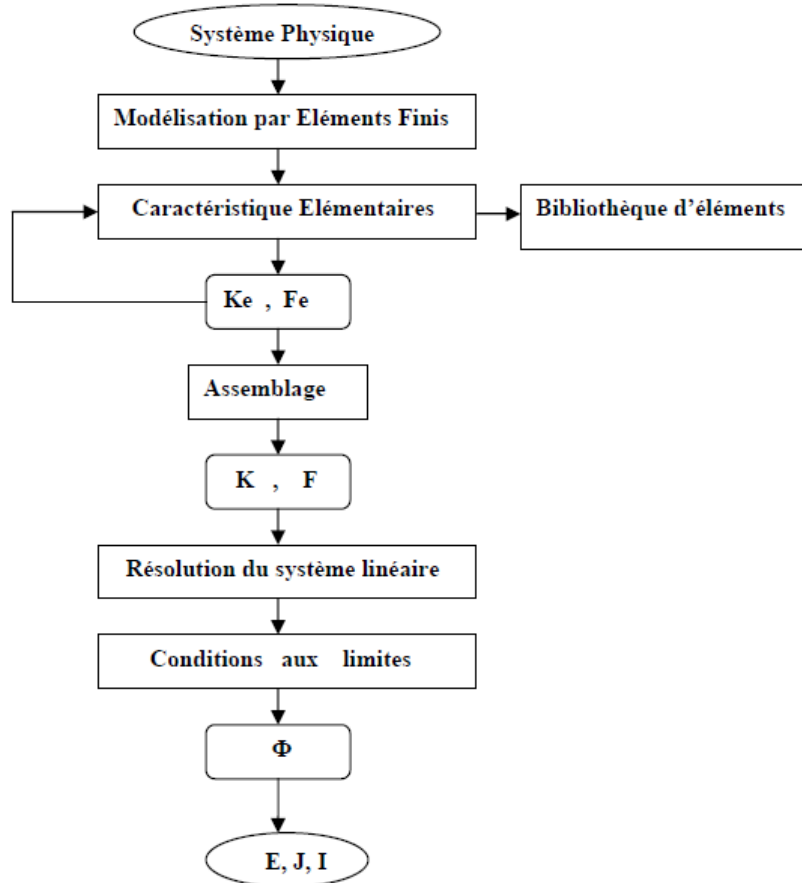


Polycopié de Cours Méthodes Numériques Appliquées 1ère Année Master: Electromécanique



SOMMAIRE

AVANT – PROPOS

CHAPITRE I: Rappels de quelques méthodes numériques

I.1. Introduction.....	01
I.2. Résolution des systèmes d'équations linéaires et non linéaires.....	01
I.2. 1. Résolution des systèmes d'équations linéaires.....	02
I.2.1.a. Méthodes de Jacobi.....	02
I.2.1.b. Méthodes de Gauss-Seidel.....	03
I.2.1.c. Méthode de Newton–Raphson.....	04
I.2. 2. Résolution des systèmes d'équations non linéaires.....	05
I.3. Interpolation et approximation.....	07
I.3.1 Rappel et définitions	07
I.3.1.a. Définitions.....	08
I.3.1.b. Méthode de Lagrange.....	08
I.3.1.c. Méthode des différences divisées.....	10

I.4. Intégration numérique.....	11
I.4.1. Méthode des Trapèzes.....	11
I.4.2. Méthode de Simpson.....	12
I.4.3. Méthode composée des Trapèzes	13
I.5. Résolution des équations différentielles ordinaires.....	14
I.5.1. Classification suivant les conditions.....	16
I.5.1.a. Equations différentielles aux conditions initiales.....	16
I.5.1.b. Equations différentielles aux conditions aux limites.....	17
I.5.2. Méthode d'Euler.....	18
I.5.2.a. Algorithme d' <i>Euler</i> général.....	18
I.5.2.b. Algorithme d'Euler modifiée.....	19
I.5.3. Méthode de Runge-Kutta.....	19
I.5.3.a. Méthode de Runge-Kutta d'ordre 2.....	19
I.5.3.b. Méthode de Runge-Kutta d'ordre 4.....	20
I.5.4. Méthode d'Adams.....	21
I.5.4.a. Méthodes d'Adams - Bashforth.....	21
I.5.4.b. Méthodes d'Adams-Moulton.....	22
I.6. Conclusion.....	23

Chapitre II : Résolution des équations aux dérivées partielles

II.1. Introduction.....	24
II.2. Définition	24
II.3. Classification des équations aux dérivées partielles.....	25
II.4. Rappels sur la classification des systèmes physiques.....	26
II.4.1. Système continu.....	26
II.4.2. Système discrets.....	26
II.4.3. Choix du maillage.....	27
II.5. Différentes méthodes de résolution des équations aux dérivées partielles.....	28
II.5.1. Méthode des Différences Finies (MDF).....	29
II.5.1.a. Avantages de la MDF.....	34
II.5.1.b. Inconvénients de la MDF.....	34
II.5.2. Méthode des éléments Finis (MEF).....	34
II.5.2.1. Principe de la method.....	39
II.5.2.2 Domaines d'application.....	43
II.6 Types de problèmes MEF	43
II. 6.1 Problèmes d'équilibre stationnaire.....	44

II. 6.2 Problèmes aux valeurs propres	44
II. 6.3 Problèmes dépendant du temps.....	45
II. 7 Type des éléments finis	45
II.8. Les différentes étapes de la résolution par la MEF.....	45
II.8.1. Définition des mailles de référence.....	48
II.8.2. Interpolation polynomiale sur les mailles de référence.....	53
II.8.3. Formulations matricielles de la MEF.....	55
II.8.4. Assemblage.....	58
II.8.5. Introduction des conditions aux limites.....	59
II.8.6. Conditions de convergence de la solution.....	59
II.9. Avantages de la MEF.....	60
II.10. Inconvénients de la MEF.....	61
II.11. Conclusion.....	61

CHAPITRE III: Techniques d'optimisation

III.1. Introduction.....	62
--------------------------	----

III.2. Définition et formulation.....	62
III.2.1. Notions de minimum, maximum, infimum, supremum.....	62
III.2.2. Description d'un problème d'optimisation.....	64
III.2.3. Types d'optimisation.....	66
III.2.4. Algorithme d'optimisation.....	66
III.3. Optimisation sans contraintes.....	69
III.4. Optimisation sous contraintes.....	71
III.5. Conclusion.....	74

Travaux pratiques

TP N° 01 : Introduction au MATLAB.....	75
TP N° 02 : Résolution d'une équation différentielle par MEF.....	78
TP N° 03 : Analyse de Structure d'un Système Electromagnétique (Cas d'une Machine à courant continu)	93
TP N° 04 : Analyse de Structure d'un Système Electromagnétique	

(Transformateur).....	97
Références.....	99

AVANT – PROPOS

Ce Module est réservé pour la présentation des différentes méthodes et techniques numériques dans le but de faire :

- Exposer et acquérir un ensemble de méthodes numériques permettant de résoudre des problèmes impossibles par des approches analytiques.*
- Savoir transposer la connaissance mathématique pure à un ordinateur aux performances finies.*
- Résoudre numériquement des problèmes dont la solution analytique est connue ou non.*
- Analyser le comportement des méthodes.*

Le contenu de ce document est destiné aux étudiants en Master de ELM ainsi à toute les personnes intéressées par l'analyse numérique au département de génie électrique à l'université de mohamed boudiaf de m'sila. Il s'inspire de nombreux ouvrages et références bien plus complets, ainsi que divers documents de collègues universitaires. Ce document est bien sur incomplet : il manque des chapitres entiers, des démonstrations, des exemples, etc. Toute remarque est la bienvenue, même en ce qui concerne les probablement nombreuses fautes d'orthographes.

Pour plus de détaille et d'aprofondement veuillez revoir les références.

I.1. Introduction

Le module « méthodes numériques appliquées », rassemble toutes les techniques de calcul qui permettent de résoudre de manière exacte ou, le plus souvent, de manière approchée un problème physique représenté par un modèle mathématique.

I.2. Résolution des systèmes d'équations linéaires et non linéaires

Soit $M_n(IR)$, l'espace vectoriel des matrices carrées d'ordre n et à coefficients réels. Soit $A \in M_n(IR)$ et $B \in IR^n$, on cherche le vecteur $X \in IR^n$, solution du système linéaire $AX = B$. Ce système admet une solution unique lorsque le déterminant de A est non nul, ce que nous supposons dans la suite.

Remarquons que la résolution de ce système à l'aide des formules de Cramer est impraticable lorsque n est 'grand', car ces formules nécessitent approximativement ' $n !.n^2$ ' opérations arithmétiques élémentaires ; par exemple si $n=15$, cela représente de l'ordre de 4 mois de calcul pour un ordinateur moyen effectuant opérations à la seconde. Ce temps de calcul évolue de manière exponentielle avec n .

On sépare généralement les problèmes en deux classes suivant les caractéristiques de la matrice A :

- La matrice A est de taille réduite et est pleine. Par taille réduite, on comprend les matrices d'ordre inférieur à 100, par matrice pleine, on signifie que A comporte peu d'éléments nuls ;
- La matrice A est éparse (creuse) et de grande taille. De grande taille nomme des matrices d'ordres égaux à plusieurs centaines ou plusieurs milliers. La matrice éparse possède peu d'éléments non nuls.

Il est à remarquer qu'il n'existe pas de règle définitive pour le choix entre les méthodes directes et les méthodes itératives.

Une méthode directe conduit à une solution en un nombre fini d'étapes, et sans les erreurs arrondie.

- Une méthode itérative fait passer d'un estimé $X^{(k)}$ de la solution à un autre estimé $X^{(k+1)}$ de cette solution. S'il ya convergence, la solution ne pourrait donc être atteinte qu'après un nombre infini d'itérations.

Ce premier chapitre est consacré aux méthodes employées pour résoudre les problèmes les plus fréquemment rencontrées pour les systèmes linéaires. Par définition, on peut écrire celles-ci comme : $A.X = b$, où A est une matrice de l'ordre de $M \times N$:

Un système de m équations à n inconnues $x_1, x_2, \dots, x_n$ s'écrit sous forme matricielle : $AX = B$ où A est une matrice comportant m lignes et n colonnes,

X est le vecteur colonne dont les composantes sont les x_i et B , le second membre, est aussi un vecteur colonne avec n composantes. Le vecteur X est appelé solution du système

$$A = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1N} \\ a_{21} & a_{22} & \dots & a_{2N} \\ \dots & \dots & \dots & \dots \\ a_{M1} & a_{M2} & \dots & a_{MN} \end{bmatrix}; \quad b = \begin{bmatrix} b_1 \\ b_2 \\ \dots \\ b_N \end{bmatrix}; \quad X = \begin{bmatrix} x_1 \\ x_2 \\ \dots \\ x_M \end{bmatrix}$$

Si $N = M$, il y a autant d'équations que d'inconnues et si aucune des équations n'est une combinaison linéaire des $N-1$ autres, la solution existe et est unique

I.2. 1. Résolution des systèmes d'équations linéaires

I.2.1.a. Méthodes de Jacobi

Pour les systèmes linéaires, cette méthode illustre bien les techniques de point fixe : si on suppose connues toutes les inconnues sauf l'inconnue i , cette dernière peut être déterminée à

l'aide de l'équation i. Procédant ainsi sur toutes les inconnues, on obtient un nouvel itéré x^k à partir de l'ancien x^{k-1} . Cependant, comme la méthode de Gauss-Seidel, elle peut présenter un intérêt dans le cas de systèmes non-linéaires. Sous forme matricielle, cela se traduit par :

$x^k = (I - D^{-1}A) x^{k-1} + D^{-1}b$ où D est la matrice diagonale, contenant la diagonale de A . La matrice d'itération est donc ici

$B = I - D^{-1}A$. On peut noter l'utilisation d'un vecteur de stockage $\bar{x}$ supplémentaire.

Algorithme :

```

- Vecteur de départ  $x^{(0)}$ 
  Tant que  $R > \varepsilon$ 
    itération sur  $k$ 
      boucle sur  $i = 1, n$ 
         $\bar{x}_i \leftarrow A_{i,1:i-1}x_{1:i-1}^{k-1} + A_{i,i+1:n}x_{i+1:n}^{k-1}$ 
         $\bar{x}_i \leftarrow (b_i - \bar{x}_i) / A_{i,i}$ 
      fin
       $x^{(k)} \leftarrow \bar{x}$ 
  Fin

```

Convergence :

Si A est symétrique, définie positive, alors :

$$\rho_J = \rho(B) < 1,$$

I.2.1.b. Méthodes de Gauss-Seidel

Le concept est semblable à celui de la méthode de Jacobi, à ceci près que l'on utilise les nouvelles valeurs estimées dès qu'elles sont disponibles, et non pas à la fin de l'itération. Sous forme matricielle, cela se traduit par:

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

Où « $-E$ » est la partie triangulaire inférieure de A , et « $-F$ » sa partie triangulaire supérieure. La matrice d'itération est ici :

$$B = (D - E)^{-1}F = (I - D^{-1}E)^{-1}D^{-1}F$$

Algorithme :

```

Vecteur de départ  $x^{(0)}$ 
Tant que  $R > \varepsilon$ 
    itération sur  $k$ 
    boucle sur  $i = 1, 2, \dots, n$ 
         $x_i^{k-1} \leftarrow 0$ 
         $\beta \leftarrow A_{i,ln} x^{(k-1)}$ 
         $x_i^k \leftarrow (b_i - \beta) / A_{i,i}$ 
    fin
Fin

```

Convergence :

Quand A est symétrique, définie positive, on est assuré de la convergence puisque $\rho_{GS} < 1$.

Mais l'estimation de la valeur du rayon spectral est difficile car il dépend en particulier de la numérotation des degrés de liberté employée.

I.2.1.c. Méthode de Newton–Raphson

En supposant la fonction f est de classe C^1 et que le zéro ξ est simple, la méthode de Newton–Raphson fait le choix

$$\alpha(x) = \frac{1}{f'(x)}$$

La relation de récurrence définissant cette méthode est alors

$$\forall k \in \mathbb{N}, x^{(k+1)} = x^{(k)} - \frac{f(x^{(k)})}{f'(x^{(k)})},$$

L'initialisation $x(0)$ étant donnée. Dans cette méthode, toute nouvelle approximation du zéro est construite au moyen d'une linéarisation de l'équation $f(x) = 0$ autour de l'approximation précédente. En effet, si l'on remplace $f(x)$ au voisinage du point $x^{(k)}$ par l'approximation affine obtenue en tronquant au premier ordre le développement de Taylor de f en $x^{(k)}$ et qu'on résout l'équation linéaire résultante :

$$f(x^{(k)}) + (x - x^{(k)})f'(x^{(k)}) = 0,$$

en notant sa solution $x^{(k+1)}$, et il en résulte que, géométriquement parlant, le point $x^{(k+1)}$ est l'abscisse du point d'intersection entre la tangente à la courbe de f au point $(x^{(k)}; f(x^{(k)}))$ et l'axe des abscisses

Par rapport à toutes les méthodes introduites jusqu'à présent, on pourra remarquer que la méthode de Newton nécessite à chaque itération l'évaluation des deux fonctions f et f' au point courant $x^{(k)}$. Cet effort est compensé par une vitesse de convergence accrue, puisque cette méthode est d'ordre deux si le zéro recherché est simple.

I.2. 2. Résolution des systèmes d'équations non linéaires

La méthode Newton peut s'appliquer à la résolution d'un système de plusieurs équations non linéaires :

$$\begin{cases} f(x, y) = 0 \\ g(x, y) = 0 \end{cases}$$

A partir d'un couple de valeurs approchées (x_1, y_1) d'une solution du système, on peut déterminer deux accroissements h et k à donner à x_1 et y_1 de manière à ce que :

$$\begin{cases} f(x_1 + h, y_1 + k) = 0 \\ g(x_1 + h, y_1 + k) = 0 \end{cases}$$

En développant en premier ordre, on obtient :

$$\begin{cases} f(x_1 + h, y_1 + k) = f(x_1, y_1) + hf'(x_1) + kf'(y_1) = 0 \\ g(x_1 + h, y_1 + k) = g(x_1, y_1) + hg'(x_1) + kg'(y_1) = 0 \end{cases}$$

Où l'on a posé :

$$\begin{cases} f'(x_1) = \frac{\partial f(x_1, y_1)}{\partial x} \\ f'(y_1) = \frac{\partial f(x_1, y_1)}{\partial y} \end{cases} \text{ et } \begin{cases} g'(x_1) = \frac{\partial g(x_1, y_1)}{\partial x} \\ g'(y_1) = \frac{\partial g(x_1, y_1)}{\partial y} \end{cases}$$

Les quantités h et k s'obtiennent donc, en résolvant le système linéaire suivant :

$$\begin{cases} hf'(x_1) + kf'(y_1) = -f(x_1, y_1) \\ hg'(x_1) + kg'(y_1) = -g(x_1, y_1) \end{cases} \text{ avec } \begin{cases} h = x_{n+1} - x_n \\ k = y_{n+1} - y_n \end{cases}$$

Le calcul est alors relancé jusqu'à ce que h et k deviennent inférieurs à une valeur ε que l'on se donne (selon la précision voulue pour le calcul).

Ainsi, l'algorithme correspondant est :

$$\begin{cases} x_{n+1} = x_n - \left(\frac{f(x_n, y_n) \cdot g'(y_k) - g(x_n, y_n) \cdot f'(y_k)}{\Delta} \right) \\ y_{n+1} = y_n - \left(\frac{g(x_n, y_n) \cdot f'(y_k) - f(x_n, y_n) \cdot g'(y_k)}{\Delta} \right) \end{cases}$$

Avec :

$$\Delta = f'(x_n) \cdot g'(y_n) - f'(y_n) \cdot g'(x_n)$$

I.3. Interpolation et approximation

L'intérêt de remplacer une fonction quelconque par un polynôme l'approchant aussi précisément que voulu sur un intervalle donné est évident d'un point de vue numérique et informatique, puisqu'il est très aisé de stocker et de manipuler, c'est-à-dire additionner, multiplier, dériver ou intégrer, des polynômes dans un calculateur. Pour ce faire, il semble naturel de chercher à utiliser un polynôme d'interpolation de Lagrange associé aux valeurs prises par la fonction en des nœuds choisis.

I.3.1 Rappel et définitions

Soit $P_n(x)$ l'espace des polynômes de degré inférieur ou égal à n .

On rappelle que

$$\left\{ 1, x, x^2, \dots, x^n \right\} \text{ est une base de } P_n(x) \text{ (dim } P_n(x) = n + 1)$$

On note δ_{ij} le symbole de Kronecker : $\delta_{ij} = 1$ si $i = j$; $\delta_{ij} = 0$ si $i \neq j$.

Soient f une fonction continue sur $[a, b]$, $x_1, \dots, x_n$ n points de $[a, b]$ et $g_1, \dots, g_n$ des réels de même signe.

Alors on a :

$$\sum_{i=1}^{i=n} f(x_i)g_i = f(c) \sum_{i=1}^{i=n} g_i \quad \text{où } c \in [a, b]$$

I.3.1.a. Définitions

Soit f une fonction réelle définie sur un intervalle $[a, b]$ contenant $n + 1$ points distincts

$x_0, x_1, \dots, x_n$. Soit P_n un polynôme de degré inférieur ou égal à n .

On dit que P_n est un interpolant de f ou interpole f en $x_0, x_1, \dots, x_n$ si :

$$P_n(x_i) = f(x_i) \text{ pour } 0 \leq i \leq n$$

I.3.1.b. Méthode de Lagrange

Soient $x_0, x_1, \dots, x_n$ $(n + 1)$ points deux à deux distincts d'un intervalle $[a, b]$ de $\mathbb{R}$

On appelle interpolant de Lagrange les polynômes L_i définis pour $i = 0, \dots, n$ par :

$$L_i(x) = \prod_{j=0, j \neq i}^{j=n} \frac{(x - x_j)}{(x_i - x_j)} = \frac{(x - x_0)(x - x_1) \dots (x - x_{i-1})(x - x_{i+1}) \dots (x - x_n)}{(x_i - x_0)(x_i - x_1) \dots (x_i - x_{i-1})(x_i - x_{i+1}) \dots (x_i - x_n)}$$

On a en particulier :

$$L_0(x) = \prod_{j=1}^{j=n} \frac{(x - x_j)}{(x_0 - x_j)} = \frac{(x - x_1)(x - x_2) \dots (x - x_i) \dots (x - x_n)}{(x_0 - x_1)(x_0 - x_2) \dots (x_0 - x_i) \dots (x_0 - x_n)}$$

$$L_n(x) = \prod_{j=0}^{j=n-1} \frac{(x - x_j)}{(x_n - x_j)} = \frac{(x - x_0)(x - x_1).....(x - x_i).....(x - x_{n-1})}{(x_n - x_0)(x_n - x_1)....(x_n - x_i)....(x_n - x_{n-1})}$$

Si on prend :

$$P_n(x) = L_0(x)f(x_0) + L_1(x)f(x_1) + + L_n(x)f(x_n)$$

Alors :

$$P_n(x_i) = f(x_i) \text{ pour } 0 \leq i \leq n$$

Exemple :

Si $x_0 = -1$, $x_1 = 0$ et $x_2 = 1$,

$f(x_0) = 2$, $f(x_1) = 1$, $f(x_2) = -1$, on obtient :

$$\begin{aligned} L_0(x) &= \frac{(x - x_1)(x - x_2)}{(x_0 - x_1)(x_0 - x_2)} = \frac{x(x - 1)}{(-1)(-1 - 1)} = \frac{x(x - 1)}{2} \\ L_1(x) &= \frac{(x - x_0)(x - x_2)}{(x_1 - x_0)(x_1 - x_2)} = \frac{(x - (-1))(x - 1)}{-(-1)(-1)} = \frac{(x + 1)(x - 1)}{-1} \\ L_2(x) &= \frac{(x - x_0)(x - x_1)}{(x_2 - x_0)(x_2 - x_1)} = \frac{(x + 1)x}{(1 - (-1))(1 - 0)} = \frac{(x + 1)x}{2} \end{aligned}$$

Alors :

$$\begin{aligned} P_2(x) &= L_0(x)f(x_0) + L_1(x)f(x_1) + L_2(x)f(x_2) \\ &= \frac{x(x - 1)}{2}f(x_0) + \frac{(x + 1)(x - 1)}{-1}f(x_1) + \frac{(x + 1)x}{2}f(x_2) \\ &= 2\frac{x(x - 1)}{2} + \frac{(x + 1)(x - 1)}{-1} - \frac{(x + 1)x}{2} \\ &= -\frac{1}{2}x^2 - \frac{3}{2}x + 1 \end{aligned}$$

On peut facilement vérifier que :

$$\begin{aligned}
 P_2(x_0) &= P_2(-1) = \frac{(-)(-1-1)}{2} = 2 = f(x_0) \\
 P_2(x_1) &= P_2(0) = \frac{(1)(-1)}{-1} = 1 = f(x_1) \\
 P_2(x_2) &= P_2(1) = -\frac{(1+1)1}{2} = -1 = f(x_2)
 \end{aligned}$$

En déduisant donc, que : Les polynômes de Lagrange ont les propriétés suivantes :

- ✓ $L_j(x)$ est un polynôme de degré n ; $\forall j = 0, \dots, n$
- ✓ $L_j(x_j) = 1 \forall j = 0, \dots, n$ et $L_j(x_i) = 0$ pour tout $j \neq i$.
- ✓ la famille $\{L_0(x), L_1(x), \dots, L_n(x)\}$ est une base de $P_n(x)$.

I.3.1.c. Méthode des différences divisées

Cette méthode consiste d'abord à établir des propriétés de continuité et de dérivabilité pour la fonction de la variable réelle x définie par $[x_0; x_1, \dots, x_n; x]f$, où les points $x_0, \dots, x_n$ sont distincts et contenus dans un intervalle $[a; b]$ borné de $\mathbb{R}$ et x appartient à $[a; b]$.

On définit les différences divisées d'ordre i de f aux points (x_i) comme suit :

$$\begin{aligned}
 [f(x_0)] &= f(x_0) \\
 [f(x_0), f(x_1)] &= \frac{f(x_1) - f(x_0)}{x_1 - x_0} \\
 [f(x_0), f(x_1), \dots, f(x_i)] &= \frac{[f(x_1), f(x_2), \dots, f(x_i)] - [f(x_0), f(x_1), \dots, f(x_{i-1})]}{x_i - x_0}
 \end{aligned}$$

Pour $i \geq 2$

Exemple :

Si $x_0 = -1$, $x_1 = 0$ et $x_2 = 1$, $f(x_0) = 2$, $f(x_1) = 1$, $f(x_2) = -1$, on obtient :

$$[f(-1)] = 2$$

$$[f(-1), f(0)] = \frac{1-2}{0-(-1)} = -1$$

$$[f(0)] = 1 \quad [f(-1), f(0), f(1)] = \frac{-2-(-1)}{1-(-1)} = \frac{-1}{2}$$

$$[f(0), f(1)] = \frac{-1-1}{1-0} = -2$$

$$[f(1)] = -1$$

En déduisant donc, que La valeur d'une différence divisée est indépendante de l'ordre des x_i

On a ainsi :

$$[f(x_0), f(x_1)] = \frac{f(x_1) - f(x_0)}{x_1 - x_0} = \frac{f(x_1)}{x_1 - x_0} + \frac{f(x_0)}{x_0 - x_1} = [f(x_1), f(x_0)]$$

$$\begin{aligned} [f(x_0), f(x_1), f(x_2)] &= \frac{f(x_2)}{(x_2 - x_1)(x_2 - x_0)} + \frac{f(x_1)}{(x_1 - x_2)(x_1 - x_0)} + \frac{f(x_0)}{(x_0 - x_1)(x_0 - x_2)} \\ &= [f(x_2), f(x_1), f(x_0)] = [f(x_1), f(x_0), f(x_2)] \end{aligned}$$

D'une manière générale on a :

$$[f(x_0), f(x_1), \dots, f(x_k)] = \sum_{i=0}^{i=k} \frac{f(x_i)}{(x_i - x_0) \dots (x_i - x_{i-1})(x_i - x_{i+1}) \dots (x_i - x_k)}$$

I.4. Intégration numérique

I.4.1. Méthode des Trapèzes :

Soit $x_0 < x_1 < \dots < x_i < x_{i+1} < \dots < x_n$ une subdivision uniforme et $P_1(x)$ un polynôme de degré 1 interpolant f aux points x_i et x_{i+1} de chaque intervalle $[x_i, x_{i+1}]$.

$$[P_1(x_i) = f(x_i) \text{ et } P_1(x_{i+1}) = f(x_{i+1})]$$

En approchant sur chaque sous intervalle $[x_i, x_{i+1}]$, $f(x)$ par $P_1(x)$ on obtient :

$$\begin{aligned} f(x) &\simeq P_1(x) = f(x_i) + [f(x_i), f(x_{i+1})](x - x_i) \\ f(x) &\simeq P_1(x) = f(x_i) + \frac{f(x_{i+1}) - f(x_i)}{x_{i+1} - x_i}(x - x_i) \end{aligned}$$

Et en conséquence :

$$I(f) = \int_a^b f(x)dx \simeq \int_a^b P_1(x)dx = \frac{h}{2}[f(x_i) + f(x_{i+1})]$$

.

