

# Remerciements

---

Le bon déroulement de cette thèse, jusqu'à son dénouement, est en grande partie imputable à Monsieur Abdelkrim Chakib. Je le remercie du fond du cœur, aussi bien pour avoir dirigé mes travaux avec talent, que pour m'avoir accompagné amicalement et fraternellement dans ce cheminement. Son énergie, ses compétences et sa constante disponibilité sont autant de qualités sans faille que j'apprécie chez lui et qui m'ont beaucoup aidé pour mener à bien ce travail. Ce fut un réel plaisir de l'avoir comme encadrant de thèse, et c'est une chance de l'avoir comme ami. J'espère encore apprendre à son contact.

Je voudrais exprimer ici ma vive reconnaissance à Monsieur Abdeljalil Nachaoui, mon oncle, pour l'aide permanente, le soutien moral et les conseils précieux, qu'il n'a cessé de m'apporter depuis mon enfance. Son hospitalité, sa disponibilité, ses compétences et ses suggestions ont été un grand atout pour mener à bien ce travail. Je voudrais le remercier du fond du cœur pour tout ce qu'il a fait pour moi. Dieu Seul peut le récompenser.

Je tiens aussi à exprimer ma gratitude à Messieurs François Jauberteau et Ahmed Zeghal pour la confiance qu'ils m'ont témoignée en acceptant de diriger cette thèse dans le cadre d'une cotutelle entre l'Université Sultan Moulay Slimane de Beni-Mellal et l'Université de Nantes. Je les remercie vivement pour leur disponibilité, leurs conseils, et leur patience, en particulier face aux tâches administratives souvent lourdes.

Ce fut un grand honneur que Monsieur Olivier Goubet, Monsieur Jaroslav Haslinger et Madame Zoubida Mghazli aient accepté de rapporter sur cette thèse. Je les remercie du temps qu'ils ont consacré à la lecture du manuscrit, et pour l'intérêt qu'ils ont accordé à mon travail.

Je remercie respectueusement M. Michel Pierre, pour le temps qu'il a consacré à examiner cette thèse et pour l'honneur qu'il me fait en présidant ce jury.

M. Christophe Berthon a bien voulu s'intéresser à ce travail et a accepté de faire partie du jury. Qu'il trouve ici mes vifs remerciements.

Faire partie du Laboratoire de Mathématiques Jean Leray, même en tant que thésard, est sans doute un avantage, non seulement pour les conditions offertes, mais aussi grâce à la convivialité

qui y règne. Que tous ses membres, qu'ils soient professeurs, thésards ou secrétaires, trouvent ici mes sincères remerciements.

Au Maroc j'ai eu la chance de faire partie du laboratoire de mathématiques et applications ainsi que l'UFR de méthodes mathématiques et applications de l'université Sultan Moulay Sliman de Beni-Mellal. L'ambiance conviviale qui y règne est le fruit du dévouement de ses membres. Je les remercie tous. Une pensée particulière va à Abdelkader Stouti et Ahmad Zeghal, Directeurs du laboratoire, ainsi qu'à Abdelhakim Maaden, Directeur de l'UFR. Je remercie également Hamid Maaroufi pour son accueil chaleureux et son amitié.

Je tiens à remercier la famille Nachaoui de Nantes, et tout particulièrement Madame Fouzia Nachaoui, pour son hospitalité, sa gentillesse et sa patience.

Le financement de ma thèse est assuré par une bourse de recherche (CNRST du Maroc) et en partie par des coopérations franco-marocaines (AI-Volubilis MA/09/202 et Convention CNRS/CNRST 21585). Je remercie les responsables scientifiques et administratifs.

Je remercie également ma famille, en particulier ma grand-mère et mes parents, ainsi que mes amis et tous ceux qui m'ont supporté, de près ou de loin.

Je remercie ma femme pour son soutien, son aide précieuse et sa patience dans la relecture de ma thèse, et pour avoir supporté mon absence. Que Dieu la récompense.

# Table des matières

|                                                                           |           |
|---------------------------------------------------------------------------|-----------|
| <b>Remerciements</b>                                                      | <b>i</b>  |
| <b>Introduction générale</b>                                              | <b>i</b>  |
| <b>1 Résultats préliminaires</b>                                          | <b>1</b>  |
| 1.1 Notations . . . . .                                                   | 1         |
| 1.2 Espaces fonctionnels . . . . .                                        | 2         |
| 1.3 Résultats fondamentaux . . . . .                                      | 4         |
| 1.4 Propriété du prolongement uniforme . . . . .                          | 5         |
| 1.4.1 Propriété du cône . . . . .                                         | 6         |
| 1.4.2 Prolongement uniforme . . . . .                                     | 6         |
| 1.5 Quelques estimations sur les espaces de Sobolev . . . . .             | 7         |
| 1.6 Degré topologique . . . . .                                           | 8         |
| 1.6.1 Degré topologique de Brouwer . . . . .                              | 8         |
| 1.6.2 Degré topologique de Leray-Schauder . . . . .                       | 9         |
| <b>2 Modélisation et formulation du problème</b>                          | <b>11</b> |
| 2.1 Généralités sur le problème de soudage . . . . .                      | 11        |
| 2.2 Modélisation du problème . . . . .                                    | 13        |
| 2.2.1 Conditions aux limites . . . . .                                    | 14        |
| 2.2.2 Système d'équations dans la partie solide . . . . .                 | 16        |
| 2.2.3 Formulation du problème en régime quasi-stationnaire . . . . .      | 17        |
| 2.3 Formulation en optimisation de forme . . . . .                        | 17        |
| <b>3 Étude de l'existence d'une solution optimale</b>                     | <b>23</b> |
| 3.1 Introduction à l'étude d'existence en optimisation de forme . . . . . | 24        |
| 3.2 Étude du problème d'état . . . . .                                    | 25        |
| 3.2.1 Formulation variationnelle . . . . .                                | 25        |
| 3.2.2 Existence et unicité de la solution d'équation d'état . . . . .     | 26        |
| 3.3 Étude de l'existence d'une solution optimale . . . . .                | 30        |

---

|          |                                                                            |           |
|----------|----------------------------------------------------------------------------|-----------|
| 3.3.1    | Topologie sur $\mathcal{F}$ . . . . .                                      | 30        |
| 3.3.2    | Continuité du problème d'état . . . . .                                    | 35        |
| 3.3.3    | Continuité de la fonctionnelle coût . . . . .                              | 48        |
| <b>4</b> | <b>Étude de l'approximation du problème</b>                                | <b>51</b> |
| 4.1      | Approximation du problème par la méthode des éléments finis . . . . .      | 51        |
| 4.1.1    | Discrétisation de la famille des domaines admissibles . . . . .            | 53        |
| 4.1.2    | Approximation du problème d'état . . . . .                                 | 55        |
| 4.2      | Étude de l'existence d'une solution du problème discret . . . . .          | 57        |
| 4.3      | Étude de la convergence . . . . .                                          | 62        |
| 4.3.1    | Résultat abstrait . . . . .                                                | 62        |
| 4.3.2    | Résultat de convergence . . . . .                                          | 64        |
| <b>5</b> | <b>Étude numérique du problème</b>                                         | <b>73</b> |
| 5.1      | Algorithmes déterministes . . . . .                                        | 75        |
| 5.1.1    | Problème matriciel d'optimisation de forme . . . . .                       | 76        |
| 5.1.2    | Calcul du gradient discret . . . . .                                       | 77        |
| 5.1.3    | Validation de la méthode vis-à-vis d'une solution exacte . . . . .         | 78        |
| 5.1.4    | Validation de la méthode vis-à-vis du modèle physique . . . . .            | 82        |
| 5.2      | Algorithmes Génétiques . . . . .                                           | 88        |
| 5.2.1    | Introduction . . . . .                                                     | 88        |
| 5.2.2    | Les grandes familles d'Algorithmes Génétiques . . . . .                    | 89        |
| 5.2.3    | Algorithmes Génétiques Simples (GAs) . . . . .                             | 90        |
| 5.2.4    | Algorithme Génétique utilisé (GA) . . . . .                                | 91        |
| 5.2.5    | Résultats obtenus sur des fonctions tests . . . . .                        | 96        |
| 5.3      | Résultats obtenus pour le problème d'optimisation de forme . . . . .       | 100       |
| 5.3.1    | Validation de l'algorithme (GA) vis-à-vis d'une solution exacte . . . . .  | 101       |
| 5.3.2    | Validation de l'algorithme (GA) vis-à-vis du modèle physique . . . . .     | 104       |
| 5.4      | Développement d'un algorithme évolutionnaire . . . . .                     | 105       |
| 5.4.1    | Logique floue . . . . .                                                    | 105       |
| 5.4.2    | Combinaison des (GA) avec les contrôleurs de logique floue (CLF) . . . . . | 120       |

---

|       |                                                                                                                    |            |
|-------|--------------------------------------------------------------------------------------------------------------------|------------|
| 5.4.3 | Comparaison de l'algorithme (GA) et l'algorithme (GA-CLF) vis-à-vis du modèle physique . . . . .                   | 125        |
| 5.5   | Parallélisation de l'algorithme (GA) . . . . .                                                                     | 127        |
| 5.5.1 | Analyse du temps d'exécution des (GA) Maître-esclaves . . . . .                                                    | 127        |
| 5.5.2 | Analyse expérimentale du temps d'exécution d'une implémentation parallèle de l'algorithme (GA-CLF) . . . . .       | 128        |
| 5.6   | Comparaison de l'algorithme de type gradient avec les algorithmes évolutionnaires sur le modèle physique . . . . . | 131        |
| 5.6.1 | Comparaison qualitative entre l'algorithme (GA) et l'algorithme du gradient                                        | 132        |
| 5.6.2 | Comparaison quantitative entre l'algorithme (GA) et l'algorithme du gradient . . . . .                             | 132        |
|       | <b>Conclusion</b>                                                                                                  | <b>135</b> |
|       | <b>Bibliographie</b>                                                                                               | <b>137</b> |



# Table des figures

|      |                                                                                                         |     |
|------|---------------------------------------------------------------------------------------------------------|-----|
| 2.1  | Schéma du procédé de soudage . . . . .                                                                  | 14  |
| 2.2  | Description du domaine $\Omega_{tot} = \Omega_s(t) \cup \Omega_l(t)$ à l'instant $t$ . . . . .          | 15  |
| 2.3  | configuration 2-D du procédé de soudage . . . . .                                                       | 18  |
| 5.1  | Exemple 1 : décroissance de la fonction coût. . . . .                                                   | 80  |
| 5.2  | Exemple 1 : convergence des frontières. . . . .                                                         | 80  |
| 5.3  | Exemple 2 : décroissance de la fonction coût. . . . .                                                   | 81  |
| 5.4  | Exemple 2 : convergence des frontières. . . . .                                                         | 82  |
| 5.5  | Température mesurée sur $\Gamma_0$ . . . . .                                                            | 83  |
| 5.6  | Maillage du domaine initial . . . . .                                                                   | 83  |
| 5.7  | Maillage du domaine optimal . . . . .                                                                   | 84  |
| 5.8  | Conteur des isovaleurs . . . . .                                                                        | 85  |
| 5.9  | Remplissage des isovaleurs . . . . .                                                                    | 86  |
| 5.10 | Variation de la température en fonction de $x$ et $y$ . . . . .                                         | 86  |
| 5.11 | Exemple 3 : décroissance de la fonction coût. . . . .                                                   | 87  |
| 5.12 | Exemple 3 : convergence des frontières. . . . .                                                         | 87  |
| 5.13 | Schéma général d'un algorithme évolutionnaire. . . . .                                                  | 89  |
| 5.14 | Principe de croisement simple . . . . .                                                                 | 91  |
| 5.15 | Principe de Mutation simple . . . . .                                                                   | 92  |
| 5.16 | Principe de croisement multi-points . . . . .                                                           | 93  |
| 5.17 | La fonction $\Delta(t, y)$ à deux instants $t_1$ et $t_2$ (resp. (a) et (b)) avec $t_1 < t_2$ . . . . . | 94  |
| 5.18 | Evolution de l'amplitude de la mutation en fonction du temps. . . . .                                   | 95  |
| 5.19 | Minimisation de la fonction sphère par (GAs) et (GA) . . . . .                                          | 96  |
| 5.20 | Minimisation de la fonction elliptique par (GAs) et (GA) . . . . .                                      | 97  |
| 5.21 | Minimisation de la fonction Rosenbrock par (GAs) et (GA) . . . . .                                      | 98  |
| 5.22 | Fonction de Shekel . . . . .                                                                            | 99  |
| 5.23 | Minimisation de la fonction Shekel par GAs et GA . . . . .                                              | 100 |
| 5.24 | Décroissance du coût en fonction des itérations. . . . .                                                | 102 |
| 5.25 | Évolution de la frontière libre au cours des itérations. . . . .                                        | 102 |

|      |                                                                                                                                        |     |
|------|----------------------------------------------------------------------------------------------------------------------------------------|-----|
| 5.26 | Décroissance du coût en fonction des itérations. . . . .                                                                               | 103 |
| 5.27 | Évolution de la frontière libre au cours des itérations. . . . .                                                                       | 104 |
| 5.28 | Décroissance du coût en fonction des itérations. . . . .                                                                               | 104 |
| 5.29 | Évolution de la frontière libre au cours des itérations. . . . .                                                                       | 105 |
| 5.30 | Fonctions caractéristiques d'un sous-ensemble classique (a) et d'un sous-ensemble<br>flou (b) pour l'exemple 5.1 . . . . .             | 107 |
| 5.31 | Fonctions d'appartenance de la variable température. . . . .                                                                           | 109 |
| 5.32 | Fonctions d'appartenance de la variable humidité. . . . .                                                                              | 110 |
| 5.33 | Fonctions d'appartenance de la variable vitesse du ventilateur. . . . .                                                                | 110 |
| 5.34 | Structure d'un SIF . . . . .                                                                                                           | 111 |
| 5.35 | Sortie floue obtenue . . . . .                                                                                                         | 116 |
| 5.36 | Surface de contrôle du contrôleur flou avec des règles d'inférence comme dans (5.4).<br>La défuzzification MOM est appliquée. . . . .  | 117 |
| 5.37 | Surface de contrôle du contrôleur flou avec des règles d'inférence comme dans (5.4).<br>La défuzzification COG est appliquée . . . . . | 118 |
| 5.38 | Surface de contrôle du contrôleur flou avec des règles d'inférence comme dans (5.4).<br>La défuzzification WABL est utilisée. . . . .  | 119 |
| 5.39 | Comparaison de différentes tailles du cycle (GA) . . . . .                                                                             | 122 |
| 5.40 | Variable d'entrée $ED$ . . . . .                                                                                                       | 123 |
| 5.41 | Variable d'entrée $CV$ . . . . .                                                                                                       | 123 |
| 5.42 | Variable de sortie $Jv$ . . . . .                                                                                                      | 124 |
| 5.43 | Comparaison entre les résultats obtenus par l'algorithme 5.2 et ceux obtenus par<br>l'algorithme 5.4 . . . . .                         | 126 |
| 5.44 | Évolution de temps d'exécution, vitesse de calcul $S_P$ et efficacité $E_P$ en fonction<br>du nombre de processeurs . . . . .          | 131 |
| 5.45 | Comparaison entre la frontière optimale obtenue par l'algorithme 5.1 et celle ob-<br>tenue par l'algorithme 5.2 . . . . .              | 132 |

# Liste des tableaux

|      |                                                                                                                                                                    |     |
|------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----|
| 5.1  | Paramètres des (GA) pour les résultats des figures 5.24 et 5.25 . . . . .                                                                                          | 101 |
| 5.2  | Résultats de convergence de l'algorithme 5.2 pour l'exemple 1 . . . . .                                                                                            | 103 |
| 5.3  | Variables linguistiques . . . . .                                                                                                                                  | 114 |
| 5.4  | Matrice des règles d'inférence . . . . .                                                                                                                           | 115 |
| 5.5  | Valeurs de vérité pour $T = 17\text{ }^{\circ}\text{C}$ et $H = 32\%$ . . . . .                                                                                    | 115 |
| 5.6  | Degré d'appartenance de la sortie floue de la vitesse du ventilateur qui résulte des règles 5.4 . . . . .                                                          | 116 |
| 5.7  | Variables linguistiques . . . . .                                                                                                                                  | 122 |
| 5.8  | Tableau de règles d'inférence . . . . .                                                                                                                            | 124 |
| 5.9  | Données de performances . . . . .                                                                                                                                  | 130 |
| 5.10 | Comparaison du temps d'exécution pour l'algorithme 5.5 basé sur la combinaison (GA-CLF) et calcul parallèle et l'algorithme 5.1 basé sur le calcul du gradient . . | 133 |



# Introduction générale

---

Dans le but d'améliorer la productivité, de réduire le coût et de maximiser le profit, les industriels de l'aéronautique, du nucléaire et de l'automobile s'orientent de plus en plus vers des outils de simulation lorsqu'il s'agit, en particulier, des opérations de soudage. En général, les procédés de soudage amènent localement les pièces à souder à une température supérieure à la température de fusion, et l'assemblage par liaisons métalliques s'effectue lors de la solidification. Pendant le procédé de soudage, les bords des deux pièces sont fondus en créant le joint soudé. Ceci est obtenu en utilisant localement une source intense d'énergie formant ainsi un bain de fusion, qui à son tour, crée le cordon de soudure au cours de la solidification. L'un des enjeux industriels majeurs de l'assemblage par soudage concerne la prédiction des effets mécaniques (distorsions, déformations résiduelles, contraintes résiduelles, etc.) directement dépendants des évolutions de température imposées par le procédé, appelés chargement thermique du procédé.

Afin d'accéder à ce chargement thermique, deux approches sont couramment utilisées par les thermiciens : l'approche multi-physique [83] complète et l'approche dite de " source équivalente " [48, 63].

La première approche consiste à résoudre un problème magnétohydrodynamique dans toutes les parties de la structure soudée : la zone fondue (le liquide) et la zone solide. On résout alors les équations de Navier Stokes [7] et les équations de conservation d'énergie en tenant compte des interactions entre la source d'énergie (arc, faisceau, etc.) et la pièce, ainsi qu'entre le liquide et le solide. Cette approche est complète et permet d'accéder au chargement thermique du procédé. Cependant, elle reste très lourde, demande de fortes hypothèses et un temps de calcul souvent prohibitif.

La deuxième approche consiste à résoudre un problème à frontière libre de conduction de la chaleur qui ne s'occupe que de la partie solide. Elle consiste à simplifier les phénomènes physiques [56] apparaissant entre la torche de soudage et la plaque ainsi que ceux du bain liquide en introduisant une condition de température imposée sur le front de fusion, frontière liquide/solide à déterminer.

Dans ce travail, nous considérons une étude qui consiste à identifier l'interface liquide/solide

et le flux thermique qui la traverse à partir de températures mesurées dans le solide. La démarche consiste alors à résoudre un problème à frontière libre. Dans ce même contexte, nous signalons que les problèmes d'identification des frontières ou les problèmes à frontière libre constituent un type important de problèmes inverses qui sont en liaison avec une variété de phénomènes dans différents secteurs scientifiques ([2, 16, 18, 19, 28],[44],[81]...). Jusqu'à présent, il n'y a pas d'études théorique et numérique qui permettent de couvrir toute cette classe de problèmes à frontière libre, bien qu'il y ait eu des progrès sur l'étude de chaque problème appréhendé à part. Si quelques travaux de simulation numérique ont été réalisés par des équipes de thermique pour la résolution de ce problème à frontière libre de soudage [11, 33], cependant aucune étude d'existence n'a été effectuée. D'où l'idée de s'intéresser à l'étude théorique et à l'approximation numérique de ce problème reformulé en un problème d'optimisation de forme [44].

La méthode d'optimisation de forme est désormais couramment abordée pour résoudre les problèmes à frontière libre [4, 9, 42, 43]. C'est une approche variationnelle, qui consiste à minimiser une fonctionnelle coût intégrale par rapport à une famille de domaines, dont la fonctionnelle dépend de la solution d'un problème aux limites sur un domaine de cette famille appelé "problème d'état". L'étude de l'existence d'une solution d'un problème d'optimisation de forme nécessite le choix d'une topologie adéquate sur la famille des domaines qui doit assurer sa compacité, ainsi que la semi-continuité inférieure de la fonctionnelle coût.

La difficulté principale de l'étude d'existence d'une solution optimale réside dans le fait que le problème d'état associé à la formulation en optimisation de forme considérée est régi par un opérateur non coercif, ce qui rend l'étude compliquée. Pour pallier cette difficulté, nous utilisons le degré topologique de Leray Schauder [32] et des estimations uniformes basées sur des résultats récents sur l'inégalité uniforme de Poincaré [12] et quelques inégalités de Sobolev [51].

Nous étudions ensuite l'approximation numérique de ce problème d'optimisation de forme en considérant une discrétisation basée sur les éléments finis linéaires [28]. Nous montrons alors que le problème d'optimisation discret admet une solution en utilisant le degré topologique de Brouwer [32], puis nous étudions la convergence d'une suite de solutions du problème approché vers la solution du problème continu. Le problème discret correspondant est résolu en utilisant deux méthodes.

L'une est basée sur les algorithmes évolutionnaires [58] qui sont des outils d'optimisation très robustes, efficaces lorsque les fonctions à optimiser sont fortement irrégulières, et dépendantes

de paramètres variés en nature et en type. Ils sont actuellement largement employés dans des domaines d'applications extrêmement variés. Cependant, ils ne sont viables, pour des problèmes d'optimisation numérique lourds, que si on les intègre à des méthodes de résolution performantes, en particulier hiérarchiques, et au calcul parallèle [39].

L'autre méthode repose sur les algorithmes déterministes à savoir les algorithmes classiques de gradient [3] modernisés qui sont performants en vitesse de convergence asymptotique, même s'ils sont délicats à mettre en œuvre puisqu'ils exigent le calcul du gradient de forme qui est souvent délicat surtout dans le cas continu [87].

Dans la suite nous allons décrire brièvement le contenu de cette thèse.

Dans le premier chapitre, nous rappelons quelques outils de base et nous présentons quelques résultats préliminaires essentiels pour ce travail. Nous donnons en particulier quelques résultats fondamentaux sur la propriété du prolongement uniforme, l'inégalité uniforme de Poincaré [12] et le degré topologique [32].

Dans le deuxième chapitre, nous présentons une description physique du problème qui modélise l'analyse des transferts de chaleur dans une opération de soudage. Il s'agit d'estimer le champ de température dans les pièces soudées afin de prévoir et maîtriser les effets mécaniques engendrés par le procédé sur ces pièces (contraintes résiduelles, distorsions). L'approche que nous considérons ne s'occupe que de la partie solide de la plaque. Elle consiste à simplifier les phénomènes physiques apparaissant entre la torche de soudage et la plaque ainsi que ceux du bain liquide en introduisant une condition de température imposée sur le front de fusion. Nous proposons ensuite une formulation en optimisation de forme de ce problème.

Dans le troisième chapitre, nous donnons un résultat d'existence de la solution du problème d'optimisation de forme. Comme souvent les hypothèses physiques prises sur les données du problème ne permettent pas d'avoir la coercivité du problème d'état qui est nécessaire pour avoir l'existence et l'unicité de la solution en utilisant le lemme de Lax-Milgram (l'outil classique pour l'étude des problèmes elliptiques linéaires [13]), nous utilisons alors les degrés topologiques de Leray Schauder [32], pour montrer l'existence et l'unicité d'une solution du problème d'état écrit sous sa forme variationnelle. Ensuite, nous utilisons des techniques d'estimations uniformes basées sur l'inégalité de Poincaré uniforme [12] et quelques inégalités de Sobolev [51], pour montrer la continuité du problème d'état qui présente l'une des difficultés principales de l'étude de l'existence d'une solution optimale.

Dans le quatrième chapitre, nous approchons la famille des domaines admissibles par des fonctions quadratiques de Bézier [34]. Nous proposons, ensuite, une approximation du problème d'optimisation de forme par la méthode des éléments finis triangulaires. En utilisant les degrés topologiques de Brouwer, nous montrons l'existence d'une solution du problème d'optimisation de forme discret. Pour conclure nous nous basons sur un résultat de convergence abstrait dû à [42], pour montrer la convergence d'une suite de solutions des problèmes d'optimisation de forme discrets vers celle du problème continu, lorsque le pas de discrétisation tend vers zéro.

Dans le dernier chapitre, nous présentons les différentes méthodes d'optimisation utilisées pour la résolution numérique du problème d'optimisation de forme. Nous faisons alors le point sur deux familles d'algorithmes à savoir, les algorithmes évolutionnaires et les algorithmes déterministes du type gradient. Nous commençons par une analyse du calcul du gradient de forme dans la cas discret. Puis nous donnons des résultats numériques obtenus en utilisant la méthode déterministe du type gradient. En ce qui concerne les algorithmes évolutionnaires, nous optons d'abord pour une étude comparative basée sur des fonctions tests "benchmark functions" [31], entre les algorithmes génétiques standards et ceux développés dans le cadre de cette thèse. Puis nous donnons des résultats numériques obtenus en utilisant ces derniers algorithmes pour la résolution de notre problème d'optimisation de forme. Ensuite, nous proposons un développement de ces algorithmes génétiques en utilisant deux techniques. La première consiste à combiner les algorithmes génétiques avec des systèmes de contrôleurs basés sur la logique floue [80, 84, 85, 86]. La deuxième consiste à faire une analyse sur les différentes méthodes de parallélisation des algorithmes génétiques, afin de motiver la méthode de parallélisation choisie. Enfin, nous proposons une étude comparative des deux méthodes d'optimisation utilisant les algorithmes déterministes du type gradient et les algorithmes génétiques améliorés par la logique floue, au niveau qualité de la solution aussi bien qu'au niveau temps de calcul, lorsque les algorithmes génétiques sont parallélisés.

## Liste des publications parues ou à paraître

- A1 A shape optimization formulation of weld pool determination.** Applied Mathematics Letters (Accepté 2011). (En collaboration avec A. Chakib, A. Ellabib et A. Nachaoui)
- A2 Finite element approximation of an optimal design problem.** Applied and Computational Mathematics. (Accepté 2011) ( en collaboration avec A. Chakib et A. Nachaoui )
- A3 Hybrid defuzzification method for fuzzy controllers.** Dokl. Nats. Akad. Nauk Azerb. (Accepté 2011) ( en collaboration avec A. R. Y. Shikhinskaya )
- A4 Approximation and numerical realization of an optimal design welding problem.** Numerical Methods for Partial Differential Equations (en révision), ( en collaboration avec A. Chakib et A. Nachaoui )

## Liste des communications dans des conférences internationales avec comité de lecture

**C1 Parallel evolutionary algorithms, for solving a free boundary problem.** International Conference on Applied Mathematics, Modeling & Computational Science (AMMCS-2011), July 25-29 Waterloo, Canada.

**C2 Shape sensitivity analysis for a free boundary welding problem.** 4th International Conference on Approximation Methods and Numerical Modelling in Environment and Natural Resources Saidia (Morocco), May 23-26, 2011.

**C3 Numerical simulation of a heat transfer problem, via a shape optimization method.** 10th International Conference on Computational and Mathematical Methods in Science and Engineering (CMMSE), Almeria Spain.

## Liste des communications dans des conférences internationales sans comité de lecture

- C1 A enhanced combined genetic algorithm-fuzzy logic controller (GA-FLC) : application to free boundary problem.** 10th IMACS International Symposium on Iterative Methods in Scientific Computing, May 18-21, 2011 Marrakech , Morocco.
- C2 The residual velocity method applied to a free boundary problem.** Workshop & Summer School on numerical methods for interactions between sediments and water, Paris13, France.
- C3 Finite element approximation of a shape optimization for a welding problem.** First Spring School of Numerical Methods for Partial Differential Equation (NumPDEs 2010), Tetouan Morroco.
- C4 A shape optimization method for the simulation of welding problem.** Numerical Analysis and Scientific Computing with Application (NASCA09), Agadir, Morocco.
- C5 Détermination d'un front de fusion par une méthode d'optimisation de forme.** Colloque International sur les Problèmes Inverses, Contrôle et Optimisation de formes PICO08 Marrakech 16-18 avril 2008



# Résultats préliminaires

---

## Sommaire

|            |                                                        |          |
|------------|--------------------------------------------------------|----------|
| <b>1.1</b> | <b>Notations</b>                                       | <b>1</b> |
| <b>1.2</b> | <b>Espaces fonctionnels</b>                            | <b>2</b> |
| <b>1.3</b> | <b>Résultats fondamentaux</b>                          | <b>4</b> |
| <b>1.4</b> | <b>Propriété du prolongement uniforme</b>              | <b>5</b> |
| 1.4.1      | Propriété du cône                                      | 6        |
| 1.4.2      | Prolongement uniforme                                  | 6        |
| <b>1.5</b> | <b>Quelques estimations sur les espaces de Sobolev</b> | <b>7</b> |
| <b>1.6</b> | <b>Degré topologique</b>                               | <b>8</b> |
| 1.6.1      | Degré topologique de Brouwer                           | 8        |
| 1.6.2      | Degré topologique de Leray-Schauder                    | 9        |

---

Dans ce chapitre, nous rappelons quelques outils de base et nous présentons quelques résultats préliminaires essentiels pour ce travail. Nous donnons en particulier quelques résultats fondamentaux sur la propriété du prolongement uniforme, l'inégalité uniforme de Poincaré et le degré topologique.

## 1.1 Notations

1.  $\langle ., . \rangle$  désigne le produit scalaire dans  $\mathbb{R}^n$  ( $n \in \mathbb{N}$ ) défini par :

$$\langle x, y \rangle = \sum_{i=1}^n x_i y_i \text{ pour tout } x, y \in \mathbb{R}^n \quad (1.1)$$

la norme associée à ce produit scalaire est définie par :

$$|x| = \langle x, x \rangle^{\frac{1}{2}}. \quad (1.2)$$

2. Le gradient de  $\varphi$  :

$$\nabla\varphi = \left( \frac{\partial\varphi}{\partial x_1}, \dots, \frac{\partial\varphi}{\partial x_n} \right). \quad (1.3)$$

3. Le laplacien de  $\varphi$  :

$$\Delta\varphi = \sum_{i=1}^n \frac{\partial^2\varphi}{\partial x_i^2}. \quad (1.4)$$

4. La divergence d'un vecteur  $\psi$  :

$$\nabla \cdot \psi = \sum_{i=1}^n \frac{\partial\psi_i}{\partial x_i}. \quad (1.5)$$

5. La dérivée normale extérieure :

$$\nabla\varphi \cdot \nu = \frac{\partial\varphi}{\partial\nu}. \quad (1.6)$$

$\nu$  étant la normale extérieure.

6. Soit  $\alpha = (\alpha_1, \dots, \alpha_n)$  un multi-indice dans  $\mathbb{N}^n$  tel que  $|\alpha| = \alpha_1 + \dots + \alpha_n$ , la dérivée partielle d'ordre  $|\alpha|$  de  $\varphi$  notée  $D^\alpha\varphi$ , est définie par :

$$D^\alpha\varphi(x) = \frac{\partial^{|\alpha|}\varphi(x)}{\partial x_1^{\alpha_1} \dots \partial x_n^{\alpha_n}} \quad (1.7)$$

7. Soit  $\Omega$ , un ensemble mesurable de  $\mathbb{R}^n$ . On note par  $|\Omega|$  le mesure de Lebesgue de  $\Omega$ .

## 1.2 Espaces fonctionnels

Dans cette section nous rappelons quelques espaces fonctionnels, pour plus de détails à propos de ces espaces et leurs propriétés, nous renvoyons le lecteur par exemple à [13, 30]. Soit  $\Omega$  un ouvert borné de  $\mathbb{R}^n$  de frontière notée par  $\Gamma = \partial\Omega$ , et soit  $p \in [1, \infty[$ .

- l'espace  $L^p(\Omega)$  est un espace de fonctions dont la puissance  $p^{ime}$  de la fonction est intégrable, au sens de Lebesgue. L'espace  $L^p(\Omega)$  muni de la norme

$$\|\varphi\|_{p,\Omega} = \left( \int_{\Omega} |\varphi|^p dx \right)^{\frac{1}{p}}, \quad (1.8)$$

est un espace de Banach.

- En particulier, pour  $p = 2$ ,  $L^2(\Omega)$  est l'espace de Hilbert muni de la norme :

$$\|\varphi\|_{0,\Omega} = \left( \int_{\Omega} |\varphi|^2 dx \right)^{\frac{1}{2}}, \quad (1.9)$$

associée au produit scalaire :

$$(\varphi, \psi)_{0,\Omega} = \int_{\Omega} \varphi(x)\psi(x)dx. \quad (1.10)$$

- $L^\infty(\Omega)$  est l'ensemble de toutes les fonctions mesurables  $\varphi$ , dont la norme :

$$\|\varphi\|_{L^\infty(\Omega)} = \inf_{\substack{M \subset \Omega \\ \text{mes } M = 0}} \sup_{x \in \Omega \setminus M} |\varphi(x)| \equiv \text{ess sup}_{x \in \Omega} |\varphi(x)| \quad (1.11)$$

est finie.  $L^\infty(\Omega)$  muni de la norme (1.11) est un espace de Banach.

- $C(\Omega)$  désigne l'espace des fonctions continues sur  $\Omega$ , muni de la norme suivante :

$$\|\varphi\|_\infty = \sup_{x \in \Omega} |\varphi(x)|. \quad (1.12)$$

- $C^m(\Omega)$  est l'espace des fonctions,  $\varphi$ , qui ont des dérivées partielles  $D^\alpha \varphi$  continues sur  $\Omega$  (pour  $|\alpha| \leq m$ ).
- $C^m(\bar{\Omega})$  est le sous-espace de  $C^m(\Omega)$  des fonctions  $\varphi$  telles que  $D^\alpha \varphi$  soit bornée et uniformément continue sur  $\Omega$  pour  $|\alpha| \leq m$ . Il est muni de la norme :

$$\|\varphi\|_{C^m(\bar{\Omega})} = \max_{0 \leq |\alpha| \leq m} \sup_{x \in \Omega} |D^\alpha \varphi(x)|. \quad (1.13)$$

- Pour  $0 < \lambda \leq 1$ ,  $C^{m,\lambda}(\bar{\Omega})$  est le sous-espace de  $C^m(\bar{\Omega})$ , des fonctions  $\varphi$  qui sont telles que  $D^\alpha \varphi$  satisfait la condition de Hölder d'exposant  $\lambda$ . Pour tout  $|\alpha| = m$ , i.e. , il existe  $c > 0$  tel que

$$|D^\alpha \varphi(x) - D^\alpha \varphi(y)| \leq c|x - y|^\lambda \quad \forall x, y \in \Omega \quad (1.14)$$

$C^{m,\lambda}(\bar{\Omega})$  est muni de la norme

$$\|\varphi\|_{C^{m,\lambda}(\bar{\Omega})} = \|\varphi\|_{C^m(\bar{\Omega})} + \max_{|\alpha|=m} \sup_{\substack{x, y \in \Omega \\ x \neq y}} \frac{|D^\alpha \varphi(x) - D^\alpha \varphi(y)|}{|x - y|^\lambda}. \quad (1.15)$$

- $D(\Omega)$  est l'espace des fonctions indéfiniment dérivables  $C^\infty$  à support compact dans  $\Omega$ , dont le dual topologique  $D'(\Omega)$  est l'espace des distributions sur  $\Omega$ .
- $D(\bar{\Omega})$  est l'espace des restrictions à  $\Omega$  de fonctions de  $D(\mathbb{R}^n)$ .
- On définit également l'espace de Sobolev  $H^m(\Omega)$  ( $m \in \mathbb{N}$ ) par

$$H^m(\Omega) = \{\varphi \in L^2(\Omega) / D^\alpha \varphi \in L^2(\Omega), \text{ pour tout } |\alpha| \leq m\} \quad (1.16)$$

où  $D^\alpha$  est la dérivée au sens des distributions d'ordre  $|\alpha|$  de  $\varphi$ . L'espace de Hilbert  $H^m(\Omega)$  est muni de la norme :

$$\|\varphi\|_{m,\Omega} = \left( \sum_{|\alpha| \leq m} \|D^\alpha \varphi\|_{0,\Omega}^2 \right)^{\frac{1}{2}}. \quad (1.17)$$

- Pour  $s = m + \sigma$  avec  $m \in \mathbb{N}$  et  $0 < \sigma < 1$ ,  $H^s(\Omega)$  est l'ensemble des éléments  $\varphi \in H^m(\Omega)$  qui vérifient :

$$\int \int_{\Omega \times \Omega} \frac{|D^\alpha \varphi(x) - D^\alpha \varphi(y)|}{|x - y|^{m+2\sigma}} dx dy < +\infty \quad \text{pour } |\alpha| = m. \quad (1.18)$$

$H^s(\Omega)$  muni de la norme

$$\|\varphi\|_{s,\Omega} = \left( \|\varphi\|_{m,\Omega}^2 + \sum_{|\alpha|=m} \int \int_{\Omega \times \Omega} \frac{|D^\alpha \varphi(x) - D^\alpha \varphi(y)|}{|x - y|^{m+2\sigma}} dx dy \right)^{\frac{1}{2}} \quad (1.19)$$

est un espace de Banach.

- $H_0^1(\Omega)$  est l'adhérence de  $D(\Omega)$  dans  $H^1(\Omega)$ . Si la frontière  $\Gamma = \partial\Omega$  est assez régulière (cf. [71]),  $H_0^1(\Omega)$  est défini par

$$H_0^1(\Omega) = \{\varphi \in H^1(\Omega) / \varphi|_\Gamma = 0\} \quad (1.20)$$

où  $\varphi|_\Gamma$  désigne la trace de  $\varphi$  sur  $\Gamma$  (cf. [30]). L'espace de Hilbert  $H_0^1(\Omega)$  est muni de la norme :

$$|\varphi|_{1,\Omega} = \|\nabla \varphi\|_{0,\Omega}. \quad (1.21)$$

On note par  $\Gamma_1$  une partie de  $\Gamma$ , telle que  $\text{mes } \Gamma_1 > 0$ ,

$$H_{\Gamma_1}^1(\Omega) = \{\varphi \in H^1(\Omega) / \varphi|_{\Gamma_1} = 0\} \quad (1.22)$$

est l'espace de Hilbert muni de la norme (1.21).

### 1.3 Résultats fondamentaux

Dans cette partie, nous allons donner des résultats fondamentaux intervenant dans la reformulation du problème de soudage, que nous présentons dans le deuxième chapitre, en un problème d'optimisation de forme, et qui seront aussi utiles pour l'étude de l'existence d'une solution de ce problème. Dans toute la suite de cette partie,  $\Omega$  désigne un ouvert borné de  $\mathbb{R}^n$ .

**Définition 1.1** *On dit qu'une suite  $(f_n)_n$ , de fonctions à valeurs réelles dans  $\Omega$ , est uniformément bornée dans  $\Omega$ , s'il existe une constante  $C > 0$  telle que  $|f_n(x)| < C$ , pour tout  $n \in \mathbb{N}$  et  $x \in \Omega$ .*

**Définition 1.2** Soit  $F$  une famille de fonctions  $f$  définies sur  $\Omega$  et à valeurs réelles. On dit que  $F$  est équicontinue en  $x_0 \in \Omega$  si pour tout  $\varepsilon > 0$ , il existe  $\delta(\varepsilon, x_0)$  tel que :

$$\|x - x_0\| < \delta(\varepsilon, x_0) \Rightarrow |f(x) - f(x_0)| \leq \varepsilon, \quad \forall f \in F. \quad (1.23)$$

La famille  $F$  est dite équicontinue dans  $\Omega$  si elle est équicontinue pour tout  $x_0 \in \Omega$ .

Nous allons maintenant énoncer deux résultats, dont le premier est démontré dans ([36]) et le deuxième est dû à Ascoli-Arzelà ([42, 74]).

**Théorème 1.1** Soit  $u$  une fonction de  $C^{0,1}(\bar{\Omega})$  (l'espace des fonctions lipschitziennes dans  $\bar{\Omega}$ ) et  $C$  sa constante de Lipschitz, alors la dérivée partielle d'ordre 1

$$\frac{\partial u}{\partial x_i}, \quad i = 1, \dots, n \quad (1.24)$$

existe pour presque tout  $x \in \Omega$ , et de plus

$$\left| \frac{\partial u}{\partial x_i} \right| \leq C \text{ p.p dans } \Omega, \quad i = 1, \dots, n. \quad (1.25)$$

**Théorème 1.2** (Ascoli-Arzelà) Soit  $(f_n)_n$  une suite de fonctions réelles uniformément bornées et équicontinues dans  $\bar{\Omega}$ . Alors il existe une sous-suite de  $(f_n)_n$  notée  $(f_{n_k})_k$  et une fonction  $f$  continue dans  $\bar{\Omega}$  telles que  $(f_{n_k})_k$  converge uniformément vers  $f$  dans  $\bar{\Omega}$ .

Nous énonçons aussi le résultat suivant (cf. [79])

**Théorème 1.3** Soit  $h : \mathbb{R} \mapsto \mathbb{R}$  une fonction uniformément lipschitzienne telle que  $h(0) = 0$ . Alors pour tout  $u \in H_{\Gamma_1}^1(\Omega)$ , nous avons  $h(u) \in H_{\Gamma_1}^1(\Omega)$ . De plus,

$$\nabla h(u) = h'(u) \nabla u \text{ p.p dans } \Omega.$$

## 1.4 Propriété du prolongement uniforme

Soit  $\Omega$  un ouvert borné et connexe de  $\mathbb{R}^n$ , de frontière continue et lipschitzienne (cf [62]). Il existe plusieurs façons de construire un prolongement linéaire et continu d'un élément de  $H^1(\Omega)$  dans  $H^1(\mathbb{R}^n)$ , qu'on définit comme suit :

$$P_\Omega : H^1(\Omega) \rightarrow H^1(\mathbb{R}^n)$$

$$u \mapsto P_\Omega u = \tilde{u}$$

avec

$$\tilde{u}|_{\Omega} = u \text{ presque partout dans } \Omega.$$

En général, la norme de l'opérateur de prolongement  $P_{\Omega}$  dépend de  $\Omega$  (cf. [62]). Dans la suite, nous allons définir une classe de domaines pour lesquels la norme de l'opérateur prolongement est bornée indépendamment de  $\Omega$ .

#### 1.4.1 Propriété du cône

**Définition 1.3** Soit  $h > 0$  et  $\theta \in (0, \frac{\pi}{2})$  deux nombres donnés, et  $\xi \in \mathbb{R}^n$  tel que  $|\xi| = 1$ . L'ensemble :

$$C(\xi, \theta, h) = \{x \in \mathbb{R}^n ; \langle x, \xi \rangle \geq |x| \cos \theta, |x| < h\} \quad (1.26)$$

est appelé le cône d'angle  $\theta$ , de hauteur  $h$  et d'axe  $\xi$ .

**Définition 1.4** Soient  $\theta \in (0, \frac{\pi}{2})$ ,  $h > 0$  et  $r > 0$  ( $2r \leq h$ ) trois nombres donnés. On dit qu'un ensemble  $\Omega$  de  $\mathbb{R}^n$  satisfait la propriété du cône, si pour tout  $x \in \Omega$ , il existe  $C_x = C(\xi_x, \theta, h)$ , tel que pour tout  $y \in B(x, r) \cap \Omega$ , on ait

$$y + C_x \subset \Omega, \quad (1.27)$$

$B(x, r)$  étant la boule ouverte de rayon  $r$  et de centre  $x$  dans  $\mathbb{R}^n$ .

Soit  $D$  un domaine donné dans  $\mathbb{R}^n$ . Nous notons par  $\Pi(\theta, h, r)$  l'ensemble de tous les domaines contenus dans  $D$  et satisfaisant la propriété du cône.

#### 1.4.2 Prolongement uniforme

Nous allons maintenant énoncer un résultat qui montre l'existence d'un prolongement "uniforme" (dont la norme est bornée indépendamment de  $\Omega$ ) pour les domaines satisfaisant la propriété du cône (cf. [25]).

**Théorème 1.4** Soient  $\theta \in (0, \frac{\pi}{2})$ ,  $h > 0$  et  $r > 0$  ( $2r \leq h$ ) trois nombres donnés et  $m \in \mathbb{N}$ .

Il existe une constante positive  $K(\theta, h, r)$ , dépendant seulement de  $\theta, h$  et  $r$ , telle que pour tout  $\Omega \in \Pi(\theta, h, r)$ , il existe

$$P_{\Omega} : H^1(\Omega) \rightarrow H^1(\mathbb{R}^n)$$

un prolongement linéaire et continu tel que

$$\|P_\Omega\| \leq K(\theta, h, r). \quad (1.28)$$

La propriété du prolongement uniforme (1.28) est aussi vérifiée, si le domaine  $\Omega$  possède une frontière  $\partial\Omega$  continue et lipschitzienne. Ceci d'après le résultat cité dans les références suivantes (cf. [25] et [69]).

**Proposition 1.1** *Un ouvert  $\Omega$  satisfait la propriété du cône si et seulement si la frontière  $\partial\Omega$  est continue et lipschitzienne.*

## 1.5 Quelques estimations sur les espaces de Sobolev

Nous commençons par énoncer un résultat dû à Ladyženskaja [51]

**Théorème 1.5** *Soit  $\Omega$  un ouvert borné connexe de frontière lipschitzienne. Pour tout  $u \in H_{\Gamma_1}^1(\Omega)$ , on a l'inégalité suivante :*

$$\|u\|_{L^{\frac{2q}{q-2}}(\Omega)} \leq c(q)|\Omega|^{\frac{1}{n}-\frac{1}{q}}\|\nabla u\|_{0,\Omega}, \quad (1.29)$$

où  $q \geq n$  pour  $n > 2$ , et  $q > 2$  pour  $n = 2$ ,  $|\Omega|$  désigne la mesure de  $\Omega$  et la constante  $c(q)$  ne dépend que de  $q$  et de  $n$

Nous aurons aussi besoin des inégalités uniformes de Poincaré démontrées dans la référence [12].

**Théorème 1.6** *(Inégalité uniforme de Poincaré)*

*Il existe une constante  $C > 0$  telle que*

$$\|u - \frac{1}{mes \Gamma_1} \int_{\Gamma_1} u d\sigma\|_{L^2(\Omega)} \leq C \sum_{i=1}^{i=n} \|\partial_i u\|_{L^2(\Omega)}, \quad \forall u \in H^1(\Omega) \quad (1.30)$$

*pour tout  $\Omega$  ouvert connexe qui vérifie la propriété du cône.*

Comme conséquence de ce résultat nous avons :

**Corollaire 1.1** *Il existe une constante  $C > 0$  telle que*

$$\|u\|_{L^2(\Omega)} \leq C \sum_{i=1}^{i=n} \|\partial_i u\|_{L^2(\Omega)}, \quad \forall u \in H_{\Gamma_1}^1(\Omega) \quad (1.31)$$

*pour tout  $\Omega$  ouvert connexe qui vérifie la propriété du cône.*

## 1.6 Degré topologique

Dans cette section nous citons quelques éléments sur le degré topologique qui est une notion géométrique qui s'avère utile en analyse fonctionnelle non linéaire, bien qu'elle ne donne que des résultats qualitatifs généraux. Pour plus de détails le lecteur peut consulter par exemple [32, 29, 64]

### 1.6.1 Degré topologique de Brouwer

Soit  $x_0 \in \Omega \subset \mathbb{R}^n$ , si  $f$  est différentiable en  $x_0$ , on note par  $J_f(x_0) = \det f'(x_0)$  le jacobien de  $f$  en  $x_0$ .

