

Ginzburg-Landau et champs de croix

Apports liés aux applications de maillage

Mémoire présenté par
Alexis MACQ

en vue de l'obtention du grade de Master
ingénieur civil en mathématiques appliquées

Promoteurs
Jean-François REMACLE, Jean VAN SCHAFTINGEN

Lecteurs
Pierre-Antoine ABSIL, Pierre-Alexandre BEAUFORT, Rémi RODIAC

Année académique 2017-2018

Remerciements

Rendre ce mémoire suscite des sentiments contrastés. Il y a d'une part la fierté d'avoir, par l'expérience, acquis une certaine compréhension de problèmes captivants. Et il y a d'autre part la déception face au peu de chemin parcouru d'autant que les ambitions originelles de ce mémoire étaient élevées.

L'expression selon laquelle les hommes de sciences ont la chance de pouvoir agrandir leur vision en montant sur les épaules des géants les ayant précédés serait due à Bernard de Chartres. Elle met en évidence l'importance dans toute entreprise intellectuelle, et particulièrement scientifique, de s'appuyer sur les travaux des grands penseurs du passé.

Je n'ai certainement pas la prétention d'avoir pu me hisser sur les épaules de géants, mais j'ai été guidé par des personnes dont l'intuition et les perspectives manifestent la profondeur des connaissances. Cela m'a permis d'éviter de nombreuses écueils et d'emprunter des chemins parfois ardues mais menant à des notions et résultats passionnants. Mes remerciements leur sont principalement adressés.

Je commencerai par M. Van Schaftingen qui m'a donné l'opportunité de faire ce mémoire. Son soutien, sa confiance dès les premières heures de travail, le temps qu'il m'a consacré et sa pédagogie m'ont été précieux. J'aimerais aussi remercier Rémi Rodiac pour ses conseils, son écoute, son rôle de guide et son indulgence lorsque je pataugeai dans des domaines mathématiques qui m'étaient obscures. Grâce à eux, j'ai vraiment l'impression d'avoir compris des éléments avancés de mathématiques.

Je voudrai ensuite remercier Pierre-Alexandre Beaufort qui m'a guidé vers des expériences fructueuses et qui m'a permis de pouvoir utiliser rapidement des outils informatiques complexes toujours en cours de développement ; merci aussi de m'avoir patiemment reçu de nombreuses fois pour répondre à mes incessantes interrogations.

Merci à M. Remacle pour m'avoir mis sur de bons rails et entre de bonnes mains. Sans sa confiance, je n'aurais jamais pu m'attaquer au problème passionnant dont ce travail de fin d'études est l'objet.

Je remercie aussi Élisabeth Jeannot pour son intérêt et sa relecture orthographique. Et je terminerai en remerciant ma chère et tendre Charlotte pour sa relecture orthographique et pour son soutien de tous les instants.

Un grand et sincère merci à tous.

La connaissance s'acquiert par
l'expérience, tout le reste n'est que de
l'information.

Albert Einstein

Table des matières

1	Introduction	7
2	Motivations	10
2.1	Caractéristique d'Euler-Poincaré	10
2.2	Applications en maillages	10
2.3	Applications en cristallographie et virologie	12
3	Théorème de la boule chevelue	14
3.1	Motivation du théorème	14
3.2	Variétés différentielles et sphère de Riemann	14
3.3	Relèvement et revêtement	15
3.4	Démonstration	17
4	Existence de minimiseurs	22
4.1	Introduction	22
4.2	Dérivée faible	22
4.2.1	Cas multidimensionnel	23
4.3	Espace de Sobolev	24
4.4	Espace H^1	26
4.5	Convergence faible	27
4.6	Dérivée covariante	27
4.6.1	Définition d'un changement de base covariant	27
4.6.2	Introduction à la dérivée covariante	27
4.6.3	Projection sur l'espace tangent	29
4.7	Théorème de Rellich	30
4.8	Compacité faible	30
4.9	Existence	31
5	Équations d'Euler-Lagrange	35
5.1	Introduction	35
5.2	Minimisation et équation d'Euler-Lagrange	35
5.3	Comment trouver les équations d'Euler-Lagrange	35
5.4	Dérivation des équations d'Euler-Lagrange	36
6	Régularité des solutions	40
6.1	Introduction	40
6.2	Injections de Sobolev	40
6.3	Régularité des solutions	41

7	Convergence dans le cas du tore vers une application harmonique	43
7.1	Introduction	43
7.2	Convergence des minimiseurs vers des applications harmoniques dans le cas du tore	43
7.2.1	Existence de solutions	43
7.2.2	Convergence des minimiseurs	44
7.2.3	Équation d'Euler-Lagrange et caractère harmonique	45
7.2.4	Convergence des gradients covariants	45
8	Etude de l'apparition de singularités dans un cas simple	47
8.1	Introduction	47
8.2	Variété lisse	47
8.3	Principe du maximum	48
8.4	Un modèle simple	48
8.4.1	Existence de minimiseurs	50
8.4.2	Forme des minimiseurs pour $\epsilon \rightarrow 0$	51
8.4.3	Analyse asymptotique	51
8.5	Apparition de tourbillons	53
8.5.1	Borne sur l'énergie pour une sphère	55
8.5.2	Résultats exacts	56
9	Formulation d'un champ de croix	58
9.1	Polynômes homogènes	58
9.2	Polynômes harmoniques	58
9.3	Définition et propriétés d'un champ de croix	59
9.4	Exemple de définition de champ de croix	59
9.4.1	Norme d'un champ de croix	60
9.5	Conclusion	60
10	Passage numérique en régime asymptotique	61
10.1	Introduction	61
10.1.1	Rappels concernant les éléments finis	62
10.1.2	Description rapide de l'outil de maillage de GMSH	63
10.2	Description du solveur numérique	64
10.2.1	Exemple d'application possible	65
10.3	Expériences	65
10.3.1	Première expérience	66
10.3.2	Deuxième expérience	75
10.4	Conclusion des expériences	84
10.5	Remarques supplémentaires	85
11	Conclusion	86
A	Le degré de Brouwer	88
A.1	Rappels sur l'orientation des espaces vectoriels	88
A.2	Orientation des variétés	88
B	Analyse statistique	90
B.1	Graphes concernant le rejet de certaines données dans la première expérience . .	90
B.2	Données supplémentaires concernant l'expérience 1	91
B.3	Un des scripts bash utilisé pour l'analyse statistique	92

Chapitre 1

Introduction

Lorsque vous voulez résoudre de façon approximée une équation faisant intervenir des dérivées spatiales sur une surface fermée, une technique classique consiste à fabriquer une mosaïque sur la surface et à imposer une valeur approximative de la solution de l'équation étudiée en un point de fragment de mosaïque via l'approximation des valeurs de l'équation aux sommets du fragment auquel appartient ce point. Ces méthodes sont appelées méthodes de maillage.

La première étape consiste à voir quels sont les différents types de maillages existants. Il est de coutume de choisir un maillage quadrangulaire (en carrés déformés) ou triangulaire, mais rien n'interdit de choisir des maillages plus complexes. Il n'est pas non plus nécessaire que le maillage soit régulier et il y a une tendance qui vise à resserrer le maillage près des endroits d'intérêt (par exemple aux endroits où nous pensons que la solution va beaucoup varier) ; cependant, il faut veiller à avoir des éléments faiblement distordus (c'est-à-dire dont la forme se rapproche d'un polygone régulier). Plus ce maillage est resserré, plus la solution que nous obtenons par la méthode des éléments finis sera précise et proche de la "vraie" solution de l'équation aux dérivées partielles. Il existe donc de nombreux type de maillages, mais il faut savoir lequel choisir dans quel contexte et pourquoi. Pour effectuer ce type de choix, il existe au moins deux écoles.

La première cherche à résoudre l'équation en cours d'étude sur la forme voulue à partir d'un maillage quelconque et de raffiner ce maillage aux endroits qui impliquent le plus d'erreurs. L'avantage de cette méthode est qu'elle permet une relative maîtrise de l'erreur moyenne qu'elle implique, mais le maillage dépend de l'équation et ne présente a priori pas de régularité particulière.

La deuxième école va, quant à elle, mailler la forme voulue avec un maillage le plus régulier possible. Bien que cela ne permette pas de maîtriser l'erreur pour un nombre de fragments de mosaïque donné, la régularité du maillage permet d'avoir des indices d'éléments ayant une grande cohérence sur la forme. De plus, les maillages réguliers peuvent se voir appliquer de manière efficace des opérateurs liés aux équations à résoudre sur la forme (cette remarque est particulièrement vraie pour des éléments de type quadrangle). Notons encore que cette méthode permet d'utiliser un même maillage pour différentes équations. Les avantages notables de l'utilisation de quadrangles sont les suivants : nous avons moins d'éléments par nœud (il y a environ la moitié moins de quadrangles que de triangles pour un même nombre de sommets comme nous pouvons le déduire de la formule de la caractéristique d'Euler-Poincaré dont nous parlerons plus tard) ; les quadrangles sont faciles à raffiner ; enfin, ils sont très adaptés aux calculs tensoriels et donc à des généralisations simples d'opérateurs 1D classiques.

Ce mémoire s'intéresse à la deuxième méthode évoquée ci-dessus pour la production de maillages quadrangulaires (un exemple d'un tel maillage pour un tore est donné à la figure 1.1 ci-dessous).

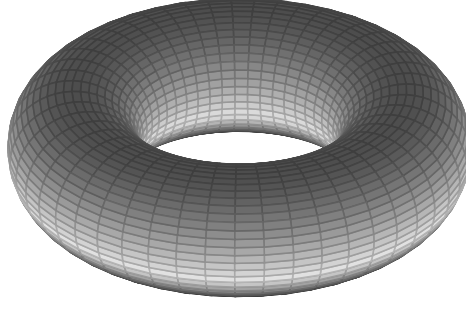


FIGURE 1.1 – Un découpe relativement fine en quadrangles du tore T^2 .

Des outils telle que la caractéristique d'Euler-Poincaré mettent l'accent sur les limites de cette méthode. En effet, comme nous le verrons, la caractéristique d'Euler-Poincaré de la sphère nous permet de dire par exemple que celle-ci ne peut pas être maillée en n'utilisant que des quadrangles (la figure 1.2 ci-dessous présente un exemple de maillage de la sphère).

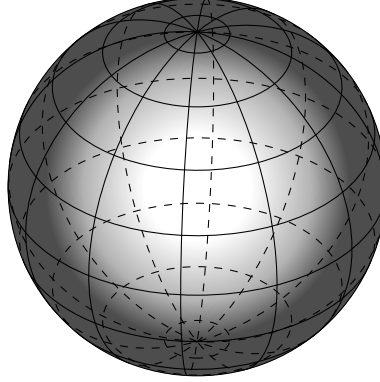


FIGURE 1.2 – Un découpage grossier en polygones curvilignes (triangles et quadrangles) de la sphère S^2 .

Dès lors, pour les formes qui ne peuvent pas être maillées uniquement avec des quadrangles, il faut accepter de placer dans le maillage des éléments qui ne correspondront pas à des quadrangles et donc des intersections à plus ou moins de quatre branches dans le maillage. Ce mémoire a pour objet d'apporter une contribution à la question suivante liée à ce type de maillage. Comment placer les intersections singulières (dont le nombre de branches n'est pas égale à quatre) d'un maillage de type quadrangulaire le plus régulier possible sur une forme qui ne peut pas, à cause d'arguments topologiques, être maillée par des quadrangles. Dans ce cadre, une famille d'équations empruntée à une description phénoménologique due à Ginzburg et Landau de la physique des supraconducteurs sera utilisée. Le but de ces équations étant de décrire des brisures de symétrie dans des supraconducteurs subissant des champs magnétiques variables. Notons toutefois que la famille d'équations utilisée ne correspond pas directement à celle décrite par Ginzburg et Landau mais correspond plutôt à une version simplifiée (il n'est notamment pas question de conditions au bord de type "Dirichlet" lorsqu'un supraconducteur subi un champ magnétique alors que nous utiliserons de telles conditions dans la suite).

Ce mémoire a donc pour objet l'étude de fonctionnelles dite "de Ginzburg-Landau" et particulièrement, mais pas uniquement, de la suivante :

$$E_\epsilon(u) = \min_{u \in H^1(\Omega) \text{ et } u \in T_x(\Omega)} \int_{\Omega} \frac{1}{2} |\nabla_c u|^2 + \frac{1}{4\epsilon^2} (1 - |u|^2)^2 \quad (1.1)$$

sur des surfaces ou sur des variétés Ω de dimension deux plongées dans $\mathbb{R}^3$ (nous parlerons

notamment dans la suite du tore T^2 et de la sphère S^2). Il sera aussi question de l'étude du comportement asymptotique de l'équation (1.1), c'est-à-dire son comportement lorsque $\epsilon \rightarrow 0$. La fonctionnelle (1.1) est composée de deux termes. Le premier pénalise les variations de u qui ne proviennent pas de la courbure de Ω (d'où l'utilisation d'un gradient covariant) et le second pénalise les points de la variété sur lesquels la norme de u est différente de 1. Quand $\epsilon \rightarrow 0$, $|u_\epsilon| \rightarrow 1$ sauf en un nombre fini ou nul (si notre forme peut être maillée par des quadrangles) de points dont la position peut être décrite par une énergie renormalisée. Les positions de ces points correspondent à des positions intéressantes pour placer des nœuds singuliers de façon à limiter les zones non régulières de maillages quadrangulaires sur les surfaces ne se laissant pas mailler qu'avec des quadrangles. Les notions utiles à la compréhension de cette fonctionnelle et des espaces de fonctions utilisés seront développés tout au long de ce mémoire.

L'étude de fonctionnelles dite de Ginzburg-Landau avec un terme pénalisant les changements d'orientation et un terme pénalisant les éléments n'ayant pas une norme égale à 1, particulièrement lorsque les fonctions manipulées correspondent à des champs de croix, peut donc être mise en rapport avec les nœuds irréguliers et les changements d'orientation de maillages quadrangulaires que nous souhaitons les plus réguliers possibles.

Ce mémoire est composé de plusieurs parties. Nous commencerons par la présentation de deux domaines dans lesquelles des fonctionnelles analogues à (1.1) et les maillages sont intéressants. Il sera ensuite question des limitations topologiques des découpes de formes via une démonstration basée sur un maillage de la sphère S^2 du théorème dit "de la boule chevelue". Nous démontrerons ensuite l'existence de minimiseurs de (1.1) pour tout type de forme. Cette démonstration sera suivie de la détermination des équations d'Euler-Lagrange associée à (1.1). Nous étudierons ensuite la régularité des solutions. Il sera alors question de la convergence des minimiseurs de (1.1) vers des applications analogues en un certain sens à des applications harmoniques dans le cas où la surface sur laquelle nous travaillons est un tore. Nous étudierons ensuite l'apparition de singularités sur des surfaces 2D ayant un bord pour une fonctionnelle légèrement plus simple que (1.1). Nous essayerons ensuite d'introduire la notion de champ de croix avec laquelle il serait intéressant d'adapter légèrement les parties précédentes de ce mémoire. Nous terminerons avec l'analyse d'une méthode numérique basée sur une fonctionnelle du même type, mais légèrement différente de (1.1). Nous verrons alors que le paramètre ϵ critique mathématiquement peut être contrebalancé numériquement (notamment de par le fait que l'infiniment petit n'est pas atteignable numériquement) par d'autres paramètres semblant pourtant de moindre importance mathématiquement.

Notons que les figures de ce mémoire (sauf celles pour lesquelles il en est fait explicitement mention) ont été réalisées via les packages TikZ de LaTeX.

Chapitre 2

Motivations

Nous commencerons ce chapitre en revenant rapidement sur une notion qui fait pont entre la géométrie et la topologie algébrique, la caractéristique d'Euler-Poincaré. Nous donnerons ensuite deux domaines dans lesquels l'étude de fonctionnelles de Ginzburg-Laudau peut être utile.

2.1 Caractéristique d'Euler-Poincaré

La caractéristique d'Euler-Poincaré est un invariant topologique qui permet de classer des surfaces. Euler avait remarqué que les polyèdres convexes respectaient l'égalité suivante : $n - n_e + n_f = 2$ où n est le nombre de sommets, n_e le nombre d'arêtes et n_f le nombre de faces (voir [13]). Lorsque cette égalité est utilisée sur des découpes de surfaces, il a été remarqué que $n - n_e + n_f$ n'était pas toujours égale à 2. Cette somme de termes est alors devenu un élément permettant de caractériser des surfaces qui s'est avéré avoir un nombre inattendu de propriétés intéressantes. Cette invariant topologique est alors noté $\chi = n - n_e + n_f$ (voir [22]). Des découpages très basiques de la sphère et du tore nous indiquent que la caractéristique d'Euler-Poincaré du tore vaut 0 tandis que celle de la sphère vaut 2. En effet, nous pouvons, avec 2 sommets antipodaux que nous rejoignons via 4 arêtes, former 4 faces sur une sphère. De façon analogue, n'importe quelle découpe basique du tore nous donne une caractéristique d'Euler-Poincaré de 0. Cette caractéristique étant un invariant topologique, elle est indépendante de la découpe utilisée.

2.2 Applications en maillages

Comme expliqué dans un article de P.-A. Beaufort (voir [5]), construire des maillages quadrangulaires comporte de nombreux avantages. Par exemple, les éléments quadrangulaires de maillage supportent les opérations tensorielles. Pour construire ses maillages, l'approche indirecte consiste à placer des points sur la surface à mailler, les triangulariser et finalement combiner les triangles en quadrangles. La qualité finale du maillage dépend alors fortement de l'adéquation de la position des points engendrés avec une grille quadrangulaire potentielle. Une stratégie pour placer les points repose sur un champ de croix précalculé. Le champ de croix idéal doit être lisse et cohérent sur le plan topologique, et aligné avec les limites de la surface.

Pour ce faire, la théorie de Ginzburg-Landau peut-être utilisée. Dans son article, P.-A. Beaufort utilise une fonctionnelle de Ginzburg-Landau, légèrement différente de (1.1), qui consiste en un terme de régularité qui minimise le gradient - et non le gradient covariant - du champ de croix et un terme de pénalité qui assure que sa norme reste proche de l'unité.

Un résultat fondamental est indiqué dans l'article : seules les surfaces de caractéristique d'Euler-Poincaré nulle peuvent être pavées avec un maillage régulier. Si $\chi \neq 0$, des sommets irréguliers seront nécessairement présents dans le maillage. Comme indiqué plus haut la caractéristique

d'Euler-Poincaré est un invariant topologique qui se calcule comme suit :

$$\chi = n - n_e + n_f,$$

où n est le nombre de sommets, n_e le nombre d'arêtes et n_f le nombre de faces. L'impossibilité de recouvrir une surface de quadrangles, si $\chi \neq 0$ se comprend aisément. En effet, un ensemble de quadrangles est tel que le nombre de nœuds est égale au nombre de faces puisque chaque nœud est entouré de 4 faces et que chaque face est entourée de 4 nœuds. De plus, chaque nœud et chaque face sont entourés de 4 arêtes, mais chaque arête est entourée de 2 nœuds et de 2 faces. Dès lors, nous avons 2 fois plus d'arêtes que de nœuds et que de faces. Ainsi, ce raisonnement intuitif nous indique que si nous avons construit un maillage avec que des quadrangles sur une surface fermée, nous avons forcément

$$n - n_e + n_f = 2n - 2n = 0 = \chi.$$

L'article de P.-A. Beaufort met en lien le nombre, l'ordre et la position de nœuds irréguliers dans des maillages quadrangulaires avec le nombre, l'ordre et la position de singularités de la fonctionnelle de Ginzburg-Landau utilisée dans son article lorsque $\epsilon \rightarrow 0$.

Toutefois, de nombreuses questions restent en suspens. Un des éléments problématiques est notamment le fait que numériquement, faire $\epsilon \rightarrow 0$ est impossible. Dès lors, en fonction des limites numériques des machines et des résultats recherchés, les nœuds irréguliers seront placés en lieu et place de points singuliers de la fonctionnelle de Ginzburg-Landau utilisée qui ne correspondent pas toujours aux configurations minimisantes du problème asymptotique (lorsque $\epsilon \rightarrow 0$) en fonction de la surface étudiée. Toutefois, ces maillages gardent d'excellentes propriétés numériques et il n'est pas clair qu'il soit moins performant que des maillages qui seraient obtenus via des minimiseurs du problème asymptotique. Par exemple, dans l'image ci-dessous tirée de l'article de P.-A. Beaufort, la figure du milieu peut-être mise en correspondance avec un minimiseur asymptotique tandis que celle de gauche correspondrait à un minimiseur pour une valeur de ϵ plus éloignée de 0. Juger lequel des deux est le meilleur d'un point de vue numérique n'est certainement pas évident.

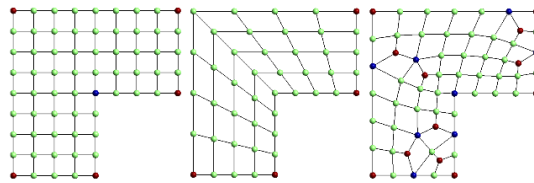


FIGURE 2.1 – Différentes quadrangulations d'un domaine en forme de L. Les sommets irréguliers d'indice $1/4$ (ils ont un voisin sur 4 en moins que prévu (il en faudrait idéalement 4 pour des nœuds centraux et 3 pour des nœuds appartenant au bord)) sont affichés en rouge, alors que ceux d'indice $-1/4$ (ils ont un voisin sur 4 de plus que prévu) sont affichés en bleu. La somme des indices des sommets irréguliers est égale à $\chi = 1$ dans tous les cas. Cette image est issue d'un article de P.-A. Beaufort qui parle d'une fonctionnelle de Ginzburg-Landau et de ses liens avec des techniques indirectes de maillage quadrangulaire (voir [5]).

Toutefois, nous ne nous aventurerons pas plus loin sur la comparaison entre les qualités des maillages issus du problème asymptotique et celles des maillages issus de minimiseurs non asymptotiques. Cela dépasse largement l'objet de ce travail de fin d'études.

Dans ce contexte, il est espéré que l'étude de l'équation (1.1) permettra d'aider au développement de ce domaine de recherche.

2.3 Applications en cristallographie et virologie

Cette section se base sur un article du cristallographe William F.Harris (voir [17]).

Bien que les cristaux soient souvent considérés comme des alignements parfaitement réguliers d'atomes ou de molécules, ils contiennent en réalité souvent des défauts. Ces imperfections sont responsables de beaucoup des propriétés physiques et chimiques des cristaux. Toutefois, elles ne peuvent prendre que quelques formes.

Une de ses imperfections qui peut facilement être reliée avec la première partie de la fonctionnelle énergétique de Ginzburg-Landau consiste en une rotation de la structure du cristal. Cela s'appelle une disclination. Les disclinations apparaissent principalement dans les structures orientées des cristaux liquides. Qui plus est, elles sont des éléments structurels importants dans de nombreux autres matériaux ordonnés que les cristaux classiques, tels que les enveloppes protéiques des virus. Les disclinations peuvent même être observées dans le modèle d'empreintes digitales et dans les peaux d'animaux rayés comme les zèbres.

Un élément notoire des disclinations est que les induire artificiellement sur des éléments sphériques nécessiterait une contrainte infinie dans une sphère et c'est donc impossible. En revanche, elles peuvent être artificiellement induites sur des tores.

Les disclinations peuvent amener des éléments de la structure à voir leur nombre de voisins changer. Cela peut être mis en lien avec le deuxième terme de la fonctionnelle de Ginzburg-Landau (voir [5]).

Notons aussi que des rotations de structures quadrangulaires coïncident avec la structure de base pour toute rotation multiple de $\frac{\pi}{2}$ ce qui correspond à l'idée intuitive derrière la définition d'une croix d'un champ de croix. En effet, un réseau cubique vient coïncider avec lui-même après une rotation de 90 degrés ; puisque quatre de ces points d'équivalence sont rencontrés dans une révolution complète, le treillis a une symétrie de rotation quadruple. Nous retrouvons aussi l'idée intuitive que si un champ de croix sur la surface d'un cube troué est déformé continument de manière à ce que la somme des déformations soit de $\frac{\pi}{2}$, en fonction du sens de rotation, le cube troué va perdre un coin ou gagner un coin dans la direction de cette rotation.

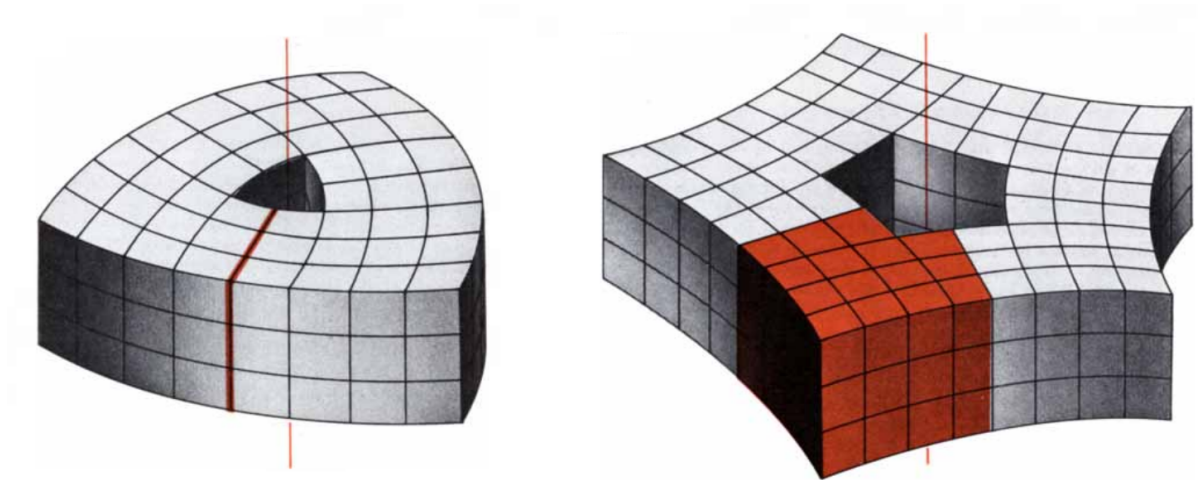


FIGURE 2.2 – Un exemple de disclination par rotation de plus 90° à gauche de l'image et à droite un exemple de disclination par rotation de moins 90° (cette image est issue de l'article de Harris [17]).

Notons d'ailleurs que dans ces structures cubiques, des rotations de n'importe quel angle autre que 90 degrés (ou des multiples de 90 degrés) entraîneraient une discontinuité : les treillis des côtés opposés de la surface coupée ne correspondraient pas. Parce que les rotations de cette magnitude introduisent de grandes contraintes, ces défauts ne peuvent pas se former dans les cristaux ordinaires.

Via la caractéristique d'Euler des polyèdres, nous pouvons montrer qu'une sphère doit inclure des disclinations avec une rotation totale de 720 degrés (voir [13]). Les 720 degrés de disclination peuvent être distribués de plusieurs façons. Une distribution particulièrement répandue est dans 12 disclinations de 60 degrés chacune. Par exemple, il n'est pas possible de carreler une sphère avec des hexagones mais le pavage peut être accompli si 12 pentagones sont inclus parmi les hexagones. Si les pentagones sont considérés comme des hexagones défectueux, alors chacun représente une disclination de 60 degrés. Dans les enveloppes protéiques de certains virus sphériques, 12 disclinations de 60 degrés chacune peuvent être identifiées.

Notons un dernier élément. Les singularités aux points et aux lignes associées aux imperfections du réseau ne sont pas nécessairement fixées dans le réseau, les disclinations sont connues depuis longtemps pour être mobiles.

Nous voyons donc que l'étude de la fonctionnelle de Ginzburg-Landau et des champs de croix trouve aussi des échos en cristallographie et en virologie.

Chapitre 3

Théorème de la boule chevelue

3.1 Motivation du théorème

Nous cherchons à savoir s'il existe des champs de vecteurs tangents en tout point de certaines surfaces et qui respectent certaines propriétés. Il existe déjà de nombreux résultats concernant l'existence de vecteurs tangents dont beaucoup sont liés à la caractéristique d'Euler-Poincaré (voir [13]).

Comme expliqué dans la partie motivation, nous cherchons via la fonctionnelle de Ginzburg-Landau et des champs de croix à créer des maillages qui respectent la forme à mailler. Comme première étape, nous allons montrer qu'il n'existe pas de champ de vecteurs continu, non nul et tangent en tout point à la sphère S^2 . Son existence aurait été une première étape pour montrer qu'il est possible d'épouser les contours de la sphère avec des éléments lui étant tangents en tout point. Pour démontrer cette impossibilité tout en se rapprochant au maximum d'un travail de maillage, nous allons travailler avec une démonstration se basant sur une découpe fine de la sphère en polygones curvilignes.

3.2 Variétés différentielles et sphère de Riemann

Les descriptions ci-dessous sont issues d'un livre de M. Berger et R. Gostiaux (voir [6]) et d'un syllabus d'un cours de mathématique de l'UCL écrit par le professeur Luc Haine (voir [16]).