I.4.2. Méthode de Simpson

Soit $P_2(x)$ un polynôme de degré 2 vérifiant :

$$P_2(x_i) = f(x_i), P_2(x_{i+1}) = f(x_{i+1}) \text{ et } P_2(x_{i+2}) = f(x_{i+2})$$

En approchant sur chaque sous intervalle $[x_i, x_{i+2}]$, $f(x)$ par $P_2(x)$ on obtient :

$$\begin{aligned} f(x) &\simeq P_2(x) \\ &\simeq f(x_0) + [f(x_i), f(x_{i+1})](x - x_i) + [f(x_i), f(x_{i+1}), f(x_{i+2})](x - x_i)(x - x_{i+1}) \end{aligned}$$

Et en déduisant :

$$I(f) = \int_{x_i}^{x_{i+2}} f(x)dx \simeq \int_{x_i}^{x_{i+2}} P_2(x)dx = \frac{h}{3}[f(x_i) + 4f(x_{i+1}) + f(x_{i+2})]$$

On a

$$\int_{x_i}^{x_{i+2}} P_2(x)dx = I(P_2) + J(P_2) + K(P_2)$$

Où :

$$I(P_2) = \int_{x_i}^{x_{i+2}} f(x_i) dx$$

$$J(P_2) = \int_{x_i}^{x_{i+2}} \frac{f(x_{i+1}) - f(x_i)}{x_{i+1} - x_i} (x - x_i) dx$$

$$K(P_2) = \int_{x_i}^{x_{i+2}} [f(x_i), f(x_{i+1}), f(x_{i+2})](x - x_i)(x - x_{i+1})dx$$

On fait le changement de variables suivant :

$$(x - x_i) = ht \rightarrow dx = hdt$$

Quand :

$$x = x_i \rightarrow t = 0 \text{ et si } x = x_{i+2} \rightarrow t = 2$$

On obtient donc :

$$\begin{aligned}
I(P_2) &= \int_0^2 f(x_i) h dt = 2h f(x_i) \\
J(P_2) &= \frac{f(x_{i+1}) - f(x_i)}{x_{i+1} - x_i} \int_0^2 h^2 t dt \\
&= h[f(x_{i+1}) - f(x_i)] \left[\frac{t^2}{2} \right]_0^2 \\
&= 2h[f(x_{i+1}) - f(x_i)] \\
K(P_2) &= [f(x_i), f(x_{i+1}), f(x_{i+2})] \int_0^2 (x - x_i)(x - x_{i+1}) dx \\
&= [f(x_i), f(x_{i+1}), f(x_{i+2})] \int_0^2 h^3 t(t-1) dt \\
&= \frac{f(x_{i+2}) - 2f(x_{i+1}) + f(x_i)}{2h^2} h^3 \left[\frac{t^3}{3} - \frac{t^2}{2} \right]_0^2 \\
&= [f(x_{i+2}) - 2f(x_{i+1}) + f(x_i)] \frac{h}{3}
\end{aligned}$$

Et enfin :

$$\int_{x_i}^{x_{i+2}} P_2(x) dx = I(P_2) + J(P_2) + K(P_2) = \frac{h}{3} [f(x_i) + 4f(x_{i+1}) + f(x_{i+2})]$$

I.4.3. Méthode composée des Trapèzes :

Pour chercher une approximation de l'intégrale sur tout l'intervalle [a, b], il

Suffit d'écrire :

$$\begin{aligned}
\int_{a=x_0}^{b=x_n} f(x) dx &= \sum_{i=0}^{n-1} \int_{x_i}^{x_{i+1}} f(x) dx \\
&= \sum_{i=0}^{n-1} \left[\frac{h}{2} [f(x_i) + f(x_{i+1})] + \sum_{i=0}^{n-1} -\frac{h^3}{12} f''(\eta_i) \right] \\
&= \frac{h}{2} [f(x_0) + 2f(x_1) + \dots + 2f(x_{n-1}) + f(x_n)] - \frac{h^3}{12} \sum_{i=0}^{n-1} f''(\eta_i) \\
&= \frac{h}{2} [f(x_0) + 2f(x_1) + \dots + 2f(x_{n-1}) + f(x_n)] - \frac{(b-a)}{12} h^2 f''(\eta)
\end{aligned}$$

Avec

$$\eta \in [a, b]$$

I.5. Résolution des équations différentielles ordinaires

La résolution numérique des équations différentielles est éventuellement le domaine de l'analyse numérique où les applications sont les plus nombreuses. Que ce soit en mécanique ou en génie électrique, on aboutit souvent à la résolution d'équations différentielles. On ne sait cependant pas toujours donner une expression exacte de la solution, et il faut alors l'approximer numériquement.

Les équations différentielles nous servent pour modéliser le processus physique (sous forme d'une équation devinette), puis de trouver les expressions mathématiques (les solutions) qui vont décrire le comportement futur de forme d'une équation différentielle (la devinette) et ensuite à trouver l'expression de la réponse et analyser le comportement (résultats de simulation), évitant un effondrement ultérieur du système qui porterait conséquences sur les vies humaines. Très peu d'équations différentielles sont solubles analytiquement. De plus, chaque type d'équation requiert une méthode particulière de résolution. Par suite, la résolution de la plus part des équations différentielles nécessite l'utilisation de méthodes numériques (exp : Euler, Runge-Kutta, Adams Moulton, ...).

Une équation est dite différentielle si elle comporte une ou plusieurs dérivées de la variable à résoudre.

Une équation différentielle est dite ordinaire lorsque la variable à résoudre ne dépend que d'une seule autre variable.

L'ordre d'une équation différentielle est déterminé selon le plus haut degré de dérivée apparaissant dans l'équation. Par exemple une équation différentielle du premier ordre ne contient que la dérivée première et possiblement la fonction elle-même. Une équation du deuxième ordre contient obligatoirement la dérivée seconde et éventuellement la dérivée première et/ou la fonction elle-même.

Une équation différentielle est dite linéaire si elle est formée par une combinaison linéaire de la variable et de ses dérivées.

La forme implicite d'une équation différentielle est: $F(x, y, y') = 0$, Exemple: $2x + y' + 2y = 0$.

La forme explicite d'une équation différentielle est : $F(x, y) = y'$, Exemple: $-2x - 2y = y'$.

La solution à une équation différentielle est une fonction $h(x)$, c'est à dire une expression explicite de la combinaison de la variable x (dans laquelle la dépendance envers la variable y a disparu), et qui substitué à y dans l'équation, satisfait cette équation dans une plage de valeurs de la variable : $a < x < b$.

Exemple1 : Solution explicite

La devinette (l'équation différentielle) est: $x y' = 2y$

Devinez alors une fonction $h(x)$ qui une fois substituée à y vérifie bien cette équation.

Solution :

Bien sûr, pour l'instant, on essaie de deviner un peu au hasard, mais nous verrons que nous sommes ici pour apprendre une méthode systématique de résolution de ce genre d'équations.

Donc si on essaie de substituer : $y = h(x) = x^2$

Nous allons voir que : $y' = h'(x) = 2x$

Et alors : $x(2x) = 2x^2$: Donc, nous dirons que $h(x) = x^2$ est une solution explicite de l'équation.

Exemple2: Solution implicite

Soit à trouver une solution pour l'équation suivante : $y y' = -x$

Solution :

Nous pouvons voir que l'expression suivante est solution implicite de l'équation ci-dessus :

$$h(x, y) = x^2 + y^2 - 1 = 0.$$

En effet, il suffit de réarranger sous la forme $y^2 = 1 - x^2$ et de dériver les deux membres de cette égalité par rapport à la variable x pour voir qu'elle vérifie bien l'équation de départ :

$$2y' y = 0 - 2x$$

I.5.1. Classification suivant les conditions

Les équations différentielles peuvent être classées en deux catégories suivant le types des conditions:

- les équations différentielles aux conditions initiales.
- les équations différentielles aux conditions aux limites.

I.5.1.a. Equations différentielles aux conditions initiales

Soit l'équation différentielle suivante :

$$A(t) \cdot \frac{d^2 y}{dt^2} + B(t) \cdot \frac{dy}{dt} + C(t) \cdot y = g(t)$$

Où A, B, C et g sont des fonctions connues de t. On donne en plus de cette équation différentielle, les conditions initiales suivantes:

$$y(0) = y_0 \text{ et } \left(\frac{dy}{dt} \right)_0 = V_0$$

Toute équation différentielle d'ordre n avec des conditions initiales peut être remplacée par un système de n équations différentielles couplées du premier ordre. En posant : $dy / dt = z$, l'équation précédente devient :

$$\frac{dz}{dt} = -\frac{B(t)}{A(t)} \cdot z - \frac{C(t)}{A(t)} \cdot y + \frac{g(t)}{A(t)}$$

Avec :

$$\begin{cases} y(0) = y_0 \\ z(0) = V_0 \end{cases}$$

Ainsi, une équation d'ordre n peut être ramenée à un système du premier ordre:

$$\begin{cases} \frac{d\vec{y}}{dt} = \vec{f}(\vec{y}, t) ; \vec{y} = \{y_1, y_2, \dots, y_n\} ; \vec{f} = \{f_1, f_2, \dots, f_n\} \\ \vec{y}(0) = \vec{y}_0 \end{cases}$$

Soit par exemple l'équation: $y'''' + y'' + y' + y = f(t)$ avec les Conditions initiales:

$$y(0) = y_0, y'(0) = v_0 \text{ et } y''(0) = g_0$$

Pour simplifier la suite, écrivons plutôt: $y_1'''' + y_1'' + y_1' + y_1 = f(t)$(*)

Et posons: $y_2 = y_1'$

L'équation (*) s'écrit maintenant: $y_2'' + y_2' + y_2 + y_1 = f(t)$(**)

Qui est encore du deuxième ordre, on pose alors : $y_3 = y_2'$

L'équation (**) devient : $y_3' + y_3 + y_2 + y_1 = f(t)$

Au final, (*) s'écrit comme le système de trois équations du premier degré:

$$\begin{cases} y_1' = y_2 \\ y_2' = y_3 \\ y_3' = -y_3 - y_2 - y_1 + f(t) \end{cases} \quad \text{CI: } \begin{cases} y_1(0) = y_0 \\ y_2(0) = v_0 \\ y_3(0) = g_0 \end{cases}$$

I.5.1.b. Equations différentielles aux conditions aux limites

A titre d'exemple on cite le cas de l'équation de la propagation de la chaleur le long d'une barre de longueur L définie comme suit:

$$\begin{cases} \frac{d^2 y}{dx^2} + D \cdot \frac{dy}{dx} + E \cdot y = h(x) \\ y(0) = y_0 \\ y(L) = y_L \end{cases}$$

Ce problème est à conditions aux limites.

Si une condition aux limites est donnée en terme de gradient, par exemple $v_0 = dy(0)/dt$ plutôt que $y(0)$ lui-même, on peut adopter l'une ou l'autre des techniques suivantes :

- On écrit y_0 en fonction de y_1 à partir de $(y_1 - y_0)/h = v_0$ et on réarrange les termes de la première équation pour se ramener comme précédemment à une équation matricielle.
- On introduit un point « fantôme » y_{-1} tel que $(y_{-1} - y_1)/2h = v_0$ et on écrit une nouvelle équation correspondante à y_{-1} , y_0 et y_1 en y remplaçant y_{-1} par $y_{-1} = y_1 + 2hv_0$. Après réarrangement des termes, on obtient une matrice avec une colonne et ligne supplémentaire qu'il ne reste qu'à inverser pour résoudre le système.

I.5.2. Méthode d'Euler

La plus simple et la plus connue des méthodes d'approximation des solutions des équations différentielles ordinaires est la méthode d'Euler donnée par :

$\Phi(x, y, h) = f(x, y)$, de l'ordre 1, stabilité acquise (φ et f lipchitzienne) avec :

$$f(x_n, y_n) = \frac{y_{n+1} - y_n}{h}$$

La formule s'écrit donc :

$$y_{n+1} = y_n + hf(x_n, y_n)$$

I.5.2.a. Algorithme d'Euler général

- Etant donné un pas de temps h , une condition initiale (x_0, y_0) et un nombre maximal d'itérations N ;
 - Pour $0 \leq n \leq N$
- $$Y(n+1) = y(n) + h \cdot f(t(n), y(n))$$

$$t(n+1) = t(n) + h$$

Ecrire $t(n+1)$ et $y(n+1)$

➤ Arrêt

I.5.2.b. Algorithme d'Euler modifiée

➤ Etant donné un pas de temps h , une condition initiale (x_0, y_0) et un nombre maximal d'itérations N ;

➤ Pour $0 \leq n \leq N$

$$\hat{y} = y(n) + h \cdot f(t(n), y(n))$$

$$y(n+1) = y(n) + \frac{h}{2} (f(t(n), y(n)) + f(t(n+1), \hat{y}))$$

$$t(n+1) = t(n) + h$$

Ecrire $t(n+1)$ et $y(n+1)$

➤ Arrêt.

I.5.3. Méthode de Runge-Kutta

I.5.3.a. Méthode de Runge-Kutta d'ordre 2

$$\Phi(x, y, h) = \beta f(x, y) + \alpha f(x + \lambda h, y + \mu h f(x, y)).$$

➤ L'ordre 1 impose $\alpha + \beta = 1$

➤ L'ordre 2 impose $\lambda = \mu$ et $\alpha \lambda = \alpha \mu = 1/2$.

Si $\alpha = 1$, alors $\beta = 0$ et $\lambda = \mu = 1/2$, d'où la formule :

$$y_{n+1} = y_n + hf\left(x_n + \frac{h}{2}, y_n + \frac{h}{2}f(x_n, y_n)\right)$$

Soit le prédicteur $y_{n+1/2}$, la méthode s'écrit :

$$\begin{cases} y_{n+1/2} = y_n + \frac{h}{2}f(x_n, y_n) \\ y_{n+1} = y_n + hf(x_n + h/2, y_{n+1/2}) \end{cases}$$

La méthode est d'ordre 2, et correspond à une amélioration de la méthode d'Euler : le taux d'accroissement est calculé pour le point milieu : $(x_n + h/2, y_{n+1/2})$.

I.5.3.b. Méthode de Runge-Kutta d'ordre 4

- Etant donné un pas de temps h , une condition initiale (x_0, y_0) et un nombre maximal d'itérations N ;
- Pour $0 \leq n \leq N$

$$k_1 = f(x, y)$$

$$k_2 = f(x + h/2, y + h/2 \cdot k_1)$$

$$k_3 = f(x + h/2, y + h/2 \cdot k_2)$$

$$k_4 = f(x + h, y + h \cdot k_3)$$

$$y(n+1) = y(n) + \frac{1}{6}(k_1 + 2k_2 + 2k_3 + k_4)$$

$$t(n+1) = t(n) + h$$

Ecrire $t(n+1)$ et $y(n+1)$

- Arrêt

I.5.4. Méthode d'Adams

On déduit ces méthodes de la forme intégrale, en évaluant de manière approchée l'intégrale de f entre t_n et t_{n+1} . On suppose les nœuds de discrétisation équirépartis, c'est-à-dire $t_j = t_0 + jh$, avec $h > 0$ et $j \geq 1$. On intègre alors, au lieu de f , son polynôme d'interpolation aux $\tilde{p} + \vartheta$ nœuds distincts,

Où $\vartheta = 1$ quand les méthodes sont explicites (dans ce cas $\tilde{p} \geq 0$) et $\vartheta = 2$ quand les méthodes sont implicites (dans ce cas $\tilde{p} \geq -1$). Les schémas obtenus ont la forme suivante :

$$u_{n+1} = u_n + h \sum_{j=-1}^{\tilde{p}+\vartheta} b_j f_{n-j}$$

Les nœuds d'interpolation peuvent être ou bien :

- $t_n, t_{n-1}, \dots, t_{n-\tilde{p}}$ (dans ce cas $b_{-1} = 0$ et la méthode est explicite) ;

Ou bien

- $t_{n+1}, t_n, \dots, t_{n-\tilde{p}}$ (dans ce cas $b_{-1} \neq 0$ et le schéma est implicite).

Ces schémas implicites sont appelés méthodes d'Adams-Moulton, et les explicites sont appelés

I.5.4.a. Méthodes d'Adams - Bashforth.

En prenant $\tilde{p} = 0$, on retrouve la méthode d'Euler progressive, puisque le polynôme d'interpolation de degré zéro au nœud t_n est simplement $\Pi_0 f = f_n$. Pour $\tilde{p} = 1$, le polynôme d'interpolation linéaire aux nœuds t_{n-1} et t_n est

$$\Pi_1 f(t) = f_n + (t - t_n) \frac{f_{n-1} - f_n}{t_{n-1} - t_n}$$

Comme :

$$\Pi_1 f(t_n) = f_n \text{ et } \Pi_1 f(t_{n+1}) = 2f_n - f_{n-1}$$

On obtient :

$$\int_{t_n}^{t_{n+1}} \Pi_1 f(t) dt = \frac{h}{2} [\Pi_1 f(t_n) + \Pi_1 f(t_{n+1})] = \frac{h}{2} [3f_n - f_{n-1}]$$

Le schéma d'Adams – Bashforth à deux pas est donc :

$$u_{n+1} = u_n + \frac{h}{2} [3f_n - f_{n-1}]$$

Si $\tilde{p} = 2$, on trouve de façon analogue le schéma d'Adams-Bashforth à trois pas.

$$u_{n+1} = u_n + \frac{h}{12} [23f_n - 16f_{n-1} + 5f_{n-2}]$$

Et pour $\tilde{p} = 3$, on a le schéma d'Adams-Bashforth à quatre pas.

$$u_{n+1} = u_n + \frac{h}{24} (55f_n - 59f_{n-1} + 37f_{n-2} - 9f_{n-3})$$

Remarquer que les schémas d'Adams-Bashforth utilisent $\tilde{p}+1$ nœuds et sont des méthodes à $\tilde{p}+1$ pas (avec $\tilde{p} \geq 0$). De plus, les schémas d'Adams-Bashforth à q pas sont d'ordre q .

I.5.4.a. Méthodes d'Adams-Moulton.

Si $\tilde{p} = -1$, on retrouve le schéma d'Euler rétrograde. Si $\tilde{p} = 0$, on construit le polynôme d'interpolation de degré un de f aux nœuds t_n et t_{n+1} , et on retrouve le schéma de Crank-Nicolson.

Pour la méthode à deux pas ($\tilde{p} = 1$), on construit le polynôme d'interpolation de degré 2 de f aux nœuds t_{n-1} , t_n , t_{n+1} , et on obtient le schéma suivant :

$$u_{n+1} = u_n + \frac{h}{12} [5f_{n+1} + 8f_n - f_{n-1}]$$

Les schémas correspondant à $\tilde{p} = 2$ et 3 sont respectivement donnés par :

$$u_{n+1} = u_n + \frac{h}{24} (9f_{n+1} + 19f_n - 5f_{n-1} + f_{n-2})$$

Et

$$u_{n+1} = u_n + \frac{h}{720} (251f_{n+1} + 646f_n - 264f_{n-1} + 106f_{n-2} - 19f_{n-3})$$

Les schémas d'Adams-Moulton utilisent $\tilde{p} + 2$ nœuds et sont à $\tilde{p} + 1$ pas si $\tilde{p} \geq 0$, la seule exception étant le schéma d'Euler rétrograde ($\tilde{p} = -1$) qui est à un pas et utilise un nœud. Les schémas d'Adams-Moulton à q pas sont d'ordre $q + 1$, excepté à nouveau le schéma d'Euler rétrograde qui est une méthode à un pas d'ordre un.

I.6. Conclusion

Le rappel que nous avons cité sur les méthodes numériques consiste à entamer une branche des mathématiques appliquées s'intéressant au développement d'outils et de méthodes numériques pour le calcul d'approximations de solutions de problèmes de mathématiques qu'il serait difficile, voire impossible, d'obtenir par des moyens analytiques. Son objectif est notamment d'introduire des procédures calculatoires détaillées susceptibles d'être mises en œuvre par des calculateurs (électroniques, électrotechniques, mécaniques ou humains) et d'analyser leurs caractéristiques et leurs performances.

II.1. Introduction

Ce chapitre traite des problèmes théoriques liés à la résolution d'équations aux dérivées partielles (EDP) qui sont ubiquistes dans toutes les sciences, puisqu'elles apparaissent aussi bien en dynamique des structures, mécanique des fluides que dans les théories de la gravitation ou de l'électromagnétisme (Exemple: les équations de Maxwell). Certaines de ces EDP ont été résolues analytiquement et leurs solutions sont connues. Toutefois, un nombre important d'autres existent sans solutions analytiques. C'est dans cette vision que les recherches se sont inclinées sur les méthodes numériques pour arriver à approximer les solutions de ces équations.

Les perfectionnements de l'informatique ont permis de développer des méthodes numériques de calcul afin de déterminer de façon précise la distribution des paramètres physiques dans la plus part des systèmes électriques. Les méthodes numériques les plus connues et les plus utilisées dans ce type de problème sont donc la méthode des différences finies (MDF), la méthode des éléments finis (MEF), la méthode de simulation de charges (MSC), la méthode des éléments finis de frontière (MEFF), la méthode des volumes finis (MVF), ... etc.

Dans ce chapitre, nous allons rappeler quelques méthodes numériques qui permettent de résoudre les équations différentielles gouvernantes les phénomènes physiques. Nous présenterons, en particulier, la méthode des différences finies, la méthode des éléments finis.

II.2. Définition

Les équations aux dérivées partielles (EDP) sont des équations dont les solutions sont les fonctions inconnues vérifiant certaines conditions concernant leurs dérivées partielles. Une équation (EDP), est une relation faisant intervenir une fonction inconnue u de R^n dans R , les variables $x, y, \dots$, ses dérivées partielles, $u_x, u_y, \dots, u_{xx}, u_{xy}, u_{yy}, \dots$. Elle s'écrit de façon générale : $f(x, y, \dots, u, u_x, u_y, \dots, u_{xx}, u_{xy}, \dots) = 0$

Cette équation est considérée dans un domaine Ω de $\mathbb{R}^n$. Les solutions de l'équation aux dérivées partielles sont les fonctions qui vérifient cette équation dans Ω . L'ordre d'une équation aux dérivées partielles est l'ordre de la dérivée partielle d'ordre le plus élevé intervenant dans l'équation.

Exemples :

$$u^2 u_{xy} + u_x = y$$

$$u_{xx} + 2y^2 u_{xy} + 3xu_{yy} = 1$$

$$(u_x)^2 + (u_y)^2 = 1$$

$$u_{xx} - u_{yy} = 0$$

Les fonctions $u(x, y) = (x + y)^3$ et $u(x, y) = \sin(x - y)$ sont toutes deux des solutions de cette dernière équation.

Les conditions étant moins strictes que dans le cas d'une équation différentielle ordinaire; les problèmes incluent souvent des conditions aux limites qui restreignent l'ensemble des solutions.

Pour assurer donc l'unicité de la solution, comme on le fait avec les équations différentielles ordinaires (EDO), on tiendra compte des conditions pré-données comme les conditions aux limites et les conditions initiales.

II.3. Classification des équations aux dérivées partielles

Soit $X = (x_1, x_2, \dots, x_n) \in \mathbb{R}^n$, une équation aux dérivées partielles du second ordre sera de la forme:

$$\sum_{i=1}^n \sum_{j=1}^n A_{i,j}(X) \frac{\partial^2 u}{\partial x_i \partial x_j}(X) + \sum_{i=1}^n B_i(X) \frac{\partial u}{\partial x_i}(X) + Cu = G(X)$$

avec $A_{i,j}$, B_i , C , G des fonctions indépendantes de u ne s'annulant pas toutes simultanément dans $\mathbb{R}^n$. Si nous nous limitons dans $\mathbb{R}^2$, c'est à dire $X = (x; y) \in \mathbb{R}^2$ l'égalité précédemment posée prend la forme de:

$$A \frac{\partial^2 u}{\partial x^2} + B \frac{\partial^2 u}{\partial x \partial y} + C \frac{\partial^2 u}{\partial y^2} + D \frac{\partial u}{\partial x} + E \frac{\partial u}{\partial y} + Fu = G(x, y)$$

La classe d'une telle équation est déterminée par le calcul de :

$$\Delta = B^2(x_0, y_0) - 4A(x_0, y_0)C(x_0, y_0)$$

- Si $\Delta < 0$, on parle d'une équation elliptique,
- Si $\Delta = 0$, on parle d'une équation parabolique,
- Si $\Delta > 0$, on parle d'une équation hyperbolique.

II.4. Rappels sur la classification des systèmes physiques

Un système physique est caractérisé par un ensemble de variables qui peuvent dépendre des coordonnées d'espace $X = (x, y, z)$ et du temps t .

- Le système est dit stationnaire si ses variables ne dépendent pas du temps.
- Le système est dit non stationnaire si ses variables dépendent du temps.

II.4.1. Système continu

Un système est dit continu s'il possède un nombre de degrés de liberté infini et son comportement est représenté le plus souvent par un système d'équations aux dérivées partielles (EDP) ou intégral-différentielles associés à des conditions aux limites en espace et en temps.

II.4.2. Système discrets

Un système est dit discret s'il possède un nombre de degrés de liberté fini et son comportement est représenté par un système d'équations algébriques.

Les équations algébriques des systèmes discrets peuvent être résolues par les méthodes numériques, par contre les équations des systèmes continus ne peuvent en générale pas être résolues directement. Il est nécessaire de discrétiser ces équations, c'est – à – dire de les remplacer par des équations matricielles.

II.4.3. Choix du maillage

Le choix du maillage consiste à diviser le domaine de travail Ω en parties égales ou non afin d'obtenir un espace discret. L'espace ainsi obtenu s'appellera espace d'interpolation et aura toutes les propriétés d'un sous-espace vectoriel de Ω . Les solutions héritées seront de ce fait approchées. Les sous-divisions obtenues sont appelées éléments finis. Les points de jonction entre les éléments sont les nœuds.

