pour trouver φ_n , appliquer (3.3.1) à (3.1.1) écrit comme $(\lambda - K)\varphi = f$. Cela donne le système linéaire

$$\sum_{j=1}^{k_n} c_j \langle \varphi_j, \varphi_i \rangle - \lambda \langle K \varphi_j, \varphi_i \rangle - \langle f, \varphi_i \rangle = 0 \quad i = 1, 2, \dots, k_n$$

C'est la méthode de Galerkin . Pour obtenir une solution approximative à (3.1.1) ou $((\lambda - K) \varphi = f)$. Le système a-t-il une solution ? Si c'est le cas, est-ce unique ?

La suite résultante desolutions approximatives φ_n converge-t-elle en φ dans V .

La converget-elle en $C(D)$, c'est- à-dire que un converge uniformément en φ . Notez également que la formulation ci-dessus contient des intégrales doubles $\langle k\varphi_j, \varphi_i \rangle$, Ceux-ci doivent être calculés numériquement.

Rappeler que

$$P_n h = 0 \text{ si seulement si } \langle h, \varphi_i \rangle = 0 \quad i = 1, 2, \dots, k_n$$

À l'aide de la projection orthogonale P_n , nous pouvons réécrire (3.3.1) comme|

$$P_n r_n = 0$$

Les polynômes de legendre sont utilisés comme fonctions d'essai dans la base. Pour cela,nous proposons une brève introduction des polynômes de Legendre d'abord. Ensuite, nous dérivons une formulation matricielle par la technique de la méthode Galerkin

Exemple 3.3.1 Soit l'équation intégrale de seconde espèce

$$u(x) = 1 - \frac{4}{3}x + \int_{-1}^1 (xt^2 - 1) u(t) dt$$

Prenant pour système complet sur $[-1; 1]$ un système de polynômes de legendre $\{l_i(x)\}$ $i = 1, 2, 3$ avec $l_1(x) = 1, l_2(x) = x, l_3(x) = \frac{3x^2-1}{2}$.

On cherche la solution sous la forme

$$u_3(x) = c_1 + c_2x + c_3 \frac{3x^2 - 1}{2}$$

Le résidu est égale a

$$r_3(x) = u_3(x) - 1 - \frac{4}{3}x + \int_{-1}^1 (xt^2 - 1) \left(c_1 + c_2x + c_3 \frac{3x^2 - 1}{2} \right) dt$$

En calculant l'intégrale on obtient

$$r_3(x) = \frac{4}{3}x + 3c_1 - \frac{2}{3}xc_1 + xc_2 - \frac{4}{15}xc_3 + c_3 \left(\frac{3}{2}x^2 - \frac{1}{2} \right) - 1$$

d'après l'orthogonalité du résidu r_3 au système $\left\{1, x, \frac{3x^2-1}{2}\right\}$ on a

$$\begin{aligned} \langle r_3, 1 \rangle &= 6c_1 - 2 = 0 \\ \langle r_3, x \rangle &= \frac{8}{9} + \frac{2}{3}c_2 - \frac{8}{45}c_3 - \frac{4}{9}c_1 = 0 \\ \left\langle r_3, \frac{3x^2-1}{2} \right\rangle &= \frac{2}{5}c_3 = 0 \end{aligned}$$

On obtient

$$c_1 = \frac{1}{3}, c_2 = -\frac{10}{9}, c_3 = 0$$

Donc la solution est

$$u_3(x) = \frac{1}{3} - \frac{10}{9}x$$

3.4 Exemples Numériques

Dans ce chapitre, nous étudions les performances des méthodes de collocation et de Galerkin pour la résolution numérique des équations intégrales de Fredholm du second type. L'objectif est de comparer la précision, la stabilité et le comportement numérique des deux méthodes à travers trois exemples représentatifs :

- **Exemple 1** : solution polynomiale $u(x) = 1 + \frac{9}{4}x$;
- **Exemple 2** : solution exponentielle $u(x) = e^x + \frac{3}{2}x$;
- **Exemple 3** : solution trigonométrique $u(x) = \sin(x) + \frac{3}{2}(\sin 1 - \cos 1)x$;

Pour chaque exemple, nous calculons la solution exacte, construisons les systèmes algébriques associés aux deux méthodes, résolvons numériquement et analysons les erreurs dans les normes L^2 et L^∞ .