Definition 3.2.1. Soit M un ensemble.

1. Une carte de M est un couple (U, ϕ) avec $U \subset M$ et ϕ une bijection de U sur un ouvert de $\mathbb{R}^n$. U s'appelle le domaine de la carte et n la dimension de la carte. Pour $p \in U$, nous appelons $\phi(p) = (x_1, \dots, x_n)$, les coordonnées locales de p dans la carte (U, ϕ) .
2. Soit $k \geq 1$ un entier. Deux cartes (U, ϕ) et (V, ψ) de dimensions respectives n et m , sont C^k compatibles si $U \cap V = \emptyset$ ou si $U \cap V \neq \emptyset$ et
 - $\phi(U \cap V)$ est un ouvert de $\mathbb{R}^n$,
 - $\psi(U \cap V)$ est un ouvert de $\mathbb{R}^m$,
 - $\psi \circ \phi^{-1} : \phi(U \cap V) \rightarrow \psi(U \cap V)$ est un difféomorphisme de classe C^k .

Dès que $U \cap V \neq \emptyset$, la dernière condition implique que les deux cartes sont nécessairement de même dimension $n = m$.

Remarque 3.2.1. Si $k = 0$, les changements de cartes sont seulement des homéomorphismes. Dans ce cas, le fait que deux cartes qui s'intersectent ont nécessairement la même dimension est un résultat non trivial de topologie.

Definition 3.2.2. Un ensemble $\mathcal{A} = \{(U_i, \phi_i)\}_{i \in I}$ de cartes C^k compatibles sur M , tel que $\cup_{i \in I} U_i = M$ s'appelle un atlas de classe C^k .

Proposition 3.2.1. *Tout atlas $\mathcal{A}$ (de classe C^k) de M est contenu dans un unique atlas maximal (de classe C^k) pour l'inclusion.*

Un atlas de classe C^k maximal pour l'inclusion est dit saturé.

Definition 3.2.3. *Nous appelons variété de classe C^k la donnée d'un couple $(M, \mathcal{A})$, où M est un ensemble et $\mathcal{A}$ est un atlas saturé de M , de classe C^k .*

Terminologie. Quand $k = 0$, nous parlons de variété topologique, quand $k = 1, 2, \dots$ de variété différentielle de classe C^k , quand $k = \infty$, nous parlons de variété de classe C^∞ ou de variété lisse. La dimension d'un point p de la variété est la dimension de toute carte locale en p . Si la dimension est la même en tout point, nous l'appelons la dimension de la variété.

Soit $S^2 = \{(u_1, u_2, u_3) \in \mathbb{R}^3 : \sum_{i=1}^3 u_i^2 = 1\}$, la sphère unité de $\mathbb{R}^3$. Nous définissons un atlas $\mathcal{A} = \{(U_1, \phi_1), (U_2, \phi_2)\}$ sur S^2 à partir des projections stéréographiques

$$\begin{aligned} \phi_1 : U_1 = S^2 \setminus \{(0, 0, 1)\} &\rightarrow \mathbb{R}^2 \\ \text{tel que } \phi_1(u_1, u_2, u_3) &= \left(\frac{u_1}{1-u_3}, \frac{u_2}{1-u_3}\right) = (x_1, x_2) \\ \text{et } \phi_2 : U_2 = S^2 \setminus \{(0, 0, -1)\} &\rightarrow \mathbb{R}^2 \\ \text{tel que } \phi_2(u_1, u_2, u_3) &= \left(\frac{u_1}{1+u_3}, \frac{u_2}{1+u_3}\right) = (y_1, y_2). \end{aligned}$$

Puisque

$$y_i = x_i \frac{1-u_2}{1+u_2} \text{ et } x_1^2 + x_2^2 = \frac{u_1^2 + u_2^2}{(1-u_3)^2} = \frac{1-u_3^2}{(1-u_3)^2} = \frac{1+u_3}{1-u_3},$$

nous obtenons le changement de cartes

$$\begin{aligned} \phi_2 \circ \phi_1^{-1} : \phi_1(U_1 \cap U_2) &= \mathbb{R}^n \setminus \{(0, 0, 0)\} \rightarrow \phi_2(U_1 \cap U_2) = \mathbb{R}^n \setminus \{(0, 0, 0)\} \\ (x_1, x_2) &\rightarrow \left(\frac{x_1}{x_1^2 + x_2^2}, \frac{x_2}{x_1^2 + x_2^2}\right), \end{aligned}$$

qui est un difféomorphisme de classe C^∞ ; l'inverse est donnée par

$$(y_1, y_2) \rightarrow \left(\frac{y_1}{y_1^2 + y_2^2}, \frac{y_2}{y_1^2 + y_2^2}\right),$$

Si nous pensons à un point $(x_1, x_2) \in \mathbb{R}^2$ comme au nombre complexe $z = x_1 + ix_2$, nous voyons que le changement de carte s'écrit

$$\bar{\phi}_2 \circ \phi_1^{-1} : \phi_1(U_1 \cap U_2) = \mathbb{C} \setminus \{0\} \rightarrow \bar{\phi}_2(U_1 \cap U_2) = \mathbb{C} \setminus \{0\} : z \rightarrow \frac{1}{z}.$$

Cet atlas définit une structure de variété complexe de dimension 1 sur S^2 car les cartes sont modelées sur des ouverts de $\mathbb{C}$ et les changements de cartes sont holomorphes. Avec cette structure, S^2 s'appelle la *sphère de Riemann*.

3.3 Relèvement et revêtement

Cette section se base sur le point 3 du chapitre 10 du livre "analyse mathématique 3" de Roger Godement (voir [15]).

Soient X et B deux espaces séparés. Soit une application continue et surjective $p : X \rightarrow B$. Nous supposons B connexe par arcs et localement connexe. Le système (X, B, p) est un revêtement s'il est localement trivial ce qui veut dire qu'il doit respecter la condition suivante.

Condition 3.3.1. *Tout $z \in B$ possède un voisinage ouvert D dont l'image réciproque $p^{-1}(D)$ est la réunion d'une famille $(D_i)_{i \in I}$ d'ouverts deux à deux disjoints que p applique homéomorphiquement sur D .*

Puisque p s'applique de X dans B de façon continue, nous pouvons construire un ouvert $] \zeta - \epsilon; \zeta + \epsilon[\subset X$ arbitraire tel que son image par p soit $]p(\zeta) - p(\epsilon); p(\zeta) + p(\epsilon)[\subset B$ et donc, p transforme tout voisinage d'un $\zeta \in X$ en un voisinage de $p(\zeta) \in B$ et p transforme tout ouvert de X en un ouvert de B . De plus, puisque p s'applique depuis une famille $(D_i)_{i \in I} \subset X$ d'ouverts disjoints, de manière bijective pour chacun des D_i , pour tout $z \in B$, la fibre $p^{-1}(\{z\})$ de z dans X est une partie discrète de X puisque chacune de ses intersections avec les ouverts D_i se réduit à un seul point.

Nous abordons maintenant la notion de relèvement d'un chemin. En plus d'être un concept fondamental dans la théorie des revêtements d'espaces topologiques, un résultat concernant les relèvements (le théorème 3.3.1) sera exploité par la suite.

Dans toute cette section, nous noterons $I = [0, 1]$. Commençons par donner la définition de relèvement ainsi que deux exemples.

Définition 3.3.1. *Soit (X, B, p) un revêtement et soit $\gamma : I \rightarrow B$ un chemin continu dans B . Un relèvement de γ est un chemin $\mu : I \rightarrow X$ tel que $p \circ \mu = \gamma$:*

$$\begin{array}{ccc} & X & \\ & \downarrow p & \\ I & \xrightarrow{\gamma} & B \end{array} \quad \begin{array}{c} \nearrow \mu \\ \end{array}$$

Exemple 3.3.1. Considérons le revêtement (X, B, p) où $X = \mathbb{C}$, $B = \mathbb{C}^*$ et $p : X \rightarrow B; w \mapsto \exp(2\pi i w)$. Prenons le chemin $\gamma : I \rightarrow B; t \mapsto \exp(\pi i t)$. Alors, un relèvement de γ est donné par $\mu : I \rightarrow X; \mu(t) = \frac{t}{2}$.

Exemple 3.3.2. Pour le revêtement (X, B, p) avec $X = \mathbb{C}^* = B$ et $p : X \rightarrow B; z \mapsto z^2$, un relèvement de $\gamma : I \rightarrow B; t \mapsto \exp(2\pi i t)$ est donné par $\mu : I \rightarrow X; t \mapsto \exp(\pi i t)$.

Il est à ce stade tout à fait légitime de se demander si, étant donné un chemin dans B , un relèvement existe toujours et, si c'est le cas, s'il y a unicité d'un tel relèvement. Nous pouvons également nous demander si les relèvements de deux chemins homotopes dans B sont à leur tour homotopes dans X . C'est à ces deux questions que répond le théorème suivant :

Théorème 3.3.1. *Soit (X, B, p) un revêtement.*

Soit $\gamma : I \rightarrow B$ un chemin dans B .

Alors, il existe un unique relèvement $\mu : I \rightarrow X$ de γ à X ayant une origine donnée.

De plus, si deux chemins γ_0 et γ_1 dans B sont homotopes avec leurs extrémités fixes et si μ_0 et μ_1 sont des relèvements de γ_0 et γ_1 (respectivement) ayant la même origine, alors μ_0 et μ_1 sont homotopes et ont la même extrémité.

Rappelons que deux chemins $\gamma_0, \gamma_1 : I \rightarrow B$ de mêmes extrémités sont homotopes (à extrémités fixes) s'il existe une application continue

$$\sigma : I \times I \rightarrow B; (s, t) \mapsto \sigma(s, t)$$

telle que

$$\begin{cases} \sigma(s, 0) = \gamma_0(0) & \text{pour tout } s \in I \\ \sigma(s, 1) = \gamma_1(1) & \text{pour tout } s \in I \\ \sigma(0, t) = \gamma_0(t) & \text{pour tout } t \in I \\ \sigma(1, t) = \gamma_1(t) & \text{pour tout } t \in I. \end{cases}$$

3.4 Démonstration

Passons maintenant à une démonstration originale du théorème dit "de la boule chevelue". L'approche de cette démonstration est la découpe de la sphère S^2 en polygones curvilignes sur sa surface. Le lien avec les applications de maillage consiste en cette découpe de la sphère.

Théorème 3.4.1 (Dit "de la boule chevelue"). *Il n'existe pas de champ de vecteurs continu sur la sphère S^2 tel que ce champ ait une norme non nulle et soit tangent à S^2 en tout point de S^2 .*

Démonstration. Cette démonstration se fera par l'absurde via une découpe fine de la sphère en polygones curvilignes. Nous utiliserons alors la variation de la différence entre l'hypothétique champ de vecteurs tangents et les arêtes des polygones curvilignes ainsi que la caractéristique d'Euler-Poincaré de la sphère pour arriver à une contradiction. Soit $n : S^2 \rightarrow \mathbb{R}^3$ tel que pour tout $x \in S^2$, $n(x)$ soit la normale unitaire à S^2 en x .

Admettons l'existence de w un champ de vecteurs continu et tel que pour tout $x \in S^2$, $n(x) \cdot w(x) = 0$ et que $\forall x \in S^2, |w(x)| > 0$. Posons $u(x) = \frac{w(x)}{|w(x)|}$.

Définissons N , un ensemble de points appartenant à S^2 défini via les intersections de courbes curvilignes quelconques se trouvant sur S^2 . Nous appellerons ces points des noeuds comme cela se fait en théorie des graphes. Soit A l'ensemble des arêtes reliant les différents noeuds via les courbes curvilignes introduites ci-dessus, de façon à ce que les arêtes $\gamma_{ij} \in A$ joignant les différents noeuds via les courbes curvilignes soient telles que que les différentes arêtes ne s'intersectent jamais les unes les autres. Définissons les arêtes $\gamma_{ij} : [0, L_{ij}] \rightarrow S^2$ où $i, j \in N$ et L_{ij} est la longueur de l'arête. Admettons qu'une découpe de S^2 en un ensemble de polygones P curvilignes ait ainsi été définie sur S^2 et que cette découpe soit très fine.

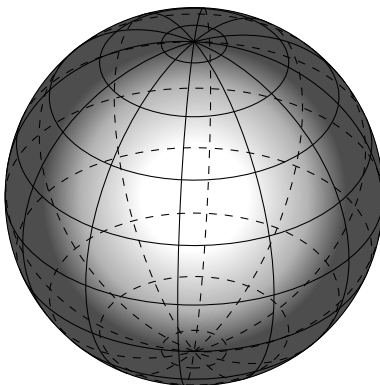


FIGURE 3.1 – Exemple de découpage grossier en polygones curvilignes (triangles et quadrangles) de la sphère S^2 .

Imposons que les arêtes $\gamma_{ij} : [0, L_{ij}] \rightarrow S^2 \in P$ soient parcourues à vitesse unitaire constante, $\gamma'_{ij}(t) = 1$. Nous avons alors $\gamma_{ij}(t) = \gamma_{ji}(L_{ij} - t)$.

Tentons maintenant de déterminer de façon unique et continue, d'arête en arête, la variation de l'angle entre le champ de vecteur $u(\gamma_{ij}(t))$ et $\gamma'_{ij}(t)$ où $t \in [0, L_{ij}]$ et ce pour chaque arête faisant partie de l'ensemble P . Appelons cet angle $\alpha_{ij}(t)$.

Nous pouvons calculer

$$\cos(\alpha_{ij}(t)) = \frac{u(\gamma_{ij}(t)) \cdot \gamma'_{ij}(t)}{\|u(\gamma_{ij}(t))\| \|\gamma'_{ij}(t)\|},$$

et

$$\sin(\alpha_{ij}(t)) = \frac{u(\gamma_{ij}(t)) \times \gamma'_{ij}(t)}{\|u(\gamma_{ij}(t))\| \|\gamma'_{ij}(t)\|},$$

où $\cdot$ désigne le produit scalaire et $\times$ dénote le produit vectoriel de deux vecteurs de $\mathbb{R}^3$. Nous pouvons constater que sur une même arête, $\cos(\alpha_{ij}(t))$ et $\sin(\alpha_{ij}(t))$ varient continûment puisque le champ de vecteurs u varie continûment.

Nous cherchons maintenant, sur base de ces sinus et cosinus, à obtenir des valeurs variant continûment avec t pour $\alpha_{ij}(t)$.

Pour le définir de façon rigoureuse, nous avons besoin des notions de revêtement et de relèvement ainsi que de résultats les concernant. Introduisons ce dont nous allons avoir besoin (voir [18]). Notons que les notions introduites le sont en toutes généralités alors que R. Godement les introduit dans le cadre d'un espace B supposé connexe par arcs et localement connexe (voir [15]). Soit un espace séparé X , un revêtement de X consiste en un espace séparé $\tilde{X}$ et une application $p : \tilde{X} \rightarrow X$ satisfaisant la condition suivante :

Condition 3.4.1. *Pour chaque $x \in X$, il y a un voisinage U de x dans X tel que $p^{-1}(U)$ est une union d'ouverts disjoints qui sont chacun appliqués homéomorphiquement dans U par p .*

Si $(X, \tilde{X}, p)$ est un revêtement, pour chaque chemin $f : I \rightarrow X$, nous appelons relèvement de f tout chemin $\tilde{f}$ de $\tilde{X}$ tel que $p \circ \tilde{f} = f$. C'est représenté sur le diagramme 3.2.

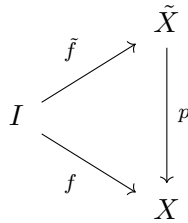


FIGURE 3.2 – Diagramme représentant le lien entre un chemin f et un de ses relevés $\tilde{f}$.

Notons aussi la propriété suivante pour un revêtement $p : \tilde{X} \rightarrow X$.

Propriété 3.4.1 (Unicité du relèvement). Pour chaque chemin $f : I \rightarrow X$ commençant en un point $x_0 \in X$ et pour chaque $\tilde{x}_0 \in p^{-1}(x_0)$, il existe un unique relèvement $\tilde{f} : I \rightarrow \tilde{X}$ commençant en $\tilde{x}_0$.

Utilisons maintenant ces notions sur notre problème. Nous pouvons maintenant définir un relèvement de S^1 par $\mathbb{R}$ via le revêtement $p : \mathbb{R} \rightarrow S^1 : a \rightarrow (\cos a, \sin a)$ (voir 3.3 pour une figure représentant ce relèvement).

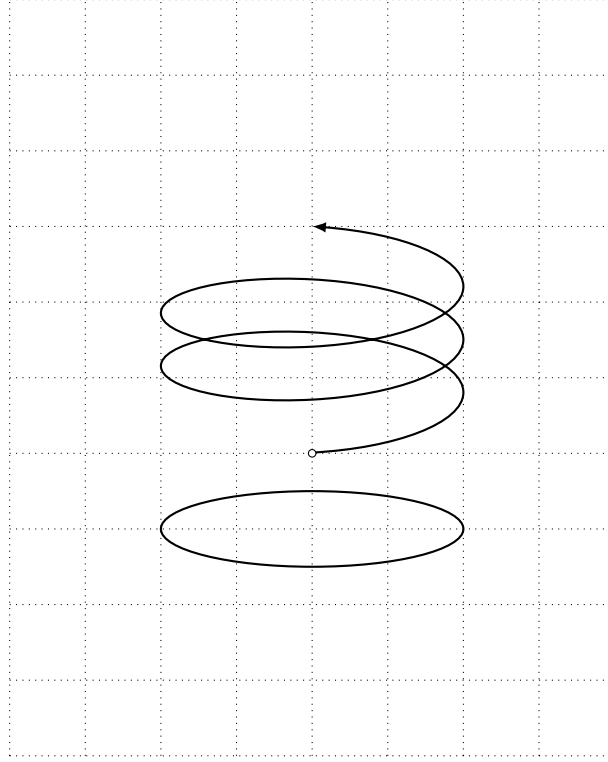


FIGURE 3.3 – Revêtement de S^1 par $\mathbb{R}$.

Notons que $\mathbb{R}$ et S^1 sont des espaces séparés simplement connexes. En parcourant n'importe quelle arête, nous pouvons donc définir en se donnant un angle arbitraire $\alpha_{ij}(0)$ tel que

$$(\cos \alpha_{ij}(0), \sin \alpha_{ij}(0)) = \left(\frac{u(\gamma_{ij}(0)) \cdot \gamma'_{ij}(0)}{\|u(\gamma_{ij}(0))\| \|\gamma'_{ij}(0)\|}, \frac{u(\gamma_{ij}(0)) \times \gamma'_{ij}(0)}{\|u(\gamma_{ij}(0))\| \|\gamma'_{ij}(0)\|} \right)$$

un unique angle $\alpha_{ij}(t)$ qui varie de façon continue sur l'arête en cours de parcours.

Définissons maintenant pour chaque polygone P appartenant à l'ensemble des polygones formés sur S^2 ($P \in S^2$) par l'ensemble des nœuds N et l'ensemble des arêtes A , une application qui lui fait correspondre un polynôme plan de l'ensemble des polygones plans $P_{\text{plan}} \in \mathbb{R}^2$. Nous avons donc

$$\psi : P_{\text{plan}} \rightarrow S^2$$

avec les propriétés suivantes. Pour tout polygone de S^2 , il existe un polygone de $\mathbb{R}^2$ tel que ψ envoie les arêtes, les sommets et l'intérieur des polygones de P_{plan} sur les arêtes, les sommets et l'intérieur des polygones P correspondants. L'existence d'une telle application est garantie par le fait que S^2 est une variété de dimension 2 pour laquelle nous pouvons trouver des cartes C^∞ de $\mathbb{R}^2$ constituant un atlas de classe C^∞ de S^2 . En prenant comme cartes locales celles issues de la projection stéréographique, nous pouvons définir ϕ un homéomorphisme entre S^2 et $\mathbb{R}^2$ avec un point rajouté à l'infini (voir la figure 3.4 pour un exemple de projection stéréographique).

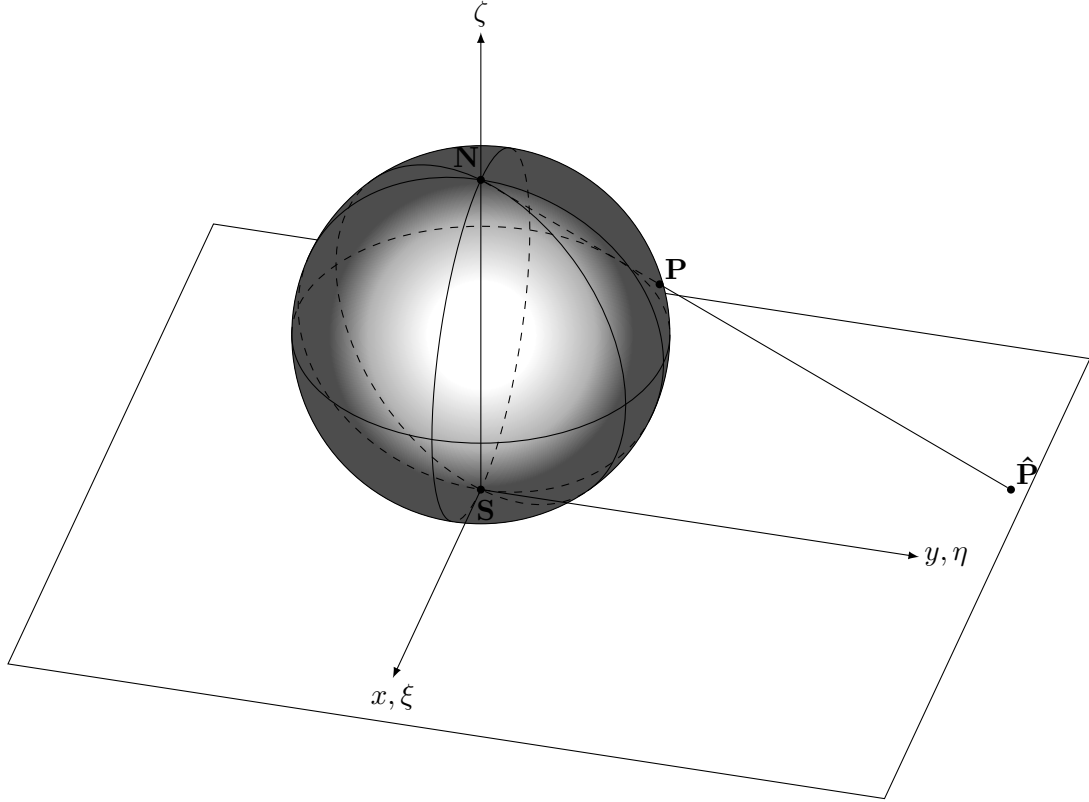


FIGURE 3.4 – Exemple de projection stéréographique.

Définissons la variation de l'angle sur une arête par

$$\delta_{ij} \triangleq \alpha_{ij}(L_{ij}) - \alpha_{ij}(0)$$

la variation de l'angle ne dépend pas de l'angle de départ choisi puisque $\alpha_{ij}(t)$ évolue continûment.

A chaque parcours d'un polygone P_l de P_{plan} , sur chaque arête la variation de l'angle va évoluer continûment (la variation ayant augmenté de 2π à chaque fois qu'il repasse sur des mêmes valeurs de sinus et cosinus) et à chaque sommet du polygone, la variation de l'angle va être augmentée de $\pi - \epsilon_{mn}$ (ϵ_{mn} étant l'angle au dessus de l'arête γ_{mn} lorsque nous arrivons à la fin de son parcours).

Nous observons que sur chaque polygone,

$$\sum_{(i,j) \in P_l} (\pi - \epsilon_{ij}) + \sum (\delta_{ij}) = k2\pi$$

avec $k \in \mathbb{Z}$ puisqu'après avoir parcouru tout le tour d'un polygone, nous retombons sur le point initial avec les mêmes valeurs $u(\gamma_{ij}(0))$ et $\gamma'_{ij}(0)$. Nous supposons notre découpage suffisamment fin pour que pour tout parcours de polygones, $k = 1$ par continuité du champ u .

Nous observons maintenant que pour chaque polygone P_i de l'ensemble P_{plan} parcouru dans le sens anti-horloger, pour des arêtes successives γ_{mn} , nous avons

$$\sum_{(m,n) \in A \cap P_i} (\pi - \epsilon_{mn}) + \sum_{(m,n) \in A \cap P_i} \delta_{mn} = 2\pi.$$

Nous avons donc sur l'ensemble des polygones

$$\sum_{(m,n) \in P_{\text{plan}}} (\pi - \epsilon_{mn}) + \sum_{(m,n) \in P_{\text{plan}}} \delta_{mn} = 2\pi|P|.$$

Observons de plus sur cet ensemble que

$$\sum_{(i,j) \in A} \delta_{ij} = 0$$

car $\delta_{ij} = -\delta_{ji}$ et que

$$\sum_{(i,j) \in P_{\text{plan}}} (\pi - \epsilon_{ij}) = \pi(2|A|) - 2\pi|S| = 2\pi|P|$$

car il y a pour chaque arête et dans chaque sens de parcours de celle-ci un saut de π moins un angle touchant un sommet or la somme des angles touchant un sommet vaut 2π .

Nous avons donc $2\pi|A| - 2\pi|S| = 2\pi|P|$ or c'est impossible puisque la caractéristique d'Euler-Poincaré de la sphère vaut $|P| - |A| + |S| = 2$. Il n'existe donc pas de champ de vecteurs continu sur la sphère S^2 tel que ce champ ait une norme non nulle et soit tangent à S^2 en tout point de S^2 .

Remarque 3.4.1 (Exemple de calcul de la caractéristique d'Euler-Poincaré de la sphère). Pour s'en convaincre aisément, nous pouvons voir que la sphère de la figure 3.1 a été découpée en 12 tranches et que pour chaque tranche nous avons 9 faces. Nous avons de plus, 12 sommets par cercle horizontal et 8 cercles horizontaux ainsi que 2 sommets antipodaux. Nous avons encore 8 arêtes horizontales et 9 verticales par tranche. Ainsi, nous avons :

$$\# \text{faces} - \# \text{arêtes} + \# \text{sommets} = 12 * 9 - (9 * 12 + 8 * 12) + (12 * 8 + 2) = 2.$$

□

Chapitre 4

Existence de minimiseurs

4.1 Introduction

Avant de s'attaquer à l'existence de solutions d'énergie finie pour la fonctionnelle (1.1), il convient de rappeler quelques éléments d'analyse fonctionnelle et de géométrie différentielle qui permettront d'appréhender avec clarté la fonctionnelle (1.1). Ainsi, nous verrons que nous travaillerons avec un espace de Sobolev de fonctions, l'espace H^1 et avec une dérivée covariante. Ces notions méritent de s'y attarder un peu d'autant qu'elles font déjà intervenir quelques éléments intéressants d'analyse et de géométrie fonctionnelle.

4.2 Dérivée faible

Si une fonction $u : \mathbb{R} \rightarrow \mathbb{R}$ est dérivable, alors, via la formule d'intégration par parties, nous savons que pour toute fonction dérivable $\phi : \mathbb{R} \rightarrow \mathbb{R}$ telle que $\phi(x) = 0$ pour $|x|$ suffisamment grand, l'égalité suivante est respectée :

$$\int_{-\infty}^{\infty} u\phi' dx = u\phi|_{-\infty}^{+\infty} - \int_{-\infty}^{+\infty} u'\phi dx = - \int_{-\infty}^{+\infty} u'\phi dx.$$

L'idée de la différentiabilité faible est que nous demandons à la formule d'intégration par partie d'être respectée pour toute fonction dérivable $\phi : \mathbb{R} \rightarrow \mathbb{R}$ telle que $\phi(x) = 0$ pour $|x|$ suffisamment grand. Cela n'implique pas forcément que la dérivabilité classique soit possible.

Les fonctions qui ne sont pas dérivables peuvent toujours être faiblement dérivables. Par exemple,

$$u(x) = |x|$$

n'est pas dérivable à cause de l'angle formé en $x = 0$, mais elle possède bien une dérivée faible.

$$u'(x) = \begin{cases} -1, & x < 0, \\ 0, & x = 0, \\ 1, & x > 0. \end{cases}$$

Cette fonction faiblement dérivables a l'air dérivable sauf pour un ensemble de mesure nulle. Toutefois, il existe des fonctions qui ne sont pas dérivable qu'en un ensemble de mesure nulle et qui ont une dérivée nulle partout en dehors de cette ensemble et qui pourtant n'ont pas une dérivée faible nulle. Un exemple de ceci est que la dérivée faible d'une fonction échelon est un delta de Dirac. Les fonctions faiblement dérivables qui sont construites via des intégrales de Lebesgue (intégrales ne tenant pas compte d'ensemble de mesure nulle) peuvent aussi être construites via la théorie des distributions. Dès lors, nous parlons parfois de dérivée faible en termes de dérivée-distribution ou de dérivée au sens des distributions (notons que la notion de dérivée au sens des distribution est plus large que celle de dérivée faible).