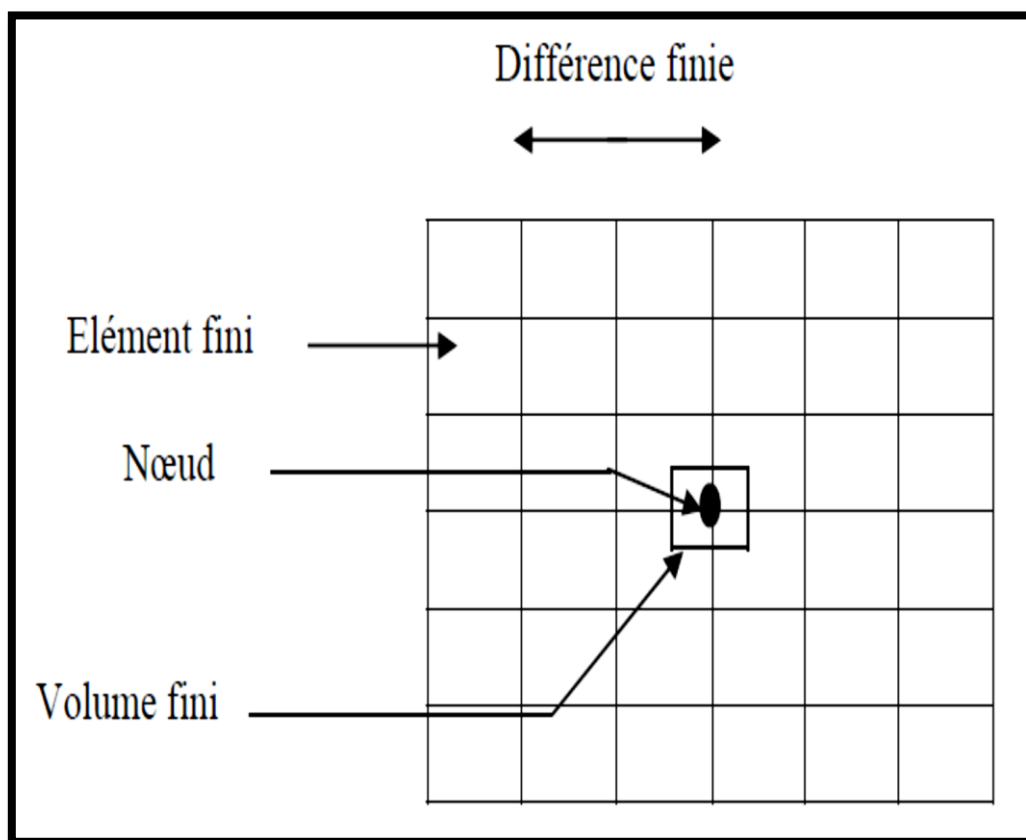


Fig.II.1. Maillage du domaine d'étude (Ω)

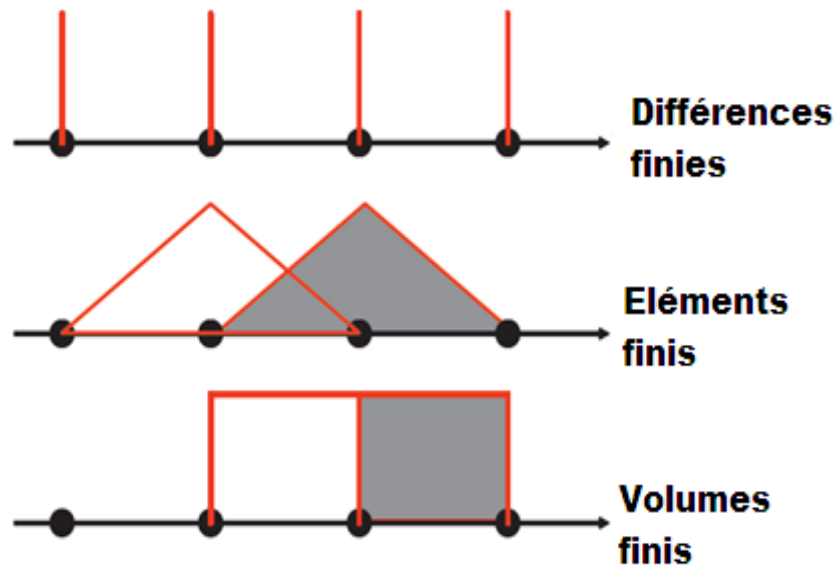


Fig. II.2. Exemple unidimensionnel de fonction de projection pour différentes méthodes numériques (différences finies, éléments finis et volumes finis)

II.5. Différentes méthodes de résolution des équations aux dérivées partielles

Il existe deux grandes catégories de méthodes de résolution des équations aux dérivées partielles mathématiques caractérisant les problèmes physiques, lorsqu'il s'agit de calculer des effets dont les causes (densité du courant (tension), densité de la puissance dissipée) sont connues à l'avance. Ces méthodes sont :

- Les méthodes analytiques.
- Les méthodes numériques.

Les méthodes analytiques, s'avèrent d'applications très difficiles dès que la complexité de la géométrie s'accroît et que certains phénomènes, dans des conditions de fonctionnement optimales, présentent des non linéarités physiques, donc mathématiques.

L'apparition des ordinateurs, de grandes puissances, a mis en valeur l'intérêt des méthodes dites numériques. Celles ci font appel à des techniques de discrétisation.

Ces méthodes numériques transforment les équations aux dérivées partielles (EDP) à des systèmes d'équations algébriques dont la solution fournit une approximation de l'inconnue en

différent points situés aux nœuds du réseau géométrique correspondant à la discrétisation. Parmi ces méthodes, nous citons la méthode des différences finies, la méthode des éléments finis, la méthode des volumes finis, la méthode des intégrales de frontières et la méthode des circuits couplés,... etc.

II.5.1. Méthode des Différences Finies (MDF) :

C'est la méthode la plus ancienne, connue depuis Gauss. Le principe fondamental de cette méthode consiste à appliquer au domaine d'étude un maillage en nœuds dont la finesse permet de donner une approximation des contours du domaine. Ensuite, en appliquant le développement limité en série de Taylor de la fonction à déterminer dans chaque nœud du maillage, ce qui permet d'obtenir un nombre d'équations algébriques égales au nombre des valeurs d'inconnues des grandeurs étudiées.

La générale de la méthode va nous permettre d'aborder les notions de consistance, de stabilité et de convergence.

1- Consistance :

Une méthode numérique est consistante si l'erreur de discrétisation de l'équation tend vers zéros lorsque le pas de discrétisation tend vers zéros.

2- Stabilité :

Une fois le schéma discret choisi, il va être nécessaire de résoudre les équations. Le processus de résolution, à la vue des équations sera la plupart du temps itératif. On calculera les valeurs de « Φ » de proche en proche : une valeur donnée de « Φ » sera donc calculée en utilisant le résultat du calcul d'autres valeurs de « Φ ».

Les erreurs d'arrondi être inévitables sur machine, la méthode sera stable si ces erreurs ne s'amplifient pas (trop) au cours du calcul.

3- Convergence :

Après s'être assuré que le schéma discret tend vers l'équation, que ce schéma conduit à un calcul stable, on dispose d'une solution. Le schéma utilisé est dit convergent si sa solution ainsi obtenue tend vers la solution de l'équation de départ, lorsque les pas de discrétisation tendent vers zéros.

Donc, la consistance et la stabilité sont nécessaires et suffisantes pour assurer la convergence.

4- Discrétisation de l'EDP :

Soit :

$$\frac{\partial^2 u}{\partial x^2} + \frac{\partial^2}{\partial x^2} = 0, \forall (x, y) \in [a, b] \times [c, d]$$

On posera h_x et h_y les pas de discrétisation des intervalles $[a, b]$ et $[c, d]$

➤ Discrétisation de l'intervalle $[a, b]$

$$h_x = \frac{b-a}{n_x}$$

(n_x étant le nombre d'intervalles dans $[a, b]$)

$$\Rightarrow x(i) = x_i = a + i \times h_x, i = 0, 1, \dots, n_x.$$

➤ Discrétisation de l'intervalle $[c, d]$

$$h_y = \frac{d-c}{n_y}$$

(n_y étant le nombre d'intervalles dans $[c, d]$)

$$\Rightarrow y(j) = y_j = c + j \times h_y, j = 0, 1, \dots, n_y.$$

Remarquant que

$$x_{i+1} = a + (i+1)h_x = (a + ih_x) + h_x = x_i + h_x$$

Dans la suite, nous remplacerons chaque fois :

$$x_i + h_x, x_i - h_x, y_j + h_y, y_j - h_y.$$

Successivement par : x_{i+1} , x_{i-1} , y_{j+1} et y_{j-1} .

5- Application de MDF:

La MDF consiste à approximer les dérivées partielles d'une équation au moyen des développements de Taylor et ceci se déduit directement de la définition de la dérivée.

Soit $f(x, y)$ une fonction continue et dérivable de classe C^∞ , alors la dérivée partielle première de f par rapport à x est calculée par la formule:

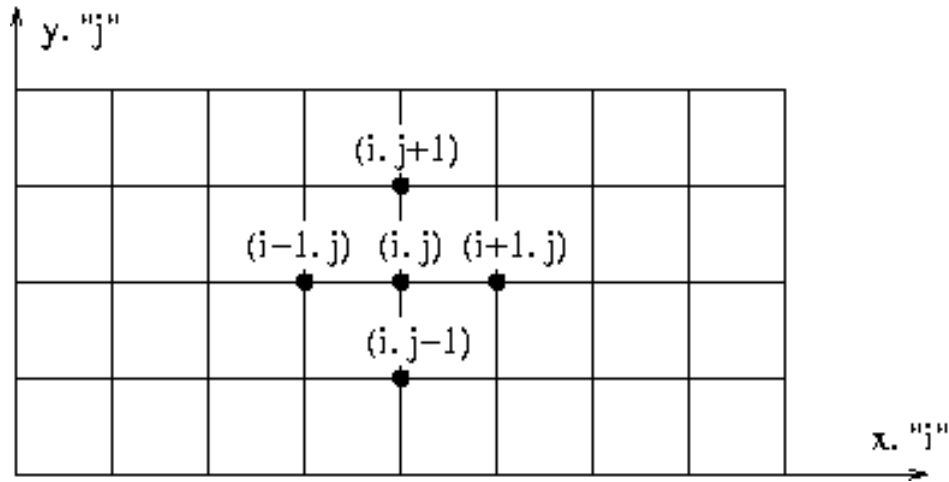


Fig.II. 3. Grille de discrétisation régulière $h_x = \Delta x = h_y = \Delta y$.

$$f'_x(x, y) = \lim_{h_x \rightarrow 0} \frac{f(x + h_x, y) - f(x, y)}{h_x}$$

Si $h_x \ll 1$, le développement de Taylor au voisinage de 0 de $f(x + h_x, y)$ donne:

$$f(x + h_x, y) = f(x, y) + h_x \frac{\partial f}{\partial x} + \theta(h_x) \simeq f(x, y) + h_x \frac{\partial f}{\partial x}$$

Avec une erreur de l'ordre de h_x .

$$\Rightarrow \frac{\partial f(x, y)}{\partial x} \simeq \frac{f(x + h_x, y) - f(x, y)}{h_x}$$

Ceci est appelé le schéma avant (différences finies en avant).

De la même manière, nous pouvons aussi donner le schéma arrière (différences finies en arrière) qui est de la forme:

$$\frac{\partial f(x, y)}{\partial x} = \lim_{h_x \rightarrow 0} \frac{f(x, y) - f(x - h_x, y)}{h_x}$$

Avec la formule de Taylor, ceci nous donne :

$$f(x, y) = f(x - h_x, y) + h \frac{\partial f(x, y)}{\partial x} + \theta(h_x) \simeq f(x - h_x, y) + h \frac{\partial f(x, y)}{\partial x}$$

$$\Rightarrow \frac{\partial f(x, y)}{\partial x} \simeq \frac{f(x, y) - f(x - h_x, y)}{h_x}$$

Avec une erreur de l'ordre de h_x .

La somme de ces deux schémas nous donne le schéma centré suivant :

$$\frac{\partial f(x, y)}{\partial x} \simeq \frac{f(x + h_x, y) - f(x - h_x, y)}{2h_x}.$$

En résumé, on a les trois approximations suivantes pour la première dérivée de $f(x, y)$ par rapport à x avec le développement limité de Taylor:

$$f'_x(x, y) = \lim_{h_x \rightarrow 0} \frac{f(x + h_x, y) - f(x, y)}{h_x} \approx \begin{cases} \frac{f(x + h_x, y) - f(x, y)}{h_x} & \text{schéma avant} \\ \frac{f(x, y) - f(x - h_x, y)}{h_x} & \text{schéma arrière} \\ \frac{f(x + h_x, y) - f(x - h_x, y)}{2h_x} & \text{schéma centré} \end{cases}$$

La dérivée seconde f''_x de $f(x, y)$ sera alors de la forme:

$$\frac{\partial^2 f}{\partial x^2} \approx \frac{\frac{f(x_{i+1}, y_j) - f(x_i, y_j)}{h_x} - \frac{f(x_i, y_j) - f(x_{i-1}, y_j)}{h_x}}{h_x}$$

$$\frac{\partial^2 f}{\partial x^2} \simeq \frac{f(x_{i+1}, y_j) - 2f(x_i, y_j) + f(x_{i-1}, y_j)}{h_x^2}$$

Exemple : prenons le cas d'une l'équation de Laplace:

$$\Delta u = 0 \Leftrightarrow \frac{\partial^2 u}{\partial x^2} + \frac{\partial^2 u}{\partial y^2} = 0 \dots\dots\dots(*)$$

Posons $u(x_i, y_j) = u_{i,j}$ (en notation indicielle). Compte tenu de la formule de dérivée seconde f''_x donnée au paragraphe précédent on a :

$$\frac{\partial^2 u}{\partial x^2} \approx \frac{u_{i+1,j} - 2u_{i,j} + u_{i-1,j}}{h_x^2}$$

Puisque x_i et y_j jouent un rôle symétrique dans l'équation du potentiel (de Laplace), un raisonnement analogue à celui de l'approximation de f_x'' nous donne:

$$\frac{\partial^2 u}{\partial y^2} \approx \frac{u_{i,j+1} - 2u_{i,j} + u_{i,j-1}}{h_y^2}$$

Rapportons ces approximations dans l'EDP de Laplace (*) on trouve :

$$\Leftrightarrow \Delta u \approx \frac{u_{i+1,j} - 2u_{i,j} + u_{i-1,j}}{h_x^2} + \frac{u_{i,j+1} - 2u_{i,j} + u_{i,j-1}}{h_y^2} = 0$$

Dans ce cas particulier où $h_x = h_y = h$, donc, nous avons finalement:

$$\begin{cases} \Delta u = 0 \Leftrightarrow \frac{u_{i+1,j} + u_{i,j+1} - 4u_{i,j} + u_{i-1,j} + u_{i,j-1}}{h^2} = 0 \\ i = 0, 1, \dots, n_x \text{ et } j = 0, 1, \dots, n_y \end{cases}$$

A chaque étape, nous remarquons que pour calculer la valeur de $u_{i,j}$ au point (x_i, y_j) nous avons besoin de connaître les points :

$$u_{i-1,j}, u_{i,j-1}, u_{i+1,j} \text{ et } u_{i,j+1}$$

Comme l'indique le dessin suivant:

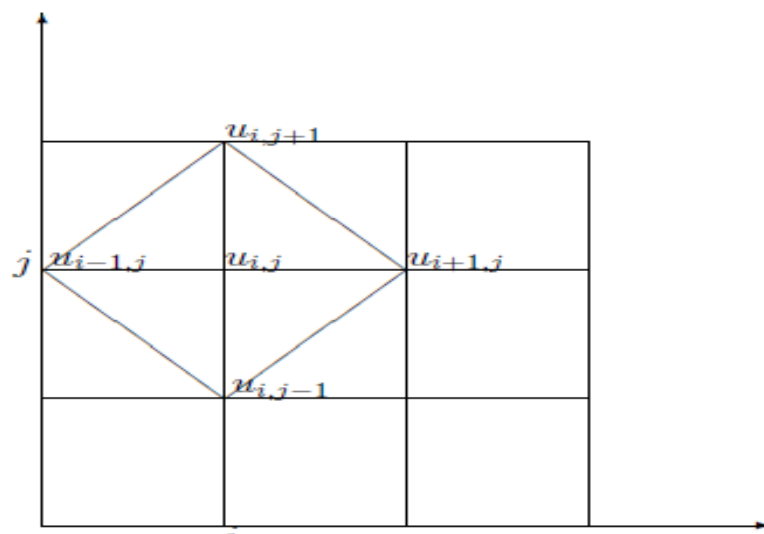


Fig. II.4. Molécule de l'équation de Laplace u .

C'est pour cela que nous appelons cette formule la formule à 5 points qui peut être représentée

Comme suit:

$$\Delta u = 0 \Rightarrow \frac{1}{h^2} \begin{pmatrix} & 1 & \\ 1 & -4 & 1 \\ & 1 & \end{pmatrix} u_{i,j} = 0$$

II.5.1.a. Avantages de la MDF

- La méthode des différences finies est une méthode simple à appliquer lorsque la géométrie le permet et c'est une méthode raisonnablement exacte.
- De plus, elle se programme facilement et nécessite peu de mémoire pour le stockage des données.

II.5.1.b. Inconvénients de la MDF

- Lorsque la géométrie est de frontière courbe, le schéma ne peut s'appliquer près des frontières irrégulières et donc cette méthode devient difficilement applicable. On doit alors rechercher à la place une méthode qui est valide indépendamment de la géométrie.
- Elle n'est pas applicable pour des problèmes en 3 dimensions.
- Cette méthode nécessite la connaissance, sur toute la frontière entourant le domaine étudié, du potentiel, ce qui n'est pas toujours le cas en général.

II.5.2. Méthode des éléments Finis (MEF)

Cette méthode, utilisée depuis longtemps en mécanique. Elle a été introduite en électromagnétisme par P. Silvestre et M.V.K. Chari en 1970.

Elle a connu depuis, un développement considérable dans ce domaine, grâce aux rapports successifs des équipes universitaires de MC Gill au Canada, Rut Herford en grande Bretagne et Grenoble en France et par quelques grands laboratoires industriels de recherches.

La méthode des éléments finis est très puissante pour la résolution des équations aux dérivées partielles (EDP) sur tout dans les géométries complexes et quelques soient les conditions physiques de fonctionnements.

A la différence avec la MDF, la MEF consiste à utiliser une approximation simple de l'inconnue pour transformer les EDP en équations algébriques.

Toute fois, cette méthode ne s'applique pas directement aux EDP, mais à une formulation intégrale qui est équivalente au problème à résoudre, en utilisant l'une des deux approches suivantes:

-La méthode variationnelle qui consiste à minimiser une fonctionnelle qui représente généralement, l'énergie du système étudié. Cette méthode n'est donc applicable que si on connaît une fonctionnelle équivalente au problème différentiel que l'on veut résoudre.

-La méthode des résidus pondérés ou méthode projective qui consiste à minimiser le résidu induit par l'approximation de la fonction inconnue.

A l'une ou à l'autre des deux méthodes, on associe une subdivision du domaine d'étude, en éléments simples, appelés éléments finis, comme il est indiqué sur la figure (5), et à approximer la fonction inconnue sur chaque élément par des fonctions d'interpolation. Ces fonctions sont généralement des polynômes de Lagrange de degré un, ou deux.

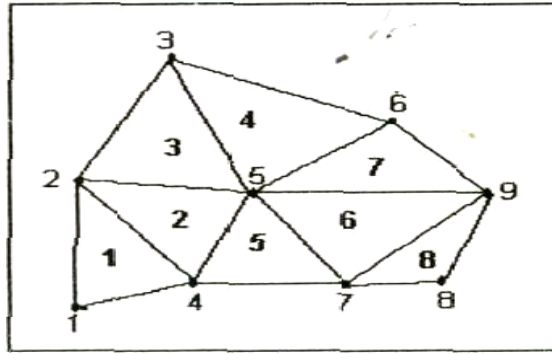


Fig.II.5. Un domaine d'étude discrétisé en éléments finis

Comme on le sait, cette méthode est une technique mathématique d'intégration des équations aux dérivées partielles mises sous forme vibrationnelle.

Elle fait appel aux trois domaines suivants:

- sciences de l'ingénieur pour construire les équations aux dérivées partielles.
- méthodes numériques pour construire et résoudre les équations algébriques.
- programmation et informatique pour exécuter efficacement les calculs sur ordinateur.

Cette méthode a connu plusieurs développements depuis son apparition, plus exactement par l'apparition des publications de Turner, Clough, Martin et Topp qui ont introduit le concept d'élément fini. Soulignons également le travail d'Argyris et Kelesy qui systématise l'utilisation de la notion d'énergie dans l'analyse des structures.

Dès 1960 cette méthode subit un développement rapide dans plusieurs directions la méthode des éléments finis est reconnue comme un outil général de résolution d'équations aux dérivées partielles. Elle est donc utilisée pour résoudre des problèmes non stationnaires, non linéaires dans le domaine des ainsi que dans d'autres domaines une base mathématique de la méthode des éléments finis est construite à partir de l'analyse fonctionnelle. La méthode est reformulée, à partir de considérations énergétiques et vibrationnelles sous la forme générale des résidus pondérés ainsi on assiste aux développements de nouveaux éléments tels que poutre, plaques, coques et l'établissement de nouvelles formulations.

La méthode des éléments finis permet donc de résoudre de manière discrète une équation dérivée partielle dont on cherche une solution approchée « suffisamment » fiable. De manière générale, cette EDP porte sur une fonction définie sur un domaine qui comporte des conditions aux bords permettant d'assurer existence et unicité d'une solution.

Sauf cas particulier, la discrétisation passe par une redéfinition et une approximation de la géométrie, on considère donc le problème posé sur la géométrie approchée par un domaine polygonal ou polyédrique par morceaux. Une fois la géométrie approchée, il faut choisir un espace d'approximation de la solution du problème, dans la MEF cet espace est défini à l'aide du maillage du domaine (ce qui explique aussi pourquoi il est nécessaire d'approcher la géométrie). Un élément fini est la donnée d'une cellule élémentaire et de fonctions de base de l'espace d'approximation dont le support est l'élément, définies de manière à être interpolantes.

Bien qu'il existe de nombreux logiciels exploitant cette méthode et permettant de « résoudre » des problèmes dans divers domaines, il est important que l'utilisateur ait une bonne idée de ce qu'il fait, notamment quant au choix du maillage et du type d'éléments qui doivent être adaptés au problème posé, aucun logiciel ne fera tout pour l'utilisateur, et il faut toujours garder un monde critique vis-à-vis de solutions approchées. Pour la solution trouvée, il reste cependant à déterminer les caractéristiques de la méthode ainsi développée, notamment l'unicité de l'éventuelle solution ou encore la stabilité numérique du schéma de résolution. Il est essentiel de trouver une estimation juste de l'erreur liée à la discrétisation et montrer que la méthode ainsi écrite converge, c'est-à-dire que l'erreur tend vers zéro si la finesse du maillage tend elle aussi vers zéro. Dans le cas d'une EDP linéaire avec opérateur symétrique, il s'agit finalement de résoudre une équation algébrique linéaire, inversible dans le meilleur des cas.

Dans cette partie, nous allons rappeler quelques propriétés de la méthode des éléments finis.

Une illustration de cette présentation est donnée en prenant l'équation de Poisson suivante en deux dimensions définie dans le domaine donné sur la figure (II.6).

$$\frac{\partial^2 \Phi}{\partial x^2} + \frac{\partial^2 \Phi}{\partial y^2} = G(x, y)$$

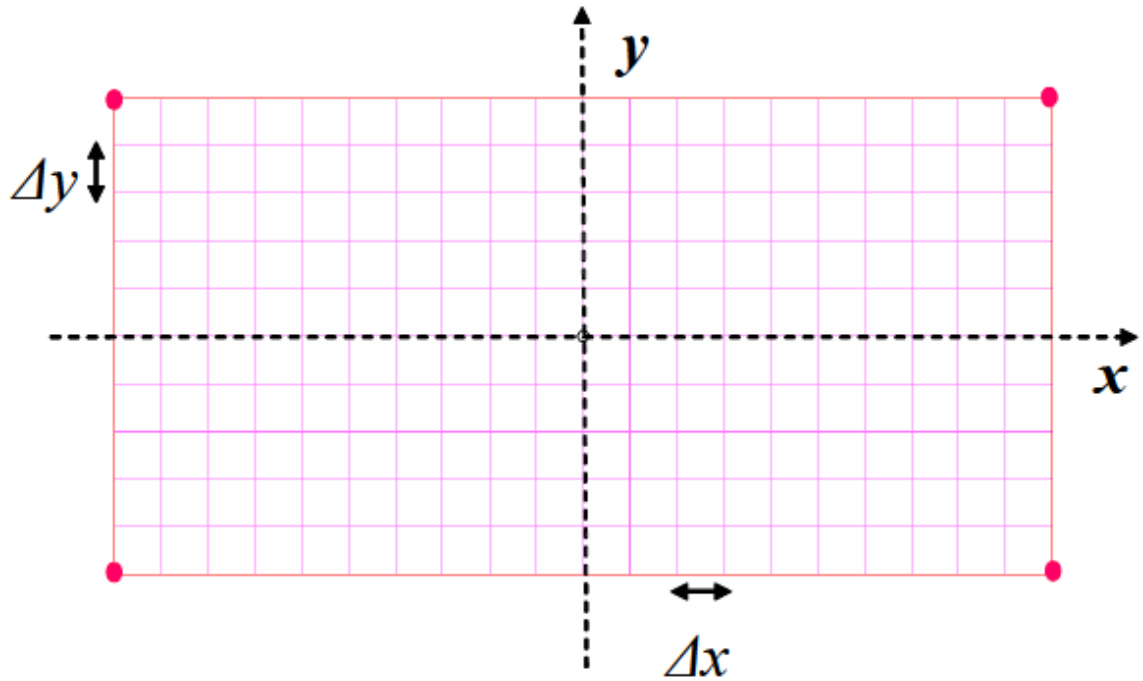


Fig.II.6. Maillage carré.

Chaque point est situé sur un des sommets d'un carré. Il existe d'autres maillages réguliers comme le maillage en triangles équilatéraux et en hexagones équi-angulaires.

La MEF, outil numérique très puissant, est beaucoup utilisé dans la résolution des problèmes à domaine spatial fini, surtout en mécanique où elle a connu son plus fort développement. Cette méthode a été appliquée avec succès dans les problèmes de calcul de potentiel et de champs électriques.

L'idée de cette méthode est de chercher une solution approchée à une équation différentielle après une reformulation sous forme d'identité intégrale appelée forme faible ou variationnelle. Au lieu de rechercher à satisfaire l'équation aux nœuds, nous décomposons ici le domaine en sous domaines appelés éléments finis, et nous imposons la satisfaction des équations par sous domaine. L'introduction d'une approximation locale par sous domaine (dit élément fini), permet de

contourner le problème de complexité des géométries car il suffit alors de choisir une approximation ou une décomposition (maillage) qui respecte la géométrie.