**Définition 1.5** Soit  $\Omega \subset \mathbb{R}^n$  un ouvert borné et  $f \in C^1(\bar{\Omega})$ . On désigne par

$$S := \{x \in \Omega; \quad J_f(x) = 0\},$$

l'ensemble des points singuliers de  $f$  et on suppose que  $p \notin f(\partial\Omega) \cup f(S)$ . On définit alors le degré topologique par :

$$\deg(f, \Omega, p) = \sum_{x \in f^{-1}(p)} \operatorname{sgn} J_f(x),$$

où  $\operatorname{sgn} J_f(x)$  est le signe de  $J_f(x)$ ; avec  $\deg(f, \Omega, p) = 0$  si  $f^{-1}(p) = \emptyset$ .

Nous avons alors le résultat cité dans la référence suivante [32], qui rappelle quelques propriétés du degré topologique de Brouwer.

**Théorème 1.7** Soit  $\Omega \subset \mathbb{R}^n$  un ouvert borné et  $f : \bar{\Omega} \rightarrow \mathbb{R}^n$  une application continue. Si  $p \notin f(\partial\Omega)$ , alors il existe un entier  $\deg(f, \Omega, p)$  satisfaisant les propriétés suivantes :

1. (Normalité)  $\deg(I, \Omega, p) = 1$  si et seulement si  $p \in \Omega$ , où  $I$  note l'application identité ;
2. (Solvabilité) Si  $\deg(f, \Omega, p) \neq 0$ , alors  $f(x) = p$  admet au moins une solution dans  $\Omega$  ;
3. (Invariance par homotopie)  $\forall t \in [0, 1]$   $f_t : \bar{\Omega} \rightarrow \mathbb{R}^n$  est continue et  $p \notin \bigcup_{t \in [0, 1]} f_t(\partial\Omega)$ , alors  $\deg(f_t, \Omega, p)$  ne dépend pas de  $t \in [0, 1]$  ;
4. (Additivité) Supposons que  $\Omega_1$  et  $\Omega_2$  sont deux sous-ensembles disjoints et ouverts de  $\Omega$  et  $p \notin f(\bar{\Omega} - \Omega_1 \cup \Omega_2)$ . Alors

$$\deg(f, \Omega, p) = \deg(f, \Omega_1, p) + \deg(f, \Omega_2, p);$$

5.  $\deg(f, \Omega, p)$  est constant sur toute composante connexe de  $\mathbb{R}^n \setminus f(\partial\Omega)$ .

En 1934, Leray et Schauder (cf. [53]) ont généralisé le degré topologique de Brouwer à des espaces de Banach de dimension infinie. Ainsi, ils ont défini ce qu'on appelle le degré topologique de Leray-Schauder. Ce résultat est devenu un outil très efficace pour montrer différents résultats d'existence pour les équations aux dérivées partielles non linéaires.

### 1.6.2 Degré topologique de Leray-Schauder

Pour construire le degré de Leray-Schauder, nous avons besoin de quelques résultats et définitions [26].

**Lemme 1.1** *Soit  $E$  un espace de Banach,  $\Omega \subset E$  un ouvert borné et  $T : \bar{\Omega} \mapsto E$  une application continue compacte. Alors, pour tout  $\varepsilon > 0$ , il existe un espace de dimension finie noté  $F$  et une application continue  $T_\varepsilon : \bar{\Omega} \mapsto F$  telle que*

$$\|T_\varepsilon x - Tx\| < \varepsilon \quad \text{pour tout } x \in \bar{\Omega}.$$

**Définition 1.6** *Soit  $E$  un espace de Banach,  $\Omega \subset E$  un ouvert borné et  $T : \bar{\Omega} \mapsto E$  une application continue compacte. Supposons maintenant que  $0 \notin (I - T)(\partial\Omega)$ . Il existe  $\varepsilon_0 > 0$  tel que pour  $\varepsilon \in (0, \varepsilon_0)$ , le degré de Brouwer  $\deg(I - T_\varepsilon, \Omega \cap F_\varepsilon, 0)$  est bien défini, où  $T_\varepsilon$  est défini comme dans le lemme 1.1. Par conséquent, nous définissons le degré de Leray-Schauder par*

$$\deg(I - T, \Omega, 0) = \deg(I - T_\varepsilon, \Omega \cap F_\varepsilon, 0).$$

**Théorème 1.8** *Le degré de Leray-Schauder possède les propriétés suivantes :*

1. (Normalité)  $\deg(I, \Omega, 0) = 1$  si et seulement si  $0 \in \Omega$  ;
2. (Solvabilité) Si  $\deg(I - T, \Omega, 0) \neq 0$  alors  $Tx = x$  a une solution dans  $\Omega$  ;
3. (Invariance par homotopie) Soit  $T_t : [0, 1] \times \bar{\Omega} \mapsto E$  continu compact et  $T_t x \neq x$  pour tout  $(t, x) \in [0, 1] \times \partial\Omega$ . Alors  $\deg(I - T_t, \Omega, 0)$  ne dépend pas de  $t \in [0, 1]$  ;
4. (Additivité) Soit  $\Omega_1$  et  $\Omega_2$  deux sous-ensembles disjoints ouverts de  $\Omega$  et

$$0 \notin (I - T)(\bar{\Omega} - \Omega_1 \cup \Omega_2).$$

Alors

$$\deg(I - T, \Omega, 0) = \deg(I - T, \Omega_1, 0) + \deg(I - T, \Omega_2, 0).$$



# Modélisation et formulation du problème

---

## Sommaire

|            |                                                      |           |
|------------|------------------------------------------------------|-----------|
| <b>2.1</b> | <b>Généralités sur le problème de soudage</b>        | <b>11</b> |
| <b>2.2</b> | <b>Modélisation du problème</b>                      | <b>13</b> |
| 2.2.1      | Conditions aux limites                               | 14        |
| 2.2.2      | Système d'équations dans la partie solide            | 16        |
| 2.2.3      | Formulation du problème en régime quasi-stationnaire | 17        |
| <b>2.3</b> | <b>Formulation en optimisation de forme</b>          | <b>17</b> |

---

Dans ce chapitre, nous présentons une description physique du problème qui modélise l'analyse des transferts de chaleur dans une opération de soudage. Il s'agit d'estimer le champ de température dans les pièces soudées afin de prévoir et maîtriser les effets mécaniques engendrés par le procédé sur ces pièces (contraintes résiduelles, distorsions). L'approche que nous considérons ne s'occupe que de la partie solide de la plaque. Elle consiste à simplifier les phénomènes physiques apparaissant entre la torche de soudage et la plaque ainsi que ceux du bain liquide en introduisant une condition de température imposée sur le front de fusion. Nous proposons ensuite une formulation en optimisation de forme de ce problème.

## 2.1 Généralités sur le problème de soudage

Les procédés de soudage occupent une place importante dans l'univers de la construction des bateaux, trains, avions, fusées, automobiles, ponts, réservoirs et tant d'autres choses qui ne sauraient être assemblées autrement. Le soudage permet d'assurer la continuité métallique entre les éléments à assembler par l'établissement de forces de liaison interatomiques de type métallique.

En général, les procédés de soudage amènent localement les pièces à souder à une température supérieure à la température de fusion. Ainsi, l'assemblage par liaisons métalliques s'effectue lors de la solidification. Ces procédés induisent des conséquences métallurgiques et mécaniques qu'il est aussi difficile qu'important de maîtriser et dont l'évaluation nécessite la modélisation des interactions complexes entre plusieurs phénomènes physiques comme la thermique, la métallurgie et la mécanique. Les différents procédés de soudage sont classés surtout selon des critères pratiques :

- Le mode d'apport de l'énergie nécessaire : faisceau, flamme, arc électrique, plasma, effet joule, étincelage, induction, frottement, explosion, etc.
- Le mode de protection du métal chaud : gaz ou laitier.

Ces procédés conduisent à des modifications de microstructure et à des contraintes et distorsions résiduelles qui jouent un rôle important sur la tenue mécanique des assemblages ou encore la faisabilité d'un procédé. Ces contraintes et distorsions proviennent principalement des gradients de température et des éventuelles transformations de phase susceptibles de se produire au cours du procédé.

Les apports de la modélisation du soudage se situent donc au niveau :

- des études de faisabilité d'un procédé en termes de défaut d'alignement des structures ou de séquence de soudage.
- de l'évaluation de la tenue mécanique des assemblages soudés en termes de résistance, de tenue en fatigue ou encore de stabilité.

La faisabilité d'un procédé se juge en termes de distorsions résiduelles. Elle nécessite des simulations portant sur la structure dans sa totalité et comportant l'ensemble des joints soudés. L'objectif peut être ici de prédire les éventuels défauts d'alignement en vue de dimensionner les conditions de bridage ou encore de déterminer une séquence de soudage optimale comme étant celle conduisant à minimiser le coût des outillages de bridage. L'évaluation de la tenue mécanique repose principalement sur la connaissance de la microstructure et des contraintes résiduelles. Le risque de rupture fragile peut ainsi être fortement augmenté par la présence de contraintes de traction, de fortes triaxialités et de phases dures comme celles obtenues dans les Zones Affectées Thermiquement (ZAT). Des analyses détaillées peuvent être réalisées dans le cadre de la mécanique linéaire élastique de la rupture en postulant la présence d'un défaut correspondant à une situation pénalisante.

Aux contraintes résultant du chargement (dans une hypothèse d'élasticité linéaire) s'ajoutent les contraintes résiduelles, qui contribuent ainsi à modifier cet état moyen. Si ce sont des contraintes de compression, elles influent positivement sur la durée de vie en fatigue. Au contraire, dans le cas de contraintes de traction, l'effet est négatif. Afin d'évaluer les conditions d'amorçage, il convient de recourir à des critères locaux comme le critère de Dang Van [82] très utilisé, par exemple, dans l'industrie automobile.

La fissuration à froid désigne l'apparition de fissures à des températures inférieures à 200°C. Elle résulte d'un phénomène de fragilisation par l'hydrogène. Les fissures apparaissent principalement dans la ZAT et sont favorisées par la présence d'une forte concentration en hydrogène et de contraintes de traction. L'évaluation du risque de fissuration à froid nécessite donc de simuler les phénomènes de diffusion et piégeage de l'hydrogène et de disposer d'un critère local d'amorçage. Des fissures peuvent également apparaître à chaud, dans la zone pâteuse, pendant la solidification du bain fondu. Pour évaluer la fissuration à chaud, il faut disposer d'un critère adéquat faisant intervenir les contraintes, les déformations et la vitesse de déformation viscoplastique [49]. Des problèmes d'instabilité peuvent apparaître dans les structures minces. Ils sont favorisés par des contraintes résiduelles de compression et d'éventuelles imperfections géométriques éventuellement induites par le soudage. L'étude de ces problèmes peut se faire soit par la recherche de chargements critiques de flambage en résolvant un problème aux valeurs propres, soit en réalisant une analyse incrémentale incluant les grands déplacements et grandes déformations de la structure. L'analyse de la tenue mécanique d'un assemblage soudé nécessite donc des études complémentaires qui dépendent du risque considéré. Mais toutes demandent en entrée la microstructure, les contraintes et déformations résiduelles de soudage.

Les effets mécaniques qui ont été présentés dans ce paragraphe sont directement dépendants des évolutions de températures imposées par le procédé de soudage i.e. le chargement thermique du procédé.

Dans la section suivante, nous allons décrire l'approche que nous avons considérée pour pouvoir accéder à ce chargement thermique.

## 2.2 Modélisation du problème

Nous commençons par définir un système d'équations mathématiques qui modélise le procédé de soudage. Nous utilisons un système de coordonnées lié à la torche (source de la chaleur

figure 2.1) qui se déplace avec une vitesse constante  $\vec{v}_{torche} = v_{torche} \vec{e}_x$ , où  $v_{torche}$  est la vitesse de la torche et  $\vec{e}_x$  est le vecteur unitaire aligné avec l'axe  $(ox)$ . L'utilisation d'un référentiel lié à la torche de soudage nous permettra de décrire l'évolution du bain de fusion, par le mouvement de l'interface de changement de phase selon trois régimes :

- Transitoire : un bain de fusion est établi et l'interface liquide/solide se développe au cours du temps.
- Quasi-stationnaire : la forme du bain de fusion est considérée constante.
- Solidification transitoire : fin de l'application de la source d'énergie sur la pièce. La zone fondue commence à se solidifier.

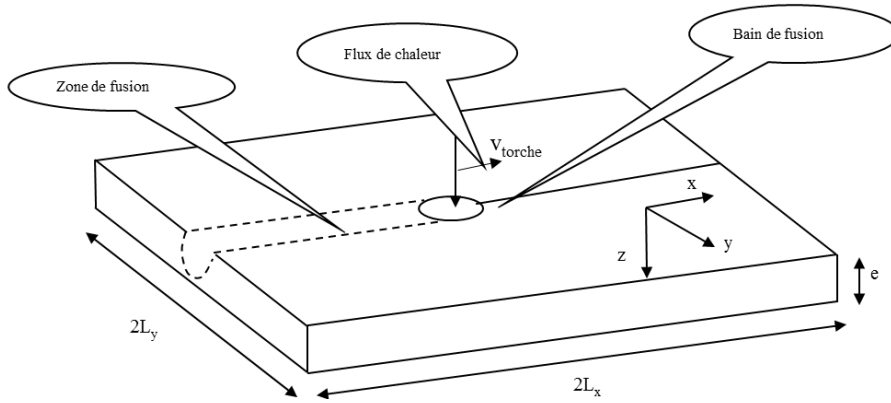


FIGURE 2.1 : Schéma du procédé de soudage

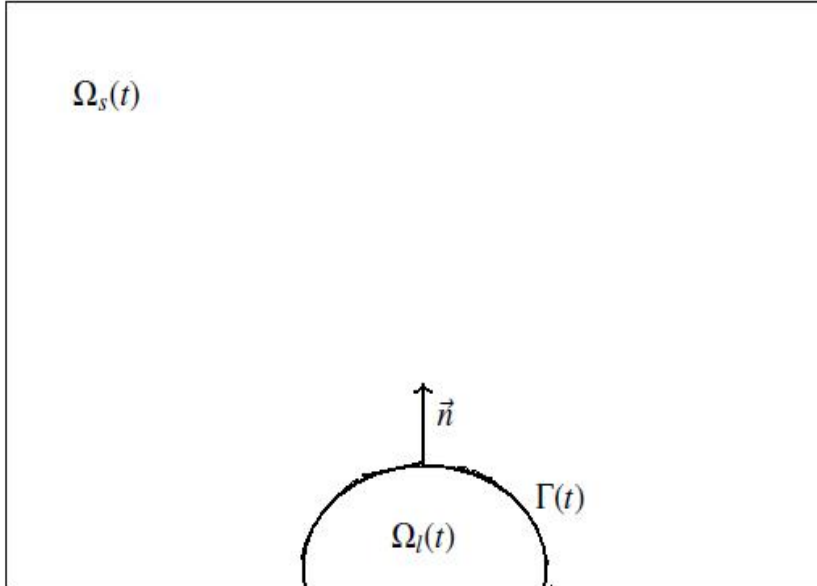
Notre approche consiste à restreindre la modélisation de transfert de la chaleur sur la partie solide  $\Omega_s(t)$ .

Nous divisons la frontière  $\partial\Omega_s$  en trois parties  $\partial\Omega_D$ ,  $\partial\Omega_N$ ,  $\partial\Omega_M$ , sur lesquelles nous appliquons respectivement les conditions aux limites de type Dirichlet, Neumann et mixte.

### 2.2.1 Conditions aux limites

Le domaine  $\Omega_s$  est un parallélépipède tronqué de la zone fondue. Ses dimensions sont définies d'une manière à satisfaire des conditions aux limites spécifiques.

(b)

FIGURE 2.2 : Description du domaine  $\Omega_{tot} = \Omega_s(t) \cup \Omega_l(t)$  à l'instant  $t$ 

### 2.2.1.1 Conditions aux limites sur l'interface liquide/solide $\Gamma(t)$

L'interface liquide/solide est caractérisée respectivement par la donnée de la température de fusion et par une condition de type Stefan qui décrit le bilan énergétique à l'interface liquide/solide et qui exprime l'absorption d'énergie due au changement de phase.

$$\begin{aligned} T &= T_f \\ \lambda \frac{\partial T}{\partial n} &= \vec{\phi}_l \vec{n} - \rho L (\vec{v} + \vec{v}_{torch}) \vec{n} \end{aligned}$$

où  $\lambda$  est la conductivité thermique,  $\phi_l$  est le flux provenant du liquide,  $\vec{n}$  est la normale extérieure à la frontière  $\Gamma(t)$ ,  $\rho$  est la densité de la plaque,  $L$  est la longueur de la plaque,  $\vec{v}$  est la vitesse de déformation et  $\vec{v}_{torch}$  est la vitesse de la torche.

### 2.2.1.2 Condition aux limites sur la surface du haut $z = e$ et celle du bas $z = 0$ ( $\partial\Omega_M$ )

La condition aux limites sur la surface du haut  $z = e$  et celle du bas  $z = 0$  est donnée par

$$\lambda \frac{\partial T}{\partial n} + hT = hT_a$$

où  $T_a$  est la température ambiante et  $h$  est le coefficient de transfert de la chaleur.

### 2.2.1.3 Condition aux limites sur la frontière de Dirichlet $\partial\Omega_D$

La condition de Dirichlet est donnée par

$$T = T_{imp},$$

où  $T_{imp}$  est une température imposée sur la frontière  $\partial\Omega_D$ .

### 2.2.1.4 Condition aux limites sur la frontière de Neumann

La condition de Neumann est décrite par

$$\lambda \frac{\partial T}{\partial n} = \phi_N$$

où  $\phi_N$  est le flux de la chaleur imposé.

## 2.2.2 Système d'équations dans la partie solide

Les équations régissant le transfert de la chaleur dans le solide doivent prendre en compte le flux de chaleur  $\vec{\phi}_l$  provenant du liquide. Le transfert de chaleur dans  $\Omega_s(t)$  est donc régi par le système d'équations suivant :

$$\left\{ \begin{array}{ll} \rho C_p \left( \frac{\partial T}{\partial t} + \vec{v}_{torch} \nabla T \right) - \nabla \cdot (\lambda_s \nabla T) & = f \quad \text{dans } \Omega_s(t) \times [0, t_{fin}] \\ T & = T_f \quad \text{sur } \Gamma(t) \\ \lambda_s \frac{\partial T}{\partial n} & = \vec{\phi}_l \vec{n} - \rho L (\vec{v} + \vec{v}_{torch}) \vec{n} \quad \text{sur } \Gamma(t) \\ T(x, y, z, 0) & = T_0(x, y, z) \quad \text{pour } t = 0 \\ T & = T_{imp} \quad \text{sur } \partial\Omega_D \times [0, t_{fin}] \\ \lambda_s \frac{\partial T}{\partial n} & = \phi_N \quad \text{sur } \partial\Omega_N \times [0, t_{fin}] \\ \lambda_s \frac{\partial T}{\partial n} + hT & = hT_a \quad \text{sur } \partial\Omega_M \times [0, t_{fin}] \end{array} \right. \quad (2.1)$$

où  $h$  est le coefficient de transfert global intégrant les échanges de chaleur par convection et par rayonnement  $h = h_{cv} + h_{ray}$  et  $C_p$  est la capacité thermique.

### 2.2.3 Formulation du problème en régime quasi-stationnaire

En régime quasi-stationnaire, la forme du front  $\Gamma(t)$  est établie et ne dépend pas de  $t$  (la vitesse de déformation  $\vec{v}$  est nulle). Cependant, sa position est a priori inconnue. Dans ce cadre, le champ thermique dans le solide est déterminé en considérant seulement la condition  $T = T_f$  sur l'interface liquide/solide  $\Gamma(t)$ , où  $T_f$  désigne la température de fusion. Nous résoudrons donc un problème à frontière libre de conduction de la chaleur dans  $\Omega_s$ .

La configuration en 2-D est obtenue en supposant que la variation de la température dans la direction  $(oz)$  est négligeable par rapport à celle observée dans les deux autres directions  $(Ox)$  et  $(Oy)$ . Pour les conditions de Neumann sur la frontière  $\partial\Omega_N = \Gamma_0 \cup \Gamma_1 \cup \Gamma_2 \cup \Gamma_3$ , nous considérons que la plaque est isolée des effets extérieurs ( $\phi_N = 0$ ).

Dans la suite de notre travail, la partie solide  $\Omega_s$  et le front de fusion  $\Gamma(t)$  seront notés respectivement par  $\Omega$  et  $\Gamma$ .

Le problème consiste alors à déterminer  $(T, \Gamma)$ , le champ de température  $T$  qui règne dans la partie solide  $\Omega$  (voir figure 2.3) et  $\Gamma$  le front de fusion, solution du problème :

$$(P) \left\{ \begin{array}{ll} \mathcal{V}_0 \cdot \nabla T - \nabla \cdot (\lambda \nabla T) &= f \quad \text{dans } \Omega \\ T &= T_f \quad \text{sur } \Gamma \\ T &= T_d \quad \text{sur } \Gamma_4 \\ \lambda \frac{\partial T}{\partial n} &= 0 \quad \text{sur } \Gamma_0 \cup \Gamma_1 \cup \Gamma_2 \cup \Gamma_3 \\ T &= T_0 \quad \text{sur } \Gamma_0 \end{array} \right. \quad (2.2)$$

Où  $\mathcal{V}_0 = C_p \rho \vec{v}_{torch}$  et  $f$  est un terme source donné. On note par  $T_d$  la température imposée, généralement la température du milieu extérieur. Quant à  $T_0$ , il désigne la mesure de température prise sur la frontière  $\Gamma_0$ .

## 2.3 Formulation en optimisation de forme

Un problème d'optimisation de forme est un problème de contrôle optimal dans lequel la variable de contrôle est la forme géométrique d'un domaine de  $\mathbb{R}^n$ .

L'optimisation de forme joue un rôle important dans la résolution de plusieurs problèmes d'ingénierie. Plusieurs travaux, utilisant la méthode d'optimisation de forme, se sont intéressés à l'étude théorique et à l'approximation numérique de différents phénomènes physiques. Nous citons par exemple : [41, 44, 69]

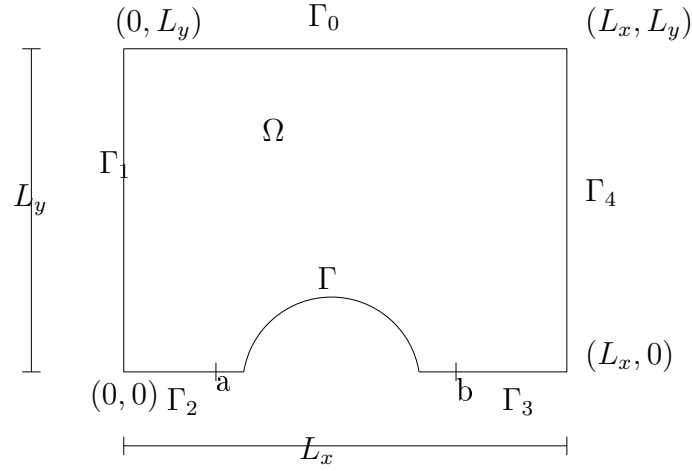


FIGURE 2.3 : configuration 2-D du procédé de soudage

L'objet principal de la méthode d'optimisation de forme est de mieux concevoir la forme d'un système physique, afin de minimiser le coût et maximiser le profit. Il s'agit en fait de déterminer un domaine optimal réalisant le minimum d'une fonction coût donnée, définie sur un ensemble de domaines admissibles. En général, un problème d'optimisation de forme se présente de la façon suivante :

$$\text{Trouver} \quad \Omega^* = \arg \min_{\Omega \in \Theta_{ad}} J(\Omega, u(\Omega)), \quad (2.3)$$

où

- $\Theta_{ad}$  est un ensemble des domaines admissibles  $\Omega \subset \mathbb{R}^n$ ,
- $u(\Omega)$  est l'état d'un système, solution d'un problème aux limites du type :

$$\begin{cases} \mathcal{A}(u(\Omega)) = 0 & \text{dans } \Omega, \\ \mathcal{B}(u(\Omega)) = 0 & \text{dans } \Gamma, \end{cases} \quad (2.4)$$

où  $\mathcal{A}, \mathcal{B}$  sont des opérateurs aux dérivées partielles linéaires ou non, donnés dans  $\mathbb{R}^n$ , et  $\Gamma$  est la frontière de  $\Omega$ ,

- la fonction  $J(\Omega, u(\Omega))$  est la fonctionnelle coût définie comme combinaison des intégrales :

$$\int_{\Omega} \mathcal{C}_1(u(\Omega))dx, \quad \int_{\Gamma} \mathcal{C}(u(\Omega))dx, \quad \int_{\Omega_0} \mathcal{C}_2(u(\Omega))dx$$

où  $\Omega_0$  est un ouvert fixe de  $\mathbb{R}^n$ ,  $\Omega_0 \subset \Omega$ , et  $\mathcal{C}$ ,  $\mathcal{C}_1$  et  $\mathcal{C}_2$  sont des opérateurs aux dérivées partielles, données dans  $\mathbb{R}^n$ .

Dans les paragraphes suivants, nous allons donner une formulation en optimisation de forme du problème (2.2)

### Formulation en optimisation de forme du problème de soudage

Pour donner une formulation en optimisation de forme du problème (2.2), nous utilisons un paramétrage de la frontière libre  $\Gamma$  de la manière suivante :

$$\Gamma = \{(x, \varphi(x)) \mid x \in [a, b] \text{ et } 0 \leq \varphi(x) \leq L_y\} \quad (2.5)$$

où  $\varphi$  est une fonction. Nous notons par  $\Omega$  l'ouvert défini par :

$$\Omega(\varphi) = ]0, a[ \times ]0, L_y[ \cup \{(x, y) \in \mathbb{R}^2 \mid a \leq x \leq b, \varphi(x) \leq y \leq L_y\} \cup ]b, L_x[ \times ]0, L_y[ \quad (2.6)$$

de frontière  $\partial\Omega$  définie par :

$$\partial\Omega = \Gamma_0 \cup \Gamma_1 \cup \Gamma_2 \cup \Gamma \cup \Gamma_3 \cup \Gamma_4 \quad (2.7)$$

Notre formulation en optimisation de forme du problème (2.2) est alors donnée par

$$(PO) \left\{ \begin{array}{l} \text{trouver } \Omega^* \in \Theta_{ad} \text{ solution de} \\ J(\Omega^*) = \inf_{\Omega \in \Theta_{ad}} J(\Omega) \\ \text{où } J(\Omega) = \frac{1}{2} \int_{\Gamma_0} |T(\Omega(x, y)) - T_0|^2 d\sigma \\ \text{et } T(\Omega) \text{ la solution de} \\ \begin{array}{l} (PE) \left\{ \begin{array}{l} \mathcal{V}_0 \cdot \nabla T = \nabla \cdot (\lambda \nabla T) + f \text{ dans } \Omega \\ \lambda \frac{\partial T}{\partial \nu} = 0 \text{ sur } \Gamma_0 \cup \Gamma_1 \cup \Gamma_2 \cup \Gamma_3 \\ T = T_d \text{ sur } \Gamma_4 \\ T = T_f \text{ sur } \Gamma, \end{array} \right. \end{array} \right. \end{array} \right. \quad (2.8)$$

où  $\Theta_{ad}$  est défini par

$$\Theta_{ad} = \{\Omega(\varphi) \mid \varphi \in U_{ad}\}$$

avec

$$U_{ad} = \left\{ \varphi \in C([a, b]) \mid \exists a_\varphi \text{ et } b_\varphi \in [a, b], \varphi|_{[a, a_\varphi]} = 0, \varphi|_{[b_\varphi, b]} = 0 \text{ et } \exists L_0 > 0 / \right. \\ \left. |\varphi(x) - \varphi(x')| \leq L_0 |x - x'| \quad \forall x, x' \in [a, b], \quad 0 \leq \varphi(x) \leq L_y \quad \forall x \in [a, b] \right\}.$$

Soit  $\Gamma_d = \Gamma \cup \Gamma_4$ . Nous définissons l'espace

$$H_{\Gamma_d}^1(\Omega) = \{u \in H^1(\Omega) \mid u|_{\Gamma_d} = 0\}$$

où  $H^1(\Omega)$  est un espace de Sobolev muni de la norme,  $\|\cdot\|_{1,\Omega}$ . et  $H_{\Gamma_d}^1(\Omega)$  est muni de la norme  $\|\cdot\|_{1,\Omega}$

### Equivalence entre (P) et (PO)

Nous avons le résultat d'équivalence entre le problème à frontière libre (P) et le problème d'optimisation de forme (PO)

**Proposition 2.1** *Si le problème (P) admet une solution  $(u_p, \Omega_p)$  avec  $\Omega_p \in \Theta_{ad}$ , alors le problème (P) est équivalent au problème (PO).*

**Démonstration :**

Soit  $(u_p, \Omega_p)$  solution de  $(P)$ , telle que  $\Omega_p \in \Theta_{ad}$ . Comme  $u_p = 0$  sur  $\Gamma_0$ , alors

$$J(u_p, \Omega_p) = 0 \leq J(u, \Omega), \quad \forall \Omega \in \Theta_{ad}.$$

Par suite  $(u_p, \Omega_p)$  est solution de  $(PO)$ .

Inversement, soit  $(u^*, \Omega^*)$  solution de  $(PO)$ , alors

$$\begin{aligned} 0 \leq J(u^*, \Omega^*) &\leq J(u, \Omega) \quad \forall \Omega \in \Theta_{ad} \\ &\leq J(\Omega_p, u_p) = 0. \end{aligned}$$

Alors  $(u^*, \Omega^*)$  est solution de  $(P)$ .

■

**Remarque 2.1** Dans le cas où  $(PO)$  n'est pas équivalent à  $(P)$ , le problème  $(PO)$  représente un moyen efficace d'approximation du problème à frontière libre  $(P)$ .



# Étude de l'existence d'une solution optimale

---

## Sommaire

|            |                                                                              |           |
|------------|------------------------------------------------------------------------------|-----------|
| <b>3.1</b> | <b>Introduction à l'étude d'existence en optimisation de forme . . . . .</b> | <b>24</b> |
| <b>3.2</b> | <b>Étude du problème d'état . . . . .</b>                                    | <b>25</b> |
| 3.2.1      | Formulation variationnelle . . . . .                                         | 25        |
| 3.2.2      | Existence et unicité de la solution d'équation d'état . . . . .              | 26        |
| <b>3.3</b> | <b>Étude de l'existence d'une solution optimale . . . . .</b>                | <b>30</b> |
| 3.3.1      | Topologie sur $\mathcal{F}$ . . . . .                                        | 30        |
| 3.3.2      | Continuité du problème d'état . . . . .                                      | 35        |
| 3.3.3      | Continuité de la fonctionnelle coût . . . . .                                | 48        |

---

Dans ce chapitre, nous donnons un résultat d'existence de la solution du problème d'optimisation de forme. Comme souvent les hypothèses physiques prises sur les données du problème ne permettent pas d'avoir la coercivité du problème d'état qui est nécessaire pour avoir l'existence et l'unicité de la solution en utilisant le lemme de Lax-Milgram (l'outil classique pour l'étude des problèmes elliptiques linéaires [13]). Nous utilisons alors les degrés topologiques de Leray Schauder [32], pour montrer l'existence et d'une solution du problème d'état écrit sous sa forme variationnelle. Quant à l'unicité, elle s'obtient facilement grâce à la linéarité du problème d'état. Ensuite, nous utilisons des techniques d'estimations uniformes basées sur l'inégalité de Poincaré uniforme [12] et quelques inégalités de Sobolev [51], pour montrer la continuité du problème d'état qui présente l'une des difficultés principales de l'étude de l'existence d'une solution optimale. Ce résultat a fait l'objet d'un article accepté [15] et d'une communication au "Numerical Analysis and Scientific Computing with Application (NASCA09), Agadir, Morocco".

### 3.1 Introduction à l'étude d'existence en optimisation de forme

Un problème d'optimisation de forme peut être vu comme étant un problème de minimisation d'une fonction à valeurs réelles, la fonction coût notée  $J(\Omega, u(\Omega))$ , dont les arguments sont un ouvert  $\Omega$ , élément d'un certain ensemble  $\Theta_{ad}$  d'ouverts admissibles, et  $u(\Omega)$  la solution d'un problème aux limites posé sur  $\Omega$ . Du point de vue mathématique, la question qui se pose est : existe-t-il une solution optimale ?

Pour pouvoir répondre à cette question, il faut définir une topologie sur  $\Theta_{ad}$ , qui doit assurer la la compacité de  $\Theta_{ad}$ , ainsi que la semi-continuité inférieure de  $J$  sur  $\Theta_{ad}$ .

Soit  $D$  un ouvert borné connexe de  $\mathbb{R}^n$  tel que  $\forall \Omega \in \Theta_{ad}, \Omega \subset D$ . Considérons une suite  $(\Omega_n)_n$ ,  $\Omega_n \in \Theta_{ad}$  et  $\Omega \in \Theta_{ad}$ . Il nous est nécessaire de définir la règle qui nous permette de dire que  $(\Omega_n)_n$  tend vers  $\Omega$  i.e.

$$\Omega_n \xrightarrow{\Theta_{ad}} \Omega, \quad n \longrightarrow \infty.$$

Ainsi, nous associons à chaque  $\Omega \in \Theta_{ad}$  un espace fonctionnel  $W(\Omega)$  qui contient  $u(\Omega)$ , la solution du problème (2.4). Nous introduisons alors la convergence de  $(u(\Omega_n))_n$ ,  $u(\Omega_n) \in W(\Omega_n)$ ,  $\Omega_n \in \Theta_{ad}$  vers  $(u(\Omega)) \in W(\Omega)$ ,  $\Omega \in \Theta_{ad}$ , de la manière suivante :

$$u(\Omega_n) \xrightarrow{n \rightarrow +\infty} u(\Omega), \quad \text{dans } W(D).$$

Soit  $\mathcal{F}$  le graphe de l'application  $u(\cdot)$  définie sur l'ensemble  $\Theta_{ad}$  i.e.

$$\mathcal{F} = \{(\Omega, u(\Omega)), \Omega \in \Theta_{ad}\}.$$

L'existence d'une solution optimale du problème (2.3) est fournie par le théorème suivant (cf. [41]) par exemple.

**Théorème 3.1** *Sous les hypothèses suivantes :*

$(H_1)$  (compacité de  $\mathcal{F}$ ) :

Pour chaque suite  $\{(\Omega_n, u(\Omega_n))\}_n$ ,  $(\Omega_n, u(\Omega_n)) \in \mathcal{F}$ , il existe une sous suite  $\{(\Omega_{n_k}, u(\Omega_{n_k}))\}_k$  et un élément  $(\Omega, u(\Omega)) \in \mathcal{F}$ , tels que

$$\Omega_{n_k} \xrightarrow{\Theta_{ad}} \Omega \tag{3.1}$$

$$u(\Omega_{n_k}) \xrightarrow{k \rightarrow +\infty} u(\Omega), \quad \text{dans } W(D).$$

( $H_2$ ) (*Semi-continuité de  $J$* ) :

$$\left. \begin{array}{l} \Omega_n \xrightarrow{\Theta_{ad}} \Omega \\ u(\Omega_n) \xrightarrow[n \rightarrow +\infty]{} u(\Omega), \quad \text{dans } W(D). \end{array} \right\} \Rightarrow \liminf_{n \rightarrow +\infty} J(\Omega_n, u(\Omega_n)) \geq J(\Omega, u(\Omega)), \quad (3.2)$$

le problème (2.3) admet au moins une solution  $\Omega^* \in \Theta_{ad}$ .

Avant de s'intéresser à l'étude de l'existence d'une solution optimale, il faudra tout d'abord vérifier si le problème d'optimisation de forme est bien posé. Cela signifie que pour tout  $\Omega \in \Theta_{ad}$  le problème d'état ( $PE$ ) admet une solution unique.

## 3.2 Étude du problème d'état

Dans cette section nous donnerons une formulation variationnelle du problème d'état ( $PE$ ). Puis nous démontrerons l'existence et l'unicité de la solution de ce problème.

### 3.2.1 Formulation variationnelle

Le problème d'état ( $PE$ ) s'écrit :

$$(PE) \left\{ \begin{array}{l} \nu_0 \cdot \nabla T = \nabla \cdot (\lambda \nabla T) + f \quad \text{dans } \Omega, \\ \lambda \frac{\partial T}{\partial \nu} = 0 \quad \text{sur } \Gamma_0 \cup \Gamma_1 \cup \Gamma_2 \cup \Gamma_3, \\ T = T_d \quad \text{sur } \Gamma_4, \\ T = T_f \quad \text{sur } \Gamma. \end{array} \right. \quad (3.3)$$

En utilisant la surjectivité de l'application trace de  $H^1(D)$  dans  $H^{\frac{1}{2}}(\partial D)$  (cf. [30]), on montre qu'il existe  $V \in H^1(D)$  tel que

$$V = \left\{ \begin{array}{l} v \quad \text{dans } ]b, L_x[ \times ]0, L_y[ \\ T_f \quad \text{dans } ]0, b[ \times ]0, L_y[, \end{array} \right.$$

où  $v \in H^1([b, L_x[\times]0, L_y])$  sachant que

$$v = \begin{cases} T_d & \text{dans } \Gamma_4 \\ T_f & \text{dans } \{b\} \times [0, L_y]. \end{cases}$$

Ainsi la formulation variationnelle du problème d'état (PE) est la suivante :

$$\begin{cases} \text{trouver } u \in H_{\Gamma_d}^1(\Omega) \\ \int_{\Omega} \lambda \nabla u \cdot \nabla \psi + \int_{\Omega} \mathcal{V}_0 \cdot \nabla u \psi = \langle \ell, \psi \rangle_{((H_{\Gamma_d}^1(\Omega))', H_{\Gamma_d}^1(\Omega))} \quad \forall \psi \in H_{\Gamma_d}^1(\Omega), \end{cases} \quad (3.4)$$

où  $\ell$  est l'opérateur défini par :

$$\langle \ell, \psi \rangle_{((H_{\Gamma_d}^1(\Omega))', H_{\Gamma_d}^1(\Omega))} = \int_{\Omega} f \psi - \int_{\Omega} \lambda \nabla V \cdot \nabla \psi - \int_{\Omega} \mathcal{V}_0 \cdot \nabla V \psi.$$

### 3.2.2 Existence et unicité de la solution d'équation d'état

Dans toute la suite nous allons supposer que les paramètres de notre problème vérifient les hypothèses suivantes :

(H<sub>1</sub>)  $\lambda \in L^\infty(D)$  et  $\exists \lambda_0 > 0$  tel que

$$\lambda(x) \xi \cdot \xi \geq \lambda_0 |\xi|^2 \quad \text{a.e } x \in D$$

(H<sub>2</sub>)  $\mathcal{V}_0 \in (L^\infty(D))^2$

(H<sub>3</sub>)  $f \in L^2(D), \quad T_0 \in L^2(\Gamma_0)$

**Remarque 3.1** On note que si  $\mathcal{V}_0$  satisfait seulement l'hypothèse (H<sub>2</sub>), alors le problème (3.4) est en général non-coercif, au sens du théorème de Lax-Milgram (l'outil classique pour l'étude des problèmes elliptiques linéaires).

**Proposition 3.1** Sous les hypothèses (H<sub>1</sub>) – (H<sub>3</sub>), il existe une constante  $\delta > 0$  indépendante de  $\Omega$  telle que :

$$\forall \psi \in H_{\Gamma_d}^1(\Omega) \quad \langle \ell, \psi \rangle_{((H_{\Gamma_d}^1(\Omega))', H_{\Gamma_d}^1(\Omega))} \leq \delta \|\psi\|_{1,\Omega}. \quad (3.5)$$

**Démonstration :**

Soit  $\psi \in H_{\Gamma_d}^1(\Omega)$ , on sait que l'opérateur  $\ell$  est défini par :

$$\langle \ell, \psi \rangle_{((H_{\Gamma_d}^1(\Omega))', H_{\Gamma_d}^1(\Omega))} = \int_{\Omega} f \psi - \int_{\Omega} \lambda \nabla V \cdot \nabla \psi - \int_{\Omega} \mathcal{V}_0 \cdot \nabla V \psi.$$

En utilisant l'inégalité de Hölder, nous obtenons

$$\begin{aligned} |\langle \ell, \psi \rangle_{((H_{\Gamma_d}^1(\Omega))', H_{\Gamma_d}^1(\Omega))}| &\leq \int_{\Omega} |f \psi| + \int_{\Omega} |\lambda \nabla V \cdot \nabla \psi| + \int_{\Omega} |\mathcal{V}_0 \cdot \nabla V \psi| \\ &\leq \|f\|_{0,\Omega} \|\psi\|_{0,\Omega} + \|\lambda\|_{\infty,D} \|\nabla V\|_{0,D} \|\nabla \psi\|_{0,\Omega} + \|\mathcal{V}_0\|_{0,\infty} \|\nabla V\|_{0,D} \|\psi\|_{0,\Omega} \\ &\leq (\|f\|_{0,\Omega} + \|\lambda\|_{\infty,D} \|\nabla V\|_{0,D} + \|\mathcal{V}_0\|_{0,\infty} \|\nabla V\|_{0,D}) \|\psi\|_{1,\Omega} \\ &\leq \delta \|\psi\|_{1,\Omega}. \end{aligned}$$

■

Nous avons le résultat d'existence et d'unicité de la solution du problème (3.4) dans  $H_{\Gamma_d}^1(\Omega)$  muni de la norme  $|\cdot|_{1,\Omega}$

**Théorème 3.2** *Soit  $\Omega \in \Theta_{ad}$  alors, sous les hypothèses  $(H_1) - (H_3)$ , le problème d'état (3.4) admet une solution unique dans  $H_{\Gamma_d}^1(\Omega)$ .*

**Démonstration :**

Comme il a été signalé dans la remarque 3.1, le problème avec les hypothèses  $(H_1) - (H_3)$  est en général non coercif. Pour y remédier, nous utilisons les degrés topologiques de Leray Schauder.

Pour démontrer ce théorème, nous considérons l'application  $\mathcal{G}_t$  pour  $t \in [0, 1]$ , définie de la façon suivante :

$$\begin{aligned} \mathcal{G}_t : H_{\Gamma_d}^1(\Omega) &\mapsto H_{\Gamma_d}^1(\Omega) \\ \bar{u} &\mapsto u \end{aligned}$$

où  $u$  est la solution unique du problème suivant (obtenue grâce au lemme de Lax-Milgram) :

$$\int_{\Omega} \lambda \nabla u \cdot \nabla \psi = t \langle \ell, \psi \rangle_{((H_{\Gamma_d}^1(\Omega))', H_{\Gamma_d}^1(\Omega))} - t \int_{\Omega} \mathcal{V}_0 \cdot \nabla \bar{u} \psi \quad \forall \psi \in H_{\Gamma_d}^1(\Omega). \quad (3.6)$$

Nous montrons que l'opérateur  $\mathcal{G}_t$  est continu et compact, en effet :

Soit  $(\bar{u}_n)_n$  une suite dans  $H_{\Gamma_d}^1(\Omega)$ , telle que

$$\bar{u}_n \longrightarrow \bar{u} \quad \text{dans} \quad H_{\Gamma_d}^1(\Omega).$$

Ceci implique que

$$\nabla \bar{u}_n \longrightarrow \nabla \bar{u} \quad \text{dans} \quad L^2(\Omega).$$

Ce qui montre que

$$\mathcal{G}_t(\bar{u}_n) \longrightarrow \mathcal{G}_t(\bar{u}) \quad \text{dans} \quad H_{\Gamma_d}^1(\Omega).$$

D'autre part, supposons que  $(\bar{u}_n)_n$  est une suite bornée dans  $H_{\Gamma_d}^1(\Omega)$ , donc la suite  $(\nabla \bar{u}_n)_n$  est bornée dans  $L^2(\Omega)$ . En utilisant  $\psi = \mathcal{G}_t(\bar{u}_n) = u_n$  comme fonction test dans le problème (3.6), nous obtenons

$$\begin{aligned} \int_{\Omega} \lambda \nabla u_n \cdot \nabla u_n &= t \langle \ell, u_n \rangle_{((H_{\Gamma_d}^1(\Omega))', H_{\Gamma_d}^1(\Omega))} - t \int_{\Omega} \mathcal{V}_0 \cdot \nabla \bar{u}_n u_n \\ \lambda_0 |u_n|_{1,\Omega}^2 &\leq | \langle \ell, u_n \rangle_{((H_{\Gamma_d}^1(\Omega))', H_{\Gamma_d}^1(\Omega))} | + \int_{\Omega} |u_n| |\mathcal{V}_0 \cdot \nabla \bar{u}_n|. \end{aligned}$$

En utilisant l'inégalité de Poincaré, il existe une constante  $C_0 > 0$  telle que :

$$\begin{aligned} C_0 \|u_n\|_{1,\Omega}^2 &\leq \lambda_0 |u_n|_{1,\Omega}^2 \\ &\leq \|\ell\|_{(H_{\Gamma_d}^1(\Omega))'} \|u_n\|_{1,\Omega} + \|\mathcal{V}_0\|_{\infty,D} \|u_n\|_{0,\Omega} \|\nabla \bar{u}_n\|_{0,\Omega} \\ &\leq (\delta + \|\mathcal{V}_0\|_{\infty,D} \|\nabla \bar{u}_n\|_{0,\Omega}) \|u_n\|_{1,\Omega}. \end{aligned}$$

Nous constatons alors que  $(u_n)_n$  est une suite bornée dans  $H^1(\Omega)$ . Ainsi la suite  $(u_n)_n$  converge faiblement dans  $H^1(\Omega)$ . D'après l'injection compacte de  $H^1(\Omega)$  dans  $L^2(\Omega)$ , nous pouvons extraire une sous suite qui converge dans  $L^2(\Omega)$ . Nous soustrayons alors l'équation (3.6) satisfaite

par  $u_n$  et celle satisfaite par  $u_m$ , puis nous utilisons  $\psi = u_n - u_m$  comme fonction test. Nous obtenons alors :

$$\begin{aligned} C_0 \|u_n - u_m\|_{1,\Omega}^2 &\leq \int_{\Omega} |\mathcal{V}_0| |u_n - u_m| |\nabla \bar{u}_n - \nabla \bar{u}_m| \\ &\leq 2 \|\mathcal{V}_0\|_{\infty,D} \sup_{k \geq 1} \|\bar{u}_k\|_{1,\Omega} \|u_n - u_m\|_{0,\Omega} \end{aligned}$$

Ce qui montre que  $(u_n)_n$  est une suite de Cauchy dans  $H^1(\Omega)$ , et qu'elle converge dans cet espace. Par conséquent  $\forall t \in [0, 1]$ ,  $\mathcal{G}_t$  est un opérateur compact.