Les erreurs sont définies par :

$$E_{L^2} = \|u - u_h\|_{L^2[0,1]} = \left(\int_0^1 |u(x) - u_h(x)|^2 dx \right)^{\frac{1}{2}}, \quad E_{L^\infty} = \|u - u_h\|_\infty = \max_{x \in [0,1]} |u(x) - u_h(x)|.$$

3.4.1 Exemple 1 : Solution polynomiale

Énoncé

Résoudre l'équation intégrale de Fredholm du second type :

$$u(x) = x + 1 + \int_0^1 xtu(t)dt, \quad x \in [0, 1]. \quad (1)$$

Calcul détaillé de la solution exacte

Posons $A = \int_0^1 tu(t)dt$. L'équation (1) devient :

$$u(x) = 1 + (1 + A)x.$$

En substituant dans la définition de A :

$$\begin{aligned} A &= \int_0^1 t [1 + (1 + A)t] dt = \int_0^1 t dt + (1 + A) \int_0^1 t^2 dt \\ &= \frac{1}{2} + \frac{1 + A}{3} + \frac{1}{2} + \frac{1}{3} + \frac{A}{3}. \end{aligned}$$

Donc $A - \frac{A}{3} = \frac{5}{6}$, soit $\frac{2A}{3} = \frac{5}{6}$, d'où $A = \frac{5}{4}$

Solution exacte :

$$u(x) = 1 + \frac{9}{4}x.$$

Construction du système de collocation

On cherche $u_h(x) = c_0 + c_1x + c_2x^2$ Dans $V = \text{span}\{1, x, x^2\}$

Points de collocation : $x_0 = 0, x_1 = \frac{1}{2}, x_2 = 1$

On calcule d'abord :

$$\int_0^1 x t u_h(t) dt = x \int_0^1 t(c_0 + c_1 t + c_2 t^2) dt = x \left(\frac{c_0}{2} + \frac{c_1}{3} + \frac{c_2}{4} \right).$$

Le résidu est :

$$R(x) = c_0 + c_1 x + c_2 x^2 - 1 - x - 1 - x \left(\frac{c_0}{2} + \frac{c_1}{3} + \frac{c_2}{4} \right).$$

La condition $R(x_i) = 0$ donne $c_0(1 - \frac{x_i}{2}) + c_1(x_i - \frac{x_i}{3}) + c_2(x_i^2 - \frac{x_i}{4}) = 1 + x_i$, soit le système:

pour $i = 0, (x_0 = 0)$:

$$c_0 = 1$$

pour $i = 1, (x_1 = \frac{1}{2})$:

$$\frac{3}{4}c_0 + \frac{1}{3}c_1 + \frac{3}{16}c_2 = \frac{3}{2} \Rightarrow \left(\frac{c_1}{3} + \frac{3c_2}{16} \right) = \frac{3}{4}. \quad (E_1)$$

pour $i = 2, (x_2 = 1)$:

$$\frac{1}{2}c_0 + \frac{2}{3}c_1 + \frac{3}{4}c_2 = 2. \Rightarrow \frac{2}{3}c_1 + \frac{3}{4}c_2 = \frac{3}{2}. \quad (E_2)$$

De (E_1) : $c_1 = \frac{9}{4} - \frac{9}{16}c_2$. En substituant dans (E_2) : $\frac{3}{2} - \frac{3}{8}c_2 + \frac{3}{4}c_2 = \frac{3}{2}$, donc $\frac{3}{8}c_2 = 0$, d'où $c_2 = 0$ et $c_1 = \frac{9}{4}$.