La méthode des éléments finis utilise des approximations simples à variables inconnues dans chaque élément pour transformer les équations aux dérivées partielles en équations algébriques. Les nœuds et les éléments n'ont pas forcément de signification physique particulière, mais sont basés sur des considérations de précision de l'approximation.

L'approximation peut fournir une solution approchée en tout point (x, y) d'une fonction difficile à évaluer ou connue seulement en certains points et une solution approchée d'une équation différentielle ou aux dérivées partielles. Lorsque la frontière du domaine est constituée par des courbes ou des surfaces plus complexes que celles qui définissent les frontières des éléments, une erreur est inévitable. Cette erreur est appelée "erreur de discrétisation géométrique". Elle peut être réduite en diminuant la taille des éléments.

II.5.2.1. Principe de la méthode

La MEF repose sur deux principes : d'une part, la formulation d'un problème approché par la méthode de Galerkin, qui permet de remplacer un problème posé en dimension infinie par un système linéaire. D'autre part, la construction d'un espace d'approximation (de dimension finie) à l'aide d'un maillage, de fonctions polynomiales par morceaux et de degré de liberté sur chaque maille.

La première étape dans la construction de l'espace d'approximation S_b , consiste à mailler l'intervalle Ω en une dimension d'espace (Fig II.7). Un maillage $\Omega =]a_k, b_k[$ est une collection indexée d'intervalle.

$$\{I_i = [x_{1,i}, x_{2,i}] \}_{1 \leq i \leq N_{ma}}$$

Tous de mesure non nulle et formant une partition de Ω .

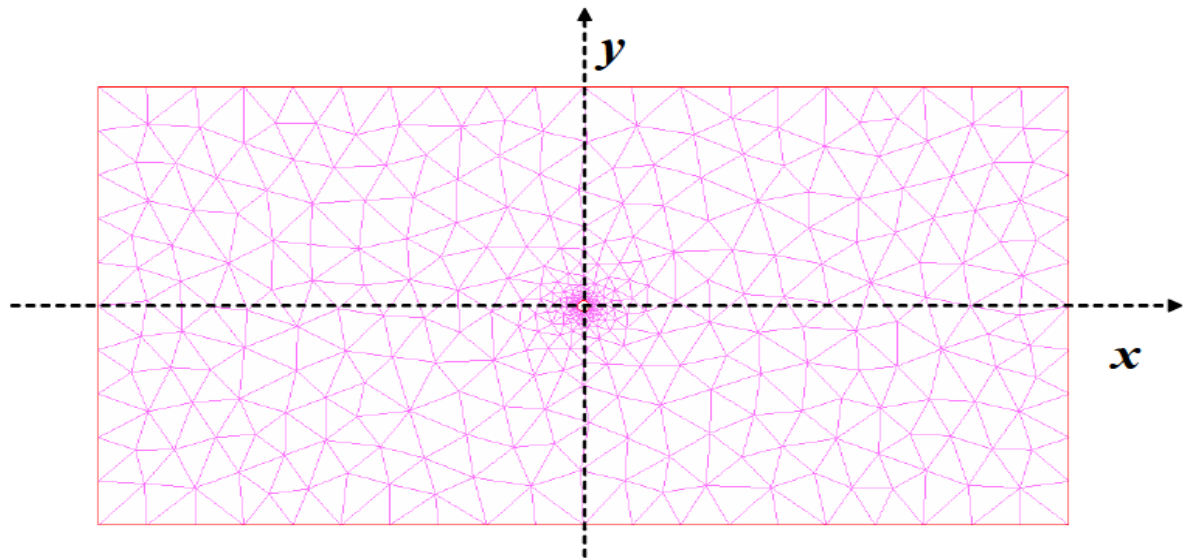


Fig.II.7. Le maillage en éléments finis triangulaires.

En d'autres termes nous avons :

$$]a_k, b_k[= \bigcup_{i=1}^{N_{ma}} [x_{1,i}, x_{2,i}] \quad \text{et} \quad]x_{1,i}, x_{2,i}[\cap]x_{1,j}, x_{2,j}[= \emptyset \quad \text{pour } i \neq j$$

Les intervalles I_i sont appelés les mailles (ou les éléments ou les cellules de maillage), et l'entier N_{ma} désigné le nombre totale de mailles.

La façon la plus simple de construire un maillage est de choisir $(N_{ma}+1)$ points distincts de Ω tels que :

$$a_k = x_1 < x_2 < \dots < x_{N_{ma}} < x_{N_{ma}+1} = b_k$$

Et de poser :

$$x_{1,i} = x_i \quad \text{et} \quad x_{1,j} = x_j$$

Pour tout :

$$i \in \{1, \dots, N_{ma}\}$$

Les points de l'ensemble :

$$\{x_1, \dots, x_{N_{ma}+1}\}$$

Sont appelés les sommets du maillage. Nous désignons par N_{s0} le nombre de sommets du maillage en une dimension d'espace. On a donc :

$$N_{s0} = N_{ma} + 1$$

Le maillage est à priori de pas variable, on pose pour tout :

$$i \in \{1 \dots N_{ma}\}$$

$$b_i = x_{i+1} - x_i \quad \text{et} \quad b_k = \text{MAX}_{1 \leq i \leq N_{ma}} b_i$$

On dit que le maillage est uniforme lorsque $b_k = b_i$ pour tout $i \in \{1 \dots N_{ma}\}$ par la suite, le maillage est désigné sous la forme :

$$T_b = \{I_i\}_{1 \leq i \leq N_{ma}}$$

La deuxième étape dans la construction de l'espace d'approximation consiste à choisir des fonctions de forme sur chaque maille. Ces fonctions d'interpolation permettent alors de donner une approximation du potentiel Φ , notée $\tilde{\Phi}$, sur chaque élément en fonction de ses valeurs aux nœuds de l'élément comme suit :

$$\tilde{\Phi} = \sum_{i=1}^{N_e} N_i \Phi_i$$

Avec N_e le nombre de nœuds d'interpolation, N_i les fonctions d'interpolation et Φ_i les valeurs nodales du potentiel.

Pour illustrer le principe de la MEF, on reprend l'exemple de l'équation de Poisson du paragraphe précédent en deux dimensions définie dans le domaine donné. Et on recherche à minimiser la quantité R_M telle que :

$$R_M = \left(\nabla^2 \tilde{\Phi} \right) \quad (1)$$

Parmi toutes les méthodes qui permettent d'annuler une grandeur dans un domaine Ω , la méthode des résidus pondérés est bien connue et souvent utilisée. Elle consiste à choisir un ensemble de fonctions linéairement indépendantes W_n , appelées fonctions de projection, et annuler ainsi toutes les intégrales (1) sur chacun des éléments finis.

$$I_n = \int_{\Omega} W_n R_M d\Omega \quad (2)$$

L'intégration par partie de l'équation (2) donne :

$$I_n = - \int_{\Omega_e} \text{grad} \tilde{\Phi} \cdot \text{grad} W_n d\Omega + \int_{S_e} (\text{grad} \tilde{\Phi} \cdot \vec{n}) W_n dS_e \quad (3)$$

Pour chaque élément, on annule les n intégrales I_n (3) qui correspondent aux n fonctions de projection. On obtient un ensemble de n équations à n inconnues formant ainsi un système élémentaire pouvant s'écrire sous la forme matricielle suivante :

$$[A_e] \{\Phi_e\} = \{b_e\}$$

Avec :

$[A_e]$: est la matrice associée à l'élément considéré ;

$\{\Phi_e\}$: ses composantes sont les inconnues du potentiel aux nœuds du même élément ;

$\{b_e\}$: est le vecteur qui tient compte des éventuelles conditions aux limites présentées sur certains nœuds de l'élément considéré.

La résolution du système final est simple puisque les équations obtenues sont linéaires et les matrices ainsi formées sont symétriques. Pour déterminer la distribution du champ électrique, il faut calculer la dérivée du potentiel par une méthode numérique adaptée.

II.5.2.2. Domaines d'application

La méthode des éléments finis est appliquée dans la majorité des domaines de la physique comme montre la figure (II.8). A titre d'exemple, mais sans s'y limiter : la mécanique, l'électromagnétique, thermodynamique, chimique, météorologie, etc,.. Dans tous ces cas, la formulation reste quasiment identique, mais la nature des champs et les lois de comportement sont adaptées au domaine d'application.

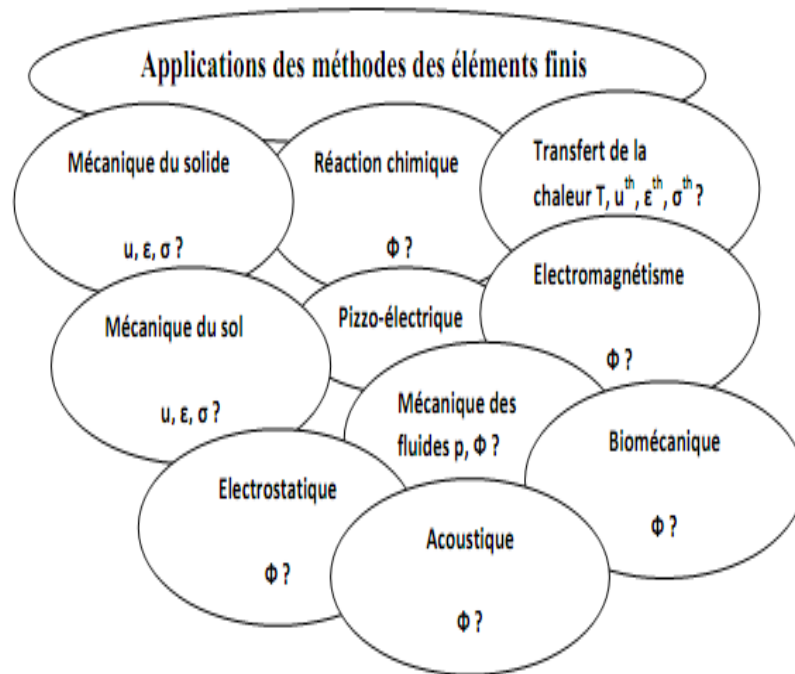


Fig.II. 8 : Domaines d'applications de la méthode des éléments finis.

II. 6 Types de problèmes MEF

La méthode des éléments finis permet la résolution de trois types de problème principaux: comme la figure suivante illustre.

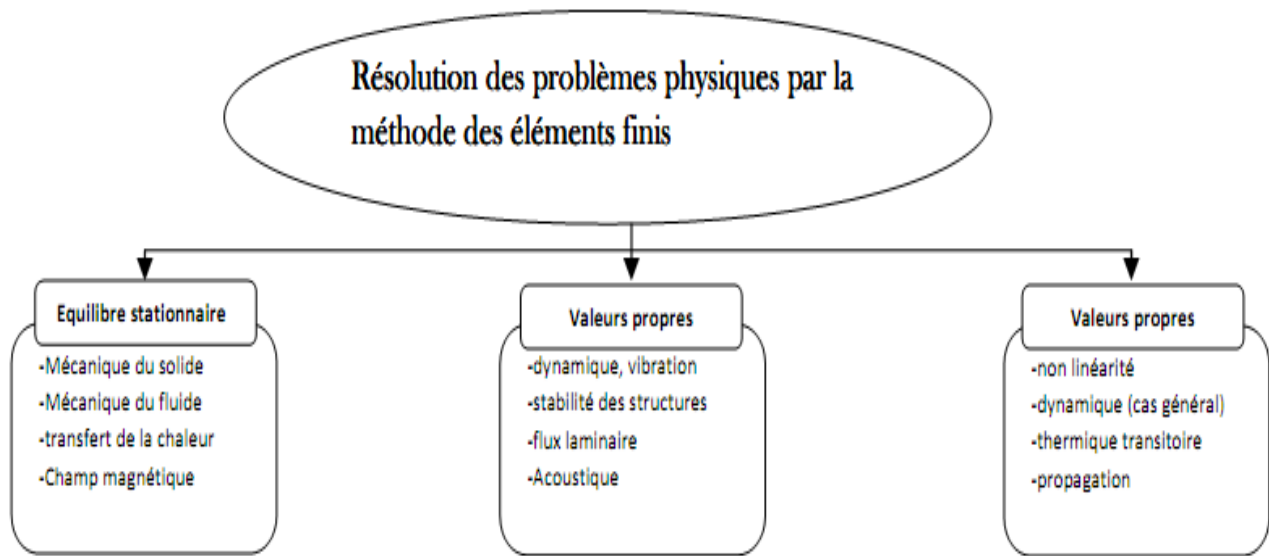


Fig. II. 9 : Différents types des problèmes physiques en éléments finis.

II. 6.1 Problèmes d'équilibre stationnaire

Dans ce type de problèmes, le comportement est défini en fonction de l'état du système, de la géométrie, du chargement et des conditions aux limites, sous forme d'un système d'équations linéaires en fonction des variables nodales. On trouve dans cette catégorie, l'équilibre statique et les régimes stationnaires d'écoulement, de transfert de chaleur et d'électromagnétisme.

II. 6.2 Problèmes aux valeurs propres

Il s'agit des phénomènes de vibration ou d'instabilité d'un état stationnaire. Les modes propres de vibration, le flambage des structures ou l'instabilité des flux laminaires font partie de cette catégorie.

II. 6.3 Problèmes dépendant du temps

Lorsque l'état du système dépend de son histoire ou bien des paramètres de sortie, le système devient interdépendant et la résolution directe n'est plus possible. Ce cas inclut le comportement non linéaire (matériaux et géométrie), la dynamique non linéaire (amortissement, rigidité,...), les régimes transitoires et la fissuration des pièces.

II. 7 Type des éléments finis

On distingue plusieurs classes d'éléments finis suivant leur géométrie :

- Les éléments unidimensionnels (1D) : Sont utilisés de façon individuelle ou associée à des plaques pour modéliser les raidisseurs. Exemple : barre, poutre rectiligne ou courbe-
- Les éléments bidimensionnels (2D) : Élasticité plane : (déformation ou contrainte plane). Exemple : plaque en flexion, coques courbes, de forme triangulaire ou quadrangulaire.
- Les éléments tridimensionnels (3D) : élément de volume, ou coques épaisses. Les éléments axisymétriques : qui constituent une classe bien particulière.

II.8. Les différentes étapes de la résolution par la MEF

Le choix de la structure en éléments finis par un maillage constitue de lignes ou de surfaces imaginaires. Les éléments sont supposés reliés en un nombre fini de points nodaux situés sur leurs frontières.

Les déplacements de ces points nodaux seront les inconnues de base du problème. Il est apparent que la méthode des éléments finis est applicable pour des structures des matériaux de propriétés hétérogènes ou de forme géométrique compliquées et irrégulière (bords courbes, trous...).

On choisit une fonction de déplacement permettant de définir de manière unique le champ des déplacements à l'intérieure de chaque « élément fini » en fonction des déplacements de ces nœuds. On se basant sur cette fonction de déplacement, nous déduisons- la matrice de rigidité de l'élément qui lie les forces nodales avec les déplacements nodaux et la matrice masse en utilisant le principe des travaux virtuels ou le principe de l'énergie potentielle total minimum.

Analyse de la structure idéalisée de l'assemblage des éléments. Cette analyse procède de la manière classique qui a été décrite par la méthode des rigidités.

En fin la solution de ces équations nous permet d'évaluer les déplacements et les efforts internes dans la structure (contrainte, déformation).

La méthode des éléments finis est extrêmement puissante puisqu'elle permet d'étudier correctement des structures continues ayant des propriétés géométriques et des conditions de charge compliquées ; elle nécessite un grand nombre de calculs qui, à cause de leur nature répétitive, s'adaptent parfaitement à la programmation numérique et à la résolution par ordinateur.

Donc pour résoudre un problème par la méthode des éléments finis, on procède par les étapes successives suivantes (figure II. 10) :

1. On pose un problème physique sous la forme d'une équation différentielle ou aux dérivées partielles à satisfaire en tout point d'un domaine Ω , avec des conditions aux limites sur le bord $\partial\Omega$ nécessaires et suffisantes pour que la solution soit unique.
2. On construit une formulation intégrale du système différentiel à résoudre et de ses conditions aux limites : C'est la formulation variationnelle du problème.
3. On divise Ω en sous domaines : C'est le *maillage*. Les sous domaines sont appelés *mailles*.
4. On choisit la famille de champs locaux, c'est à dire à la fois la position des nœuds dans les sous domaines et les polynômes (ou autres fonctions) qui définissent le champ local en

fonction des valeurs aux nœuds (et éventuellement des dérivées). La maille complétée par ces informations est alors appelée *élément*.

5. On ramène le problème à un problème discret : C'est la *discrétisation*. En effet, toute solution approchée est complètement déterminée par les valeurs aux nœuds des éléments. Il suffit donc de trouver les valeurs à attribuer aux nœuds pour décrire une solution approchée.
6. On résout le problème discret: C'est la *résolution*
7. On peut alors construire la solution approchée à partir des valeurs trouvées aux nœuds et en déduire d'autres grandeurs : C'est le *post-traitement*.
8. On visualise et on exploite la solution pour juger de sa qualité numérique et juger si elle satisfait les critères du cahier des charges : C'est *l'exploitation des résultats*.

Les étapes 1, 2, 3,4 et 5 sont souvent rassemblées sous le nom de *prétraitement*.

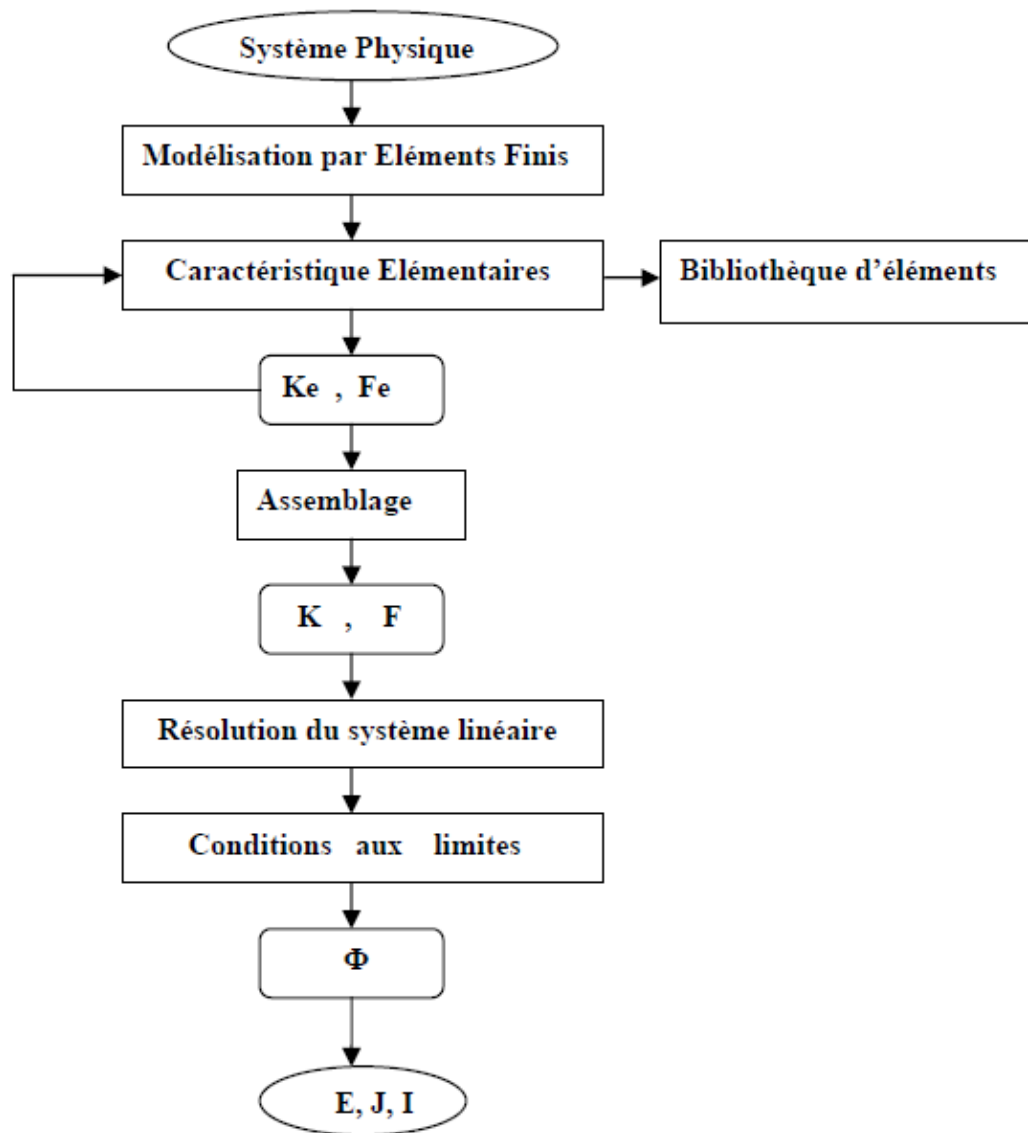


Fig.II. 10.Organigramme simplifié de l'analyse par la méthode des éléments finis.

II.8.1. Définition des mailles de référence

De manière à simplifier la définition analytique des éléments de forme complexe, introduisons la notion d'élément de référence (fig.II.11) : un élément de référence Ω^r est un élément de forme

très simple, repéré dans un espace de référence, qui peut être transformé en chaque élément réel Ω^e par une transformation géométrique. Cette transformation dépend de la forme et de la position de l'élément réel.

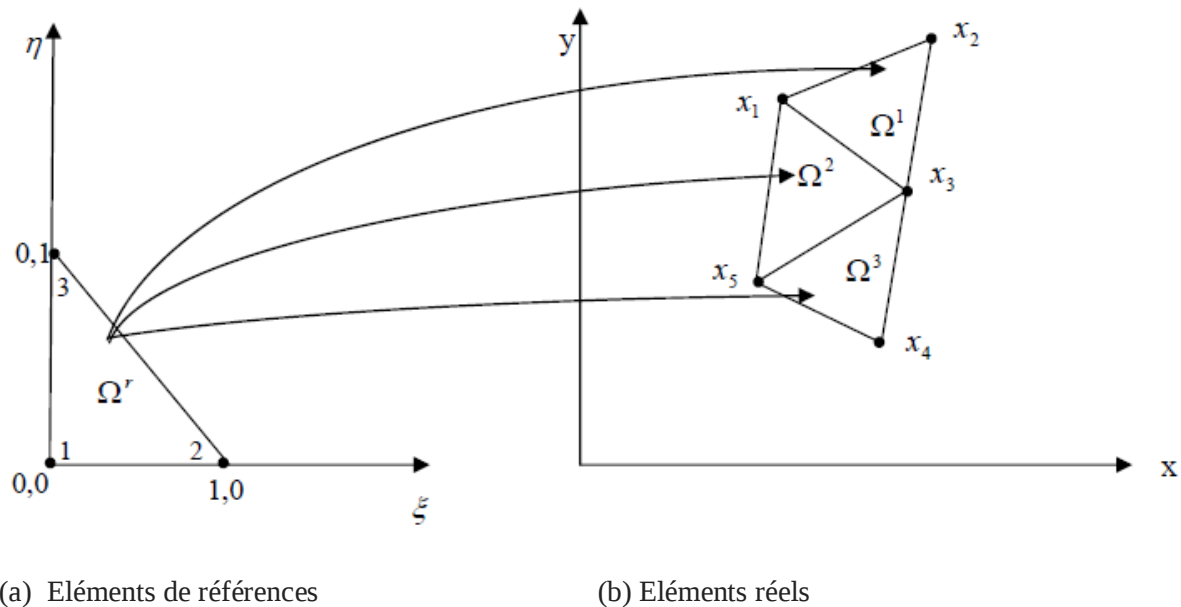


Fig. II.11. Relation entre l'élément réel et l'élément de référence

Chaque transformation est choisie de manière à présenter les propriétés suivantes :

- Elle est bijective en tout point situé sur l'élément de référence ou sur sa frontière : à tout point de Ω' .
- Les nœuds géométriques de l'élément de référence correspondent aux nœuds de l'élément réel.
- Chaque portion de frontière de l'élément de référence, définie par les nœuds géométriques de cette frontière, correspond à la portion de frontière de l'élément réel définie par nœuds correspondants.

Soulignons que même l'élément de référence Ω' (par exemple un triangle à 3 nœuds) se transforme en éléments réels Ω^e de même type par des transformations différentes.

Nous présentons ci-dessous la forme et la définition analytique des éléments de référence correspondant aux éléments classiques. Dans un maillage, toutes les mailles ont des formes et des dimensions différentes. On trouve des mailles linéiques, des mailles surfaciques et des mailles volumiques de toutes formes et de toutes tailles. Dans le but d'uniformiser et d'automatiser les calculs, on introduit la notion de maille de référence. Nous ne donnons ici que les mailles les plus classiques. Par convention généralement admise.

- La maille de référence linéique est le segment (fig.II.12) :

$$x_1 \in [-1,1]$$

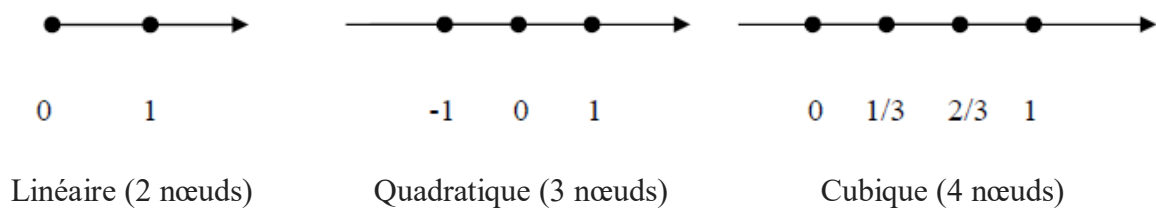


Fig. II.12. Éléments de référence à une dimension

- La maille de référence surfacique triangulaire est le triangle (fig. II.13.):

$$x_1 \geq 0 ; \quad x_2 \geq 0 ; \quad x_1 + x_2 \leq 1$$

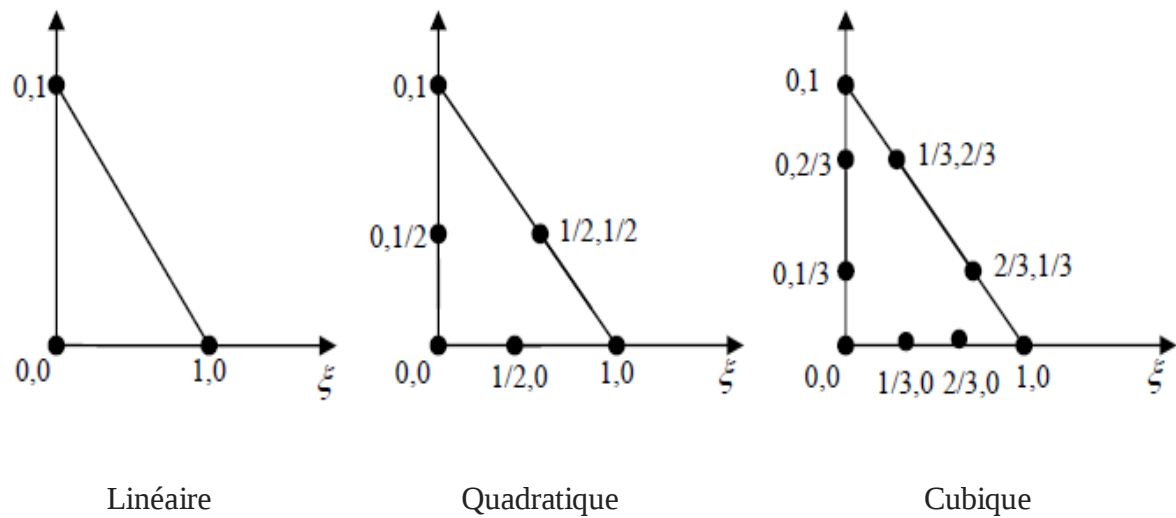


Fig. II.13. Eléments de referens triangulaires.