4.2.1 Cas multidimensionnel

Nous venons de voir la notion de dérivée faible de façon intuitive et non rigoureuse dans le cadre de fonctions à une seule dimension. Nous pourrions en tirer l'intuition que les fonctions faiblement dérivables sont continues, ce qui serait une hypothèse erronée. En effet, des fonctions faiblement dérivables peuvent avoir des comportements très particuliers. Par exemple, nous pouvons construire une fonction faiblement dérivable sur la boule unité $B(0, 1) \subset \mathbb{R}^n$ qui est non bornée sur chaque sous-ensemble ouvert de $B(0, 1)$.

Notons tout d'abord que nous indiquons par $C^\infty(\bar{\Omega})$ l'espace des restrictions à $\bar{\Omega}$ des fonctions C^∞ sur $\mathbb{R}^N$ et par $C_c^\infty(\Omega)$ l'espace des fonctions C^∞ à support compact sur U (celles-ci sont donc nulles aux voisinages de $+\infty$ et $-\infty$ comme nous le demandons aux fonctions multipliant une fonction et sa dérivée faible pour la définition des dérivées faibles par la formule d'intégration par parties).

Définissons et caractérisons de façon rigoureuse les dérivées faibles dans le cadre plus général d'espaces multidimensionnels. Les éléments qui suivent sont issus d'un ouvrage de Grégoire Allaire (voir [1]).

Définition 4.2.1 (Dérivation faible dans $L^2(\Omega)$ avec Ω un ouvert de $\mathbb{R}^N$). *Soit v une fonction de $L^2(\Omega)$. La fonction v est dite dérivable au sens faible dans $L^2(\Omega)$ s'il existe des fonctions $w_i \in L^2(\Omega)$, pour $i \in \{1, \dots, N\}$, telles que, pour toute fonction $\phi \in C_c^\infty(\Omega)$, nous avons*

$$\int_{\Omega} v(x) \frac{\partial \phi}{\partial x_i}(x) dx = - \int_{\Omega} w_i(x) \phi(x) dx.$$

Chaque w_i est appelée la i -ème dérivée partielle faible de v et notée $\frac{\partial v}{\partial x_i}$.

Bien sûr, si v est dérivable au sens usuel et que ses dérivées partielles appartiennent à $L^2(\Omega)$, alors les dérivées usuelle et faible de v coïncident. Donnons tout de suite un critère simple et pratique pour déterminer si une fonction est dérivable au sens faible.

Lemme 4.2.1. *Soit v une fonction de $L^2(\Omega)$. S'il existe une constante $C > 0$ telle que, pour toute fonction $\phi \in C_c^\infty(\Omega)$ et pour tout indice $i \in \{1, \dots, N\}$, nous avons*

$$\left| \int_{\Omega} v(x) \frac{\partial \phi}{\partial x_i}(x) dx \right| \leq C \|\phi\|_{L^2(\Omega)},$$

alors v est dérivable au sens faible.

Proposition 4.2.1. *Soit v une fonction de $L^2(\Omega)$ dérivable au sens faible et telle que toutes ses dérivées partielles faibles $\frac{\partial v}{\partial x_i}$, pour $1 \leq i \leq N$, sont nulles. Alors, pour chaque composante connexe de Ω , il existe une constante C telle que $v(x) = C$ presque partout dans cette composante connexe.*

La notion de dérivée faible se généralise à certains opérateurs différentiels qui ne font intervenir que certaines combinaisons de dérivées partielles (et non pas toutes) comme le gradient par exemple.

Remarque 4.2.1. La notion de dérivation faible et tous les résultats de cette sous-section s'étendent aux espaces $L^p(\Omega)$ pour $1 \leq p \leq +\infty$. Pour $p \neq 2$, le critère de dérivation faible 4.2.1 doit alors être remplacé par

$$\left| \int_{\Omega} v(x) \frac{\partial \phi}{\partial x_i}(x) dx \right| \leq C \|\phi\|_{L^{p'}(\Omega)} \text{ avec } \frac{1}{p} + \frac{1}{p'} = 1 \text{ et } 1 < p \leq +\infty.$$

4.3 Espace de Sobolev

Les éléments suivants sont tirés d'un livre d'analyse fonctionnelle de H.Brezis (voir [11]).

Soit $\Omega \subset \mathbb{R}^N$ un ouvert et soit $p \in \mathbb{R}$ avec $1 \leq p \leq \infty$.

Definition 4.3.1. *L'espace de Sobolev $W^{1,p}(\Omega)$ est défini par*

$$W^{1,p} = \left\{ u \in L^p(\Omega) \left| \begin{array}{l} \exists g_1, g_2, \dots, g_N \in L^p(\Omega) \text{ tels que} \\ \int_{\Omega} u \frac{\partial \phi}{\partial x_i} = - \int_{\Omega} g_i \phi \quad \forall \phi \in C_c^\infty(\Omega) \quad \forall i = 1, 2, \dots, N \end{array} \right. \right\}.$$

Nous posons

$$H^1(\Omega) = W^{1,2}(\Omega).$$

Pour $u \in W^{1,p}(\Omega)$, nous notons

$$\frac{\partial u}{\partial x_i} = g_i \text{ et } \nabla u = \left(\frac{\partial u}{\partial x_1}, \frac{\partial u}{\partial x_2}, \dots, \frac{\partial u}{\partial x_N} \right),$$

ces dérivées définies via des intégrales, des fonctions $C_c^\infty(\Omega)$ et la formule de dérivation par parties sont appelées des dérivées faibles ou dérivées-distributions comme nous l'avons vu précédemment. Dans la suite, c'est avec ce type de dérivées et de gradients que nous travaillerons.

L'espace $W^{1,p}(\Omega)$ est muni de la norme

$$\|u\|_{W^{1,p}} = \|u\|_{L^p} + \sum_{i=1}^N \left\| \frac{\partial u}{\partial x_i} \right\|_{L^p}$$

ou parfois de la norme équivalente

$$\left(\|u\|_{L^p}^p + \sum_{i=1}^N \left\| \frac{\partial u}{\partial x_i} \right\|_{L^p}^p \right)^{1/p} \quad \text{si } 1 \leq p < \infty.$$

L'espace $H^1(\Omega)$ est muni du produit scalaire

$$\langle u, v \rangle_{H^1} = \langle u, v \rangle_{L^2} + \sum_{i=1}^N \left\langle \frac{\partial u}{\partial x_i}, \frac{\partial v}{\partial x_i} \right\rangle_{L^2};$$

et de la norme associée

$$\|u\|_{H^1} = \left(\|u\|_{L^2}^2 + \sum_{i=1}^N \left\| \frac{\partial u}{\partial x_i} \right\|_{L^2}^2 \right)^{1/2}$$

qui est équivalente à la norme de $W^{1,2}$.

Proposition 4.3.1. *L'espace H^1 est un espace de Hilbert (espace vectoriel normé complet pour la norme issue de son produit scalaire) séparable.*

Remarque 4.3.1. Dans la définition de $W^{1,p}$, nous pouvons utiliser indifféremment $C_c^1(\Omega)$ ou $C_c^\infty(\Omega)$ comme ensemble de fonctions test.

Remarque 4.3.2. Il est clair que si $u \in C^1(\Omega) \cap L^p(\Omega)$ et si $\frac{\partial u}{\partial x_i} \in L^p(\Omega)$ pour tout $i = 1, 2, \dots, N$ (ici $\frac{\partial u}{\partial x_i}$ désigne la dérivée partielle au sens de $W^{1,p}$), alors $u \in W^{1,p}(\Omega)$; de plus les dérivées partielles au sens usuel coïncident avec les dérivées partielles au sens $W^{1,p}$. En particulier, si Ω est borné, alors $C^1(\bar{\Omega}) \subset W^{1,p}(\Omega)$ pour tout $1 \leq p \leq \infty$. Inversement, il est possible de démontrer que si $u \in W^{1,p}(\Omega) \cap C(\Omega)$ avec $1 \leq p \leq \infty$ et si $\frac{\partial u}{\partial x_i} \in C(\Omega)$ pour tout $i = 1, 2, \dots, N$ ($\frac{\partial u}{\partial x_i}$ désigne ici la dérivée partielle au sens de $W^{1,p}$), alors $u \in C^1(\Omega)$.

Remarque 4.3.3. Utilisant le langage des distributions, nous pouvons dire que $W^{1,p}$ est l'ensemble des fonctions $u \in L^p(\Omega)$ telles que toutes les dérivées partielles $\frac{\partial u}{\partial x_i}$, $1 \leq i \leq N$ (au sens des dérivées-distributions) appartiennent à $L^p(\Omega)$.

Remarque 4.3.4. Soit $(u_n)_{n \in \mathbb{N}}$ une suite de $W^{1,p}$ telle que $u_n \rightarrow u$ dans L^p et (∇u_n) converge vers une limite dans $(L^p)^N$, alors $u \in W^{1,p}$ et $\|u_n - u\|_{W^{1,p}} \rightarrow 0$. Lorsque $1 < p \leq \infty$, il suffit de savoir que $u_n \rightarrow u$ dans L^p et que (∇u_n) reste borné dans $(L^p)^N$ pour conclure que $u \in W^{1,p}$.

Notons aussi les éléments suivants.

Definition 4.3.2. Une espace E s'injecte continûment dans un espace F si $E \subset F$ et si la restriction à E d'une application continue définie sur F est encore continue. Il existe alors une constante C telle que pour tout $u \in E$, $\|u\|_F \leq C\|u\|_E$.

Definition 4.3.3. Une injection est dite compacte de E dans F si la boule unité de E , vue dans F , est relativement compacte, c'est-à-dire si $\overline{B_E(0,1)}$ est compact dans F .

Théorème 4.3.1. Soit Ω un ouvert de $\mathbb{R}^N$. Nous supposons que Ω est de classe C^1 avec Γ borné (ou bien $\Omega = \mathbb{R}_+^N$). Alors il existe un opérateur de prolongement

$$P : W^{1,p}(\Omega) \rightarrow W^{1,p}(\mathbb{R}^N)$$

linéaire, tel que pour tout $u \in W^{1,p}(\Omega)$

- (i) $Pu|_\Omega = u$
- (ii) $\|P\|_{L^p(\mathbb{R}^N)} \leq C\|u\|_{L^p(\Omega)}$
- (iii) $\|Pu\|_{W^{1,p}(\mathbb{R}^N)} \leq C\|u\|_{W^{1,p}(\Omega)}$

où C dépend seulement de Ω .

Nous pouvons utiliser les résultats obtenus sur $\mathbb{R}^N$ sur nos Ω puisque nos Ω peuvent être construits sur base d'un recouvrement fini d'ouverts de $\mathbb{R}^N$.

Théorème 4.3.2 (Sobolev, Gagliardo, Nirenberg). Soit $1 \leq p < N$, alors

$$W^{1,p}(\mathbb{R}^N) \subset L^{p^*}(\mathbb{R}^N) \text{ où } p^* \text{ est donné par } \frac{1}{p^*} = \frac{1}{p} - \frac{1}{N},$$

et il existe une constante C telle que

$$\|u\|_{L^{p^*}} \leq C\|\nabla u\|_{L^p} \quad \forall u \in W^{1,p} \subset W^{1,p}(\mathbb{R}^N).$$

Ce n'est pas tellement le théorème précédent mais bien le corollaire suivant que nous utiliserons souvent dans la suite de ce mémoire. Ce dernier caractérise le cas où $p = N$ ce qui nous concerne puis que nous travaillons avec des champs appartenant à $H^1(\Omega) = W^{1,2}(\Omega)$ de vecteurs à trois dimensions avec des surfaces Ω qui sont des variétés de dimension 2.

Corollaire 4.3.1 (Le cas limite $p = N$). Nous avons

$$W^{1,N}(\mathbb{R}^N) \subset L^q(\mathbb{R}^N) \quad \forall q \in [N, +\infty[$$

avec injection continue.

Dans notre cas nous aurons donc l'injection continue de $H^1(\Omega)$ dans $L^p(\Omega)$ et ce $\forall p \in [2, +\infty[$.

4.4 Espace H^1

L'espace de Sobolev $H^1(\Omega)$ est un espace de fonctions de carré intégrable et dont le gradient (au sens des dérivées faibles) est aussi de carré intégrable

$$H^1(\Omega) = \{u \in L^2(\Omega), \nabla u \in L^2(\Omega)\},$$

$$\|u\|_{H^1(\Omega)} = (\|u\|_{L^2(\Omega)}^2 + \|\nabla u\|_{L^2(\Omega)}^2)^{1/2}.$$

Un livre de Grégoire Allaire (voir [1]) nous permet d'énoncer quelques propriétés de l'espace H^1 .

L'espace H^1 est parfois appelé en physique quantique un espace d'énergie dans le sens où il est constitué de fonctions qui peuvent être associée à des fonctions d'onde d'énergie finie (c'est-à-dire de norme $\|u\|_{H^1(\Omega)}$ finie). Les fonctions d'énergie finie peuvent éventuellement être "singulières" ce qui a un sens physique possible (concentration ou explosion locale de certaines grandeurs).

Proposition 4.4.1 (H^1 est un espace de Hilbert). *Muni du produit scalaire*

$$\langle u, v \rangle = \int_{\Omega} (u(x)v(x) + \nabla u(x) \cdot \nabla v(x)) dx$$

et de la norme

$$\|u\|_{H^1(\Omega)} = \left(\int_{\Omega} (|u(x)|^2 + |\nabla u(x)|^2) dx \right)^{1/2}$$

l'espace de Sobolev $H^1(\Omega)$ est un espace de Hilbert.

Démonstration. Ceci est une preuve partielle. Nous allons nous contenter de prouver que $H^1(\Omega)$ est complet pour la norme associée.

Soit $(u_n)_{n \in \mathbb{N}}$ une suite de Cauchy dans $H^1(\Omega)$. Par définition de la norme de $H^1(\Omega)$, $(u_n)_{n \in \mathbb{N}}$ ainsi que $(\frac{\partial u_n}{\partial x_i})_{n \in \mathbb{N}}$ pour $i \in 1, \dots, \dim x$ sont des suites de Cauchy dans $L^2(\Omega)$. Comme $L^2(\Omega)$ est complet, il existe des limites u et w_i telles que u_n converge vers u et $\frac{\partial u_n}{\partial x_i}$ converge vers w_i dans $L^2(\Omega)$. Or, par définition de la dérivée faible de u_n , pour toute fonction $\phi \in C_c^\infty(\Omega)$, nous avons

$$\int_{\Omega} u_n(x) \frac{\partial \phi}{\partial x_i}(x) dx = - \int_{\Omega} \frac{\partial u_n}{\partial x_i}(x) \phi(x) dx.$$

Passant à la limite $n \rightarrow +\infty$, nous obtenons

$$\int_{\Omega} u(x) \frac{\partial \phi}{\partial x_i}(x) dx = - \int_{\Omega} w_i(x) \phi(x) dx,$$

ce qui prouve que u est dérivable au sens faible et que w_i est la i -ème dérivée partielle faible de u $\frac{\partial u}{\partial x_i}$. Donc, u appartient bien à $H^1(\Omega)$ et $(u_n)_{n \in \mathbb{N}}$ converge vers u dans $H^1(\Omega)$. □

Il est très important en pratique de savoir si les fonctions régulières sont denses dans l'espace de Sobolev $H^1(\Omega)$ (notamment pour établir des résultats depuis l'un de ces ensembles vers l'autre). Cela justifie en partie la notion d'espace de Sobolev qui apparaît ainsi très simplement comme l'ensemble des fonctions régulières complété par les limites de suites de fonctions régulières tout en gardant la norme de leur énergie $\|u\|_{H^1(\Omega)}$ finie.

Théorème 4.4.1 (densité de $C_c(\bar{\Omega})$ dans $H^1(\Omega)$). *Si Ω est un ouvert régulier de classe C^1 , ou bien si $\Omega = \mathbb{R}_+^N$, ou encore si $\Omega = \mathbb{R}^N$, alors $C_c^\infty(\bar{\Omega})$ est dense dans $H^1(\Omega)$.*

Remarque 4.4.1. L'espace $C_c^\infty(\bar{\Omega})$ qui est dense dans $H^1(\Omega)$ est constitué des fonctions régulières de classe C^∞ à support borné (ou compact) dans le fermé de $\bar{\Omega}$. En particulier, si Ω est borné, toutes les fonctions de $C^\infty(\bar{\Omega})$ ont nécessairement un support borné, et donc $C_c^\infty(\bar{\Omega}) = C^\infty(\bar{\Omega})$.

4.5 Convergence faible

La description qui va suivre de la convergence faible est issue du livre "Analyse numérique et optimisation" de Grégoire Allaire (voir [1]).

Definition 4.5.1 (Convergence faible). *Soit V un espace de Hilbert (complet - toute suite de Cauchy y converge - pour la distance issue de sa norme qui elle-même découle d'un produit scalaire). Une suite $(u_n)_{n \in \mathbb{N}}$ de V converge faiblement vers $u \in V$ si*

$$\forall v \in V, \lim_{n \rightarrow \infty} \langle u_n, v \rangle = \langle u, v \rangle.$$

La convergence faible est équivalente à la convergence de toutes les suites des composantes de $(u_n)_{n \in \mathbb{N}}$. L'intérêt de la convergence faible vient du résultat suivant.

Lemme 4.5.1. *De toute suite $(u_n)_{n \in \mathbb{N}}$ bornée dans V , nous pouvons extraire une sous-suite qui converge faiblement.*

4.6 Dérivée covariante

Certains ensembles, comme les variétés, peuvent être plongés dans un espace global d'une certaine dimension alors que localement ces ensembles peuvent être décrits via un système de coordonnées de dimension moindre. Il est alors fréquent que ce système de coordonnées que nous dirons "local" varie lorsque nous nous déplaçons dans notre ensemble. C'est dans ce contexte que la notion de covariance prend tout son sens.

L'adjectif covariant est utilisé pour décrire la manière avec laquelle une grandeur varie lors d'un changement de base. Des grandeurs sont dites covariantes lorsqu'elles varient comme les vecteurs de la base.

4.6.1 Définition d'un changement de base covariant

La définition suivante est basée sur un ouvrage de Jean Hladik (voir [19]).

Soit un espace vectoriel $\mathcal{V}$ de dimension finie n , ainsi que deux bases $(\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n)$ et $(\mathbf{e}'_1, \mathbf{e}'_2, \dots, \mathbf{e}'_n)$ telles que le changement de base de $\mathbf{e}$ vers $\mathbf{e}'$ s'écrit :

$$\mathbf{e}'_i = A_i^j \mathbf{e}_j.$$

Soit une famille de fonctions $(X(i))_{i=1\dots n}$ chacune de $\mathcal{V}^n$ vers un espace vectoriel de même corps que $\mathcal{V}$.

Les familles de vecteurs $(X(i)(\mathbf{e}'))_{i=1\dots n}$ et $(X(i)(\mathbf{e}))_{i=1\dots n}$ sont alors notées respectivement $(x'(i))_{i=1\dots n}$ et $(x(i))_{i=1\dots n}$.

X est dite "covariante" lorsque $x'(i) = \sum_{j=1}^n A_i^j x(j)$.

L'indice est alors noté en bas et la convention d'Einstein peut être utilisée, de telle sorte que nous notons

$$x'_i = A_i^j x_j.$$

Un élément est donc dit covariant lorsqu'il change via la matrice de changement de base sans changer de signe.

4.6.2 Introduction à la dérivée covariante

Le description ci-dessous de la dérivée covariante se base sur un ouvrage de Claude Semay et de Bernard Silvestre-Brac (voir [23]).

Lorsque nous fixons un système de coordonnées, cette dérivée se transforme lors d'un changement de base "de la même manière" que le vecteur lui-même (selon une transformation

covariante), d'où le nom de dérivée covariante. Nous avons besoin de termes décrivant la rotation et les déformations des coordonnées elles-mêmes.

La notion de dérivée covariante peut être motivée par la constatation suivante. Soit un vecteur e défini au pôle Nord d'une sphère et dirigé selon la tangente en ce point à un grand cercle (correspondant à un méridien) passant par lui (cette situation est dessinée à la figure 4.1). Supposons qu'on transporte parallèlement tout d'abord le vecteur le long de ce grand cercle et dans le sens choisie pour le vecteur tangent jusqu'à un point P situé sur l'équateur, puis (toujours grâce au transport parallèle le long de l'équateur) vers le point Q et enfin à nouveau jusqu'au pôle Nord le long d'un autre méridien. Nous remarquons alors que le transport parallèle de ce vecteur le long d'un circuit fermé ne retourne pas le même vecteur. À la place, il a pris une autre orientation. Ceci ne pourrait arriver dans un espace euclidien et est causé par la courbure de la surface du globe. Le même effet peut être observé si nous déplaçons le vecteur le long d'une surface infinitésimale. Le changement infinitésimal du vecteur est une mesure de la courbure. Pour tenir compte de la variation d'un élément, il faut donc, sur une surface de courbure non nulle, tenir compte d'une variation due à la variation du système de coordonnées local. C'est ce type d'ajout qu'apporte la dérivée covariante à la dérivée classique.

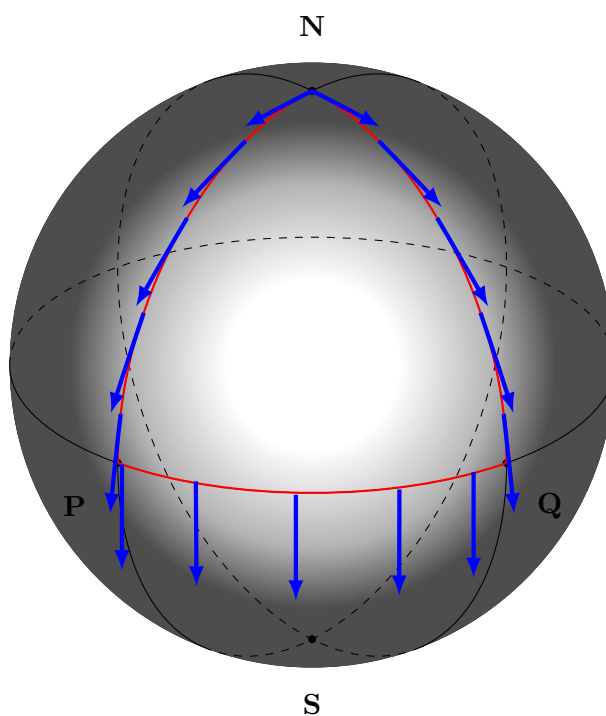


FIGURE 4.1 – Exemple de transport parallèle sur une sphère.

La dérivée covariante peut être décrite comme un tenseur dans un système de coordonnées donné, mais ce n'est pas un tenseur dans la mesure où elle n'est pas invariante par changement de coordonnées.

L'information sur le voisinage du point P dans la dérivée covariante peut être utilisée pour définir le transport parallèle d'un vecteur (le point d'application du vecteur se déplace sur des géodésiques successives de telle façon que le module de ce vecteur et l'angle qu'il fait avec cette géodésique restent constants). De même, la courbure et les géodésiques peuvent être définies uniquement en termes de dérivée covariante.

Une dérivée covariante ∇ d'un champ de vecteurs u est définie par les propriétés suivantes, valables pour tous champs de vecteurs v et w et toutes fonctions f et g :

1. $\nabla_v u$ est un champ de vecteurs, appelé dérivée covariante du champ de vecteurs u dans la direction v ;
2. $\nabla_v u$ est linéaire en v d'où $\nabla_{fv+gw} u = f\nabla_v u + g\nabla_w u$;
3. $\nabla_v u$ est additive en u d'où $\nabla_v(u+w) = \nabla_v u + \nabla_v w$;
4. $\nabla_v u$ obéit aux règles générales de la dérivation concernant le produit, c'est-à-dire $\nabla_v fu = f\nabla_v u + u\nabla_v f$.

Notons que $\nabla_v u$ au point P dépend de la valeur de v en P , ainsi que des valeurs de u dans un voisinage de P . Ceci signifie que la dérivée covariante n'est pas un tenseur.

Soit des coordonnées x^i , $i = 0, 1, 2, \dots$, tout vecteur tangent peut être décrit à l'aide de ses composantes dans la base e_i correspondant aux dérivations par rapport à x^i . La dérivée covariante est un vecteur et peut ainsi être exprimée comme la combinaison linéaire $\Gamma^k_{ij} e_k$ de tous les vecteurs de base, où Γ^k_{ij} représente chacune des composantes du vecteur (en notation d'Einstein). Pour décrire la dérivée covariante, il suffit de décrire celle de chacun des vecteurs de base e_j le long de e_i

$$\nabla_{e_i} e_j = \Gamma^k_{ij} e_k.$$

Les coefficients Γ^k_{ij} sont appelés les symboles de Christoffel. La donnée d'une dérivée covariante définit les coefficients de Christoffel correspondants. Réciproquement, la donnée de coefficients de Christoffel définit la dérivée covariante. En effet, en utilisant ces coefficients pour des vecteurs $\mathbf{v} = v^i e_i$ et $\mathbf{u} = u^i e_i$ nous avons :

$$\nabla_{\mathbf{v}} \mathbf{u} = \left(v^i u^j \Gamma^k_{ij} + v^i \frac{\partial u^k}{\partial x^i} \right) e_k,$$

Le premier terme de cette formule décrit la "déformation" du système de coordonnées par rapport à la dérivée covariante, et le second, les changements de coordonnées du vecteur $\mathbf{u}$. En particulier

$$\nabla_{e_j} \mathbf{u} = \nabla_j \mathbf{u} = \left(\frac{\partial u^i}{\partial x^j} + u^k \Gamma^i_{jk} \right) e_i$$

La dérivée covariante est la dérivée selon les coordonnées de départ à laquelle nous ajoutons des termes correctifs décrivant l'évolution des coordonnées. Donnons comme dernière équation celle du gradient covariant en utilisant les conventions d'Einstein

$$\nabla_c \mathbf{u} = \left(\frac{\partial u^i}{\partial x^j} + u^k \Gamma^i_{jk} \right) e_i e_j = \nabla u + \left(u^k \Gamma^i_{jk} \right) e_i e_j.$$

Nous prenons donc le gradient auquel nous rajoutons pour chaque coordonnée la façon dont les axes évoluent dans cette coordonnée.

Parfois la dérivée covariante le long d'une courbe est qualifiée de dérivée intrinsèque.

4.6.3 Projection sur l'espace tangent

Le gradient covariant peut aussi être définie comme étant en chaque point la projection du gradient sur l'espace tangent (l'espace tangent dépendant du point considéré). Nous pouvons alors définir le gradient covariant comme

$$\nabla_c u = \nabla u - n(n \cdot \nabla u). \quad (4.1)$$

Ce qui va nous être particulièrement utile c'est que nous allons travailler avec des espaces dont les éléments sont pris tangents à Ω considéré. Nous aurons alors, par intégration par parties,

$$n \cdot \nabla u = \nabla(n \cdot u) - \nabla n \cdot u$$

or, nous aurons alors $n \cdot u = 0$ et donc

$$n \cdot \nabla u = -\nabla n \cdot u$$

et l'équation (4.1) devient alors

$$\nabla_c u = \nabla u + n(\nabla n \cdot u). \quad (4.2)$$

4.7 Théorème de Rellich

Toujours dans le livre de G. Allaire (voir [1]), nous pouvons trouver les informations suivantes concernant le théorème de Rellich.

Nous consacrerons donc cette sous-section à l'étude d'une propriété de compacité connue sous le nom de théorème de Rellich. Rappelons tout d'abord que, dans un espace de Hilbert de dimension infinie, il n'est pas vrai que, de toute suite bornée, nous pouvons extraire une sous-suite convergente.

Théorème 4.7.1 (de Rellich). *Si Ω est un ouvert borné régulier de classe C^1 , alors de toute suite bornée de $H^1(\Omega)$ nous pouvons extraire une sous-suite convergente dans $L^2(\Omega)$ (nous disons que l'injection canonique de $H^1(\Omega)$ dans $L^2(\Omega)$ est compacte).*

4.8 Compacité faible

Les informations de cette sous-section sont issues d'un livre de Franck Boyer et de Pierre Fabrie (voir [10]).

Un espace est compact si, chaque fois qu'il est recouvert par des ouverts, il est recouvert par un nombre fini d'entre eux. Ce qui revient à dire que de toute suite d'un espace compact, nous pouvons extraire une sous-suite convergente. Dans $\mathbb{R}^n$, les compacts sont les fermés bornés.

Soit E un espace de Banach (espace vectoriel normé et complet pour la distance issue de sa norme). Dès que l'espace sur lequel nous travaillons est de dimension infinie (des espaces fonctionnels de type $L^p(\Omega)$ par exemple), les fermés bornés de cet espace ne sont pas compacts pour la topologie de la norme sur E . En revanche, le résultat suivant montre une propriété de compacité faible des fermés bornés.