Solution par collocation :

$$c_0 = 1, \quad c_1 = \frac{9}{4}, \quad c_2 = 0, \quad u_h^{\text{col}}(x) = 1 + \frac{9}{4}x.$$

Construction du système de Galerkin

La condition de Galerkin $\langle R, \varphi_i \rangle = 0$ avec $\varphi_i(x) = x^i$ donne :

$$\sum_{j=0}^2 c_j \underbrace{\left[\frac{1}{i+j+1} - \frac{1}{(i+2)(j+2)} \right]}_{A_{ij}} = \underbrace{\int_0^1 (1+x)x^i dx}_{B_i}.$$

Calcul des éléments de A :

$$A_{ij} = \frac{1}{i+j+1} - \frac{1}{(i+2)(j+2)}, \quad i, j = 0, 1, 2$$

$$\begin{pmatrix} 1 - \frac{1}{4} & \frac{1}{2} - \frac{1}{6} & \frac{1}{3} - \frac{1}{8} \\ \frac{1}{2} - \frac{1}{6} & \frac{1}{3} - \frac{1}{9} & \frac{1}{4} - \frac{1}{12} \\ \frac{1}{3} - \frac{1}{8} & \frac{1}{4} - \frac{1}{12} & \frac{1}{5} - \frac{1}{16} \end{pmatrix} = \begin{pmatrix} \frac{3}{4} & \frac{1}{3} & \frac{5}{24} \\ \frac{1}{3} & \frac{2}{9} & \frac{1}{6} \\ \frac{5}{24} & \frac{1}{6} & \frac{11}{80} \end{pmatrix}.$$

Calcul du second membre :

$$B_i = \int_0^1 (1+x)x^i dx = \frac{1}{i+1} + \frac{1}{i+2}.$$

$$B = \begin{pmatrix} 1 + \frac{1}{2} \\ \frac{1}{2} + \frac{1}{3} \\ \frac{1}{3} + \frac{1}{4} \end{pmatrix} = \begin{pmatrix} \frac{3}{2} \\ \frac{5}{6} \\ \frac{7}{12} \end{pmatrix}.$$

La résolution de $AC = B$ donne (la solution exacte étant dans V_2) : $c_0 = 1, c_1 = \frac{9}{4}, c_2 = 0$.

Solution par Galerkin :

$$u_h^{Gal}(x) = 1 + \frac{9}{4}x.$$

Analyse des erreurs

Puisque la solution exacte $u(x) = 1 + \frac{9}{4}x$ appartient à V_2 , les deux méthodes la reproduisent exactement :

$$U_h^{Col}(x) = u_h^{Gal}(x) = u(x).$$

$$E_{L^2}^{Col} = E_{L^2}^{Gal} = 0, \quad E_{L^\infty}^{Col} = E_{L^\infty}^{Gal} = 0.$$

Interprétation. Cet exemple constitue un **test de validation** : lorsque la solution exacte appartient à l'espace d'approximation V_n , les deux méthodes convergent sans erreur. Cela confirme la consistance des deux formulations.

3.4.2 Exemple 2 : Solution exponentielle

Énoncé

Résoudre l'équation intégrale de Fredholm du second type :

$$u(x) = e^x + \int_0^1 xtu(t)dt, x \in [0, 1]. \quad (2)$$

Calcul détaillé de la solution exacte

Posons $A = \int_0^1 tu(t)dt$. L'équation (2) devient $u(x) = e^x + Ax$. En substituant :

$$A = \int_0^1 t(e^t + At)dt = \int_0^1 te^t dt + A \int_0^1 t^2 dt.$$

Calcul de $\int_0^1 te^t dt$ par intégration par parties ($u = t, dv = e^t dt$):

$$\int_0^1 te^t dt = [te^t]_0^1 - \int_0^1 e^t dt = e - [e^t]_0^1 = e - (e - 1) = 1.$$

Donc :

$$A = 1 + \frac{A}{3} \Rightarrow \frac{2A}{3} = 1 \Rightarrow A = \frac{3}{2}.$$

Solution exacte :

$$u(x) = e^x + \frac{3}{2}x.$$

Construction du système de collocation

On cherche $u_h(x) = c_0 + c_1x + c_2x^2$ dans V_2 .

Points de collocation : $x_0 = 0, x_1 = \frac{1}{2}, x_2 = 1$.