- La maille de référence surfacique quadrangulaire est le carré (fig. II.14.):

$$x_1 \in [-1,1] ; \quad x_2 \in [-1,1]$$

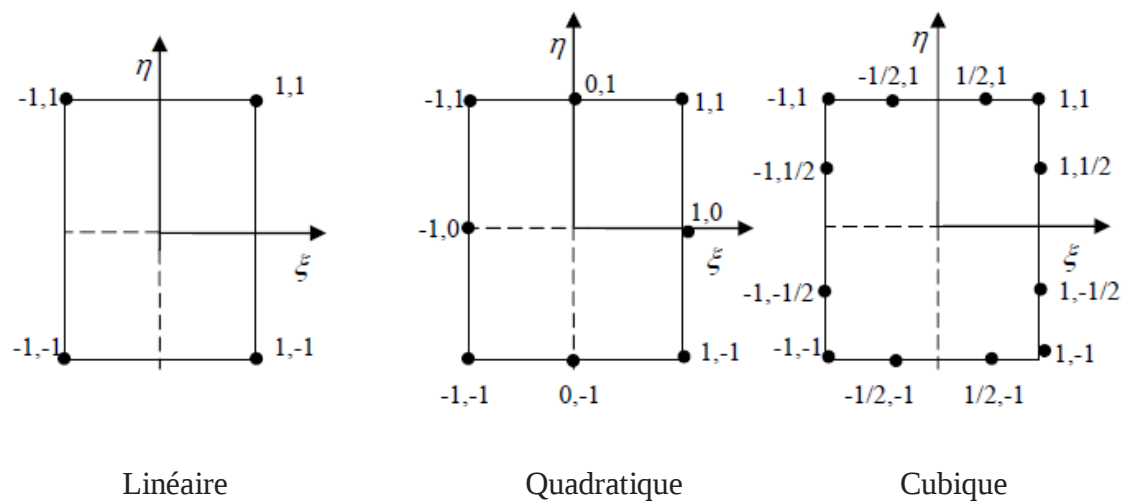


Fig. II.14. Eléments de referens carrés.

- La maille de référence volumique tétraédrique est le tétraèdre (fig. II.15):

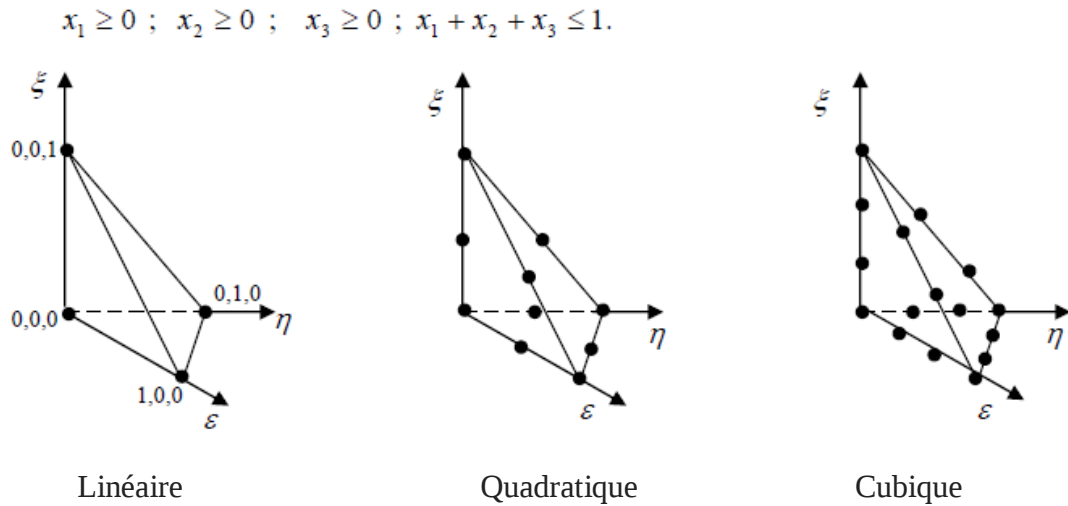


Fig. II.15. Eléments de référence tétraédrique

- La maille de référence volumique hexaédrique est le cube (fig.II. 16):

Soit :

$$x_1 \in [-1,1] ; \quad x_2 \in [-1,1] ; \quad x_3 \in [-1,1]$$

Où x_1, x_2 et x_3 sont les coordonnées d'un point courant m de la maille de référence.

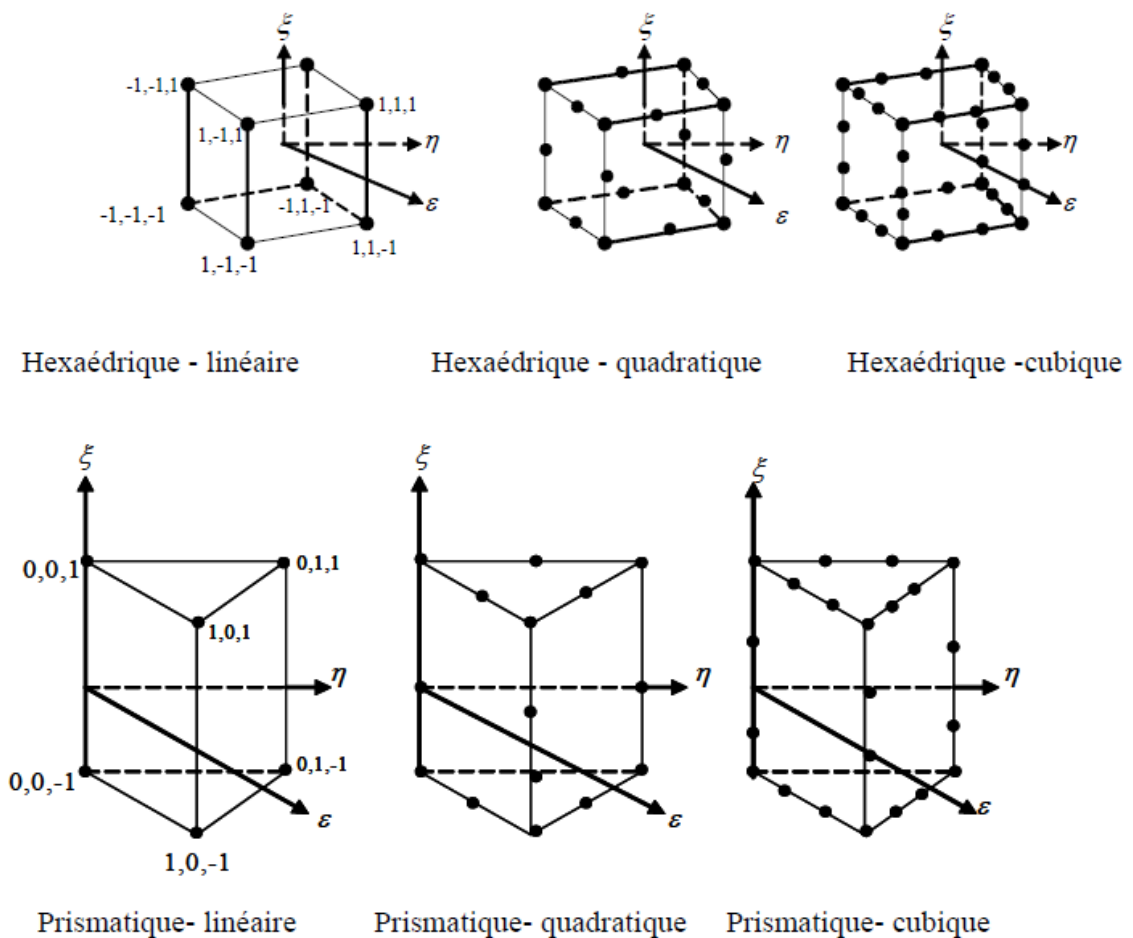


Fig. II.16. La maille de référence volumique hexaédrique.

II.8.2. Interpolation polynomiale sur les mailles de référence

Les concepts de transformation géométrique et d'élément de référence simplifient la construction des fonctions d'interpolation pour des éléments de formes compliquées. Les fonctions $\Phi(x, y, a_1, a_2, \dots, a_n)$ sont souvent choisies de manière à être faciles à évaluer sur ordinateur, à intégrer ou dériver explicitement. Ainsi l'approximation peut fournir :

- Une expression approchée en tout point (x, y) d'une fonction difficile à évaluer ou connue seulement en certains points ;

- Une solution approchée d'une équation différentielle ou aux dérivées partielles.

Le polynôme d'interpolation d'ordre n à deux dimensions, est donné sous la forme:

$$\phi(x, y) = a_1 + a_2x + a_3y + a_4x^2 + a_5y^2 + a_6xy + \dots + a_{2n+1}y^n$$

Pour $n=1$, nous avons l'équation d'interpolation linéaire suivante (fig. 15) :

$$\phi(x, y) = a_1 + a_2x + a_3y$$

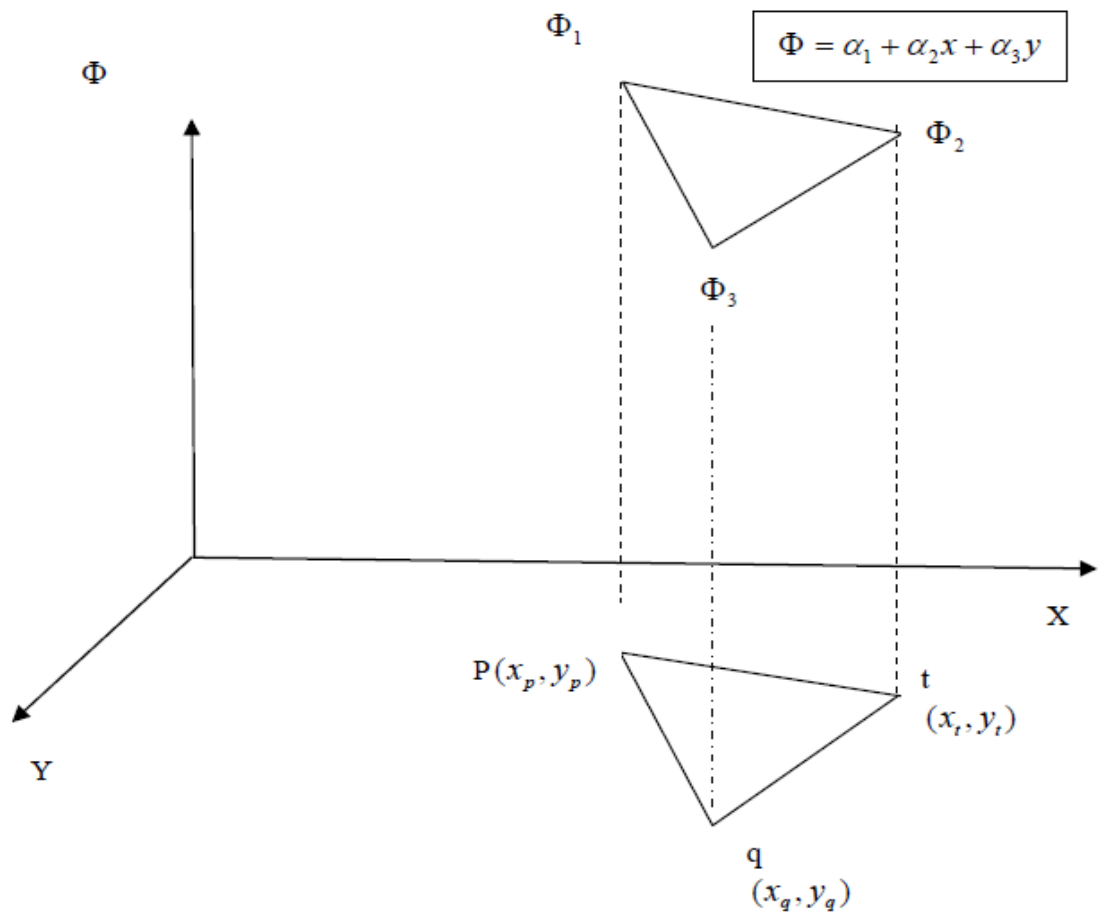


Fig.II.15. Paramètres de l'élément triangulaire linéaire.

Pour $n=2$, nous avons l'équation d'interpolation quadratique sous la forme :

$$\phi(x, y) = a_1 + a_2x + a_3y + a_4x^2 + a_5y^2 + a_6xy.$$

Pour notre cas, nous utilisons l'élément triangulaire comme élément de référence.

Alors nous pouvons utiliser une équation d'interpolation linéaire de premier ordre.

II.8.3. Formulations matricielles de la MEF

Lorsque nous utilisons les techniques matricielles, nous sommes amenés successivement à s'intéresser à deux niveaux de formulation :

- La formulation élémentaire au niveau de l'élément fini ;
- La formulation globale au niveau de la structure complète.

L'équation de Poisson pour deux dimensions, est donnée sous la forme:

$$\frac{\partial^2 \Phi(x, y)}{\partial x^2} + \frac{\partial^2 \Phi(x, y)}{\partial y^2} - G(x, y) = 0$$

D'après le théorème d'Euler, la résolution de cette équation différentielle est équivalente à minimiser la fonctionnelle d'énergie W qui s'écrit :

$$W = \iint_{\Omega} \left[\frac{1}{2} \left[\frac{\partial^2 \Phi}{\partial x^2} + \frac{\partial^2 \Phi}{\partial y^2} \right] - \frac{\rho}{\epsilon_0} \right] dx dy \quad (4)$$

L'intégrale est étendue à un ensemble du domaine de calcul, que l'on subdivise en petits éléments par maillage triangulaire où chaque triangle est repéré par ses trois sommets (nœuds).

Les polynômes d'interpolations du potentiel aux sommets du triangle sont donnés par le système d'équations suivant :

$$\begin{cases} \phi_1 = \alpha_1 + \alpha_2 x_1 + \alpha_3 y_1 \\ \phi_2 = \alpha_1 + \alpha_2 x_2 + \alpha_3 y_2 \\ \phi_3 = \alpha_1 + \alpha_2 x_3 + \alpha_3 y_3 \end{cases}$$

La résolution de ce système, nous donne les coefficients α_i :

$$\begin{cases} \alpha_1 = \frac{1}{2\Delta_e} [a_1 \cdot \phi_1 + a_2 \cdot \phi_2 + a_3 \cdot \phi_3] \\ \alpha_2 = \frac{1}{2\Delta_e} [b_1 \cdot \phi_1 + b_2 \cdot \phi_2 + b_3 \cdot \phi_3] \\ \alpha_3 = \frac{1}{2\Delta_e} [c_1 \cdot \phi_1 + c_2 \cdot \phi_2 + c_3 \cdot \phi_3] \end{cases}$$

Avec :

$$\begin{cases} a_1 = x_2 \cdot y_3 - x_3 y_2 & b_1 = y_2 - y_3 & c_1 = x_3 - x_2 \\ a_2 = x_3 \cdot y_1 - x_1 y_3 & ; & b_2 = y_3 - y_1 & ; & c_2 = x_1 - x_3 \\ a_3 = x_1 \cdot y_2 - x_2 y_1 & b_3 = y_1 - y_2 & c_3 = x_2 - x_1 \end{cases}$$

$$2.\Delta_e = \begin{vmatrix} 1 & x_1 & y_1 \\ 1 & x_2 & y_2 \\ 1 & x_3 & y_3 \end{vmatrix} = b_1 c_2 - b_2 c_1$$

$2.\Delta_e$; est l'aire de l'élément triangulaire.

Le potentiel en tout point d'un élément (e) est exprimé par la relation:

$$\Phi(x, y) = \sum_{i=1}^3 \phi_i \cdot N_i(x, y) = [N_1 N_2 N_3] \begin{bmatrix} \phi_1 \\ \phi_2 \\ \phi_3 \end{bmatrix}$$

Avec :

$$N_i(x, y) = \frac{1}{2 \cdot \Delta_e} (a_i + b_i x + c_i y) \quad ; i=1, 2, 3 .$$

Les fonctions N_i sont appelées fonctions de formes. Elles doivent vérifier la condition suivante :

$$N_i(x_j, y_j) = \delta_{ij} = \begin{cases} 0 & \text{Si } i \neq j \\ 1 & \text{si } i = j \end{cases}$$

La minimisation de la fonctionnelle W est réalisée quand les dérivées partielles par rapport à tous les potentiels des nœuds sont nulles, c'est-à-dire :

$$\frac{\partial W}{\partial \{\phi\}} = 0$$

Avec : $\{\phi\}$ est le vecteur potentiel pour tous les nœuds du système.

L'équation (4) montre que nous pouvons restreindre l'intégrale double au domaine formé des n éléments $(e_i)_n$ comportant le sommet i . Nous déduisons:

$$\frac{\partial W}{\partial \phi_i} = \sum_n \iint_{(e_i)_n} \left[\left(\frac{\partial \phi_{(e)}}{\partial x} \frac{\partial}{\partial \phi_i} \left(\frac{\partial \phi_{(e)}}{\partial x} \right) + \frac{\partial \phi_{(e)}}{\partial y} \frac{\partial}{\partial \phi_i} \left(\frac{\partial \phi_{(e)}}{\partial y} \right) \right) - \frac{\rho}{\varepsilon_0} \frac{\partial \phi_{(e)}}{\partial \phi_i} \right] dx \cdot dy \quad (5)$$

Ce qui devient, compte tenu de (4) à (5) :

$$\frac{\partial W}{\partial \phi_i} = \sum_n \iint_{(\epsilon_i)_n} \left[\left(\frac{b_i}{2\Delta} \sum_j \frac{b_j \phi_j}{2\Delta} + \frac{c_i}{2\Delta} \sum_j \frac{c_j \phi_j}{2\Delta} \right) - \frac{\rho}{\epsilon_0} N_i \right] dx.dy$$

Nous aboutissons à un système linéaire de forme matricielle :

$$K.\Phi = F$$

Avec : K : Matrice de raideur ;

Φ : Matrice des potentiels aux nœuds ;

F : Second membre du système.

Le terme K_{ij} se calcule par intégration sur les triangles comportant les nœuds i et j , soit :

$$K_{ij} = \sum_n \iint_{(\epsilon_i)_n} \frac{b_i b_j + c_i c_j}{4\Delta^2} dx.dy = \frac{1}{4\Delta} \sum_n (b_i b_j + c_i c_j) = \iint_{\Delta_e} \left[\frac{\partial N_i}{\partial x} \frac{\partial N_j}{\partial y} + \frac{\partial N_i}{\partial y} \frac{\partial N_j}{\partial x} \right] dx dy$$

Tandis que le second membre F est déterminé par :

$$F_i = \sum_n \iint_{(\epsilon_i)_n} \frac{\rho}{\epsilon_0} N_i dx.dy = \sum_n \frac{\rho}{\epsilon_0} \frac{\Delta}{3}$$

La matrice K comporte un fort pourcentage de termes nuls, car K_{ij} n'est calculé que si le nœud i est relié au nœud j .

Pour le calcul de ces intégrales, nous utilisons la transformation suivante:

$$\iint_{\Delta_e} (N_1)^l (N_2)^m (N_3)^n dx dy = \frac{l!m!n!}{(l+m+n+2)} \quad ; \quad i,j = 1, 2, 3.$$

II.8.4. Assemblage

La phase d'assemblage consiste à construire la matrice globale K et le vecteur global F de la structure complète à partir de la matrice et de vecteur caractéristiques des différents éléments K_e , F_e préalablement calculés. L'assemblage s'effectue en additionnant bloc à bloc les sous – matrices nodales de chaque élément, les indices de ligne et colonne correspondant à la numérotation des nœuds de cet élément dans la matrice globale $[K]$ et le vecteur global $\{F\}$.

Avec :

$$\begin{cases} [K] = \sum_{e=1}^{N=\text{nbre. d'éléments}} [K_e] \\ \{F\} = \sum_{e=1}^{N=\text{nbre. d'éléments}} \{F_e\} \end{cases}$$

Une fois K et F construites, les potentiels inconnus sont obtenus en multipliant l'inverse de K par F . Dans le cas de l'équation de Laplace la même procédure est à suivre en prenant $\rho = 0$.

II.8.5. Introduction des conditions aux limites

Après l'assemblage, la forme intégrale globale s'écrit :

$$W = \langle \delta\Phi \rangle ([K] \{ \Phi \} - \{ F \}) = 0$$

Le problème consiste à trouver $\{ \Phi \}$ qui annule W pour tout $\langle \delta\Phi \rangle$ en satisfaisant les conditions aux limites.

Donc le système algébrique : $[K] \{ \Phi \} = \{ F \}$ doit être modifié en introduisant les conditions aux limites.

II.8.6. Conditions de convergence de la solution

La méthode des éléments finis fournit une solution approchée qui converge vers la solution exacte lorsque l'on diminue la taille des éléments, si l'approximation de Φ satisfait aux conditions suivantes :

- **Base polynomiale complète :**

Pour que la solution approchée tende vers la solution exacte lorsque la taille h des éléments tend vers zéro, il faut que l'erreur d'approximation de tous les termes We soit d'ordre h_n avec $n \geq 1$. Nous avons vu que l'approximation de Φ doit utiliser au moins une base polynomiale complète jusqu'à l'ordre m pour assurer la convergence des dérivées de Φ d'ordre m .

- **Continuité :**

A la condition locale précédente, il faut ajouter une condition globale concernant la continuité des approximations de Φ et de ses dérivées entre les éléments de manière à pouvoir écrire la solution du problème sans la discontinuité.

- **La fonction Φ** et toutes ses dérivées jusqu'à l'ordre m qui apparaissent dans les équations doivent être bornées et continues sur les frontières entre éléments ; dans ce cas un élément est dit conforme.

II.9. Avantages de la MEF

- La flexibilité est l'un des plus importants avantages de la MEF. Les éléments peuvent avoir plusieurs formes variées et peuvent donc s'adapter facilement à n'importe quelles formes géométriques complexes et aussi tenir compte des propriétés inhomogènes et non linéaires des matériaux.
- La programmation de la méthode est assez simple surtout lorsqu'il s'agit de tenir compte de l'introduction des conditions aux limites.
- La MEF a fait ses preuves dans beaucoup de domaine en ingénierie. De plus, avec son développement important, il existe de très bons logiciels commerciaux qui sont basés sur

cette méthode et qui la rendent très accessible. Et par conséquent, elle est applicable à beaucoup de problèmes sans que nous connaissions nécessairement la MEF en détail.

- Les matrices formant le système final d'équations sont symétriques ce qui simplifie grandement la résolution de celui ci.

II.10. Inconvénients de la MEF

- L'utilisation de la MEF pour la résolution d'un problème donné, nécessite la connaissance parfaite de la géométrie du problème mais aussi des conditions aux limites, ce qui n'est pas toujours le cas.
- Une fois le potentiel connu en chaque nœud, il faut procéder à un autre calcul numérique pour déterminer le champ électrique en tout point. Ce qui peut engendrer d'autres erreurs.
- Il a été dit que la MEF était une méthode flexible car elle s'adapte facilement aux différentes géométries, mais ce n'est pas le cas du maillage car celui ci doit être entièrement refait si une modification sur une partie de la géométrie du problème considéré intervient.

II.11. Conclusion

La méthode des éléments finis est l'un des outils les plus efficaces et plus généraux de la simulation numérique. L'application des conditions d'équilibre et des lois de comportement de la physique permet de construire des équations approchées dont les inconnues sont les valeurs de la solution en un ensemble bien choisi de points, appelés les nœuds de la discrétisation.

Une simulation réaliste peut exiger des centaines de milliers de nœuds et d'éléments.

III.1. Introduction

L'optimisation et particulièrement l'optimisation numérique a connu un départ important ces dernières années avec l'apparition de l'ordinateur. Elle est souvent la dernière étape de l'analyse numérique où, après avoir étudié un phénomène physique, l'avoir mis en équation, avoir étudié ces équations et avoir montré que l'on pouvait calculer les solutions avec un ordinateur, on commence à optimiser le système en changeant certains paramètres pour changer la solution dans un sens désiré.

III.2. Définition et formulation

On distingue trois étapes pour la modélisation mathématique de problème d'optimisation :

- Identification des variables de décisions : ce sont les paramètres sur lesquels l'utilisateur peut agir pour faire évoluer le système considéré.
- Définition d'une fonction coût ou fonction objectif : permettant d'évaluer l'état du système (ex : rendement, performance,...).
- Description des contraintes imposées aux variables de décision.

Le problème d'optimisation consiste alors à déterminer les variables de décision conduisant aux meilleures conditions de fonctionnement du système (ce qui revient à minimiser ou maximiser la fonction coût), tout en respectant les contraintes d'utilisation définies à l'étape 3.

III.2.1. Notions de minimum, maximum, infimum, supremum

On classera les notions de minimum et de maximum des notions d'infimum et de supremum.

Ces notions sont des prérequis pour les démonstrations des résultats d'existence et d'unicité d'extrema d'une fonction donnée.

- **Minorant / Majorant** : Soit E un sous-ensemble de $\mathbb{R}$, $m \in \mathbb{R} \cup \{-\infty, +\infty\}$ est un minorant de E si m est inférieur ou égal à tous les éléments de E tandis que $M \in \mathbb{R} \cup \{-\infty, +\infty\}$ est un majorant de E si il est supérieur ou égal à tous les éléments de E . Ainsi :

(m minorant $\Leftrightarrow \forall x \in E, m \leq x$) et (M majorant $\Leftrightarrow \forall x \in E, M \geq x$):

Si E admet un minorant (respectivement. majorant) fini alors il est dit minoré (respectivement. majoré).