Théorème 4.8.1. *Soit E un espace de Hilbert, alors de toute suite bornée d'éléments de E , nous pouvons extraire une sous-suite qui converge faiblement dans E .*

Nous avons aussi une propriété de semi-continuité inférieure de la norme pour la topologie faible sur un espace de Banach.

Proposition 4.8.1. *Soit E un espace de Banach et $(u_n)_{n \in \mathbb{N}}$ une suite d'éléments de E qui converge faiblement vers $u \in E$. Alors la suite $(u_n)_{n \in \mathbb{N}}$ est bornée dans E et nous avons*

$$\|u\|_E \leq \liminf_{n \rightarrow \infty} \|u_n\|_E.$$

Nous avons aussi dans le cadre des espaces de Hilbert, la proposition importante suivante qui nous donne un critère de convergence forte pour les suites faiblement convergentes dans un espace de Hilbert.

Proposition 4.8.2. Soient H un espace de Hilbert et $(u_n)_{n \in \mathbb{N}}$ une suite d'éléments de H qui converge vers u dans H . Nous supposons que

$$\lim_{n \rightarrow \infty} \sup \|u_n\|_H \leq \|u\|_H,$$

alors la suite $(u_n)_{n \in \mathbb{N}}$ converge fortement vers u dans H .

Notons aussi que la notion de convergence faible, bien que plus aisée à utiliser, ne permet pas en général de passer à la limite dans des termes non linéaires.

4.9 Existence

Nous cherchons à prouver l'existence de minimiseurs de la fonctionnelle de Ginzburg-Landau (1.1) sur des surfaces compactes et pour un ϵ quelconque.

La démonstration d'existence se fera selon les étapes indiquées ci-dessous.

1. Nous montrerons que l'ensemble de fonctions avec lequel nous travaillons est non vide et nous montrerons que $\inf E_\epsilon(u)$, avec u appartenant à cet ensemble, est borné.
2. Nous définirons une suite minimisante et nous utiliserons un argument de coupure pour montrer que nous pouvons prendre une suite minimisante dont chaque fonction est uniformément bornée.
3. Nous utiliserons la compacité faible dans $H^1(\Omega, \mathbb{R}^3)$ pour montrer la convergence vers un certain u_ϵ d'une sous-suite extraite de notre suite minimisante.
4. Nous montrerons que le premier terme de $E_\epsilon(u_n)$ est semi-continu inférieurement.
5. Nous utiliserons le théorème de Rellich pour montrer que, quand n tend vers l'infini, le deuxième terme de $E_\epsilon(u_n)$ tend fortement vers le deuxième terme de $E_\epsilon(u_\epsilon)$.
6. Via la semi-continuité inférieure de $E_\epsilon(u_n)$, nous montrerons que $\inf E(u) = u_\epsilon = \lim_{n \rightarrow \infty} E(u_n)$.

Théorème 4.9.1 (Existence). Soit Ω une surface fermée plongée dans $\mathbb{R}^3$. Il existe au moins une fonction u_ϵ , tangente en tout point à Ω , minimisant la fonctionnelle suivante :

$$E_\epsilon(u) = \frac{1}{2} \int_{\Omega} |\nabla_c u|^2 + \frac{1}{4\epsilon^2} \int_{\Omega} (1 - |u|^2)^2$$

dans l'espace $H^1(\Omega, \mathbb{R}^3)$.

Démonstration. Notons tout d'abord que les contraintes à respecter pour être dans $H^1(\Omega, \mathbb{R}^3)$ ne dépendent que d'intégrales (de Lebesgue) (la définition de H^1 a été donnée précédemment), nous négligerons donc les ensembles de mesure nulle (autrement dit, nous travaillerons presque partout avec nos fonctions de $H^1(\Omega, \mathbb{R}^3)$ dans ce qui suit).

Notons ensuite que l'ensemble des fonctions appartenant à $H^1(\Omega, \mathbb{R}^3)$ et étant tangent à Ω en tout point est non nul. En effet, la fonction $u_0(x) = (0, 0, 0) \forall x \in \Omega$ appartient à cette espace ($u_0(x) \cdot n(x) = 0 \forall x \in \Omega$ avec $n(x)$ la normale à Ω en tout x). Notons aussi que son énergie $E_\epsilon(u_0)$ est proportionnelle à $\frac{1}{4\epsilon^2} S(\Omega)$ où $S(\Omega)$ donne l'aire de la surface Ω . Remarquons aussi que $E_\epsilon(u) \geq 0$ puisque la fonctionnelle considérée est la somme des intégrales de deux carrés. Nous savons donc que $m = \inf E_\epsilon(u)$ est borné puisque $0 \leq m \leq \frac{1}{4\epsilon^2} S(\Omega)$.

Soit $(u_n)_{n \in \mathbb{N}}$ une suite minimisante de $E_\epsilon(u)$ dans l'espace $H^1(\Omega, \mathbb{R}^3)$ et telle que chacun de ses éléments soit tangent à Ω . Notons que remplacer n'importe quelle fonction v par

$$\tilde{v} = \begin{cases} v/|v| & \text{si } |v| > 1 \\ v & \text{sinon,} \end{cases}$$

implique que $E_\epsilon(\tilde{v}) \leq E_\epsilon(v)$. En effet, nous avons,

$$E_\epsilon(\tilde{v}) = \frac{1}{2} \int_{\Omega} |\nabla_c(\tilde{v})|^2 + \frac{1}{4\epsilon^2} \int_{\Omega} (1 - |\tilde{v}|^2)^2$$

or,

$$\int_{\Omega} |\nabla_c(\tilde{v})|^2 = \frac{1}{2} \int_{T^2} \left| \nabla_c \left(\frac{v}{|v|} \right) \right|^2 \mathbb{1}_{|v|>1} + \frac{1}{2} \int_{T^2} |\nabla_c(v)|^2 \mathbb{1}_{|v|\leq 1}$$

où $\mathbb{1}_E$ est la fonction indicatrice de l'ensemble E .

Développons $|\nabla_c(\tilde{v})|^2$ dans le cas où $|v| > 1$. Notons d'abord que

$$\nabla(|v|) = \frac{v}{|v|} \nabla v,$$

et nous avons donc

$$\begin{aligned} \nabla \left(\frac{v}{|v|} \right) &= \frac{|v| \nabla v - \frac{v^2 \nabla v}{|v|}}{|v|^2} \\ &= \frac{\nabla v}{|v|} - \frac{v^2 \nabla v}{|v|^3}. \end{aligned}$$

Nous avons alors

$$\begin{aligned} \nabla_c \left(\frac{v}{|v|} \right) &= \nabla \left(\frac{v}{|v|} \right) + n \left(\nabla n \cdot \frac{v}{|v|} \right) \\ &= \frac{\nabla v}{|v|} - \frac{v^2 \nabla v}{|v|^3} + n \left(\nabla n \cdot \frac{v}{|v|} \right). \end{aligned}$$

Notons que nous travaillons avec des opérateurs nabla sur des fonctions et que nous travaillons donc avec des jacobiniennes; autrement dit $\nabla v = Dv$ où Dv est la jacobienne de v . Les opérations entre nos objets sont donc des opérations de calcul matriciel. Les utiliser comme tels nous permettra d'avancer de façon rigoureuse.

Nous avons

$$\begin{aligned} \left| \nabla_c \left(\frac{v}{|v|} \right) \right|^2 &= \left| \left(\frac{\nabla v}{|v|} + n \left(\nabla n \cdot \frac{v}{|v|} \right) \right) - \frac{v^2 \nabla v}{|v|^3} \right|^2 \\ &= \left| \left(\frac{\nabla v}{|v|} + n \left(\nabla n \cdot \frac{v}{|v|} \right) \right) \right|^2 - 2 \left(\frac{\nabla v}{|v|} + n \left(\nabla n \cdot \frac{v}{|v|} \right) \right) \cdot \frac{v^2 \nabla v}{|v|^3} + \left| \frac{v^2 \nabla v}{|v|^3} \right|^2 \\ &= \left| \frac{\nabla_c v}{|v|} \right|^2 - 2 \frac{\nabla v}{|v|} \cdot \frac{v^2 \nabla v}{|v|^3} - 2n \left(\nabla n \cdot \frac{v}{|v|} \right) \cdot \frac{v^2 \nabla v}{|v|^3} + \left| \frac{v^2 \nabla v}{|v|^3} \right|^2 \\ &= \left| \frac{\nabla_c v}{|v|} \right|^2 - 2tr \left(\frac{(Dv)^t v v^t Dv}{|v|^4} \right) - 2tr \left(\frac{(Dv)^t v v^t n v^t Dn}{|v|^4} \right) + tr \left(\frac{(Dv)^t v v^t v v^t Dv}{|v|^6} \right) \end{aligned}$$

où tr est l'opérateur qui nous permet de prendre la trace d'une matrice. Or, $v \cdot n = v^t n = 0$ et $v \cdot v = v^t v = |v|^2$ donc nous obtenons finalement

$$\begin{aligned} \left| \nabla_c \left(\frac{v}{|v|} \right) \right|^2 &= \left| \frac{\nabla_c v}{|v|} \right|^2 - tr \left(\frac{(Dv)^t v v^t Dv}{|v|^4} \right) \\ &= \left| \frac{\nabla_c v}{|v|} \right|^2 - tr \left(\frac{(v^t Dv)^t v^t Dv}{|v|^4} \right) \\ &\leq |\nabla_c v|^2 \end{aligned}$$

car nous sommes dans le cas $|v| > 1$ et que $\text{tr} \left(\frac{(v^t Dv)^t v^t Dv}{|v|^4} \right) > 0$.

En conclusion,

$$\int_{\Omega} |\nabla_c(\tilde{v})|^2 = \frac{1}{2} \int_{T^2} \left| \nabla_c \left(\frac{v}{|v|} \right) \right|^2 \mathbb{1}_{|v|>1} + \frac{1}{2} \int_{T^2} |\nabla_c(v)|^2 \mathbb{1}_{|v|\leq 1} \leq |\nabla_c v|^2.$$

Notons que nous avons aussi

$$\begin{aligned} \frac{1}{4\epsilon^2} \int_{\Omega} (1 - |\tilde{v}|^2)^2 &= \frac{1}{4\epsilon^2} \int_{\Omega} (1 - \left| \frac{v}{|v|} \right|^2)^2 \mathbb{1}_{|v|>1} + \frac{1}{4\epsilon^2} \int_{\Omega} (1 - |v|^2)^2 \mathbb{1}_{|v|\leq 1} \\ &= \frac{1}{4\epsilon^2} \int_{\Omega} (1 - |v|^2)^2 \mathbb{1}_{|v|\leq 1} \\ &\leq \frac{1}{4\epsilon^2} \int_{\Omega} (1 - |v|^2)^2 \end{aligned}$$

ce qui implique bien que

$$E_{\epsilon}(\tilde{v}) \leq E_{\epsilon}(v).$$

Cela s'appelle une coupure. Admettons que notre suite minimisante soit le résultat de la coupure d'une autre suite minimisante. Nous avons alors que $\forall n$ et $\forall x \in \Omega$, $|u_n(x)| \leq 1$ et les fonctions successives de notre suite minimisante sont uniformément bornées.

Nous savons donc que les u_n sont uniformément bornés et il nous reste à montrer que ∇u est borné pour pouvoir utiliser la compacité faible dans $H^1(\Omega)$. Notons tout d'abord que le gradient covariant des u_n est borné puisque $(u_n)_{n \in \mathbb{N}}$ est une suite minimisante et que donc

$$\int_{\Omega} |\nabla_c u_n|^2 \leq E(u_n) \leq M.$$

De plus, nous avons, grâce à l'identité $(a + b)^2 \leq 2(a^2 + b^2)$,

$$\begin{aligned} \int_{\Omega} |\nabla u|^2 &\leq 2 \int_{\Omega} |\nabla u + |n(u^t \nabla n)|^2 \\ &\leq \int_{\Omega} |\nabla u|^2 + |n(n \cdot \nabla u)|^2, \end{aligned}$$

puisque nous pouvons supposer que la normale à Ω est bornée ainsi que ses dérivées.

Nous pouvons maintenant utiliser la compacité faible dans $H^1(\Omega, \mathbb{R}^3)$ pour dire qu'il existe une sous-suite de notre suite minimisante qui converge faiblement vers un certain u_{ϵ} .

Par ailleurs, nous avons les résultats suivants.

$$\begin{aligned} \int_{\Omega} |\nabla_c u_{\epsilon}|^2 &= \int_{\Omega} |\nabla_c(u_{\epsilon} - u_n + u_n)|^2 \\ &= \int_{\Omega} |\nabla_c(u_{\epsilon} - u_n) + \nabla_c(u_n)|^2 \\ &= \int_{\Omega} |\nabla_c(u_{\epsilon} - u_n)|^2 + 2 \int_{\Omega} \nabla_c(u_{\epsilon} - u_n) \cdot \nabla_c(u_n) + \int_{\Omega} |\nabla_c(u_n)|^2, \end{aligned}$$

or, par convergence faible de $u_n \rightharpoonup u_{\epsilon}$ dans $H^1(\Omega, \mathbb{R}^3)$ et par linéarité des gradients et des symboles de Christoffel (ceci a été introduit précédemment dans ce chapitre),

$$\lim_{n \rightarrow \infty} \int_{\Omega} \nabla_c(u_{\epsilon} - u_n) \cdot \nabla_c(u_n) = 0;$$

ainsi,

$$\lim_{n \rightarrow \infty} \int_{\Omega} |\nabla_c(u_n)|^2 \leq \inf \int_{\Omega} |\nabla_c(u_\epsilon)|^2$$

et le premier terme de $E_\epsilon(u_n)$ est donc semi-continu inférieurement.

Le développement du deuxième terme de $E_\epsilon(u)$ nous donne $\int_{\Omega} 1 - 2|u|^2 + |u|^4$ à un facteur multiplicatif près. Grâce au théorème de Rellich, nous savons que $(u_n) \rightarrow u_\epsilon$ fortement dans L^2 . Nous avons

$$\int_{\Omega} |u_n - u_\epsilon|^2 = \int_{\Omega} |u_n|^2 - 2 \int_{\Omega} u_n \cdot u_\epsilon + \int_{\Omega} |u_\epsilon|^2$$

ce qui nous donne, via la convergence faible de u_n vers u_ϵ ,

$$\lim_{n \rightarrow \infty} \int_{\Omega} |u_n - u_\epsilon|^2 = \lim_{n \rightarrow \infty} \int_{\Omega} |u_n|^2 - \lim_{n \rightarrow \infty} \int_{\Omega} |u_\epsilon|^2$$

et donc $\int_{\Omega} |u_n|^2$ tend vers $\int_{\Omega} |u_\epsilon|^2$. Par le corollaire 4.3.1, nous savons que $H^1(\Omega) \subset L^p(\Omega) \quad \forall [1, \infty[$ avec injection continue. Donc, via l'inégalité triangulaire dans L^4 , nous savons que

$$|\int_{\Omega} |u_n|^4 - \int_{\Omega} |u_\epsilon|^4| \leq \int_{\Omega} (u_n - u_\epsilon)^4$$

et donc nous savons que $\int_{\Omega} |u_n|^4 \rightarrow \int_{\Omega} |u_\epsilon|^4$. Ainsi, nous avons montré la convergence forte du deuxième terme.

Nous avons donc montré via la semi-continuité inférieure de $E_\epsilon(u_n)$ que

$$\lim_{n \rightarrow \infty} E_\epsilon(u_n) = E_\epsilon(u_\epsilon) \leq \inf E_\epsilon(u) \geq E_\epsilon(u_\epsilon)$$

ainsi, notre suite minimisante tend bien vers un minimiseur de $E_\epsilon(u)$.

Notons finalement que puisque $u_n \rightarrow u$ dans $L^2(\Omega)$ nous avons

$$u_n \cdot n \rightarrow u \cdot n$$

et donc notre minimiseur est bien tangent à Ω en tout point puisque les u_n le sont. □

Chapitre 5

Équations d'Euler-Lagrange

5.1 Introduction

L'équation d'Euler-Lagrange permet d'obtenir une EDP qui peut nous renseigner sur les points fixes d'un problème de minimisation. Numériquement, une stratégie possible est d'attaquer un problème de minimisation via son équation d'Euler-Lagrange plutôt que d'attaquer directement le problème de minimisation.

5.2 Minimisation et équation d'Euler-Lagrange

Nous pouvons tirer d'un ouvrage de Claude Le Bris (voir [12]) les informations suivantes.

Pour résoudre numériquement un problème de minimisation énergétique, nous pouvons ou bien minimiser directement la fonctionnelle d'énergie, ou bien résoudre les équations d'Euler-Lagrange associées à ce problème de minimisation.

Le plus souvent, nous avons intérêt pour optimiser le temps de calcul, à résoudre les équations d'Euler-Lagrange plutôt qu'à minimiser directement la fonctionnelle d'énergie. Pour ce faire, en présence d'éléments non-linéaires, il faut utiliser une procédure itérative. La difficulté vient alors du fait que rien n'assure alors a priori la décroissance de l'énergie et que des difficultés de convergence peuvent apparaître.

Notons que minimiser directement la fonctionnelle d'énergie et résoudre les équations d'Euler-Lagrange ne sont pas théoriquement deux stratégies équivalentes (sauf dans le cas de problèmes convexes, ce qui n'est le cas d'aucune des fonctionnelles de Ginzburg-Landau étudiées dans ce mémoire). D'une part, minimiser directement la fonctionnelle, c'est prendre le risque de rester "bloqué" dans un minimum local, non global. D'autre part, résoudre les équations d'Euler-Lagrange, c'est aussi prendre le risque de déterminer un point critique qui n'est pas un minimum global puisque l'équation d'Euler-Lagrange d'un problème est satisfaite pour tout point stationnaire de ce problème. Respecter les équations d'Euler-Lagrange ne nous garantit donc même pas d'être à un minimum.

5.3 Comment trouver les équations d'Euler-Lagrange

Admettons que nous ayons connaissance d'une fonction stationnaire u pour une certaine fonctionnelle énergétique E , c'est-à-dire une fonction u telle que

$$\frac{d}{du}(E(u)) = 0.$$

Nous pouvons alors définir une équation $t \rightarrow E(u + t\phi)$ avec $\phi \in C_c^\infty$ qui est une équation en t qui atteint son minimum quand $t = 0$.

Nous savons alors que quand $t = 0$,

$$\frac{d}{dt}\bigg|_{t=0} E(u + t\phi) = 0$$

et ce $\forall \phi \in C_c^\infty$.

Notons aussi que $(u(x)\phi(x))' = u'(x)\phi(x) + u(x)\phi'(x)$ et que $\int (u(x)\phi(x))' = [u(x)\phi(x)]_{-\infty}^{+\infty} = 0$ car $\phi \in C_c^\infty$ ce qui permet de simplifier des calculs lorsque nous cherchons à obtenir des équations d'Euler-Lagrange.

5.4 Dérivation des équations d'Euler-Lagrange

Soit Ω une surface fermée régulière et bornée plongée dans $\mathbb{R}^3$. Nous allons considérer des cartes complexes sur Ω , c'est-à-dire des cartes de Ω dans $\mathbb{R}^2$ et la fonctionnelle de Ginzburg-Landau

$$E_\epsilon(v) = \frac{1}{2} \int_\Omega |\nabla_c v|^2 + \frac{1}{4\epsilon^2} \int_\Omega (1 - |v|^2)^2 \quad (5.1)$$

où ϵ est un paramètre strictement positif et homogène à une longueur. Nous cherchons à dériver l'équation d'Euler-Lagrange associée à la minimisation de la fonctionnelle énergétique (5.1).

Nous définissons maintenant l'espace de Sobolev suivant, dans lequel nous voulons trouver nos minimiseurs de (5.1). Nous cherchons un $u \in H^1(\Omega; \mathbb{R}^3)$ et parallèle à Ω en tout point de Ω tel que u minimise (5.1). Autrement dit, nous cherchons des fonctions $u \in H^1(\Omega; \mathbb{R}^3)$ telles que $u(x) \cdot n(x) = 0 \quad \forall x \in \Omega$ où $n(x)$ désigne la normale à Ω au point x de Ω .

Si nous trouvons une fonction $u \in H^1(\Omega, \mathbb{R}^3)$ et parallèle à Ω dans Ω telle que $E_\epsilon(u) \leq E_\epsilon(v)$ pour tout $v \in H^1(\Omega, \mathbb{R}^3)$ parallèle à Ω , alors nous savons que

$$\forall \phi \in \{C_c^\infty(\Omega) \text{ et } T_x(\Omega)\}, \quad E_\epsilon(u) \leq E_\epsilon(u + t\phi)$$

où $T_x(\Omega)$ est l'espace des fonctions tangentes en tout point à la normale de Ω . Or,

$$E_\epsilon(u + t\phi) = \frac{1}{2} \int_\Omega |\nabla_c(u + t\phi)|^2 + \frac{1}{4\epsilon^2} \int_\Omega (1 - |u + t\phi|^2)^2.$$

Nous savons, grâce à ce que nous avons vu précédemment, que

$$\begin{aligned} \nabla_c u &= \left(\frac{\partial u_i}{\partial x^j} + u_k \Gamma_{jk}^i \right) e_i e_j \\ &= \nabla u + (u_k \Gamma_{jk}^i) e_i e_j. \end{aligned}$$

Nous avons donc

$$\begin{aligned} E_\epsilon(u + t\phi) &= \frac{1}{2} \int_\Omega |\nabla_c(u + t\phi)|^2 + \frac{1}{4\epsilon^2} \int_\Omega (1 - |u + t\phi|^2)^2 \\ &= \frac{1}{2} \int_\Omega \left| \nabla(u + t\phi) + (u_k + t\phi_k) \Gamma_{jk}^i e_i e_j \right|^2 + \frac{1}{4\epsilon^2} \int_\Omega (1 - |u + t\phi|^2)^2 \\ &= \frac{1}{2} \int_\Omega \left| \nabla u + t \nabla \phi + (u_k + t\phi_k) \Gamma_{jk}^i e_i e_j \right|^2 + \frac{1}{4\epsilon^2} \int_\Omega (1 - |u + t\phi|^2)^2. \end{aligned}$$

Définissons, par souci de concision,

$$\begin{aligned} A(u + t\phi) &\triangleq (u_k + t\phi_k) \Gamma_{jk}^i e_i e_j \\ &= \sum_{j=1}^3 \sum_{i=1}^3 \sum_{k=1}^3 (u_k(x) + t\phi_k(x)) \Gamma_{jk}^i(x). \end{aligned}$$

Notons que $\frac{d}{dt}A(u+t\phi) = A(\phi)$, que $A(u+t\phi) = A(u) + tA(\phi)$ et que $A(0) = 0$.

Nous savons par hypothèse que lorsque $t \rightarrow 0$, $\frac{d}{dt}E_\epsilon(u+t\phi) = 0$. De plus, nous avons

$$\begin{aligned}\frac{d}{dt}E_\epsilon(u+t\phi) &= \frac{1}{2} \int_{\Omega} 2(\nabla\phi + A(\phi)) \cdot (t\nabla\phi + \nabla u + A(u+t\phi)) - \frac{1}{4\epsilon^2} \int_{\Omega} 4\phi \cdot (u+t\phi)(1-|u+t\phi|^2) \\ &= \int_{\Omega} (\nabla\phi + A(\phi)) \cdot (t\nabla\phi + \nabla u + A(u+t\phi)) - \frac{1}{\epsilon^2} \int_{\Omega} \phi \cdot (u+t\phi)(1-|u+t\phi|^2)\end{aligned}$$

et donc

$$\frac{d}{dt}|_{t=0}E_\epsilon(u+t\phi) = \int_{\Omega} \nabla_c\phi \cdot \nabla_c u - \frac{1}{\epsilon^2} \int_{\Omega} \phi \cdot u(1-|u|^2) = 0.$$

Notons que

$$\int_{\Omega} (\nabla\phi + A(\phi)) \cdot (\nabla u + A(u)) = \int_{\Omega} \nabla\phi \cdot \nabla u + \nabla\phi \cdot A(u) + A(\phi) \cdot \nabla u + A(\phi) \cdot A(u).$$

Or, par intégration par parties et par linéarité des symboles de Christoffel, puisque $\phi \in C_c^\infty(\Omega)$, nous avons

$$\begin{aligned}\int_{\Omega} \nabla\phi \cdot \nabla u &= - \int_{\Omega} \Delta u \cdot \phi \\ \int_{\Omega} \nabla\phi \cdot A(u) &= - \int_{\Omega} \phi \cdot \nabla A(u),\end{aligned}$$

Calculons maintenant ce que nous donne $\nabla(A(u))$. Nous avons

$$\begin{aligned}\nabla A(u) &= \sum_{m=1}^3 \frac{\partial}{\partial x_m} \sum_{j=1}^3 \sum_{i=1}^3 \sum_{k=1}^3 u_k(x) \Gamma_{jk}^i(x) \\ &= \sum_{j=1}^3 \sum_{i=1}^3 \sum_{k=1}^3 \frac{\partial u_k(x)}{\partial x_k} \Gamma_{jk}^i(x) + \sum_{j=1}^3 \sum_{i=1}^3 \sum_{k=1}^3 u_k(x) \frac{\partial \Gamma_{jk}^i(x)}{\partial x} \\ &= A(\nabla u) + u \nabla \Gamma_{jk}^i.\end{aligned}$$

Il nous reste à réussir à isoler ϕ dans l'expression $A(\phi) \cdot \nabla u$ et dans l'expression $A(\phi) \cdot A(u)$. Nous avons d'abord

$$\begin{aligned}A(\phi) \cdot \nabla u &= \left(\sum_{j=1}^3 \sum_{i=1}^3 \sum_{k=1}^3 \phi_k(x) \Gamma_{jk}^i(x) \right) \cdot \left(\sum_{m=1}^3 \frac{\partial u_m(x)}{\partial x_m} \right) \\ &= \sum_{j=1}^3 \sum_{i=1}^3 \sum_{k=1}^3 \frac{\partial u_k(x)}{\partial x_k} \phi_k(x) \Gamma_{jk}^i(x) \\ &= A(\nabla u) \cdot \phi.\end{aligned}$$

Nous avons ensuite

$$\begin{aligned}A(\phi) \cdot A(u) &= \left(\sum_{j=1}^3 \sum_{i=1}^3 \sum_{k=1}^3 \phi_k(x) \Gamma_{jk}^i(x) \right) \cdot \left(\sum_{m=1}^3 \sum_{n=1}^3 \sum_{p=1}^3 u_p(x) \Gamma_{mp}^n(x) \right) \\ &= \sum_{j=1}^3 \sum_{i=1}^3 \sum_{k=1}^3 \sum_{m=1}^3 \sum_{n=1}^3 u_k \phi_k \Gamma_{jk}^i \Gamma_{mk}^n \\ &= \left(\sum_{j=1}^3 \sum_{i=1}^3 \sum_{k=1}^3 \sum_{m=1}^3 \sum_{n=1}^3 \Gamma_{jk}^i \Gamma_{mk}^n u_k \right) \cdot \phi \\ &\triangleq B(u) \cdot \phi.\end{aligned}$$

Nous avons donc

$$\begin{aligned}
\int_{\Omega} \nabla_c \phi \cdot \nabla_c u &= \int_{\Omega} \nabla \phi \cdot \nabla u + \nabla \phi \cdot A(u) + A(\phi) \cdot \nabla u + A(\phi) \cdot A(u) + B(u) \cdot \phi \\
&= \int_{\Omega} -\Delta u \cdot \phi - \phi \cdot (A(\nabla u) + u \nabla \Gamma_{jk}^i) + A(\nabla u) \cdot \phi + B(u) \cdot \phi \\
&= \int_{\Omega} -\Delta u \cdot \phi - u \nabla \Gamma_{jk}^i \cdot \phi + B(u) \cdot \phi.
\end{aligned}$$

Posons

$$\tilde{C}u \triangleq u \nabla \Gamma_{jk}^i.$$

nous obtenons donc que

$$\frac{d}{dt} \Big|_{t=0} E_{\epsilon}(u + t\phi) = - \int_{\Omega} \phi \cdot (\Delta u - B(u) + \tilde{C}u) - \frac{1}{\epsilon^2} \int_{\Omega} \phi \cdot v(1 + |v|^2) = 0$$

pour u et $\phi \in \{H^1(\Omega) \text{ et } T_x(\Omega)\}$.