D'autre part, dans le paragraphe suivant, nous allons montrer qu'il existe une constante  $C > 0$ , telle que  $\|u\|_{1,\Omega} \leq C$ . Puisque  $\mathcal{G}_t$  est un opérateur continu et compact, pour prouver que  $\mathcal{G}_t$  admet un point fixe, nous considérons la boule ouverte  $B$ , définie par :

$$B = \{u \in H_{\Gamma_d}^1(\Omega), \|u\|_{1,\Omega} \leq R\}$$

avec  $R = C + 1$ . Il est clair que l'opérateur  $\mathcal{G}_t$  n'admet aucun point fixe sur  $\partial B$ , la frontière de  $B$ . Par conséquent la fonction  $\deg[I - \mathcal{G}_t, B, 0]$  est définie et indépendant de  $t$ . De plus elle vérifie les propriétés du théorème 1.8. Puisque  $\mathcal{G}_0$  correspond au problème

$$\int_{\Omega} \lambda \nabla u \cdot \nabla \psi = 0 \quad \forall \psi \in H_{\Gamma_d}^1(\Omega), \quad (3.7)$$

qui bien évidemment admet une solution unique  $u$ , alors

$$\deg[I - \mathcal{G}_0, B, 0] = +1.$$

Ainsi

$$\deg[I - \mathcal{G}_1, B, 0] = +1.$$

Par conséquent, il existe  $u \in H_{\Gamma_d}^1(\Omega)$  tel que  $\mathcal{G}_1(u) = u$  i.e l'opérateur  $\mathcal{G}_1$  admet un point fixe à l'intérieur de  $B$ .

Pour l'unicité de la solution  $u$ , puisque l'opérateur est linéaire, il suffit de prouver que la solution unique de (3.4) avec  $\ell = 0$  est la fonction nulle. En effet, soit  $u$  une solution du problème (3.4) avec  $\ell = 0$ , i.e

$$\int_{\Omega} \lambda \nabla u \cdot \nabla \psi + \int_{\Omega} \mathcal{V}_0 \cdot \nabla u \psi = 0, \quad (3.8)$$

comme  $\frac{u}{|u|} \in L^2(\Omega)$ , il est clair que  $\frac{u}{|u|} \in (H_{\Gamma_d}^1(\Omega))'$  soit  $v$  une solution du problème suivant :

$$\int_{\Omega} \lambda \nabla v \cdot \nabla \psi + \int_{\Omega} \mathcal{V}_0 \cdot \nabla \psi v = \left\langle \frac{u}{|u|}, \psi \right\rangle_{((H_{\Gamma_d}^1(\Omega))', H_{\Gamma_d}^1(\Omega))} \quad (3.9)$$

nous utilisons  $\psi = v$  comme fonction test dans la formulation faible (3.8) et  $\psi = u$  dans la formulation (3.9). C'est ainsi que nous obtenons

$$\int_{\Omega} |u| dx = 0 \quad \text{donc} \quad u = 0.$$

■

### 3.3 Étude de l'existence d'une solution optimale

L'étude de l'existence d'une solution optimale nécessite le choix d'une topologie sur l'ensemble  $\mathcal{F}$ . Dans le paragraphe suivant nous allons définir une topologie adéquate, qui va nous permettre d'établir cette existence.

#### 3.3.1 Topologie sur $\mathcal{F}$

Nous définissons la famille des domaines  $\Theta_{ad}$  de la façon suivante

$$\Theta_{ad} = \{\Omega(\varphi) \mid \varphi \in U_{ad}\}$$

avec

$$U_{ad} = \left\{ \varphi \in C([a, b]) \mid \exists a_{\varphi} \text{ et } b_{\varphi} \in [a, b], \varphi|_{[a, a_{\varphi}]} = 0, \varphi|_{[b_{\varphi}, b]} = 0 \text{ et } \exists L_0 > 0 / \right.$$

$$\left. |\varphi(x) - \varphi(x')| \leq L_0 |x - x'| \forall x, x' \in [a, b], 0 \leq \varphi(x) \leq L_y \forall x \in [a, b] \right\}.$$

où  $L_0$  est une constante donnée, strictement positive.

Une topologie naturelle sur  $\mathcal{F}$  est alors construite à partir de la convergence des domaines de  $\Theta_{ad}$  et de la convergence faible des solutions associées à ces domaines.

##### 3.3.1.1 Topologie sur $\Theta_{ad}$

La topologie que nous considérons sur  $\Theta_{ad}$  est définie par la convergence uniforme des fonctions paramétrisantes des domaines de  $\Theta_{ad}$ .

**Définition 3.1** Soit  $\Omega_n = \Omega(\varphi_n)$  une suite dans  $\Theta_{ad}$  et  $\Omega = \Omega(\varphi)$  un élément de  $\Theta_{ad}$ . On dit que  $\Omega_n$  converge vers  $\Omega$  si et seulement si  $\varphi_n$  converge uniformément vers  $\varphi$  dans l'intervalle

$[a, b]$ . On écrit alors :

$$\Omega_n \xrightarrow{n \rightarrow \infty} \Omega \iff \varphi_n \xrightarrow{n \rightarrow \infty} \varphi \text{ uniformément sur } [a, b]. \quad (3.10)$$

À présent, nous énonçons un résultat de convergence qui nous sera utile pour la suite.

**Proposition 3.2** Soient  $(\Omega_n)_n$  une suite de  $\Theta_{ad}$  et  $\Omega$  un élément de  $\Theta_{ad}$  tels que :

$$\Omega_n \xrightarrow{n \rightarrow \infty} \Omega,$$

alors

$$\chi_{\Omega_n} \xrightarrow{n \rightarrow \infty} \chi_{\Omega}, \text{ faible}^* \text{ dans } L^\infty(D),$$

De plus, nous avons :

$$\lim_{n \rightarrow \infty} \int_D (\chi_{\Omega_n} - \chi_{\Omega})^2 f = 0, \quad \forall f \in L^1(\Omega)$$

Notons que  $\chi_A$  désigne la fonction caractéristique d'un ensemble mesurable  $A$ .

**Démonstration :**

Soit  $(\Omega_n)_n$  une suite de  $\Theta_{ad}$  et  $\Omega$  un élément de  $\Theta_{ad}$  tels que :

$$\Omega_n \xrightarrow{n \rightarrow \infty} \Omega$$

$$\text{i.e. } \varphi_n \xrightarrow{n \rightarrow \infty} \varphi \text{ uniformément dans } [a, b].$$

Ce qui implique que pour tout  $x \in \partial\Omega$ , il existe  $x_n \in \partial\Omega_n$  tel que  $x_n$  converge vers  $x$  dans  $\mathbb{R}^2$ .

Cette assertion est exactement la définition de la convergence ponctuelle des frontières  $\partial\Omega_n$  vers  $\partial\Omega$  (cf. [69]). Donc, d'après le lemme 4 du chapitre 3 de [69], nous avons

$$\chi_{\Omega_n} \xrightarrow{n \rightarrow \infty} \chi_{\Omega} \text{ faible}^* \text{ dans } L^\infty(D)$$

Donc

$$\lim_{n \rightarrow \infty} \int_D (\chi_{\Omega_n} - \chi_{\Omega}) f = 0 \quad \forall f \in L^1(D). \quad (3.11)$$

Or nous avons

$$\begin{aligned} (\chi_{\Omega_n} - \chi_{\Omega})^2 &= \chi_{\Omega_n}^2 + \chi_{\Omega}^2 - 2\chi_{\Omega_n}\chi_{\Omega} \\ &= \chi_{\Omega_n}^2 - \chi_{\Omega}^2 + 2\chi_{\Omega}^2 - 2\chi_{\Omega_n}\chi_{\Omega} \\ &= (\chi_{\Omega_n} - \chi_{\Omega}) + 2\chi_{\Omega}(\chi_{\Omega_n} - \chi_{\Omega}). \end{aligned}$$

Ainsi nous avons :

$$\int_D (\chi_{\Omega_n} - \chi_\Omega)^2 f = \int_D (\chi_{\Omega_n} - \chi_\Omega) f - 2 \int_D (\chi_{\Omega_n} - \chi_\Omega) \chi_\Omega f.$$

En utilisant la relation (3.11) et le fait que  $\chi_\Omega f \in L^1(D)$ , nous obtenons :

$$\lim_{n \rightarrow \infty} \int_D (\chi_{\Omega_n} - \chi_\Omega) \chi_\Omega f = 0.$$

D'où le résultat. ■

### 3.3.1.2 Topologie sur l'espace des solutions

Les domaines de la famille  $\Theta_{ad}$  ont des frontières lipschitziennes (cf. [62]).

D'après la Proposition 1.1, les domaines de  $\Theta_{ad}$  satisfont la propriété du cône. Donc pour toute fonction  $u$  dans  $H_{\Gamma_d}^1(\Omega)$ , il existe  $\tilde{u}$  une extension uniforme de  $u$  dans  $H^1(D)$ . En nous inspirant du Théorème 1.4, nous obtenons le résultat suivant :

**Lemme 3.1** *Il existe une constante  $C$ , telle que :*

$\forall \Omega \in \Theta_{ad}, \forall u \in H_{\Gamma_d}^1(\Omega)$ , il existe  $\tilde{u} \in H^1(D)$  vérifiant :

$$\|\tilde{u}\|_{1,D} \leq C \|u\|_{1,\Omega}, \quad (3.12)$$

avec  $\tilde{u}|_\Omega = u$  presque partout dans  $\Omega$ .

Pour une suite  $(\Omega_n)_n$  de  $\Theta_{ad}$ , nous associons la suite de solution  $u_n = u(\Omega_n)$  de (3.4) sur  $\Omega_n$  pour tout  $n$ .

Nous définissons la convergence de  $u_n$  vers  $u = u(\Omega)$ , comme étant la convergence faible de l'extension uniforme de  $u_n$ , vers l'extension uniforme de  $u$  dans  $H^1(D)$ , et nous écrivons :

$$u_n \xrightarrow[n \rightarrow \infty]{} u \Leftrightarrow \tilde{u}_n \xrightarrow[n \rightarrow \infty]{} \tilde{u} \text{ dans } H^1(D) - \text{faible}. \quad (3.13)$$

**Définition 3.2** *Soit  $\{(\Omega_n, u_n)\}_n$  une suite de  $\mathcal{F}$  et  $(\Omega, u)$  un élément de  $\mathcal{F}$ . Nous définissons alors la convergence de  $(\Omega_n, u_n)$  vers  $(\Omega, u)$  par*

$$(\Omega_n, u_n) \xrightarrow[n \rightarrow \infty]{} (\Omega, u) \Leftrightarrow \begin{cases} \Omega_n \xrightarrow[n \rightarrow \infty]{} \Omega \\ u_n \xrightarrow[n \rightarrow \infty]{} u. \end{cases} \quad (3.14)$$

La topologie que nous allons considérer sur  $\mathcal{F}$  est définie par la convergence (3.14).

### 3.3.1.3 Compacité de $\mathcal{F}$

La compacité de l'ensemble  $\mathcal{F}$  nécessite l'étude de la compacité de  $\Theta_{ad}$  pour la convergence (3.10), ainsi que l'étude de la continuité de l'application qui, à un domaine  $\Omega$  de  $\Theta_{ad}$ , associe  $u(\Omega)$  la solution de (3.4) dans  $H_{\Gamma_d}^1(\Omega)$ .

Nous allons démontrer que  $U_{ad}$  est compact dans  $C([a, b])$ . Ainsi, nous aurons montré, d'une manière indirecte, que  $\Theta_{ad}$  est compact pour la topologie définie par la convergence (3.10).

**Proposition 3.3**  *$U_{ad}$  est compact dans  $C([a, b])$ .*

**Démonstration :**

Soit  $(\varphi_n)_n$  une suite de  $U_{ad}$ . D'après le théorème d'Ascoli-Arzelà (cf. chapitre 1), il existe une sous-suite qu'on notera encore  $(\varphi_n)_n$  et une fonction  $\varphi$  continue dans  $[a, b]$  telles que

$$\varphi_n \xrightarrow{n \rightarrow \infty} \varphi \text{ uniformément dans } [a, b].$$

Il reste alors à prouver que  $U_{ad}$  est fermé dans  $C([a, b])$ , ce qui fait l'objet du lemme 3.2.

■

**Lemme 3.2**  *$U_{ad}$  est fermé dans  $C([a, b])$ .*

**Démonstration :**

Soient  $(\varphi_n)_n$  une suite de  $U_{ad}$  et  $\varphi$  une fonction dans  $[a, b]$  telles que

$$\varphi_n \xrightarrow{n \rightarrow \infty} \varphi \text{ uniformément dans } [a, b].$$

$$\text{i.e.} \quad \sup_{x \in [a, b]} |\varphi_n(x) - \varphi(x)| \xrightarrow{n \rightarrow \infty} 0.$$

Montrons que  $\varphi \in U_{ad}$ .

Soient  $x$  et  $x'$  deux éléments de  $[a, b]$ . Puisque  $\varphi_n \in U_{ad}$ , alors nous avons :

$$|\varphi_n(x) - \varphi_n(x')| \leq C_0 |x - x'|.$$

En utilisant l'inégalité triangulaire, nous obtenons :

$$\begin{aligned} |\varphi(x) - \varphi(x')| &\leq |\varphi(x) - \varphi_n(x)| + |\varphi_n(x) - \varphi_n(x')| + |\varphi_n(x') - \varphi(x')| \\ &\leq |\varphi(x) - \varphi_n(x)| + C_0|x - x'| + |\varphi_n(x') - \varphi(x')|. \end{aligned}$$

En faisant tendre  $n$  vers  $\infty$ , nous obtenons :

$$|\varphi(x) - \varphi(x')| \leq C_0|x - x'|.$$

Comme  $\varphi_n \in U_{ad}$ , alors il existe  $a_n, b_n \in [a, b]$  tels que

$$\varphi_n|_{[a, a_n]} = 0 \text{ et } \varphi_n|_{[b_n, b]} = 0.$$

Il existe deux sous suites  $(a_{n_k})_k$  et  $(b_{n_k})_k$  telles que

$$a_{n_k} \xrightarrow{k \rightarrow \infty} a^* \text{ et } b_{n_k} \xrightarrow{k \rightarrow \infty} b^*.$$

Montrons que

$$\varphi|_{[a, a^*]} = 0 \text{ et } \varphi|_{[b^*, b]} = 0.$$

En effet, nous avons :

$$\begin{aligned} \int_a^{a^*} |\varphi(x)| dx &= \int_a^b \chi_{[a, a^*]} |\varphi(x)| dx - \int_a^b \chi_{[a, a^*]} |\varphi_n(x)| dx + \int_a^b \chi_{[a, a^*]} |\varphi_n(x)| dx \\ &\quad - \int_a^b \chi_{[a, a_n]} |\varphi_n(x)| dx \\ &\leq \int_a^b \chi_{[a, a^*]} |\varphi(x) - \varphi_n(x)| dx + \int_a^b |\varphi_n(x)| |\chi_{[a, a^*]} - \chi_{[a, a_n]}| dx \\ &\leq \int_a^b \chi_{[a, a^*]} |\varphi(x) - \varphi_n(x)| dx + \sup_{x \in [a, b]} |\varphi_n(x)| \int_a^b |\chi_{[a, a^*]} - \chi_{[a, a_n]}| dx. \end{aligned}$$

En faisant tendre  $n$  vers  $\infty$  nous aurons :

$$\int_a^{a^*} |\varphi(x)| dx = 0.$$

Ainsi on a  $\varphi = 0$  sur  $[a, a^*]$ , de même on montre que  $\varphi = 0$  sur  $[b^*, b]$ .

Donc  $\varphi \in U_{ad}$  et, par conséquent,  $U_{ad}$  est fermé dans  $C([a, b])$ .

■

### 3.3.2 Continuité du problème d'état

Soit  $(\Omega_n)_n$  une suite dans  $\Theta_{ad}$ , et  $\Omega$  un élément de  $\Theta_{ad}$ . D'après la Proposition 3.3, nous pouvons extraire une sous-suite notée encore  $(\Omega_n)_n$  telle que

$$\Omega_n \xrightarrow{n \rightarrow \infty} \Omega.$$

Nous définissons l'espace  $H_{\Gamma_d}^1(\Omega_n)$  par

$$H_{\Gamma_d}^1(\Omega_n) = \{v \in H^1(\Omega_n), v|_{\Gamma_4 \cup \Gamma_n} = 0\}.$$

Soit  $u_n = u(\Omega_n)$  la solution du problème suivant

$$(P_n) \begin{cases} \text{trouver } u_n \in H_{\Gamma_d}^1(\Omega_n) \\ \int_{\Omega_n} \lambda \nabla u_n \cdot \nabla \psi + \int_{\Omega_n} \mathcal{V}_0 \cdot \nabla u_n \psi = \langle \ell, \psi \rangle_{((H_{\Gamma_d}^1(\Omega_n))', H_{\Gamma_d}^1(\Omega_n))} \quad \forall \psi \in H_{\Gamma_d}^1(\Omega_n), \end{cases} \quad (3.15)$$

Dans ce qui suit, nous allons montrer qu'il existe  $\tilde{u}_n$ , un prolongement uniforme de  $u_n$  dans  $H^1(D)$ , qui converge faiblement dans  $H^1(D)$  vers une limite  $W$ , tel que  $W|_{\Omega} = u$  est solution du problème (3.4). Pour cela, nous aurons besoin du résultat de prolongement uniforme suivant

**Théorème 3.3** *Il existe une constante  $C > 0$  indépendante de  $n$ , et  $\tilde{u}_n$  un prolongement de  $u_n$  dans  $H^1(D)$  tels que*

$$\|\tilde{u}_n\|_{1,D} \leq C. \quad (3.16)$$

**Démonstration :**

Soit  $u_n$  la solution du problème (3.15) dans  $\Omega_n$ . D'après le Lemme 3.1, il existe un prolongement  $\tilde{u}_n$  de  $u_n$  dans  $H^1(D)$  et une constante  $C > 0$  indépendante de  $n$  tels que

$$\|\tilde{u}_n\|_{1,D} \leq C \|u_n\|_{1,\Omega_n}. \quad (3.17)$$

Il suffit alors de prouver que  $\|u_n\|_{1,\Omega_n}$  est uniformément borné par rapport à  $\Omega_n$ . En fait, c'est l'une des difficultés principales de ce travail. Pour obtenir une estimation uniforme de la solution du problème (3.4), nous définissons le problème intermédiaire suivant

$$\begin{cases} \text{Trouver } p \in H_{\Gamma_d}^1(\Omega) \\ \int_{\Omega} \lambda \nabla p \cdot \nabla \psi + \int_{\Omega} p \mathcal{V}_0 \cdot \nabla \psi = \langle \ell, \psi \rangle_{((H_{\Gamma_d}^1(\Omega))', H_{\Gamma_d}^1(\Omega))} \quad \forall \psi \in H_{\Gamma_d}^1(\Omega) \end{cases} \quad (3.18)$$

Afin d'obtenir une estimation uniforme de  $\|u\|_{1,\Omega}$ , nous allons tout d'abord commencer par montrer que  $\|p\|_{1,\Omega}$  est uniformément borné par rapport à  $\Omega$ .

• **Estimation uniforme de  $\|p\|_{1,\Omega}$**

En utilisant  $\psi = p$  comme fonction test dans (3.18), puis en appliquant le corollaire 1.1, nous obtenons :

$$\begin{aligned} C_0 \|p\|_{H^1(\Omega)}^2 &\leq \lambda_0 \int_{\Omega} \nabla p \cdot \nabla p \\ &\leq \int_{\Omega} \lambda \nabla p \cdot \nabla p \\ &\leq |\langle \ell, p \rangle_{((H_{\Gamma_d}^1(\Omega))', H_{\Gamma_d}^1(\Omega))}| + \int_{\Omega} |\mathcal{V}_0| |\nabla p| |p|, \end{aligned}$$

On considère l'ensemble  $A_k$  défini par

$$A_k = \{(x, y) \in \Omega, |p(x, y)| > k\}.$$

En utilisant le fait que

$$\Omega = A_k \cup A_k^c,$$

Puis, en appliquant l'inégalité de Hölder et la Proposition 3.1, nous obtenons :

$$C_0 \|p\|_{1,\Omega}^2 \leq \delta \|p\|_{1,\Omega} + \int_{A_k} |\mathcal{V}_0| |p| |\nabla p| + \int_{A_k^c} |\mathcal{V}_0| |p| |\nabla p|.$$

En appliquant à nouveau l'inégalité de Hölder deux fois, et en tenant compte du fait que

$$|p| \leq k \quad \text{sur} \quad A_k^c$$

nous obtenons

$$\begin{aligned} C_0 \|p\|_{1,\Omega}^2 &\leq \delta \|p\|_{1,\Omega} + k \|\mathcal{V}_0\|_{\infty,D} \int_{A_k} |\nabla p| + \|\mathcal{V}_0\|_{\infty,D} |A_k|^{\frac{1}{4}} \left( \int_{A_k} |p|^4 \right)^{\frac{1}{4}} \|\nabla p\|_{0,\Omega} \\ &\leq \delta \|p\|_{1,\Omega} + k \|\mathcal{V}_0\|_{\infty,D} |\Omega|^{\frac{1}{2}} \|\nabla p\|_{0,\Omega} + \|\mathcal{V}_0\|_{\infty,D} |A_k|^{\frac{1}{4}} \|p\|_{L^4(\Omega)} \|\nabla p\|_{0,\Omega}. \end{aligned}$$

D'après le Théorème 1.5, il existe une constante  $C_1 > 0$  indépendante de  $\Omega$  telle que

$$\|p\|_{L^4(\Omega)} \leq C_1 |\Omega|^{\frac{1}{4}} \|\nabla p\|_{0,\Omega}.$$

Par suite, nous avons :

$$\begin{aligned} C_0 \|p\|_{1,\Omega}^2 &\leq \delta \|p\|_{1,\Omega} + k \|\mathcal{V}_0\|_{\infty,D} |D|^{\frac{1}{2}} \|\nabla p\|_{0,\Omega} + \|\mathcal{V}_0\|_{\infty,D} |A_k|^{\frac{1}{4}} |D|^{\frac{1}{4}} \|\nabla p\|_{0,\Omega}^2 \\ &\leq (\delta + k \|\mathcal{V}_0\|_{\infty,D} |D|^{\frac{1}{2}}) \|p\|_{1,\Omega} + C_1 \|\mathcal{V}_0\|_{\infty,D} |A_k|^{\frac{1}{4}} |D|^{\frac{1}{4}} \|p\|_{1,\Omega}^2 \end{aligned} \quad (3.19)$$

Afin d'éliminer le terme quadratique de  $\|p\|_{1,\Omega}$ , nous avons besoin de prouver une estimation uniforme de  $\|\ln(1 + |p|)\|_{0,\Omega}$  i.e nous devons montrer qu'il existe une constante  $C_2 > 0$  indépendante de  $\Omega$  telle que

$$\|\ln(1 + |p|)\|_{0,\Omega} \leq C_2. \quad (3.20)$$

Grâce à cette estimation uniforme, qui sera démontrée par la suite dans la Proposition 3.4, et en utilisant l'inégalité de Markov, nous obtiendrons le contrôle uniforme de la mesure de Lebesgue de  $A_k$ , i.e. nous avons

$$\begin{aligned} |A_k| &= |\{(x, y) \in \Omega / \ln(1 + |p|)^2 \geq \ln(1 + k)^2\}| \\ &\leq \frac{1}{\ln(1+k)^2} \|\ln(1 + |p|)\|_{0,\Omega}^2 \\ &\leq \frac{C_2}{\ln(1+k)^2} \end{aligned}$$

Alors il existe  $k_0$  tel que,  $\forall k \geq k_0$ , nous avons :

$$C_1 \|\mathcal{V}_0\|_{\infty,D} |A_k|^{\frac{1}{4}} |D|^{\frac{1}{4}} \leq \frac{C_0}{2} \quad (3.21)$$

En conclusion, l'inégalité (3.19) devient

$$\forall k \geq k_0$$

$$\frac{C_0}{2} \|p\|_{1,\Omega}^2 \leq (\delta + k \|\mathcal{V}_0\|_{\infty,D} |D|^{\frac{1}{2}}) \|p\|_{1,\Omega} \quad (3.22)$$

$$\|p\|_{1,\Omega} \leq \frac{2}{C_0} (\delta + k \|\mathcal{V}_0\|_{\infty,D} |D|^{\frac{1}{2}})$$

$$\|p\|_{1,\Omega} \leq C_3.$$

• **Estimation uniforme de  $\|u\|_{1,\Omega}$**

Nous considérons les deux formulations variationnelles (3.4) et (3.18). D'après l'inégalité (3.22), il existe une constante  $C_3 > 0$  indépendante de  $\Omega$  telle que

$$\forall \theta \in ((H_{\Gamma_d}^1(\Omega))'), \text{ satisfaisant } \|\theta\|_{((H_{\Gamma_d}^1(\Omega))')} \leq 1,$$

nous avons l'existence de  $p$  solution du problème (3.18) pour  $\ell = \theta$  qui vérifie l'inégalité suivante :

$$\|p\|_{1,\Omega} \leq C_3.$$

Nous prenons  $\psi = p$  comme fonction test dans le problème variationnel (3.4) et  $\psi = u$  dans le problème variationnel (3.18). Nous obtenons alors

$$\begin{aligned} \langle \theta, u \rangle_{((H_{\Gamma_d}^1(\Omega))', H_{\Gamma_d}^1(\Omega))} &= \langle \ell, p \rangle_{((H_{\Gamma_d}^1(\Omega))', H_{\Gamma_d}^1(\Omega))} \\ &\leq \delta C_3. \end{aligned} \quad (3.23)$$

Or nous savons que

$$\|u\|_{1,\Omega} = \sup_{\|\theta\|_{((H_{\Gamma_d}^1(\Omega))')} \leq 1} \frac{\langle \theta, u \rangle_{((H_{\Gamma_d}^1(\Omega))', H_{\Gamma_d}^1(\Omega))}}{\|\theta\|_{((H_{\Gamma_d}^1(\Omega))')}}.$$

Finalement, nous obtenons :

$$\|u\|_{1,\Omega} \leq \delta C_3. \quad (3.24)$$

Ce qui achève la démonstration. ■

Il reste à prouver l'estimation uniforme (3.20), qui va être obtenue par la proposition suivante

**Proposition 3.4** *Sous les hypothèses  $(H_1) - (H_3)$ , nous avons :  $\forall p \in H_{\Gamma_d}^1(\Omega)$  solution du problème (3.18) dans  $\Omega$ , il existe une constante  $C_2 > 0$  indépendante de  $\Omega \in \Theta_{ad}$  telle que :*

$$\|\ln(1 + |p|)\|_{0,\Omega} \leq C_2.$$

**Démonstration :**

Nous définissons la fonction à valeurs réelles

$$h(s) = \int_0^s \frac{dt}{(1 + |t|)^2}.$$

Comme  $h$  est lipschitzienne, alors d'après le Lemme 1.3, pour tout  $p \in H_{\Gamma_d}^1(\Omega)$ , nous avons  $h(p) \in H_{\Gamma_d}^1(\Omega)$ . Par suite, nous prenons  $h(p)$  comme fonction test dans la formulation faible (3.18).

Nous obtenons alors

$$\begin{aligned} \int_{\Omega} \lambda \frac{|\nabla p|^2}{(1 + |p|)^2} &= \langle \ell, h(p) \rangle_{((H_{\Gamma_d}^1(\Omega))', H_{\Gamma_d}^1(\Omega))} - \int_{\Omega} p \mathcal{V}_0 \cdot \nabla h(p) \\ &\leq \delta \|h(p)\|_{1,\Omega} + \int_{\Omega} |p| \frac{|\mathcal{V}_0 \cdot \nabla p|}{(1 + |p|)^2}, \end{aligned}$$

Ensuite, nous remarquons que

$$\frac{|p|}{1 + |p|} \leq 1 \quad \text{et} \quad |h(p)| < 1.$$

Ainsi, nous obtenons :

$$\|h(p)\|_{1,\Omega} \leq |\Omega|^{\frac{1}{2}} + \|\nabla h(p)\|_{0,\Omega}.$$

Par conséquent,

$$\begin{aligned} \lambda_0 \int_{\Omega} \frac{|\nabla p|^2}{(1 + |p|)^2} &\leq \int_{\Omega} \lambda \frac{|\nabla p|^2}{(1 + |p|)^2} \\ &\leq \delta (|D|^{\frac{1}{2}} + \left( \int_{\Omega} \frac{|\nabla p|^2}{(1 + |p|)^2} \right)^{\frac{1}{2}}) + \|\mathcal{V}_0\|_{\infty,D} \int_{\Omega} \frac{|\nabla p|}{(1 + |p|)}. \end{aligned}$$

En utilisant l'inégalité de Young

$$(\sqrt{\frac{\lambda_0}{2}} a \frac{b}{\sqrt{\frac{\lambda_0}{2}}} \leq \frac{\lambda_0}{4} a^2 + \frac{b^2}{\lambda_0})$$

nous obtenons :

$$\lambda_0 \int_{\Omega} \frac{|\nabla p|^2}{(1+|p|)^2} \leq \delta |D|^{\frac{1}{2}} + \frac{\delta^2}{\lambda_0} + \frac{\lambda_0}{4} \int_{\Omega} \frac{|\nabla p|^2}{(1+|p|)^2} + \|\mathcal{V}_0\|_{\infty,D}^2 \frac{|D|}{\lambda_0} + \frac{\lambda_0}{4} \int_{\Omega} \frac{|\nabla p|^2}{(1+|p|)^2},$$

$$\frac{\lambda_0}{2} \int_{\Omega} \frac{|\nabla p|^2}{(1+|p|)^2} \leq \delta |D|^{\frac{1}{2}} + \frac{\delta^2}{\lambda_0} + \|\mathcal{V}_0\|_{\infty,D}^2 \frac{|D|}{\lambda_0}.$$

Finalement il résulte que

$$\begin{aligned} \int_{\Omega} |\nabla \ln(1+|p|)|^2 &= \int_{\Omega} \frac{|\nabla p|^2}{(1+|p|)^2} \\ &\leq \frac{2\delta |D|^{\frac{1}{2}}}{\lambda_0} + 2\delta^2 + 2\|\mathcal{V}_0\|_{\infty,D}^2 |D| = C_4. \end{aligned} \tag{3.25}$$

Notons que la constante  $C_4$  est indépendante de  $\Omega \in \Theta_{ad}$ .

Nous concluons cette démonstration en utilisant l'inégalité (3.25), plus le fait que

$$\forall p \in H_{\Gamma_d}^1(\Omega) \quad \ln(1+|p|) \in H_{\Gamma_d}^1(\Omega),$$

nous obtenons à l'aide du corollaire 1.1

$$\begin{aligned} \|\ln(1+|p|)\|_{0,\Omega}^2 &\leq \|\ln(1+|p|)\|_{1,\Omega}^2 \\ &\leq \frac{1}{C_0} \int_{\Omega} |\nabla \ln(1+|p|)|^2 \\ &\leq \frac{C_4}{C_0} = C_2. \end{aligned} \tag{3.26}$$

■

En nous basant sur le théorème 3.3, nous démontrons le résultat de continuité suivant

**Théorème 3.4** *Soit  $(u_n)_n$  une suite de  $H_{\Gamma_d}^1(\Omega_n)$ . Il existe une sous suite de  $(\tilde{u}_n)_n$  qui converge faiblement vers  $\tilde{u}$  dans  $H^1(D)$ . De plus,  $\tilde{u}|_{\Omega}$  est solution de l'équation variationnelle (3.4) dans  $\Omega$ .*

**Démonstration :**

D'après le Théorème 3.3, nous savons que la suite  $(\tilde{u}_n)_n$  est uniformément bornée dans  $H^1(D)$ . En utilisant la faible compacité de la boule unité et la réflexivité de  $H^1(D)$ , nous pouvons extraire une sous suite notée encore  $(\tilde{u}_n)_n$  qui converge faiblement vers une limite que nous noterons  $\tilde{u}$ . Il suffit alors de montrer que  $u = \tilde{u}|_{\Omega}$  est solution de (3.4). Pour cela, nous allons démontrer les deux assertions suivantes

$$u = \tilde{u}|_{\Omega} \in H_{\Gamma_d}^1(\Omega),$$

et

$$\int_{\Omega} \lambda \nabla u \cdot \nabla \psi + \int_{\Omega} u \mathcal{V}_0 \cdot \nabla \psi = \langle \ell, \psi \rangle_{(H_{\Gamma_d}^1(\Omega))', H_{\Gamma_d}^1(\Omega)} \quad \forall \psi \in H_{\Gamma_d}^1(\Omega).$$

En effet, nous avons  $u = \tilde{u}|_{\Omega} \in H^1(\Omega)$ . Il suffit alors de prouver que  $u$  satisfait la condition de Dirichlet sur  $\Gamma \cup \Gamma_4$ .

En utilisant le fait que  $\tilde{u}_n$  converge faiblement vers  $\tilde{u}$  dans  $H^1(D)$ , ainsi que l'injection compacte de  $H^1(D)$  dans  $L^2(\Gamma_4)$ , nous obtenons

$$\tilde{u}_n|_{\Gamma_4} \xrightarrow{n \rightarrow \infty} \tilde{u}|_{\Gamma_4} \text{ dans } L^2(\Gamma_4) \text{ forte,}$$

Il en résulte que :

$$\tilde{u}|_{\Gamma_4} = 0.$$

Il reste alors à prouver que  $\tilde{u} = 0$  sur  $\Gamma$ , ce qui revient à montrer que

$$\int_{\Gamma} |\tilde{u}|^2 d\sigma = 0.$$

En effet, d'après la définition de  $U_{ad}$ , nous avons

$$\forall n \in \mathbb{N} \quad |\varphi'_n(x)| \leq L_0 \text{ p.p. dans } [a, b].$$

Par suite,

$$\begin{aligned} 0 \leq \int_{a_n}^{b_n} |\tilde{u}_n(x, \varphi_n(x))|^2 dx &\leq \int_{\Gamma_n} |\tilde{u}_n(x, y)|^2 d\sigma = \int_{a_n}^{b_n} |\tilde{u}_n(x, \varphi_n(x))|^2 \sqrt{1 + \varphi_n'(x)^2} dx \\ &\leq C \int_{a_n}^{b_n} |\tilde{u}_n(x, \varphi_n(x))|^2, \end{aligned}$$

où  $C = \sqrt{1 + L_0^2}$ . Il est clair que

$$u_n = 0 \quad \text{sur} \quad \Gamma_n \quad \Leftrightarrow \quad u_n(\cdot, \varphi_n(\cdot)) = 0 \quad \text{sur} \quad [a_n, b_n].$$

Puisque

$$\forall n \in \mathbb{N} \quad \text{on a} \quad \int_{a_n}^{b_n} |\tilde{u}_n(x, \varphi_n(x))|^2 = 0, \quad (3.27)$$

il suffit alors de montrer que

$$\int_{a_n}^{b_n} |\tilde{u}_n(x, \varphi_n(x))|^2 \xrightarrow{n \rightarrow \infty} \int_{a^*}^{b^*} |\tilde{u}(x, \varphi^*(x))|^2.$$

En effet, nous avons

$$\int_{a_n}^{b_n} |\tilde{u}_n(x, \varphi_n(x))|^2 - \int_{a^*}^{b^*} |\tilde{u}(x, \varphi^*(x))|^2 = I_1 + I_2,$$

avec

$$I_1 = \int_{a_n}^{b_n} |\tilde{u}_n(x, \varphi_n(x))|^2 - \int_{a_n}^{b_n} |\tilde{u}(x, \varphi(x))|^2.$$

$$I_2 = \int_{a_n}^{b_n} |\tilde{u}(x, \varphi(x))|^2 - \int_{a^*}^{b^*} |\tilde{u}(x, \varphi(x))|^2.$$

En appliquant l'inégalité de Hölder à  $I_1$ , nous avons

$$\begin{aligned} I_1 &= \int_a^b \chi_{[a_n, b_n]} (|\tilde{u}_n(x, \varphi_n(x))|^2 - |\tilde{u}(x, \varphi(x))|^2) \\ &\leq \left( \int_a^b \chi_{[a_n, b_n]} |\tilde{u}_n(x, \varphi_n(x)) + \tilde{u}(x, \varphi(x))|^2 \right)^{\frac{1}{2}} \left( \int_a^b \chi_{[a_n, b_n]} |\tilde{u}_n(x, \varphi_n(x)) - \tilde{u}(x, \varphi(x))|^2 \right)^{\frac{1}{2}} \end{aligned}$$

Il existe une constante  $C > 0$  indépendante de  $n$  telle que

$$\begin{aligned}
\int_a^b \chi_{[a_n, b_n]} |\tilde{u}_n(x, \varphi_n(x)) + \tilde{u}(x, \varphi(x))|^2 &\leq 2 \left( \int_a^b \chi_{[a_n, b_n]} |\tilde{u}_n(x, \varphi_n(x))|^2 + \int_a^b |\tilde{u}(x, \varphi(x))|^2 \right) \\
&\leq 2 \int_a^b |\tilde{u}(x, \varphi(x))|^2 \\
&\leq C.
\end{aligned}$$

D'autre part, nous avons

$$\int_a^b \chi_{[a_n, b_n]} |\tilde{u}_n(x, \varphi_n(x)) - \tilde{u}(x, \varphi(x))|^2 \leq 2(I_{11} + I_{12})$$

où

$$\begin{aligned}
I_{11} &= \int_a^b |\tilde{u}_n(x, \varphi_n(x)) - \tilde{u}_n(x, \varphi(x))|^2 \\
I_{12} &= \int_a^b |\tilde{u}_n(x, \varphi(x)) - \tilde{u}(x, \varphi(x))|^2
\end{aligned}$$

En appliquant l'inégalité de Hölder à  $I_{11}$  nous aurons

$$\begin{aligned}
I_{11} &\leq \int_a^b \left| \int_{\varphi(x)}^{\varphi_n(x)} \frac{\partial \tilde{u}_n(x, y)}{\partial y} dy \right|^2 dx \\
&\leq \max_{x \in [a, b]} |\varphi_n(x) - \varphi(x)| \|\tilde{u}_n\|_{1, D}^2.
\end{aligned}$$

En utilisant la convergence uniforme de  $\varphi_n$  vers  $\varphi$ , et le fait que  $\tilde{u}_n$  est uniformément borné dans  $H^1(D)$ , nous avons

$$\lim_{n \rightarrow \infty} I_{11} = 0.$$

D'après l'injection compacte de  $H^1(D)$  dans  $L^2(\Gamma)$ , nous avons

$$\lim_{n \rightarrow \infty} I_{12} = 0.$$

Par suite, nous obtenons

$$\lim_{n \rightarrow \infty} I_1 = 0.$$

D'autre part, nous avons

$$I_2 = \int_a^b (\chi_{[a_n, b_n]} - \chi_{[a^*, b^*]}) |\tilde{u}(x, \varphi(x))|^2.$$

D'après le théorème de convergence dominée de Lebesgue, nous avons

$$\lim_{n \rightarrow \infty} I_2 = 0.$$

Soit maintenant  $\psi \in H_{\Gamma_d}^1(\Omega)$ . Nous notons par  $\tilde{\psi} \in H^1(D)$  une extension de  $\psi$  définie par

$$\tilde{\psi} = \begin{cases} \psi & \text{dans } \Omega \\ 0 & \text{dans } D \setminus \Omega. \end{cases}$$

Ainsi, nous pouvons construire une suite  $(\psi_j)_j$ ,  $\psi_j \in \mathcal{D}(\bar{D})$ , telle que,

$$\text{dist}(\text{supp } \psi_j, \overline{\Gamma_d}) > 0 \quad \forall j \in \mathbb{N} \quad \text{et} \quad \psi_j \rightarrow \tilde{\psi} \quad \text{dans } H^1(D). \quad (3.28)$$

Soit  $j \in \mathbb{N}$ . Puisque

$$\Omega_n \xrightarrow{n \rightarrow \infty} \Omega,$$

il existe  $n_0$  tel que

$$\psi_j|_{\Omega_n} \in H_{\Gamma_d}^1(\Omega_n), \quad \forall n \geq n_0.$$

Pour tout  $n \geq n_0$ , nous avons

$$\int_{\Omega_n} \lambda \nabla u_n \cdot \nabla \psi_j + \int_{\Omega_n} \mathcal{V}_0 \cdot \nabla u_n \psi_j = \int_{\Omega_n} f \psi_j - \int_{\Omega_n} \lambda \nabla V \cdot \nabla \psi_j - \int_{\Omega_n} \psi_j \mathcal{V}_0 \cdot \nabla V. \quad (3.29)$$

En faisant tendre dans l'équation (3.29) d'abord  $n$ , puis  $j$  vers l'infini, il en résulte que  $u$  est solution de l'équation suivante :

$$\int_{\Omega} \lambda \nabla u \cdot \nabla \psi + \int_{\Omega} \mathcal{V}_0 \cdot \nabla u \psi = \int_{\Omega} f \psi - \int_{\Omega} \lambda \nabla V \cdot \nabla \psi - \int_{\Omega} \psi \mathcal{V}_0 \cdot \nabla V. \quad (3.30)$$

En effet, nous définissons  $I_1, I_2, I_3, I_4$  et  $I_5$  par

$$\begin{aligned}
 I_1 &= \int_{\Omega_n} \lambda \nabla u_n \cdot \nabla \psi_j - \int_{\Omega} \lambda \nabla u \cdot \nabla \psi \\
 I_2 &= \int_{\Omega_n} \mathcal{V}_0 \cdot \nabla u_n \psi_j - \int_{\Omega} \mathcal{V}_0 \cdot \nabla u \psi \\
 I_3 &= \int_{\Omega_n} f \psi_j - \int_{\Omega} f \psi \\
 I_4 &= \int_{\Omega_n} \lambda \nabla V \cdot \nabla \psi_j - \int_{\Omega} \lambda \nabla V \cdot \nabla \psi \\
 I_5 &= \int_{\Omega_n} \mathcal{V}_0 \cdot \nabla V \psi_j - \int_{\Omega} \mathcal{V}_0 \cdot \nabla V \psi.
 \end{aligned} \tag{3.31}$$

Il suffit alors de montrer que

$$\lim_{j \rightarrow \infty} \lim_{n \rightarrow \infty} I_i = 0, \quad i = 1, \dots, 5.$$

Nous commençons par  $I_1$ , qui peut s'écrire

$$\begin{aligned}
 I_1 &= \int_D \lambda (\chi_{\Omega_n} - \chi_{\Omega}) \nabla \tilde{u}_n \cdot \nabla \psi_j + \int_D \lambda \chi_{\Omega} (\nabla \tilde{u}_n - \nabla \tilde{u}) \cdot \nabla \psi_j + \\
 &\quad \int_D \lambda \chi_{\Omega} \nabla \tilde{u} \cdot (\nabla \psi_j - \nabla \tilde{\psi}).
 \end{aligned}$$

Donc

$$\begin{aligned}
 |I_1| &\leq \left| \int_D \lambda (\chi_{\Omega_n} - \chi_{\Omega}) \nabla \tilde{u}_n \cdot \nabla \psi_j \right| + \left| \int_D \lambda \chi_{\Omega} (\nabla \tilde{u}_n - \nabla \tilde{u}) \cdot \nabla \psi_j \right| + \\
 &\quad \left| \int_D \lambda \chi_{\Omega} \nabla \tilde{u} \cdot (\nabla \psi_j - \nabla \tilde{\psi}) \right|.
 \end{aligned}$$

Par suite, en appliquant l'inégalité de Hölder au premier terme, nous avons :

$$\begin{aligned}
\left| \int_D \lambda(\chi_{\Omega_n} - \chi_{\Omega}) \nabla \tilde{u}_n \cdot \nabla \psi_j \right| &\leq \int_D |\lambda| |\chi_{\Omega_n} - \chi_{\Omega}| |\nabla \tilde{u}_n| |\nabla \psi_j| \\
&\leq \|\lambda\|_{\infty, D} \left( \int_D |\nabla \tilde{u}_n|^2 \right)^{\frac{1}{2}} \left( \int_D (\chi_{\Omega_n} - \chi_{\Omega})^2 |\nabla \psi_j|^2 \right)^{\frac{1}{2}} \\
&\leq \|\lambda\|_{\infty, D} \|\tilde{u}_n\|_{1, D} \left( \int_D (\chi_{\Omega_n} - \chi_{\Omega})^2 |\nabla \psi_j|^2 \right)^{\frac{1}{2}}.
\end{aligned}$$

D'après la Proposition 3.2 et en utilisant le Théorème 3.3, ainsi que fait que  $\psi_j$  est borné dans  $H^1(D)$ , nous obtenons :

$$\forall j \in \mathbb{N}, \quad \lim_{n \rightarrow \infty} \int_D \lambda(\chi_{\Omega_n} - \chi_{\Omega}) \nabla \tilde{u}_n \cdot \nabla \psi_j = 0. \quad (3.32)$$

D'autre part, puisque nous avons la convergence suivante :

$$\tilde{u}_n \xrightarrow[n \rightarrow \infty]{} \tilde{u} \text{ dans } H^1(D) \text{ faible},$$

en utilisant la linéarité et la continuité de l'application gradient de  $H^1(D)$  dans  $L^2(D)$ , nous avons aussi la convergence :

$$\nabla \tilde{u}_n \xrightarrow[n \rightarrow \infty]{} \nabla \tilde{u} \text{ dans } L^2(D) \text{ faible},$$

Et comme  $\forall j \in \mathbb{N}, \lambda \chi_{\Omega} \nabla \psi_j \in L^2(D)$ , nous avons

$$\forall j \in \mathbb{N}, \quad \lim_{n \rightarrow \infty} \int_D \lambda \chi_{\Omega} (\nabla \tilde{u}_n - \nabla \tilde{u}) \cdot \nabla \psi_j = 0. \quad (3.33)$$

D'après la convergence de  $\psi_j$  vers  $\tilde{\psi}$  dans  $H^1(D)$  (3.28), et en utilisant la linéarité et la continuité de l'application gradient de  $H^1(D)$  dans  $L^2(D)$ , ainsi que fait que  $\lambda \chi_{\Omega} \nabla \tilde{u} \in L^2(D)$ , nous obtenons

$$\lim_{j \rightarrow \infty} \int_D \lambda \chi_{\Omega} \nabla \tilde{u} \cdot (\nabla \psi_j - \nabla \tilde{\psi}) = 0. \quad (3.34)$$

Par conséquent, (3.32), (3.33) et (3.34) impliquent que

$$\lim_{j \rightarrow \infty} \lim_{n \rightarrow \infty} I_1 = 0.$$

En utilisant des techniques similaires à celles utilisées pour  $I_1$ , nous obtenons :

$$\lim_{j \rightarrow \infty} \lim_{n \rightarrow \infty} I_2 = 0.$$

Pour  $I_3$ , nous avons

$$I_3 = \int_D \chi_{\Omega_n} \lambda f(\psi_j - \psi) + \int_D \lambda f \psi (\chi_{\Omega_n} - \chi_{\Omega}).$$

Donc

$$|I_3| \leq \left| \int_D \chi_{\Omega_n} \lambda f(\psi_j - \psi) \right| + \left| \int_D \lambda f \psi (\chi_{\Omega_n} - \chi_{\Omega}) \right|.$$

En appliquant l'inégalité de Hölder au deuxième terme, nous obtenons :

$$\begin{aligned} \left| \int_D \lambda f \psi (\chi_{\Omega_n} - \chi_{\Omega}) \right| &\leq \int_D \chi_{\Omega_n} |\lambda| |f| |\psi| |\chi_{\Omega_n} - \chi_{\Omega}| \\ &\leq \|\lambda\|_{\infty, D} \left( \int_D |f|^2 \right)^{\frac{1}{2}} \left( \int_D |\psi|^2 (\chi_{\Omega_n} - \chi_{\Omega})^2 \right)^{\frac{1}{2}} \\ &\leq \|\lambda\|_{\infty, D} \|f\|_{0, D} \left( \int_D |\psi|^2 (\chi_{\Omega_n} - \chi_{\Omega})^2 \right)^{\frac{1}{2}} \end{aligned}$$

D'après la Proposition 3.2, nous avons

$$\lim_{j \rightarrow \infty} \int_D |\psi|^2 (\chi_{\Omega_n} - \chi_{\Omega})^2 = 0. \quad (3.35)$$

D'autre part, la convergence de  $\psi_j$  vers  $\psi$  dans  $L^2(D)$  entraîne que :

$$\lim_{j \rightarrow \infty} \int_D \chi_{\Omega_n} \lambda f(\psi_j - \psi) = 0. \quad (3.36)$$

Par conséquent (3.35) et (3.36) impliquent que

$$\lim_{j \rightarrow \infty} \lim_{j \rightarrow \infty} I_3 = 0.$$

En procédant de la même manière que pour  $I_1$  et  $I_2$ , nous démontrons que

$$\lim_{j \rightarrow \infty} \lim_{j \rightarrow \infty} I_4 = 0.$$

et

$$\lim_{j \rightarrow \infty} \lim_{n \rightarrow \infty} I_5 = 0.$$

Ce qui achève la démonstration du Théorème 3.4. ■

Nous avons alors le résultat de compacité de  $\mathcal{F}$ .