- **Infimum / Supremum** : Soit $E \subset \mathbb{R}$. L'infimum $\inf(E) \in \mathbb{R} \cup \{-\infty, +\infty\}$ de E est le plus grand des minorants. Le supremum $\sup(E) \in \mathbb{R} \cup \{-\infty, +\infty\}$ de E est le plus petit des majorants.

On les note respectivement :

$$\inf(E) = \inf_{x \in E} (x) \quad \text{et} \quad \sup(E) = \sup_{x \in E} (x)$$

Nous avons parlé d'infimum et de supremum, nous les relierons maintenant aux définitions classiques de minimum et de maximum.

- **Minimum / maximum** : Soit $E \subset \mathbb{R}$. L'infimum de E est appelé minimum si $\inf(E) \in E$.

Le supremum de E est appelé maximum si $\sup(E) \in E$. Dans ce cas, on les note respectivement $\min(E)$ et $\max(E)$.

Exemple :

Si on se donne $E =]0, 1]$, alors l'ensemble des minorants de E est $[-\infty, 0]$ et l'ensemble des majorants est $[1, +\infty]$. On en déduit que le l'infimum vaut 0 et le supremum vaut 1. Comme 1 appartient à E alors c'est aussi le maximum de E . Cependant 0 n'appartient pas à E et donc E n'a pas de minimum. On dit souvent que l'infimum de E n'est pas atteint.

La formulation et l'analyse d'un problème d'optimisation consiste à chercher la minimisation (ou maximisation) d'une fonction objectif (de coût).

Cette fonction comporte des paramètres et est généralement soumise à des contraintes en fonction des caractéristiques des problèmes (paramètres fortement variables, domaines discrets,

types de contraintes, etc.), le problème peut être difficile à résoudre et nécessite des outils avancés pour la modélisation et la résolution comme :

- Variables (paramètres)
- Fonction objectif
- Contraintes

III.2.2. Description d'un problème d'optimisation

Comme nous l'avons vu en introduction, tous les problèmes d'optimisation que nous considérerons peuvent être exprimés de la façon suivante :

$$\min_{x \in \mathbb{R}^n} f(x)$$

Sous la contrainte : $x \in X$:

Où X est un sous-ensemble de $\mathbb{R}^n$. On pourra écrire aussi

$$\min_{x \in X} f(x)$$

Les variables $x = (x_1, \dots, x_n)$ sont appelées "variables d'optimisation" ou variables de décision (sont appelés également points admissibles).

La fonction $f : X \subset \mathbb{R}^n \rightarrow \mathbb{R}$ est appelée fonction coût, objectif ou critère et l'ensemble X est appelé ensemble des contraintes ou domaine des contraintes.

Chercher une solution du problème avec des contraintes, revient à chercher un point de minimum local de f dans l'ensemble des points admissibles, au sens de la définition suivante :

- $x \in \mathbb{R}^n$ est un point de minimum local de f sur $X \subset \mathbb{R}^n$ si : $x \in X$

Et

$$\exists r > 0 \mid \forall y \in X \cap B(x, r), \quad f(x) \leq f(y)$$

→ On dit alors que $f(x)$ est un minimum local de f sur X .

➤ $x \in \mathbb{R}^n$ est un point de minimum global de f sur X si : $x \in X$ et

$$\forall y \in X, \quad f(x) \leq f(y)$$

→ On dit alors que $f(x)$ est un minimum global de f sur X .

Les notions de maximum local et global sont définies de façon tout à fait analogue. En fait, on peut facilement démontrer que les problèmes (avec ou sans contraintes) :

$$\min_{x \in X} f(x) \quad \text{et} \quad \max_{x \in X} -f(x)$$

Sont équivalents dans le sens où ils ont même ensemble de solutions et :

$$\min_{x \in X} f(x) = - \max_{x \in X} -f(x)$$

Ou encore :

$$\max_{x \in X} f(x) = - \min_{x \in X} -f(x)$$

Ainsi la recherche d'un maximum pouvant se ramener à la recherche d'un minimum, nous supporterons une attention plus particulière à la recherche du minimum.

III.2.3. Types d'optimisation

Les problèmes d'optimisation peuvent être classés en plusieurs grandes familles :

- Optimisation numérique : $X \subset \mathbb{R}^n$.
- Optimisation discrète (ou combinatoire) : X fini ou mesurable.
- Commande optimale : X est un ensemble de fonctions.
- Optimisation stochastique : données aléatoires (à ne pas confondre avec les méthodes stochastiques d'optimisation).
- Optimisation multicritères : plusieurs fonctions objectives.

III.2.4. Algorithme d'optimisation

Un algorithme d'optimisation est aussi jugé sur les propriétés (théoriques) de $(x_k)_k$, en particulier propriété de convergence :

- Convergence globale : pour tout point initial $x_0 \in \mathbb{R}^n$.
- Convergence locale.

Nous nous intéressons dans cette partie à la conception de méthodes numériques pour la recherche des points $x \in \mathbb{R}^n$ qui réalisent le minimum d'une fonction f : où f est une fonction définie sur $\mathbb{R}^n$ à valeurs réelles supposée différentiable, voire même deux fois différentiable. Les conditions nécessaires d'optimalité du premier et du second ordre expriment le fait qu'il n'est pas possible de "descendre" à partir d'un point de minimum (local ou global). Cette observation va servir de point de départ à l'élaboration des méthodes dites de descente rappelées dans cette partie.

Partant d'un point x_0 arbitrairement choisi, un algorithme de descente va chercher à générer une suite d'itérés $(x_k)_{k \in \mathbb{N}}$ définie par :

$x_{k+1} = x_k + s_k d_k$ telle que :

$$\forall k \in \mathbb{N}, \quad f(x_{k+1}) \leq f(x_k).$$

Un tel algorithme est ainsi déterminé par deux éléments : le choix de la direction d_k appelée direction de descente, et le choix de la taille du pas s_k à faire dans la direction d_k . Cette étape est appelée recherche linéaire.

Un vecteur $d \in \mathbb{R}^n$ est une direction de descente pour f à partir d'un point $x \in \mathbb{R}^n$ si :

$t \rightarrow f(x + t d)$ est décroissante en $t = 0$, c'est-à-dire s'il existe $\eta > 0$ tel que :

$$\forall t \in]0, \eta], \quad f(x + t d) < f(x).$$

Il est donc important d'analyser le comportement de la fonction f dans certaines directions.

Lorsqu'elle existe, la dérivée directionnelle donne des informations sur la pente de la fonction dans la direction d , en déduisant donc :

- Le vecteur $d \in \mathbb{R}^n$ est une direction de descente de f au point $x \in \mathbb{R}^n$ si $f'(x; d) < 0$.
- Soit $f : \mathbb{R}^n \rightarrow \mathbb{R}$ une fonction différentiable. Soit $x \in \mathbb{R}^n$. Alors pour toute direction d de norme constante égale à :

$$\|d\| = \|\nabla f(x)\|$$

On a :

$$(-\nabla f(x))^T \nabla f(x) \leq d^T \nabla f(x)$$

Ainsi, la direction d^* est appelée "direction de plus forte descente".

$$d^* = -\nabla f(x)$$

D'après la caractérisation de la descente, il s'agit donc à chaque itération k , de trouver un point x_{k+1} dans une direction d vérifiant :

$$\nabla f(x_k)^\top d < 0.$$

- Algorithme de descente Modèle

Données: $f : \mathbb{R}^n \rightarrow \mathbb{R}$ différentiable, x_0 point initial arbitraire.

Sortie: une approximation de la solution du problème : $\min_{x \in \mathbb{R}^n} f(x)$.

1. $k := 0$
2. Tant que "test d'arrêt" non satisfait,
 - (a) Trouver une direction de descente d_k telle que : $\nabla f(x_k)^\top d_k < 0$.
 - (b) *Recherche linéaire* : Choisir un pas $s_k > 0$ à faire dans cette direction et tel que :

$$f(x_k + s_k d_k) < f(x_k).$$

- (c) Mise à jour : $x_{k+1} = x_k + s_k d_k$; $k := k + 1$;
3. Retourner x_k .

- Structure d'un algorithme d'optimisation

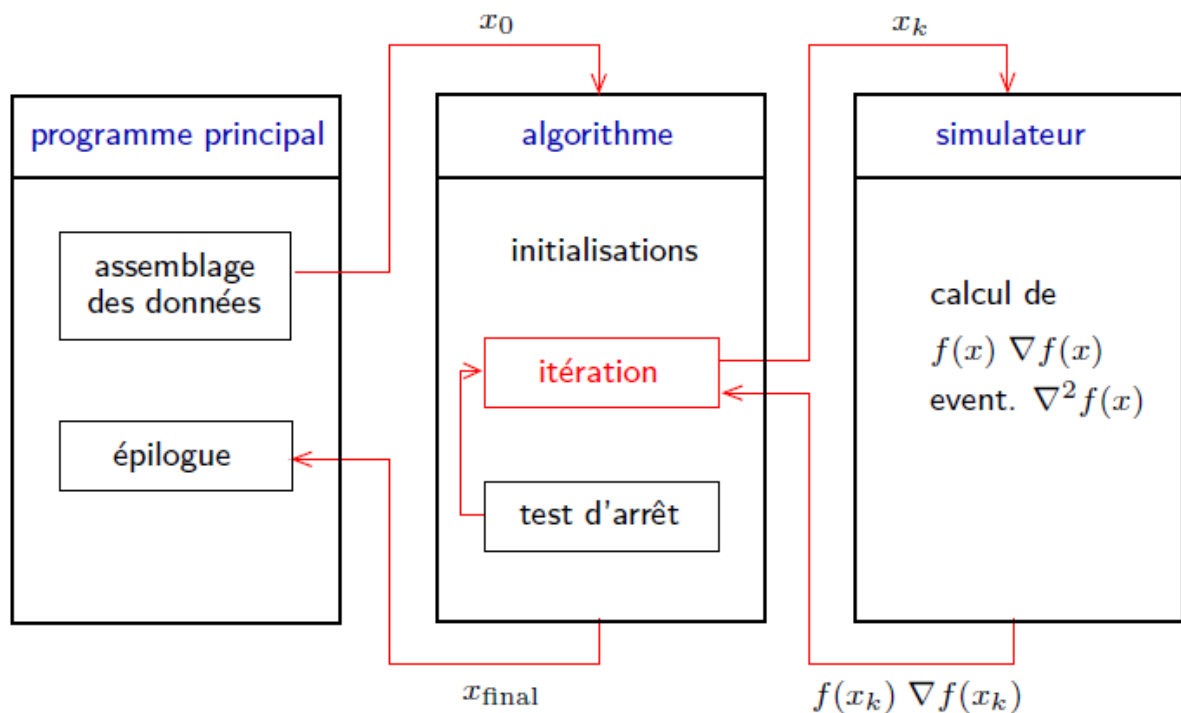


Fig. III.1. Structure d'un algorithme d'optimisation

III.3. Optimisation sans contraintes

On s'intéresse ici aux algorithmes de descente qui sont déterminés par les stratégies de choix des directions de descente successives, puis par le pas qui sera effectué dans la direction choisie. Concentrons-nous également, dans cette partie sur le choix de la direction de descente : l'idée est de remplacer f par un modèle local plus simple, dont la minimisation nous donnera une direction de descente de f . On distingue donc :

➤ **Algorithmes de type gradient**

Données: f , x_0 première approximation de la solution cherchée, $\varepsilon > 0$ précision demandée

Sortie: une approximation x^* de la solution de : $\nabla f(x) = 0$

1. $k := 0$;
2. Tant que critère d'arrêt non satisfait,
 - (a) *Direction de descente* : $d_k = -\nabla f(x_k)$.
 - (b) *Recherche linéaire* : trouver un pas s_k tel que : $f(x_k + s_k d_k) < f(x_k)$.
 - (c) $x_{k+1} = x_k - s_k \nabla f(x_k)$; $k := k + 1$;
3. Retourner x_k .

➤ **Algorithmes de type Newton locale**

Données: $f : \mathbb{R}^n \rightarrow \mathbb{R}$ de classe C^2 , x_0 première approximation de la solution cherchée, $\varepsilon > 0$ précision demandée

Sortie: une approximation x^* de la solution

1. $k := 0$;
2. Tant que $\|\nabla f(x_k)\| > \varepsilon$,
 - (a) Calculer d_k solution du système : $H[f](x_k)d_k = -\nabla f(x_k)$;
 - (b) $x_{k+1} = x_k + d_k$;
 - (c) $k := k + 1$;
3. Retourner x_k ;

➤ **Algorithmes de type Gauss-Newton**

Données: F fonction différentiable, x_0 point initial, $\varepsilon > 0$ précision demandée

Sortie: une approximation de la solution du problème de moindres carrés :

$$\min_{x \in \mathbb{R}^n} r(x) = \frac{1}{2} F(x)^\top F(x).$$

1. $k := 0$;
2. Tant que ,
 - (a) *Calcul d'une direction de recherche* : calculer d_{k+1} solution de :

$$J_F(x_k)^T J_F(x_k) d = -J_F(x_k)^\top F(x_k).$$

- (b) $x_{k+1} = x_k + d_{k+1}$;
 - (c) $k := k + 1$;
3. Retourner s_k .

➤ **Algorithmes de type Levenberg-Marquardt**

L'algorithme de Levenberg-Marquardt peut être vu comme une régularisation de l'algorithme de Gauss-Newton, en particulier lorsque la jacobienne de f n'est pas de rang plein.

D'où :

Données: F fonction différentiable, x_0 point initial, $\varepsilon > 0$ précision demandée.

Sortie: une approximation de la solution du problème de moindres carrés :

$$\min_{x \in \mathbb{R}^n} r(x) = \frac{1}{2} F(x)^\top F(x).$$

1. $k := 0$;
2. Tant que *critère d'arrêt à définir*,
 - (a) *Calcul d'une direction de recherche* : calculer d_k solution de :

$$(J_F(x_k)^T J_F(x_k) + \lambda I) d = -J_F(x_k)^T F(x_k).$$
 - (b) $x_{k+1} = x_k + d_k$;
 - (c) Mise à jour du paramètre λ .
 - (d) $k := k + 1$;
3. Retourner x_k .

III.4. Optimisation sous contraintes

Cette partie est une courte introduction à l'optimisation sous contraintes. On s'intéresse à la résolution de problèmes d'optimisation de la forme :

$$\min_{x \in X} f(x)$$

Où X est un sous-ensemble non vide de $\mathbb{R}^n$ défini par des contraintes d'égalité ou d'inégalités de fonctions :

$$X = \{x \in \mathbb{R}^n : h(x) = 0, g(x) \leq 0\}$$

Où $h : \mathbb{R}^n \rightarrow \mathbb{R}^p$ et $g : \mathbb{R}^n \rightarrow \mathbb{R}^q$ sont continues. Ici les écritures $h(x) = 0$ et $g(x) \leq 0$ signifient :

$$\begin{cases} \forall i = 1, \dots, p, h_i(x) = 0 \\ \forall j = 1, \dots, q, g_j(x) \leq 0 \end{cases}$$

L'ensemble X est appelé ensemble ou domaine des contraintes. Tout point $x \in \mathbb{R}^n$ vérifiant :

$x \in X$, est appelé point admissible du problème d'optimisation.

D'après les résultats précédents, rappelons que si g et h sont continues alors X est un ensemble fermé (mais non nécessairement borné). Dans la suite nous travaillerons toujours avec un problème sous "forme standard" c'est-à-dire un problème écrit sous la forme suivante :

$$(P) \left| \begin{array}{l} \min_{x \in \mathbb{R}^n} f(x) \\ \text{s.c.} \quad h_i(x) = 0, \quad i = 1, \dots, p \\ \quad \quad g_j(x) \leq 0, \quad j = 1, \dots, q. \end{array} \right.$$

Par exemple, la forme standard du problème :

$$\max_{(x,y) \in \mathbb{R}^2} f(x,y) \text{ s.t. : } x^2 + y^2 \geq 1 \text{ et } x + y = 5,$$

Est :

$$\min_{(x,y) \in \mathbb{R}^2} -f(x,y) \text{ s.t. : } 1 - x^2 - y^2 \leq 0 \text{ et } x + y - 5 = 0.$$

Le problème (P), écrit sous forme standard, est dit convexe si h est affine, g convexe et si la fonction objectif f est convexe sur X

Le but de cette section est d'élaborer un algorithme de résolution du problème :

Minimiser $f(x)$, $x \in \mathbb{R}^n$, sous la contrainte : $x \in X$,

On distingue donc :

➤ Algorithmes du gradient projeté

Données: f , p_X l'opérateur de projection sur X , x_0 un point initial et $\varepsilon > 0$ la précision demandée

Sortie: une approximation x^* de la solution

1. $k := 0$;
2. Tant que critère d'arrêt non satisfait,
 - (a) $y_k = x_k - s \nabla f(x_k)$ où s est le pas calculé par la méthode de gradient choisie ;
 - (b) *Projection sur X* : $x_{k+1} = p_X(y_k)$; $k := k + 1$;
3. Retourner x_k .

Il est important de remarquer que le calcul à l'étape 2(b) du projeté sur X , peut parfois être aussi difficile que le problème initial. En effet y_k est obtenu en résolvant le problème :

$$\min_{y \in \mathbb{R}^n} \frac{1}{2} \|x_k - s \nabla f(x_k) - y\|_2^2$$

Il s'agit donc de résoudre un problème d'optimisation sur un convexe, avec une fonction objectif convexe. Lorsque le domaine X des contraintes est simple (contraintes de bornes en particulier), c'est faisable. Dès que les contraintes ne sont pas des contraintes de bornes, le calcul de la projection devient beaucoup plus délicat.

Vérifions que la direction $d_k = x_{k+1} - x_k$, si elle est non nulle, est bien une direction de descente de f en x_k .

➤ **Algorithmes newtoniens – Méthode SQP avec contraintes d'égalité.**

Comme son nom l'indique, la méthode SQP (Sequential Quadratic Programming) consiste à remplacer le problème initial par une suite de problèmes quadratiques sous contraintes linéaires plus faciles à résoudre. L'algorithme est le suivant :

Données: $f : \mathbb{R}^n \rightarrow \mathbb{R}$, $h : \mathbb{R}^n \rightarrow \mathbb{R}^p$ différentiables, x_0 point initial, $\lambda_0 \in \mathbb{R}^p$ multiplicateur initial, $\varepsilon > 0$ précision demandée.

Sortie: une approximation x^* de la solution.

1. $k := 0$;

2. Tant que $\|\nabla L(x^k; \lambda^k)\| > \varepsilon$,

(a) Résoudre le sous-problème quadratique :

$$(QP_k) \quad \begin{cases} \min_{d \in \mathbb{R}^n} & \nabla f(x_k)^\top d + \frac{1}{2} d^\top H_k d \\ \text{s.t.} & h_i(x^k) + \nabla h_i(x^k)^\top d = 0. \end{cases}$$

et obtenir la solution primale d_k et le multiplicateur λ' associé à la contrainte d'égalité.

(b) $x^{k+1} = x_k + d_k$; $\lambda^{k+1} = \lambda'$; $k = k + 1$;

3. Retourner x_k .

➤ **Algorithmes newtoniens – Méthode SQP avec contraintes d'égalité et d'inégalité.**

Données: $f : \mathbb{R}^n \rightarrow \mathbb{R}$, $g : \mathbb{R}^n \rightarrow \mathbb{R}^q$, $h : \mathbb{R}^n \rightarrow \mathbb{R}^p$ différentiables, x_0 point initial, $\lambda_0 \in \mathbb{R}_+^q$ et $\mu_0 \in \mathbb{R}^p$ multiplicateurs initiaux, $\varepsilon > 0$ précision demandée.

Sortie: une approximation x^* de la solution.

1. $k := 0$;
2. Tant que $\|\nabla L(x^k; \lambda^k)\| > \varepsilon$,

(a) Résoudre le sous-problème quadratique :

$$(QP_k) \quad \left| \begin{array}{ll} \min_{d \in \mathbb{R}^n} & \nabla f(x_k)^\top d + \frac{1}{2} d^\top H_k d \\ \text{s.t.} & g_j(x^k) + \nabla g_j(x^k)^\top d = 0, \quad i = 1, \dots, q, \\ & h_i(x^k) + \nabla h_i(x^k)^\top d = 0, \quad i = 1, \dots, p. \end{array} \right.$$

et obtenir la solution primale d_k et les multiplicateurs λ' et μ' associés aux contraintes d'inégalité et d'égalité respectivement.

(b) $x^{k+1} = x_k + d_k$; $\lambda^{k+1} = \lambda'$; $\mu^{k+1} = \mu'$; $k = k + 1$;

3. Retourner x_k .

III.5. Conclusion

L'optimisation différentiable sans contraintes nous apprend que :

- Les algorithmes utilisent une information locale - mais on les globalise
- Les méthodes efficaces exigent d'utiliser de l'information du 2nd ordre

Mêmes notions et les mêmes techniques de base en :

- Optimisation avec contraintes
- Optimisation non-différentiable

Choix d'une méthode :

- dépend de la structure de la fonction – mais souvent plusieurs méthodes sont disponibles
- Critères : taille des problèmes, contraintes pratiques et surtout temps de calcul de $f(x)$, $\Delta f(x)$ et $\Delta^2 f(x)$.

Université de M'sila
Faculté de Technologie

Département de Génie électrique
1^{ère} Année Master
Option : Electromécaniques

Module : Méthodes numériques appliquées

TP N° 01 : Introduction au MATLAB

1- But du TP :

L'objectif de ce TP est de faire une introduction sur l'environnement MATLAB dans le calcul des matrices et les différentes opérations, dans ce cas on va citer quelques exemples.

L'étudiant va programmer sous l'environnement MATLAB pour : voir les solutions recherchées.

2- Travail de préparation et d'exécution

-Ouvrir l'environnement MATLAB, calculer :

```
» A=[16 3 2 13; 5 10 11 8; 9 6 7 12; 4 15 14 1]
```

```
A =
```

```
    16     3     2    13
     5    10    11     8
     9     6     7    12
     4    15    14     1
```

```
» A=[2 3 5 8;7 6 8 9;4 8 3 5;1 3 4 8]
```

```
A =
```

```
     2     3     5     8
     7     6     8     9
     4     8     3     5
     1     3     4     8
```

```
» A(1:3,3)
```

```
ans =
```

```
     5
     8
     3
```

```
» A(:,3)
```

```
ans =
```

```
     5
     8
     3
     4
```

```

» A=[3 4 5;6 7 8];
» B=[3 4 7 8;5 6 11 3;7 7 8 13];
» C=A*B
C =
    64    71   105   101
   109   122   183   173
» A=[1 2 3;4 5 6];
» B=[1 2 3;4 5 6];
» A.*B
ans =
    1     4     9
   16    25    36
» A=[3 4 5;6 7 8; 5 8 6];
» B=[3 4 7;5 6 11;7 7 8];
» inv(B)*A
ans =
    2.0833    1.5833    0.2500
   -2.4167   -0.9167   -0.2500
    0.9167    0.4167    0.7500

```

```

» det(A)
ans =
    15

```

```

» inv(A)
ans =
   -1.4667    1.0667   -0.2000
    0.2667   -0.4667    0.4000
    0.8667   -0.2667   -0.2000
» A*inv(A)

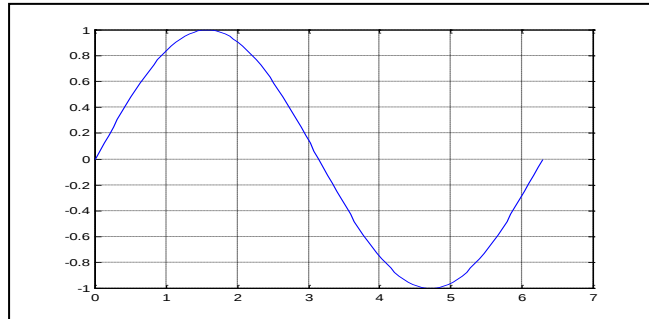
```

```

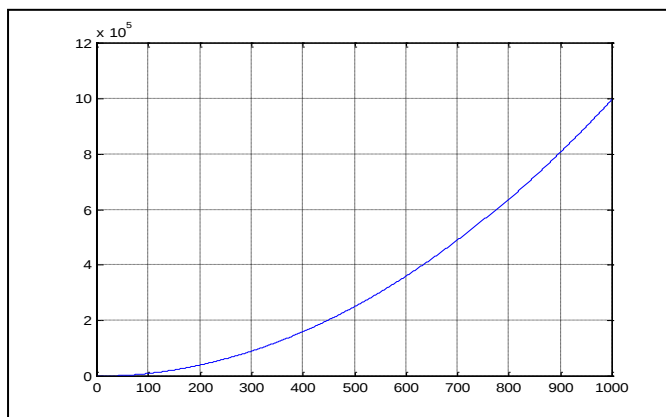
ans =
    1.0000         0   -0.0000
    0.0000    1.0000   -0.0000
   -0.0000    0.0000    1.0000
» inv(A)*A
ans =
    1.0000   -0.0000   -0.0000
         0    1.0000    0.0000
   -0.0000   -0.0000    1.0000

```

```
» t=0:pi/50:2*pi;  
» y=sin(t);  
» plot(t,y);  
» grid;  
» shg
```



```
» x=0:0.2:1000;  
» y=x.^2+3;  
» plot(x,y)  
» grid  
» shg
```



3- Travail Demandé

Soit la fonction suivante :

```
f2(x,y) = y*sin(x)+x*cos(y)  
en utilisant fplot, dessiner;  
fplot('f',[-5 5])
```

Université de M'sila
Faculté de Technologie

Département de Génie électrique
1^{ère} Année Master
Option : Electromécaniques

Module : Méthodes numériques appliquées

TP N° 02 : RESOLUTION D'EQUATIONS DIFFERENTIELLES PAR MEF

But du TP :

L'objectif de ce TP est de présenter un exemple sur la résolution d'une équation différentielle par la méthode des éléments finis pour cela, l'étudiant va programmer sous l'environnement MATLAB pour : voir les solutions recherchées.

Etape 1 : Formulation des équations

Le problème revient à trouver une fonction inconnue u définie dans le domaine $\Omega = [0, 2]$

et régie par une équation différentielle avec une valeur à la limite inférieure de Ω nulle.

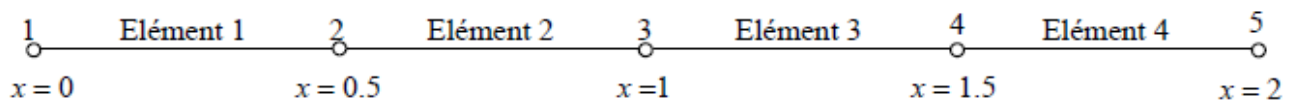
$$\frac{du}{dx} + 2x(u - 1) = 0 \quad \text{avec} \quad u(0) = 0$$

Dans ce problème, Ω est un domaine de dimension 1 ; sa frontière $\partial\Omega$ se réduit à deux points : 0 et 2.