Essayons d'obtenir le même genre de résultats pour des u et $\phi \in H^1(\Omega)$. Nous avons, puisque nous travaillons avec des éléments tangents à Ω , les résultats suivants où le fait que u et ϕ soient tangents à Ω est directement utilisé dans l'équation ;

$$\begin{aligned}
-\Delta u + B(u) - \tilde{C}u - n(n \cdot (-\Delta u + B(u) - \tilde{C}u)) &= -\Delta u + B(u) - \tilde{C}u + n(n \cdot (\Delta u - B(u) + \tilde{C}u)) \\
&= \frac{1}{\epsilon^2} u(1 + |u|^2) - (n \cdot (\frac{1}{\epsilon^2} u(1 + |u|^2))).
\end{aligned}$$

Or, puisque $n \cdot u = 0$, nous obtenons

$$n(n \cdot (\Delta u - B(u) + \tilde{C}u)) = n(n \cdot \Delta u)$$

et

$$n \cdot (\frac{1}{\epsilon^2} u(1 + |u|^2)) = 0.$$

Nous avons alors

$$-(\Delta u - n(n \cdot \Delta u)) + B(u) - \tilde{C}u = \frac{1}{\epsilon^2} u(1 - |u|^2).$$

Posons

$$\Delta_c = \Delta u - n(n \cdot \Delta u) - B(u) + \tilde{C}u$$

et remarquons que via des petites astuces mathématiques que $n(n \cdot \Delta u)$ ne dépend pas du laplacien de u mais seulement de dérivées inférieures. En effet, nous avons, en utilisant les conventions de sommation d'Einstein et puisque $n \cdot u = 0$,

$$\begin{aligned}
0 &= \partial_k \partial_k (n_j u_j) \\
&= \partial_k (n_j \partial_k u_j + u_j \partial_k n_j) \\
&= n_j \partial_k \partial_k u_j + 2 \partial_k n_j \partial_k u_j + u_j \partial_k \partial_k n_j \\
&= n \cdot \Delta u + 2 \nabla n \cdot \nabla u + u \cdot \Delta n,
\end{aligned}$$

nous avons donc

$$n \cdot \Delta u = -2 \nabla n \cdot \nabla u - u \cdot \Delta n$$

et

$$n(n \cdot \Delta u) = -2n \nabla n \cdot \nabla u - n(\Delta n \cdot u).$$

Ainsi, nous pouvons trouver une équation pour le laplacien covariant

$$\Delta_c = \Delta u + 2n \nabla n \cdot \nabla u + n(\Delta n \cdot u) - B(u) + \tilde{C}u$$

qui ne dépend que du laplacien classique et de termes d'ordre inférieur. Cela sera primordial dans l'étude de la régularité de nos solutions. Pour simplifier l'écriture de nos calculs ultérieurs, posons

$$\begin{aligned} -C(u) &= n(\Delta n \cdot u) - B(u) + \tilde{C}u \\ -D(\nabla u) &= 2n \nabla n \cdot \nabla u. \end{aligned}$$

Ainsi,

$$\Delta_c u = \Delta u - D(\nabla u) - C(u).$$

Nous obtenons finalement que

$$\frac{d}{dt} \Big|_{t=0} E_\epsilon(u + t\phi) = - \int_{\Omega} \phi \cdot \Delta_c u - \frac{1}{\epsilon^2} \int_{\Omega} \phi \cdot v(1 + |v|^2) = 0$$

pour u et $\phi \in H^1(\Omega)$. Nous avons donc

$$\langle \phi, -\Delta_c u \rangle = \langle \phi, \frac{1}{\epsilon^2} u(1 + |u|^2) \rangle$$

dans $H^1(\Omega, \mathbb{R}^2)$ ce qui n'est vrai que si

$$-\Delta_c u = \frac{1}{\epsilon^2} u(1 + |u|^2)$$

dans $H^1(\Omega, \mathbb{R}^2)$.

Cette dernière égalité n'est rien de moins que l'équation d'Euler-Lagrange de notre problème de minimisation.

Chapitre 6

Régularité des solutions

6.1 Introduction

Nous allons maintenant utiliser l'équation d'Euler-Lagrange obtenue précédemment pour obtenir la classe de régularité des solutions de notre problème de minimisation (5.1) via l'utilisation de théorèmes d'injection (dont le théorème de Rellich déjà introduit précédemment).

Notons que le fait de travailler avec des variétés sans bord va nous simplifier la tâche. En effet, pour démontrer un résultat de régularité sur un ensemble, nous étudions la régularité d'un ensemble l'englobant sauf lorsque cet ensemble touche un bord auquel cas, il faut faire intervenir la régularité du bord. Dès lors, puisque intuitivement (le travail avec des cartes locales reste toutefois laborieux) nous pouvons trouver pour tout sous-ensemble inclus dans nos variétés sans bord un ensemble l'englobant, l'obtention de résultats de régularité ne risque pas de se voir complexifier ou compromettre par la présence d'un bord plus ou moins régulier.

Nous commencerons par définir des résultats utiles puis nous les utiliserons avec l'équation d'Euler-Lagrange pour trouver la régularité des solutions de (5.1).

6.2 Injections de Sobolev

Rappelons que les espaces $W^{m,p}$ sont donc des espaces dans lesquels chaque élément possède des dérivées faibles au moins jusque l'ordre m et dans lesquels chaque fonction et ses dérivées successives jusqu'à l'ordre m ont toutes leurs puissances jusqu'à la puissance p qui sont intégrables.

Rappelons aussi qu'un espace E s'injecte continûment dans un espace F si $E \subset F$ et si la restriction à E d'une application continue définie sur F est encore continue. Il existe alors une constante C telle que pour tout $u \in E$, $\|u\|_F \leq C\|u\|_E$.

Rappelons finalement qu'une injection est dite compacte de E dans F si la boule unité de E , vue dans F , est relativement compacte, c'est-à-dire si $\overline{B_E(0,1)}$ est compacte dans F .

Les résultats d'injection concernant des espaces de Sobolev permettent de "comparer", au sens de l'inclusion, les espaces de Sobolev entre eux mais aussi de comparer certains espaces de Sobolev à des espaces de fonctions continues.

Les résultats suivants sont issus du livre "Analyse fonctionnelle" de H.Brezis ([11]).

Rappelons que nous pouvons utiliser les résultats obtenus sur $\mathbb{R}^N$ sur nos Ω puisque nos Ω peuvent être construits sur base d'un recouvrement fini d'ouverts de $\mathbb{R}$. Il peut être montré que si Ω est de dimension 1, alors $W^{1,p} \subset L^\infty(\Omega)$ avec injection continue. En dimension $N \geq 2$ cette inclusion reste vraie seulement pour $p > N$; lorsque $p \leq N$ des exemples de fonctions de $W^{1,p}$ qui n'appartiennent pas à L^∞ peuvent être construits. Néanmoins un résultat important, dû essentiellement à Sobolev, affirme que si $1 \leq p < N$ alors $W^{1,p}(\Omega) \subset L^{p^*}(\Omega)$ avec injection continue pour un certain $p^* \in]p, +\infty[$. Nous avons en fait déjà introduit de tels résultats nous

concernant précédemment (4.3.2 et surtout 4.3.1).

Continuons à énoncer des théorèmes remarquables de régularité.

Soit $I =]a, b[$ un intervalle quelconque et soit $p \in \mathbb{R}$ avec $1 \leq p \leq \infty$.

Théorème 6.2.1. *Il existe une constante C (dépendant seulement de $|I| \leq \infty$) telle que*

$$\|u\|_{L^\infty(I)} \leq C \|u\|_{W^{1,p}(I)} \quad \forall u \in W^{1,p}(I), \quad \forall 1 \leq p \leq \infty,$$

autrement dit $W^{1,p}(I) \subset L^\infty(I)$ avec injection continue pour tout $1 \leq p \leq \infty$.

De plus, lorsque I est borné, nous avons

1. *l'injection $W^{1,p}(I) \subset C(\bar{I})$ est compacte pour $1 < p \leq \infty$*
2. *l'injection $W^{1,1}(I) \subset L^q(I)$ est compacte pour $1 \leq q < \infty$.*

Théorème 6.2.2 (Morrey). *Soit $p > N$, alors*

$$W^{1,p}(\mathbb{R}^N) \subset L^\infty(\mathbb{R}^N)$$

avec injection continue.

De plus, pour tout $u \in W^{1,p}(\mathbb{R}^N)$, nous avons

$$|u(x) - u(y)| \leq C |x - y|^\alpha \|\nabla u\|_{L^p} \quad p.p. \quad x, y \in \mathbb{R}^N \quad (6.1)$$

avec $\alpha = 1 - \frac{N}{p}$ et C est une constante (qui dépend seulement de p et N).

L'inégalité (6.1) implique l'existence d'une fonction $\tilde{u} \in C(\mathbb{R}^N)$ telle que $u = \tilde{u}$ p.p sur $\mathbb{R}^N$. Autrement dit toute fonction $u \in W^{1,p}$, $p > N$, admet un représentant continu.

Théorème 6.2.3 (Régularité pour le problème de Dirichlet). *Soit Ω une variété de classe C^2 ou bien $\Omega = \mathbb{R}^N$. Soit $f \in L^2(\Omega)$ et soit $u \in H^1(\Omega)$ vérifiant*

$$\int_{\Omega} \nabla u \nabla \phi + \int_{\Omega} u \phi = \int_{\Omega} f \phi \quad \forall \phi \in H_0^1(\Omega).$$

Alors $u \in H^2(\Omega)$ et $\|u\|_{H^2} \leq C \|f\|_{L^2}$ où C est une constante qui dépend seulement de Ω . De plus, si Ω est de classe C^{m+2} et si $f \in H^m(\Omega)$, alors

$$u \in H^{m+2} \text{ avec } \|u\|_{H^{m+2}} \leq C \|f\|_{H^m};$$

en particulier si $m > \frac{N}{2}$, alors $u \in C^2(\bar{\Omega})$.

Enfin si Ω est de classe C^∞ et si $f \in C^\infty(\bar{\Omega})$, alors $u \in C^\infty(\bar{\Omega})$.

6.3 Régularité des solutions

Admettons que nous travaillons sur un domaine Ω borné et régulier.

Commençons par nous rappeler que nos solutions se trouvent dans l'ensemble $H^1(\Omega) = W^{1,2}(\Omega)$. Rappelons-nous aussi que le théorème de Rellich nous indiquait que si Ω est un ouvert borné régulier de classe C^1 , alors de toute suite bornée de $H^1(\Omega)$ nous pouvons extraire une sous-suite convergente dans $L^2(\Omega)$.

Notons que nous avons via l'équation d'Euler-Lagrange du problème (1.1) et la définition que nous avons choisie pour le laplacien covariant

$$\Delta u = C(u) + D(\nabla u) - \frac{1}{\epsilon^2} u(1 + |u|^2) \triangleq f(u). \quad (6.2)$$

Notons que le second terme est borné puisque $|u| \leq 1$, que les symboles de Christoffel sont bornés et que donc ∇u est borné aussi.

Nous savons par hypothèse que $u \in W^{1,2}$, nous avons de plus, par le corollaire 4.3.1 que $u \in L^p \forall p \in [2, +\infty]$ et puisque $|u| \leq 1$ et que notre domaine est borné, $(1 + |u|^2) \in L^p, \forall p > 1$. Retirons maintenant des informations sur la régularité de nos solutions via l'équation (6.2). Nous pouvons nous ramener, en changeant la balance des termes de l'égalité (6.2), au cadre d'application du théorème 6.2.3 et nous savons donc que $u \in H^2(\Omega) = W^{2,2}(\Omega)$. Montrons maintenant que nous avons $D(\nabla u)$, $C(u)$ et $\frac{1}{\epsilon^2}u(1 + |u|^2) \in W^{1,2}(\Omega)$.

Nous avons d'abord

$$-D(\nabla u) = 2n\nabla n \cdot \nabla u \in L^2(\Omega)$$

par régularité de n sur une surface régulière et car $\nabla u \in L^2(\Omega)$; de plus,

$$-\nabla(D(\nabla u)) = (2\nabla^2 n + 2n\Delta n) \cdot \nabla u + 2n\nabla n \cdot \Delta u \in L^2(\Omega)$$

par régularité de n et puisque ∇u et $\Delta u \in L^2(\Omega)$. Nous avons donc bien $D(\nabla u) \in W^{1,2}(\Omega)$.

Nous avons ensuite

$$\begin{aligned} C(u) &= B(u) - n(\Delta n \cdot u) - \tilde{C}u \\ &= \left(\sum_{j=1}^3 \sum_{i=1}^3 \sum_{k=1}^3 \sum_{m=1}^3 \sum_{n=1}^3 \Gamma_{jk}^i \Gamma_{mk}^n u_k \right) - n(\Delta n \cdot u) - u \nabla \Gamma_{jk}^i \in W^{1,2} \end{aligned}$$

par régularité de la normale et des symboles de Christoffel sur une surface régulière et puisque $u \in W^{1,2}$.

Et nous avons finalement

$$\nabla \left(\frac{1}{\epsilon^2} u(1 + |u|^2) \right) = \frac{1}{\epsilon^2} \left(\nabla u(1 + |u|^2) + 2u \frac{\nabla u}{|u|} \right) \in L^2(\Omega)$$

car ∇u , $(1 + |u|^2)$, $|u|$ et $u \in L^2(\Omega)$ et donc, $\frac{1}{\epsilon^2}u(1 + |u|^2) \in W^{1,2}$.

Via le théorème 6.2.3, nous avons donc, en supposant Ω de classe C^3 ,

$$u \in H^3(\Omega) = W^{3,2}(\Omega)$$

et donc avec des raisonnements analogues à précédemment, $D(\nabla u)$, $C(u)$ et $\frac{1}{\epsilon^2}u(1 - |u|^2) \in W^{2,2}(\Omega)$ et donc, $u \in H^4(\Omega) = W^{4,2}(\Omega)$.

En répétant en boucle ce même schéma de raisonnements, nous obtenons, via le théorème 6.2.3 et en supposant Ω de classe C^p que $u \in W^{p,2} \forall p \in [1, \infty[$ et que $u \in C^2(\bar{\Omega})$.

Chapitre 7

Convergence dans le cas du tore vers une application harmonique

7.1 Introduction

Dans cette section, nous allons montrer que lorsque nous travaillons sur un tore plongé dans $\mathbb{R}^3$ ($\Omega = T^2$), les minimiseurs de notre fonctionnelle (1.1) sont des applications harmoniques. Cette convergence vers des applications harmoniques est rendue possible de par le fait que le tore T^2 , contrairement à la sphère S^2 , peut être peigné, c'est-à-dire que nous pouvons trouver un champ de vecteurs tangents partout non nuls et partout tangents à la surface du tore T^2 (voir par exemple [21]).

7.2 Convergence des minimiseurs vers des applications harmoniques dans le cas du tore

Soit $v_* \in H^1(T^2)$ telle que

$$E(v_*) = \min E(v) = \min \left\{ \frac{1}{2} \int_T |\nabla_c v|^2 \right\} \text{ tel que } v \in H^1(T^2), v \in T_x(T^2) \text{ et } |v| = 1. \quad (7.1)$$

L'équation d'Euler-Lagrange associée au problème de minimisation ci-dessus est $\Delta_c v = 0$. Nous avons donc que les minimiseurs de (7.1) sont des applications dont le laplacien covariant est nul et ce sont donc des applications harmoniques pour chaque système de coordonnées local du tore T^2 .

Nous allons montrer que lorsque $\epsilon \rightarrow 0$, les minimiseurs de (1.1) tendent vers les minimiseurs de (7.1).

7.2.1 Existence de solutions

Nous avons déjà démontré précédemment que des minimiseurs existent pour le problème

$$E_\epsilon(u) = \min_{u \in H^1(\Omega)} \int_\Omega \frac{1}{2} |\nabla_c u|^2 + \frac{1}{4\epsilon^2} (1 - |u|^2)^2,$$

ces minimiseurs étant tangents à Ω en tout point et appartenant à $H^1(\Omega)$.

Voyons maintenant pourquoi le problème

$$\kappa_* = \inf \left\{ E_*(v) = \frac{1}{2} \int_T |\nabla_c v|^2 \text{ tel que } v \in H^1(T^2, \mathbb{R}^3), v \in T_x(T^2) \text{ et } |v| = 1 \right\}$$

admet au moins un minimiseur.

Nous avons $\kappa_* \geq 0$ puisqu'il correspond à l'infimum de l'intégrale d'un carré sur un domaine non nul. De plus, pour toute suite minimisante, nous savons que les éléments de la suite minimisante sont bornés par hypothèse puisque nous travaillons avec des fonctions de norme égale à un presque partout. Or, nous savons que de toute suite bornée dans un espace de Hilbert, nous pouvons extraire une sous-suite qui converge faiblement. Nous avons donc une suite $(v_n) \rightharpoonup v_*$ avec $\kappa_* = E_*(v_*)$.

Par ailleurs, nous avons les résultats suivants :

$$\begin{aligned} \int_{\Omega} |\nabla_c v_*|^2 &= \int_{\Omega} |\nabla_c(v_* - v_n + v_n)|^2 \\ &= \int_{\Omega} |\nabla_c(v_* - v_n) + \nabla_c(v_n)|^2 \\ &= \int_{\Omega} |\nabla_c(v_* - v_n)|^2 + 2 \int_{\Omega} \nabla_c(v_* - v_n) \cdot \nabla_c(v_n) + \int_{\Omega} |\nabla_c(v_n)|^2, \end{aligned}$$

or, par convergence faible de $v_n \rightharpoonup v_*$ dans $H^1(T^2, \mathbb{R}^3)$ et par linéarité des gradients et des symboles de Christoffel,

$$\lim_{n \rightarrow \infty} \int_{\Omega} \nabla_c(v_* - v_n) \cdot \nabla_c(v_n) = 0;$$

ainsi,

$$\lim_{n \rightarrow \infty} \int_{\Omega} |\nabla_c(v_n)|^2 \leq \inf \int_{\Omega} |\nabla_c(v_*)|^2$$

et $E_*(v_n)$ est donc semi-continu inférieurement.

Notons de plus, que $|v_*| = 1$ presque partout en tant que limite de la suite v_n puisque chaque $|v_n| = 1$ presque partout et que donc

$$\int_{\Omega} \left| |v_n|^2 - 1 \right|^2 = 0,$$

et que donc

$$\lim_{n \rightarrow 0} \int_{\Omega} \left| |v_n|^2 - 1 \right|^2 = \int_{\Omega} \left| |v_*|^2 - 1 \right|^2 = 0$$

ce qui signifie bien que $|v_*| = 1$ presque partout sur Ω .

Nous avons alors

$$\lim_{n \rightarrow \infty} E_*(v_n) = E_*(u_*) \leq \inf E_*(u) \geq u_*$$

ainsi, notre suite minimisante tend bien vers un minimiseur de $E_*(v)$.

7.2.2 Convergence des minimiseurs

Notons tout d'abord que sur le tore T^2 , l'ensemble suivant n'est pas vide :

$$\Upsilon(T^2) = \{u \in H^1(T^2, \mathbb{R}^3), u(x) \in T_x T^2 \text{ tel que } |u| = 1\}.$$

De plus, Υ est un sous-ensemble de $H^1(T^2, \mathbb{R}^3)$. Dès lors, nous avons

$$\min_{u \in H^1(T^2)} \int_{T^2} \frac{1}{2} |\nabla_c u|^2 + \frac{1}{4\epsilon^2} (1 - |u|^2)^2 \leq \min_{u \in \Upsilon(T^2)} \int_{T^2} \frac{1}{2} |\nabla_c u|^2 + \frac{1}{4\epsilon^2} (1 - |u|^2)^2 = \min_{u \in \Upsilon(T^2)} \int_{T^2} \frac{1}{2} |\nabla_c u|^2.$$

En posant, $E_{\Upsilon}(v_*) = \min_{u \in \Upsilon(T^2)} \int_{T^2} \frac{1}{2} |\nabla_c u|^2$ et $E_{\epsilon}(u_{\epsilon}) = \min_{u \in H^1(T^2)} \int_{T^2} \frac{1}{2} |\nabla_c u|^2 + \frac{1}{4\epsilon^2} (1 - |u|^2)^2$, nous avons

$$E_{\epsilon}(u_{\epsilon}) \leq E_*(v_*) = E_{\epsilon}(v_*)$$

avec $|u_\epsilon| \leq 1$.

Notons que $E_*(v_*)$ et $\frac{1}{2}|\nabla_c u_\epsilon|^2$ sont bornés puisque nous travaillons dans $H^1(T^2)$ et donc, par hypothèse,

$$\int_{T^2} \frac{1}{2} |\nabla u_*|^2 < \infty \text{ et } \int_{T^2} \frac{1}{2} |\nabla u_\epsilon|^2 < \infty;$$

$A(u)$ et $A(u_*)$ sont bornés puisque u et u_* sont bornées et que les symboles de Christoffel d'une variété différentielle régulière sont bornés. Dès lors, puisque nous avons,

$$\int_{T^2} \frac{1}{2} |\nabla_c u_\epsilon|^2 + \frac{1}{4\epsilon^2} (1 - |u_\epsilon|^2)^2 \leq E_*(v_*) < \infty,$$

nous avons

$$\int_{T^2} (1 - |u_\epsilon|^2)^2 \leq 4\epsilon^2 \left(E_*(v_*) - \int_{T^2} \frac{1}{2} |\nabla_c u_\epsilon|^2 \right)$$

ce qui tend vers 0 lorsque ϵ tend vers 0. Nous avons donc

$$\lim_{\epsilon \rightarrow 0} \int_{T^2} (1 - |u_\epsilon|^2)^2 = 0$$

ce qui revient à dire que

$$\int_{T^2} (1 - |u_*|^2)^2 = \lim_{\epsilon \rightarrow 0} \int_{T^2} (1 - |u_\epsilon|^2)^2 = 0$$

et donc que $|u_*| = |u_\epsilon| = 1$ presque partout. Nous avons donc $u_\epsilon \in \Upsilon(T^2)$ et nous avons donc

$$E_*(u_*) = \frac{1}{2} |\nabla_c u_*|^2 \leq \lim_{\epsilon \rightarrow 0} \frac{1}{2} |\nabla_c u_\epsilon|^2 \leq \lim_{\epsilon \rightarrow 0} E_\epsilon(u_\epsilon) \leq E_*(v_*).$$

Or, à une extraction de sous-suite près, $u_n \rightharpoonup u_*$ et donc u_* est un minimiseur.

7.2.3 Équation d'Euler-Lagrange et caractère harmonique

Nous avons donc vu que les minimiseurs de notre fonctionnelle de Ginzburg-Landau s'avèrent être dans le cas du tore T^2 des fonctions qui minimisent l'énergie suivante

$$\frac{1}{2} \int_T |\nabla_c u|^2 \text{ avec } u \in H^1(T^2, \mathbb{R}^3) \text{ et } u(x) \in T_x T^2$$

qui peut être vue comme une énergie covariante de Dirichlet. Dès lors, nos minimiseurs sont des analogues non-linéaires d'applications harmoniques étant donnés que les minimiseurs de l'énergie de Dirichlet (sans gradient covariant) sont des applications harmoniques.

7.2.4 Convergence des gradients covariants

Notons aussi le résultat suivant puisque dans un espace de Hilbert, la convergence faible associée à la convergence en norme implique la convergence forte.

Puisque $u_n \rightharpoonup u_*$, nous avons

$$\lim_{n \rightarrow \infty} \int_\Omega |\nabla_c u_n|^2 = \int_\Omega |\nabla_c u_*|^2.$$

Or,

$$\int_\Omega |\nabla_c u_n - \nabla_c u_*|^2 = \int_\Omega |\nabla_c u_n|^2 + \int_\Omega |\nabla_c u_*|^2 - 2 \int_\Omega \langle \nabla_c u_n, \nabla_c u_* \rangle$$

et nous avons donc

$$\lim_{n \rightarrow \infty} \int_{\Omega} |\nabla_c u_n - \nabla_c u_*|^2 = E(u_*) - E(u_*) = 0$$

et nous avons donc convergence forte entre les gradients covariants (dont nous avons déjà expliqué précédemment pourquoi ils étaient bornés (le gradient et les symboles de Christoffel sont bornés)).

Chapitre 8

Etude de l'apparition de singularités dans un cas simple

8.1 Introduction

Cette partie est basée sur l'article de F.Bethuel (voir [7]) tout en donnant des raisonnements plus complets et parfois légèrement différents ainsi que quelques résultats et remarques supplémentaires.

Dans cette section, nous allons étudier l'apparition de singularités sur une surface Ω possédant un bord $\partial\Omega$ en imposant une fonction g telle que

$$g : \partial\Omega \rightarrow S^1$$

avec $\deg(g, \partial\Omega) = 1$ où $\deg$ est le degré de Brouwer (si en savoir plus sur le degré de Brouwer vous intéresse, rendez-vous en annexes A) de g et correspond dans notre cas au nombre de tours équivalents sur le cercle.

Nous verrons que notre champ de vecteurs, lorsque nous travaillons sur un domaine simplement connexe, ne peut "se répandre" de façon continue en gardant une norme égale à 1 dans tout le domaine, puisque cela impliquerait de devoir tourner autour d'un point avec une norme égale à 1 ce qui nécessiterait une accélération infinie. Ce qui est impossible. Ce sera la raison de l'apparition de ce que nous appellerons des tourbillons qui sont des points autour desquels la rotation oblige la norme de notre champ de vecteur à être égale à 0.

Nous commencerons par étudier le modèle en cours dans la section "un modèle simple". Ce qui correspondra relativement au type d'étude effectuée précédemment dans ce mémoire. Après, dans la section "apparition de tourbillons", nous étudierons l'apparition de tourbillons dans notre cas de surface à bord et nous expliquerons intuitivement ce que cela implique pour la surface sans bord qu'est la sphère.

8.2 Variété lisse

En topologie différentielle, une variété lisse est un ensemble muni d'une structure qui lui permet d'être localement difféomorphe à un espace localement convexe. Les variétés localement de dimension finie (localement difféomorphes à un espace $\mathbb{R}^n$) en sont un cas particulier.

Les variétés lisses de dimension m sont des ensembles qui admettent pour chacun de leur point un voisinage qui peut s'appliquer via une bijection C^∞ d'inverse C^∞ dans un ouvert de $\mathbb{R}^m$.

8.3 Principe du maximum

L'énoncé du théorème 8.3.1 ci-dessous est basé sur l'énoncé du même théorème fait dans le livre [2].

Théorème 8.3.1 (Principe du maximum de Hopf). *Soit $u = u(x)$, $x = (x_1, \dots, x_n)$, une fonction C^2 qui satisfait l'inégalité différentielle*

$$Lu \approx \sum_{i,j} a_{ij}(x) \partial_{x_i x_j}^2 u + \sum_i \partial_{x_i} u \geq 0$$

dans un domaine Ω . Supposons que la matrice symétrique $[a_{ij}] = [a_{ij}(x)]$ est localement uniformément définie positive dans Ω (c'est-à-dire que, pour n'importe quel sous-ensemble Ω' de Ω , la forme quadratique

$$\sum_{i,j} a_{ij}(x) \eta_i \eta_j$$

est positive et uniformément bornée pour tout $x \in \Omega'$ et pour tout vecteur $\boldsymbol{\eta} \in \mathbb{R}^n$ avec $|\boldsymbol{\eta}| = 1$), et les coefficients a_{ij} , $b_i = b_i(x)$ sont localement bornés dans Ω .

Si u prend sa valeur maximale M dans Ω , alors $u \approx M$ dans Ω .

8.4 Un modèle simple

Soit Ω un domaine régulier et borné de $\mathbb{R}^2$. Nous allons considérer des cartes complexes sur Ω , c'est-à-dire des cartes de Ω dans $\mathbb{R}^2$ et la fonctionnelle de Ginzburg-Landau

$$E_\epsilon(v) = \frac{1}{2} \int_{\Omega} |\nabla v|^2 + \frac{1}{4\epsilon^2} \int_{\Omega} (1 - |v|^2)^2 \quad (8.1)$$

où ϵ est un paramètre strictement positif et homogène à une longueur. La fonctionnelle énergétique de Ginzburg-Landau contient un potentiel non convexe $V(u) = (1 - |u|^2)^2$ dont les minimums sont $\{+1, -1\}$ dans $\mathbb{R}$ et le cercle unité S^1 dans $\mathbb{R}^2$. Les topologies de ces variétés minimales du potentiel $V(u)$ vont jouer un rôle crucial dans l'étude de la fonctionnelle de Ginzburg-Landau. En effet, ces topologies vont induire des défauts sur les surfaces dont nous cherchons des minimums de (8.1). Ces défauts seront appelés des tourbillons. En effet, pour des petits ϵ , le potentiel $V(x) = \frac{1}{4\epsilon^2} (1 - |v|^2)^2$ impose $|v|$ d'être proche de 1 et les valeurs prises par v appartiennent alors presque toutes à S^1 . Cependant, certains points de v peuvent devoir avoir une norme nulle, ce qui introduit des défauts de nature topologique que nous appellerons des tourbillons. L'étude du comportement des minimiseurs de la fonctionnelle lorsqu' ϵ tendra vers zéro sera au centre de l'analyse qui va suivre.

Puisque nous travaillons avec un domaine possédant un bord, il nous faut imposer une condition sur $\partial\Omega$, le bord de Ω . Soit g une application lisse de $\partial\Omega$ vers S^1 , nous imposons à v d'être égal à g sur $\partial\Omega$.