Le résidu est :

$$R(x) = c_0 + c_1x + c_2x^2 - e^x - x\left(\frac{c_0}{2} + \frac{c_1}{3} + \frac{c_2}{4}\right).$$

La condition $R(x_i) = 0$ donne le système $A_{col}C = B_{col}$:

$$A_{col} = \begin{pmatrix} 1 & 0 & 0 \\ \frac{3}{4} & \frac{1}{3} & \frac{3}{16} \\ \frac{1}{2} & \frac{2}{3} & \frac{3}{4} \end{pmatrix}, \quad B_{col} = \begin{pmatrix} e^0 \\ e^{\frac{1}{2}} \\ e^1 \end{pmatrix} = \begin{pmatrix} 1 \\ 1.648721 \\ 2.718282 \end{pmatrix}.$$

Pour $i = 0$ ($x_0 = 0$) :

$$c_0 = 1$$

Pour $i = 1$ ($x_1 = \frac{1}{2}$) :

$$\frac{3}{4}(1) + \frac{c_1}{3} + \frac{3c_2}{16} = e^{1/2} \Rightarrow \frac{c_1}{3} + \frac{3c_2}{16} = 0.898721. \quad (E_1)$$

Pour $i = 2$ ($x_2 = 1$) :

$$\frac{1}{2}(1) + \frac{2c_1}{3} + \frac{3c_2}{4} = e \Rightarrow \frac{2c_1}{3} + \frac{3c_2}{4} = 2.218282. \quad (E_2)$$

De (E_1) : $c_1 = 3(0.898721 - \frac{3c_2}{16}) = 2.696163 - \frac{9c_2}{16}$.

En substituant dans (E_2) :

$$\frac{2}{3} \left(2.696163 - \frac{9c_2}{16} \right) + \frac{3c_2}{4} = 2.218282.$$

$$1.797442 - \frac{3c_2}{8} + \frac{3c_2}{4} = 2.218282 \Rightarrow \frac{3c_2}{8} = 0.420840 \Rightarrow c_2 = 1.122240.$$

(après arrondi numérique : $c_2 \approx 1.122240$, $c_1 \approx 2.064903$.)

Solution par collocation :

$$c_0 = 1.000000, \quad c_1 = 2.064903, \quad c_2 = 1.122240.$$

$$u_h^{Col}(x) \approx 1.000000 + 2.064903x + 1.122240x^2.$$

Construction du système de Galerkin

La matrice A est la même que dans l'Exemple 1 :

$$A = \begin{pmatrix} \frac{3}{4} & \frac{1}{3} & \frac{5}{24} \\ \frac{1}{3} & \frac{2}{9} & \frac{1}{6} \\ \frac{5}{24} & \frac{1}{6} & \frac{11}{80} \end{pmatrix}.$$

Le second membre change :

$$B_i = \int_0^1 e^x x^i dx.$$

Calcul de B_0 :

$$B_0 = \int_0^1 e^x dx = e - 1 \approx 1.718282..$$

Calcul de B_1 (intégration par parties) :

$$B_1 = \int_0^1 x e^x dx = [x e^x - e^x]_0^1 = 1 \approx 1.000000.$$

Calcul de B_2 (intégration par parties deux fois) :

$$B_2 = \int_0^1 x^2 e^x dx = [x^2 e^x]_0^1 - 2 \int_0^1 x e^x dx = e - 2 \approx 0.718282.$$

Donc :

$$B = \begin{pmatrix} e - 1 \\ 1 \\ e - 2 \end{pmatrix} \approx \begin{pmatrix} 1.718282 \\ 1.000000 \\ 0.718282 \end{pmatrix}.$$

La résolution de $AC = B$ donne :

Solution par Galerkin :

$$c_0 = 1.00937, c_1 = 2.38182, c_2 = 0.80311.$$

$$u_h^{Gal}(x) \approx 1.00937 + 2.38182x + 0.80311x^2.$$

Analyse des erreurs

La solution exacte $u(x) = e^x + \frac{3}{2}x \notin V_2$, donc une erreur d'approximation apparaît.

Erreur de collocation :

$$u(x) - u_h^{Col}(x) = (e^x + \frac{3}{2}x) - (1.000 + 2.064903x + 1.122240x^2).$$

$$E_{L^2}^{Col} \approx 6.790 \times 10^{-2}, \quad E_{L^\infty}^{Col} \approx 8.571 \times 10^{-2}.$$

Erreur de Galerkin

$$u(x) - u_h^{Gal}(x) = (e^x + \frac{3}{2}x) - (1.00937 + 2.38182x + 0.80311x^2).$$

$$E_{L^2}^{Gal} \approx 1.342 \times 10^{-3}, \quad E_{L^\infty}^{Gal} \approx 2.398 \times 10^{-2}.$$

Méthode	c_0	c_1	c_2
Solution exacte $(e^x + \frac{3}{2}x)$	1.000	2.500	0.500(<i>Taylor</i>)
Collocation	1.000000	2.064903	1.122240
Galerkin	1.00937	2.38182	0.80311

Interprétation.