L'équation à résoudre est une équation différentielle ordinaire dont la solution exacte est : $u_e = 1 - e^{-x^2}$

Etape 2 : Discrétisation du domaine

Le domaine Ω est divisé en n segments (appelés éléments) de taille $1/n$. Chaque élément contient deux nœuds sur lesquelles la fonction u est interpolée.



La division du domaine Ω en plusieurs éléments est appelée maillage.

On utilise deux tableaux pour la description du maillage : tableau de connectivités des éléments et tableau des coordonnées des nœuds. Pour un exemple de trois éléments on obtient les deux tableaux comme suit :

Tableau des connectivités		
Elément	Nœud 1 (début)	Nœud 2 (fin)
1	1	2
2	2	3
3	3	4
4	4	5

Tableau des coordonnées	
Nœud	Coordonnée (x)
1	0.0
2	0.5
3	1.0
4	1.5
5	2.0

Les lignes de commandes MATLAB qui permettent d'obtenir les deux tableaux sont :

```
n = 4; % nombre d'éléments
x = 0:2/n:2; % coordonnées des nœuds

for i = 1:n % début de boucle sur les éléments
    t(i,:) = [i, 1+i]; % connectivite de chaque élément
end; % fin de boucle
```

Ou bien, sous forme plus compacte :

```
x = [0:2/n:2] '
t([1:n], :) = [[1:n] ', [2:n+1] ']
```

Remarque

Pour un nombre plus élevé d'éléments il suffit d'augmenter la valeur de n .

$x(i)$ donne la coordonnée du nœud i ; exemple $x(2)$ affiche : 0.5

$t(i, :)$ donne les deux nœuds de l'élément i ; exemple $t(2, :)$ affiche : 2 3

$t(:, i)$ donne la colonne i du tableau t ; exemple $t(:, 2)$ affiche : [2 3 4 5]^T

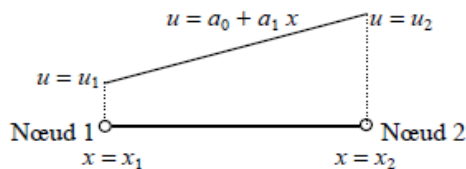
Etape 3a : Discrétisation – Interpolation sur l'élément

On peut interpoler la fonction u recherchée dans un élément par un polynôme. L'ordre du polynôme conditionne la précision de la solution approchée. Pour un élément à deux nœuds on peut prendre :

$$u = a_0 + a_1 x$$

soit sous forme vectorielle :

$$u = \begin{bmatrix} 1 & x \end{bmatrix} \begin{Bmatrix} a_0 \\ a_1 \end{Bmatrix} \equiv u = p a_n$$



avec p vecteur ligne contenant les monômes x^n

et a_n vecteur colonne contenant les facteurs du polynôme.

Cette approximation de la fonction inconnue u est appelée interpolation polynomiale, elle est fonction de a_0 et a_1 qui sont des coefficients sans valeur physique. Pour utiliser les valeurs de u aux nœuds on cherche une interpolation en fonction de u_1 et u_2 . L'interpolation polynomiale aux nœuds s'écrit :

$$\begin{Bmatrix} u_1 \\ u_2 \end{Bmatrix} = \begin{bmatrix} 1 & x_1 \\ 1 & x_2 \end{bmatrix} \begin{Bmatrix} a_0 \\ a_1 \end{Bmatrix} \equiv u_n = P_n a_n$$

L'inverse de ce système d'équations donne les paramètres a_n .

$$a_n = P_n^{-1} U_n \equiv \begin{Bmatrix} a_0 \\ a_1 \end{Bmatrix} = \frac{1}{x_2 - x_1} \begin{bmatrix} x_2 & -x_1 \\ -1 & 1 \end{bmatrix} \begin{Bmatrix} u_1 \\ u_2 \end{Bmatrix}$$

En remplaçant les a_n on peut maintenant approcher la fonction u par :

$$u = \begin{bmatrix} 1 & x \end{bmatrix} \frac{1}{x_2 - x_1} \begin{bmatrix} x_2 & -x_1 \\ -1 & 1 \end{bmatrix} \begin{Bmatrix} u_1 \\ u_2 \end{Bmatrix} = \begin{bmatrix} \frac{x_2 - x}{x_2 - x_1} & \frac{x - x_1}{x_2 - x_1} \end{bmatrix} \begin{Bmatrix} u_1 \\ u_2 \end{Bmatrix} \equiv u = N U_n$$

avec N est un vecteur ligne contenant des fonctions de x appelées fonctions de forme.

Cette interpolation est appelée interpolation nodale puisqu'elle dépend des valeurs aux nœuds de la fonction inconnue u .

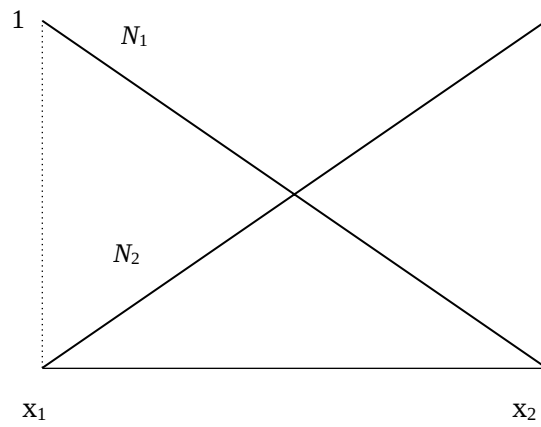
Propriétés des fonctions de forme

Il est intéressant de relever les propriétés suivantes pour les fonctions de forme N :

$$N_1(x) = \frac{x_2 - x}{x_2 - x_1} = \begin{cases} 1 & x = x_1 \\ 0 & x = x_2 \end{cases} ; \quad N_2(x) = \frac{x - x_1}{x_2 - x_1} = \begin{cases} 1 & x = x_2 \\ 0 & x = x_1 \end{cases}$$

$$N_1(x) + N_2(x) = 1 \quad \forall x \in [x_1 ; x_2]$$

- 1) Elles prennent la valeur unité aux nœuds de même indice et la valeur nulle aux autres nœuds.
- 2) Leur somme est égale à l'unité sur tout l'intervalle de l'élément :

*Remarque*

Les deux fonctions de forme peuvent s'écrire sous forme des polynômes de Jacobi :

$$N_i(x) = \prod_{j \neq i} \frac{x - x_j}{x_i - x_j}$$

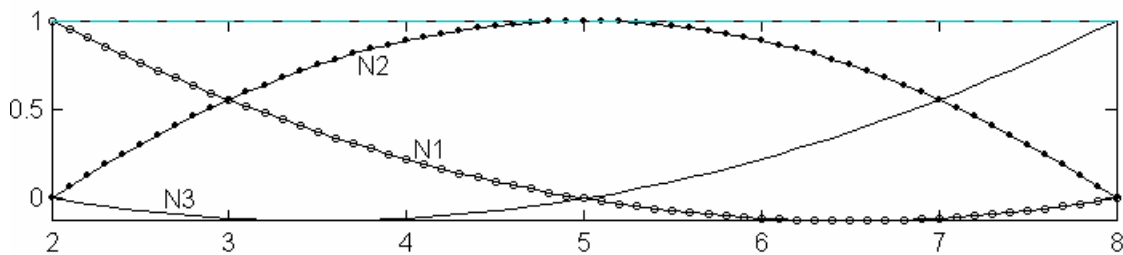
Elément de type Jacobi

Ainsi les éléments pour lesquels les fonctions de forme sont des polynômes de Jacobi sont appelés éléments de type Jacobi. On peut construire directement, sans passer par l'interpolation polynomiale, des fonctions de forme pour des éléments d'ordre supérieur (plus de deux nœuds).

Par exemple les fonctions de forme de l'élément à trois nœuds sont :

$$N_1(x) = \frac{(x-x_2)(x-x_3)}{(x_1-x_2)(x_1-x_3)} ; N_2(x) = \frac{(x-x_1)(x-x_3)}{(x_2-x_1)(x_2-x_3)} ; N_3(x) = \frac{(x-x_1)(x-x_2)}{(x_2-x_1)(x_3-x_2)}$$

On peut vérifier aisément sur la figure ci-dessous les propriétés aux nœuds des trois fonctions ainsi que leur somme sur l'élément.



$$N_1(x) = \begin{cases} 1 & x = x_1 \\ 0 & x = x_2 \\ 0 & x = x_3 \end{cases} ; N_2(x) = \begin{cases} 0 & x = x_1 \\ 1 & x = x_2 \\ 0 & x = x_3 \end{cases} ; N_3(x) = \begin{cases} 0 & x = x_1 \\ 0 & x = x_2 \\ 1 & x = x_3 \end{cases}$$

Le scripte MATLAB qui permet de calculer et de tracer ces trois fonctions est le suivant :

```
x1 = 2;
x2 = 5;
x3 = 8;
x = x1:0.1:x3;
N1 = (x-x2).*(x-x3)/(x1-x2)/(x1-x3);
N2 = (x-x1).*(x-x3)/(x2-x1)/(x2-x3);
N3 = (x-x1).*(x-x2)/(x3-x1)/(x3-x2); S =
N1+N2+N3;
plot(x,N1,x,N2,x,N3,x,S)
```

Noter qu'il est préférable d'utiliser un fichier fonction d'usage plus général.

Etape 3b : Discrétisation – Matrices élémentaires

Le calcul des matrices élémentaires passe par la réécriture du problème sous forme intégrale :

$$\int_{\Omega} \delta u \left[\frac{du}{dx} + 2x(u-1) \right] d\Omega = 0$$

avec δu est une fonction de pondération prise égale à une perturbation de la fonction inconnue u .

Le domaine Ω comprend l'intervalle 0 à 2, $d\Omega = dx$ et avec l'interpolation nodale on a :

$$\frac{du}{dx} = \frac{dN}{dx} U_n ; \text{ et } \delta u = N \delta U_n$$

Puisque seules les fonctions N dépendent de x et les perturbations ne touchent que les valeurs de u .

Pour commodité on écrit : $\delta u = (\delta N_n)^T N^T$

L'intégrale de 0 à 2 peut être remplacée par la somme des intégrales de x_i à x_{i+1} (ou bien : l'intégrale sur Ω est la somme des intégrales sur Ω_e , avec Ω_e est le domaine de chaque élément).

$$\int_{\Omega} = \sum_{\text{éléments}} \int_{\Omega_e} \equiv \int_0^2 = \sum_{i=1,n} \int_{x_i}^{x_{i+1}} = \sum_{\text{éléments}} \int_{x_1}^{x_2}$$

La forme intégrale de l'équation différentielle devient alors pour chaque élément :

$$\int_{x_1}^{x_2} \left[\delta u \frac{du}{dx} \right] dx + \int_{x_1}^{x_2} [\delta u 2x u] dx - \int_{x_1}^{x_2} [\delta u] dx = 0$$

$$\int_{x_1}^{x_2} \left[\delta U_n^T N^T \frac{dN}{dx} U_n \right] dx + \int_{x_1}^{x_2} [\delta U_n^T N^T 2x N U_n] dx - \int_{x_1}^{x_2} [\delta U_n^T N^T] 2x dx = 0$$

Cette écriture discrétisée est valable pour tous les types d'éléments. Dans le cas particulier d'un élément linéaire à deux nœuds, elle s'écrit comme suit :

$$\langle \delta u_1 \quad \delta u_2 \rangle \int_{x_1}^{x_2} \begin{Bmatrix} N_1 \\ N_2 \end{Bmatrix} \langle dN_1 \quad dN_2 \rangle dx \begin{Bmatrix} u_1 \\ u_2 \end{Bmatrix} + \langle \delta u_1 \quad \delta u_2 \rangle \int_{x_1}^{x_2} \begin{Bmatrix} N_1 \\ N_2 \end{Bmatrix} 2x \langle N_1 \quad N_2 \rangle dx \begin{Bmatrix} u_1 \\ u_2 \end{Bmatrix} -$$

$$\langle \delta u_1 \quad \delta u_2 \rangle \int_{x_1}^{x_2} \begin{Bmatrix} N_1 \\ N_2 \end{Bmatrix} 2x dx = 0$$

On voit qu'il est possible de simplifier $(\delta u_n)^T$ puisqu'il n'est pas nul et revient à chaque terme.

Au fait c'est les autres termes qui doivent être nuls ! on le voit bien avec cet exemple :

$$\langle \delta u_1 \quad \delta u_2 \rangle \begin{Bmatrix} v_1 \\ v_2 \end{Bmatrix} + \langle \delta u_1 \quad \delta u_2 \rangle \begin{Bmatrix} w_1 \\ w_2 \end{Bmatrix} = 0$$

Donne :

$$\delta u_1 v_1 + \delta u_2 v_2 + \delta u_1 w_1 + \delta u_2 w_2 = 0 ;$$

$$\text{soit : } v_1 + w_1 = 0 \text{ et } v_2 + w_2 = 0 ; \text{ ou : } \begin{Bmatrix} v_1 \\ v_2 \end{Bmatrix} + \begin{Bmatrix} w_1 \\ w_2 \end{Bmatrix} = 0 .$$

Finalement l'équation intégrale discrétisée se met sous la forme matricielle :

$$\left(\int_{x_1}^{x_2} \begin{Bmatrix} N_1 \\ N_2 \end{Bmatrix} < dN_1 \quad dN_2 > dx + \int_{x_1}^{x_2} \begin{Bmatrix} N_1 \\ N_2 \end{Bmatrix} 2x < N_1 \quad N_2 > dx \right) \begin{Bmatrix} u_1 \\ u_2 \end{Bmatrix} - \int_{x_1}^{x_2} \begin{Bmatrix} N_1 \\ N_2 \end{Bmatrix} 2x dx = 0$$

$$\left(\int_{x_1}^{x_2} \begin{bmatrix} N_1 dN_1 & N_1 dN_2 \\ N_2 dN_1 & N_2 dN_2 \end{bmatrix} dx + \int_{x_1}^{x_2} \begin{bmatrix} N_1 2x N_1 & N_1 2x N_2 \\ N_2 2x N_1 & N_2 2x N_2 \end{bmatrix} dx \right) \begin{Bmatrix} u_1 \\ u_2 \end{Bmatrix} = \int_{x_1}^{x_2} \begin{Bmatrix} 2x N_1 \\ 2x N_2 \end{Bmatrix} dx$$

Qui est un système d'équations linéaires $Ke Ue = Fe$; avec Ke et Fe sont appelés matrice et vecteur élémentaires du système d'équation. Dans le cas de la présente équation différentielle K est la somme de deux matrices : $Ke = Ke_1 + Ke_2$ tel que :

$$Ke_1 = \int_{x_1}^{x_2} \begin{bmatrix} N_1 dN_1 & N_1 dN_2 \\ N_2 dN_1 & N_2 dN_2 \end{bmatrix} dx ; \quad Ke_2 = \int_{x_1}^{x_2} \begin{bmatrix} N_1 2x N_1 & N_1 2x N_2 \\ N_2 2x N_1 & N_2 2x N_2 \end{bmatrix} dx$$

En remplaçant les fonctions de forme et leurs dérivées par leurs expressions respectives on obtient :

$$Ke_1 = \int_{x_1}^{x_2} \begin{bmatrix} \frac{x-x_2}{x_1-x_2} \frac{1}{x_1-x_2} & \frac{x-x_2}{x_1-x_2} \frac{1}{x_2-x_1} \\ \frac{x-x_1}{x_2-x_1} \frac{1}{x_1-x_2} & \frac{x-x_1}{x_2-x_1} \frac{1}{x_2-x_1} \end{bmatrix} dx ; \quad Ke_2 = \int_{x_1}^{x_2} \begin{bmatrix} \frac{x-x_2}{x_1-x_2} 2x \frac{x-x_2}{x_1-x_2} & \frac{x-x_2}{x_1-x_2} 2x \frac{x-x_1}{x_2-x_1} \\ \frac{x-x_1}{x_2-x_1} 2x \frac{x-x_2}{x_1-x_2} & \frac{x-x_1}{x_2-x_1} 2x \frac{x-x_1}{x_2-x_1} \end{bmatrix} dx$$

Soit après intégration des composantes des deux matrices :

$$Ke_1 = \frac{1}{2} \begin{bmatrix} -1 & 1 \\ -1 & 1 \end{bmatrix} \text{ et } Ke_2 = \frac{x_2 - x_1}{6} \begin{bmatrix} 3x_1 + x_2 & x_1 + x_2 \\ x_1 + x_2 & x_1 + 3x_2 \end{bmatrix}$$

Le vecteur F est donné par :

$$Fe = \int_{x_1}^{x_2} \begin{Bmatrix} 2x (x-x_2)/(x_1-x_2) \\ 2x (x-x_1)/(x_2-x_1) \end{Bmatrix} dx \text{ soit : } Fe = \frac{x_2 - x_1}{3} \begin{Bmatrix} 2x_1 + x_2 \\ x_1 + 2x_2 \end{Bmatrix}$$

Remarque

Il est possible d'utiliser MATLAB pour intégrer analytiquement les matrices élémentaires, Le script peut être le suivant :

```
syms x x1 x2 real % déclaration de variables symboliques
N = [(x-x2)/(x1-x2) (x-x1)/(x2-x1)] % fonctions de forme
dN = simple(diff(N,x)) % dérivées des fonctions de forme
Ke1 = simple(int(N' * dN , x, x1, x2) ) % matrice Ke1
Ke2 = simple(int(N' * 2*x * N , x, x1, x2) ) % matrice Ke2
Fe = simple(int(N' * 2*x , x, x1, x2) ) % vecteur Fe
```

Etape 4 : Assemblage

Le calcul des matrices élémentaires permet d'obtenir pour les quatre éléments les systèmes d'équations élémentaires suivants :

Elément 1 : $x_1 = 0$; $x_2 = \frac{1}{2}$;

$$Ke_1^{(1)} = \begin{bmatrix} -\frac{1}{2} & \frac{1}{2} \\ -\frac{1}{2} & \frac{1}{2} \end{bmatrix} ; Ke_2^{(1)} = \frac{\frac{1}{2}-0}{6} \begin{bmatrix} 3 \times 0 + \frac{1}{2} & 0 + \frac{1}{2} \\ 0 + \frac{1}{2} & 0 + 3 \times \frac{1}{2} \end{bmatrix} ;$$

$$Ke^{(1)} = Ke_1^{(1)} + Ke_2^{(1)} = \frac{1}{24} \begin{bmatrix} -11 & 13 \\ -11 & 15 \end{bmatrix} = \begin{bmatrix} -0.4583 & 0.5417 \\ -0.4583 & 0.6250 \end{bmatrix}$$

$$Fe^{(1)} = \frac{\frac{1}{2}-0}{3} \begin{Bmatrix} 2 \times 0 + \frac{1}{2} \\ 0 + 2 \times \frac{1}{2} \end{Bmatrix} = \frac{1}{12} \begin{Bmatrix} 1 \\ 2 \end{Bmatrix} = \begin{Bmatrix} 0.0833 \\ 0.1667 \end{Bmatrix}$$

$$Ke^{(1)} Ue^{(1)} = Fe^{(1)} ; \frac{1}{24} \begin{bmatrix} -11 & 13 \\ -11 & 15 \end{bmatrix} \begin{Bmatrix} u_1 \\ u_2 \end{Bmatrix} = \frac{1}{12} \begin{Bmatrix} 1 \\ 2 \end{Bmatrix}$$

On notant par U les valeurs de la fonction u aux cinq nœuds, les valeurs du vecteur élémentaire $Ue^{(1)}$ de l'élément 1 correspondent aux composantes u_1 et u_2 vecteur global U , celles de $Ue^{(2)}$ de l'élément 2 correspondent à la seconde et troisième du vecteur global U , celles de $Ue^{(3)}$ à la troisième et quatrième et celles de $Ue^{(4)}$ à la quatrième et cinquième.

Elément 2 : $x_1 = \frac{1}{2}$; $x_2 = 1$;

$$Ke^{(2)} = \frac{1}{24} \begin{bmatrix} -7 & 15 \\ -9 & 19 \end{bmatrix} = \begin{bmatrix} -0.2917 & 0.6250 \\ -0.3750 & 0.7917 \end{bmatrix} ; Fe^{(2)} = \frac{1}{12} \begin{Bmatrix} 4 \\ 5 \end{Bmatrix} = \begin{Bmatrix} 0.3333 \\ 0.4167 \end{Bmatrix}$$

$$Ke^{(2)} Ue^{(2)} = Fe^{(2)} ; \frac{1}{24} \begin{bmatrix} -7 & 15 \\ -9 & 19 \end{bmatrix} \begin{Bmatrix} u_2 \\ u_3 \end{Bmatrix} = \frac{1}{12} \begin{Bmatrix} 4 \\ 5 \end{Bmatrix}$$

Elément 3 : $x_1 = 1$; $x_2 = 3/2$;

$$Ke^{(3)} = \frac{1}{24} \begin{bmatrix} -3 & 17 \\ 1 & 19 \\ -5 & 27 \end{bmatrix} = \begin{bmatrix} -0.1250 & 0.7083 \\ 0.0417 & 0.7917 \\ -0.2083 & 1.1250 \end{bmatrix} ; Fe^{(3)} = \frac{1}{12} \begin{Bmatrix} 7 \\ 10 \\ 11 \end{Bmatrix} = \begin{Bmatrix} 0.5833 \\ 0.8333 \\ 0.9167 \end{Bmatrix}$$

$$Ke^{(4)} Ue^{(4)} = Fe^{(4)} ; \frac{1}{24} \begin{bmatrix} 1 & 19 \\ -5 & 27 \end{bmatrix} \begin{Bmatrix} u_4 \\ u_5 \end{Bmatrix} = \frac{1}{12} \begin{Bmatrix} 10 \\ 11 \end{Bmatrix}$$

Elément 4 : $x_1 = 3/2$; $x_2 = 2$;

$$Ke^{(4)} = \frac{1}{24} \begin{bmatrix} 1 & 19 \\ -5 & 27 \end{bmatrix} = \begin{bmatrix} 0.0417 & 0.7917 \\ -0.2083 & 1.1250 \end{bmatrix} ; Fe^{(4)} = \frac{1}{12} \begin{Bmatrix} 10 \\ 11 \end{Bmatrix} = \begin{Bmatrix} 0.8333 \\ 0.9167 \end{Bmatrix}$$

$$Ke^{(4)} Ue^{(4)} = Fe^{(4)} ; \frac{1}{24} \begin{bmatrix} 1 & 19 \\ -5 & 27 \end{bmatrix} \begin{Bmatrix} u_4 \\ u_5 \end{Bmatrix} = \frac{1}{12} \begin{Bmatrix} 10 \\ 11 \end{Bmatrix}$$

En réécrivant les systèmes élémentaires en fonction de toutes les composantes de U on obtient :

$$\text{Elément 1 : } \frac{1}{24} \begin{bmatrix} -11 & 13 & 0 & 0 & 0 \\ -11 & 15 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \end{bmatrix} \begin{Bmatrix} u_1 \\ u_2 \\ u_3 \\ u_4 \\ u_5 \end{Bmatrix} = \frac{1}{12} \begin{Bmatrix} 1 \\ 2 \\ 0 \\ 0 \\ 0 \end{Bmatrix}$$

$$\text{Elément 2 : } \frac{1}{24} \begin{bmatrix} 0 & 0 & 0 & 0 & 0 \\ 0 & -7 & 15 & 0 & 0 \\ 0 & -9 & 19 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \end{bmatrix} \begin{Bmatrix} u_1 \\ u_2 \\ u_3 \\ u_4 \\ u_5 \end{Bmatrix} = \frac{1}{12} \begin{Bmatrix} 0 \\ 4 \\ 5 \\ 0 \\ 0 \end{Bmatrix}$$

$$\text{Elément 3 : } \frac{1}{24} \begin{bmatrix} 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & -3 & 17 & 0 \\ 0 & 0 & -7 & 23 & 0 \\ 0 & 0 & 0 & 0 & 0 \end{bmatrix} \begin{Bmatrix} u_1 \\ u_2 \\ u_3 \\ u_4 \\ u_5 \end{Bmatrix} = \frac{1}{12} \begin{Bmatrix} 0 \\ 0 \\ 7 \\ 8 \\ 0 \end{Bmatrix}$$

$$\text{Elément 4 : } \frac{1}{24} \begin{bmatrix} 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 19 \\ 0 & 0 & 0 & -5 & 27 \end{bmatrix} \begin{Bmatrix} u_1 \\ u_2 \\ u_3 \\ u_4 \\ u_5 \end{Bmatrix} = \frac{1}{12} \begin{Bmatrix} 0 \\ 0 \\ 0 \\ 10 \\ 11 \end{Bmatrix}$$

En prenant maintenant la somme ($\equiv$ somme des intégrales), le système global s'écrit enfin :

$$\frac{1}{24} \begin{bmatrix} -11 & 13 & 0 & 0 & 0 \\ -11 & 15-7 & 15 & 0 & 0 \\ 0 & -9 & 19-3 & 17 & 0 \\ 0 & 0 & -7 & 23+1 & 19 \\ 0 & 0 & 0 & -5 & 27 \end{bmatrix} \begin{Bmatrix} u_1 \\ u_2 \\ u_3 \\ u_4 \\ u_5 \end{Bmatrix} = \frac{1}{12} \begin{Bmatrix} 1 \\ 2+4 \\ 5+7 \\ 8+10 \\ 11 \end{Bmatrix} \equiv K U = F$$

On appelle cette opération *assemblage*. En pratique, on ne procède pas de cette manière pour des raisons d'économie de mémoire et de temps de calcul mais on fait un assemblage des matrices élémentaires en utilisant les connectivités des éléments. La matrice globale K est d'abord initialisée à la matrice nulle, ensuite à chaque construction de matrice élémentaire, on localise avec une table là où il faut l'ajouter à la matrice globale. Dans le cas d'un degré de liberté par nœud (ce cas présent) la table de localisation correspond à la table des connectivités.

Exemple ; pour l'élément 1 la table de localisation est $t = [1 \ 2]$ on ajoute donc la matrice $Ke^{(1)}$ à $K(1,1)$, $K(1,2)$, $K(2,1)$ et $K(2,2)$; et pour l'élément 2 la table de localisation est $t = [2 \ 3]$ on ajoute la matrice $Ke^{(2)}$ à $K(2,2)$, $K(2,3)$, $K(3,2)$ et $K(3,3)$. Et ainsi de suite.