Nous introduisons maintenant l'espace de Sobolev suivant, associé à cette condition au bord :

$$H_g^1(\Omega; \mathbb{R}^2) = \{v \in H^1(\Omega, \mathbb{R}^2), v = g \text{ sur } \partial\Omega\}$$

dans lequel nous chercherons nos minimums de la fonctionnelle (8.1).

Si nous trouvons une fonction $u \in H_g^1(\Omega, \mathbb{R}^2)$ dans Ω telle que $E_\epsilon(u) \leq E_\epsilon(v)$ pour tout $v \in H_g^1(\Omega, \mathbb{R}^2)$, alors nous savons que

$$\forall \phi \in C_c^\infty(\Omega), E(u) \leq E(u + t\phi)$$

or,

$$\begin{aligned} E_\epsilon(u + t\phi) &= \frac{1}{2} \int_{\Omega} |\nabla(u + t\phi)|^2 + \frac{1}{4\epsilon^2} \int_{\Omega} (1 - |u + t\phi|^2)^2 \\ &= \frac{1}{2} \int_{\Omega} |t\nabla\phi + \nabla u|^2 + \frac{1}{4\epsilon^2} \int_{\Omega} (1 - |u + t\phi|^2)^2. \end{aligned}$$

Nous savons par hypothèse que lorsque $t \rightarrow 0$, $\frac{d}{dt}E_\epsilon(u + t\phi) = 0$. De plus, nous avons

$$\begin{aligned} \frac{d}{dt}E_\epsilon(u + t\phi) &= \frac{1}{2} \int_{\Omega} 2\nabla\phi(t\nabla\phi + \nabla u) - \frac{1}{4\epsilon^2} \int_{\Omega} 4\phi(u + t\phi)(1 - |u + t\phi|^2) \\ &= \int_{\Omega} \nabla\phi(t\nabla\phi + \nabla u) - \frac{1}{\epsilon^2} \int_{\Omega} \phi(u + t\phi)(1 - |u + t\phi|^2) \end{aligned}$$

et donc

$$\frac{d}{dt}|_{t=0}E_\epsilon(u + t\phi) = \int_{\Omega} \nabla\phi\nabla v - \frac{1}{\epsilon^2} \int_{\Omega} \phi u(1 - |u|^2) = 0.$$

Or, par intégration par parties et puisque $\phi \in C_c^\infty(\Omega)$, nous avons

$$\int_{\Omega} \nabla\phi\nabla v = - \int_{\Omega} \Delta v \phi$$

et donc

$$\frac{d}{dt}|_{t=0}E_\epsilon(u + t\phi) = - \int_{\Omega} \phi \Delta u - \frac{1}{\epsilon^2} \int_{\Omega} \phi v(1 + |v|^2) = 0.$$

Nous avons donc

$$\langle \phi, -\Delta u \rangle = \langle \phi, \frac{1}{\epsilon^2} u(1 + |u|^2) \rangle$$

dans $H_g^1(\Omega, \mathbb{R}^2)$ ce qui n'est vrai que si

$$-\Delta u = \frac{1}{\epsilon^2} v(1 + |v|^2)$$

dans $H_g^1(\Omega, \mathbb{R}^2)$.

De plus, sur $\partial\Omega$, même pour un u critique, la condition au bord doit rester d'application. Nous obtenons donc que les fonctions critiques de $E_\epsilon(v)$ doivent respecter le système d'équation suivant :

$$\begin{cases} -\Delta v = \frac{1}{\epsilon^2} v(1 - |v|^2) & \text{dans } \Omega \\ v = g & \text{sur } \partial\Omega. \end{cases} \quad (8.2)$$

Nous passons maintenant à une première proposition concernant les solutions de (8.2).

Proposition 8.4.1. *Toute solution de (8.2) vérifie*

$$|v| \leq 1 \text{ sur } \Omega, \quad (8.3)$$

et

$$|\nabla v| \leq \frac{C}{\epsilon} \quad (8.4)$$

où la constante C ne dépend que de Ω et g .

Démonstration. L'inégalité (8.3) est une conséquence de la règle de Leibniz pour un opérateur différentiel d'ordre 2 et du principe du maximum. Pour rappel, la règle de Leibniz pour un opérateur différentiel d'ordre 2 nous indique que

$$\Delta(fg) = (\Delta f)g + 2 \cdot (\nabla f) \cdot (\nabla g) + f(\Delta g).$$

Nous avons donc

$$\Delta(vv) = (\Delta v)v + 2\nabla v \cdot \nabla v + v\Delta v,$$

ce qui nous mène à

$$\begin{aligned} \frac{1}{2}\Delta|v|^2 &= v \cdot \Delta v + |\nabla v|^2 \\ &= \frac{1}{\epsilon^2}|v|^2(|v|^2 - 1) + |\nabla v|^2 \\ &\geq \frac{1}{\epsilon^2}|v|^2(|v|^2 - 1). \end{aligned}$$

Posons $w = |v|^2 - 1$ et notons que

$$\frac{1}{\epsilon^2}|v|^2(|v|^2 - 1) = \frac{1}{\epsilon^2}w|v|^2.$$

Nous avons donc

$$0 \geq \frac{2}{\epsilon^2}|v|^2w - \Delta|v|^2.$$

De plus, puisque le laplacien est linéaire

$$\Delta w = \Delta(|v|^2 - 1) = \Delta|v|^2 - \Delta(1) = \Delta|v|^2.$$

Nous avons donc, en posant $a(x) = \frac{2}{\epsilon^2}|v|^2$,

$$\Delta w - a(x)w \geq 0 \text{ sur } \Omega.$$

De plus, nous savons que sur le bord $\partial\Omega$, $v = g$ et donc $w = |g|^2 - 1 = 0$ car g est une application lisse de $\partial\Omega$ dans le cercle unité S^1 .

En appliquant le principe du maximum (voir 8.3.1), nous savons que si $\Delta w - a(x)w$ atteint son maximum dans Ω , alors $\Delta w - a(x)w$ est constant et vaut se maximum sur tout Ω . Or, $w = 0$ sur $\partial\Omega$ donc $w \leq 0$.

Pour 8.4, nous notons via 8.3 que

$$|\Delta u| \leq \frac{1}{\epsilon^2}$$

puisque $|\Delta u| = |\frac{1}{\epsilon^2}u(1 - |u|^2)^2|$. Nous stoppons la démonstration ici. La conclusion peut être obtenue via des estimations elliptiques similaires à celles faites dans [8] qui dépasse le cadre de ce mémoire.

□

8.4.1 Existence de minimiseurs

Nous avons l'existence de minimiseurs de (8.1) dans H_g^1 . En effet, puisque E_ϵ est positive, nous savons

$$\kappa_\epsilon = \inf\{E_\epsilon(v), v \in H_g^1 > 0\}$$

et donc, n'importe quelle suite minimisante de κ_ϵ est bornée dans H_g^1 ce qui implique, par compacité faible, que nous pouvons, à une extraction de sous-suite près, trouver une suite convergeant vers un minimiseur u_ϵ . Il ne nous reste plus qu'à montrer la semi-continuité inférieure pour la topologie faible de E_ϵ pour avoir $u_\epsilon \in H_g^1$.

Nous avons semi-continuité inférieure si $\forall \epsilon > 0, \exists n \in \mathbb{N}$ tel que $E_\epsilon(u_\epsilon) - E_\epsilon(v) \leq \epsilon$. Pour montrer que E_ϵ est semi-continue inférieure, notons les calculs ci-dessous

$$\begin{aligned}
\int_{\Omega} |\nabla u_\epsilon|^2 &= \int_{\Omega} |\nabla(u_\epsilon - u_n + u_n)|^2 \\
&= \int_{\Omega} |\nabla(u_\epsilon - u_n) + \nabla(u_n)|^2 \\
&= \int_{\Omega} |\nabla(u_\epsilon - u_n)|^2 + 2 \int_{\Omega} \nabla(u_\epsilon - u_n) \nabla(u_n) + \int_{\Omega} |\nabla(u_n)|^2, \\
\int_{\Omega} (1 - |u_\epsilon|^2)^2 &= \int_{\Omega} (1 - |(u_\epsilon - u_n) + u_n|^2)^2 \\
&= \int_{\Omega} |u_\epsilon - u_n|^4 + 4|u_\epsilon - u_n|^3|u_n| + 6|u_\epsilon - u_n|^2|u_n|^2 \\
&\quad - 2|u_\epsilon - u_n|^2 + 4|u_\epsilon - u_n||u_n|^3 - 4|u_\epsilon - u_n||u_n| + |u_n|^4 - 2|u_n|^2 + 1 \\
&= \int_{\Omega} (1 - |u_n|^2)^2 + \int_{\Omega} |u_\epsilon - u_n|^4 + 4|u_\epsilon - u_n|^3|u_n| \\
&\quad + 6|u_\epsilon - u_n|^2|u_n|^2 - 2|u_\epsilon - u_n|^2 + 4|u_\epsilon - u_n||u_n|^3 - 4|u_\epsilon - u_n||u_n|.
\end{aligned}$$

Nous pouvons donc choisir, pour tout $\epsilon > 0$, un $n \in \mathbb{N}$ tel que

$$\begin{aligned}
&\int_{\Omega} |\nabla(u_\epsilon - u_n)|^2 + 2 \int_{\Omega} \nabla(u_\epsilon - u_n) \nabla(u_n) + \int_{\Omega} |u_\epsilon - u_n|^4 + 4|u_\epsilon - u_n|^3|u_n| \\
&\quad + 6|u_\epsilon - u_n|^2|u_n|^2 - 2|u_\epsilon - u_n|^2 + 4|u_\epsilon - u_n||u_n|^3 - 4|u_\epsilon - u_n||u_n| \leq \epsilon
\end{aligned}$$

puisque tous ces termes tendent vers 0 lorsque $n \rightarrow \infty$.

8.4.2 Forme des minimiseurs pour $\epsilon \rightarrow 0$

Nous allons maintenant étudier la forme des minimiseurs lorsque ϵ tend vers 0.

Puisque dans notre problème la norme de u_ϵ est imposée comme étant égale à 1 presque partout, le degré (dans le sens du nombre de tours) d de l'application au bord va jouer un rôle crucial puisque tourner avec une vitesse finie, par prolongement continu autour d'un point central avec une vitesse non nulle est impossible. Cela impliquera pour les cas $d \neq 0$ l'apparition de tourbillons isolés lorsque $\epsilon \rightarrow 0$.

8.4.3 Analyse asymptotique

Nous allons d'abord rappeler le contexte. Nous travaillons avec des applications de Ω dans $\mathbb{R}^2$ et une énergie qui prend la forme suivante :

$$E_\epsilon(v) = \frac{1}{2} \int_{\Omega} |\nabla v|^2 + \frac{1}{4\epsilon^2} \int_{\Omega} (1 - |v|^2)^2.$$

Dans ce contexte, la variété nulle pour le potentiel V ($V = \frac{1}{4\epsilon^2} \int_{\Omega} (1 - |v|^2)^2$) de notre énergie est le cercle S^1 lorsque nous travaillons dans $\mathbb{R}^2$.

Pour les applications ayant une faible énergie et pour des petits ϵ , le potentiel V force $|v|$ à être proche de 1 et ces applications sont donc presque des applications à valeurs dans S^1 .

Nous imposons des conditions aux limites de Dirichlet via une fonction $g : \partial\Omega \rightarrow S^1$. Nous imposons donc $v = g$ sur $\partial\Omega$.

Nous cherchons des solutions dans l'espace de Sobolev suivant :

$$H_g^1(\Omega; \mathbb{R}^2) = \{v \in H^1(\Omega; \mathbb{R}^2), v = g \text{ sur } \partial\Omega\}.$$

Nous allons maintenant analyser le comportement asymptotique de minimiseurs de $E_\epsilon(v)$ sur base du lemme (lemme 8.4.1) et de la proposition (proposition 8.4.2) suivants (qui ne seront pas démontrés).

Lemme 8.4.1. *Soit Ω un domaine simplement connexe. Soit v une application dans $H^1(\Omega; S^1)$, où*

$$H_g^1(\Omega; S^1) = \{v \in H_g^1(\Omega; \mathbb{R}^2), |v| = 1\}.$$

Alors, il existe une application à valeurs réelles $\phi \in H^1(\Omega; S^1)$ telle que

$$v = \exp i\phi.$$

De plus,

$$|\nabla v| = |\nabla \phi| \text{ presque partout.} \quad (8.5)$$

Proposition 8.4.2. *Soit Ω un domaine simplement connexe. Soit l'ensemble*

$$H_g^1(\Omega; S^1) = \{v \in H_g^1(\Omega; \mathbb{R}^2), |v| = 1\}.$$

Alors, $H_g^1(\Omega; S^1)$ est non vide si et seulement si $d = 0$.

Dans la suite, nous noterons l'énergie minimale κ_ϵ ($\kappa_\epsilon \leq E_\epsilon(v)$).

Cas $d = 0$

Nous commençons par noter que κ_ϵ reste bornée indépendamment du ϵ choisi. Pour le voir, nous prenons un v_0 dans $H_g^1(\Omega; S^1)$ (qui est non vide par la proposition 8.4.2). Nous avons

$$\kappa_\epsilon \leq E_\epsilon(v_0) = \frac{1}{2} \int_{\Omega} |\nabla v_0|^2 \quad (8.6)$$

(puisque $(|v_0| - 1)^2 = 0$). Notons qu'en plus du fait que κ_ϵ reste bornée indépendamment de ϵ , les minimiseurs u_ϵ restent bornés dans H^1 par la proposition 8.4. Nous pouvons donc, par compacité faible de H^1 , pour toute séquence u_ϵ de minimiseurs ($\epsilon \rightarrow 0$), extraire une sous-suite u_{ϵ_n} , $\epsilon_n \rightarrow 0$, telle que u_{ϵ_n} converge faiblement dans H^1 vers une application u_* . De plus, grâce à l'équation (8.6), nous avons

$$\begin{aligned} \kappa_{\epsilon_n} &= \frac{1}{4\epsilon_n^2} \int_{\Omega} (1 - |u_{\epsilon_n}|^2)^2 + \frac{1}{2} \int_{\Omega} |\nabla u_{\epsilon_n}|^2 \leq \frac{1}{2} \int_{\Omega} |\nabla v_0|^2 \\ &\Leftrightarrow \frac{1}{4\epsilon_n^2} \int_{\Omega} (1 - |u_{\epsilon_n}|^2)^2 \leq \frac{1}{2} \int_{\Omega} |\nabla v_0|^2 - \frac{1}{2} \int_{\Omega} |\nabla u_{\epsilon_n}|^2 \\ &\Rightarrow \int_{\Omega} (1 - |u_{\epsilon_n}|^2)^2 \leq C\epsilon_n^2 \rightarrow 0, \end{aligned}$$

ce qui implique $|u_*| = 1$ et donc, $u_* \in H_g^1(\Omega; S^1)$. Par l'équation (8.6) et la semi-continuité inférieure de l'énergie de Dirichlet, nous avons, pour n'importe quel $v_0 \in H_g^1(\Omega; S^1)$

$$\int_{\Omega} |\nabla u_*|^2 \leq \liminf_{\epsilon_n \rightarrow 0} \kappa_{\epsilon_n} \leq \int_{\Omega} |\nabla v_0|^2,$$

ce de quoi nous déduisons que

$$u_{\epsilon_n} \rightarrow u_* \text{ fortement dans } H^1,$$

et u_* est une solution du problème de minimisation

$$\inf\left\{\int_{\Omega} |\nabla v|^2, v \in H_g^1(\Omega; S^1)\right\}. \quad (8.7)$$

Nous clôturons le cas $d \neq 0$ par une proposition et quelques remarques.

Sans surprise, si nous avons une solution de type "étoilée" où le bord ne prend que des valeurs dans S^1 et que donc notre solution n'implique pas de rotation dans le domaine Ω qui ne pourrait pas être propagée dans le centre du domaine en gardant une norme égale à 1 partout, nous n'avons pas d'apparition de singularités.

Proposition 8.4.3. *Il existe une unique solution au problème de minimisation (8.7). Cette solution u_* s'écrit*

$$u_* = \exp i\phi_*$$

avec $\phi \in H^1(\Omega; \mathbb{R})$ qui est solution de

$$\begin{aligned} \Delta\phi_* &= 0 \text{ dans } \Omega \\ \exp i\phi_* &= g \text{ sur } \partial\Omega. \end{aligned}$$

Une conséquence du fait que l'application limite u_* est unique est que la suite u_ϵ converge fortement vers u_* (voir [8]). Dans le livre [8], les estimations suivantes du taux de convergence de u_ϵ vers u_* sont établies :

$$\|u_\epsilon - u_*\|_{L^\infty(\Omega)} \leq C\epsilon^2, \|\nabla(u_\epsilon - u_*)\|_{L^\infty(\Omega)} \leq C\epsilon,$$

et pour n'importe quel sous-ensemble compact K de Ω et n'importe quel entier k ,

$$\|u_\epsilon - u_*\|_{C^k(K)} \leq C(k, K)\epsilon^2.$$

Nous allons maintenant passer au cas plus ardu correspondant à $d \neq 0$.

8.5 Apparition de tourbillons

Cas $d \neq 0$

Supposons à partir de maintenant que $d > 0$. Nous allons faire face à une nouvelle difficulté puisque dorénavant

$$\kappa_\epsilon \rightarrow +\infty, \text{ lorsque } \epsilon \rightarrow 0.$$

En effet, montrons-le par l'absurde et supposons que κ_{ϵ_n} reste borné pour certaines sous-suites ϵ_n tendant vers 0. Alors, de façon similaire au cas $d = 0$, nous pouvons affirmer que $u_{\epsilon_n} \rightarrow u_*$ dans H^1 et que $u_* \in H_g^1(\Omega; S^1)$ ce qui est en contradiction avec la proposition 8.4.2.

Le problème devient ainsi un problème de limite singulière. Puisque u_ϵ est lisse, le fait que le degré soit différent de zéro force u_ϵ à disparaître quelque part dans Ω . Les points où u_ϵ tombe à zéro joueront un rôle important : l'énergie se concentre dans leurs voisinages. Afin d'avoir un aperçu du problème, nous allons commencer par donner une borne supérieure de κ_ϵ .

Proposition 8.5.1 (Borne supérieure pour κ_ϵ). *Il existe une constante C dépendant de g et Ω telle que*

$$\kappa_\epsilon \leq \pi d |\log \epsilon| + C.$$

Démonstration. Nous allons d'abord prouver l'estimation dans un cas particulier.

- (A) Nous imposons ici que notre domaine soit un disque ($\Omega = D$) et que $g(\exp(i\theta)) = \exp(i\theta)$. Nous construisons maintenant une application v_ϵ à laquelle nous comparerons les minimiseurs u_ϵ . Nous posons pour les coordonnées polaires (r, θ) sur D :

$$v_\epsilon(r, \theta) = f_\epsilon(r) \exp(i\theta)$$

avec f_ϵ une fonction quelconque sur $[0, 1]$ vérifiant les conditions suivantes :

$$\begin{aligned} f_\epsilon(r) &= 0 && \text{pour } 0 < r < \frac{\epsilon}{2}, \\ f_\epsilon(r) &= 1 && \text{pour } \epsilon < r \leq 1, \\ \text{et } |f'_\epsilon(r)| &\leq \frac{4}{\epsilon} && \text{pour tout } 0 \leq r \leq 1. \end{aligned}$$

Nous calculons maintenant l'énergie de v_ϵ via les calculs suivants. Notons que nos domaines d'intégration d'intérêt sont représentés sur la figure suivante (8.1).

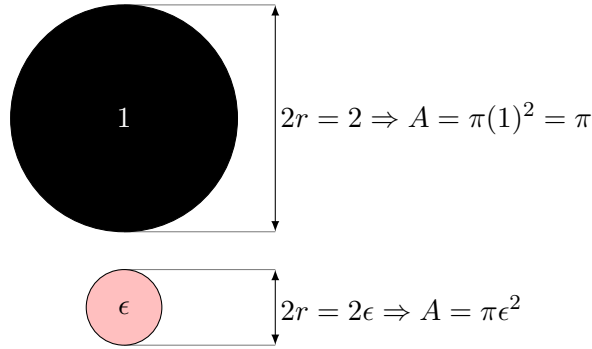


FIGURE 8.1 – Domaines d'intégration.

Nous avons

$$\frac{1}{4\epsilon^2} \int_D (1 - |v_\epsilon|^2)^2 = \frac{1}{4\epsilon^2} \int_{B(\epsilon)} (1 - |v_\epsilon|^2)^2 \leq \frac{\pi\epsilon^2}{4\epsilon^2} = \frac{\pi}{4},$$

et

$$\int_D |\nabla v_\epsilon|^2 = I + II,$$

avec, grâce notamment à la proposition 8.4,

$$\begin{aligned} I &= \int_{B(\epsilon)} |\nabla v_\epsilon|^2 \leq \frac{C}{\epsilon^2} \int_{B(\epsilon)} dx \leq \frac{C}{\epsilon^2} \pi \epsilon^2 = \pi C, \\ II &= \int_{D \setminus B(\epsilon)} |\nabla v_\epsilon|^2 = \int_{D \setminus B(\epsilon)} |\nabla(\exp i\theta)|^2 = \int_{D \setminus B(\epsilon)} \frac{1}{r^2} = 2\pi \int_\epsilon^1 \frac{1}{r^2} r dr. \end{aligned}$$

Remarquons que cette dernière égalité correspond bien à l'intuition, nous devons garder une vitesse instantanée égale à 1 en tout parcours de tout cercle et donc, plus le rayon du cercle en cours de parcours est petit et plus l'accélération nécessaire lors de son parcours sera grande. Nous avons en fait $II = 2\pi |\log \epsilon|$.

Nous concluons en constatant que

$$\kappa_\epsilon = \frac{1}{4\epsilon^2} \int_\Omega (1 - |u_\epsilon|^2)^2 + \frac{1}{2} \int_\Omega |\nabla u_\epsilon|^2 \leq \pi |\log \epsilon| + \pi C + \frac{\pi}{4} = \pi d |\log \epsilon| + C.$$

(B) Attaquons maintenant le cas général. Nous allons construire une application analogue à celle construite au point précédent. Soient d points $a_1, \dots, a_d$ dans Ω . Considérons l'application

$$w_\epsilon = \prod_{i=1}^d f_\epsilon(|z - a_i|) \frac{z - a_i}{|z - a_i|}$$

qui généralise celle du point précédent. Pour un ϵ suffisamment petit, nous avons sur le bord

$$w_\epsilon = \prod_{i=1}^d \frac{z - a_i}{|z - a_i|}$$

de manière à ce que w_ϵ soit à valeurs dans S^1 sur $\partial\Omega$ et soit de degré d . Par conséquent, il existe une application lisse $\phi : \partial\Omega \rightarrow \mathbb{R}$ telle que

$$g = \prod_{i=1}^d \frac{z - a_i}{|z - a_i|} \exp i\phi.$$

Nous étendons maintenant la définition ϕ dans Ω via l'égalité suivante

$$v_\epsilon = \prod_{i=1}^d f_\epsilon(|z - a_i|) \frac{z - a_i}{|z - a_i|} \exp i\phi.$$

Cette application est lisse, égale à g sur $\partial\Omega$ (à condition que ϵ soit suffisamment petit). Par conséquent, v_ϵ peut être utilisée comme application de comparaison. Calculer l'énergie de v_ϵ comme dans le cas A nous conduit à

$$E_\epsilon(v_\epsilon) \leq \pi d |\log \epsilon| + C.$$

□

Essayons maintenant de construire de façon non rigoureuse, via une application telle que celle construite dans la preuve de la proposition 8.5.1 que nous ne donnerons pas explicitement, une borne supérieure sur l'énergie des minimas de la fonctionnelle de Ginzburg-Landau (8.1) dans le cas d'une sphère de rayon unité.

8.5.1 Borne sur l'énergie pour une sphère

Soit une sphère S^2 de rayon 1 tel que représentée sur la figure 8.2.

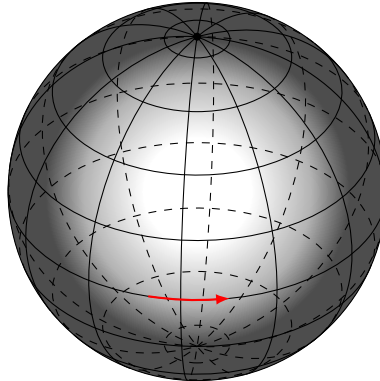


FIGURE 8.2 – Sphère avec une représentation d'un vecteur correspondant à l'application $\exp(i\theta)$ sur l'équateur.

Admettons que nous travaillons sur des cercles correspondants aux différentes latitudes. Nous construisons pour la sphère des applications du même type que précédemment. Rappelons que l'aire d'une sphère est de $\int_0^{2\pi} \int_0^\pi r^2 \sin \theta d\theta d\varphi = 4\pi r^2 = 4\pi$ et donc de 4π pour une sphère ayant un rayon égale à 1. L'aire de la sphère dont nous avons retiré un bout au pôle Sud et au pôle Nord selon des coupure parallèle aux latitudes est donnée par

$$\int_0^{2\pi} \int_\epsilon^{\pi-\epsilon} r^2 \sin \theta d\theta d\varphi = 2\pi(-\cos(\pi-\epsilon) + \cos(\epsilon)) = 4\pi \cos(\epsilon).$$

Comme précédemment, borner la première intégrale revient à borner l'aire des demies sphères dont nous avons mis la norme à 0. Or, puisque nous travaillons avec des petits ϵ , nous avons

$$\begin{aligned} \int_{S^2(\epsilon)} (1 - |v_\epsilon|^2)^2 &= 4\pi(1 - \cos(\epsilon)) \\ &\leq 4\pi(1 - \cos^2(\epsilon)) \\ &= 4\pi \sin^2(\epsilon) \\ &\leq 4\pi\epsilon^2 \end{aligned}$$

et nous avons donc

$$\frac{1}{4\epsilon^2} \int_{S^2(\epsilon)} (1 - |v_\epsilon|^2)^2 \leq \pi.$$

Pour ce qui est des champs de vitesses, puisque nous les avons supposer parallèles aux latitudes et que nous pouvons reconstruire deux disques avec les cercles ainsi construits sur la sphère, nous avons (comme précédemment notre accélération est inversement proportionnelle à notre rayon)

$$\begin{aligned} \int_{S^2(\epsilon)} |\nabla v_\epsilon|^2 &\leq 2 \frac{C}{\epsilon^2} \int_{B(\epsilon)} dx \leq \frac{2C}{\epsilon^2} \pi \epsilon^2 = 2\pi C, \\ \int_{S^2 \setminus S^2(\epsilon)} |\nabla v_\epsilon|^2 &= 2 \int_{D \setminus B(\epsilon)} \frac{1}{r^2} = 4\pi \int_\epsilon^1 \frac{1}{r} dr = 4\pi |\log \epsilon|. \end{aligned}$$

Ainsi, nous avons

$$\kappa_\epsilon = \frac{1}{4\epsilon^2} \int_\Omega (1 - |u_\epsilon|^2)^2 + \frac{1}{2} \int_\Omega |\nabla u_\epsilon|^2 \leq 2\pi |\log \epsilon| + 2\pi C + \pi.$$

8.5.2 Résultats exacts

Il s'avère que l'application de comparaison v_ϵ construite dans la preuve de la proposition 8.5.1 est assez similaire aux vrais minimas u_ϵ (bien sûr, pour un choix approprié des points a_i) et que la limite de la proposition 8.5.1 est forte. Plus précisément dans [9], le théorème suivant est prouvé (initialement pour des domaines Ω étoilé mais cette restriction a été levée par M.Struwe dans [24]). Nous énoncerons donc ce théorème en toute généralité bien que ce ne soit pas exactement ce qui est fait dans [9].

Théorème 8.5.1. *Il existe une constante $C > 0$ dépendant uniquement de g telle que*

$$|\kappa_\epsilon - \pi d| \log \epsilon \leq C, \text{ pour } 0 < \epsilon < 1.$$

- L'application u_ϵ possède exactement d zéros pour un ϵ suffisamment petit.
- Il y a exactement d points $a_1, \dots, a_d$ dans Ω tels que, à une sous-suite ϵ_n tendant vers 0 près,

$$u_{\epsilon_n} \rightarrow u_* \text{ sur n'importe quel sous-ensemble compact de } \Omega \setminus \bigcup_{i=1}^d \{a_i\},$$

où u_* est l'application à valeurs dans S^1 , égale à g sur $\partial\Omega$, donnée par

$$u_* = \prod_{i=1}^d \frac{z - a_i}{|z - a_i|} \exp i\phi,$$

où ϕ est une fonction harmonique.