**Théorème 3.5** *L'espace  $\mathcal{F}$  est compact pour la topologie définie par (3.14)*

**Démonstration :**

Soit  $(\Omega_n, u_n)_n$  une suite dans  $\mathcal{F}$ . D'après le théorème 3.3 et le théorème 3.4, il existe une sous suite, notée encore  $(\Omega_n, u_n)$ , et un couple  $(\Omega, W) \in \Theta_{ad} \times H^1(D)$ , tels que

$$(\Omega_n, u_n) \xrightarrow{n \rightarrow \infty} (\Omega, W|_{\Omega}),$$

$$\text{et } (\Omega, W|_{\Omega}) \in \mathcal{F}. \quad \text{■}$$

### 3.3.3 Continuité de la fonctionnelle coût

Dans ce paragraphe, nous allons démontrer la continuité de la fonctionnelle coût  $j$  sur  $\mathcal{F}$ . Ce qui revient en fait à montrer que  $J$  :

$$J(\varphi) = j(\Omega, u(\Omega)) = \frac{1}{2} \int_{\Gamma_0} (T(\varphi) - T_0)^2 d\sigma.$$

est continue sur  $U_{ad}$

Les termes  $\Omega$  et  $u$  qui interviennent dans l'expression de  $J$  sont tels que  $\Omega = \Omega(\varphi) \in \Theta_{ad}$  et  $T(\varphi) = u + V$  est la solution de l'équation (3.3).

Considérons maintenant  $(\varphi_n)_n$  une suite de  $U_{ad}$ , d'après le paragraphe précédent, il existe une sous-suite notée encore  $(\varphi_n)_n$  et un élément  $\varphi \in U_{ad}$  tels que

$$\varphi_n \xrightarrow{n \rightarrow \infty} \varphi \text{ uniformément dans } [a, b].$$

$$\text{i.e. } \Omega_n \xrightarrow{n \rightarrow \infty} \Omega.$$

D'après le Théorème 3.4, nous savons que la limite faible  $\tilde{u}$  de la suite  $\tilde{u}_n$  (solution de (3.4) sur  $\Omega_n$ ) est solution de (3.4) sur  $\Omega$ . Nous pouvons alors écrire :

$$\tilde{u}_n \xrightarrow{n \rightarrow \infty} \tilde{u} \quad \text{dans } H^1(D) \text{ faible,}$$

Par conséquent

$$\tilde{T}_n = \tilde{u}_n + V \xrightarrow{n \rightarrow \infty} \tilde{T} = \tilde{u} + V \quad \text{dans } H^1(D) \text{ faible,}$$

où  $\tilde{T}_n|_{\Omega_n}$  (respectivement  $\tilde{T}|_{\Omega}$ ) est la solution du problème (3.3) sur  $\Omega_n$  (respectivement sur  $\Omega$ ).

Nous avons alors le résultat de la continuité de  $J$  sur  $U_{ad}$ .

**Théorème 3.6** *La fonction  $J$  est continue sur  $U_{ad}$*

$$i.e \quad \lim_{n \rightarrow \infty} J(\varphi_n) = J(\varphi).$$

**Démonstration :**

Considérons la différence :

$$J(\varphi_n) - J(\varphi) = \frac{1}{2} \int_{\Gamma_0} (T_n - T_0)^2 d\sigma - \frac{1}{2} \int_{\Gamma_0} (T - T_0)^2 d\sigma.$$

Nous avons

$$\begin{aligned} J(\varphi_n) - J(\varphi) &= \frac{1}{2} \int_{\Gamma_0} (T_n - T_0)^2 - (T - T_0)^2 \\ &= \frac{1}{2} \int_{\Gamma_0} (T_n - T)(T_n + T - 2T_0) d\sigma. \end{aligned} \tag{3.37}$$

En appliquant l'inégalité de Hölder à 3.37, nous obtenons

$$\begin{aligned} |J(\varphi_n) - J(\varphi)| &\leq \frac{1}{2} \int_{\Gamma_0} |(T_n - T)| |(T_n + T - 2T_0)| d\sigma \\ &\leq \frac{1}{2} \left( \int_{\Gamma_0} (T_n - T)^2 d\sigma \right)^{\frac{1}{2}} \left( \int_{\Gamma_0} (T_n + T - 2T_0)^2 d\sigma \right)^{\frac{1}{2}}. \end{aligned}$$

D'une part, nous avons

$$\int_{\Gamma_0} (T_n + T - 2T_0)^2 d\sigma \leq 3 \left( \int_{\Gamma_0} (T_n)^2 d\sigma + \int_{\Gamma_0} (T)^2 d\sigma + 4 \int_{\Gamma_0} (T_0)^2 d\sigma \right)$$

D'après la continuité de l'application trace de  $H^1(D)$  dans  $L^2(\Gamma_0)$ , il existe une constante  $c_1 > 0$  indépendante de  $n$  telle que

$$\|T_n\|_{0,\Gamma_0} \leq c_1 \|\tilde{T}_n\|_{1,D} = c_1 \|\tilde{u}_n + V\|_{1,D} \tag{3.38}$$

$$\leq c_1 (\|\tilde{u}_n\|_{1,D} + \|V\|_{1,D}) \leq M_1.$$

D'autre part, d'après l'injection compacte de  $H^1(D)$  dans  $L^2(\Gamma_0)$ , nous avons

$$\int_{\Gamma_0} (T_n - T)^2 d\sigma = \|T_n - T\|_{0,\Gamma_0}^2 \xrightarrow{n \rightarrow \infty} 0. \quad (3.39)$$

En utilisant les inégalités (3.38), l'équation (3.39) et le fait que  $T_0$  est fixé dans  $L^2(\Gamma_0)$ , nous avons

$$\lim_{n \rightarrow \infty} (J(\varphi_n) - J(\varphi)) = 0.$$

■

# Étude de l'approximation du problème

---

## Sommaire

|            |                                                                              |           |
|------------|------------------------------------------------------------------------------|-----------|
| <b>4.1</b> | <b>Approximation du problème par la méthode des éléments finis . . . . .</b> | <b>51</b> |
| 4.1.1      | Discrétisation de la famille des domaines admissibles . . . . .              | 53        |
| 4.1.2      | Approximation du problème d'état . . . . .                                   | 55        |
| <b>4.2</b> | <b>Étude de l'existence d'une solution du problème discret . . . . .</b>     | <b>57</b> |
| <b>4.3</b> | <b>Étude de la convergence . . . . .</b>                                     | <b>62</b> |
| 4.3.1      | Résultat abstrait . . . . .                                                  | 62        |
| 4.3.2      | Résultat de convergence . . . . .                                            | 64        |

---

Dans ce chapitre, nous allons proposer une approximation du problème d'optimisation de forme par la méthode des éléments finis triangulaires. Cette méthode est basée sur des maillages variables réguliers qui dépendent continuellement du pas et de la frontière libre approchée. Celle-ci sera paramétrée par les courbes B-splines de Bézier. Ensuite, nous montrerons l'existence d'une solution optimale du problème discret. Puis, nous donnerons, sans démonstration, un résultat abstrait de la convergence. Enfin, nous utiliserons ce résultat abstrait pour étudier la convergence de la solution discrète vers la solution du problème d'optimisation de forme continu, lorsque le pas de discrétisation tend vers zéro.

Le travail présenté dans ce chapitre a donné lieu à une communication [21] et deux articles [22, 23].

## 4.1 Approximation du problème par la méthode des éléments finis

Nous allons discrétiser la famille des domaines admissibles  $\Theta_{ad}$ , ainsi que le problème d'état  $(PE)$ . Pour cela, le plus simple est de chercher à approcher les frontières libres par des courbes

linéaires par morceaux. Cependant, ce type d'approximation présente plusieurs inconvénients qui résident en particulier dans le fait que la frontière est déterminée via un grand nombre de paramètres auxquels il faut imposer des contraintes afin de préserver la régularité de la frontière. Les ingénieurs préfèrent discrétiser la famille des domaines admissibles, par des domaines qui sont suffisamment lisses et en même temps définis par un nombre fini de paramètres [41]. Pour toutes ces raisons, nous allons approcher  $\Gamma(\varphi)$  par les courbes Splines, qui sont localement déterminées par les fonctions quadratiques de Bézier.

Dans le paragraphe suivant, nous présentons brièvement quelques propriétés des courbes de Bézier dans  $\mathbb{R}^2$ .

#### 4.1.0.1 Courbes de Bézier

Notons par  $B_i^{(n)}$ ,  $i = 0, \dots, n$ , le polynôme de Bernstein défini sur  $[0, 1]$  par

$$B_i^{(n)}(t) = C_i^n t^i (1-t)^{n-i}, \quad t \in [0, 1]. \quad (4.1)$$

Le polynôme de Bernstein vérifie les propriétés suivantes :

$$B_i^{(n)}(t) \geq 0 \quad \forall t \in [0, 1]; \quad (4.2)$$

$$B_i^{(n)}(0) = \begin{cases} 0, & i \neq 0 \\ 1, & i = 0, \end{cases} \quad (4.3)$$

$$B_i^{(n)}(1) = \begin{cases} 0, & i \neq n \\ 1, & i = n, \end{cases} \quad (4.4)$$

$$\sum_{j=0}^n B_j^n(t) = 1 \quad \forall t \in [0, 1]; \quad (4.5)$$

$$\frac{d}{dt} B_i^n(t) = n(B_{i-1}^{n-1}(t) - B_i^{n-1}(t)) \quad \forall t \in [0, 1]; \quad (4.6)$$

Soit  $b_0, b_1, \dots, b_n \in \mathbb{R}^2$  un ensemble de points de contrôles.

**Définition 4.1** On appelle courbe de Bézier de degré  $n \geq 1$ , associée aux points de contrôle  $b_0, b_1, \dots, b_n \in \mathbb{R}^2$ , la courbe  $b(t)$  donnée par la paramétrisation suivante :

$$b(t) = \sum_{j=0}^n b_j B_j^n(t). \quad (4.7)$$

D'après (4.1) – (4.6) les courbes de Bézier possèdent les propriétés suivantes

1.  $b(0) = b_0, \quad b(1) = b_n;$
2.  $b(t) \in \text{conv} \{b_0, b_1, \dots, b_n\} \quad \forall t \in [0, 1];$
3. Une courbe de Bézier se trouve à l'intérieur de l'enveloppe convexe de ses points de contrôle.
4.  $\frac{db}{dt}(0) = n(b_1 - b_0), \quad \frac{db}{dt}(t) = n(b_n - b_{n-1});$

Dans notre cas nous considérons  $\varphi : [a, b] \rightarrow \mathbb{R}$  telle que

$$\varphi(s) = \sum_{i=0}^n \varphi_i B_i^{(n)}\left(\frac{s-a}{b-a}\right), \quad \varphi_i \in \mathbb{R}, t = \frac{s-a}{b-a}, s \in [a, b], t \in [0, 1]. \quad (4.8)$$

Où  $(a + j\frac{b-a}{n}, \varphi_j) \in \mathbb{R}^2, \quad j = 0, \dots, n$  sont les points de contrôle de  $\varphi$ , et  $(\varphi_j)_{j=0}^n$  sont les variables de formes qui caractérisent la forme de la frontière  $\Gamma(\varphi)$ . Pour que  $\varphi$  soit admissible, les coefficients  $\varphi_j$  dans (4.8) doivent satisfaire certaines conditions. Si  $0 \leq \varphi_j \leq L_y, \quad j = 0, \dots, n$ , alors, d'après la propriété 3 l'enveloppe convexe des points  $(a + j\frac{b-a}{n}, \varphi_j), \quad j = 0, \dots, n$ , est contenue dans  $[a, b] \times [0, L_y]$ . Par suite  $(a + (b-a)t, \varphi(t)) \in [a, b] \times [0, L_y] \quad \forall t \in [0, 1]$ .

En utilisant (4.2) et (4.6), nous avons le résultat suivant

**Lemme 4.1** Si  $\frac{1}{n} \sum_{j=0}^n |\varphi_j| \leq L_0$ , alors pour tout  $t, t' \in [0, 1]$  nous avons

$$\begin{aligned} |\varphi(a + (b-a)t) - \varphi(a + (b-a)t')| &= |\sum_{j=0}^n \varphi_j (B_j^n(t) - B_j^n(t'))| \\ &\leq \sum_{j=0}^n |\varphi_j| |B_j^n(t) - B_j^n(t')| \\ &\leq \sum_{j=0}^n |\varphi_j| \left| \frac{d}{dt} B_j^n(t'') \right| |t - t'| \quad / \quad t'' \in (\min(t, t'), \max(t, t')) \\ &\leq \sum_{j=0}^n |\varphi_j| |n(B_{j-1}^{n-1}(t'') - B_j^{n-1}(t''))| |t - t'| \quad / \quad t'' \in (\min(t, t'), \max(t, t')) \\ &\leq \sum_{j=0}^n |\varphi_j| |t - t'| \\ &\leq L_0 |t - t'| \end{aligned}$$

#### 4.1.1 Discrétisation de la famille des domaines admissibles

Considérons une partition uniforme  $(a_i)_{i=0}^d$  de l'intervalle  $[a, b]$ , telle que

$$a = a_0 < a_1 < \dots < a_d = b, \quad a_i = i\mu + a, \quad \mu = (b-a)/d, \quad i = 0, \dots, d;$$

et  $a_{i+1/2}$  est le point milieu de  $[a_i, a_{i+1}]$ .

On note par  $A_i = (a_i, \varphi_i), \quad \varphi_i \in \mathbb{R}, \quad i = 0, \dots, d$ , les noeuds et  $A_{i+1/2} = \frac{1}{2}(A_i + A_{i+1})$  sont les milieux des segments  $\overline{A_i A_{i+1}}, \quad i = 0, \dots, d-1$ . De plus, on pose  $a_{-\frac{1}{2}} = a - \frac{\mu}{2}, \quad a_{d+\frac{1}{2}} = b + \frac{\mu}{2},$

$A_{-\frac{1}{2}} = (a_{-\frac{1}{2}}, \frac{1}{2}(\varphi_0 + \varphi_1))$  et  $A_{d+\frac{1}{2}} = (a_{d+\frac{1}{2}}, \frac{1}{2}(\varphi_{d-1} + \varphi_d))$ .

Nous introduisons les ensembles suivants :

$\tilde{U}_\mu = \{\tilde{s}_\mu \in C^1([a - \frac{\mu}{2}, b + \frac{\mu}{2}]) \mid \tilde{s}_\mu|_{[a_{i-\frac{1}{2}}, a_{i+\frac{1}{2}}]}$  est une fonction quadratique de Bézier  
déterminée par  $\{A_{i-\frac{1}{2}}, A_i, A_{i+\frac{1}{2}}\}, i = 0, \dots, d\}$ ,

$$U_\mu = \tilde{U}_\mu|_{[a,b]} = \{s_\mu \in C^1([a,b]) \mid \exists \tilde{s}_\mu \in \tilde{U}_\mu : s_\mu = \tilde{s}_\mu|_{[a,b]}\}.$$

**Remarque 4.1** Les points du triplet  $\{A_{i-\frac{1}{2}}, A_i, A_{i+\frac{1}{2}}\}$  sont ce qu'on appelle les points de contrôle de la fonction de Bézier.

Afin de définir une famille de formes admissibles, localement déterminées par des fonctions de Bézier, il est nécessaire de préciser les  $\varphi_i \in \mathbb{R}$ , qui définissent les positions de  $A_i$ ,  $i = 0, \dots, d$ . À la partition  $(a_i)_{i=0}^d$ , nous associons l'ensemble  $\mathcal{Q}_\mu^{ad} \subset U_{ad}$  des fonctions continues, linéaires par morceaux sur  $(a_i)_{i=0}^d$  :

$$\mathcal{Q}_\mu^{ad} = \{\varphi_\mu \in C([a,b]) \mid \varphi_\mu|_{[a_{i-1}, a_i]} \in P_1([a_{i-1}, a_i]) \forall i = 1, \dots, d\} \cap U_{ad}. \quad (4.9)$$

Ainsi, nous pouvons identifier  $\mathcal{Q}_\mu^{ad}$  avec le sous ensemble convexe, compact

$\mathcal{U} \subset \mathbb{R}^{d+1}$  :

$$\mathcal{U} = \{ \varphi = (\varphi_0, \dots, \varphi_d) \in \mathbb{R}^{d+1} \mid 0 \leq \varphi_i \leq L_y, i = 0, \dots, d; \exists i_0, i_d \in \{0, \dots, d\} \mid \varphi_k = 0, \quad (4.10)$$

$$k = 0, \dots, i_0; \varphi_{d-k} = 0, k = 0, \dots, i_d \text{ et } |\varphi_{i+1} - \varphi_i| \leq L_0\mu, i = 0, \dots, d-1 \}$$

où  $\varphi_i = \varphi_\mu(a_i)$ ,  $i = 0, \dots, d$  et  $\varphi_\mu \in \mathcal{Q}_\mu^{ad}$ .

La famille des domaines admissibles discrétisée est alors représentée par

$$\Theta_{ad}^\mu = \{ \Omega(s_\mu) \mid s_\mu \in U_{ad}^\mu \} \quad (4.11)$$

où

$$U_{ad}^\mu = \{s_\mu \in U_\mu \mid \text{les noeuds de forme } A_i = (a_i, \varphi_i), i = 0, \dots, d \quad (4.12)$$

sont tels que  $\varphi = (\varphi_0, \dots, \varphi_d) \in \mathcal{U} \}$

### 4.1.2 Approximation du problème d'état

À présent, nous commençons l'approximation du problème d'état ( $PE$ ). Nous utilisons la méthode des éléments finis triangulaires, sur  $\Omega(s_\mu) \in \Theta_{ad}^\mu$ , qui désigne une approximation appropriée de  $\Omega \in \Theta_{ad}$ . Nous introduisons une autre famille de partitions  $(b_i)_{i=0}^q$  de  $[a, b]$ , telle que :  $a = b_0 < b_1 < \dots < b_q = b$  (n'est pas nécessairement équidistant). La longueur de cette partition sera notée par  $h$ . Nous supposons que  $h \rightarrow 0^+$  si  $\mu \rightarrow 0^+$ .

Notons par  $r_h s_\mu$ , l'interpolation linéaire par morceaux de Lagrange, de  $s_\mu$  en  $(b_i)_{i=0}^q$  :

$$r_h s_\mu(b_i) = s_\mu(b_i) \quad \forall i = 0, \dots, q;$$

$$r_h s_\mu|_{[b_{i-1}, b_i]} \in P_1([b_{i-1}, b_i]) \quad \forall i = 1, \dots, q;$$

Ainsi, le domaine de calcul  $\Omega(s_\mu)$  sera représenté par  $\Omega(r_h s_\mu)$ . Dans ce qui suit, l'ensemble de tous les  $\Omega(r_h s_\mu)$ ,  $s_\mu \in U_{ad}^\mu$ , sera noté par  $\Theta_{ad}^{\mu h}$  :

$$\Theta_{ad}^{\mu h} = \{\Omega(r_h s_\mu) \mid s_\mu \in U_{ad}^\mu\}. \quad (4.13)$$

Puisque  $\Omega(r_h s_\mu)$  est déjà polygonal, nous pouvons construire sa triangulation notée  $\mathcal{T}(h, s_\mu)$ , qui dépend de  $h > 0$  et de  $s_\mu \in U_{ad}^\mu$ . Nous supposons que pour  $h > 0$  fixe, les triangulations  $\mathcal{T}(h, s_\mu)$  sont topologiquement équivalentes pour tout  $s_\mu \in U_{ad}^\mu$ , i.e :

**(HT1)**  $\mathcal{T}(h, s_\mu)$  ont le même nombre de noeuds, et ces noeuds ont les mêmes voisinages pour tout  $s_\mu \in U_{ad}^\mu$ .

**(HT2)** La position des noeuds de  $\mathcal{T}(h, s_\mu)$  dépend continuellement des variations des points de forme  $\{A_i\}_{i=1}^d$ .

**(HT3)** Il existe  $\eta_0 > 0$  tel que tous les angles intérieurs  $\eta(h, s_\mu)$  vérifient :

$$\eta(h, s_\mu) \geq \eta_0 \quad \forall h > 0, \forall s_\mu \in U_{ad}^\mu$$

**(HT4)** Le segment  $\Gamma_i(s_\mu) = \{(x, y) / y = r_h s_\mu(x), x \in [b_i, b_{i+1}]\}$  est côté d'un triangle

$$T \in \mathcal{T}(h, s_\mu), \quad i = 0, \dots, q-1.$$

Le domaine  $\Omega(r_h s_\mu)$  muni de la triangulation  $\mathcal{T}(h, s_\mu)$  sera noté  $\Omega_h(s_\mu)$  ou  $\Omega_h$ . Quant à l'approché de  $\Gamma$ , il sera noté par  $\Gamma_h$ . Nous définissons alors l'espace de dimension finie associé à  $H^1(\Omega)$

$$H_h(\Omega_h) = \{v_h \in C(\overline{\Omega}_h) \mid v_h|_T \in P1(T), T \in \mathcal{T}(h, s_\mu)\}$$

puis l'espace de dimension finie associé à  $H_{\Gamma_d}^1(\Omega)$  :

$$H_{\Gamma_d}^h(\Omega_h) = \{v_h \in H_h(\Omega_h) \mid v_h|_{\Gamma_{d,h}} = 0\}.$$

**Remarque 4.2** Notons que la méthode des éléments finis utilisée ici est celle des éléments finis conformes [28], i.e pour tout  $h > 0$ , nous avons les inclusions suivantes :

$$H_h(\Omega_h) \subset H^1(\Omega), \quad H_{\Gamma_d}^h(\Omega_h) \subset H_{\Gamma_d}^1(\Omega), \quad U_{ad}^\mu \subset U_{ad}$$

Soit  $s_\mu \in U_{ad}^\mu$ , l'approximation  $u_h := u_h(s_\mu) \in H_{\Gamma_d}^h(\Omega_h)$  de  $u \in H_{\Gamma_d}^1(\Omega)$  donnée par :

$$u_h = \sum_{i=1}^N u_h(\bar{b}_i) \psi_i, \quad (4.14)$$

où  $N$  désigne le nombre de noeuds de la triangulation  $\mathcal{T}(h, s_\mu)$  du domaine  $\overline{\Omega}_h$ , tandis que  $(\bar{b}_i)_{1 \leq i \leq N}$  sont les noeuds de la triangulation et  $(\psi)_{i=1}^N$  sont les fonctions de base de  $H_{\Gamma_d}^h(\Omega_h)$ . Soit  $F(r_h(s_\mu)) = D \setminus \Omega(r_h(s_\mu))$ . Nous allons construire une autre famille de triangulation du domaine  $\overline{F(r_h(s_\mu))}$  notée par  $\{\mathcal{T}E(h, s_\mu)\}$ , qui satisfait les hypothèses (HT1) – (HT4). L'union des deux triangulations  $\mathcal{T}(h, s_\mu)$  et  $\mathcal{T}E(h, s_\mu)$  définit une triangulation régulière de  $\overline{D}$ . Soit  $V_h$ , une interpolation, linéaire par morceaux de type Lagrange, de  $V$  dans  $\overline{D}$ . Par conséquent, le problème d'état discret s'écrit :

$$\left\{ \begin{array}{l} \text{Trouver } u_h \in H_{\Gamma_d}^h(\Omega_h) \text{ tel que} \\ \int_{\Omega_h} \lambda_h \nabla u_h \cdot \nabla v_h + \int_{\Omega_h} \mathcal{V}_0^h \cdot \nabla u_h v_h = \int_{\Omega_h} f v_h - \int_{\Omega_h} \lambda_h \nabla V_h \cdot \nabla v_h - \\ \int_{\Omega_h} \mathcal{V}_0^h \cdot \nabla V_h v_h \quad \forall v_h \in H_{\Gamma_d}^h(\Omega_h) \end{array} \right. \quad (4.15)$$

où  $\lambda_h$  est une approximation de  $\lambda$  tel que  $\lambda_h$  est uniformément borné, converge vers  $\lambda$  presque partout dans  $D$ , et satisfait l'équation suivante :

$$\exists \lambda_0 > 0 \quad \text{indépendant de } h \text{ tel que } \lambda_h(x) \xi \cdot \xi > \lambda_0 |\xi|^2 \quad \text{p.p } x \in D \quad (4.16)$$

et

$$\begin{cases} \mathcal{V}_0^h \text{ est une approximation de } \mathcal{V}_0 \text{ tel que} \\ \mathcal{V}_0^h \text{ est uniformément borné, converge vers } \mathcal{V}_0 \text{ p.p} \end{cases} \quad (4.17)$$

Nous approchons la fonctionnelle coût par la fonction discrète suivante :

$$J_h(u_h(s_\mu)) = J_h(\Omega_h) = \frac{1}{2} \int_{\Gamma_0^h} |T_h(x, y) - T_0^h(x, y)|^2 d\sigma \quad (4.18)$$

où,  $T_0^h$  est l'interpolé de  $T_0$  et  $T_h = u_h + V_h$  tel que  $u_h \in H_{\Gamma_d}^h(\Omega_h)$  solution de (4.15).

Finalement, le problème d'optimisation de forme discret est le suivant

$$(PO_h) \begin{cases} \inf_{s_\mu \in U_{ad}^\mu} J_h(u_h(s_\mu)) \\ \text{telle que } u_h(s_\mu) \text{ est solution de (4.15) on } \Omega_h \end{cases} \quad (4.19)$$

## 4.2 Étude de l'existence d'une solution du problème discret

Dans le cas d'un problème d'état régi par un opérateur coercif, pour montrer l'existence d'une solution optimale du problème d'optimisation de forme discret, le moyen qui est relativement simple serait de l'écrire sous une forme matricielle équivalente (cf. [41, 42], [14, 17]). Ce qui nous amène à utiliser les outils classiques d'optimisation en dimension finie pour étudier l'existence d'une solution du problème discret. Comme dans notre cas le problème d'état est régi par un opérateur non coercif, ces outils ne sont plus valables ; il faut alors procéder autrement. Pour cela, nous utilisons le degré topologique de Brouwer.

Nous commençons par énoncer quelques résultats qui seront utiles pour prouver le résultat d'existence. Nous énonçons tout d'abord ce lemme dont la démonstration est analogue au Théorème 3.3.

**Lemme 4.2**  $\exists C > 0$  ,  $\forall s_\mu \in U_{ad}^\mu$  et  $\forall h > 0$

$$\|u_h(s_\mu)\|_{1,\Omega_h} \leq C. \quad (4.20)$$

Nous démontrons aussi le résultat suivant

**Lemme 4.3** Soient  $(s_\mu^j)_j \subset U_{ad}^\mu$  et  $s_\mu \in U_{ad}^\mu$  tels que

$$s_\mu^j \rightarrow s_\mu \quad \text{quand } j \rightarrow \infty,$$

alors

$$\forall h > 0 \quad \Pi_h V(s_\mu^j) \rightarrow \Pi_h V(s_\mu) \quad \text{in } H^1(D).$$

### Démonstration

Soit  $(M_i)_{1 \leq i \leq N}$ , l'ensemble des noeuds intérieurs de la triangulation  $\mathcal{T}(h, s_\mu)$  avec  $M_i = (x(i), y(i))$ .

Donc, d'après les hypothèses  $(HT_1)$  et  $(HT_2)$ , nous avons

$$M_i = \phi_i(s_\mu), \quad i = 1, \dots, n, \quad \forall s_\mu \in U_{ad}^\mu, \quad (4.21)$$

où  $\phi_i : U_{ad}^\mu \mapsto \mathbb{R}^2$  sont des fonctions continues. Soit  $\psi_k \in H_h^1(D)$  les fonctions de base associées aux  $M_k$ , i.e.,  $\psi_i(M_k) = \delta_{ik}$ , où  $\delta_{ik}$  est le symbole de Kronecker. Finalement, soit  $T \in \mathcal{TE}(h, s_\mu) \cup \mathcal{T}(h, s_\mu)$  un triangle dont  $M_k$  est l'un des noeuds, et dont les autres noeuds sont  $M_i$  et  $M_l$ . Nous savons que

$$\psi_{i|T} = \left( \begin{vmatrix} x(i) & x(l) \\ y(i) & y(l) \end{vmatrix} - \begin{vmatrix} 1 & y(i) \\ 1 & y(l) \end{vmatrix} x + \begin{vmatrix} 1 & x(i) \\ 1 & x(l) \end{vmatrix} y \right) / (2\text{meas}(T))$$

où

$$\text{meas}(T) = \begin{vmatrix} 1 & x(k) & y(k) \\ 1 & x(i) & y(i) \\ 1 & x(l) & y(l) \end{vmatrix}$$

D'après ce qui précède et d'après l'équation (4.21), nous avons

$$\psi_{i|T} = a_i(s_\mu)x + b_i(s_\mu)y + c_i(s_\mu), \quad (4.22)$$

où les coefficients,  $a_i(s_\mu)$ ,  $b_i(s_\mu)$  et  $c_i(s_\mu)$  sont des fonctions continues sur  $s_\mu$  dans  $U_{ad}^\mu$ . Par suite,  $\psi_{i|T}$  et  $\nabla\psi_{i|T}$  sont des fonctions continues sur  $s_\mu$  dans  $U_{ad}^\mu$ . Nous en déduisons que  $\Pi_h V(s_\mu)$  et  $\nabla\Pi_h V(s_\mu)$  sont continues.

Puisque  $measT$  dépend continuellement de  $s_\mu$ , alors pour tout  $(s_\mu^j)_j \subset U_{ad}^\mu$ , nous avons

$$s_\mu^j \xrightarrow{j \rightarrow \infty} s_\mu \implies \begin{cases} \int_D (\Pi_h V(s_\mu^j))^2 \xrightarrow{j \rightarrow \infty} \int_D (\Pi_h V(s_\mu))^2 \\ \text{et} \\ \int_D |\nabla \Pi_h V(s_\mu^j)|^2 \xrightarrow{j \rightarrow \infty} \int_D |\nabla \Pi_h V(s_\mu)|^2. \end{cases}$$

Ceci achève la démonstration. ■

À présent, nous pouvons énoncer le théorème d'existence.

**Théorème 4.1** *Sous les hypothèses (4.16) – (4.17), le problème (4.19) admet une solution dans  $U_{ad}^\mu$ , pour tout  $h > 0$  et  $\mu > 0$ .*

**Démonstration.**

Tout d'abord, pour chaque  $s_\mu \in U_{ad}^\mu$  fixe et  $h > 0$ , nous devons montrer qu'il existe une solution unique  $u_h(s_\mu) \in H_{\Gamma_d}^h(\Omega_h)$  de (4.15). Notre outil principal est le degré topologique de Brouwer.

Pour  $t$  appartenant à  $[0, 1]$ , nous définissons l'opérateur  $F_t$  par :

$$F_t : H_{\Gamma_d}^h(\Omega_h) \rightarrow H_{\Gamma_d}^h(\Omega_h), \\ \bar{u}_h \mapsto u_h,$$

telle que  $u_h$  est la solution unique du problème suivant

$$\int_{\Omega_h} \lambda_h \nabla u_h \cdot \nabla v_h = \int_{\Omega_h} f v_h - t \int_{\Omega_h} \mathcal{V}_0^h \cdot \nabla \bar{u}_h v_h - \int_{\Omega_h} \lambda_h \nabla V_h \cdot \nabla v_h - \int_{\Omega_h} \mathcal{V}_0^h \cdot \nabla V_h v_h \quad (4.23)$$

D'après le lemme 4.2, nous avons  $\|u_h\|_{1,\Omega_h} < C$ , ce qui nous permet de construire une boule

ouverte  $B$ , telle que  $F_t$  n'admet aucun point fixe sur  $\partial B$ , la frontière de  $B$ . Par suite, la fonction  $\deg[I - F_t, B, 0]$  est définie. De plus elle est indépendante de  $t$ . Puisque  $F_0$  est trivial, nous concluons que

$$\deg[I - F_0, B, 0] = \deg[I - F_1, B, 0] = +1.$$

Par conséquent,  $F_1$  admet un point fixe à l'intérieur de  $B$  qui est une solution de (4.15). Pour montrer que la solution discrète est unique, puisque le second membre de (4.15) est linéaire, alors il suffit de montrer que la seule solution de (4.15) avec un second membre nul est la solution nulle (comme dans le cas continu).

Ensuite, d'après l'hypothèse  $(HT_1)$ , nous avons :  $\dim H_{\Gamma_d}^h(\Omega_h)$  est la même pour tout  $s_\mu \in U_{ad}^h$ . D'où le fait que l'ensemble  $U_{ad}^\mu$  peut être identifié à un sous ensemble de  $\mathbb{R}^{d+1}$ .

Soit  $(s_\mu^j)_j \subset U_{ad}^\mu$ . Nous pouvons alors extraire une sous suite notée encore  $(s_\mu^j)_j$  telle que

$$s_\mu^j \xrightarrow{j \rightarrow \infty} s_\mu \quad \text{dans } U_{ad}^\mu$$

et

$$\Omega_h^j \xrightarrow{j \rightarrow \infty} \Omega_h.$$

D'après le lemme 4.2,  $\exists C > 0$  telle que

$$\|u_h(s_\mu^j)\|_{1, \Omega_h(s_\mu^j)} \leq C. \quad (4.24)$$

D'après le résultat d'extension uniforme [25], il existe  $\tilde{u}_h(s_\mu^j)$  un prolongement uniforme de  $u_h(s_\mu^j)$  de  $\Omega_h(s_\mu^j)$  dans  $D$ , tel que

$$\exists M > 0 \quad \forall j \quad \|\tilde{u}_h(s_\mu^j)\|_{1,D} < M.$$

Ainsi, il existe une sous suite  $\tilde{u}_h(s_\mu^j)$  et un élément  $\tilde{u}_h \in H^1(D)$ , tels que

$$\tilde{u}_h(s_\mu^j) \xrightarrow{j \rightarrow \infty} \tilde{u}_h \quad \text{dans } H^1(D).$$

Enfin, montrons que  $u_h = \tilde{u}_h|_{\Omega_h}$  résout (4.15). Il est aisé de voir que  $u_h|_{\Gamma_4} = 0$ . Et, de la même manière que dans le cas continu, nous montrons que  $u_h \in H_{\Gamma_d}^1(\Omega_h)$ .

Il reste à montrer que  $u_h$  résout (4.15). En effet, soit  $\psi_h \in H_{\Gamma_d}^h(\Omega_h)$  et  $\tilde{\psi} \in H^1(D)$  est une extension de  $\psi_h$  définie par

$$\tilde{\psi} = \begin{cases} \psi_h & \text{dans } \Omega_h^* \\ 0 & \text{dans } D \setminus \Omega_h^*. \end{cases}$$

Ainsi, nous pouvons construire une suite  $(\psi_n)_n$ ,  $\psi_n \in \mathcal{D}(\bar{D})$  telle que,

$$\text{dist}(\text{supp } \psi_n, \overline{\Gamma_d}) > 0 \quad \forall n \in \mathbb{N} \quad \text{et} \quad \psi_n \rightarrow \tilde{\psi} \quad \text{dans } H^1(D), n \rightarrow \infty.$$

Soit  $n \in \mathbb{N}$ . Puisque  $\Omega_h^j \xrightarrow{j \rightarrow \infty} \Omega_h^*$ , il existe  $j_0$  tel que

$$\psi_n^h|_{\Omega_h^j} \in H_{\Gamma_d}^1(\Omega_h^j), \forall j \geq j_0,$$

où  $\psi_n^h = \pi_h \psi_n$  est l'interpolé linéaire par morceaux de  $\psi_n$  dans  $\mathcal{T}(h, s_\mu^j) \cup \mathcal{T}E(h, s_\mu^j)$ . Nous notons que

$$\pi_h \psi_n \xrightarrow{j \rightarrow \infty} \pi_h \tilde{\psi} \quad \text{dans } H^1(D),$$

avec  $\pi_h \tilde{\psi}|_{\Omega_h^*} \in H_{\Gamma_d}^1(\Omega_h^*)$ . Pour tout  $j \geq j_0$ , nous avons :

$$\begin{aligned} & \int_D \chi_{\Omega_h^j} \lambda_h \nabla \tilde{u}_h(s_\mu^j) \cdot \nabla \psi_n^h + \int_D \chi_{\Omega_h^j} \mathcal{V}_0^h \cdot \nabla \tilde{u}_h(s_\mu^j) \psi_n^h = \int_D \chi_{\Omega_h^j} f \psi_n^h \\ & - \int_D \chi_{\Omega_h^j} \lambda_h \nabla \pi_h V(s_\mu^j) \cdot \nabla \psi_n^h - \int_D \chi_{\Omega_h^j} \mathcal{V}_0^h \cdot \nabla \pi_h V(s_\mu^j) \psi_n^h. \end{aligned} \tag{4.25}$$

De la même façon que dans le cas continu, en faisant tendre  $n$ , puis  $j$  vers l'infini dans l'équation (4.25), nous obtenons ainsi que  $u_h$  est solution de l'équation (4.15).

■

### 4.3 Étude de la convergence

Nous commençons par rappeler un résultat abstrait de la convergence d'une suite de solutions discrètes vers la solution du problème d'optimisation de forme continu [41, 42]. Ensuite, nous utilisons ce résultat abstrait pour étudier la convergence de notre problème discret.

#### 4.3.1 Résultat abstrait

Dans la suite nous allons utiliser une approximation de la frontière libre discrète par des fonctions splines de type Bézier. D'autre part, le problème d'état sera discrétisé en utilisant la méthode des éléments finis  $P_1$ .

Notons par  $\mu > 0$  le paramètre de discrétisation des formes, et par  $n(\mu)$  le nombre de paramètres qui définissent la forme du domaine discret  $\Omega_\mu$ . Pour  $\mu > 0$  fixé, l'ensemble des domaines admissibles discrets est noté par  $\Theta_\mu$ . Nous supposons que le nombre  $n(\mu)$  est le même pour tout  $\Omega_\mu \in \Theta_\mu$ ,  $\mu > 0$  fixe, et que  $\Theta_\mu \subset \tilde{\Theta}$ , pour tout  $\mu > 0$  (mais non nécessairement  $\Theta_\mu \subset \Theta_{ad}$ ), où  $\tilde{\Theta}$  est une famille qui contient  $\Theta_{ad}$ .

Pour tout  $\Omega_\mu \in \Theta_\mu$ , nous associons, d'une manière unique, un domaine de calcul  $(\Omega_\mu)_h$ , où  $h = h(\mu) > 0$  est un autre paramètre de discrétisation tel que :

$$h \rightarrow 0 \Leftrightarrow \mu \rightarrow 0.$$

Autrement dit, la connaissance de  $\Omega_\mu$ , nous permet de construire  $(\Omega_\mu)_h$  et inversement. L'ensemble de tous les domaines de calculs correspondant à  $\Theta_\mu$  sera noté  $\Theta_{\mu,h}$ .

Dans ce qui suit, nous supposons que  $\Theta_{\mu,h} \subset \tilde{\Theta}$ , pour tout  $\mu > 0$ ,  $h(\mu) > 0$ .

Nous associons à  $(\Omega_\mu)_h \in \Theta_{\mu,h}$ , l'espace de dimension finie  $W_h((\Omega_\mu)_h) \subset W((\Omega_\mu)_h)$ . Nous supposons que  $\dim W_h((\Omega_\mu)_h)$  est la même pour tout  $(\Omega_\mu)_h \in \Theta_{\mu,h}$  pour  $h > 0, \mu > 0$  fixés.

Puisque les deux ensembles  $\Theta_\mu$  et  $\Theta_{\mu,h}$  sont contenus dans  $\tilde{\Theta}$ , nous pouvons utiliser la même convergence des domaines utilisée dans le chapitre 3. Nous définissons alors l'application suivante :

$$u_h : (\Omega_\mu)_h \mapsto u_h((\Omega_\mu)_h) \in W_h((\Omega_\mu)_h).$$

Nous associons à tout  $(\Omega_\mu)_h \in \Theta_{\mu,h}$ , un unique élément  $u_h((\Omega_\mu)_h)$  de  $W_h((\Omega_\mu)_h)$ . Ainsi en

commençant par  $\Omega_\mu \in \Theta_\mu$ , nous avons la chaîne des applications suivantes

$$\Omega_\mu \mapsto (\Omega_\mu)_h \mapsto u_h((\Omega_\mu)_h) \mapsto J((\Omega_\mu)_h, u_h((\Omega_\mu)_h)).$$

Par suite, le problème d'optimisation discret s'écrit

$$(P_{\mu,h}) \left\{ \begin{array}{l} \text{trouver } (\Omega_\mu^*)_h \in \Theta_{\mu,h} \text{ tel que} \\ J((\Omega_\mu^*)_h, u_h((\Omega_\mu^*)_h)) \leq J((\Omega_\mu)_h, u_h((\Omega_\mu)_h)) \quad \forall (\Omega_\mu)_h \in \Theta_{\mu,h}. \end{array} \right. \quad (4.26)$$

Nous notons par  $\Omega_\mu^*$ , le domaine correspondant à  $(\Omega_\mu^*)_h$  la solution du problème  $(P_{\mu,h})$ .

Dans ce qui suit, la solution optimale du problème  $(P_{\mu,h})$  sera notée par  $(\Omega_\mu^*, u_h((\Omega_\mu^*)_h))$ .

À présent, nous allons introduire les hypothèses qui vont nous permettre d'énoncer le théorème de convergence.

$(H_1)_\mu$  pour tout  $\Omega \in \Theta_{ad}$  il existe  $\{\Omega_\mu\}_\mu$ ,  $\Omega_\mu \in \Theta_\mu$ , tel que

$$\Omega_\mu \xrightarrow{\tilde{\Theta}} \Omega \quad \text{et} \quad (\Omega_\mu)_h \xrightarrow{\tilde{\Theta}} \Omega, \quad h, \mu \rightarrow 0+;$$

$(H_2)_\mu$  pour toute suite  $(\Omega_\mu, u_h(\Omega_\mu))_\mu \subset \Theta_\mu \times W_h(\Omega_\mu)$ , il existe une sous suite  $(\Omega_{\mu_j}, u_{h_j}(\Omega_{\mu_j}))_{\mu_j}$  et un élément  $(\Omega, u) \in \Theta_{ad} \times W(\Omega)$  tels que

$$\Omega_{\mu_j} \xrightarrow{\tilde{\Theta}} \Omega \quad \text{et} \quad (\Omega_{\mu_j})_{h_j} \xrightarrow{\tilde{\Theta}} \Omega, \quad j \rightarrow \infty,$$

$$u_{h_j}((\Omega_{\mu_j})_{h_j}) \xrightarrow{j \rightarrow \infty} u(\Omega), \quad \text{faiblement dans } W(D);$$

$(H_3)_\mu$

$$\left. \begin{array}{l} (\Omega_\mu)_h \xrightarrow{\tilde{\Theta}} \Omega, \quad (\Omega_\mu)_h \in \Theta_{\mu,h}, \quad \Omega \in \Theta_{ad} \\ u_h((\Omega_\mu)_h) \xrightarrow{\mu, h \rightarrow 0} u(\Omega), \quad \text{dans } W(D) \end{array} \right\} \Rightarrow J((\Omega_\mu)_h, u_h((\Omega_\mu)_h)) \xrightarrow{\mu, h \rightarrow 0} J((\Omega, u((\Omega))).$$

Nous avons alors le théorème de convergence suivant dû à Haslinger et al. [42] :

**Théorème 4.2** *Si les hypothèses  $(H_1)_\mu - (H_3)_\mu$  sont vérifiées et si pour tout  $\mu, h > 0$ , le problème  $(P_{\mu,h})$  admet au moins une solution notée  $(\Omega_\mu^*, u_h((\Omega_\mu^*)_h))$ . Alors, il existe une sous suite*

$(\Omega_{\mu_j}^*, u_h((\Omega_{\mu_j}^*)_{h_j}))$  et un élément  $(\Omega^*, u(\Omega^*)) \in \mathcal{F}$  tels que

$$\left\{ \begin{array}{l} \Omega_{\mu_j}^* \xrightarrow{\tilde{\Theta}} \Omega^* \quad \text{quand } j \rightarrow \infty \\ u_{h_j}((\Omega_{\mu_j}^*)_{h_j}) \xrightarrow{j \rightarrow \infty} u(\Omega^*), \quad \text{dans } W(D). \end{array} \right. \quad (4.27)$$

De plus,  $(\Omega^*, u(\Omega^*))$  est une solution optimale de (2.3).

### 4.3.2 Résultat de convergence

En nous basant sur le résultat abstrait du paragraphe précédent, nous allons montrer un théorème de convergence d'une suite de solutions du problème discret vers la solution du problème continu, sous certaines hypothèses de régularité.

Pour cela, nous commençons d'abord par donner quelques notations et propriétés, qui seront utiles par la suite.