·La méthode de Galerkin minimise l'erreur en norme $L^2(1.342 \times 10^{-3})$ par rapport à la collocation (6.790×10^{-2}) , soit un facteur ≈ 2 .

·La collocation impose l'exactitude aux seuls points $\{0, \frac{1}{2}, 1\}$, d'où une erreur plus grande entre ces points.

·La régularité de e^x permet aux polynômes de faible degré de fournir une bonne approximation dans les deux cas.

3.4.3 Exemple 3 : Solution trigonométrique**Énoncé**

Résoudre l'équation intégrale de Fredholm du second type :

$$u(x) = \sin(x) + \int_0^1 xtu(t)dt, \quad x \in [0, 1]. \quad (3)$$

Calcul détaillé de la solution exacte

Posons $A = \int_0^1 tu(t)dt$. L'équation (3) devient $u(x) = \sin(x) + Ax$.

En substituant :

$$A = \int_0^1 t[\sin(t) + At]dt = \int_0^1 t\sin(t)dt + A \int_0^1 t^2dt.$$

Calcul de $\int_0^1 t\sin(t)dt$ par intégration par parties ($u = t, dv = \sin t dt$) :

$$\int_0^1 t\sin(t)dt = [-t \cos t]_0^1 + \int_0^1 \cos(t)dt = -\cos(1) + [\sin t]_0^1 = \sin(1) - \cos(1).$$

Donc :

$$A = \sin(1) - \cos(1) + \frac{A}{3} \Rightarrow \frac{2A}{3} = \sin(1) - \cos(1) \Rightarrow A = \frac{3}{2}(\sin(1) - \cos(1)).$$

Numériquement : $\sin(1) \approx 0.841471$, $\cos(1) \approx 0.540302$, donc $A \approx \frac{3}{2}(0.841471 - 0.540302) = \frac{3}{2}(0.301169) \approx 0.451753$.

Solution exacte :

$$u(x) = \sin(x) + \frac{3}{2}(\sin(1) - \cos(1))x \approx \sin(x) + 0.451753x.$$

Construction du système de collocation

On cherche $u_h(x) = c_0 + c_1x + c_2x^2$.

La matrice A_{col} est identique à l'Exemple 2. Le second membre change :

$$B_{col} = \begin{pmatrix} \sin(0) \\ \sin(\frac{1}{2}) \\ \sin(1) \end{pmatrix} = \begin{pmatrix} 0 \\ 0.479426 \\ 0.841471 \end{pmatrix}.$$

Pour $i = 0$ ($x_0 = 0$) :

$$c_0 = \sin(0) = 0.$$

Pour $i = 1$ ($x_1 = \frac{1}{2}$) :

$$\frac{c_1}{3} + \frac{3c_2}{16} = \sin(1/2) = 0.479426. \quad (E_1)$$

Pour $i = 2$ ($x_2 = 1$) :

$$\frac{2c_1}{3} + \frac{3c_2}{4} = \sin(1) = 0.841471. \quad (E_2)$$

$$\text{De } (E_1) : c_1 = 3(0.479426 - \frac{3c_2}{16}) = 1.438279 - \frac{9c_2}{16}.$$

En substituant dans (E_2) :

$$\frac{3}{2}(1.438279 - \frac{9c_2}{16}) + \frac{3c_2}{4} = 0.841471.$$

$$0.958853 - \frac{3c_2}{8} + \frac{3c_2}{4} = 0.841471 \Rightarrow \frac{3c_2}{8} = -0.117382 \Rightarrow c_2 \approx -0.312973.$$

$$\text{Et } c_1 = 1.438279 - \frac{9(-0.312973)}{16} \approx 1.614260.$$

(Les valeurs exactes issues de la résolution numérique sont données ci-dessous.)