- La configuration $(a_1, \dots, a_d)$ n'est pas arbitraire mais minimise sur $\Omega^d \setminus \Delta$ (où Δ indique la diagonale, c'est-à-dire $\Delta = \{(x_1, \dots, x_d) \in \Omega^d, \text{ tel que } \exists i \neq j : x_i = x_j\}$) une énergie renormalisée, qui a (grossièrement) la forme suivante

$$W_g(a_1, \dots, a_d) = -\pi \sum_{i \neq j} \log |a_i - a_j| + \text{des contributions limites.}$$

- Nous avons le développement suivant

$$\kappa_\epsilon = \pi d |\log \epsilon| + W_g(a_1, \dots, a_d) + d\gamma_0 + o(1)$$

où $o(1) \rightarrow 0$ quand $\epsilon \rightarrow 0$, et où γ_0 est une constante universelle.

Notons que pour des minimiseurs, le degré de chaque singularité a_i est $+1$ ce qui n'est pas le cas pour des solutions qui ne minimisent pas $E_\epsilon(v)$.

Chapitre 9

Formulation d'un champ de croix

Pour former un champ de croix, une première idée pourrait être de travailler avec des quotients de groupes, notamment en quotientant notre espace d'intérêt par le groupe du carré. Toutefois, comme nous l'avons vu précédemment, la régularité des domaines sur lesquels nous travaillons est primordiale pour établir nos résultats. Dès lors, effectuer l'équivalent de découpage et de collage sur nos espaces fonctionnels n'est probablement pas une bonne idée. Nous essayerons donc de définir des champs de croix sur un espace $H^1(\Omega, \mathbb{R}^3)$ via la définition de polynômes. Nous verrons alors quelle forme nous pourrions donner à ces polynômes.

Nous commencerons ce chapitre par détailler quelques éléments qui seront utiles dans nos développements. Ensuite, nous essayerons de construire par un raisonnement logique des polynômes pouvant être associés à des champs de croix. Nous terminerons en détaillant une construction de champ de croix basée sur un article en pré-publication (voir [4]).

9.1 Polynômes homogènes

Lorsque tous les termes d'un polynôme sont du même degré, nous disons que le polynôme est homogène. Le degré d'un polynôme homogène est le même que les degrés des termes qui le composent.

Notons la propriété évidente suivante. Pour un polynôme homogène de degré d , nous avons

$$P(\lambda x) = \lambda^d P(x).$$

9.2 Polynômes harmoniques

Definition 9.2.1. *Nous appelons une fonction harmonique sur $\mathbb{R}^3$ toute fonction f de classe C^2 telle que*

$$\Delta f = 0.$$

Pour tout entier non négatif l nous désignons par $P^{(l)}$ l'espace vectoriel des polynômes homogènes de degré l à coefficients complexes sur $\mathbb{R}^3$. Nous considérons ensuite le sous-espace vectoriel de $P^{(l)}$ constitué de polynômes harmoniques, c'est-à-dire de polynômes ayant un Laplacien nul, que nous notons $H^{(l)}$.

Il y aurait moyen d'en dire beaucoup plus sur les polynômes harmoniques, pour cela consulter par exemple le livre de Yvette Kosmann-Schwarzbach [20].

9.3 Définition et propriétés d'un champ de croix

Soit Ω une variété lisse de dimension 2 et plongée dans $\mathbb{R}^3$.

Nous cherchons à construire via l'utilisation d'un polynôme P un champ de croix qui soit tangent en tout point à Ω . Tentons de définir un tel polynôme. Soit $n(x) : \Omega \rightarrow \mathbb{R}^3$ la normale à Ω en tout point de Ω . Nous voulons avoir la propriété suivant :

$$DP(w(x)) \cdot n(x) = DP(w(x))[n(x)] = 0.$$

Ceci revient à dire que nous voulons que notre polynôme P n'ait, en tout point de Ω aucune composante parallèle à $n(x)$ et nous voulons donc que

$$P(z + \lambda n(x)) = P(z).$$

Pour ce faire, le plus simple est de prendre des polynômes appartenant à $\mathbb{R}^2$, $P(z(x)) = (z_1(x), z_2(x), 0)$ et $n(x) = (0, 0, n)$ où les composantes sont définies selon un repère local (duquel nous alignons un axe avec la normale).

L'autre propriété essentielle que nous voulons voir être satisfaite par P est qu'une rotation d'un quart de tour nous donne le même P et que nous ne puissions pas ainsi distinguer deux éléments différents d'une rotation d'un quart de tour, nous voulons donc que

$$P(\exp(i\theta)z) = P(z) \quad \forall \theta \in k\frac{\pi}{2} \text{ avec } k \in \mathbb{Z}.$$

Notons aussi que nous ne voulons aucune symétrie supplémentaire. Des polynômes tels que

$$\begin{aligned} P(iz) &= i^4 P(z) = P(z), \\ P(-1z) &= (-1)^4 P(z) = P(z), \\ P(-iz) &= (-i)^4 P(z) = P(z), & P(z) &= 1^4 P(z) = P(z), \end{aligned}$$

sans autre symétrie, sont les polynômes harmoniques de degré 4. En effet, ceux-ci ont la propriété suivante

$$P(\lambda z) = \lambda^4 P(z)$$

et nous avons donc bien invariance par les 4 rotations qui nous intéressent.

Admettons maintenant que nous choissions ce polynôme comme étant harmonique et réel. Ces deux contraintes supplémentaires n'ont pas d'autres raisons d'être que de restreindre la famille des polynômes appartenant à $\mathbb{R}^2$ que nous considérons et de prendre une restriction dans laquelle la forme des polynômes considérés est a priori relativement avenante.

Un exemple de tel polynôme est donnée par

$$p(x) = \Re(p(z)) = \Re(\alpha(z_1 + iz_2)^4).$$

9.4 Exemple de définition de champ de croix

Voici une construction de champ de croix telle qu'effectuer dans un article en pré-publication (voir [4]).

Un champ de croix définit quatre directions orthogonales préférées à chaque point d'une variété de dimension 2. Même si ces directions peuvent être construites sur base d'un seul angle θ mesuré par rapport à une direction relative quelconque, cela ne contient pas toute l'information

que nous voulons faire passer. En effet, des directions avec des angles $\theta + k\frac{\pi}{2}$, $k \in \mathbb{Z}$ représentent une même croix. Une croix peut être représentée par un champ de deux composants réels :

$$(\cos(4\theta); \sin(4\theta)).$$

Cette définition est unique puisque

$$\cos(4[\theta + k\frac{\pi}{2}]), \forall k \in \mathbb{Z}.$$

De plus, il est aisé de définir une distance permettant de comparer deux croix :

$$\begin{aligned} D((\cos(4\theta_i); \sin(4\theta_i)), (\cos(4\theta_j); \sin(4\theta_j))) &= \int_0^{2\pi} |\cos(4[\theta_i + \alpha]) - \cos(4[\theta_j + \alpha])|^2 d\alpha \\ &= \pi((\cos(4\theta_i) - \cos(4\theta_j))^2 + (\sin(4\theta_i) - \sin(4\theta_j))^2) \end{aligned}$$

En fait, l'équation (9.4) peut être exprimée comme une seule fonction complexe $f = u + iv$, en posant $u = \cos(4\theta)$ et $v = \sin(4\theta)$. Cela correspond à la définition de l'exponentielle complexe avec un argument égale à 4θ , nous pouvons donc poser la définition suivante :

$$(\cos(4\theta); \sin(4\theta)) \equiv u + iv = \exp(4\theta) = f.$$

De facto, la relation entre les champs de croix et les fonctions complexes est profonde et riche. En effet, un champ de croix bidimensionnel peut être exprimé comme une fonction complexe $f(z) : \mathbb{C} \rightarrow \mathbb{C}$. Le domaine correspond à des points sur la variété de dimension 2, tandis que le co-domaine correspond aux composantes des vecteurs définis à ces points.

Soit $I(z^c)$ l'indice d'un point critique correspondant à un élément de norme nul du champ de croix. Cet indice nous donne le nombre de tours que le champ de croix effectue autour de ce point. Soit χ la caractéristique d'Euler-Poincaré d'une surface fermée. Nous avons le théorème suivant.

Théorème 9.4.1 (Poincaré-Hopf). *Si un champ de croix sur une variété de dimension 2 de caractéristique d'Euler-Poincaré χ a un nombre fini n de point critique z_j^c , alors*

$$\sum_{j=1}^n I(z_j^c) = \chi$$

Ce qu'il faut retenir jusqu'ici c'est qu'un champ de croix peut être représenté par une fonction complexe f qui représente un champ de vecteurs bidimensionnels et qu'un champ de croix peut être indéfini sur certains points (sa norme y sera alors nulle). Définissons donc une norme pour les croix d'un champ de croix.

9.4.1 Norme d'un champ de croix

Soit un champ de croix dont la norme n'est pas forcément égale à 1 :

$$f := r^4 \exp(i4\theta).$$

Nous aurons alors $r = |f|^{\frac{1}{4}}$ qui sera la norme de la croix au point considéré du champ f .

9.5 Conclusion

Il serait intéressant, bien que la plupart des résultats seraient alors probablement analogues à ceux développés précédemment, d'effectuer l'étude faites lors des chapitres précédents en travaillant avec des espaces de polynômes tels que nous les avons décrits dans cette section.

Chapitre 10

Passage numérique en régime asymptotique

10.1 Introduction

Cette section donnera une analyse d'une méthode numérique destinée à produire des maillages quadrangulaires sur base de l'équation de Ginzburg-Landau suivante

$$E_\epsilon(v) = \frac{1}{2} \int_{\Omega} |\nabla v|^2 + \frac{1}{4\epsilon^2} \int_{\Omega} (1 - |v|^2)^2.$$

Cette méthode est en cours de développement dans le cadre du projet ERC hextreme et dans le cadre du projet ARC Waves dans un ensemble de modules référencés par les lettres "hxt". Le code numérique dont nous parlerons dans cette section fera référence à une partie du code hxt.

Nous allons voir que bien que mathématiquement le passage en régime asymptotique (à la solution correspondant à $\epsilon \rightarrow 0$) est régi par le paramètre ϵ , ce n'est pas toujours le cas dans le problème numérique discrétisé.

Les expériences décrites ci-après ont été menées sur des anneaux de cercle interne R_0 et de cercle externe R_1 . Pour rappel, une couronne ou un anneau est la partie du plan délimitée par deux cercles concentriques. Si R_0 et R_1 sont deux nombres réels positifs, nous noterons

$$A(z_0; R_0, R_1) = \{z \in \mathbb{C} | R_0 < \|z - z_0\| < R_1\}.$$

Notons que R_0 peut être nul (cas du disque privé de son centre) et que R_1 peut être infini.

Un des éléments essentiels justifiant le choix du disque à un trou comme surface d'étude dans le cadre de ce mémoire est que le disque à trou est un élément d'un espace euclidien. Nous avons alors, puisque nous travaillons dans un espace euclidien, équivalence entre le gradient covariant et le gradient et donc, nous pouvons supposer que la même étude effectuée sur base d'une méthode numérique se basant sur l'équation (1.1) nous donnerait les mêmes résultats. En effet, puisque la normale sur un disque à trou reste la même en tout point, nous n'avons pas de variations de notre système de coordonnées dues à la variation de la normale (variation qui, qui quand elle est présente, est responsable de la différence entre gradient classique et gradient covariant). Ainsi, nous avons égalité entre l'équation (10.2) et l'équation (1.1) lorsque nous travaillons sur des disques à un trou.

Nous verrons que le rapport des rayons des anneaux que nous utiliserons dans nos expériences sera le paramètre qui influencera le plus le passage ou non en régime asymptotique des solutions numériques.

Notons finalement la remarque suivante. D'ordinaire, pour effectuer des maillages, les méthodes reposent sur l'utilisation de points déjà placés. La méthode analysée ici consiste à construire un champ de croix qui permettra par la suite de placer des points en vue de les mailler de façon idéale.

L'étude statistique qui va vous être présentée a été effectuée via l'outil JMP qui est un logiciel statistique développé par SAS. Les figures présentant des éléments statistiques ont donc été produites via ce logiciel. Les données ont été obtenues via l'utilisation du code hxt et via l'utilisation de scripts bash dont un exemple est donné en annexe (voir B.3).

Ci-dessous un exemple de champ de croix obtenu numériquement (voir 10.1).

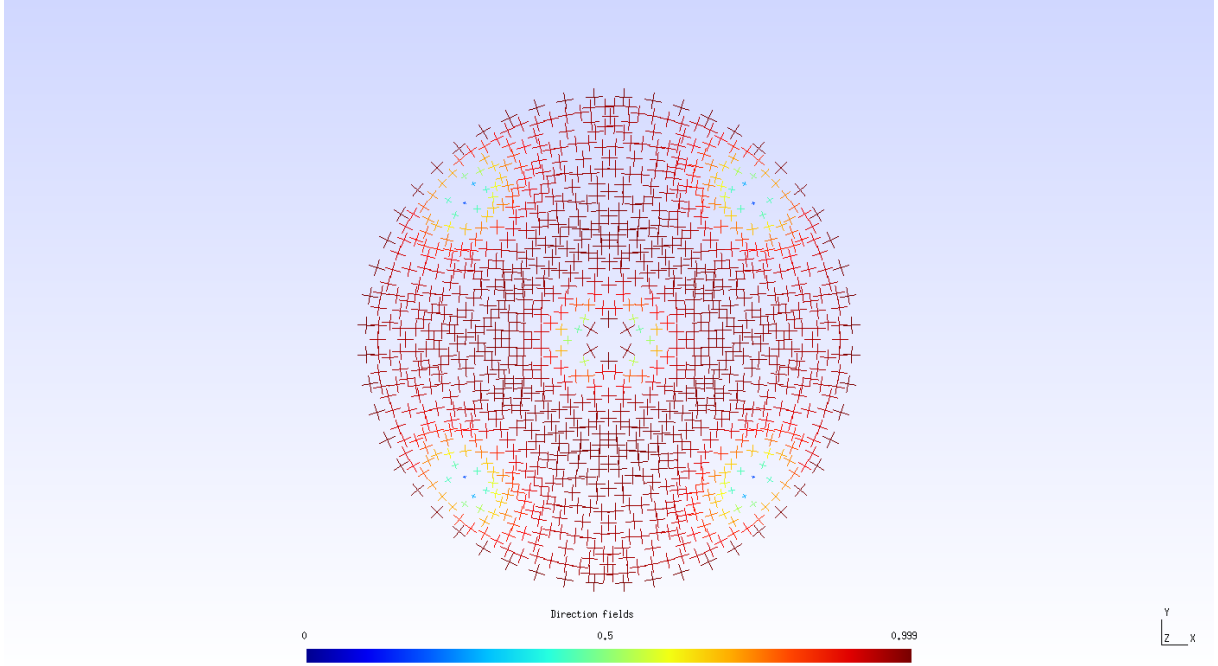


FIGURE 10.1 – Exemple de champ de croix retourné par le code hxt.

10.1.1 Rappels concernant les éléments finis

La description qui suit est principalement basée sur les deux ouvrages suivants [25] et [3].

La méthode des éléments finis est une méthode numérique destinée à résoudre des équations aux dérivées partielles. Contrairement aux schémas de différences finies qui cherchent à approcher la solution analytique inconnue d'une équation différentielle en un nombre fini de points sélectionnés (les points de maillage), la méthode des éléments finis fournit une approximation de la solution analytique via une fonction polynomiale par morceaux, définie sur l'ensemble du domaine de calcul. Cela permet d'avoir un meilleur contrôle de l'erreur globale due à nos approximations numériques. Il s'agit de mettre en place, à l'aide des principes hérités de la formulation variationnelle ou formulation faible, un algorithme discret mathématique permettant de rechercher une solution approchée d'une équation aux dérivées partielles sur un domaine compact avec, ou non, des conditions aux bords. Par exemple, dans le cadre de l'équation utilisée par hxt, nous avons l'équation d'Euler-Lagrange suivante

$$\begin{cases} -\Delta v = \frac{1}{\epsilon^2} v(1 - |v|^2) & \text{dans } \Omega \\ v = g & \text{sur } \partial\Omega \end{cases}$$

dont la formulation faible est donnée par

$$-\int_{\Omega_{\text{fractionné}}} \phi \Delta v - \frac{1}{\epsilon^2} \int_{\Omega_{\text{fractionné}}} \phi v (1 + |v|^2) = 0.$$

où ϕ est une fonction qui doit être prise dans l'espace de fonctions tests choisi et où les intégrales sont faites séparément sur chaque élément fini.

Il s'agit donc avant tout de la résolution approchée d'un problème où, grâce à la formulation variationnelle, les solutions du problème vérifient des conditions d'existence plus faibles que celles des solutions du problème de départ et où une discrétisation permet de trouver une solution approchée.

La discrétisation consiste à « découper » le domaine Ω , c'est-à-dire à chercher une solution du problème sur un domaine polygonal ou polyédrique par morceaux. Il y a donc une redéfinition de la géométrie. Une fois la géométrie approchée, il faut choisir un espace d'approximation de la solution du problème qui est défini à l'aide du maillage du domaine (ce qui explique aussi pourquoi il est nécessaire d'approcher la géométrie). Le maillage du domaine permet d'en définir un pavage dont les pavés sont les éléments finis.

Sur chacun des éléments finis, il est possible de linéariser l'EDP, c'est-à-dire de remplacer l'équation aux dérivées partielles par un système d'équations linéaires, par approximation. Ce système d'équations linéaires peut se décrire par une matrice. Il y a donc une matrice par élément fini. Cependant, les conditions aux frontières sont définies sur les frontières du système global et pas sur les frontières de chaque élément fini ; il est donc impossible de résoudre indépendamment chaque système. Les matrices sont donc réunies au sein d'une matrice globale. Le système d'équations linéaires global est résolu numériquement.

L'EDP est résolue aux nœuds du maillage, c'est-à-dire que la solution est calculée en des points donnés (résolution discrète) et non en chaque point du domaine Ω . Cela nécessite de pouvoir interpoler, c'est-à-dire déterminer les valeurs en tout point à partir des valeurs connues en certains points.

Un élément fini est la donnée d'une cellule élémentaire et de fonctions de base de l'espace d'approximation dont le support est l'élément, et définies de manière à être interpolantes.

La géométrie idéale et continue est remplacée par une géométrie discrète, et les valeurs sont interpolées entre des points ; plus les points sont espacés, plus la fonction d'interpolation risque de s'écarter de la réalité, mais à l'inverse, un maillage trop fin conduit à des temps de calculs extrêmement longs et nécessite des ressources informatiques (en particulier mémoire vive) importante, il faut donc trouver un compromis entre coût du calcul et précision des résultats.

Dans le cas d'une EDP linéaire avec opérateur symétrique (comme l'est l'opérateur laplacien), il s'agit finalement de résoudre une équation algébrique linéaire, inversible dans le meilleur des cas (comme l'équation d'onde par exemple).

Concrètement, cela permet par exemple de calculer numériquement le comportement d'objets même très complexes, à condition qu'ils soient continus et décrits par une équation aux dérivées partielles linéaire ou linéarisée.

La méthode des éléments finis est utilisée, dans le code hxt, sur des éléments triangulaires (obtenu via GMSH dont nous parlerons juste après), dont chaque sommet va être associé à une croix, dans le but d'utiliser l'ensemble des croix pour créer un maillage quadrangulaire régulier.

10.1.2 Description rapide de l'outil de maillage de GMSH

Les informations qui vont suivre sont toutes disponibles dans le manuel de l'outil GMSH (voir [14]).

Un maillage d'éléments finis est un pavage en mosaïque d'un sous-ensemble donné de l'espace tridimensionnel par des éléments géométriques élémentaires de formes diverses (dans le cas de Gmsh : lignes, triangles, quadrangles, tétraèdres, prismes, hexaèdres et pyramides), disposés de telle sorte que si deux d'entre eux se croisent, ils le font le long d'une face, d'un bord ou d'un nœud, et jamais autrement. Les éléments géométriques élémentaires ne sont définis que par une liste ordonnée de leurs nœuds mais aucune relation d'ordre prédéfinie n'est supposée entre deux éléments quelconques.

La génération de maillage est effectuée dans un flux ascendant : les lignes du bord sont discrétisées en premier ; le maillage des lignes est ensuite utilisé pour engrener les surfaces et le maillage des surfaces est alors utilisé pour mailler les volumes. Dans ce processus, le maillage d'une entité est seulement contraint par le maillage de sa frontière. Par exemple, en trois dimensions, les triangles discrétisant une surface seront forcés à être des faces de tétraèdres dans le maillage 3D final seulement si la surface fait partie de la limite d'un volume. Les éléments de ligne discrétisant une courbe seront forcés à être des arêtes de tétraèdres dans le maillage 3D final seulement si la courbe fait partie de la limite d'une surface, elle-même faisant partie de la limite d'un volume. Un seul nœud discrétisant un point au milieu d'un volume sera obligé d'être un sommet de l'un des tétraèdres dans le maillage 3D final seulement si ce point est connecté à une courbe, elle-même partie de la frontière d'une surface, elle-même partie de la limite d'un volume. Ceci assure automatiquement la conformité du maillage lorsque, par exemple, deux surfaces partagent une ligne commune. Mais cela implique aussi que la discrétisation d'une entité "isolée" de dimension $(n-1)$ à l'intérieur d'une entité de n -ième dimension ne contraint pas le maillage de n -ième dimension. Chaque étape de maillage est contrainte par un "champ de taille" (parfois appelé "champ de longueur caractéristique"), qui prescrit la taille souhaitée des éléments dans le maillage. Ce champ de taille peut être uniforme, spécifié par des valeurs associées à des points de la géométrie, ou défini par des «champs» généraux (par exemple liés à la distance à une frontière, à un champ scalaire arbitraire défini sur un autre maillage, etc.).

10.2 Description du solveur numérique

Le code numérique permettant de construire un champ de croix se voulant relativement régulier est basé sur l'équation suivante (légèrement différente de l'équation (1.1)).

$$E_\epsilon(u) = \min_{u \in H^1(\Omega)} \int_{\Omega} \frac{1}{2} |\nabla u|^2 + \frac{1}{4\epsilon^2} (1 - |u|^2)^2$$

Il part d'un maillage en triangles obtenu via GMSH (voir [14]). Les conditions aux bords imposées au champ de croix sont qu'il doit se conformer aux segments limites du maillage obtenu via GMSH. Pour chaque arête du maillage obtenu via GMSH, une orientation locale est attribuée et celle-ci est cohérente avec celle des éléments lui étant voisins (par exemple, toutes les normales aux bords de la surface sont sortantes ou toutes entrantes). Pour former les croix, les quatre racines quatrièmes de $\exp(i\gamma)$ sont calculées avec γ qui est forcée de s'aligner avec le bord et donc avec le premier axe du repère local qui correspond avec la frontière sur le bord du domaine. Dans le repère local associé au point considéré, la condition au bord s'exprime comme suit $\exp(i4\gamma) = \cos(4\gamma) + i \sin(4\gamma) = 1 + 0i = \exp(i0)$.

La norme d'une branche est donnée par $|\sqrt[4]{\exp(i4\gamma)}|$. Pour ce qui est de la résolution numérique qui permet d'associer une norme à chaque branche de chaque croix, elle se fait par méthode d'éléments finis. Nous prenons la formulation faible de l'équation d'Euler-Lagrange de notre système

$$-\int_{\Omega_{\text{fractionné}}} \phi \Delta u - \frac{1}{\epsilon^2} \int_{\Omega_{\text{fractionné}}} \phi u (1 + |u|^2) = 0$$

telle qu'elle a été décrite précédemment.

La méthode utilisée est de type Galerkin et la fonction test qui multipliera notre équation d'Euler-Lagrange est donc prise dans l'espace de nos fonctions d'essai à savoir l'espace fonctionnel auquel doit appartenir notre approximation de u sur chaque élément fini. Notons que les points où les approximations sont faites sont les bords des triangles du maillage obtenu par GMSH. Notons aussi que les méthodes d'éléments finis ne permettent pas de résoudre des équations différentielles non linéaires. La partie non linéaire de notre équation d'Euler-Lagrange ($\frac{1}{\epsilon^2}u(1 - |u|^2)$) sera donc approximé via l'algorithme de Newton-Raphson. Nous espérons alors que la solution du problème approché sera une solution approchée du problème de base. L'algorithme aura considéré qu'il a convergé lorsque $-\int_{\Omega_{\text{fractionné}}} \phi \Delta u - \frac{1}{\epsilon^2} \int_{\Omega_{\text{fractionné}}} \phi u(1 + |u|^2) = 0$ sera inférieure à une certaine tolérance (de l'ordre de 10^{-12}).

Newton-Raphson se base sur un itéré initial qui est obtenu en alignant chaque croix avec un bord d'un triangle du maillage triangulaire obtenu via GMSH (le bord extérieur pour les points sur le bord).

Une des principales limitations de cette méthode est que la vérification de l'équation d'Euler-Lagrange ne nous indique que le fait que nous avons atteint un point d'équilibre et non forcément un minimum. De plus, l'incapacité de faire descendre numériquement ϵ en dessous de la taille caractéristique de maille dans certains cas impliquera que nous n'atteignons que rarement des minimas asymptotiques (des minimas correspondant aux minima de l'équation (10.2) avec $\epsilon \rightarrow 0$).

Les croix sont comparées les unes les autres via leur puissance 4 ce qui permet de considérer tous les élément équivalents à une symétrie du carré près comme étant égaux.

10.2.1 Exemple d'application possible

L'outil final qui servira à remailler des points de façon régulière avec des quadrangles pourrait par exemple être utilisé dans l'étude de nos sols. En effet, lorsqu'une étude des sols profonds est effectuée avec des ondes sonores, les ondes réfléchies nous permettent d'avoir des informations (mais qui sont très bruitées et qui ne présentent a priori aucune régularité) concernant des points des interfaces entre les différentes couches du sol. Il est dès lors très compliqué d'utiliser directement les points de réflexions acoustiquement estimés sur les interfaces entre les couches du sol pour effectuer des calculs de type éléments finis. Il faut donc les retravailler pour que nous puissions construire des maillages quadrangulaires avec ceux-ci. Ces mesures retravaillées sont susceptibles de permettre, en utilisant l'équation d'onde dans le sens inverse du temps, de faire correspondre ces points et les caractéristiques mesurées à leur proximité à des propriétés des différentes couches du sol.

10.3 Expériences

Lors de l'utilisation du code hxt, il a été constaté que le régime asymptotique pouvait être atteint (ou parfois pas du tout) avec des valeurs très variables de ϵ selon les surfaces étudiées et ce même pour des problèmes proches.

C'est dans ce contexte, que deux expériences avec des disques à trous et avec des rapports de rayons variables ont été menées. L'étude de leurs résultats a permis de mettre en avant l'influence prédominante du rapport des rayons sur ϵ (du moins pour ceux pris en compte) pour prédire le passage ou non en régime asymptotique. Le disque à trou étant de caractéristique d'Euler-Poincaré nulle, nous nous attendons dans la cas d'une solution ayant convergé vers un minimum asymptotique de l'énergie de Ginzburg-Landau à n'avoir aucune singularité (croix dont la norme tend vers zéro) sur la surface.

Les paramètres utilisés seront la taille relative des mailles par rapport au domaine (h relatif), ϵ , le rapport des rayons. L'énergie des solutions sera aussi prise en compte dans les analyses qui vont suivre.

Les résultats seront donnés de façon à respecter les étapes effectuées lors de l'investigation des différents ensembles de données traités.

10.3.1 Première expérience

Une première expérience consistait à faire varier via une triple boucle for, le h relatif du maillage dans l'ensemble suivant : $\{0, 1; 0, 2; 0, 3; 0, 4; 0, 5; 0, 6; 0, 7; 0, 8; 0, 9\}$, le rapport des rayons dans l'ensemble : $\{0, 009; 0, 09; 0, 1; 0, 2; 0, 3; 0, 4; 0, 5; 0, 6; 0, 7; 0, 8; 0, 9\}$ et ϵ dans $\{0, 00001; 0, 0001; 0, 001; 0, 01; 0, 1; 1\}$. Commençons par analyser l'énergie des données obtenues via les simulations avec les paramètres indiqués. Cela nous permettra de nous débarrasser des données qui correspondent à une énergie trop élevée que pour que l'algorithme ait convergé (l'algorithme aura alors atteint son nombre maximal d'itération à la place).

Nous observons sur les données brutes que quand ϵ et le h relatif se retrouvent être d'ordres trop différents, l'énergie de la solution du problème de Ginzburg-Landau explose (le premier graphe qui a permis de mettre cela en évidence est donné en annexe B.1).

Il semble à ce stades que certaines des énergies obtenues ne correspondent pas à des minima de schémas de Ginzburg-Landau mais qu'ils correspondent plutôt à des énergies obtenues après non convergence de la résolution numérique de la fonctionnelle énergétique de Ginzburg-Landau sur les disques à trous correspondant et donc à l'atteinte du nombre maximal d'itérations sans convergence. Pour corroborer cette intuition, des graphes de l'énergie en fonction du logarithme de ϵ ont été produits. En voici quatre exemples (voir 10.2).