Puisque la triangulation vérifie les hypothèses de régularité  $(HT_1) - (HT_4)$ , alors d'après les résultats classiques d'approximation (cf. [28]), il existe une constante  $C > 0$  indépendante de  $h, \mu$ , telle que

$$\|V - \Pi_h V\|_{1,D} \leq Ch \quad (4.28)$$

Nous supposons que  $\lambda_h$  vérifie les conditions suivantes :

$$\left\{ \begin{array}{l} \lambda_h \text{ est uniformément borné par rapport à } h \\ \lambda_h \text{ converge presque partout vers } \lambda \text{ quand } h \rightarrow 0 \end{array} \right. \quad (4.29)$$

Puis, nous définissons la convergence de la suite  $\Omega_h(s_\mu)$  vers  $\Omega(\varphi)$  par

$$\Omega_h(s_\mu) \xrightarrow{h, \mu \rightarrow 0} \Omega(\varphi) \iff r_h(s_\mu) \xrightarrow{h, \mu \rightarrow 0} \varphi \text{ uniformément dans } [a, b]. \quad (4.30)$$

Notons par  $\tilde{\psi} \in H^1(D)$  une extension de  $\psi \in H^1(\Omega_h)$  dans  $H^1(D)$ . Par ailleurs,  $u_h$  et  $u$  désignent respectivement la solution du problème (4.15) dans  $\Omega_h$ , et celle du problème (3.4) dans  $\Omega(\varphi)$ . Nous définissons alors la convergence de  $u_h(s_\mu)$  vers  $u(\varphi)$ , par :

$$u_h(s_\mu) \xrightarrow{h, \mu \rightarrow 0} u \iff \tilde{u}_h(s_\mu) \xrightarrow{h, \mu \rightarrow 0} \tilde{u} \text{ faiblement dans } H^1(D). \quad (4.31)$$

D'après le théorème 4.2, pour montrer la convergence de la solution du problème d'optimisation de forme discret vers la solution du problème continu, il suffit de démontrer le lemme suivant :

**Lemme 4.4** (i) *Pour tout  $\varphi \in U_{ad}$ , il existe une suite  $(s_\mu)_\mu$  dans  $U_{ad}^\mu$  telle que*

$$\Omega_h(s_\mu) \xrightarrow{h, \mu \rightarrow 0} \Omega(\varphi) \quad (4.32)$$

(ii) *Soit  $((\Omega_h(s_\mu), u_h(s_\mu)))$  une suite telle que  $u_h(s_\mu)$  est la solution de (4.15) dans  $\Omega_h(s_\mu)$ . Alors il existe une sous suite de  $((\Omega_h(s_\mu), u_h(s_\mu)))$  notée encore par  $((\Omega_h(s_\mu), u_h(s_\mu)))$  et deux éléments  $\varphi \in U_{ad}$  et  $u \in H_{\Gamma_D}^1(D)$  tels que*

$$\begin{cases} \Omega_h(s_\mu) \xrightarrow{h, \mu \rightarrow 0} \Omega(\varphi), \\ u_h(s_\mu) \xrightarrow{h, \mu \rightarrow 0} u \end{cases} \quad (4.33)$$

*De plus  $u|_{\Omega(\varphi)}$  est une solution de (3.4) dans  $\Omega(\varphi)$ .*

(iii) *Si  $((\Omega_h(s_\mu), u_h(s_\mu)))$  est une suite, avec  $u_h(s_\mu)$  est une solution de (4.15) , et  $(\Omega(\varphi), u) \in \mathcal{F}$  tels que*

$$\begin{cases} \Omega_h(s_\mu) \xrightarrow{h, \mu \rightarrow 0} \Omega(\varphi), \\ u_h(s_\mu) \xrightarrow{h, \mu \rightarrow 0} u \end{cases} \quad (4.34)$$

*alors*

$$J(u_h(s_\mu)) \rightarrow J(u(\varphi)) \quad \text{quand } h, \mu \rightarrow 0, \quad (4.35)$$

**Démonstration :**

(i) Étant donné  $\bar{\varphi} \in U_{ad}$ , pour  $\mu > 0$ , nous allons définir une suite  $(\bar{\varphi}_\mu)$  qui converge vers  $\bar{\varphi}$ . En effet, soit  $(\bar{\varphi}_\mu)$ , telle que

$$\bar{\varphi}_\mu(a_i) = \frac{1}{\mu} \int_{a_{i-\frac{1}{2}}}^{a_{i+\frac{1}{2}}} \bar{\varphi}(x) dx \quad i = 0, \dots, d$$

où  $a_{i+\frac{1}{2}}$  sont les milieux des segments  $[a_i, a_{i+1}]$ , avec  $(a_{-\frac{1}{2}} = a - \frac{\mu}{2}, a_{d+\frac{1}{2}} = b + \frac{\mu}{2})$ . D'après le résultat de Begis et Glowinski (cf [8]), nous avons

$$\bar{\varphi}_\mu \rightarrow \bar{\varphi} \text{ uniformément dans } [a, b], \quad \mu \rightarrow 0+. \quad (4.36)$$

De plus, nous avons  $\bar{\varphi}_\mu \in Q_\mu^{ad}$ . Ainsi, nous associons  $\bar{\varphi}_\mu$  avec un unique  $\bar{s}_\mu \in U_{ad}$  défini par les points de forme  $\bar{A}_i = (\bar{\varphi}_\mu(a_i), i\mu), i = 0, \dots, d$ . D'après les propriétés de base des fonctions de Bézier, nous avons

$$|s'_\mu(x)| \leq L_0 \quad \forall x \in [a, b]; \quad (4.37)$$

et puisque  $\bar{\varphi}_\mu|_{[a_i, a_{i+1}]}$  est tangente à  $s_\mu$  au point  $A_{i+\frac{1}{2}}$ , alors nous avons

$$\exists C > 0 \quad \|\bar{s}_\mu - \bar{\varphi}_\mu\|_{C([a, b])} \leq C\mu \|\bar{s}_\mu\|_{C^1([a, b])} \leq C\mu. \quad (4.38)$$

Par suite, en utilisant l'inégalité triangulaire ainsi que les équations (4.36) et (4.38), nous obtenons

$$\|\bar{s}_\mu - \bar{\varphi}\|_{C([a, b])} \leq \|\bar{s}_\mu - \bar{\varphi}_\mu\|_{C([a, b])} + \|\bar{\varphi}_\mu - \bar{\varphi}\|_{C([a, b])} \xrightarrow{\mu \rightarrow 0} 0.$$

Ceci montre (i).

(ii) Soit  $((s_\mu, u_h(s_\mu)))$  une suite telle que  $u_h(s_\mu)$  est une solution de (4.15) dans  $\Omega_h(s_\mu)$ . D'après la compacité de l'ensemble  $U_{ad}^\mu$ , nous savons qu'il existe une sous suite de  $(s_\mu)_\mu$  notée encore  $(s_\mu)_\mu$  et un élément  $\varphi \in U_{ad}$ , tels que

$$\Omega(s_\mu) \xrightarrow{\mu \rightarrow 0} \Omega(\varphi). \quad (4.39)$$

Nous avons l'estimation d'erreur suivante

$$\|r_h(s_\mu) - s_\mu\|_{C[a, b]} \leq ch. \quad (4.40)$$

Par suite, en utilisant (4.39) et (4.40), nous obtenons :

$$\Omega_h(s_\mu) \xrightarrow{h, \mu \rightarrow 0} \Omega(\varphi).$$

Notons par  $\tilde{u}_h(s_\mu)$  l'extension uniforme de  $u_h(s_\mu)$  de  $H_{\Gamma_d}^1(\Omega_h)$  dans  $H^1(D)$ . D'après le lemme 4.2, nous prouvons que  $\|\tilde{u}_h(s_\mu)\|_{1,D}$  est uniformément bornée. Ainsi, il existe une sous suite telle que

$$\tilde{u}_h(s_\mu) \xrightarrow{h,\mu \rightarrow 0} \tilde{u} \quad \text{dans } H_{\Gamma_d}^1(D), \quad \text{faible} \quad (4.41)$$

De plus, nous avons  $u \in H_{\Gamma_d}^1(D)$ .

Pour conclure, il suffit de montrer que  $u(\varphi)$  satisfait l'équation suivante :

$$\begin{aligned} \int_{\Omega} \lambda \nabla u \cdot \nabla \psi + \int_{\Omega} \mathcal{V}_0 \cdot \nabla u \psi &= \int_{\Omega} f \psi - \int_{\Omega} \lambda \nabla V \cdot \nabla \psi \\ &- \int_{\Omega} \mathcal{V}_0 \cdot \nabla V \psi \quad \forall \psi \in H_{\Gamma_d}^1(\Omega). \end{aligned} \quad (4.42)$$

En effet, soit  $\psi \in H_{\Gamma_d}^1(\Omega(\varphi))$ . Notons par  $\tilde{\psi} \in H^1(D)$  une extension de  $\psi$  définie par

$$\tilde{\psi} = \begin{cases} \psi & \text{dans } \Omega(\varphi) \\ 0 & \text{dans } D \setminus \Omega(\varphi). \end{cases}$$

Ainsi, nous pouvons construire une suite  $(\psi_j)_j$ ,  $\psi_j \in \mathcal{D}(\bar{D})$ , telle que,

$$\text{dist}(\text{supp } \psi_j, \overline{\Gamma_D}) > 0 \quad \forall j \in \mathbb{N} \quad \text{et} \quad \psi_j \rightarrow \tilde{\psi} \quad \text{in } H^1(D), \quad j \rightarrow \infty. \quad (4.43)$$

Notons par  $\pi_h \psi_j$  l'interpolé linéaire par morceaux de  $\psi_j$  dans  $\mathcal{T}(h, s_\mu)$ .

Soit  $j \in \mathbb{N}$ . Puisque  $\Omega_h(s_\mu) \xrightarrow{h,\mu \rightarrow 0} \Omega(\varphi)$ , il existe  $(h_0, \mu_0)$  tel que

$$\pi_h \psi_j|_{\Omega_h} \in H_{\Gamma_d}^1(\Omega_h), \quad \forall h \leq h_0, \quad \mu \leq \mu_0.$$

Pour tout  $h \leq h_0$ ,  $\mu \leq \mu_0$ , nous avons

$$\begin{aligned} \int_{\Omega_h} \lambda_h \nabla \tilde{u}_h \cdot \nabla \pi_h \psi_j + \int_{\Omega_h} \mathcal{V}_0^h \cdot \nabla \tilde{u}_h \pi_h \psi_j &= \int_{\Omega_h} f \pi_h \psi_j - \int_{\Omega_h} \lambda_h \nabla V_h \cdot \nabla \pi_h \psi_j \\ &- \int_{\Omega_h} \mathcal{V}_0^h \cdot \nabla V_h \pi_h \psi_j \end{aligned} \quad (4.44)$$

En passant à la limite dans l'équation (4.44) quand  $h, \mu \rightarrow 0$  et  $j \rightarrow \infty$ , nous obtenons que  $u(\varphi)$  est une solution de l'équation (4.42). En effet, nous avons

$$\int_{\Omega_h} \lambda_h \nabla \tilde{u}_h \cdot \nabla \pi_h \psi_j dx dy - \int_{\Omega} \lambda \nabla \tilde{u} \cdot \nabla \tilde{\psi} dx dy = I_1 + I_2 + I_3 + I_4 \quad (4.45)$$

où

$$I_1 = \int_{\Omega_h} \nabla \tilde{u}_h \cdot (\lambda_h \nabla \pi_h \psi_j - \lambda \nabla \psi_j) dx dy, \quad I_2 = \int_{\Omega_h} \lambda \nabla \tilde{u}_h \cdot (\nabla \psi_j - \nabla \tilde{\psi}) dx dy.$$

$$I_3 = \int_{\Omega_h} \lambda (\nabla \tilde{u}_h - \nabla \tilde{u}) \cdot \nabla \tilde{\psi} dx dy, \quad I_4 = \int_D \lambda (\chi_{\Omega_h} - \chi_{\Omega}) \nabla \tilde{u} \cdot \nabla \tilde{\psi} dx dy.$$

En utilisant l'inégalité de Hölder sur  $I_1$ , nous avons

$$\begin{aligned} |I_1|^2 &\leq \int_{\Omega_h} |\nabla \tilde{u}_h|^2 \int_{\Omega_h} (\lambda_h \nabla \pi_h \psi_j - \lambda \nabla \psi_j)^2 \\ &\leq 2 \int_D |\nabla \tilde{u}_h|^2 \left( \int_D (\lambda_h)^2 (\nabla \pi_h \psi_j - \nabla \psi_j)^2 + \int_D (\lambda_h - \lambda)^2 |\nabla \psi_j|^2 \right) \end{aligned}$$

En utilisant (4.29) et le fait que pour tout  $j \in \mathbb{N}$ ,  $\pi_h \psi_j$  satisfait l'équation (4.28) et que  $\tilde{u}_h$  est uniformément borné dans  $H^1(D)$ , nous aurons

$$\lim_{\substack{h, \mu \rightarrow 0+ \\ j \rightarrow \infty}} I_1 = 0.$$

Pour  $I_2$ , en utilisant la convergence de la suite  $(\psi_j)_j$  dans  $H^1(D)$  (4.43), plus le fait que  $u_h$  est uniformément borné dans  $H^1(D)$ , nous obtenons que

$$\lim_{\substack{h, \mu \rightarrow 0+ \\ j \rightarrow \infty}} I_2 = 0.$$

En utilisant la convergence (4.41) dans  $H^1(D)$ , plus le fait que  $\tilde{\psi} \in L^2(D)$ , nous obtenons :

$$\lim_{\substack{h, \mu \rightarrow 0+ \\ j \rightarrow \infty}} I_3 = 0.$$

La convergence de  $I_4$  vers 0 découle de la convergence des fonctions caractéristiques,

$$\chi_{\mu h} \rightarrow \chi \quad \text{in } L^p(D) \quad \forall p \in [1, \infty[, \quad h, \mu \rightarrow 0+ \quad (4.46)$$

où  $\chi_{\mu h}$  et  $\chi$  sont respectivement les fonctions caractéristiques de  $\Omega(s_\mu)$  et  $\Omega(\varphi)$ .

Par conséquent, nous avons

$$\lim_{\substack{h, \mu \rightarrow 0+ \\ j \rightarrow \infty}} \int_{\Omega_h} \lambda_h \nabla \tilde{u}_h \cdot \nabla \pi_h \psi_j dx dy = \int_{\Omega} \lambda \nabla \tilde{u} \cdot \nabla \tilde{\psi} dx dy \quad (4.47)$$

En adoptant la même décomposition que celle utilisée dans (4.45) pour la différence des autres termes de l'équation (4.42) avec ceux de l'équation (4.44), et en adoptant les mêmes techniques que celles utilisées pour montrer (4.47), nous obtenons la convergence de ces termes.

Ce qui achève la preuve de l'assertion (ii).

(iii) il reste alors à montrer que

$$\lim_{h, \mu \rightarrow 0+} J_h(s_\mu) = J(\varphi).$$

En effet, nous avons

$$\begin{aligned} J_h(s_\mu) - J(\varphi) &= \frac{1}{2} \int_0^{L_x} \left( (T_h(x, L_y) - T_0^h(x, L_y))^2 - (T(x, L_y) - T_0(x, L_y))^2 \right) dx \\ &= \frac{1}{2} \int_0^{L_x} \left( T_h(x, L_y) - T(x, L_y) + T_0(x, L_y) - T_0^h(x, L_y) \right) \\ &\quad \left( T_h(x, L_y) + T(x, L_y) - T_0^h(x, L_y) - T_0(x, L_y) \right) dx \end{aligned}$$

En appliquant l'inégalité de Hölder, nous avons

$$J_h(s_\mu) - J(\varphi) \leq I_1 I_2$$

avec

$$I_1^2 = \int_0^{L_x} \left( T_h(x, L_y) - T(x, L_y) + T_0(x, L_y) - T_0^h(x, L_y) \right)^2 dx$$

et

$$I_2^2 = \int_0^{L_x} \left( T_h(x, L_y) + T(x, L_y) - T_0^h(x, L_y) - T_0(x, L_y) \right)^2 dx$$

D'une part, nous avons

$$\begin{aligned} I_1^2 &\leq 2 \left( \int_0^{L_x} (T_h(x, L_y) - T(x, L_y))^2 dx + \int_0^{L_x} (T_0(x, L_y) - T_0^h(x, L_y))^2 dx \right) \\ &\leq 2 \left( \|T_h - T\|_{0,\Gamma_0}^2 + \|T_0^h - T_0\|_{0,\Gamma_0}^2 \right) \end{aligned}$$

Par suite, en utilisant l'injection compacte de  $H^1(D)$  dans  $L^2(\Gamma_0)$ , nous avons

$$\lim_{h,\mu \rightarrow 0+} \|T_h - T\|_{0,\Gamma_0}^2 = 0$$

et

$$\lim_{h,\mu \rightarrow 0+} \|T_0^h - T_0\|_{0,\Gamma_0}^2 = 0$$

et par conséquent, nous obtenons :

$$\lim_{h,\mu \rightarrow 0+} I_1 = 0. \quad (4.48)$$

D'autre part nous avons

$$I_2^2 \leq 4 \int_0^{L_x} \left( (T_h(x, L_y))^2 + (T(x, L_y))^2 + (T_0^h(x, L_y))^2 + (T_0(x, L_y))^2 \right) dx$$

D'après le lemme 4.2 et la continuité de l'application trace de  $H^1(D)$  dans  $L^2(\Gamma_0)$ , on montre qu'il existe une constante  $c > 0$  indépendante de  $h$  et  $\mu$  telle que

$$I_2^2 \leq c. \quad (4.49)$$

Par conséquent, les relations (4.48) et (4.49) nous permettent de conclure que

$$\lim_{h, \mu \rightarrow 0+} J_h(s_\mu) = J(\varphi).$$

Ceci achève la démonstration de (iii).





# Étude numérique du problème

---

## Sommaire

---

|                                                                                                                     |            |
|---------------------------------------------------------------------------------------------------------------------|------------|
| <b>5.1 Algorithmes déterministes</b>                                                                                | <b>75</b>  |
| 5.1.1 Problème matriciel d'optimisation de forme                                                                    | 76         |
| 5.1.2 Calcul du gradient discret                                                                                    | 77         |
| 5.1.3 Validation de la méthode vis-à-vis d'une solution exacte                                                      | 78         |
| 5.1.4 Validation de la méthode vis-à-vis du modèle physique                                                         | 82         |
| <b>5.2 Algorithmes Génétiques</b>                                                                                   | <b>88</b>  |
| 5.2.1 Introduction                                                                                                  | 88         |
| 5.2.2 Les grandes familles d'Algorithmes Génétiques                                                                 | 89         |
| 5.2.3 Algorithmes Génétiques Simples (GAs)                                                                          | 90         |
| 5.2.4 Algorithme Génétique utilisé (GA)                                                                             | 91         |
| 5.2.5 Résultats obtenus sur des fonctions tests                                                                     | 96         |
| <b>5.3 Résultats obtenus pour le problème d'optimisation de forme</b>                                               | <b>100</b> |
| 5.3.1 Validation de l'algorithme (GA) vis-à-vis d'une solution exacte                                               | 101        |
| 5.3.2 Validation de l'algorithme (GA) vis-à-vis du modèle physique                                                  | 104        |
| <b>5.4 Développement d'un algorithme évolutionnaire</b>                                                             | <b>105</b> |
| 5.4.1 Logique floue                                                                                                 | 105        |
| 5.4.2 Combinaison des (GA) avec les contrôleurs de logique floue (CLF)                                              | 120        |
| 5.4.3 Comparaison de l'algorithme (GA) et l'algorithme (GA-CLF) vis-à-vis du modèle physique                        | 125        |
| <b>5.5 Parallélisation de l'algorithme (GA)</b>                                                                     | <b>127</b> |
| 5.5.1 Analyse du temps d'exécution des (GA) Maître-esclaves                                                         | 127        |
| 5.5.2 Analyse expérimentale du temps d'exécution d'une implémentation parallèle de l'algorithme (GA-CLF)            | 128        |
| <b>5.6 Comparaison de l'algorithme de type gradient avec les algorithmes évolutionnaires sur le modèle physique</b> | <b>131</b> |

---

|       |                                                                              |     |
|-------|------------------------------------------------------------------------------|-----|
| 5.6.1 | Comparaison qualitative entre l'algorithme (GA) et l'algorithme du gradient  | 132 |
| 5.6.2 | Comparaison quantitative entre l'algorithme (GA) et l'algorithme du gradient | 132 |

---

Dans ce chapitre, nous présentons les différentes méthodes d'optimisation utilisées pour la résolution numérique du problème d'optimisation de forme. Nous faisons alors le point sur deux familles d'algorithmes à savoir, les algorithmes déterministes du type gradient et les algorithmes évolutionnaires. Nous commençons par une analyse du calcul du gradient de forme dans le cas discret. Puis nous donnons des résultats numériques obtenus en utilisant la méthode déterministe du type gradient. En ce qui concerne les algorithmes évolutionnaires, nous optons d'abord pour une étude comparative basée sur des fonctions tests "benchmark functions", entre les algorithmes génétiques standards et ceux développés dans le cadre de cette thèse. Puis nous donnons des résultats numériques obtenus en utilisant ces derniers algorithmes pour la résolution de notre problème d'optimisation de forme. Ensuite, nous proposons un développement de ces algorithmes génétiques en utilisant deux techniques. La première consiste à combiner les algorithmes génétiques avec des systèmes de contrôleurs basés sur la logique floue. La deuxième consiste à faire une analyse sur les différentes méthodes de parallélisation des algorithmes génétiques, afin de motiver la méthode de parallélisation choisie. Enfin, nous proposons une étude comparative des deux méthodes d'optimisation utilisant les algorithmes déterministes du type gradient et les algorithmes génétiques améliorés par la logique floue, au niveau qualité de la solution aussi bien qu'au niveau temps de calcul, lorsque les algorithmes génétiques sont parallélisés.

Le travail présenté dans ce chapitre a donné lieu à un article [61] et deux communications [24, 20].

## Introduction

Dans cette partie, nous présentons les différentes méthodes d'optimisation utilisées pour le problème d'optimisation de forme. Ces méthodes ont pour objectif de trouver la ou les solutions  $\Omega$ , représentant la partie solide de la pièce, qui minimise(nt) la fonctionnelle coût  $J(\Omega)$ . Nous distinguons principalement deux familles d'algorithmes : les algorithmes déterministes du type gradient et les algorithmes évolutionnaires.

Les algorithmes déterministes sont des méthodes locales connues pour leur convergence rapide, mais ils présentent l'inconvénient de rester bloqués dans des minima locaux. Ces algorithmes

utilisent l'analyse de sensibilité comme moyen de calculer la pente de la fonction objectif. Cependant, cette analyse s'avère la plupart du temps une tâche très laborieuse - voire parfois non faisable, si la fonction est non dérivable ou lorsque le solveur utilisé ne permet pas d'avoir une analyse de sensibilité performante.

Les algorithmes évolutionnaires sont des méthodes globales de type stochastique et sont beaucoup moins sensibles aux minima locaux. Ces algorithmes sont dits d'ordre zéro car ils ne nécessitent que l'évaluation de la fonction à minimiser. Ce point les rend moins exigeants et plus souples, lorsque l'on souhaite modifier la fonction à minimiser par exemple. Les algorithmes évolutionnaires requièrent un grand nombre d'évaluations de fonction, ce qui peut les ralentir considérablement, surtout quand la fonction à optimiser est coûteuse en temps de calcul. Néanmoins, puisque l'étape la plus coûteuse (l'évaluation de la performance de toute une population) est constituée de calculs totalement indépendants entre eux, ceci rend ces algorithmes faciles à paralléliser. De plus, depuis leur avènement, ces algorithmes n'ont pas cessé de s'améliorer. On trouve deux volets principaux de développement : le premier se focalise sur les mécanismes internes afin d'accélérer la convergence, tandis que le second essaie de perfectionner la qualité de la solution en combinant ces algorithmes avec d'autres méthodes. On parle alors de méthodes hybrides.

Dans ce contexte, un travail commun avec R. Y. Shikhinskaya [61] nous a conduits à développer une approche qui consiste à combiner les algorithmes génétiques avec la logique floue. Cette approche sera expliquée dans la section 5.4.2.

## 5.1 Algorithmes déterministes

Dans cette partie, nous présentons le principe de la méthode déterministe de type gradient appliquée à notre problème d'optimisation de forme discret écrit sous forme matricielle [17, 41, 42]. Pour une présentation détaillée des méthodes déterministes, nous renvoyons à [70, 65].

Les algorithmes déterministes d'optimisation de type gradient consistent à minimiser une fonction  $f(x)$  avec  $x \in \mathbb{R}^n$  en cherchant un point  $\bar{x}$  tel que  $\nabla f(\bar{x}) = 0$ , ceci en construisant une suite minimisante  $\{x_k\}$  telle que  $\liminf \nabla f(x_k) = 0$ .

Il existe d'autres méthodes de calcul du gradient de forme. Nous citons par exemple la méthode due à Hadmard [1] puis développée par Murat-Simon [59] et d'autres auteurs [57, 73, 77, 87]. D'autres méthodes ont été développées récemment, parmi lesquelles le gradient topolo-

gique [78, 37], la méthode d'homogénéisation [60],[4] et la méthode des lignes de niveaux [66],[5].

### 5.1.1 Problème matriciel d'optimisation de forme

Le problème (4.15) peut être réécrit sous la forme suivante

$$\begin{cases} \text{trouver } q(\varphi) = (u_h(\bar{b}_i))_{i=1}^N \in \mathbb{R}^N \\ X(\varphi)q(\varphi) = F(\varphi) \end{cases} \quad (5.1)$$

D'après l'hypothèse (HT<sub>1</sub>), la dimension  $N = \dim H_h(\Omega_h)$  ne dépend pas de  $s_\mu \in U_\mu^{ad}$ . Les éléments  $X_{ij}(\varphi)$  et  $F_i(\varphi)$  de la matrice  $X(\varphi)$  et le vecteur  $F(\varphi)$  sont, respectivement, donnés par

$$\begin{aligned} X_{ij}(\varphi) &= \int_{\Omega_h} \lambda_h \nabla \psi_i \cdot \nabla \psi_j + \int_{\Omega_h} \mathcal{V}_0^h \cdot \nabla \psi_i \psi_j \\ F_j(\varphi) &= \int_{\Omega_h} f \psi_j - \int_{\Omega_h} \lambda_h \nabla V_h \cdot \nabla \psi_j - \int_{\Omega_h} \mathcal{V}_0^h \cdot \nabla V_h \psi_j. \end{aligned}$$

La fonctionnelle coût discrète  $J_h(\varphi)$  s'écrit :

$$J_h(\varphi) = \frac{1}{2} (\hat{X}(\varphi)(q(\varphi) - q_0), q(\varphi) - q_0),$$

où  $\hat{X}(\varphi) = (\hat{X}_{i,j}(\varphi))_{1 \leq i,j \leq N}$  est défini par

$$\hat{X}_{i,j}(\varphi) = \begin{cases} \int_{\Gamma_0^h} \psi_i \psi_j ds, & \text{si } N - d + 1 \leq i, j \leq N \\ 0, & \text{ailleurs.} \end{cases}$$

et  $q_0 = (0, \dots, 0, T_0(\bar{b}_{N-d+1}), \dots, T_0(\bar{b}_N)) \in \mathbb{R}^N$ .

La forme matricielle du problème d'optimisation de forme s'écrit :

$$(P_d) \begin{cases} \inf_{\varphi \in \mathcal{U}} J_h(\varphi), \\ s/c \ X(\varphi)q(\varphi) = F(\varphi). \end{cases} \quad (5.2)$$

La méthode du gradient utilisée pour la résolution numérique du problème (5.2) nécessite le calcul du gradient de  $J$  par rapport à  $\varphi$ . Ceci est développé dans la section suivante.

## 5.1.2 Calcul du gradient discret

Soit  $\mathcal{J} : \mathcal{U} \times \mathbb{R}^N \mapsto \mathbb{R}$  une fonction réelle définie par

$$\mathcal{J}(\varphi, q(\varphi)) = J_h(\varphi),$$

avec  $q(\varphi)$  comme solution du problème (5.1). Pour toute direction  $\beta \in \mathbb{R}^{d+1}$ , nous avons

$$J'_h(\varphi)(\beta) = (\nabla_{\varphi} \mathcal{J}(\varphi, q(\varphi)))^T \beta + (\nabla_q \mathcal{J}(\varphi, q(\varphi)))^T q'(\varphi), \quad (5.3)$$

où  $\nabla_{\varphi} \mathcal{J}$  et  $\nabla_q \mathcal{J}$  désignent, respectivement, les dérivées partielles de  $\mathcal{J}$  par rapport à  $\varphi$  et  $q$ .

Par suite, nous avons :

$$\nabla J_h(\varphi) = \left( \frac{\partial J_h(\varphi)}{\partial \varphi_i} \right)_{i=1}^d = \frac{1}{2} (\hat{X}'(\varphi)(q(\varphi) - q_0), q(\varphi) - q_0) + (\hat{X}(\varphi)(q(\varphi) - q_0), q'(\varphi)),$$

où

$$\hat{X}'(\varphi) = \left( \frac{\partial \hat{X}(\varphi)}{\partial \varphi_i} \right)_{i=1}^d, \quad q'(\varphi) = \left( \frac{\partial q(\varphi)}{\partial \varphi_i} \right)_{i=1}^d$$

Puisque

$$X(\varphi)q(\varphi) = F(\varphi)$$

nous avons

$$X'(\varphi)q(\varphi) + X(\varphi)q'(\varphi) = F'(\varphi),$$

où  $X'(\varphi) = (X'_{ij}(\varphi))_{1 \leq i, j \leq N}$  et  $F'(\varphi) = (F'_j(\varphi))_{1 \leq j \leq N}$ . Les éléments  $X'_{ij}$  et  $F'_j$  peuvent être calculés de la manière suivante :

$$\begin{aligned} X'_{ij}(\varphi) &= (\nabla_{\varphi} X_{ij}(\varphi))^T \beta \quad \forall i, j = 1, \dots, N, \\ F'_j(\varphi) &= (\nabla_{\varphi} F_j(\varphi))^T \beta \quad \forall j = 1, \dots, N, \end{aligned}$$

En utilisant le problème d'état adjoint défini par

$$X^T(\varphi)p(\varphi) = \nabla_q \mathcal{J}(\varphi, q(\varphi)) = \hat{X}(\varphi)(q(\varphi) - q_0), \quad (5.4)$$

ainsi le gradient de  $J_h$  s'écrit :

$$J'_h(\varphi)(\beta) = (\nabla_{\varphi} \mathcal{J}(\varphi, q(\varphi)))^T \beta + p(\varphi)^T (F'(\varphi) - X'(\varphi)q(\varphi)). \quad (5.5)$$

Nous utilisons alors l'algorithme suivant :

**Algorithm 5.1** Algorithme du gradient

1. Choisir un domaine initial  $\Omega(\varphi^0)$ , avec  $\varphi_0 \in U_{ad}$ , une précision désirée  $\varepsilon$ , prendre  $k = 0$ .
2. Calculer la matrice  $X(\varphi^k)$ , le vecteur  $F(\varphi^k)$  et le gradient  $\nabla_q \mathcal{J}(\varphi^k, q(\varphi^k))$ ,
3. Résoudre l'équation d'état et l'équation d'état adjoint

$$X(\varphi^k)q(\varphi^k) = F(\varphi^k)$$

$$X^T(\varphi^k)p(\varphi^k) = \nabla_q \mathcal{J}(\varphi^k, q(\varphi^k))$$

4. Calculer le gradient

$$\nabla J_h(\varphi^k) = \nabla_{\varphi^k} \mathcal{J}(\varphi^k, q(\varphi^k))^T \beta + p(\varphi^k)^T (F'(\varphi^k) - X'(\varphi^k)q(\varphi^k)).$$

5. Recherche  $\rho^k$  qui réalise le  $\min_{\rho > 0} J_h(\varphi^k - \rho^k \nabla J_h(\varphi^k))$
6. Si  $\|\nabla J_h(\varphi^k)\| \leq \varepsilon$ , aller à l'étape 8
7. Faire  $\varphi^{k+1} = \varphi^k - \rho^k \nabla J_h(\varphi^k)$  et aller à l'étape 2
8. Fin

Dans le paragraphe suivant, nous présentons des exemples où la solution analytique exacte est connue i.e. où la frontière libre  $\Gamma$  et température sont connues. Pour chaque exemple, nous allons reconstruire la frontière libre  $\Gamma$ , en utilisant la formulation d'optimisation de forme du problème (4.19), les courbes de Bézier et l'algorithme 5.1. Dans tous les exemples numériques qui suivent, nous considérons que la pièce à souder  $D$  est telle que :  $L_x = 1$ ,  $L_y = 1$ .

### 5.1.3 Validation de la méthode vis-à-vis d'une solution exacte

Dans les deux exemples suivants nous allons approcher  $T(x, y) = \exp(x + y)$  et  $\Gamma$ , solution du problème (4.19) défini avec les données suivantes :

$$\lambda = 1.0, \quad \mathcal{V}_0 = \begin{pmatrix} -1 \\ 0 \end{pmatrix}, \quad f(x, y) = -3 \exp(x + y),$$

$$\frac{\partial T}{\partial \nu} = \exp(x + y) \text{ dans } \Gamma_0, \quad \frac{\partial T}{\partial \nu} = -\exp(x + y) \text{ dans } \Gamma_1 \cup \Gamma_2 \cup \Gamma_3 \quad (5.6)$$

et  $T_f = T_0 = T_d = \exp(x + y).$

La fonctionnelle coût est définie par :

$$J(\Gamma) = \frac{1}{2} \int_0^1 |T(x, 1) - \exp(x + 1)|^2 dx.$$

Dans le premier exemple nous supposons que la frontière exacte est définie par

$$\Gamma = \{ (x, y) / y = 5(x - 0.35)(0.65 - x), \quad x \in [0.35, 0.65] \}. \quad (5.7)$$

La particularité est que cette frontière est complètement décrite par une courbe de Bézier avec les points de contrôles suivants :

$$P_1 = (0.35, 0.0), \quad P_2 = (0.35, 0.15), \quad P_3 = (0.65, 0.0), \quad P_4 = (0.65, 0.15).$$

Dans la figure 5.1, nous avons représenté la variation du coût en fonction du nombre d'itérations. Quant à la figure 5.2, elle représente l'évolution des frontières libres en fonction du nombre d'itérations.

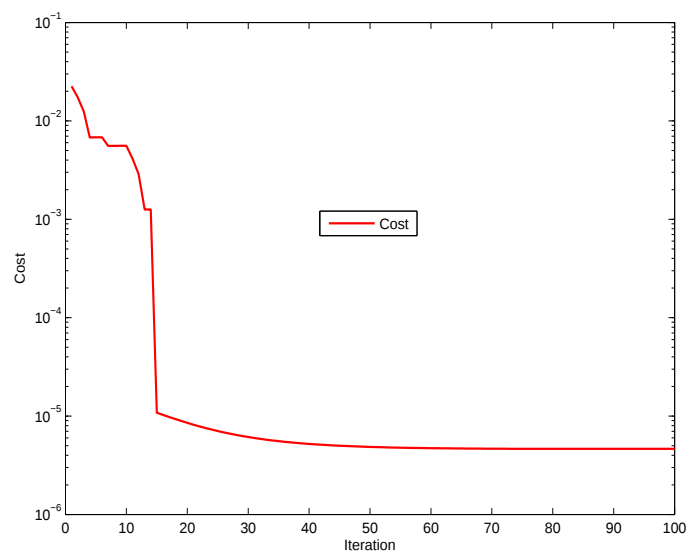


FIGURE 5.1 : Exemple 1 : décroissance de la fonction coût.

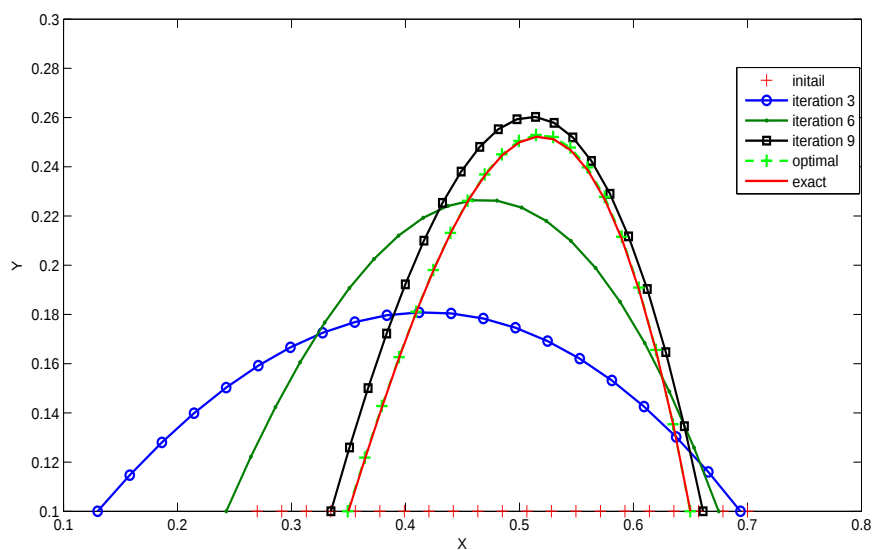


FIGURE 5.2 : Exemple 1 : convergence des frontières.

Pour le deuxième exemple, nous gardons les mêmes données (5.6) que pour le premier exemple, mais nous changeons la frontière libre. La frontière  $\Gamma$  est donnée par :

$$\Gamma = \{ (x, y) / y = 0.15 * \sin((x - 0.35) * \pi / 0.3), x \in [0.35, 0.65] \} \quad (5.8)$$

La particularité de cette frontière est qu'elle ne peut pas être représentée par une courbe de Bézier.

La variation du coût en fonction du nombre d'itérations est représentée dans la figure 5.3. Ainsi, nous avons représenté l'évolution des frontières libres en fonction du nombre d'itérations dans la figure 5.4.

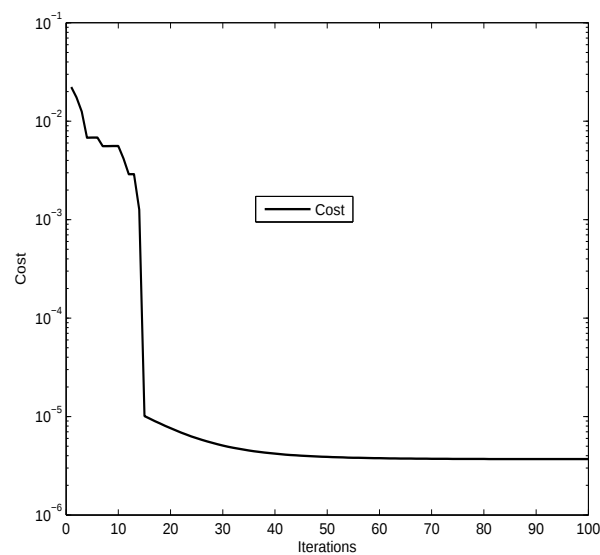


FIGURE 5.3 : Exemple 2 : décroissance de la fonction coût.

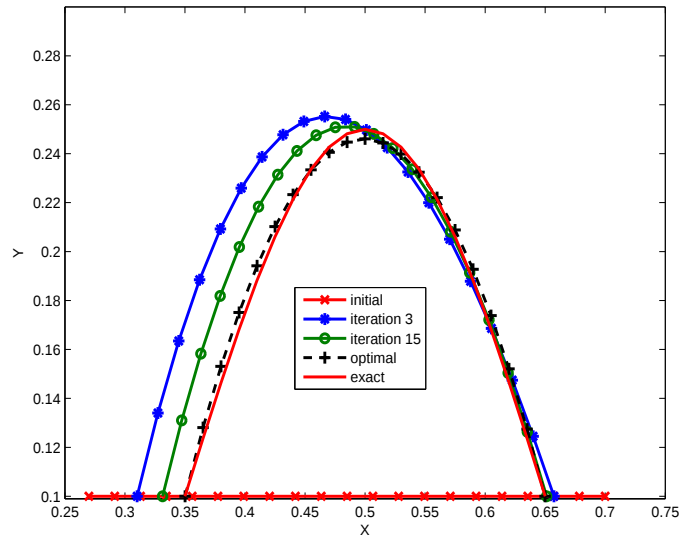


FIGURE 5.4 : Exemple 2 : convergence des frontières.

#### 5.1.4 Validation de la méthode vis-à-vis du modèle physique

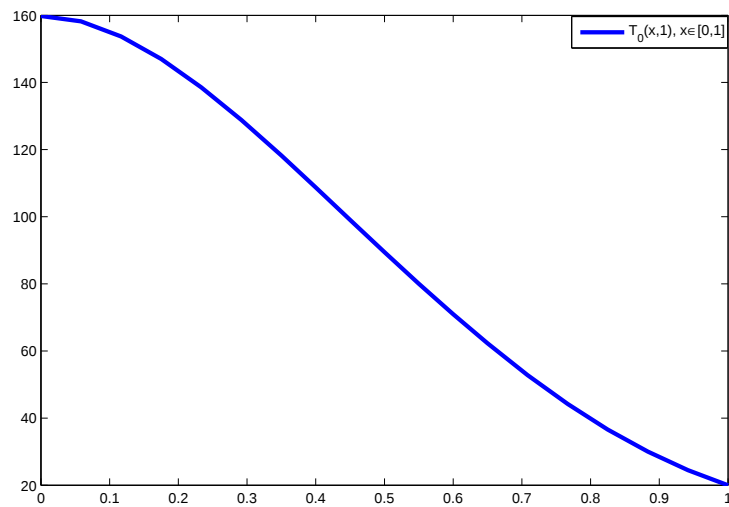
À présent, nous considérons notre modèle ( 4.19), avec les paramètres physiques suivants correspondant à une variante d'aluminium (Al-3.5 wt% Ni) [83] :

$$\lambda = 0.221 \text{ kJ} \cdot (\text{K} \cdot \text{m} \cdot \text{S})^{-1}, \quad \nu_0 = \rho C_p v_{torch}, \quad \rho = 2.37 \times 10^3 \text{ kg} \cdot \text{m}^{-3}$$

$$C_p = 0.124 \text{ kJ} \cdot (\text{kg} \cdot \text{K})^{-1}, \quad \vec{v}_{torch} = -30 \text{ mm} \cdot \text{s}^{-1} \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \quad f(x, y) = 0, \quad (5.9)$$

$$\frac{\partial T}{\partial \nu} = 0 \quad \text{on} \quad \Gamma_0 \cup \Gamma_1 \cup \Gamma_2 \cup \Gamma_3 \quad \text{et} \quad T_f = 659.25 \text{ } ^\circ\text{C}, \quad T_d = 20 \text{ } ^\circ\text{C}$$

La température  $T_0$  sur  $\gamma_0$  est représentée dans la figure 5.5.

FIGURE 5.5 : Température mesurée sur  $\Gamma_0$ 

Dans la figure 5.6 nous représentons le maillage du domaine initial constitué de 648 éléments et 361 degrés de libertés.

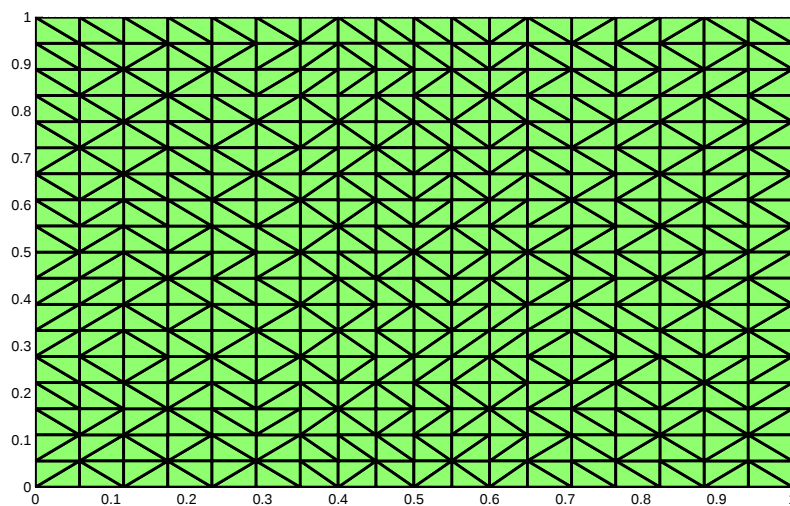


FIGURE 5.6 : Maillage du domaine initial

Le maillage du domaine optimal est représenté dans la figure 5.7. Ce domaine optimal est atteint au bout de 94 itérations avec un coût égal à  $4.810^{-8}$ .

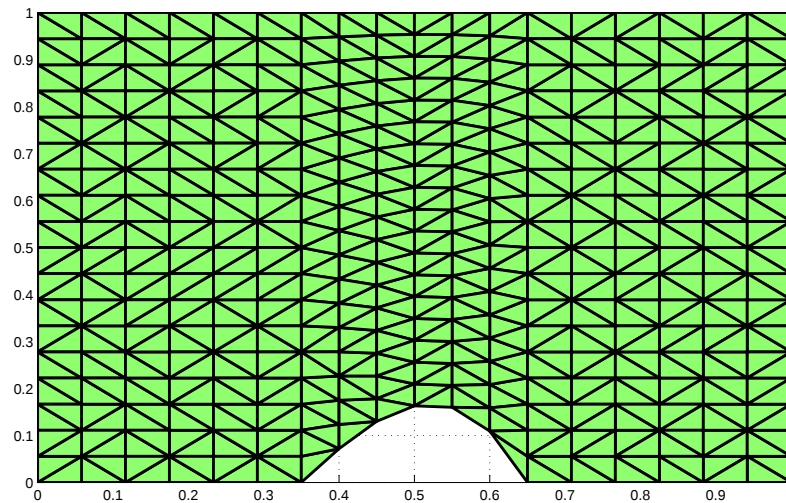


FIGURE 5.7 : Maillage du domaine optimal

Dans les figures 5.8 et 5.9 nous avons représenté les isovaleurs de la température (contour et remplissage). Leur détermination nous donne une idée sur la propagation de la température dans la plaque. Dans la figure 5.10, nous avons représenté la variation de la température en fonction de  $x$  et de  $y$ . La connaissance de la température dans chaque point de la plaque nous permet de prédire les effets mécaniques.

La variation du coût en fonction du nombre d'itérations est représentée sur la figure 5.11. Nous remarquons que la courbe devient presque stationnaire à partir de la vingtième itération.