Solution par collocation :

$$c_0 = 0.000000, c_1 = 1.614260, c_2 = -0.312973.$$

$$u_h^{Col}(x) \approx 1.614260x - 0.312973x^2.$$

Construction du système de Galerkin

La matrice A est identique. Le second membre :

$$B_i = \int_0^1 \sin(x) x^i dx.$$

Calcul de B_0 :

$$B_0 = \int_0^1 \sin(x) dx = [-\cos x]_0^1 = 1 - \cos(1) \approx 0.459698.$$

Calcul de B_1 (intégration par parties) :

$$B_1 = \int_0^1 x \sin(x) dx = [-x \cos x]_0^1 + \int_0^1 \cos(x) dx = -\cos(1) + \sin(1) = \sin(1) - \cos(1) \approx 0.301169.$$

Calcul de B_2 (intégration par parties deux fois) :

$$B_2 = \int_0^1 x^2 \sin(x) dx = [-x^2 \cos x]_0^1 + 2 \int_0^1 x \cos(x) dx.$$

$$\int_0^1 x \cos(x) dx = [x \sin x]_0^1 - \int_0^1 \sin(x) dx = \sin(1) - (1 - \cos(1)) = \sin(1) + \cos(1) - 1.$$

$$B_2 = -\cos(1) + 2(\sin(1) + \cos(1) - 1) = 2\sin(1) + \cos(1) - 2 \approx 0.222845.$$

Donc :

$$B \approx \begin{pmatrix} 0.459698 \\ 0.301169 \\ 0.222845 \end{pmatrix}.$$

Solution par Galerkin :

$$c_0 = -0.00653, c_1 = 1.53736, c_2 = -0.22972.$$

$$u_h^{Gal}(x) \approx -0.00653 + 1.53736x - 0.22972x^2.$$

Analyse des erreurs

Erreur de collocation :

$$u(x) - u_h^{Col}(x) = \sin(x) + 0.451753x - (1.614260x - 0.312973x^2).$$

$$E_{L^2}^{Col} \approx 1.800 \times 10^{-2}, E_{L^\infty}^{Col} \approx 2.368 \times 10^{-2}.$$

Erreur de Galerkin :

$$u(x) - u_h^{Gal}(x) = \sin(x) + 0.451753x - (-0.00653 + 1.53736x - 0.22972x^2).$$

$$E_{L^2}^{Gal} \approx 4.900 \times 10^{-3}, E_{L^\infty}^{Gal} \approx 7.890 \times 10^{-3}.$$

Méthode	c_0	c_1	c_2
Solution exacte (approx.)	0	1.292	$-0.167(Taylor)$
Collocation	0.00000	1.614260	-0.312973
Galerkin	-0.00653	1.53736	-0.22972

Interprétation.

·Le terme trigonométrique $\sin(x)$ rend l'approximation polynomiale plus délicate que dans l'exemple exponentiel.

·La méthode de **Galerkin** conserve sa supériorité en norme $L^2(4.900 \times 10^{-3} vs 1.800 \times 10^{-2})$, soit un facteur ≈ 2 .

·La **collocation** peut produire une erreur oscillatoire entre les points choisis, phénomène réduit par la formulation globale de Galerkin.

3.4.4 Comparaison globale

Tableau comparatif des erreurs

Exemple	Solution	Collocation		Galerkin	
		E_{L^2}	E_{L^∞}	E_{L^2}	E_{L^∞}
Ex. 1	$1 + \frac{9}{4}x(polynomiale)$	0	0	0	0
Ex. 2	$e^x + \frac{3}{2}x(exponentielle)$	6.790×10^{-2}	8.571×10^{-2}	1.342×10^{-2}	2.398×10^{-2}
Ex. 3	$\sin(x) + Ax(trigonometrique)$	1.800×10^{-2}	2.368×10^{-2}	4.900×10^{-3}	7.890×10^{-3}

Discussion générale

Comportement numérique des deux méthodes.

1. Lorsque la solution appartient à V_n (Exemple 1) : les deux méthodes reproduisent la solution exacte sans aucune erreur. Cela valide la cohérence des deux formulations.