Le script MATLAB qui permet de faire ce type d'assemblage est simple :

```
K = zeros(n+1);           % initialisation de la matrice globale (n+1 nœuds)
F = zeros(n+1,1) ;       % initialisation du vecteur global (remarquer ,1)
for i = 1:n               % boucle sur les éléments
    t = [i i+1];          % table de localisation de chaque élément
    x1 = x(i); x2 = x(i+1); % coordonnées des deux nœuds de l'élément
    [Ke,Fe] = MatElt2Nd(x1,x2); % Fonction pour calculer la mat. et vect. élémentaires
    K(t,t) = K(t,t) + Ke;   % Assemblage de la matrice globale
    F(t) = F(t) + Fe;       % Assemblage du vecteur global
end;                      % fin de boucle
```

Etape 5a : Résolution – Application des CAL

Avant de résoudre le système il faut appliquer les conditions aux limites de la fonction u . Au nœud 1 de coordonnée $x = 0$, $u = 0$; ce qui se traduit par $u_1 = 0$. Ceci induit la réduction du nombre d'équations total à résoudre. Le système global devient alors :

$$\frac{1}{24} \begin{bmatrix} 15-7 & 15 & 0 & 0 \\ -9 & 19-3 & 17 & 0 \\ 0 & -7 & 23+1 & 19 \\ 0 & 0 & -5 & 27 \end{bmatrix} \begin{Bmatrix} u_2 \\ u_3 \\ u_4 \\ u_5 \end{Bmatrix} = \frac{1}{12} \begin{Bmatrix} 2+4 \\ 5+7 \\ 8+10 \\ 11 \end{Bmatrix}$$

La ligne et la colonne d'indice 1 qui correspondent à la valeur u_1 ont été supprimées de la matrice K et du vecteur F . Si $u(0) = a \neq 0$ ($u_1 = a$) alors on retranche de F le produit de la 1^{ère} colonne de K par a .

La commande MATLAB qui permet de supprimer une ligne ou une colonne d'une matrice consiste à lui affecter un vecteur nul (vide) :

```
A(i, :) = []           % supprime la ligne i de la matrice A
A(:, i) = []           % supprime la colonne i
V(i) = []              % supprime la composante i du vecteur V
```

Pour appliquer la condition aux limites $u_1 = 0$, il suffit d'écrire :

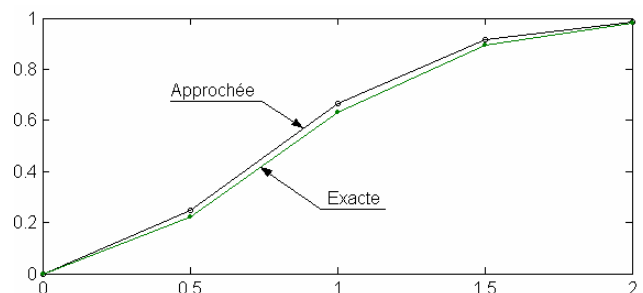
```
K(1, :) = [];          % suppression de la 1ere ligne de K
K(:, 1) = [];          % suppression de la 1ere colonne de K
F(1, :) = [];          % suppression de la 1ere composante de F
```

Remarque : pour $u(0) = a$, on insère, avant les suppressions, la commande : $F = F - K(:, 1) * a$.

Etape 5b : Résolution – Calcul de la solution

La solution peut être maintenant calculée par un des algorithmes de résolution de système linéaires. Notons que dans ce cas la matrice K est tridiagonale, mais avec MATLAB il suffit d'utiliser la division gauche $U = K \setminus F$. Les résultats sont donnés dans le tableau et la figure suivants :

Coord. x	Solution Exacte	Solution MEF	Erreur (%)
0.5000	0.2212	0.2488	12.4555
1.0000	0.6321	0.6673	5.5705
1.5000	0.8946	0.9154	2.3225
2.0000	0.9817	0.9843	0.2694



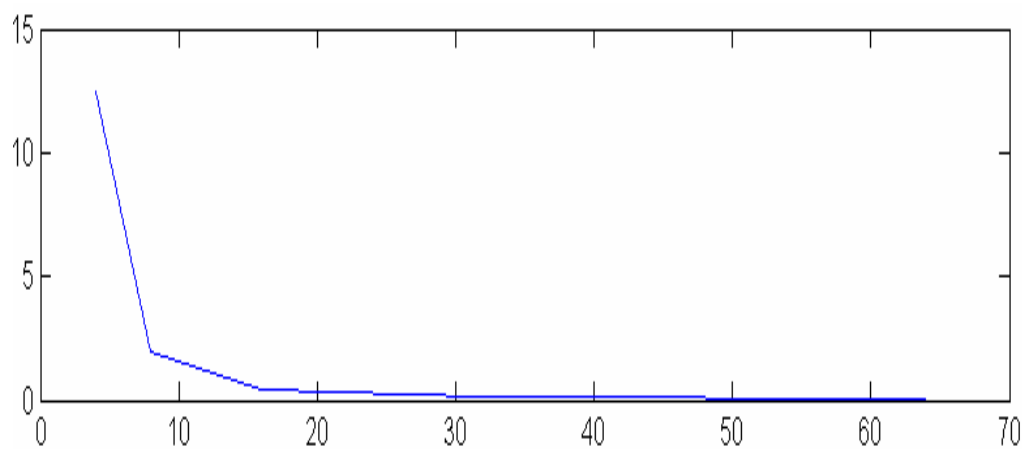
Etude de la convergence

Toute étude qui passe par une solution numérique approchée doit faire l'objet d'un examen de convergence.

Dans les études par éléments finis la convergence peut être atteinte en augmentant le degré des polynômes d'interpolation (nombre de nœuds par élément), en utilisant un nombre d'éléments assez grand avec un raffinement du maillage ou bien en cherchant une bonne adaptation du maillage au problème traité.

La figure ci après donne le pourcentage d'erreur relative au point $x = 0.5$ en fonction du nombre n d'éléments utilisés.

On voit que le décroissement de l'erreur relative est une fonction de type $1/n$ et tends à s'annuler d'une manière monotone. Un nombre d'éléments $n \approx 30$ suffit pour avoir une erreur acceptable.



Le programme MATLAB suivant regroupe toutes les étapes précédentes.

```
function [U, Ue, Err, x] = ExempleEquaDiff(n)
%
% du/dx - 2*x*(1-u) = 0
% u(0) = 0
%
%-----
% Solution avec des fonctions d'interpolation à 2Nds
%-----
x = [0:2/n:2]';           % vecteur des coordonnées
K = zeros(n+1, 1);        % initialisations
F = zeros(n+1, 1);

for i = 1:n                % boucle sur les éléments
    j = i+1;
    t = [i j];
    x1 = x(i);
    x2 = x(j);
    [Ke, Fe] = MatElt2Nd(x1,x2);
    K(t,t) = K(t,t) + Ke;
    F(t) = F(t) + Fe;
end;
K(1,:) = [];              % application des CAL
K(:,1) = [];
F(1,:) = [];

U = K\F;                  % résolution

U = [0;U];               % incorporation de la valeur initiale pour plot
Ue = 1-exp(-x.^2);       % calcul de la solution exacte
Err = 100*(U-Ue)./Ue;     % calcul de l'erreur relative
plot(x,U,x,Ue)           % trace les deux solutions
return

%-----
% Calcul de la matrice Ke et du vecteur Fe
%-----
function [Ke, Fe] = MatElt2Nd(x1,x2)           % déclaration de la fonction
    Ke1 = 1/2*[ -1    1                    % matrice Ke1
               -1    1 ];

    Ke2 = (x2-x1)/6* [ 3*x1+x2   x1+x2          % matrice Ke2
                      x1+x2   x1+3*x2];
```

```
Ke = Ke1 + Ke2; % matrice Ke
Fe = (x2-x1)/3 * [2*x1+x2 ; x1+2*x2]; % vecteur Fe
return
```

Module : Méthodes numériques appliquées

**TP N° 03 : Analyse de Structure d'un Système Electromagnétique
(Cas d'une Machine à courant continu)****Position du Problème:**

La conception des systèmes électromagnétiques, fait appel de plus en plus à l'outil informatique au moyen des modélisations et simulations, elle permet de prévoir le comportement des ces systèmes ainsi que les éventuels problèmes que peuvent rencontrer les concepteurs. Le travail suivant présente un exemple d'outil d'analyse de la structure du circuit magnétique d'une machine à courant continu. Cette analyse fait appel à un logiciel mettant en œuvre l'utilisation de la méthode des éléments finis, qui est un outil mathématique très puissant pour la résolution de tels problèmes.

But du TP :

L'objectif premier de ce TP est de montrer l'utilisation de tels outils dans la conception des machines électriques.

L'étudiant illustrera (sur ordinateur) les paramètres électromagnétiques (potentiel vecteur magnétique, induction et champ magnétiques, perméabilité, flux...) pour l'analyse d'une machine à courant continu, on utilisera l'outil « Pdetool » sous MATLAB (Partiel Differential Equation Tool)

1. Phase d'Elaboration du Modèle Géométrique

- Air
- Stator et rotor
- Inducteur

2. Phase de Définition du Modèle à Etudier

- Définition des matériaux
- Définition des conditions aux limites
- Maillage

3. Phase de Résolution du Problème

Cette phase nous permet de résoudre le problème électromagnétique posé, au moyen de « pdetool ».

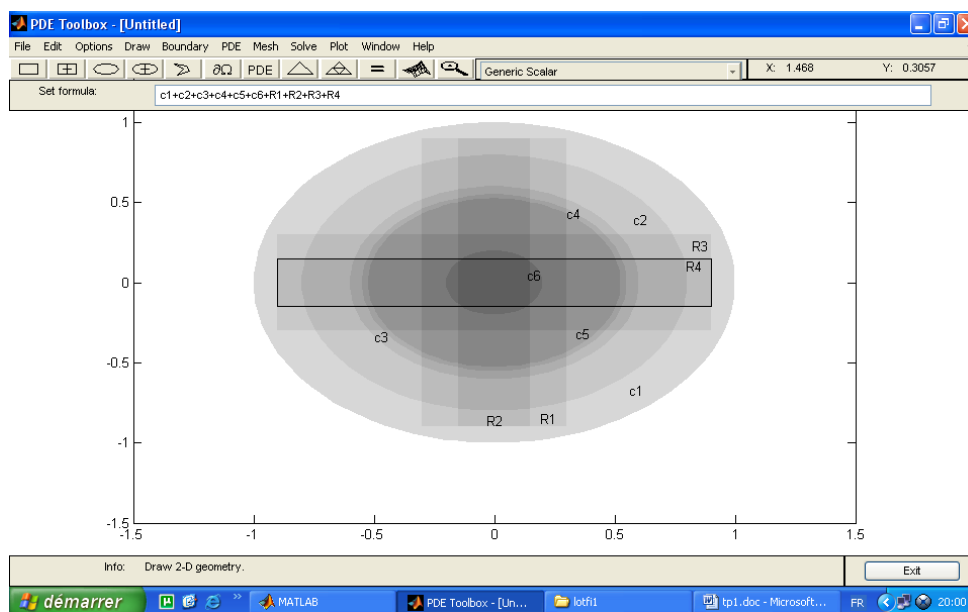
4. Phase d'Analyse du Système Développé

L'interface graphique de « pdetool » permet d'extraire et d'analyser un certain nombre de résultats.

5. Travail Demandé

Travail de préparation et d'exécution

- Ouvrir le « pdetool » sous MATLAB,
- Dessiner la géométrie

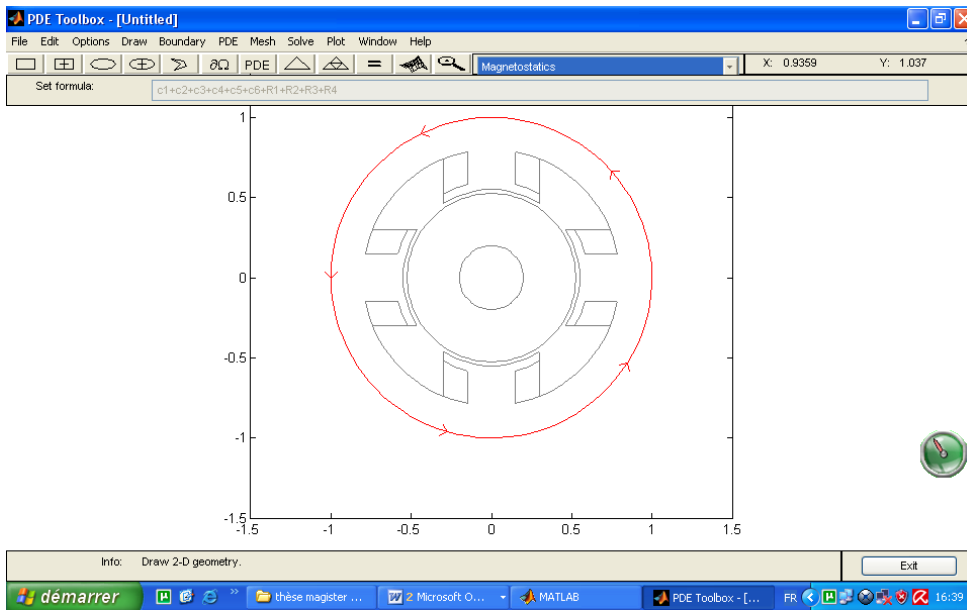


Cercle	C1 :	0	0	1	1	0	c1
Cercle	C2 :	0	0	0.8	0.8	0	c2
Cercle	C3 :	0	0	0.6	0.6	0	c3
Cercle	C4 :	0	0	0.55	0.55	0	c4
Cercle	C5 :	0	0	0.525	0.525	0	c5
Cercle	C6 :	0	0	0.2	0.2	0	c6

Rectangle	R1 :	-0.3	-0.9	0.6	1.8	r1
Rectangle	R2 :	-0.15	-0.9	0.3	1.8	r2
Rectangle	R3 :	-0.9	-0.3	1.8	0.6	r3
Rectangle	R4 :	-0.9	-0.15	1.8	0.3	r4

Afin d'obtenir la géométrie de la machine, enlever les segments de droite et portion d'arc en plus en utilisant « Boundary mode », on obtiendra la géométrie désirée comme illustré sur la figure ci-dessous.

On procédera par l'introduction des conditions aux limites du problème (la condition de Dirichlet à l'extérieur), en utilisant le bouton « specify boundary condition ».



On définira les différentes régions, en utilisant le mode « Pde » et le bouton « Pde specification »

- Stator et Rotor :

H/m $\mu = 1500 \mu_0$ - Perméabilité magnétique constante :

- Densité de courant J nulle.

- Inducteur :

H/m $\mu_0 = 4 \pi 10^{-7}$. - Perméabilité magnétique constante :

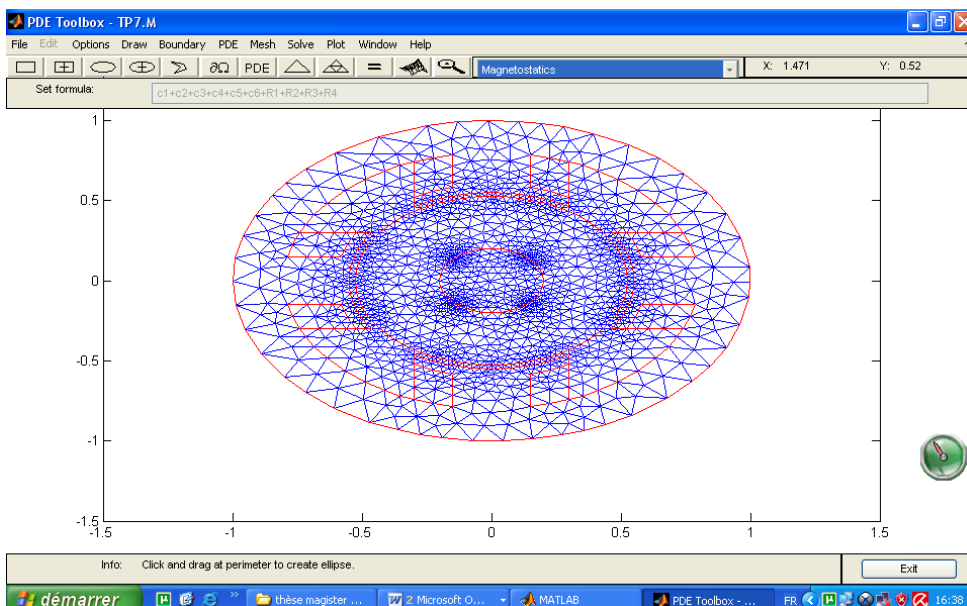
- Densité de courant $J = \pm 10^6 \text{ A/m}^2$.

- Air :

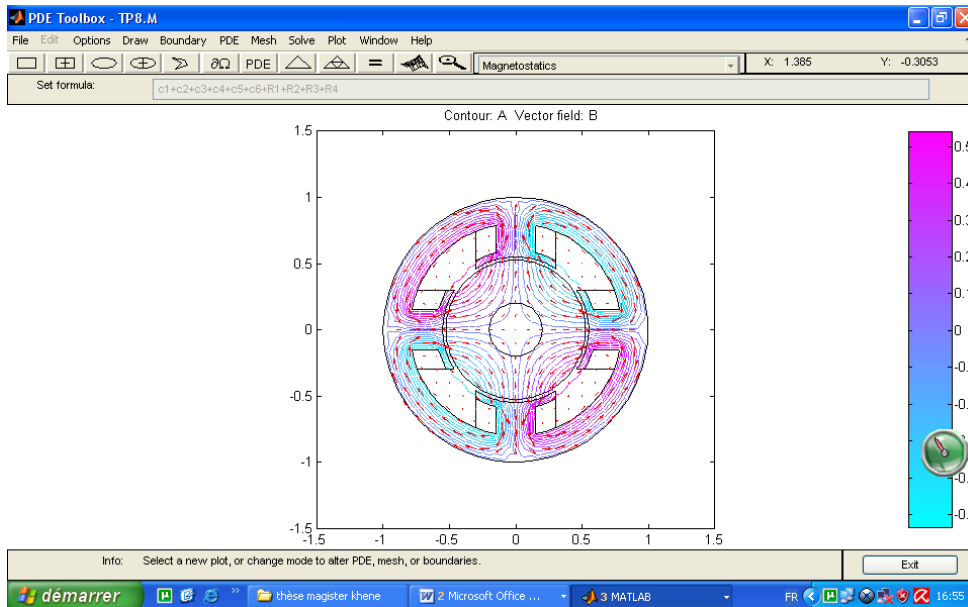
H/m $\mu_0 = 4 \pi 10^{-7}$ - Perméabilité magnétique constante :

- Densité de courant J nulle.

Ensuite on procédera au maillage de la géométrie obtenu, par l'utilisation de « Mesh mode».



On analysera ensuite le problème à l'aide « Plot » mode du menu principal.



Travail d'analyse et de résolution

Après avoir exécuté les différentes étapes du paragraphe précédent, résoudre le problème ainsi défini en utilisant le bouton « Solve » du menu principal. Cette résolution permet d'avoir un certain nombre de résultats (Potentiel vecteur magnétique, Induction magnétique....etc).

Travail à remettre

- Impression de dessin du modèle géométrique complet (avec maillage).
- Impression des résultats du modèle avec les lignes de potentiel vecteur magnétique et les flèches donnant le vecteur de l'induction magnétique.
- Analyser et interpréter les différents résultats obtenus.
- Quelle conclusion peut-on en tirer ?

Université de M'sila
Faculté de Technologie
Option : Electromécaniques

Département de Génie électrique
1^{ère} Année Master

Module : Méthodes numériques appliquées

TP N° 04 : Analyse de Structure d'un Système Electromagnétique (Transformateur)

Première Partie : Géométrie, maillage et affectation des propriétés physique du dispositif étudié (transformateur)

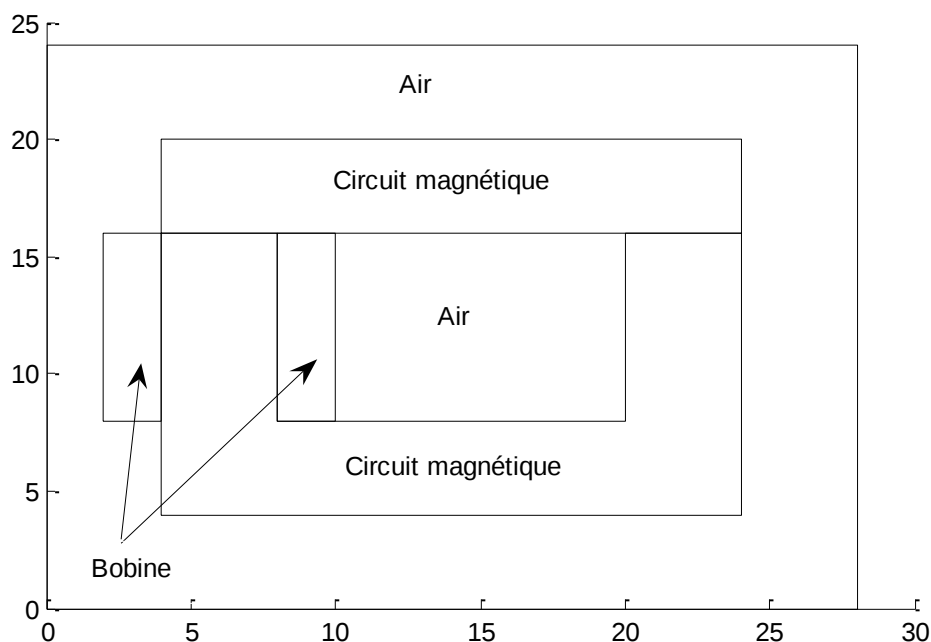
1-But du TP :

L'objectif de ce TP est de montrer l'utilisation de l'environnement MATLAB dans l'analyse de structure d'un système électromagnétique, dans ce cas on va étudier un transformateur.

L'étudiant va programmer sous l'environnement MATLAB pour : dessiner la géométrie, le maillage et l'affectation des propriétés physiques de chaque région du dispositif étudié.

2-Travail Demandé

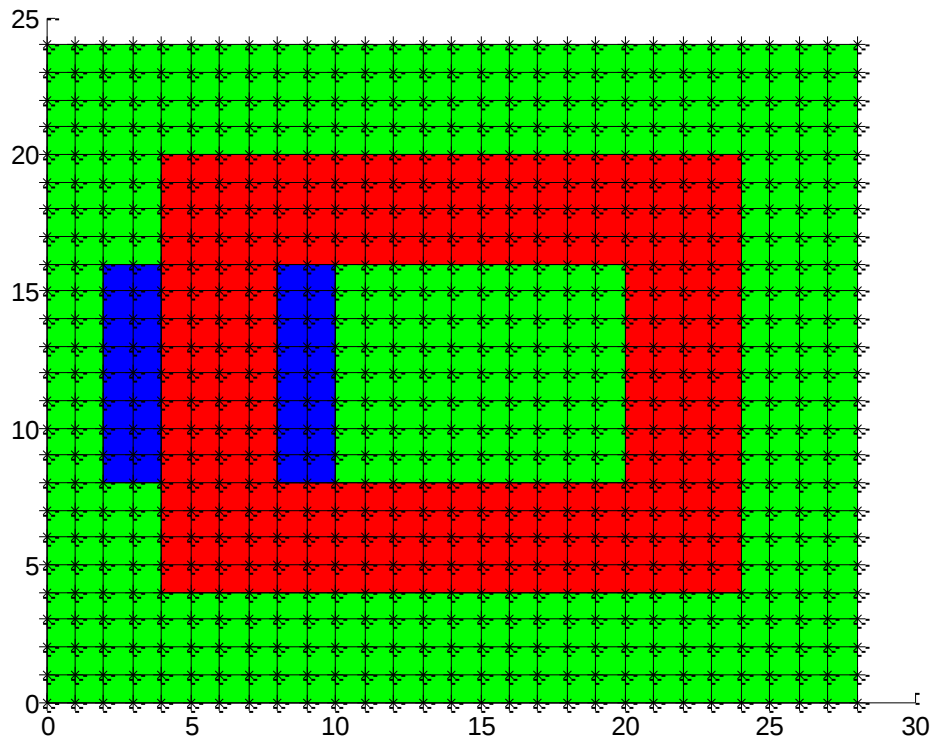
La figure ci-dessous représente la géométrie du dispositif étudié avec les différentes régions (circuit magnétique, bobine et l'air).



2-1 Travail de préparation et d'exécution

- Ouvrir l'environnement MATLAB,
- Utiliser la commande « patch » pour dessiner la géométrie du dispositif étudié et spécifier chaque région avec une couleur différentes.

La figure suivante représente le maillage du domaine d'étude



- Utiliser la commande « meshgrid » et « plot » pour mailler le domaine d'étude.
- Utiliser les boucles « for » et « end » pour affecter les propriétés physiques de chaque région
 - Circuit magnétique:
 - Perméabilité magnétique constante : $\mu = 1500 \mu_0$ H/m
 - Densité de courant J nulle.
 - Bobine :
 - Perméabilité magnétique constante : $\mu_0 = 4 \pi 10^{-7}$ H/m
 - Densité de courant $J = \pm 10^6$ A/m².
 - Air :
 - Perméabilité magnétique constante : $\mu_0 = 4 \pi 10^{-7}$ H/m
 - Densité de courant J nulle.

Références

- Méthodes numériques, Introduction à l'analyse numérique et au calcul scientifique, Guillaume Legendre. (version provisoire du 4 janvier 2017).
- Introduction aux méthodes numériques Deuxième édition Franck Jedrzejewski CEA Saclay - INSTN / UERTI 91191 Gif-sur-Yvette Cedex 2001.
- Méthode des éléments finis Hervé Oudin. 28/09/2008
- Méthodes Numériques Appliquées pour le scientifique et l'ingénieur, Jean-Philippe GRIVET.
- Méthodes Numériques Appliquées : Cours, TD, TP, NOURI HAMOU.
- Méthodes Numériques Appliquées à la conception par élément finis, DAVID Dureisseix, 2008.
- Analyse Numérique: SMA-SMI S4, Cours, exercices et examens, Boutayeb A, Derouich M, Lamlili M et Boutayeb W.
- INTRODUCTION AUX ELEMENTS FINIS V. Legat, Notes du cours MECA2120 Année académique 2005-2006.
- Cours: Equations aux Dérivées Partielles, Méthodes des Différences Finies, A. Taik, 2008.
- Méthodes et outils d'optimisation Optimisation Mathias Kleiner mathias.kleiner@ensam.eu <http://www.lsis.org/kleiner/m> Juin 2015.
- Introduction à l'Optimisation Numérique Frédéric de Gournay & Aude Rondepierre DÉPARTEMENT STPI 3ÈME ANNÉE MIC
- Optimisation : méthodes numériques pour l'optimisation différentiable Jérôme MALICK, 2009.
- Méthodes numériques pour l'optimisation non linéaire déterministe, Aude Rondepierre. Département Génie Mathématique et Modélisation 4ème année, 2017-2018. INSA, Toulouse.
- Liens sur internet.