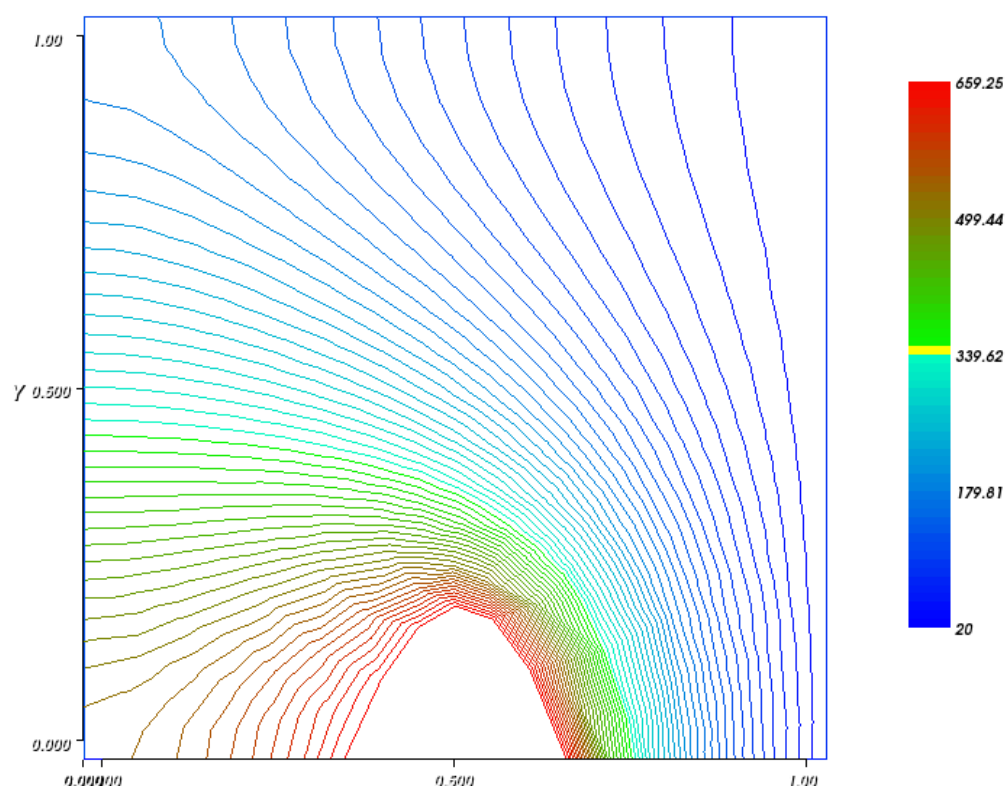


FIGURE 5.8 : Contour des isovalues

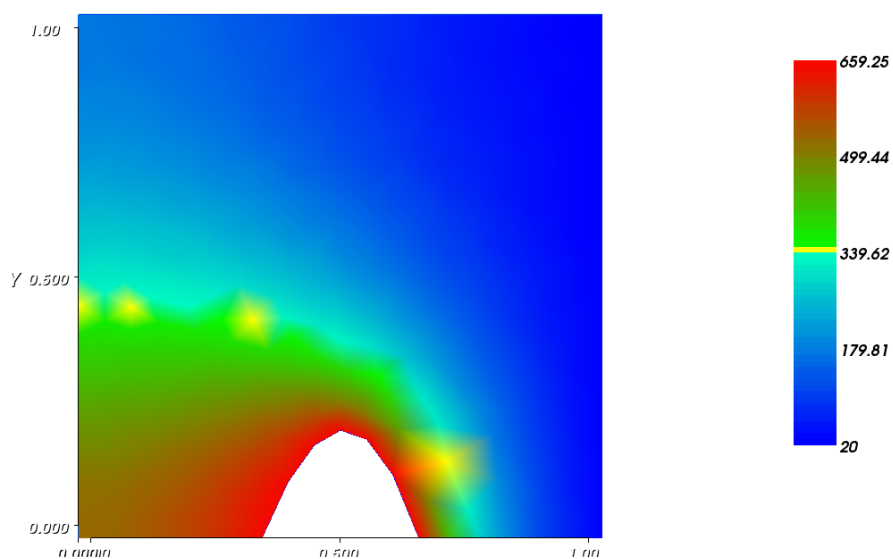
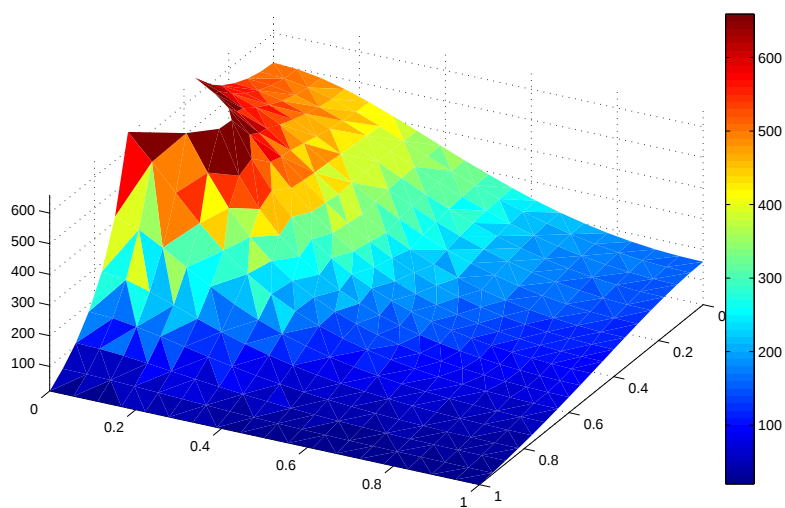


FIGURE 5.9 : Remplissage des isovaleurs

FIGURE 5.10 : Variation de la température en fonction de  $x$  et  $y$

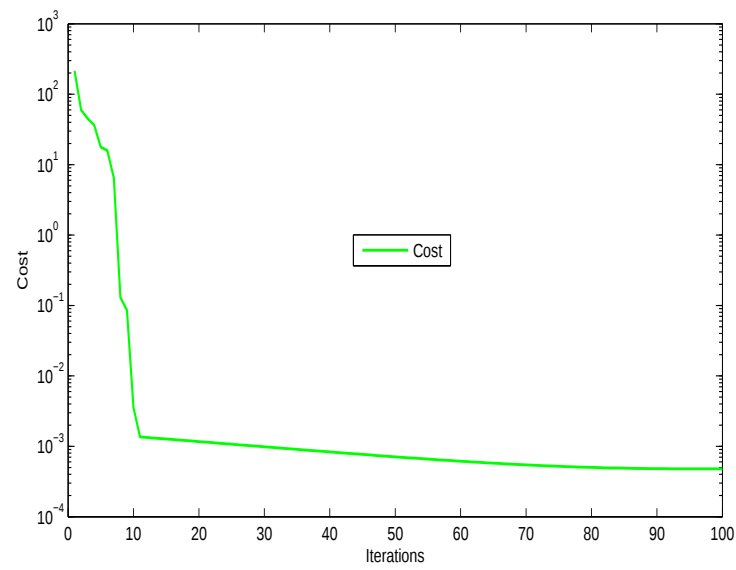


FIGURE 5.11 : Exemple 3 : décroissance de la fonction coût.

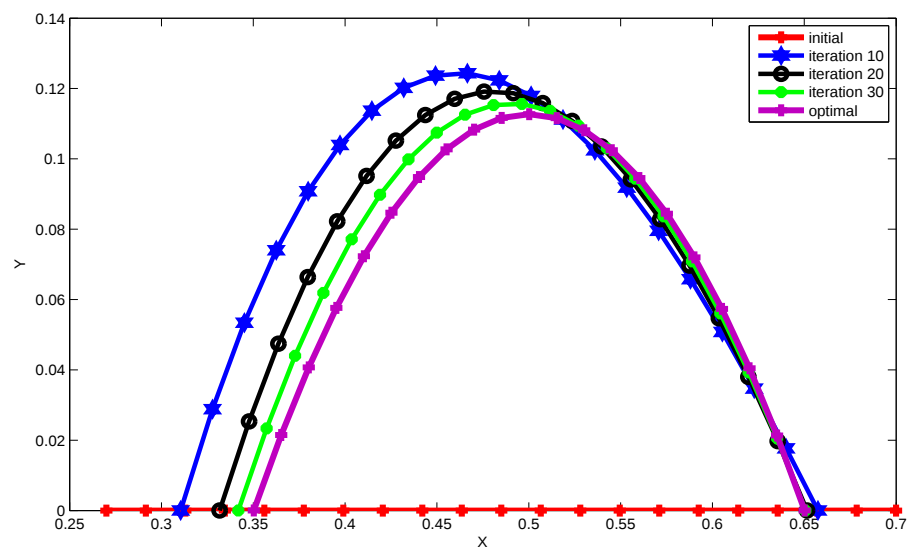


FIGURE 5.12 : Exemple 3 : convergence des frontières.

Dans la figure 5.12, nous avons représenté la frontière initiale et la forme optimale, ainsi que

la frontière aux dixième, vingtième et trentième itérations. Une grande variation de la frontière initiale a été observée au bout de dix itérations. À la trentième itération, la frontière obtenue est très proche de la frontière optimale. Le reste des itérations pour atteindre la solution est consacré à l'ajustement des points proches des extrémités libres de la frontière.

## 5.2 Algorithmes Génétiques

### 5.2.1 Introduction

Le but d'un algorithme évolutionnaire est d'optimiser une fonction  $f$  dite *fonction objectif* sur un espace de recherche. Pour cela, une population d'individus, typiquement un P-uplet de points de l'espace de recherche, évolue selon un darwinisme artificiel (reproduction, mutation, sélection naturelle) basé sur la *fitness*  $F$  de chaque individu. La fitness est directement liée à la valeur de la fonction objectif  $f$  de cet individu (exemple, la fonction  $f$  elle-même).

Des opérateurs appliqués à la population permettent de créer de nouveaux individus (croisement et mutation) et de sélectionner les individus de la population qui vont survivre (sélection et remplacement). Les opérateurs appliqués à un individu ne sont pas en général définis sur le même espace que celui sur lequel est définie la fonction fitness, appelé *espace des phénotypes*, mais sur un espace de *représentation* appelé *l'espace des génotypes*. Par exemple, pour un codage binaire les algorithmes génétiques simples utilisent un espace de génotypes de la forme  $\{0,1\}^n$  ( $n$  est la dimension de l'espace de recherche).

La Figure 5.13 illustre le schéma général d'un algorithme évolutionnaire : après l'initialisation de la population (généralement d'une façon aléatoire) l'algorithme évalue la fitness de chaque individu. La boucle de l'algorithme suit les étapes suivantes :

1. *Critère d'arrêt* : un des critères simples souvent utilisé est lorsque le nombre maximum de générations, fixé par l'utilisateur, est atteint.
2. *Sélection* : cet opérateur sélectionne parmi les parents ceux qui vont générer des enfants. Plusieurs opérateurs sont possibles qui peuvent être soit déterministes soit stochastiques. La sélection est basée sur la fitness des individus.
3. *Création de nouveaux individus* : la création de nouveaux individus se fait essentiellement à l'aide des opérateurs de croisement et de mutation. L'opérateur de croisement est un opérateur stochastique qui combine  $k$  parents pour créer un ou plusieurs enfants. L'opé-

rateur de mutation est un opérateur stochastique qui modifie un individu pour en créer un autre qui lui est généralement proche (ce qui dépend énormément de la représentation choisie).

4. *Evaluation* : calcul de la fitness de chaque enfant. C'est l'étape la plus coûteuse en temps de calcul.
5. *Remplacement* : on détermine qui, parmi la population courante, fera partie des parents de la génération suivante. Cet opérateur est basé, comme l'opérateur de sélection, sur la fitness des individus.

### 5.2.2 Les grandes familles d'Algorithmes Génétiques

Tous ces algorithmes ont en commun le fait de faire évoluer des populations d'individus. La différence entre eux est principalement d'ordre historique. On peut classer ces algorithmes en 4 grandes familles. Pour une description plus détaillée de ces algorithmes, nous renvoyons à [58].

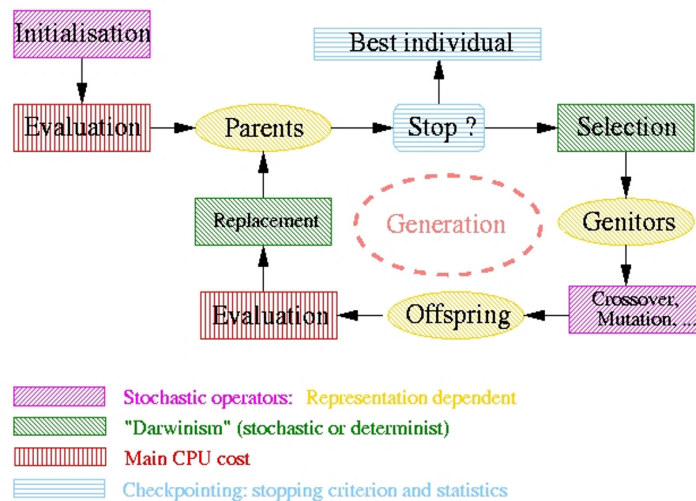


FIGURE 5.13 : Schéma général d'un algorithme évolutionnaire.

- Algorithmes Génétiques (GA) : [47, 38]. Les plus connus et les plus populaires des algorithmes évolutionnaires. Ils ont été développés pour modéliser l'adaptation des populations en biologie.
- Stratégies d'Evolution (ES) : [72, 40, 6]. Développés par des ingénieurs pour résoudre des problèmes d'optimisation paramétriques. Ces algorithmes sont les plus efficaces pour ce type de problèmes.
- Programmation Evolutionnaire (EP) : L. J. Fogel (1966). Développée à l'origine pour la découverte d'automates à états finis (cf. [35]).
- Programmation Génétique (GP) : J. Koza (1990). Apparue initialement comme un sous-domaine des (GAs), la programmation génétique est devenue une branche à part entière. La spécificité de ces algorithmes est de représenter des individus par des arbres (cf. [50]).

Dans la section suivante, nous donnons d'abord l'exemple d'un Algorithme Génétique Simple (GAs) utilisant des opérateurs basiques pour le croisement, la mutation et la sélection. Ensuite nous présentons un algorithme génétique (GA), développé dans le cadre de cette thèse.

### 5.2.3 Algorithmes Génétiques Simples (GAs)

Les (GAs) standards utilisent un codage binaire avec un espace de génotype de la forme  $\{0, 1\}^n$ . Plusieurs opérateurs de sélection ont été développés par différents auteurs. Nous citons ici :

- la sélection à roulette
- la sélection stochastique universelle.

Nous notons par  $P_{X_p}$  la probabilité de sélectionner un individu  $X_p$ .

Pour la roulette,  $P_{X_p}$  est donnée par :

$$P_{X_p} = \frac{F(X_p)}{\sum_{i \in Population} F(X_i)}.$$

Pour la sélection stochastique universelle, la probabilité  $P_{X_p}$  est proportionnelle à sa *fitness*  $F(X_p)$ .

Le mécanisme de croisement le plus simple consiste à échanger les gènes de chaque parent entre le site sélectionné et la position finale  $n$  des deux chaînes, comme le montre la figure 5.14. Par exemple, deux parents

$$X_1 = (x_1^1, x_1^2, \dots, x_1^n) \quad \text{et} \quad X_2 = (x_2^1, x_2^2, \dots, x_2^n),$$

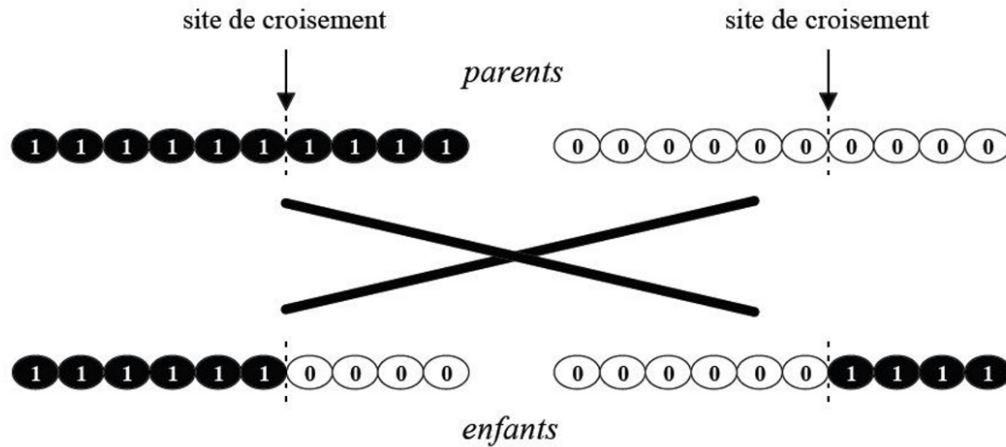


FIGURE 5.14 : Principe de croisement simple

permettent de générer deux enfants

$$Y_1 = (x_1^1, x_1^2, \dots, x_1^q, x_2^{q+1}, \dots, x_2^n) \text{ et } Y_2 = (x_2^1, x_2^2, \dots, x_2^q, x_1^{q+1}, \dots, x_1^n),$$

où l'entier  $q$  est choisi aléatoirement dans  $[1, n]$ .

L'un des points faibles du codage binaire est qu'un tel croisement peut "mélanger" des variables de types différents. La mutation dans le cas d'un codage binaire consiste à changer un 0 par 1, ou inversement (voir Figure 5.15). Dans le cas d'un codage réel, la mutation peut se faire en remplaçant une variable  $x_i$  par  $x_i + \delta x_i$  où  $\delta x_i$  est une "petite" variation de la variable  $x_i$ .

#### 5.2.4 Algorithme Génétique utilisé (GA)

Pour résoudre le problème d'optimisation de forme, nous avons développé un algorithme qui utilise un codage réel, écrit en Fortran 90. Pour la sélection, après un certain nombre d'essais, nous avons choisi de nous restreindre aux trois types qui sont connus pour leurs succès : la sélection à roulette, la sélection à roulette par reste stochastique et la sélection par tournoi. Les croisements employés sont le croisement multi-points et le croisement barycentrique.

Dans le cas d'un croisement multi-points avec deux points (voir figure 5.16), le croisement

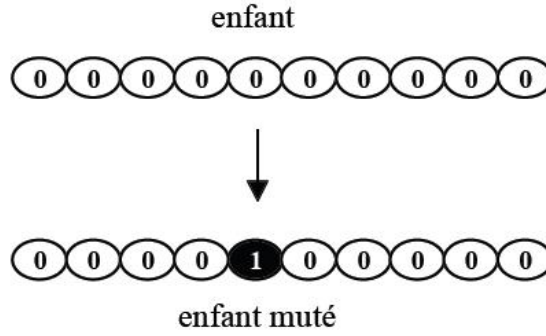


FIGURE 5.15 : Principe de Mutation simple

se fait de la manière suivante : deux parents

$$X_1 = (x_1^1, x_1^2, \dots, x_1^n) \text{ et } X_2 = (x_2^1, x_2^2, \dots, x_2^n)$$

génèrent deux enfants

$$Y_1 = (x_1^1, x_1^2, \dots, x_1^{q_1}, x_2^{q_1+1}, \dots, x_2^{q_2}, x_1^{q_2+1}, \dots, x_1^n)$$

et

$$Y_2 = (x_2^1, x_2^2, \dots, x_2^{q_1}, x_1^{q_1+1}, \dots, x_1^{q_2}, x_2^{q_2+1}, \dots, x_2^n),$$

où les entiers  $q_1$  et  $q_2$  sont choisis aléatoirement dans  $[1, n]$ .

Avec un croisement barycentrique, les enfants créés sont donnés par  $Y_1 = (y_1^1, \dots, y_1^n)$  et  $Y_2 = (y_2^1, \dots, y_2^n)$  avec pour tout  $i$ ,  $y_1^i = \alpha x_1^i + (1 - \alpha)x_2^i$  et  $y_2^i = \alpha x_2^i + (1 - \alpha)x_1^i$  où  $\alpha$  est un réel dans  $[0, 1]$ . Ce nombre  $\alpha$  est soit choisi par l'utilisateur, soit choisi aléatoirement dans le cas d'un croisement barycentrique aléatoire.

Deux types de mutations sont possibles. La mutation gaussienne avec une variance constante ou une variance décroissante au cours des itérations, et la mutation non-uniforme [58] où une variable  $x_i \in [x_{\min_i}, x_{\max_i}]$  prend la nouvelle valeur  $x'_i$

$$x'_i = \begin{cases} x_i + \Delta(t, x_{\max_i} - x_i) & \text{si } s \leq 0.5, \\ x_i - \Delta(t, x_i - x_{\min_i}) & \text{si } s \geq 0.5, \end{cases}$$

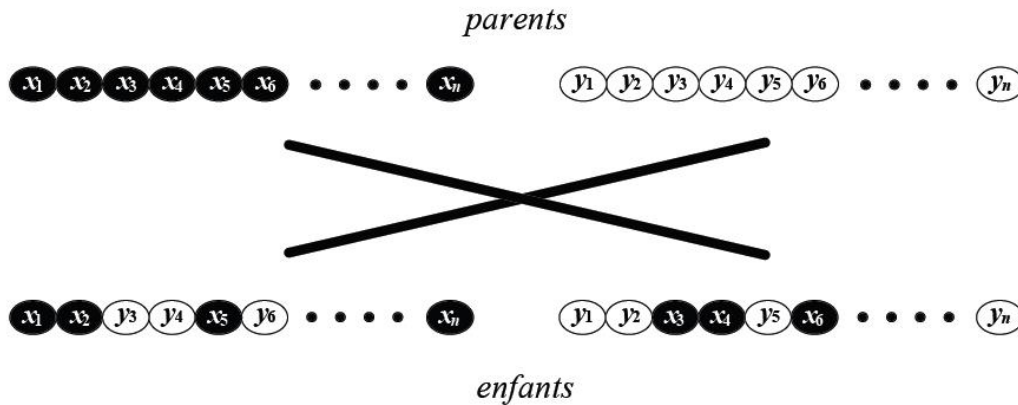


FIGURE 5.16 : Principe de croisement multi-points

où  $t$  désigne le nombre de générations,  $s$  un nombre aléatoire dans  $[0, 1]$  et où la fonction  $\Delta(t, x)$  est définie de la manière suivante :

$$\Delta(t, x) = x \cdot r \cdot \left(1 - \frac{t}{T}\right)^b,$$

telle que  $r$  est un nombre aléatoire dans  $[0, 1]$ ,  $T$  le nombre maximal de générations et  $b$  un paramètre de raffinement. Le graphe de la fonction  $\Delta(t, x)$  est représenté sur la figure 5.17. Ainsi, nous remarquons qu'au début l'amplitude maximale de la mutation est grande, tandis que qu'elle très petite vers la fin de l'algorithme : nous passons alors d'une recherche globale à une recherche locale au cours de l'algorithme. Pour déterminer la stratégie de compromis entre l'exploration et l'exploitation, nous nous servons du paramètre  $b$ , qui permet d'ajuster l'amplitude de la mutation. En fait, pour une grande valeur de  $b$  (Figure 5.18 (a)), c'est l'exploration de l'espace qui est favorisée, alors que pour une petite valeur de  $b$  (Figure 5.18 (b)), c'est la phase d'exploitation et de recherche locale.

L'un des points fondamentaux du processus des algorithmes génétiques est la diversité. Grâce à elle, ces algorithmes réussissent à échapper à des minima locaux.

Pour assurer la décroissance de la fonction à minimiser, génération après génération, on tient toujours à sélectionner l'individu (soit un enfant, soit un parent) qui lui donne une valeur minimale. C'est ce qu'on appelle l'élitisme.

Pour le bon fonctionnement de l'algorithme, nous utilisons ce qu'on appelle la mise en échelle

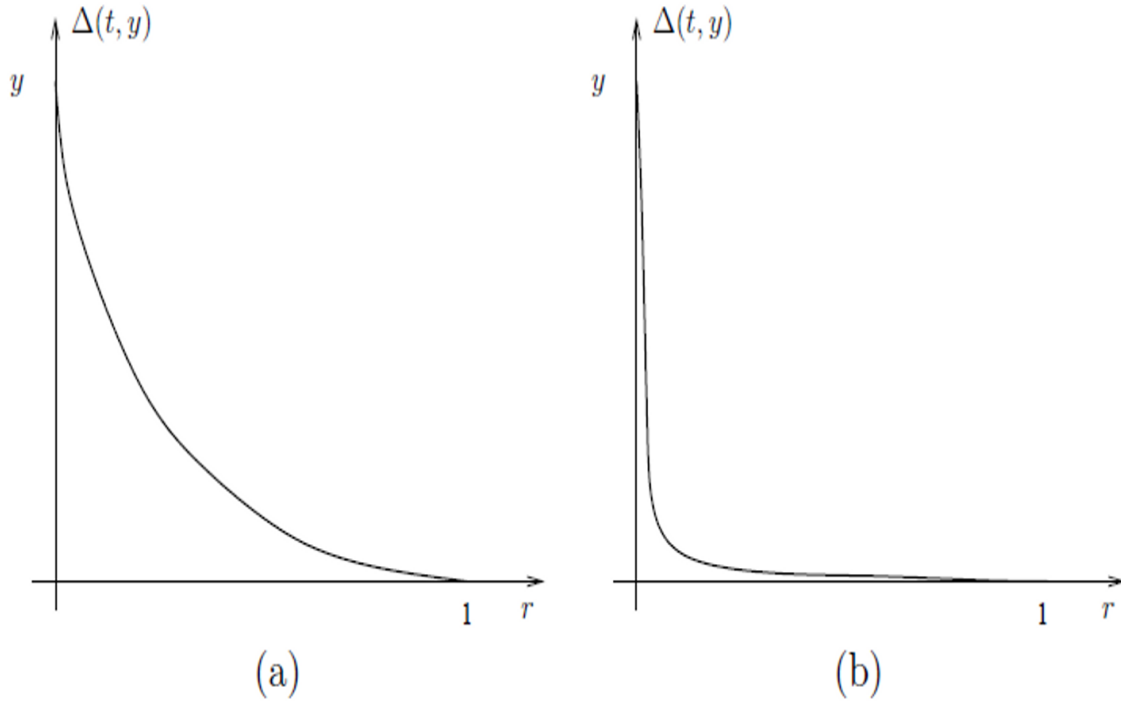


FIGURE 5.17 : La fonction  $\Delta(t, y)$  à deux instants  $t_1$  et  $t_2$  (resp. (a) et (b)) avec  $t_1 < t_2$ .

de la fitness (scaling). Cette technique est imposée par le choix de la méthode de sélection basée sur la méthode de la roulette. Pour cette méthode, la probabilité  $P_{X_p}$  de sélectionner un individu  $X_p$  est proportionnelle à sa fitness  $F(X_p)$  :

$$P_{X_p} = \frac{F(X_p)}{\sum_{i \in Population} F(X_i)}.$$

Cette formule oblige à avoir une fitness positive pour tous les individus. Ce qui suggère d'utiliser un scaling qui consiste à prendre pour fitness

$$F(X_p) = f(X_p) - \min_{i \in Population} f(X_i).$$

Cependant, ce choix reste insuffisant. En effet, si on considère la fonction  $g = f + C$  où  $C$  est une constante plus grande que les valeurs de  $f$ , alors les individus ne sont plus réellement différents au vu de leur fitness donnée avec la fonction  $g$ . Dans ce cas la sélection se fait différemment pour les fonctions  $f$  et  $g$ , et devient sensible aux translations. Pour éviter ce problème une autre condition est imposée pour le choix du scaling. Cette condition consiste à empêcher que la sélection pour

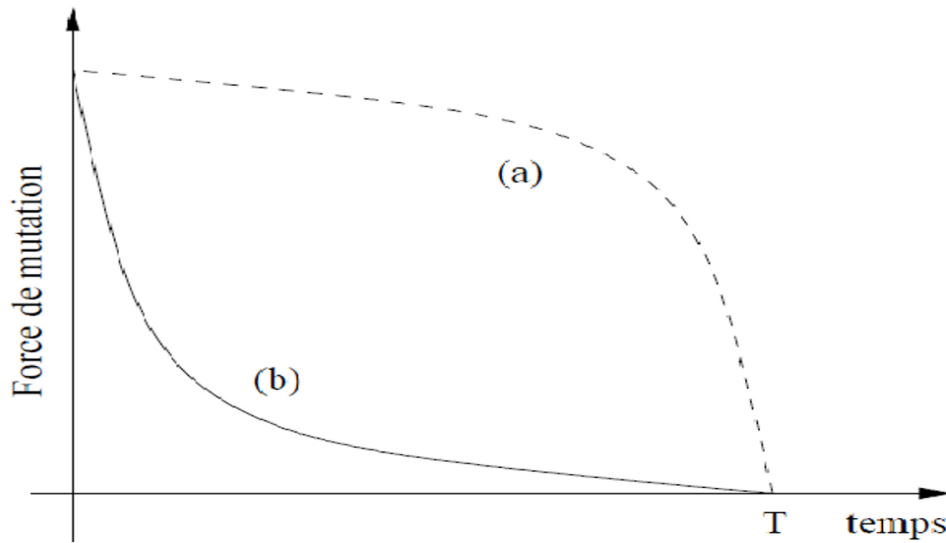


FIGURE 5.18 : Evolution de l'amplitude de la mutation en fonction du temps.

la génération suivante soit restreinte au super-individu (celui qui donne la meilleure valeur dans une génération). Plusieurs stratégies de scaling satisfaisant à ces exigences ont été testées. La meilleure pour nos tests est le *sigma scaling tronqué* :

$$F(X_p) = f(X_p) + (\bar{f} - c\sigma),$$

où  $\bar{f}$  et  $\sigma$  représentent respectivement la moyenne et la variance de la fonction  $f$  sur la population, et où  $c$  est une constante généralement prise entre 1 et 5. Les valeurs négatives éventuelles de  $F(X_p)$  sont tronquées et ramenées à zéro.

Les techniques discutées ci-dessus ont permis d'accroître les performances de l'algorithme développé. L'inconvénient est que le nombre de paramètres à ajuster a augmenté, ce qui induit un réglage plus lent. Le temps de calcul requis par l'algorithme est essentiellement dû au temps d'évaluation d'une fonction (voir Figure 5.13). Pour le problème étudié (problème d'optimisation de forme), ce temps est assez important. En effet, chaque évaluation de la fonction  $f$  correspond à la résolution d'un système linéaire. L'étape de l'évaluation de la fonction se fait indépendamment pour les différents individus de la population, ce qui la rend facilement parallélisable. Ainsi, une version parallèle de l'algorithme a été développée en utilisant la bibliothèque MPI [67]. A

chaque génération, l'évaluation de la fitness des nouveaux individus est répartie sur plusieurs processeurs. Le reste des tâches est géré par le programme principal (ou programme maître) sur un seul processeur. Ceci a permis un gain considérable en temps de calculs. Nous reviendrons en détails sur la parallélisation dans la section 5.5.

### 5.2.5 Résultats obtenus sur des fonctions tests

Cette partie s'attache à présenter les résultats des simulations obtenues par les algorithmes évolutionnaires sur quatre fonctions tests issues de la littérature [31]. Nous avons choisi ces exemples en vue d'illustrer le fonctionnement des algorithmes présentés, mais aussi de manière à dégager les caractéristiques propres à chacun. Précisons que, pour chaque méthode, nous nous sommes évertués à ajuster les paramètres utilisés. Cet effort a été utile a posteriori pour l'ajustement des paramètres sur le problème d'optimisation de forme. Nous avons opté pour trois fonctions unimodales (c'est-à-dire avec un seul minimum) et une fonction multimodale. Pour chacune de ces fonctions, nous fournissons les résultats de deux algorithmes stochastiques :

- un algorithme génétique simple (GAS),
- un algorithme génétique développé (GA).

#### 5.2.5.1 La fonction Sphère

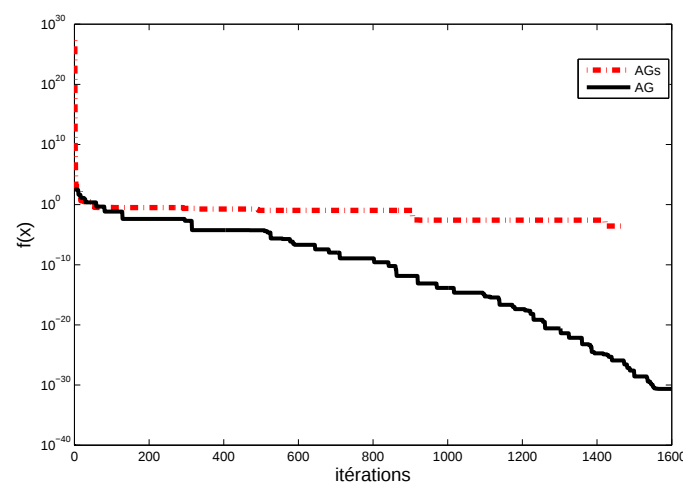


FIGURE 5.19 : Minimisation de la fonction sphère par (GAs) et (GA)

Dans un premier temps, nous allons considérer la minimisation de la fonction Sphère définie de la façon suivante :

$$f(x) = \sum_{i=1}^{30} x_i^2.$$

Ce cas présente un intérêt à la fois théorique et numérique pour les algorithmes évolutionnaires. Sur la Figure 5.19 il apparaît que l'algorithmes (GA) d'un côté, et l'algorithme (GAs) de l'autre côté, affichent deux comportements différents. Ainsi, nous remarquons que, a contrario de l'algorithme (GAs), l'algorithme (GA) est capable d'augmenter la précision de la solution cherchée. Ceci est dû à la variation de la force de mutation en fonction du temps pour l'algorithme (GA).

### 5.2.5.2 La fonction Elliptique

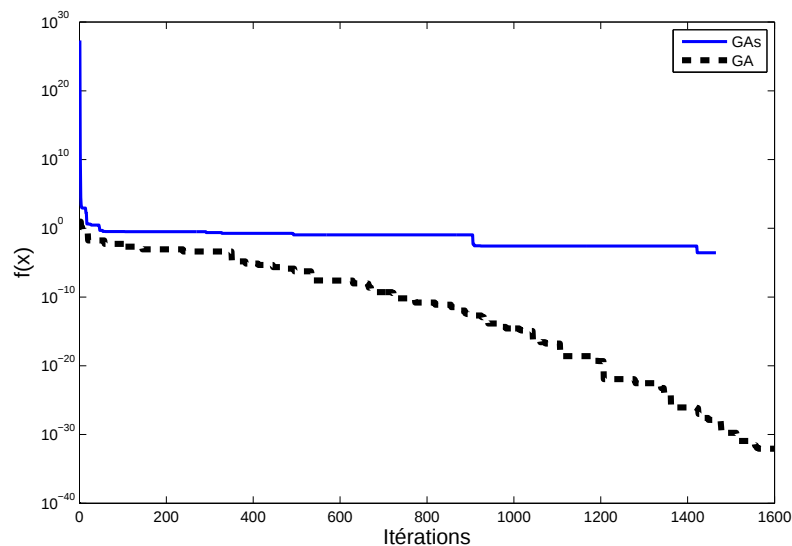


FIGURE 5.20 : Minimisation de la fonction elliptique par (GAs) et (GA)

La fonction elliptique est une variante de la fonction Sphère telle que les variables contribuent d'une façon inégale à la valeur de la fonction :

$$f(x) = \sum_{i=1}^6 1.5^{i-1} x_i^2.$$

Nous constatons sur la Figure 5.20 que l'algorithme (GA) surmonte la difficulté supplémentaire et arrive à converger vers l'optimum global. Tandis que l'algorithme (GAs) est incapable de

trouver l'optimum global (convergence prématurée).

### 5.2.5.3 La fonction de Rosenbrock

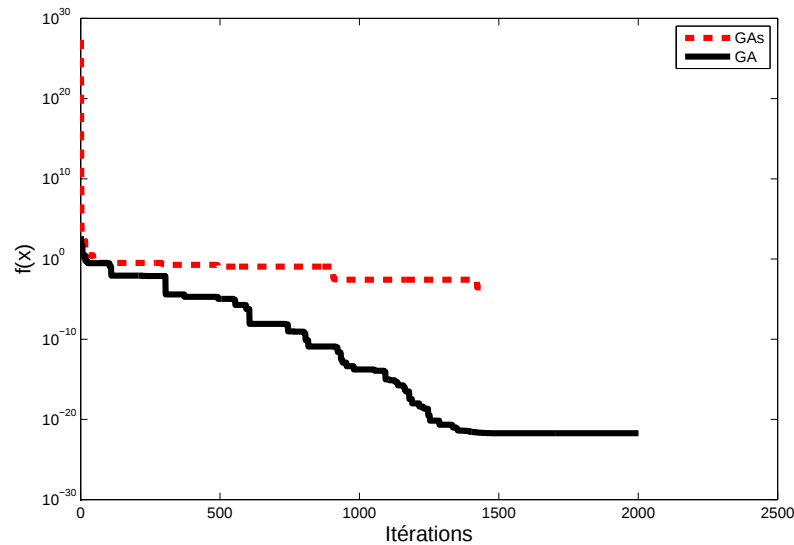


FIGURE 5.21 : Minimisation de la fonction Rosenbrock par (GAs) et (GA)

La fonction de Rosenbrock est donnée par :

$$f(x) = 100(x_1^2 - x_2)^2 + (1 - x_1)^2,$$

et son optimum unique est le point  $(1, 1)$ . Elle présente une "large vallée" autour de ce minimum. Cette fonction offre aux algorithmes évolutionnaires un exemple idéal où un équilibre entre la phase d'exploration et la phase d'exploitation doit être trouvé : l'algorithme doit d'abord explorer la vallée, pour ensuite converger localement vers l'optimum. Par ailleurs, la Figure 5.21 nous révèle que, contrairement à l'algorithme (GAs), l'algorithme (GA) est capable d'améliorer constamment sa performance.

## 5.2.5.4 La fonction de Shekel

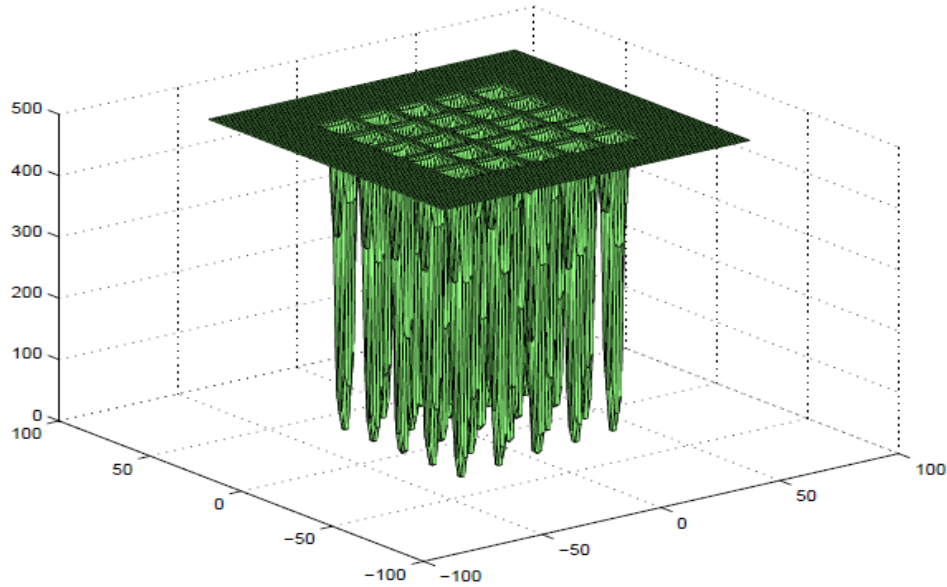


FIGURE 5.22 : Fonction de Shekel

En dimension 2, la fonction de Shekel est définie de la façon suivante :

$$f(x) = \frac{1}{0.002 + \sum_{j=1}^{25} \frac{1}{j + \sum_{i=1}^2 (x_i - a_{ij})^6}},$$

et présente 25 minima dans le carré  $[-64, 64]^2$ . Les valeurs de la fonction en ces minima s'étalent entre 1 et 25 (voir Figure 5.22). Si l'on veut optimiser cette fonction, un algorithme évolutionnaire doit posséder trois propriétés :

1. être capable de trouver la bonne vallée contenant l'optimum global (c'est la phase d'exploitation).
2. pouvoir y rester en conservant l'individu qui s'y trouve d'une génération à la génération suivante.
3. avoir la capacité d'effectuer efficacement la recherche locale dans cette vallée pour arriver jusqu'au minimum global.

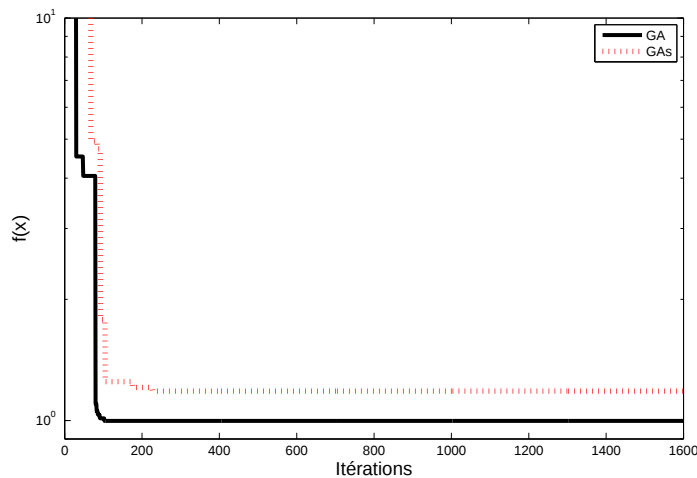


FIGURE 5.23 : Minimisation de la fonction Shekel par GAS et GA

La Figure 5.23 montre que l'algorithme (GA) arrive à atteindre l'optimum global, alors que (GAS) est incapable d'accéder à la bonne vallée.

Selon l'ensemble des résultats obtenus sur les cas tests, l'algorithme (GA) est plus efficace qu'un simple (GAS) : il converge là où un simple (GAS) ne converge pas. Bien plus, dans les cas où ce dernier converge, l'algorithme (GA) offre une meilleure précision.

Dans la section suivante nous présentons les résultats obtenus, en appliquant l'algorithme (GA) pour résoudre le problème d'optimisation de forme.

### 5.3 Résultats obtenus pour le problème d'optimisation de forme

Dans cette section, nous présentons quelques résultats numériques pour le problème d'optimisation de forme en utilisant l'algorithme (GA). Pour ce faire, nous considérons les mêmes exemples que ceux traités dans la section 5.1. Pour chaque exemple, nous allons reconstruire la frontière libre  $\Gamma$ , en utilisant la formulation d'optimisation de forme du problème (4.19), les courbes de Bézier et l'algorithme 5.2.

**Algorithm 5.2** GA

- 
1. Choisir une précision désirée  $\varepsilon$
  2. Choisir les bornes  $x_i$  et  $x_f$ .
  3.  $t \leftarrow 0$   
Générer aléatoirement une population initiale  $P(t)$  dans l'intervalle  $[x_i, x_f]$
  4.  $t \leftarrow t + 1$   
pour chaque élément  $P_i$  de la population  $P(t)$  avec  $i = 1, \dots, N$  faire :
    - (a) Construire  $\Omega_i(t)$
    - (b) Résoudre le problème d'état dans le domaine  $\Omega_i(t)$
    - (c) Calculer  $J(\Omega_i(t), u_i(t))$
  5. Opération génétique
    - (a) Sélection
    - (b) Croisement
    - (c) Mutation
  6. Si  $|J(\Omega_{best}(t), u_{best}(t))| < \varepsilon$  stop. Sinon aller à l'étape 4
  7. Fin
- 

|                                 |                                       |
|---------------------------------|---------------------------------------|
| taille de population $N$        | 16                                    |
| type de sélection               | à roulette                            |
| type de croisement              | barycentrique à coefficient aléatoire |
| probabilité de croisement $P_c$ | 0.6                                   |
| type de mutation                | non uniforme                          |
| probabilité de mutation $P_m$   | 5%                                    |

TABLE 5.1 : Paramètres des (GA) pour les résultats des figures 5.24 et 5.25

**5.3.1 Validation de l'algorithme (GA) vis-à-vis d'une solution exacte**

Pour le premier exemple nous résoudrons le problème (4.19) avec les paramètres cités dans (5.6) et la frontière exacte définie dans (5.7). Nous lançons l'algorithme 5.2 avec les paramètres

cités dans le tableau 5.1, et nous obtenons ainsi les résultats représentés sur les figures suivantes 5.24, 5.25 et le tableau 5.2. Nous remarquons que sur cet exemple la frontière obtenue représente une très bonne approximation de la frontière exacte. Ceci est dû au fait que la frontière exacte est paramétrisée par une courbe de Bézier.

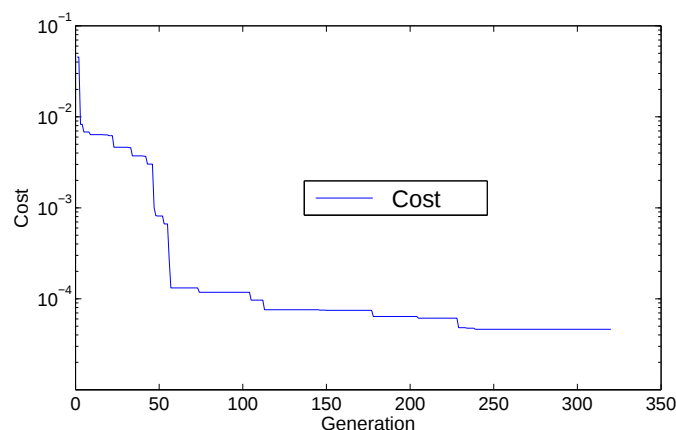


FIGURE 5.24 : Décroissance du coût en fonction des itérations.

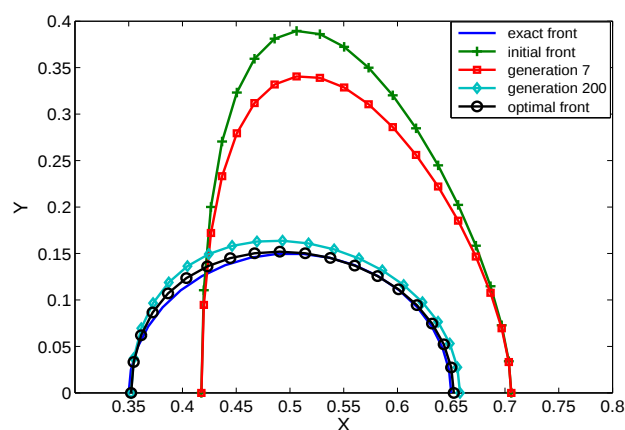


FIGURE 5.25 : Évolution de la frontière libre au cours des itérations.

|                           |                      |
|---------------------------|----------------------|
| nombre de générations     | 227                  |
| coût minimal              | $4.2 \times 10^{-5}$ |
| pas de discrétisation $h$ | 1/36                 |

TABLE 5.2 : Résultats de convergence de l'algorithme 5.2 pour l'exemple 1

Pour le deuxième exemple nous résoudrons le problème (4.19) avec les paramètres cités dans (5.6) et la frontière exacte définie dans (5.8). Nous implémentons l'algorithme 5.2 avec les mêmes paramètres que ceux cités dans le tableau 5.1. Ainsi nous obtenons les résultats présentés dans les figures 5.26 et 5.27.

Nous remarquons que pour atteindre la même valeur minimale que dans l'exemple précédent, l'algorithme a besoin d'aller jusqu'à la génération 335. Ceci est dû au fait que la frontière  $\Gamma$  recherchée est plus compliquée, étant donné qu'elle ne peut être représentée par une courbe de Bézier.

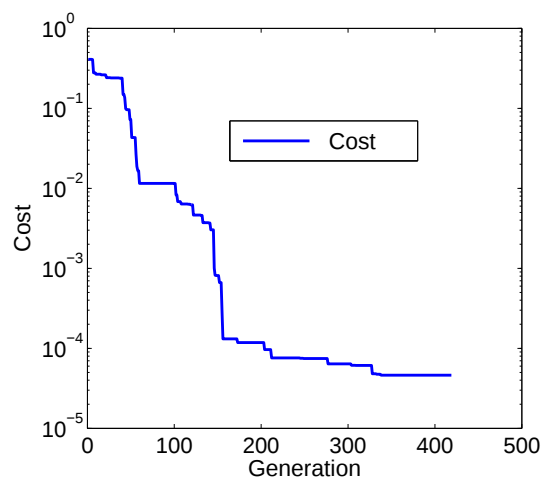


FIGURE 5.26 : Décroissance du coût en fonction des itérations.