2. Pour les solutions non polynomiales (Exemples 2 et 3) :

·La méthode de **Galerkin** fournit une erreur L^2 environ **deux fois plus petite** que la collocation, grâce à la condition d'orthogonalité du résidu qui distribue l'erreur de façon optimale sur tout le domaine.

·La méthode de **collocation** impose l'exactitude aux seuls points de collocation ; l'erreur peut être plus grande entre ces points .

3. Stabilité : Galerkin est plus stable car sa formulation variationnelle minimise globalement le résidu, tandis que la collocation est sensible au choix des points.

4. Coût de calcul : la collocation est plus simple à mettre en oeuvre (pas de calcul d'intégrales doubles), ce qui en fait un choix attractif pour des problèmes de grande taille.

5. Convergence : les deux méthodes convergent lorsque le degré n de l'espace V_n augmente ; les erreurs diminuent progressivement vers zéro.

Conclusion

Dans ce mémoire, nous nous sommes intéressés à l'étude comparative entre la méthode de collocation et la méthode de Galerkin pour la résolution des équations intégrales. Après avoir présenté les notions théoriques nécessaires concernant les équations intégrales et certains espaces fonctionnels, nous avons étudié les principes fondamentaux de chacune des deux méthodes numériques.

La méthode de collocation s'est révélée simple à mettre en œuvre et efficace pour obtenir des solutions approchées avec un coût de calcul relativement réduit. De son côté, la méthode de Galerkin présente une base théorique solide et offre généralement une meilleure stabilité et une précision satisfaisante dans l'approximation des solutions.

L'étude comparative réalisée à travers différents exemples nous a permis de mettre en évidence les avantages et les limites de chaque méthode. Le choix de la méthode la plus adaptée dépend essentiellement de la nature du problème étudié, de la précision recherchée ainsi que des moyens de calcul disponibles.

Bibliographie

- [1] LAKHAL. Aissa. Etude du point fixe de Kannan sur les équations fonctionnelles, Doctorat, université de M'sila, 2025.
- [2] R. Kress, Linear Integral Equations, 3rd Edition, Springer, New York, 2014.
- [3] K. E. Atkinson, The Numerical Solution of Integral Equations of the Second Kind, Cambridge University Press, Cambridge, 1997.
- [4] E. Kreyszig, Introductory Functional Analysis with Applications, John Wiley & Sons, New York, 1978.
- [5] A. Quarteroni, R. Sacco and F. Saleri, Numerical Mathematics, 2nd Edition, Springer, Berlin, 2007.
- [6] M. NADIR .Cours d'analyse fonctionnelle ,université de M'sila Algérie, 2004.
- [7] M. N. NADIR, Sur la solution numérique des équations intégrales de Volterra- Fredholm en utilisant les polynômes de Chebyshev, Master, université de M'sila, 2022.
- [8] P. J. Davis and P. Rabinowitz, Methods of Numerical Integration, 2nd Edition, Academic Press, New York, 1984.
- [9] C. T. H. Baker, The Numerical Treatment of Integral Equations, Oxford University Press, Oxford, 1977.
- [10] F. B. Hildebrand, Introduction to Numerical Analysis, 2nd Edition, Dover Publications, New York, 1987.

- [11] Wazwaz, A. M. (2011). Linear and nonlinear integral equations (Vol. 639, pp. 35-36). Berlin: Springer.

Résumé

Ce mémoire présente une étude comparative entre deux méthodes numériques importantes utilisées dans la résolution des équations intégrales : la méthode de collocation et la méthode de Galerkin. Après avoir introduit les notions fondamentales des équations intégrales et des espaces fonctionnels nécessaires, nous présentons les bases théoriques des deux méthodes, puis nous réalisons une comparaison numérique sur des exemples tests. Les résultats montrent les avantages et limites de chaque méthode selon la nature du problème étudié.

Mots clés : Polynômes orthogonaux , équation intégrale de Volterra-Fredholm , méthode de collocation , méthode de Galerkin

Abstract

This work presents a comparative study between the collocation method and the Galerkin method for solving integral equations. After introducing the theoretical framework, numerical experiments are conducted to compare accuracy and computational efficiency.

Keywords : orthogonal Polynomials , Volterra-Fredholm integral equation , collocation method , Galerkin method