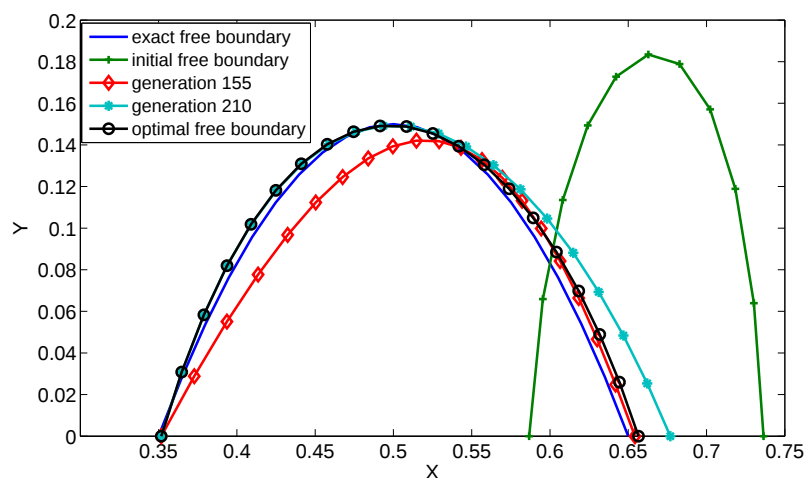


FIGURE 5.27 : Évolution de la frontière libre au cours des itérations.

### 5.3.2 Validation de l'algorithme (GA) vis-à-vis du modèle physique

Pour le troisième exemple, nous résoudrons le problème (4.19) avec les données physiques du paragraphe 5.1.4. Nous implémentons l'algorithme 5.2 avec les mêmes paramètres que ceux mentionnés dans le tableau 5.1. Les résultats obtenus sont présentés dans les figures 5.28 et 5.29.

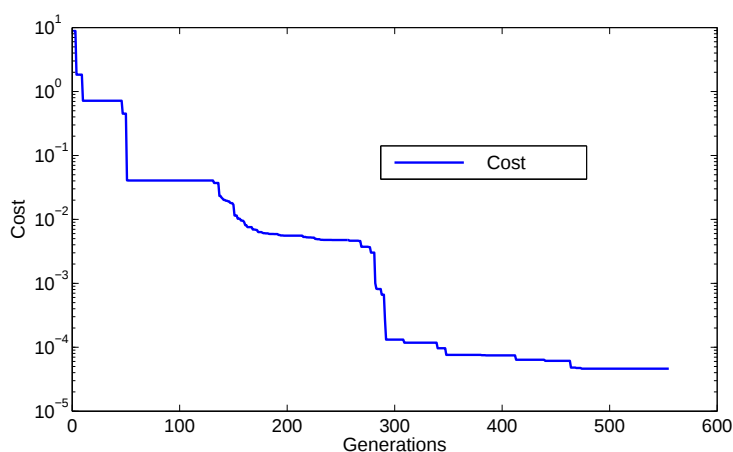


FIGURE 5.28 : Décroissance du coût en fonction des itérations.

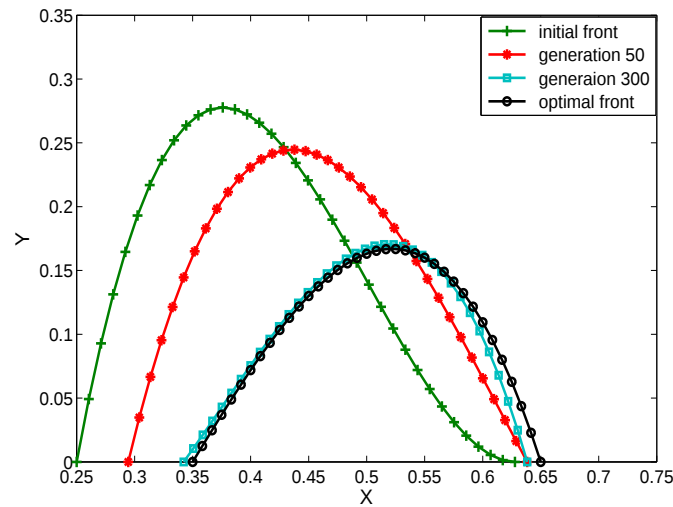


FIGURE 5.29 : Évolution de la frontière libre au cours des itérations.

## 5.4 Développement d'un algorithme évolutionnaire

En ce qui concerne les algorithmes génétiques standards, les bornes supérieures et inférieures des régions de recherche doivent être données par le décideur avant de commencer le processus d'optimisation. Si l'on fait le choix d'un petit intervalle, on risque de perdre une bonne solution qui serait située au-delà de l'intervalle donné. D'autre part, le choix d'un grand intervalle peut conduire au même problème. En effet, un nombre limité d'individus de la population sera dispersé aléatoirement sur un grand intervalle. Par ailleurs, si le choix d'un grand nombre d'individus peut couvrir plus de régions de l'intervalle, cependant cela peut ralentir l'optimisation sans garantir l'optimum global.

Pour trouver le meilleur choix de l'intervalle d'initialisation de la population, nous proposons l'utilisation d'un contrôleur de logique floue. Ce contrôleur prend avantage de la manière dont chaque variable évolue au cours de l'optimisation et ajuste les intervalles d'initialisation de sorte que le prochain round d'exécution de l'algorithme donnera un meilleur résultat.

### 5.4.1 Logique floue

La logique floue a été développée par L. A. Zadeh en 1965 à partir de sa théorie des sous-ensembles flous [84]. Les sous-ensembles flous sont une manière mathématique de représenter

l'imprécision de la langue naturelle ; ils peuvent être considérés comme une généralisation de la théorie des ensembles classiques. La logique floue est aussi appelée "logique linguistique" car ses valeurs de vérité sont des mots du langage courant : "plutôt vrai, presque faux, loin, si loin, près de, grand, petit...". La logique floue a pour objectif d'étudier la représentation des connaissances imprécises des raisonnements approchés et de chercher à modéliser les notions vagues du langage naturel pour pallier l'inadéquation de la théorie des ensembles classiques dans ce domaine.

#### 5.4.1.1 Sous-ensembles flous

En théorie des ensembles classiques, l'appartenance d'un élément à un sous-ensemble est booléenne. Les sous-ensembles flous permettent en revanche de connaître le degré d'appartenance d'un élément au sous-ensemble. Un sous-ensemble flou  $A$  d'un univers du discours  $U$  est caractérisé par une fonction d'appartenance [84] :

$$\mu_A : U \rightarrow [0, 1]$$

où  $\mu_A$  est le niveau ou degré d'appartenance d'un élément de l'univers du discours  $U$  dans le sous-ensemble flou. On peut aussi définir un sous-ensemble flou  $\bar{A}$  dans l'univers du discours  $U$  comme suit :

$$\bar{A} = \{(x, \mu_{\bar{A}}(x)) | x \in U\},$$

avec  $\mu_{\bar{A}}(x)$  comme le degré d'appartenance de  $x$  dans  $\bar{A}$ .

**Exemple 5.1** Soit  $U$  un sous-intervalle de  $\mathbb{R}$  et  $A$  un sous-ensemble classique pour représenter les nombres réels supérieurs ou égaux à 5 ; alors, nous avons :

$$A = \{(x, \mu_A(x)) | x \in U\}$$

où la fonction caractéristique est définie par :

$$\mu_A(x) = \begin{cases} 0, & \text{si } x < 5, \\ 1, & \text{si } x \geq 5. \end{cases}$$

Cette fonction est présentée dans la Figure 5.30(a).

À présent, nous considérons un sous-ensemble flou  $\bar{A}$  qui représente les nombres réels proches de 5 et qui est défini de la façon suivante :

$$\bar{A} = \{(x, \mu_{\bar{A}}(x)) | x \in U\}$$

où la fonction caractéristique est définie par exemple par :

$$\mu_A(x) = \frac{1}{1 + 10(x - 5)^2}.$$

Cette fonction est présentée dans la Figure 5.30(b).

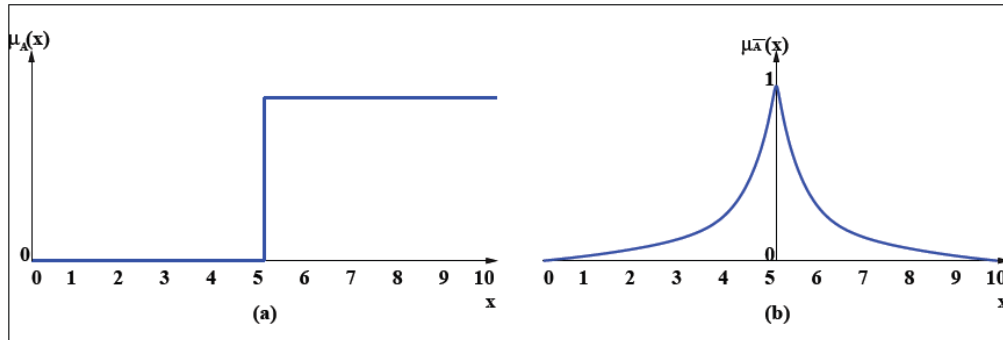


FIGURE 5.30 : Fonctions caractéristiques d'un sous-ensemble classique (a) et d'un sous-ensemble flou (b) pour l'exemple 5.1

#### 5.4.1.2 Variables linguistiques

En logique floue les concepts des systèmes sont normalement représentés par des variables linguistiques. Une variable linguistique est une variable dont les valeurs sont des mots ou des phrases utilisés couramment dans une langue naturelle ou un langage artificiel. Une variable linguistique est définie par :

$$(X, U, T(X), \mu_x)$$

où  $X$  désigne le nom de la variable,  $U$  est l'univers du discours associé à la variable  $X$  (appelé aussi référentiel),  $T(X) = \{T_1, T_2, \dots, T_n\}$  est l'ensemble des valeurs linguistiques de la variable  $X$  (appelées également termes linguistiques ou étiquettes linguistiques), et finalement  $\mu_x$  sont les fonctions d'appartenance associées à l'ensemble de termes linguistiques.

**Exemple 5.2** On considère un contrôleur flou pour la régulation de la vitesse d'un ventilateur en fonction de la température et l'humidité à l'intérieur d'une pièce. Les variables linguistiques

de température et d'humidité (dites variables d'entrée) seront définies comme suit :

$$(Température, U^t = \{0, 50\}, T(X) = \{T_1, T_2, T_3\}, \mu = \{\mu_1^t, \mu_2^t, \mu_3^t\})$$

$$(Humidité, U^h = \{0, 100\}, H(X) = \{H_1, H_2, H_3\}, \mu = \{\mu_1^h, \mu_2^h, \mu_3^h\})$$

où  $U^t$  et  $U^h$  sont respectivement l'univers des températures et l'univers des humidités. Les ensembles  $T$  et  $H$  sont constitués respectivement de trois étiquettes linguistiques :  $T_1 = \text{froid}$ ,  $T_2 = \text{moyenne}$  et  $T_3 = \text{chaud}$ , (resp  $H_1 = \text{faible}$ ,  $H_2 = \text{moyenne}$  et  $H_3 = \text{grande}$ ) et les fonctions d'appartenances définies par chaque terme linguistique sont :

$$\mu_{\text{froid}}^t = \text{trapézoïde}(x, 0, 0, 15, 25)$$

$$\mu_{\text{moyenne}}^t = \text{triangulaire}(x, 15, 25, 35)$$

$$\mu_{\text{chaud}}^t = \text{trapézoïde}(x, 25, 35, 50, 50)$$

$$\mu_{\text{faible}}^h = \text{trapézoïde}(x, a', 0, 20, 50)$$

$$\mu_{\text{moyenne}}^h = \text{triangulaire}(x, 20, 50, 80)$$

$$\mu_{\text{grande}}^h = \text{trapézoïde}(x, 50, 80, 100, 100)$$

Les variables linguistiques de la vitesse du ventilateur (dites variables de sortie) seront définies comme suit :

$$(Vitesse, U^p = \{0, 100\}, P(X) = \{P_1, P_2, P_3\}, \mu = \{\mu_1^p, \mu_2^p, \mu_3^p\})$$

où  $U^p$  est l'univers des vitesses. L'ensemble  $P$  est constitué des trois étiquettes linguistiques :  $P_1 = \text{faible}$ ,  $P_2 = \text{moyenne}$  et  $P_3 = \text{maximale}$ , et les fonctions d'appartenances définies par chaque terme linguistique sont :

$$\mu_{\text{faible}}^t = \text{trapézoïde}(x, 0, 0, 10, 40)$$

$$\mu_{\text{moyenne}}^t = \text{triangulaire}(x, 10, 40, 70)$$

$$\mu_{\text{maximale}}^t = \text{trapézoïde}(x, 40, 70, 100, 100)$$

La fonction d'appartenance triangulaire est définie comme suit :

$$\text{triangulaire}(x; a, b, c) = \max \left( \min \left( \frac{x-a}{b-a}, \frac{c-x}{c-b} \right), 0 \right)$$

avec ( $a < b < c$ ) où  $b$  est le sommet du triangle tandis que  $a$  et  $c$  imposent la largeur du domaine de la valeur à fuzzifier. La fonction d'appartenance trapézoïde est définie comme suit :

$$\text{trapézoïde}(x; a, b, c, d) = \max \left( \min \left( \frac{x-a}{b-a}, 1, \frac{d-x}{d-b} \right), 0 \right)$$

avec ( $a < b \leq c < d$ ) où  $b$  et  $c$  sont le sommet du trapézoïde tandis que  $a$  et  $d$  fournissent la largeur du domaine de la valeur à fuzzifier.

Nous illustrons l'exemple 5.2 dans les Figures 5.31, 5.32 et 5.33 où l'on représente respectivement les variables linguistiques de la température, l'humidité et la vitesse du ventilateur. La définition de chaque sous-ensemble flou repose sur l'intuition. Si la température dans la pièce est de 20 °C et l'humidité est de 62%, cela se traduira par différents degrés d'appartenance à chacun des sous-ensembles flous :

$$\mu_{froid}^t(20) = 0.5, \quad \mu_{moyenne}^t(20) = 0.5, \quad \mu_{chaud}^t(20) = 0.$$

$$\mu_{faible}^h(62) = 0, \quad \mu_{moyenne}^h(62) = 0.6, \quad \mu_{grande}^h(62) = 0.4$$

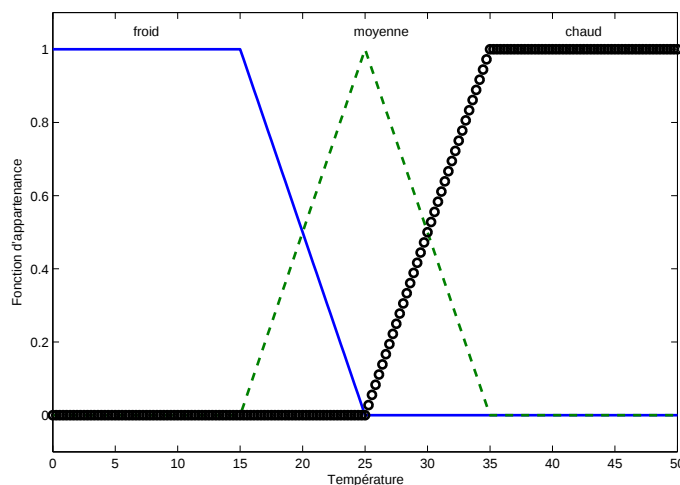


FIGURE 5.31 : Fonctions d'appartenance de la variable température.

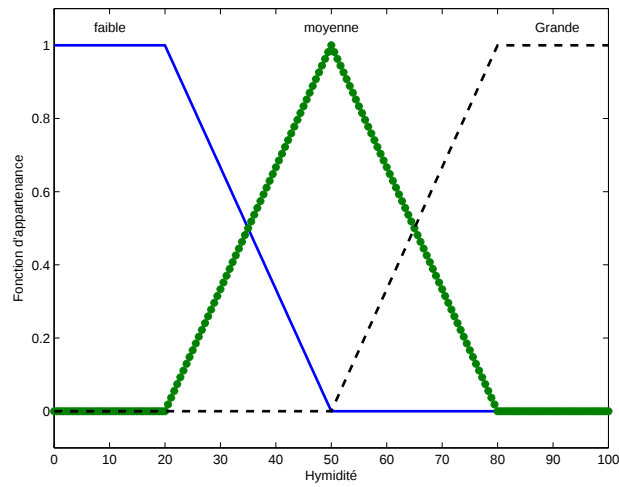


FIGURE 5.32 : Fonctions d'appartenance de la variable humidité.

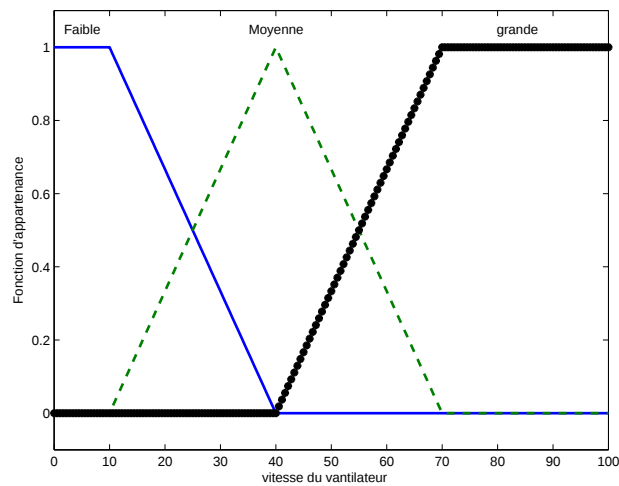


FIGURE 5.33 : Fonctions d'appartenance de la variable vitesse du ventilateur.

## 5.4.1.3 Système d'Inférence Floue

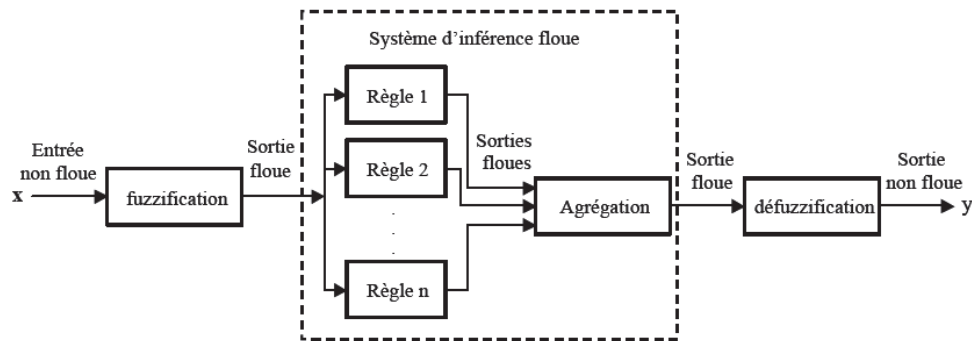


FIGURE 5.34 : Structure d'un SIF

Un Système d'Inférence Floue (SIF) a pour but de transformer les données d'entrée en données de sortie à partir de l'évaluation d'un ensemble de règles. Les entrées sont issues du processus de fuzzification et l'ensemble de règles est normalement défini par le savoir-faire de l'expert. Un SIF (voir Figure 5.34) est constitué de trois étapes :

- a) Fuzzification
- b) Inférence
- c) Défuzzification

La première étape est la fuzzification, qui consiste à caractériser les variables linguistiques utilisées dans le système. Il s'agit donc d'une transformation des entrées réelles en une partie floue définie sur un espace de représentation lié à l'entrée. Cet espace de représentation est normalement un sous-ensemble flou. Durant l'étape de la fuzzification, chaque variable d'entrée et de sortie est associée à des sous-ensembles flous.

La deuxième étape est le moteur d'inférence, qui est un mécanisme permettant de condenser l'information d'un système à travers un ensemble de règles définies pour la représentation d'un problème quelconque. Chaque règle délivre une conclusion partielle qui est ensuite agrégée aux autres règles pour fournir une conclusion (agrégation). Les règles constituent le système d'inférence floue.

Dans la suite de cette section nous donnons une description de règles floues dans un cadre plus formel.

La troisième étape est la défuzzification. Cette opération est l'inverse de la fuzzification et permet de transformer les sorties floues de l'inférence en une valeur non floue comme réponse finale du SIF.

#### 5.4.1.4 Définition des règles floues

Le nombre de règles dans un SIF dépend du nombre de variables (d'entrée et de sortie). Les règles floues sont généralement du type "SI ... ALORS" et permettent de représenter les relations entre les variables d'entrée et de sortie. Plus précisément, une règle floue  $R_j$  ( $0 \leq j \leq N$ ) est définie de la manière suivante [61] :

$$R_j : \text{ Si } x \text{ est } A_j \text{ et } y \text{ est } B_j \text{ Alors } z \text{ est } C_j \quad (5.10)$$

où  $(A_j)_{0 \leq j \leq N}$ ,  $(B_j)_{0 \leq j \leq N}$  et  $(C_j)_{0 \leq j \leq N}$  sont des variables linguistiques définies respectivement dans des univers du discours  $X$ ,  $Y$  et  $S$ . La première partie de la règle " $x$  est  $A_j$  et  $y$  est  $B_j$ " est l'antécédent, et la deuxième partie " $z$  est  $C_j$ " est le conséquent.

#### 5.4.1.5 Inférence à partir de règles floues

Le but de l'inférence floue est de déterminer les sorties du système à partir des entrées floues issues de la fuzzification des entrées réelles. Considérons une collection de règles ; le mécanisme d'inférence consiste alors à dériver un ensemble flou de sorties à partir de l'agrégation des conclusions de l'ensemble des règles floues.

#### Inférence avec une seule règle

Dans le cas où une seule règle floue serait activée, l'inférence repose sur la valeur d'appartenance ( $\mu$ ) (appelée poids) associée à la variable linguistique d'entrée. La définition pour ce cas est comme suit :

$$\text{Règle 1 : Si } x \text{ est } A \text{ et Si } y \text{ est } B \text{ Alors } z \text{ est } C$$

Par ailleurs, le degré d'appartenance de la variable linguistique de sortie ( $C$ ) est défini comme

suit :

$$\mu_C(z) = \text{poids de la règle } 1 = \min(\mu_A(x), \mu_B(y))$$

### Inférence avec plusieurs règles

Dans le cas où plusieurs règles floues seraient activées, l'inférence repose sur les différentes valeurs d'appartenance ( $\mu$ ) associées aux variables linguistiques d'entrée. D'une manière générale, l'inférence des règles (5.10) peut être représentée par la relation suivante :

$$R = \bigcup_{i=1}^N [(A_i \cap B_i) \times C_i]. \quad (5.11)$$

Le symbole  $\bigcup$  représente l'opérateur "OU" introduit entre les règles. Le symbole  $\cap$  représente l'opérateur "ET" utilisé dans la partie antécédente des règles et  $\times$  représente l'opérateur d'implication floue ("ALORS"). Le sous-ensemble obtenu correspond aux variables d'entrée  $(x_0, y_0)$ , et est donné par :

$$C(z) = R(x_0, y_0, z).$$

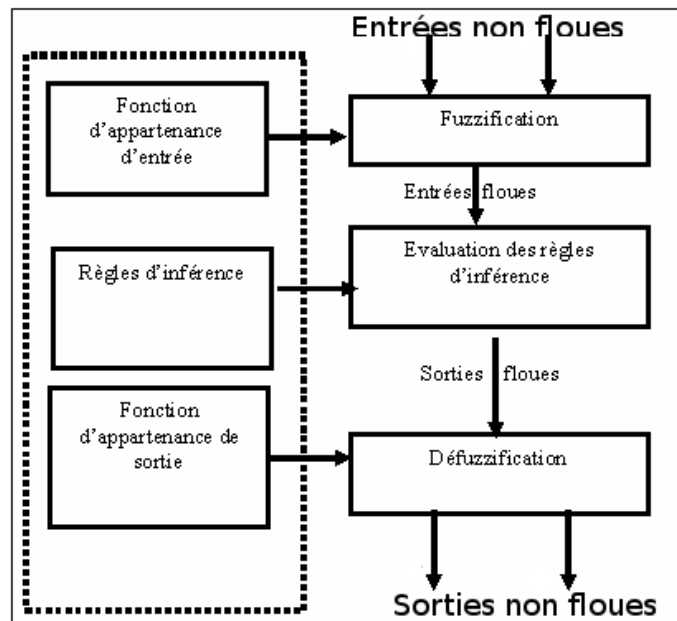
Si nous adoptons le minimum ("min" ou  $\wedge$ ) pour les opérateurs "et" et "Alors", et le maximum (max) pour l'opérateur "OU", nous pouvons représenter  $C(z)$  comme suit :

$$C(z) = \max_{1 \leq i \leq N} [A_i(x_0) \wedge B_i(y_0) \wedge C_i(z)]. \quad (5.12)$$

#### 5.4.1.6 Défuzzification

La défuzzification permet d'avoir un résultat numérique non flou à partir de la sortie de l'inférence. Il y a différentes techniques de défuzzification telles que (COG) centre de gravité, (MOM) moyenne de maximas et d'autres.

La configuration des contrôleurs de logique floue est capsulée dans l'algorithme suivant.

**Algorithm 5.3** Contrôleur de logique floue

Dans ce travail nous utilisons la méthode WABL proposée dans [76] et généralisée dans le cadre de cette thèse [61]. Dans le but d'illustrer l'importance et la nécessité d'avoir une méthode de défuzzification fiable, nous considérons alors l'exemple 5.2 et nous supposons que la vitesse est proportionnelle aux valeurs linguistiques de température  $T$  et d'humidité  $H$ .

|    |                        |
|----|------------------------|
| F  | faible                 |
| B  | moyenne                |
| M  | Moyenne                |
| G  | Grande                 |
| C  | Chaud                  |
| Ht | Haute                  |
| V  | Vitesse du ventilateur |

TABLE 5.3 : Variables linguistiques

Nous supposons aussi que le but des contrôleurs flous est d'avoir une décroissance de la vitesse du ventilateur quand la température décroît et l'humidité reste fixe et vice versa. En

gardant cette objectif à l'esprit, nous pouvons adopter les règles d'inférence suivantes :

| V               | H (Humidité) |   |    |
|-----------------|--------------|---|----|
|                 | F            | M | Ht |
| T (Température) | F            | F | M  |
|                 | B            | M | G  |
|                 | C            | G | G  |

TABLE 5.4 : Matrice des règles d'inférence

|         | température | humidité |
|---------|-------------|----------|
|         | 17 °C       | 32 %     |
| froid   | 0.8         |          |
| moyenne | 0.2         |          |
| chaud   | 0           |          |
| faible  |             | 0.6      |
| moyenne |             | 0.4      |
| grande  |             | 0        |

TABLE 5.5 : Valeurs de vérité pour  $T = 17\text{ °C}$  et  $H = 32\%$

Pour mieux comprendre, nous supposons par exemple que la température ambiante est  $T = 17\text{ °C}$  et que le pourcentage de l'humidité est  $H = 32\%$ . Alors les valeurs de vérité sont citées dans le tableau 5.5. Si nous utilisons la convention "**min**" pour "**et**" alors la sortie floue du ventilateur est présentée dans le tableau 5.6

|                            | degré d'appartenance |         |        |
|----------------------------|----------------------|---------|--------|
|                            | Faible               | Moyenne | Grande |
| Règle 1 : $\min(0.8, 0.6)$ | 0.6                  |         |        |
| Règle 2 : $\min(0.8, 0.4)$ | 0.4                  |         |        |
| Règle 3 : $\min(0.8, 0.0)$ |                      | 0.0     |        |
| Règle 4 : $\min(0.2, 0.6)$ |                      | 0.2     |        |
| Règle 5 : $\min(0.2, 0.4)$ |                      | 0.2     |        |
| Règle 6 : $\min(0.2, 0.0)$ |                      |         | 0.0    |
| Règle 7 : $\min(0.0, 0.6)$ |                      |         | 0.0    |
| Règle 8 : $\min(0.0, 0.4)$ |                      |         | 0.0    |
| Règle 9 : $\min(0.0, 0.0)$ |                      |         | 0.0    |

TABLE 5.6 : Degré d'appartenance de la sortie floue de la vitesse du ventilateur qui résulte des règles 5.4

En appliquant la relation (5.12) nous obtenons ainsi la sortie floue présentée dans la figure 5.35

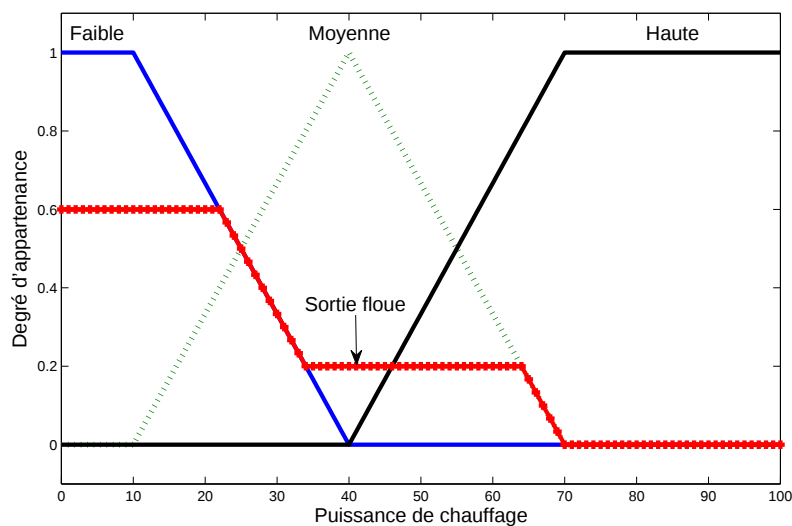


FIGURE 5.35 : Sortie floue obtenue

### 1. Moyenne de Maxima

Nous appliquons la méthode MOM à la sortie floue  $C(z)$  en prenant la moyenne des valeurs  $z$  qui maximisent  $C(z)$ , et nous écrivons :

$$MOM[(C(z))] = \frac{\int_a^b z dz}{\int_a^b dz} = \frac{1}{2}(a + b) \quad (5.13)$$

Pour notre exemple, nous avons  $MOM[C(17, 32)] = 11\%$ .

La figure 5.36 présente la variation de la vitesse du ventilateur en fonction de la température et de l'humidité en utilisant 5.13. Nous remarquons que la technique MOM ne prend en compte que les règles qui sont agrégées dans le maximum de la fonction d'appartenance.

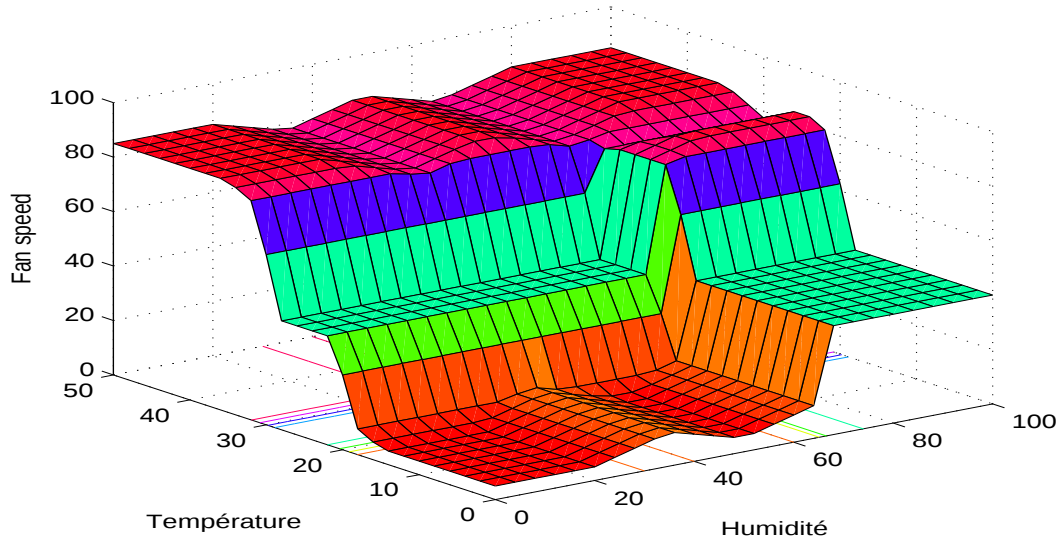


FIGURE 5.36 : Surface de contrôle du contrôleur flou avec des règles d'inférence comme dans (5.4). La défuzzification MOM est appliquée.

## 2. Centre de gravité

Il s'agit de trouver le centre de gravité de la fonction d'appartenance de la variable de sortie  $C(z)$ . Nous écrivons alors :

$$COG[C(z)] = \frac{\int_{-\infty}^{\infty} z C(z) dz}{\int_{-\infty}^{\infty} C(z) dz} \quad (5.14)$$

En appliquant la méthode (5.14) à l'exemple 5.2, nous obtenons

$COG[(17, 32)] = 24.73\%$ . Ainsi, dans la figure 5.37 nous représentons les variations de la vitesse du ventilateur en fonction des variations de la température et de l'humidité.

Alors que la méthode MOM ne prend en compte que les règles qui agrègent les valeurs maximales, la méthode COG prend aussi en compte celles qui agrègent les valeurs qui sont en dessous de la fonction d'appartenance. En revanche, elle a comme inconvénient le fait de restreindre l'action de contrôle dans la région désirée. De plus, la figure contient des parties indésirables où la vitesse du ventilateur décroît quand la température croît et l'humidité reste fixe, et vice versa.

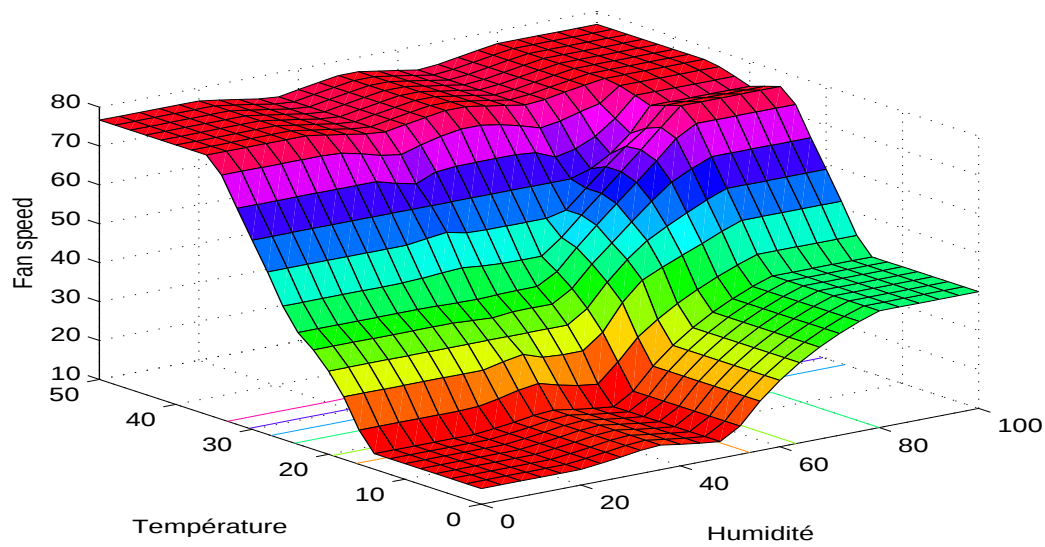


FIGURE 5.37 : Surface de contrôle du contrôleur flou avec des règles d'inférence comme dans (5.4). La défuzzification COG est appliquée

### 3. La méthode de WABL (weighted averaging based on the levels)

La forme mathématique de la méthode WABL est donnée par :

$$WABL[c(t_0, h_0)] = \sum_{i=1}^3 \sum_{j=1}^3 \mu_{ij} I_w(c_{ij}) \quad (5.15)$$

où :  $\mu_{ij} = \mu_{lt_i}(t_0) \times \mu_{lh_j}(h_0)$ ,  $lt_1 = \text{froid}$ ,  $lt_2 = \text{moyenne}$ ,  $lt_3 = \text{chaud}$ ,

$lh_1 = \text{faible}$ ,  $lh_2 = \text{moyenne}$ ,  $lh_3 = \text{grande}$  et  $c_{ij}$  est le résultat correspondant à la règle

"SI Température est  $lt_i$  et Humidité est  $lh_j$  ALORS  $c_{ij}$ ".

$$I_w(c_{ij}) = c_l \int_0^1 L_{c_{ij}}(\xi) P(\xi) d\xi + (1 - c_l) \int_0^1 R_{c_{ij}}(\xi) P(\xi) d\xi \quad (5.16)$$

avec  $c_l \in [0, 1]$  indique le degré d'importance du côté gauche du nombre flou,  $L_{c_{ij}}(\xi) = \mu_{\uparrow}^{-1}, R_{c_{ij}}(\xi) = \mu_{\downarrow}^{-1}$ .  $\mu_{\uparrow}^{-1}$  et  $\mu_{\downarrow}^{-1}$  sont respectivement l'inverse de la fonction du côté gauche et du côté droit de la fonction d'appartenance sortie. Quant à  $P(\xi)$  c'est une fonction distribution de degré d'importance définie par :

$$P : [0, 1] \mapsto [0, +\infty] \text{ vérifie} \quad \int_0^1 P(\xi) d\xi = 1 \quad (5.17)$$

Nous appliquons la méthode "WABL" à l'exemple 5.2. Nous obtenons alors

$WABL[(17, 32)] = 19\%$ , avec  $cl = 0.35$  et  $P(\xi) = (k + 1)\xi^k$ ,  $k = 1$ .

L'avantage principal de la méthode "WABL" est l'utilisation du paramètre libre  $c_l$  et de la fonction  $P$  qui permettent l'adaptation de la défuzzification pour obtenir un résultat fiable. La figure 5.38 présente les variations de la vitesse du ventilateur en fonction des

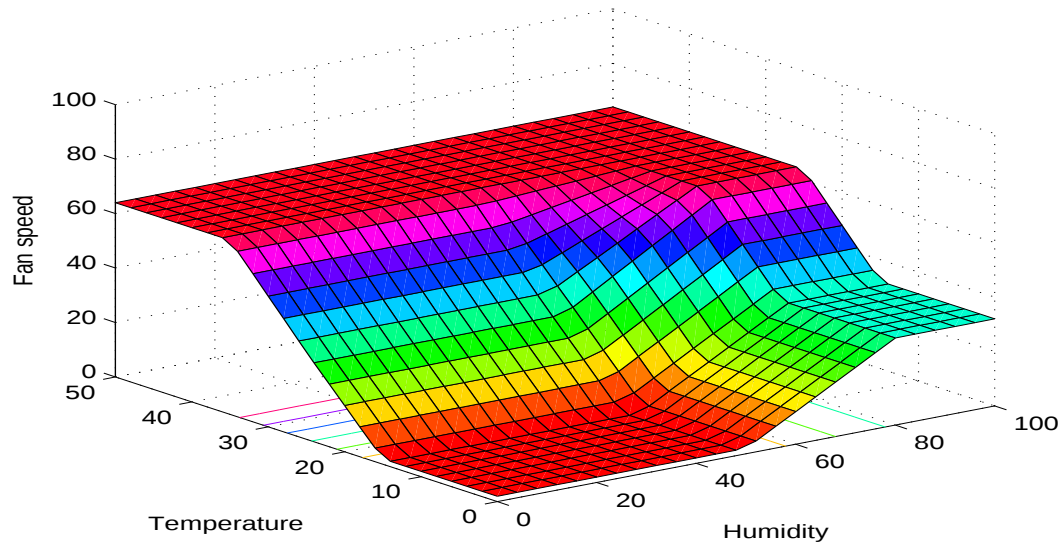


FIGURE 5.38 : Surface de contrôle du contrôleur flou avec des règles d'inférence comme dans (5.4). La défuzzification WABL est utilisée.

variations de la température et de l'humidité en appliquant la méthode WABL. Cette méthode parvient à surmonter les inconvénients des méthodes "MOM" et "COG". Elle prend en compte les règles agrégées à tous les niveaux de la fonction d'appartenance. De plus, elle permet l'action de contrôle au-delà de l'intervalle d'action.

Il existe d'autres méthodes pour faire la défuzzification : MOM (moyenne de Maxima), WAF

(formule moyenne pondérée), QM (Méthode Qualité), Wabl (moyenne pondérée en fonction des niveaux) et l'ÉTS. Le problème est que chacune donne une valeur différente pour le même problème. La question la plus importante est de savoir laquelle est la bonne. Dans le but de répondre à cette question, un travail commun a été réalisé avec R. Y. Shikhinskaya [61].

Dans le paragraphe suivant, nous allons présenter un algorithme qui combine les (GA) avec un système d'inférence.

#### 5.4.2 Combinaison des (GA) avec les contrôleurs de logique floue (CLF)

Pour trouver le meilleur choix de l'intervalle de sélection, nous proposons l'utilisation d'un contrôleur de logique floue. Ce contrôleur prend avantage de la manière dont chaque variable évolue au cours de l'optimisation et ajuste les intervalles de sélection de sorte que la prochaine exécution de l'optimisation donnera un meilleur résultat. Le processus, consiste à lancer l'algorithme génétique décrit avant, avec une population initiale choisie au hasard dans un intervalle d'admission initiale. La (CLF) surveille l'évolution des différentes variables au cours du processus d'optimisation et ajuste les intervalles de sélection de chaque variable pour le prochain cycle du processus d'optimisation.

Ces nouveaux intervalles sont ensuite utilisés pour lancer un second cycle de l'algorithme. En analysant le résultat final après chaque cycle. Un paramètre  $Jv$  est déterminé et l'intervalle de la sélection pour chaque variable sera corrigé selon le critère suivant

$$\begin{aligned} x_i^* &= x_m - Jv/2(x_f - x_i) \\ x_f^* &= x_m + Jv/2(x_f - x_i) \end{aligned} \quad (5.18)$$

où  $x_m$  est la valeur moyenne de tous les individus de la dernière génération et  $[x_i, x_f]$  est l'intervalle de sélection de la variable  $x$ . Le coefficient  $Jv$  (la variable de sortie) est obtenu à partir de la connaissance de l'entrée à deux variables  $ED$  et  $CV$ , où  $ED$  est un critère pour mesurer la diversité de l'algorithme génétique. Nous proposons l'une de ces mesures, qui est basée sur les distances euclidiennes entre les chromosomes de la population et le meilleur individu.

$$ED = \frac{\bar{d} - d_{min}}{d_{max} - d_{min}}$$

où

$$\begin{aligned} \bar{d} &= \frac{1}{N} \sum_{i=1}^{i=N} d(C_{best}, C_i), \\ d_{min} &= \min\{d(C_{best}, C_i) / C_i \in \mathcal{P}\} \end{aligned}$$

et

$$d_{max} = \max\{d(C_{best}, C_i) / C_i \in \mathcal{P}\}$$

où  $C_{best}$  est le meilleur individu de la population  $\mathcal{P}$ . La variable  $CV$  est un compteur de variations de chaque paramètre durant chaque cycle de l'algorithme.

$ED$  évolue dans l'intervalle  $[0, 1]$ . Une petite valeur de  $ED$  signifie que la plupart des individus de la population  $\mathcal{P}$  sont concentrés autour du meilleur individu, et dans ce cas la convergence est achevée. Une valeur proche de 1 signifie que les individus de la population sont dispersés, cela signifie qu'un autre cycle est nécessaire.

Le compteur  $CV$  est un indicateur du nombre de fois où les individus ont changé. Un indicateur faible signifie que les individus de l'actuelle population sont proches de leur optimum. Par conséquent nous n'avons pas besoin de changer les bornes de l'intervalle initial. Un indicateur élevé nécessite l'ajustement des bornes de l'intervalle initial afin d'améliorer la solution finale. Cet ajustement est conditionné aussi par l'amélioration de la solution obtenue. En effet, si nous arrivons à une solution précise au premier cycle de l'algorithme (GA) il n'est pas nécessaire d'effectuer de changement, même si nous obtenons des indicateurs élevés pour certaines variables. Ce raisonnement est mieux mis en œuvre en utilisant un contrôleur de logique floue.

La première étape dans la conception d'un contrôleur de logique floue est la fuzzification des variables d'entrée. Cette étape implique une quantification des univers en un certain nombre d'ensembles flous. Les variables de sortie doivent également être quantifiées de façon similaire. La quantification consiste à briser une variable d'entrée floue (et aussi de sortie) en plusieurs sous-ensembles flous. Les variables linguistiques sont utilisées pour étiqueter les sous-ensembles flous (tableau 5.7) afin que les règles puissent être écrites dans un langage naturel (tableau 5.8).

Afin de trouver la bonne taille ( $N_r$ ) d'un cycle d'algorithme 5.2, nous avons fait plusieurs essais. La figure 5.39 représente les résultats obtenus avec différentes valeurs de  $N_r$ . Nous remarquons que la meilleure accélération est obtenue avec  $N_r = 50$ .

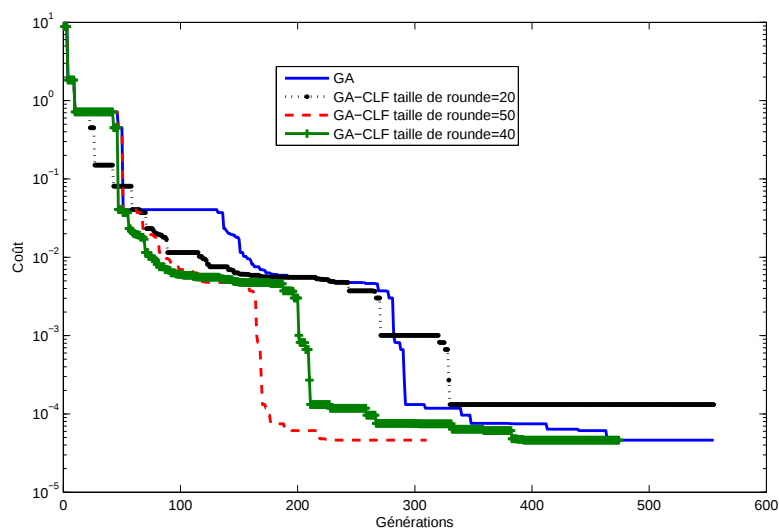


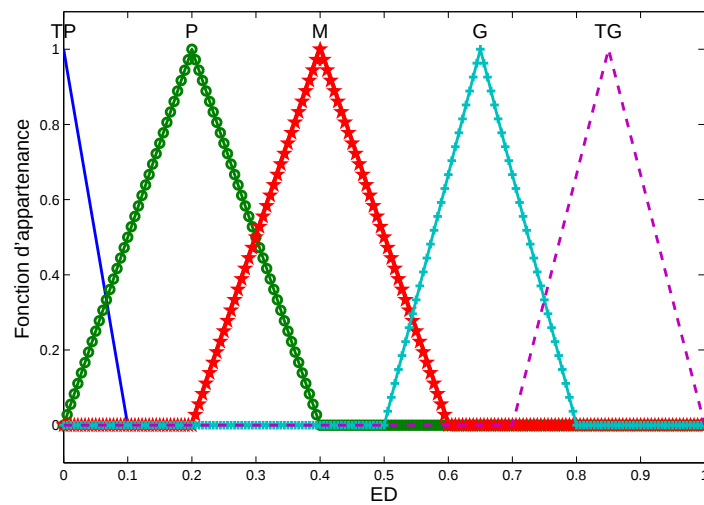
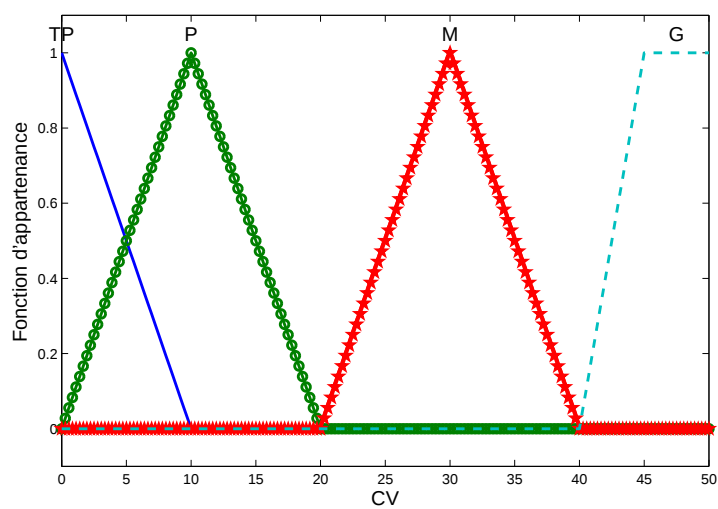
FIGURE 5.39 : Comparaison de différentes tailles du cycle (GA)

Pour plus de simplicité, nous pouvons choisir les mêmes types de fonction d'appartenance (triangulaire) pour chacun des sous-ensembles flous de l'ensemble des variables floues. Les figures 5.40, 5.41 et 5.42 montrent les fonctions d'appartenance choisies pour les deux variables d'entrée floue et la variable de sortie. Le tableau 5.7 contient la définition des paramètres linguistiques.

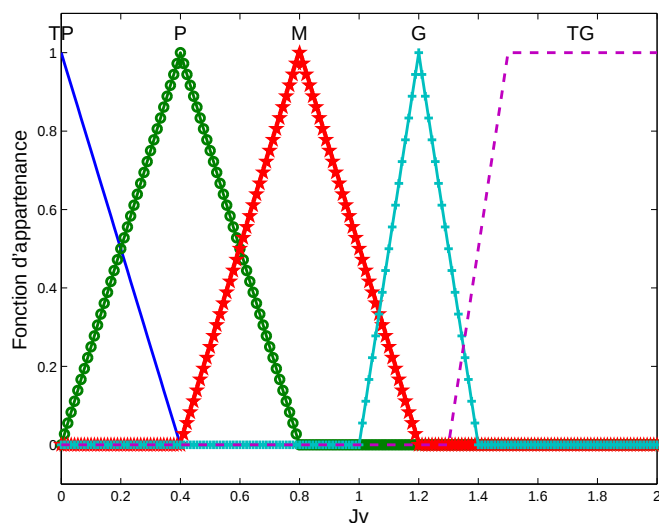
|    |            |
|----|------------|
| TP | Très petit |
| P  | Petit      |
| M  | Moyen      |
| G  | Grand      |
| TG | Très grand |

TABLE 5.7 : Variables linguistiques

Les fonctions d'appartenances qui définissent les variables d'entrée  $ED$  et  $CV$  sont représentées dans les figures 5.40 et 5.41

FIGURE 5.40 : Variable d'entrée  $ED$ FIGURE 5.41 : Variable d'entrée  $CV$ 

Les fonctions d'appartenances qui définissent la variable de sortie  $Jv$  sont représentées dans la figure 5.42

FIGURE 5.42 : Variable de sortie  $Jv$ 

Pour simplifier nous représentons les règles d'inférence sous la forme d'une matrice appelée matrice des règles floues. Dans le tableau 5.8, nous exposons la matrice des règles floues utilisée par le mécanisme d'inférence.

| JV                       | CV (Compteur) |    |   |   |    |   |
|--------------------------|---------------|----|---|---|----|---|
|                          | TP            | P  | M | G | TG |   |
| ED (mesure de diversité) | TP            | TP | P | P | M  | / |
|                          | P             | TP | P | M | G  | / |
|                          | M             | P  | P | G | M  | G |
|                          | G             | M  | M | M | G  | G |
|                          | TG            | G  | G | M | M  | M |

TABLE 5.8 : Tableau de règles d'inférence

L'algorithme 5.4 résume le processus (GA-CFL) qui combine les algorithmes génétiques et les contrôleurs de la logique floue.

**Algorithm 5.4** GA-CLF

- 
1. Donner une précision  $\varepsilon$  et la taille d'un cycle  $N_r$
  2. Choisir les bornes  $x_i$  et  $x_f$ .
  3.  $t \leftarrow 0$   
Générer aléatoirement une population initiale  $P(t)$  dans l'intervalle  $[x_i, x_f]$
  4.  $t \leftarrow t + 1$   
pour chaque élément  $P_i$  de la population  $P(t)$  avec  $i = 1, \dots, N$  faire :
    - (a) Construire  $\Omega_i(t)$
    - (b) Résoudre le problème d'état dans le domaine  $\Omega_i(t)$
    - (c) Calculer  $J(\Omega_i(t), u_i(t))$
  5. Si  $|J(\Omega_{best}(t), u_{best}(t))| < \varepsilon$  aller à 8.
  6. Faire les opérations génétiques
    - (a) Sélection
    - (b) Croisement
    - (c) Mutation
  7. Si  $\text{mod}(t, N_r) \neq 0$  aller à l'étape 4. Sinon appliquer l'algorithme 5.3 pour recalculer  $x_f$  et  $x_i$  en utilisant l'équation 5.18, puis aller à l'étape 3.
  8. Fin
- 

### 5.4.3 Comparaison de l'algorithme (GA) et l'algorithme (GA-CLF) vis-à-vis du modèle physique

Nous considérons le modèle physique présenté dans la section 5.1.4 et défini avec les paramètres (5.9). Nous appliquons l'algorithme 5.4, avec les paramètres cités dans le tableau 5.1, les fonctions d'appartenances présentées dans les figures 5.40, 5.41 et 5.42, et la matrice de règle d'inférence 5.8.

Dans la figure 5.43, nous faisons une comparaison entre les résultats obtenus par l'algorithme 5.2 et ceux obtenus par l'algorithme 5.4.

Nous remarquons que, jusqu'à l'itération 50, les deux courbes sont identiques, ce qui est normal étant donné que les contrôleurs flous n'interviennent qu'au second cycle, qui commence à

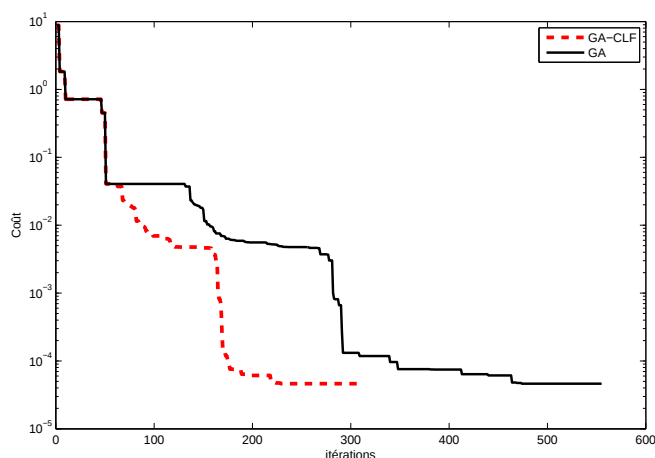


FIGURE 5.43 : Comparaison entre les résultats obtenus par l'algorithme 5.2 et ceux obtenus par l'algorithme 5.4

l'itération 51. Normalement, nous n'avons aucune garantie que le résultat obtenu à l'itération 51 sera meilleur que celui obtenu à l'itération 50. Cependant, grâce à la notion d'élitisme utilisée par l'algorithme 5.2, qui consiste à garder le meilleur individu, nous garantissons une décroissance permanente de la fonctionnelle coût. Nous constatons alors qu'avec les cycles successifs, la courbe obtenue par l'algorithme 5.4 converge plus vite que celle obtenue par l'algorithme 5.2. Ceci nous permet de gagner un nombre important d'itérations et de réduire considérablement le temps de calcul. En ce qui concerne la frontière optimale, les deux algorithmes convergent vers la même solution. En effet, ils sont quasi identiques, puisqu'ils tournent notamment avec les mêmes paramètres. La seule différence, c'est que l'algorithme 5.4 change l'intervalle de sélection à chaque cycle.

Ainsi, nous avons développé l'algorithme 5.4 qui nous a permis d'accélérer considérablement la vitesse de convergence et de réduire le temps de calcul. Cependant, celui-ci reste encore relativement grand en comparaison avec le temps d'exécution de l'algorithme 5.1 utilisant le gradient.

C'est la raison pour laquelle nous avons mis en œuvre une parallélisation de l'algorithme 5.4, que nous présentons dans la section suivante.

## 5.5 Parallélisation de l'algorithme (GA)

La façon la plus simple pour paralléliser les algorithmes (GA) est de distribuer l'évaluation de la fitness sur plusieurs processeurs esclaves, tandis que le processeur maître exécute les opérations d'algorithmes (GA) (sélection, croisement et mutation). Les algorithmes (GA) maître-esclave ont l'avantage d'explorer l'espace de recherche exactement de la même manière qu'une série d'algorithmes (GA), de plus ils apportent des améliorations significatives à la performance.

À présent, nous présentons une analyse qui examine une génération d'algorithmes (GA) maître-esclave, et montre comment on peut réduire au minimum le temps d'exécution. On prend en compte le coût des communications, mais on ignore le temps utilisé par la sélection, le croisement et la mutation, car nous supposons que ce temps est beaucoup plus petit que le temps utilisé pour l'évaluation de la fonction objective et le temps de communication entre les processeurs de calcul. L'analyse suppose également que le nombre d'individus affectés à chaque processeur est constant, et que le temps d'évaluation est le même pour chaque individu.

### 5.5.1 Analyse du temps d'exécution des (GA) Maître-esclaves

Le processeur maître envoie les paramètres à chaque processeur esclave en utilisant un temps de communication  $T_c$ . Chaque processeur évalue sa portion de population. Le temps d'évaluation d'une population est  $\frac{nT_f}{P}$ , où  $T_f$  désigne le temps d'évaluation d'un individu,  $n$  est la taille de la population et  $P$  est le nombre de processeurs utilisés. Le temps total d'exécution d'une seule génération est

$$T_p = \frac{nT_f}{P} + PT_c$$

On sait que lorsque le nombre de processeurs augmente, le temps de calcul décroît, tandis que le temps de communication croît. Pour trouver l'optimum, nous résoudrons l'équation suivante :

$$\frac{\partial T_p}{\partial P} = 0,$$

donc l'optimum  $P^*$  est  $P^* = \sqrt{\frac{nT_f}{T_c}}$ .

Une remarque importante concernant l'implémentation des algorithmes maître-esclave est que le temps de communication peut absorber tout gain possible. Notons par  $T_s = nT_f$  le temps de calcul d'une génération avec une implémentation séquentielle. L'algorithme parallèle (GA)

maître-esclave est plus performant qu'une implémentation séquentielle d'un algorithme (GA) si la condition suivante est vérifiée

$$S_p = \frac{T_s}{T_p} > 1.$$

Nous avons remarqué que l'implémentation parallèle du (GA) maître-esclave n'apporte pas de gain pour des problèmes avec un très petit temps d'évaluation. En outre, il faut que l'implémentation garantisse une forte utilisation des processeurs. Formellement, l'efficacité du programme parallèle est définie par le rapport de  $S_p$  divisé par le nombre de processeurs utilisés ( $P$ ) :

$$E_f = \frac{S_p}{P}$$

L'idéal est que  $S_p$  soit égale au nombre de processeurs utilisés et que l'efficacité  $E_f$  soit égale à 1. En réalité, le coût de communication fait diminuer l'efficacité lorsqu'on augmente les processeurs utilisés.

Dans notre cas, l'évaluation de la fonctionnelle coût correspond à la résolution d'un problème aux limites. Par conséquent, le temps  $T_f$  est a priori suffisamment grand devant le temps de communication  $T_c$  tant que la taille de la population reste raisonnable (dans notre cas on a pris 16). Ceci nous garantit une très grande performance au niveau du temps de calcul.

Dans le paragraphe suivant, nous exposons les résultats expérimentaux d'une implémentation parallèle de type maître-esclave de l'algorithme 5.4.

### 5.5.2 Analyse expérimentale du temps d'exécution d'une implémentation parallèle de l'algorithme (GA-CLF)

Pour réduire le temps d'exécution de l'algorithme 5.4 nous utilisons une simulation en parallèle sur une machine multiprocesseurs à mémoire distribuées. Il existe plusieurs façons d'exploiter le parallélisme de l'algorithme 5.4. Nous optons pour l'implémentation maître-esclave expliquée dans la section 5.5. La communication entre les processeurs est faite via le standard MPI, (Message Passing Interface) [67].

**Algorithm 5.5** GA-CLF parallèle

- 
1. Initialiser l'environnement MPI et la configuration des paramètres.
  2. Calculer NP, le nombre d'individus à évaluer par chaque processeur.
  3. Choisir les bornes  $x_i$  et  $x_f$ .
  4.  $t \leftarrow 0$   
 Le processeur maître exécute la tâche suivante : générer aléatoirement une population initiale  $P(t)$  dans l'intervalle  $[x_i, x_f]$
  5.  $t \leftarrow t + 1$   
 Chaque processeur exécute la tâche suivante :  
 pour chaque élément  $P_i$  de la population  $P(t)$  avec  $i = 1, \dots, NP$  faire :
    - (a) Construire  $\Omega_i(t)$
    - (b) Résoudre le problème d'état dans le domaine  $\Omega_i(t)$
    - (c) Calculer  $J(\Omega_i(t), u_i(t))$
 Les données calculées sont échangées par le protocole MPI.
  6. Si  $|J(\Omega_{best}(t), u_{best}(t))| < \varepsilon$  aller à 9.
  7. Exécuter l'opération génétique par le processeur maître :
    - (a) Sélection
    - (b) Croisement
    - (c) Mutation
  8. Si  $\text{mod}(t, N_r) \neq 0$  aller à l'étape 4. Sinon appliquer l'algorithme 5.3 pour recalculer  $x_f$  et  $x_i$  en utilisant l'équation 5.18, puis aller à l'étape 3.
  9. Collecter les résultats de calcul par le processeur maître, puis stopper l'environnement MPI.
- 

**5.5.2.1 Message passing interface MPI**

La communication de données est un élément essentiel de toute formulation parallèle. Cette communication interprocesseur qui est nécessaire pour l'évaluation de la fitness de la population a été obtenue grâce à MPI. Dans cette technique, chaque processeur peut accéder directement

à sa mémoire et doit explicitement communiquer avec d'autres processeurs pour accéder aux données dans leur mémoire [67]. L'objectif de MPI est de fournir un standard pour l'écriture des programmes de passage des messages. Cette norme définit la syntaxe et la sémantique d'un noyau d'une bibliothèque de routines utiles pour écrire des programmes parallèles.

Diverses ressources sur MPI peuvent être trouvées sur<sup>1</sup>.

### 5.5.2.2 Résultats numériques

Nous lançons l'algorithme 5.5 avec les paramètres cités dans le tableau 5.1, la variation du temps de calcul en fonction du nombre de processeurs utilisés ainsi que la vitesse de calcul et l'efficacité sont résumées dans le tableau suivant :

| P  | Taille de problème/proc | Temps     | Vitesse  | Efficacité |
|----|-------------------------|-----------|----------|------------|
| 1  | $16 \times 1296$        | 44.826014 |          |            |
| 2  | $8 \times 1296$         | 22.846483 | 1.980876 | 0.981220   |
| 4  | $4 \times 1296$         | 11.722477 | 3.823937 | 0.955984   |
| 8  | $2 \times 1296$         | 05.829671 | 7.689267 | 0.961160   |
| 16 | $1 \times 1296$         | 08.11634  | 5.522934 | 0.345183   |

TABLE 5.9 : Données de performances

Dans le tableau 5.9, la première ligne représente le temps résultant pour une exécution séquentielle. Ainsi nous remarquons que le temps d'exécution se réduit considérablement. En effet, le seul fait de passer d'un seul processeur à deux processeurs nous permet d'économiser plus que la moitié du temps de calcul. Bien évidemment, l'idéal serait d'avoir la première et la quatrième colonnes identiques. Cependant, dans la pratique, la différence entre les éléments des deux colonnes augmente avec le nombre de processeurs. Comme il a été signalé précédemment, ceci est dû à l'apparition du temps de communication, qui augmente avec le nombre de processeurs. Ceci est confirmé dans la dernière ligne du tableau. En effet, au lieu d'avoir une réduction totale du temps de calcul du fait que chaque processeur va faire une seule évaluation, on remarque en réalité que ce temps a augmenté légèrement. Ceci s'explique par le fait que le temps de communication est devenu plus important que le gain obtenu par la distribution des tâches.

1. <http://www.mpi-forum.org>

Par ailleurs, l'augmentation du temps de communication diminue aussi l'efficacité.

Dans la figure 5.44 nous représentons l'évolution de temps d'exécution, vitesse de calcul  $S_P$  et efficacité  $E_P$  en fonction du nombre de processeurs.

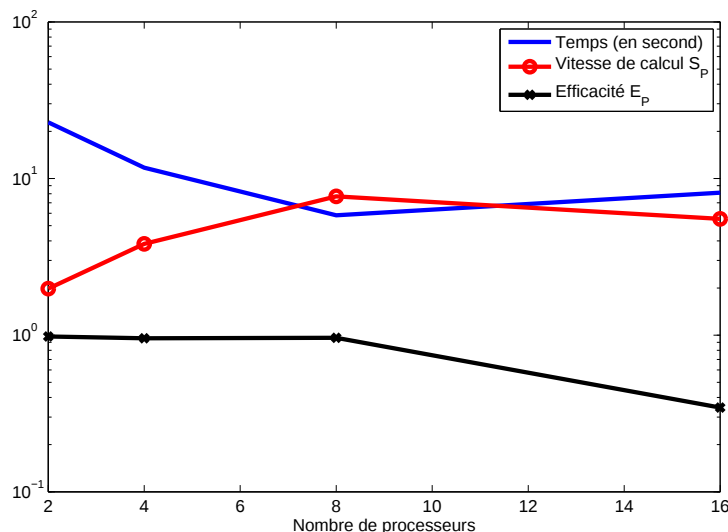


FIGURE 5.44 : Évolution de temps d'exécution, vitesse de calcul  $S_P$  et efficacité  $E_P$  en fonction du nombre de processeurs

## 5.6 Comparaison de l'algorithme de type gradient avec les algorithmes évolutionnaires sur le modèle physique

Dans cette section, nous proposons une comparaison quantitative (en ce qui concerne le temps de calcul) et qualitative (en ce qui concerne la qualité de la solution). Nous précisons que les deux algorithmes sont codés en fortrant 90 et ils sont exécutés sous la même machine<sup>2</sup>, en utilisant le même solveur.

Nous commençons d'abord par la comparaison qualitative.

2. Calculateur SGI, de type grappe de calcul, est composé de 888 cœurs Xeon cadencés à 2.66 GHz avec 2 Go de mémoire par cœur (<http://www.ccipl.univ-nantes.fr>).

### 5.6.1 Comparaison qualitative entre l'algorithme (GA) et l'algorithme du gradient

Nous considérons les résultats obtenus par l'algorithme 5.1 et ceux obtenus par l'algorithme 5.2. La comparaison se restreint à la qualité de la solution approchée par chaque algorithme, étant donné que ces deux algorithmes sont incomparables au niveau temps de calcul.

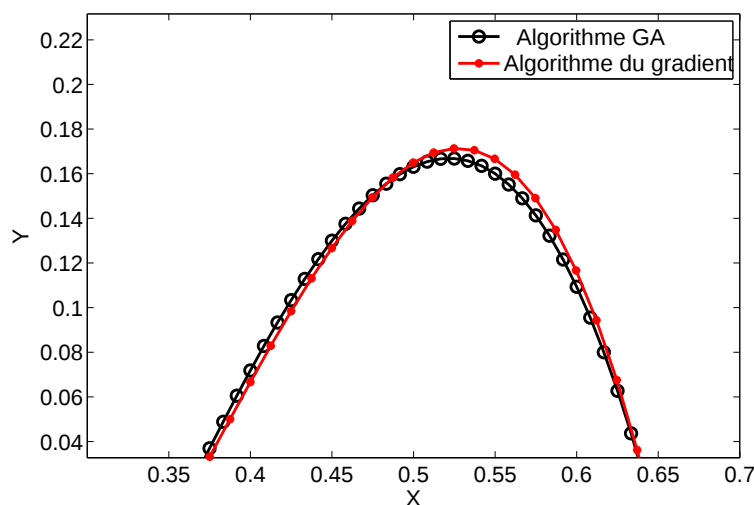


FIGURE 5.45 : Comparaison entre la frontière optimale obtenue par l'algorithme 5.1 et celle obtenue par l'algorithme 5.2

Dans la figure 5.45, nous constatons que les deux frontières obtenues par les deux algorithmes sont de bonnes approximations de la frontière exacte.

### 5.6.2 Comparaison quantitative entre l'algorithme (GA) et l'algorithme du gradient

Nous lançons l'algorithme 5.5 et l'algorithme 5.1 dans la même machine de configuration parallèle. L'algorithme 5.5 est tourné avec 8 processeurs, tandis que l'algorithme 5.1, qui est séquentiel, ne nécessite qu'un seul processeur. Le temps d'exécution de chaque algorithme est cité dans le tableau suivant :

|                                               |           |
|-----------------------------------------------|-----------|
| Temps d'exécution de l'algorithme 5.2         | 87.813993 |
| Temps d'exécution de l'algorithme 5.4         | 44.826014 |
| Temps d'exécution de l'algorithme 5.5 $P = 8$ | 05.829671 |
| Temps d'exécution de l'algorithme 5.1         | 23.625476 |

TABLE 5.10 : Comparaison du temps d'exécution pour l'algorithme 5.5 basé sur la combinaison (GA-CLF) et calcul parallèle et l'algorithme 5.1 basé sur le calcul du gradient

Nous remarquons que, bien que le temps d'exécution de l'algorithme 5.4 demeure plus élevé que celui de l'algorithme 5.1, la parallélisation a permis de le réduire considérablement.

En conclusion, le développement que nous avons apporté à l'algorithme génétique nous a permis d'obtenir un gain très important au niveau du temps d'exécution. Notre but était plutôt de savoir si les algorithmes génétiques pouvaient donner un résultat comparable au niveau de la qualité des solutions, et avec un temps d'exécution relativement raisonnable, ce qui a été démontré par le présent travail. Bien évidemment, cette comparaison ne peut pas avoir lieu sur un problème où l'on ne peut pas appliquer les algorithmes du gradient à cause de la non différentiabilité de la fonctionnelle coût et auquel cas seuls les algorithmes génétiques peuvent être utilisés.



# Conclusion

---

Dans ce travail, nous nous sommes consacrés à l'étude d'un problème d'identification de frontière modélisant un procédé de soudage.

L'approche que nous avons considérée ne s'occupe que de la partie solide de la plaque. Elle consistait à simplifier les phénomènes physiques apparaissant entre la torche de soudage et la plaque, ainsi que ceux du bain liquide, en introduisant une condition de température imposée sur le front de fusion, frontière liquide/solide, que nous devons déterminer. En optant pour une fonctionnelle coût adéquate, nous avons formulé ce problème en un problème d'optimisation de forme.

Ensuite, nous avons démontré l'existence d'une solution optimale. Puis, nous avons étudié l'approximation numérique de ce problème en considérant une discrétisation basée sur les éléments finis. Nous avons montré que le problème discret admet une solution, et nous avons prouvé la convergence d'une suite de solutions du problème approché vers la solution du problème continu.

La difficulté résidait dans le fait que le problème d'état associé à la formulation en optimisation de forme considérée est régi par un opérateur non coercif. Ceci rend l'étude de la continuité du problème d'état compliquée. Pour pallier cette difficulté, nous avons utilisé le degré topologique de Leray Schauder [32], ainsi que des estimations uniformes basées sur des résultats récents sur l'inégalité uniforme de Poincaré [12] et quelques inégalités de Sobolev [51].

Pour la résolution numérique, nous avons utilisé deux méthodes. La première est basée sur les algorithmes déterministes nécessitant le calcul du gradient de forme, tandis que la deuxième méthode repose sur les algorithmes évolutionnaires. Les solutions obtenues par ces deux méthodes ont pratiquement la même allure. Ceci montre que les deux méthodes nous offrent des solutions de même qualité. Bien évidemment, à ce niveau, le temps d'exécution des deux algorithmes est incomparable. Dans le but d'améliorer les algorithmes évolutionnaires afin de les rendre plus compétitifs par rapport aux algorithmes déterministes du type gradient, nous avons proposé un développement des algorithmes évolutionnaires qui s'appuie sur deux apports. Le premier consiste à combiner les algorithmes génétiques avec les contrôleurs de logique floue, alors que

le second réside dans la parallélisation de ce dernier algorithme. Ces apports nous ont permis d'obtenir des résultats de qualité comparable aux autres algorithmes avec un temps d'exécution meilleur que celui obtenu par l'algorithme déterministe du type gradient.

En conclusion, les algorithmes évolutionnaires développés dans la cadre de cette thèse nous offrent un autre moyen efficace pour résoudre les problèmes d'optimisation de forme, surtout lorsque l'analyse de sensibilité pour les méthodes déterministes du type gradient s'avère une tâche non faisable, par exemple lorsque la fonctionnelle coût est non dérivable où lorsque le solveur utilisé ne permet pas d'avoir une meilleure performance.

Les perspectives immédiates de ce travail sont relatives à la complexification des modèles. Il s'agit de les traiter en dimension 3 ou bien de garder la configuration en deux dimensions, en ajoutant un terme correctif. Ce terme contient en général une puissance de la température, ce qui rend le problème non linéaire.

# Bibliographie

- [1] J. Hadamard (1908). Mémoire sur un problème d'analyse relatif à l'équilibre des plaques elastiques encastrées. *Mémoire des savants étrangers.*, Oeuvres de J. Hadamard, CNRS, Paris, 1968. (Cité en page 75.)
- [2] J. Abouchabaka, R. Aboulaich, O. Guennoun, A. Nachaoui, and A. Souissi. Shape optimization for a simulation of a semiconductor problem. *Math. Comput. Simulation*, 56(1) :1–16, 2001. (Cité en page ii.)
- [3] G. Allaire. *Analyse numérique et optimisation : une introduction à la modélisation mathématique et à la simulation numérique*. Mathématiques appliquées. École polytechnique, 2005. (Cité en page iii.)
- [4] G. Allaire. *Conception optimale de structures*, volume 58 of *Mathématiques & Applications (Berlin) [Mathematics & Applications]*. Springer-Verlag, Berlin, 2007. With the collaboration of Marc Schoenauer (INRIA) in the writing of Chapter 8. (Cité en pages ii et 76.)
- [5] G. Allaire, F. de Gournay, F. Jouve, and A.-M. Toader. Structural optimization using topological and shape sensitivity via a level set method. *Control Cybernet.*, 34(1) :59–80, 2005. (Cité en page 76.)
- [6] A. Auger and N. Hansen. Theory of evolution strategies : a new perspective. In A. Auger and B. Doerr, editors, *Theory of Randomized Search Heuristics : Foundations and Recent Developments*, chapter 10, pages 289–325. World Scientific Publishing, 2011. (Cité en page 90.)
- [7] G. K. Batchelor. *An introduction to fluid dynamics*. Cambridge Mathematical Library. Cambridge University Press, Cambridge, paperback edition, 1999. (Cité en page i.)
- [8] D. Begis and R. Glowinski. Application de la méthode des éléments finis à l'approximation d'un problème de domaine optimal. Méthodes de résolution des problèmes approchés. *Appl. Math. Optim.*, 2(2) :130–169, 1975/76. (Cité en page 66.)
- [9] P. Beremlijski, J. Haslinger, M. Kočvara, R. Kučera, and J. V. Outrata. Shape optimization in three-dimensional contact problems with Coulomb friction. *SIAM J. Optim.*, 20(1) :416–444, 2009. (Cité en page ii.)
- [10] J. M. Bergheau. Numerical simulation of welding. *Revue européenne des éléments finis*, 13(3-4), 2004.

- 
- [11] J. M. Bergheau and R. Fortunier. *Finite element simulation of heat transfer*. ISTE, London, 2008. Translated from the 2004 French original by Robert Meillier. (Cité en page ii.)
- [12] A. Boulkhemair and A. Chakib. On the uniform Poincaré inequality. *Comm. Partial Differential Equations*, 32(7-9) :1439–1447, 2007. (Cité en pages ii, iii, 7, 23 et 135.)
- [13] H. Brezis. *Analyse fonctionnelle*. Collection Mathématiques Appliquées pour la Maîtrise. [Collection of Applied Mathematics for the Master’s Degree]. Masson, Paris, 1983. Théorie et applications. [Theory and applications]. (Cité en pages iii, 2 et 23.)
- [14] A. Chakib. *Etude de quelques problèmes d’écoulement dans les milieux poreux homogènes et non homogène par le méthode d’optimisation de forme*. PhD thesis, Université Mohammed V-Agdal, Faculté des Sciences RABAT, 2000. (Cité en page 57.)
- [15] A. Chakib, A. Ellabib, A. Nachaoui, and M. Nachaoui. A shape optimization formulation of weld pool determination. *Applied Mathematics Letters*, 2011. <http://dx.doi.org/10.1016/j.aml.2011.09.017>. (Cité en page 23.)
- [16] A. Chakib, T. Ghemires, and A. Nachaoui. An optimal shape design formulation for inhomogeneous dam problems. *Math. Methods Appl. Sci.*, 25(6) :473–489, 2002. (Cité en page ii.)
- [17] A. Chakib, T. Ghemires, and A. Nachaoui. A numerical study of filtration problem in inhomogeneous dam with discontinuous permeability. *Appl. Numer. Math.*, 45(2-3) :123–138, 2003. (Cité en pages 57 et 75.)
- [18] A. Chakib and A. Nachaoui. Nonlinear programming approach for a transient free boundary flow problem. *Appl. Math. Comput.*, 160(2) :317–328, 2005. (Cité en page ii.)
- [19] A. Chakib and A. Nachaoui. Simulation of a free boundary problem using boundary element method and restarted re-analysis technique. *Appl. Math. Comput.*, 177(2) :824–838, 2006. (Cité en page ii.)
- [20] A. Chakib, A. Nachaoui, and M. Nachaoui. Shape sensitivity analysis for a free boundary welding problem. In *4th International Conference on Approximation Methods and Numerical Modelling in Environment and Natural Resources*, pages 274–277. (Cité en page 74.)
- [21] A. Chakib, A. Nachaoui, and M. Nachaoui. Numerical simulation of a heat transfer problem, via a shape optimization method. In *10th International Conference on Computational and*

- Mathematical Methods in Science and Engineering (CMMSE)*., pages 359–362, Almeria Spain., 26-30 june 2010. (Cité en page 51.)
- [22] A. Chakib, A. Nachaoui, and M. Nachaoui. Approximation and numerical realization of an optimal design welding problem. *Numerical Methods for Partial Differential Equations*, 2011. (Cité en page 51.)
- [23] A. Chakib, A. Nachaoui, and M. Nachaoui. Finite element approximation of an optimal design problem. *Appl. Comput. Math.*, 2011. (Cité en page 51.)
- [24] A. Chakib, A. Nachaoui, and M. Nachaoui. Parallel evolutionary algorithms, for solving a free boundary problem. In *International Conference on Applied Mathematics, Modeling & Computational Science (AMMCS-2011)*, Waterloo, Canada, July 25-29 2011. (Cité en page 74.)
- [25] D. Chenaïs. On the existence of a solution in a domain identification problem. *J. Math. Anal. Appl.*, 52(2) :189–219, 1975. (Cité en pages 6, 7 et 60.)
- [26] Donal O'Regan Yeol Je Cho and Yu-Qing Chen. *TOPOLOGICAL DEGREE THEORY AND APPLICATIONS*, volume 10. Chapman & Hall/CRC Taylor & Francis Group, 6000 Broken Sound Parkway NW, Suite 300 Boca Raton, FL 33487-2742, 2006. (Cité en page 9.)
- [27] P. G. Ciarlet. *Introduction à l'analyse numérique matricielle et à l'optimisation*. Collection Mathématiques Appliquées pour la Maîtrise. [Collection of Applied Mathematics for the Master's Degree]. Masson, Paris, 1982.
- [28] P. G. Ciarlet. *The finite element method for elliptic problems*, volume 40 of *Classics in Applied Mathematics*. Society for Industrial and Applied Mathematics (SIAM), Philadelphia, PA, 2002. Reprint of the 1978 original [North-Holland, Amsterdam ; MR0520174 (58 #25001)]. (Cité en pages ii, 56 et 64.)
- [29] J. Cronin. *Fixed points and topological degree in nonlinear analysis*. Mathematical Surveys, No. 11. American Mathematical Society, Providence, R.I., 1964. (Cité en page 8.)
- [30] R. Dautray and J. L. Lions. *Analyse mathématique et calcul numérique pour les sciences et les techniques. Tome 3*. Collection du Commissariat à l'Énergie Atomique : Série Scientifique. [Collection of the Atomic Energy Commission : Science Series]. Masson, Paris, 1985. With the collaboration of Michel Artola, Claude Bardos, Michel Cessenat, Alain Kavenoky,

- Hélène Lanchon, Patrick Lascaux, Bertrand Mercier, Olivier Pironneau, Bruno Scheurer and Rémi Sentis. (Cité en pages 2, 4 et 25.)
- [31] K. Deep and M. Thakur. A new mutation operator for real coded genetic algorithms. *Appl. Math. Comput.*, 193(1) :211–230, 2007. (Cité en pages iv et 96.)
- [32] K. Deimling. *Nonlinear functional analysis*. Springer-Verlag, Berlin, 1985. (Cité en pages ii, iii, 8, 23 et 135.)
- [33] D. D. Doan, F. Gabriel, Y. Jarny, and P. Le-Masson. Estimation de la forme du bain fondu en quasi-stationnaire dans un procédé de soudage tig- résultats numériques et expérimentaux. In *Actes du congrès SFT 2007, Ile des Embiez*, 2007. (Cité en page ii.)
- [34] G. Farin. *Curves and surfaces for computer-aided geometric design*. Computer Science and Scientific Computing. Academic Press Inc., San Diego, CA, fourth edition, 1997. A practical guide, Chapter 1 by P. Bézier ; Chapters 11 and 22 by W. Boehm, With 1 IBM-PC floppy disk (3.5 inch ; HD). (Cité en page iv.)
- [35] David B. Fogel and Lawrence J. Fogel. An introduction to evolutionary programming. In *Artificial Evolution'95*, pages 21–33, 1995. (Cité en page 90.)
- [36] A. Friedman. *Foundations of modern analysis*. Holt, Rinehart and Winston, Inc., New York, 1970. (Cité en page 5.)
- [37] Stéphane Garreau, Philippe Guillaume, and Mohamed Masmoudi. The topological asymptotic for PDE systems : the elasticity case. *SIAM J. Control Optim.*, 39(6) :1756–1778 (electronic), 2001. (Cité en page 76.)
- [38] D. E. Goldberg. Genetic algorithms and Walsh functions. I. A gentle introduction. *Complex Systems*, 3(2) :129–152, 1989. (Cité en page 90.)
- [39] L. Groupe. *Calcul Parallèle : Calcul Distribué, Grappe de Serveurs, Parallélisme, Altivec, Thread Local Storage*. General Books LLC, 2010. (Cité en page iii.)
- [40] Hans-Paul and P. Schwefel. *Evolution and Optimum Seeking : The Sixth Generation*. John Wiley & Sons, Inc., New York, NY, USA, 1993. (Cité en page 90.)
- [41] J. Haslinger and R. A. E. Mäkinen. *Introduction to shape optimization*, volume 7 of *Advances in Design and Control*. Society for Industrial and Applied Mathematics (SIAM), Philadelphia, PA, 2003. Theory, approximation, and computation. (Cité en pages 17, 24, 52, 57, 62 et 75.)

- [42] J. Haslinger and P. Neittaanmäki. *Finite Element Approximation for Optimal Shape Design : Theory and Applications*. John Wiley and Sons, Chichester, 1988. (Cité en pages [ii](#), [iv](#), [5](#), [57](#), [62](#), [63](#) et [75](#).)
- [43] J. Haslinger and J. Stebel. Shape optimization for Navier-Stokes equations with algebraic turbulence model : numerical analysis and computation. *Appl. Math. Optim.*, 63(2) :277–308, 2011. (Cité en page [ii](#).)
- [44] A. Henrot and M. Pierre. *Variation et optimisation de formes*, volume 48 of *Mathématiques & Applications (Berlin) [Mathematics & Applications]*. Springer, Berlin, 2005. Une analyse géométrique. [A geometric analysis]. (Cité en pages [ii](#) et [17](#).)
- [45] M. Hinze and S. Ziegenbalg. Optimal control of the free boundary in a two-phase stefan problem. *J. Comput. Phys.*, 223 :657–684, May 2007.
- [46] I. Hlaváček and J. Nečas. Optimization of the domain in elliptic unilateral boundary value problems by finite element method. *RAIRO Anal. Numér.*, 16(4) :351–373, 1982.
- [47] J. H. Holland. *Adaptation in natural and artificial systems*. University of Michigan Press, Ann Arbor, Mich., 1975. An introductory analysis with applications to biology, control, and artificial intelligence. (Cité en page [90](#).)
- [48] Y. F. Hsu and B. Rubinsky. An inverse finite element method for the analysis of stationary arc welding processes. *J. Heat Transfer*, 108 :734–742, November 1986. (Cité en page [i](#).)
- [49] N. Kerrouault. *Fissuration à chaud en soudage d'un acier inoxydable austénitique*. PhD thesis, Ecole Centrale de Paris, 2001. (Cité en page [13](#).)
- [50] J. R. Koza. Genetic programming : automatic synthesis of topologies and numerical parameters. In *Handbook of metaheuristics*, volume 57 of *Internat. Ser. Oper. Res. Management Sci.*, pages 83–104. Kluwer Acad. Publ., Boston, MA, 2003. (Cité en page [90](#).)
- [51] O. A. Ladyženskaja and N. N. Ural'ceva. *Équations aux dérivées partielles de type elliptique*. Traduit par G. Roos. Monographies Universitaires de Mathématiques, No. 31. Dunod, Paris, 1968. (Cité en pages [ii](#), [iii](#), [7](#), [23](#) et [135](#).)
- [52] J. Leray and J. L. Lions. Quelques résultats de Višik sur les problèmes elliptiques nonlinéaires par les méthodes de Minty-Browder. *Bull. Soc. Math. France*, 93 :97–107, 1965.
- [53] J. Leray and J. Schauder. Topologie et équations fonctionnelles. *Ann. Sci. École Norm. Sup. (3)*, 51 :45–78, 1934. (Cité en page [9](#).)

- [54] J. L. Lions and E. Magenes. *Problèmes aux limites non homogènes et applications. Vol. 1.* Travaux et Recherches Mathématiques, No. 17. Dunod, Paris, 1968.
- [55] K. Maarten, J. J. Merelo, G. Romero, and M. Schoenauer. Evolving objects : A general purpose evolutionary computation library. *Artificial Evolution*, 2310 :829–888, 2002.
- [56] H. Makhlouf. *Modélisation numérique du soudage à l’arc des aciers.* PhD thesis, École Nationale Supérieure des Mines de Paris, 2008. (Cité en page i.)
- [57] M. Masmoudi. *Outils pour la conception optimale de forme.* PhD thesis, Université de Nice, 1987. (Cité en page 75.)
- [58] Z. Michalewicz. *Genetic algorithms + data structures = evolution programs.* Springer-Verlag, Berlin, second edition, 1994. (Cité en pages ii, 89 et 92.)
- [59] F. Murat and J. Simon. *Sur le contrôle par un domaine géométrique.* Thèse d’état, Université Pierre et Marie Curie, Paris, 1976. (Cité en page 75.)
- [60] F. Murat and L. Tartar. Calcul des variations et homogénéisation. In *Homogenization methods : theory and applications in physics (Bréau-sans-Nappe, 1983)*, volume 57 of *Collect. Dir. Études Rech. Élec. France*, pages 319–369. Eyrolles, Paris, 1985. (Cité en page 76.)
- [61] M. Nachaoui and R. Y. Shikhlinakaya. Hybrid defuzzification method for fuzzy controllers. *Dokl. Nats. Akad. Nauk Azerb.* (Cité en pages 74, 75, 112, 114 et 120.)
- [62] J. Nečas. *Les méthodes directes en théorie des équations elliptiques.* Masson et Cie, Éditeurs, Paris, 1967. (Cité en pages 5, 6 et 32.)
- [63] A. Nehad. Application of optimisation methods for solving inverse phase-change problems. *J. Heat Transfer*, Part B,31 :477–497, 1997. (Cité en page i.)
- [64] L. Nirenberg. *Topics in nonlinear functional analysis.* Courant Institute of Mathematical Sciences New York University, New York, 1974. With a chapter by E. Zehnder, Notes by R. A. Artino, Lecture Notes, 1973–1974. (Cité en page 8.)
- [65] J. Nocedal and S. J. Wright. *Numerical optimization.* Springer Series in Operations Research. Springer-Verlag, New York, 1999. (Cité en page 75.)
- [66] S. J. Osher and F. Santosa. Level set methods for optimization problems involving geometry and constraints. I. Frequencies of a two-density inhomogeneous drum. *J. Comput. Phys.*, 171(1) :272–288, 2001. (Cité en page 76.)

- 
- [67] P. S. Pacheco. *Parallel programming with MPI*. Morgan Kaufmann Publishers Inc., San Francisco, CA, USA, 1996. (Cité en pages 95, 128 et 130.)
  - [68] P. C. Paris. *The fracture mechanics approach to fatigue, Fatigue, an interdisciplinary approach*. Syracuse University Press, Syracuse, 1964.
  - [69] O. Pironneau. *Optimal shape design for elliptic systems*. Springer Series in Computational Physics. Springer-Verlag, New York, 1984. (Cité en pages 7, 17 et 31.)
  - [70] E. Polak. *Optimization*, volume 124 of *Applied Mathematical Sciences*. Springer-Verlag, New York, 1997. Algorithms and consistent approximations. (Cité en page 75.)
  - [71] P.-A. Raviart and J.-M. Thomas. *Introduction à l'analyse numérique des équations aux dérivées partielles*. Collection Mathématiques Appliquées pour la Maîtrise. [Collection of Applied Mathematics for the Master's Degree]. Masson, Paris, 1983. (Cité en page 4.)
  - [72] I. Rechenberg. Evolutionsstrategie : Optimierung technischer systeme nach prinzipien der biologischen evolution. *Leice Transactions*, 1994. (Cité en page 90.)
  - [73] B. Rousselet. *Quelques résultats en optimisation de domaine*. PhD thesis, Université de Nice, 1982. (Cité en page 75.)
  - [74] W. Rudin. *Functional analysis*. International Series in Pure and Applied Mathematics. McGraw-Hill Inc., New York, second edition, 1991. (Cité en page 5.)
  - [75] M. Sefrioui, J. Périaux, and J.-G. Ganascia. Fast convergence thanks to diversity. In *Evolutionary Programming'96*, pages 313–321, 1996.  
:fast :fast :fast :fast
  - [76] R. Y. Shikhlinakaya. *Analysis of the WABL Method for Fuzzy Numbers and its Applications to Fuzzy Models*. Dr.-Ing. Thesis, Baku city university Azerbaijan, 2007. (Cité en page 114.)
  - [77] J. Sokolowski and J.P. Zolésio. *Introduction to shape optimisation*. Springer Devices in Computational Mathematics, Berlin, 1992. (Cité en page 75.)
  - [78] J. Sokołowski and A. Żochowski. On the topological derivative in shape optimization. *SIAM J. Control Optim.*, 37(4) :1251–1272 (electronic), 1999. (Cité en page 76.)
  - [79] G. Stampacchia. Le problème de Dirichlet pour les équations elliptiques du second ordre à coefficients discontinus. *Ann. Inst. Fourier (Grenoble)*, 15(fasc. 1) :189–258, 1965. (Cité en page 5.)

- 
- [80] T. Takagi and M. Sugeno. Fuzzy identification of systems and its applications to modeling and control. *IEEE Transactions on Systems, Man and Cybernetics*, 15(1) :116–132, 1985. (Cité en page [iv](#).)
- [81] J. I. Toivanen, J. Haslinger, and R. A. E. Mäkinen. Shape optimization of systems governed by Bernoulli free boundary problems. *Comput. Methods Appl. Mech. Engrg.*, 197(45–48) :3803–3815, 2008. (Cité en page [ii](#).)
- [82] K. Dang Van. Sur la résistance à la fatigue des métaux. *Science et technique de l’armement*, 3(47) :641–722, 1973. (Cité en page [13](#).)
- [83] T. Zacharia, J. M. Vitek, J. A. Goldak, T. A. DebRoy, M. Rappaz, and K. H. Bhadeshia. Modeling of fundamental phenomena in welds. *Modelling and Simulation in Materials Science and Engineering*, 3(2) :265, 1995. (Cité en pages [i](#) et [82](#).)
- [84] L. A. Zadeh. Fuzzy sets. *Information and Control*, 8 :338–353, 1965. (Cité en pages [iv](#), [105](#) et [106](#).)
- [85] L. A. Zadeh. Outline of a new approach to the analysis of complex systems and decision processes. *Information Science*, 9 :43–80, 1973. (Cité en page [iv](#).)
- [86] L. A. Zadeh. The concept of a linguistic variable and its application to approximate reasoning, i–ii. *Information Science*, 8 :199–249, 301–357, 1975. (Cité en page [iv](#).)
- [87] J.P. Zolésio. The material derivative (or speed) method for shape optimisation, in optimisation of distributed parameter structures. *E.J Haug and J.Cea eds.*, II :1098–1151, 1981. (Cité en pages [iii](#) et [75](#).)

# Résumé

Nous nous intéressons à l'étude d'un problème d'analyse des transferts de chaleur qui modélise une opération de soudage.

L'approche que nous considérons ne s'occupe que de la partie solide de la plaque. Elle consiste à résoudre un problème à frontière libre. Pour cela, nous proposons une formulation en optimisation de forme. Le problème d'état est gouverné par un opérateur qui, pour certaines données, n'est pas coercif. Cela complique l'étude de la continuité du problème d'état. Nous surmontons cette difficulté en utilisant le degré topologique de Leray-Schauder, ainsi nous montrons l'existence d'un domaine optimal.

Ensuite, nous considérons une discrétisation de ce problème basée sur les éléments finis linéaires. Nous prouvons alors que le problème discret admet une solution et nous montrons qu'une sous-suite des solutions de ce problème converge vers la solution du problème continu.

Enfin, nous présentons des résultats numériques réalisés par deux méthodes : la méthode déterministe basée sur le calcul du gradient de forme, et les algorithmes génétiques combinés avec la logique floue et le calcul parallèle. Ainsi une étude comparative de ces deux méthodes aux niveaux qualitatif et quantitatif a été présentée.



# Abstract

We are interested in studying a heat transfer problem which modeling a welding process. The approach that we consider deals only with the solid part of the plate. It consists in solving a free boundary problem. For this, we propose a shape optimization formulation. The state problem governed by an operator which for some data is not coercif. This complicates the study of the continuity of the state problem. We overcome this difficulty using the topological degree of Leray-Schauder and we show the existence of an optimal domain. Next, we consider a discretization of this problem based on linear finite elements. We prove that the approximate problem is solvable and we show that a subsequence of the solution of this approximate problem converges to the solution of the continuous problem.

Finally, we present numerical results achieved by two methods : the deterministic method based on the gradient-likes method and genetic algorithms combined with fuzzy logic and parallel computing. A comparative study of two methods for qualitative and quantitative levels was presented.