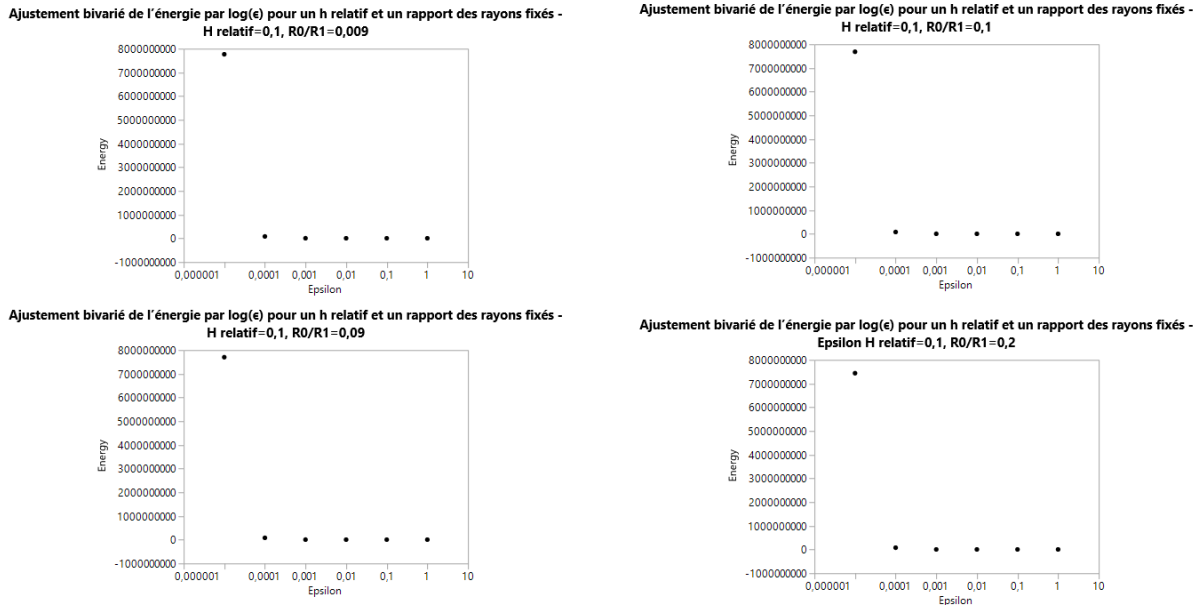


FIGURE 10.2 – Graphes des énergies en fonctions du logarithme de ϵ pour différentes paires de h relatif et de rapport des rayons.

Comme nous l'a indiqué le théorème 8.5.1 vu précédemment, nous nous attendons à avoir

une énergie qui évolue de façon linéaire et inversement proportionnellement avec le log de ϵ or ce n'est pas ce que nous observons. Sur base de ces graphes (voir 10.2), nous allons donc considérer que certaines énergies obtenues ne correspondent pas à des minima de notre fonctionnelle de Ginzburg-Landau. Les données associées à une énergie supérieure ou égale à 1000 ont donc été écartées. Ce choix est subjectif toutefois, dans le cadre d'une analyse exploratoire, nous considérons que quelques choix partiels peuvent être effectués sans nuire au propos qui se veut plus qualitatif que quantitativement précis. Il a aussi été décidé de ne garder que les données pour lesquelles le nombre de singularités est pair (car en cas de présence de singularités, la somme de leurs indices (qui valent soit 1 soit -1) doit être nulle).

Analysons maintenant nos données filtrées. Tout d'abord, les ϵ ne prennent plus que les valeurs 0,1 et 1. Ensuite, le nombre de singularités ne dépasse plus 12. L'étude de la distribution de l'énergie nous donne les résultats suivants.

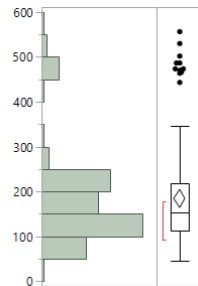


FIGURE 10.3 – Distribution de l'énergie parmi les données filtrées.

Nous constatons que nous n'atteignons plus d'énergies affreusement élevées. Les données des quantiles de l'énergie sous forme de tableau ainsi que des statistiques la décrivant (avec notamment un intervalle de confiance pour sa moyenne) sont données en annexes (voir B.1 et B.2).

Cette première analyse nous a permis de nous débarrasser de données inintéressantes mais ne nous apprend pas grand chose. Nous allons donc maintenant effectuer un modèle linéaire pour essayer de dégager un candidat principal comme paramètre critique dans l'apparition de singularités.

L'utilisation d'un modèle linéaire, nous donne les résultats suivants.

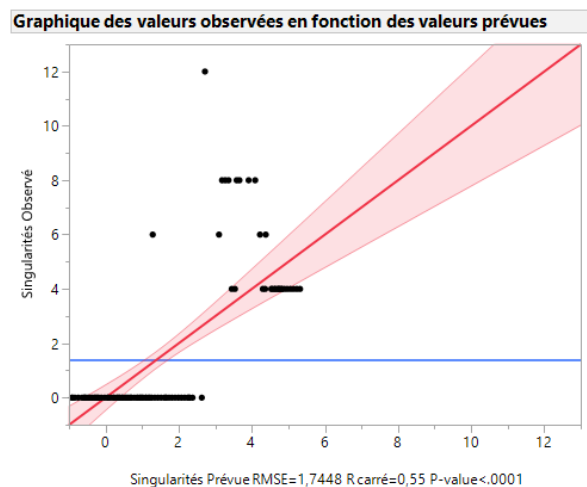


FIGURE 10.4 – Modèle linéaire censé expliquer le nombre de singularités via l'ensemble de nos paramètres.

Nous constatons aisément sur la figure précédente (voir 10.4) que la ligne rouge censée expliquer nos données (les points noirs) ne les explique pas bien. Notons qu'en bleu, nous avons la moyenne globale du nombre de singularités (qui est donc inférieure à deux). Notons aussi que la zone rouge claire est un intervalle de confiance, nous savons donc que la vraie régression linéaire (celle obtenue via l'ensemble des données observables) à un certain nombre de chance (95 pourcent) d'être dans cette intervalle. Cela se confirme en analysant un résumé de l'ajustement (voir 10.1).

R carré	0,554108
Racine de l'erreur quadratique moyenne	1,744834
Moyenne de la réponse	1,396825
Observations (ou sommes pondérées)	126

TABLE 10.1 – Résumé d'ajustement du modèle linéaire de la figure 10.4.

Notons tout d'abord que le R carré est une statistique qui nous permet de mesurer la qualité de la prédiction d'une régression linéaire. Il est donnée par la formule suivante

$$R^2 = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2} \in [0, 1]$$

où n est le nombre de mesures, y_i la valeur de la i ème mesure, $\hat{y}_i$ la valeur prédite correspondante et $\bar{y}$ la moyenne des mesures.

Ici les données obtenues ne sont que descriptibles de façon linéaire à 55 pourcents par nos paramètres. Cela s'explique notamment par le fait que la variable que nous tentons d'expliquer (le nombre de singularités) ne prend que 4 valeurs différentes parmi nos 126 observations et il est dès lors compliqué d'avoir un modèle linéaire performant.

Regardons maintenant les effets dont la p-valeur va nous permettre de nous rendre compte de la probabilité que l'effet observé soit exagéré (voir 10.2). Dans notre cas, la p-valeur est la probabilité sous l'hypothèse de l'existence d'une relation linéaire entre les nombres de singularités obtenues dans nos données et nos paramètres d'obtenir le même effet ou un effet plus faible dans l'ensemble des données observables.

Source	LogWorth	P-value
R_0/R_1	18,404	0,00000
H relatif	0,545	0,28527
Epsilon	0,404	0,39466
Energie	0,055	0,88019

TABLE 10.2 – Résumé des effets du modèle linéaire de la figure 10.4.

Nous constatons via la table 10.2 que le seul terme qui permet d'expliquer linéairement et de façon significative le nombre de singularités est le rapport R_0/R_1 .

Nous allons maintenant nous intéresser à l'explication linéaire de l'apparition de singularités par le rapport des rayons et ce par valeur de ϵ .

Nous voyons, dans le graphique qui suit, que dans le cas d'un $\epsilon = 0, 1$, le rapport R_0/R_1 permet de relativement bien expliquer le nombre de singularités via une relation linéaire alors que quand $\epsilon = 1$, le nombre de points est trop faible pour avoir une explication linéaire avec un certain poids (R carré = 0,065). Cette corrélation linéaire est négative. Une hypothèse pour l'expliquer serait une répulsion due aux bords qui empêcherait l'apparition de singularités sur un anneau trop fin. Toutefois, une répulsion du bord n'a pas de raison claire de faire disparaître des singularités, il y a donc probablement d'autres raisons que nous tenterons d'expliciter plus tard.

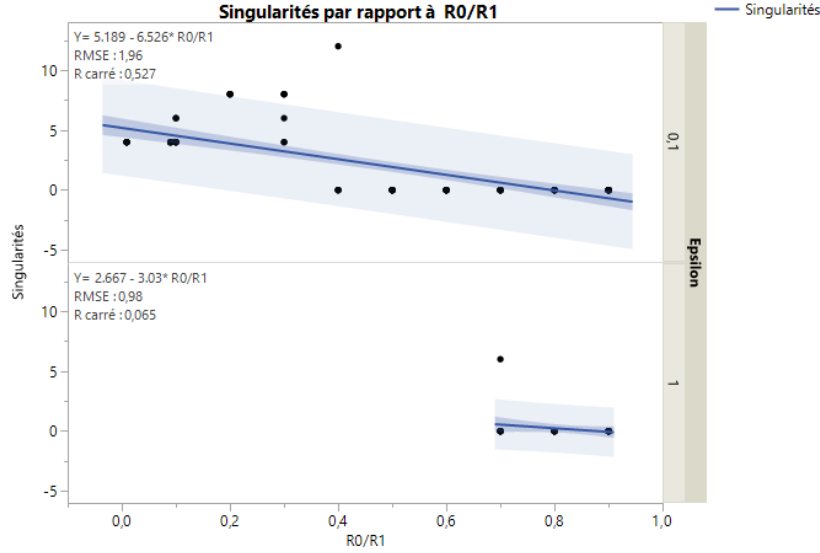


FIGURE 10.5 – Tentative d’explication linéaire du nombre de singularités par le rapport R_0/R_1 pour les différentes valeurs de ϵ .

Notons que, sur la figure précédente (voir 10.5), la zone bleue claire est un intervalle de confiance, nous avons donc que la vraie régression linéaire (celle obtenue via l’ensemble des données observables) à un certain nombre de chance (95 pourcent) d’être dans cette intervalle.

Modèle logarithmique

Nous allons maintenant utiliser un modèle logarithmique qui va nous permettre d’effectuer une régression logistique. L’avantage dans notre cadre d’étude où nous voulons faire apparaître une zone de passage en régime est que la régression logistique fait le passage entre deux niveaux avec une zone plus ou moins grande qui les lie. Si nous obtenons une régression logistique précise avec un passage clair entre les deux niveaux, nous pourrions estimer pour le paramètre associé, le niveau critique de celui-ci pour le passage ou non en régime asymptotique. Parmi les modèles logarithmiques, les modèles logistiques ont la forme suivante

$$\ln \left(\frac{P(\text{niveau1}|\text{données})}{1 - P(\text{niveau1}|\text{données})} \right) = b_0 + b_1 X_1 + \dots b_i X_i$$

pour un vecteur aléatoire $\mathbf{X}$ dont nous supposons que nos données sont des réalisations.

Nous allons maintenant tenter de donner des probabilités sur les nombres de singularités via une régression logistique qui ne dépendra que de notre rapport des rayons. Cette utilisation de la régression logistique est censée nous pénaliser puisque la régression logistique ne permet que de montrer deux niveaux et est donc habituellement utilisée pour expliquer des variables binaires.

Voici le modèle logistique à trois paramètres que nous allons utiliser

$$\frac{c}{1 + \exp(-a(R_0/R_1 - b))}$$

où a = est le taux de croissance, b = le point d’inflexion c = l’asymptote de la fonction logistique utilisée.

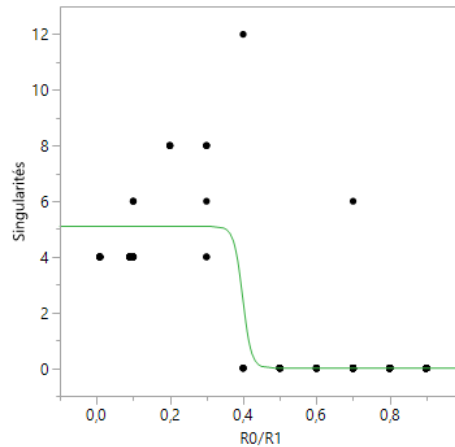


FIGURE 10.6 – Modèle tentant d'approximer le nombre de singularités de nos données via une régression logistique basée sur le rapport R_0/R_1 .

TABLE 10.3 – Estimations des coefficients du modèle logistique de la figure 10.6

Coefficient	Estimation	Erreur standard	Inférieur à 95 %	Supérieur à 95 %
Taux de croissance	-85,22281	3770,3263	-7474,927	7304,481
Point d'inflexion	0,398644	0,0606336	0,2798044	0,5174837
Asymptote	5,0968393	0,2632692	4,5808412	5,6128375

Nous constatons que bien que nous ne soyons pas dans le cadre classique de l'utilisation d'une régression logistique, nous arrivons assez bien à expliquer le passage en mode asymptotique par le rapport des rayons.

Nous allons maintenant définir une nouvelle variable "en régime asymptotique" qui vaudra 1 ou "oui" lorsque nous n'avons pas de singularités et qui vaudra 0 ou "non" lorsqu'il y en aura. Regardons comment un modèle logistique classique permet d'expliquer le passage en régime par le rapport des rayons.

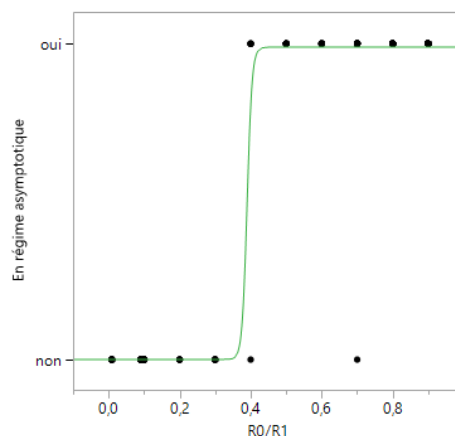


FIGURE 10.7 – Modèle tentant d'approximer une variable qualitative "en régime" issue de nos données via une régression logistique basée sur le rapport R_0/R_1 .

TABLE 10.4 – Estimations des coefficients du modèle logistique de la figure 10.7.

Coefficient	Estimation	Erreur standard	Inférieur à 95 %	Supérieur à 95 %
Taux de croissance	143,42231	198644,9	-389193,4	389480,27
Point d'inflexion	0,398644	0,0606336	0,2798044	0,5174837
Asymptote	0,9888889	0,0127151	0,9639678	1,01381

Nous constatons que le rapport des rayons permet d'expliquer le passage en régime asymptotique pour la grande majorité de nos données. Ces dernières semblent donc nous indiquer que contrairement à ce qui devrait se passer mathématiquement dans le problème de base, nous savons que numériquement c'est le rapport des rayons plus que tout autre paramètre (et donc plus qu' ϵ) qui permet d'expliquer le passage en régime asymptotique.

Notons que nous avons ici fait une analyse quantitative. Si nous voulions effectuer une analyse qualitative afin d'essayer d'identifier, par exemple, un rapport des rayons critique pour le passage en régime asymptotique, nous pourrions construire différents modèles logistiques sur base de différents ensembles de données que nous tenterions de valider les uns les autres via les données non utilisées pour ceux-ci (ce qui est appelé de la validation croisée en machine learning). Nous pourrions alors dégager un ou des candidats pour la valeur critique de rapports des rayons.

Présentation de résultats

Nous allons montrer maintenant via des images associées aux données traitées comment l'augmentation du rapport des rayons fait passer le champ de croix sur un disque à trou en régime asymptotique.

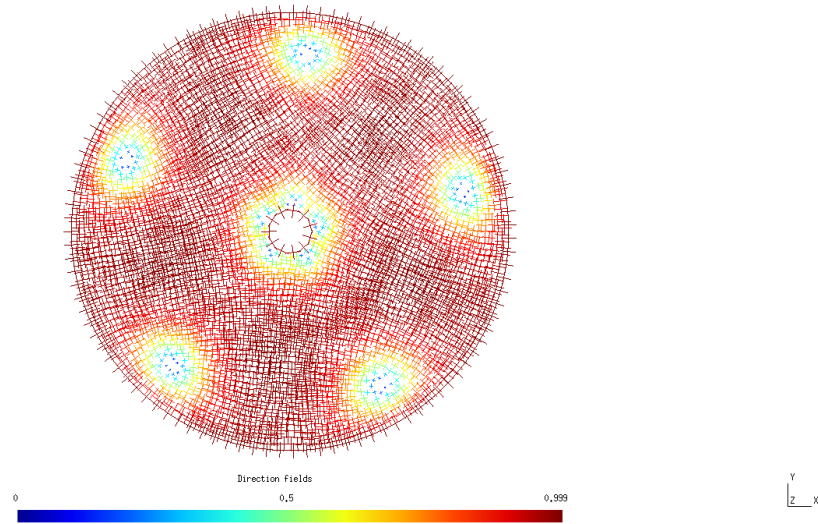


FIGURE 10.8 – Champ de croix avec 6 singularités et ayant une énergie valant 222,1583764 obtenu avec les paramètres suivants : h relatif = 0,2, $R_0/R_1 = 0,1$ et $\epsilon = 0,1$.

Le nombre de singularités correspond au nombre de zones séparées où des croix voient leur norme tendre vers 0. Nous constatons que ce choix algorithmique n'est pas parfait puisqu'il nous donne un nombre de singularités égale à 6 sur la figure 10.8 alors que nous pouvons constater visuellement que le nombre de singularités y est plutôt de 10. Les 5 singularités proches du bord extérieur sont des singularités de degré $+1/4$ puisqu'elles semblent correspondre à des éléments nécessitant 3 voisins plutôt que 4 alors que les singularités proches du trou semblent être de

degré $-1/4$ puisqu'elles paraissent nécessiter 5 voisins plutôt que 4. Notons que nous avons donc peut-être supprimé "trop" de données en ne gardant que celles correspondant à un nombre paire de singularités. Toutefois, les erreurs de décompte du nombre de singularités ne concernent que des résultats pour lesquels un relativement grand nombre de singularités se retrouvent autour d'un suffisamment petit cercle interne et cela n'influence pas les résultats qualitatifs obtenus puisque pour ces disques ayant un petit trou central, nous avons toujours la présence de singularités.

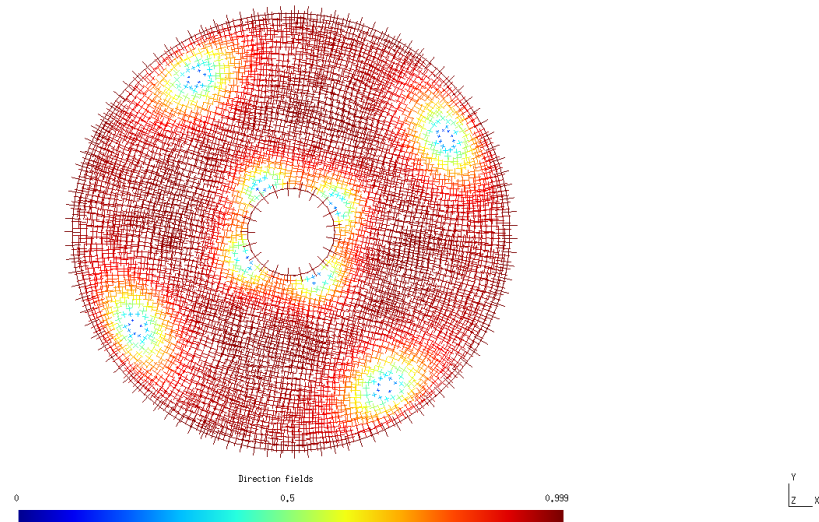


FIGURE 10.9 – Champ de croix avec 8 singularités et ayant une énergie valant 232,2110465 obtenu avec les paramètres suivants : h relatif = 0,2, $R_0/R_1 = 0,2$ et $\epsilon = 0,1$.

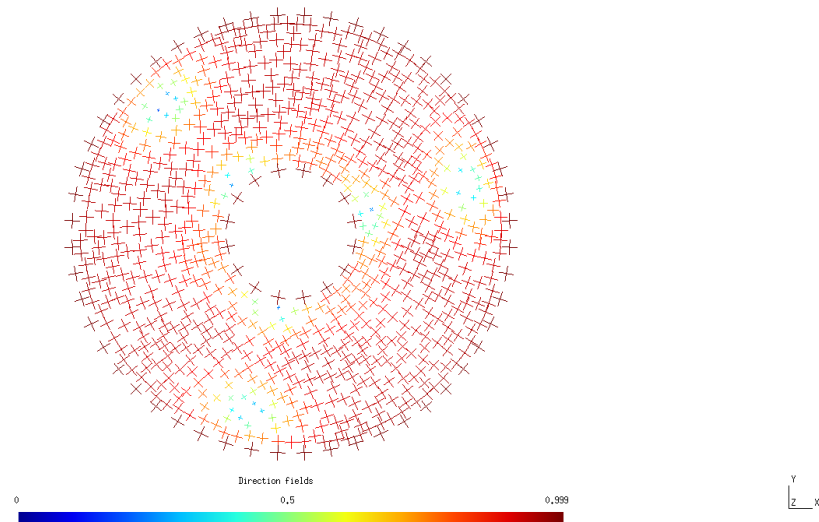


FIGURE 10.10 – Champ de croix avec 6 singularités et ayant une énergie valant 160,39000843 obtenu avec les paramètres suivants : h relatif = 0,4, $R_0/R_1 = 0,3$ et $\epsilon = 0,1$.

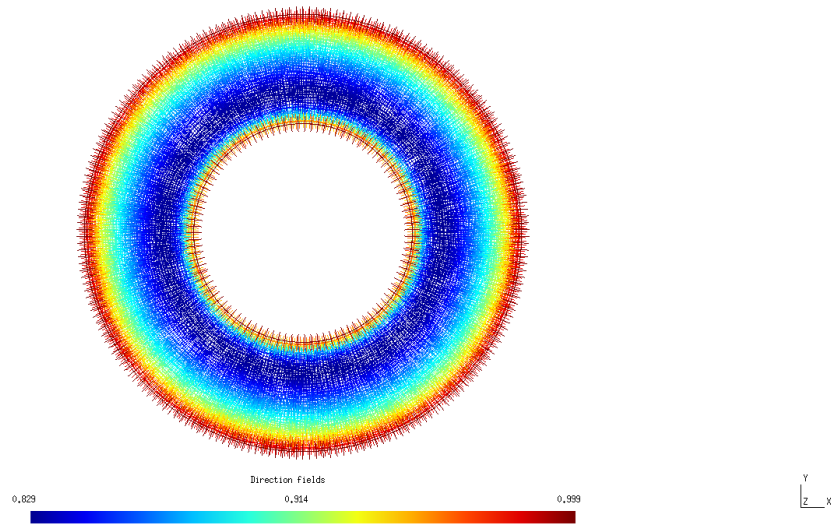


FIGURE 10.11 – Champ de croix avec 0 singularité et ayant une énergie valant 465,9363285 obtenu avec les paramètres suivants : h relatif = 0,1, $R_0/R_1 = 0,5$ et $\epsilon = 0,1$.

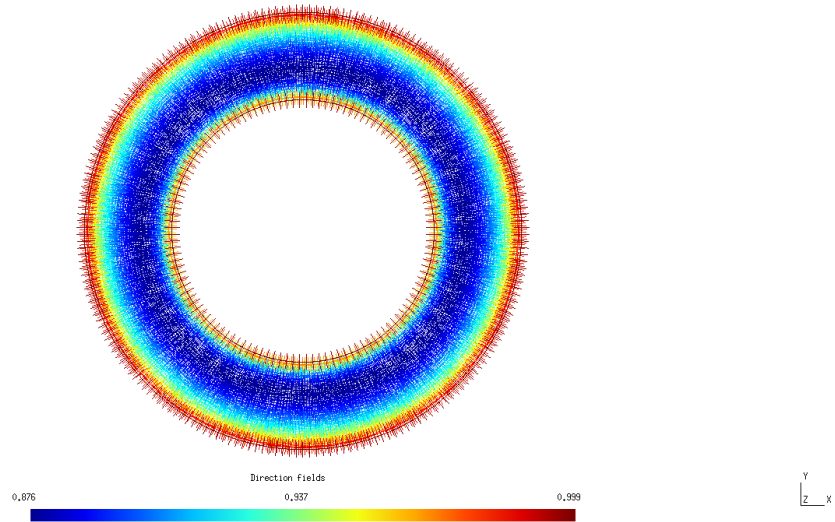


FIGURE 10.12 – Champ de croix avec 0 singularité et ayant une énergie valant 470,2773556 obtenu avec les paramètres suivants : h relatif = 0,1, $R_0/R_1 = 0,6$ et $\epsilon = 0,1$.

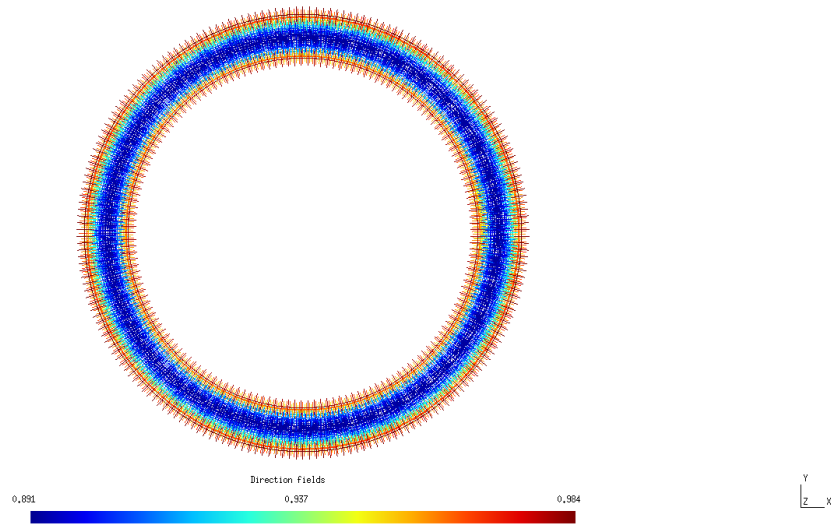


FIGURE 10.13 – Champ de croix avec 0 singularité et ayant une énergie valant 529,8027879 obtenu avec les paramètres suivants : h relatif = 0,1, $R_0/R_1 = 0,8$ et $\epsilon = 1$.

Nous terminons l'analyse de cette première expérience par un graphique qui semble montrer que les changements d'épsilon (pour ce qui concerne les données filtrées) ne change pas qualitativement le résultat et que c'est bien le rapport des rayons qui permet de passer en régime (voir 10.14).

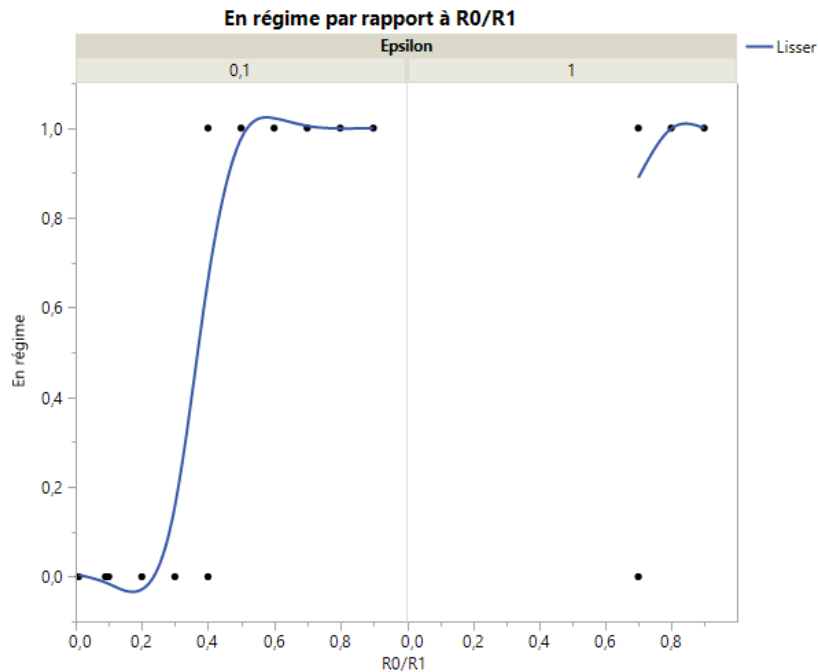


FIGURE 10.14 – Présentation graphique du passage en régime via le rapport de rayons par valeurs de ϵ .

Nous allons maintenant passer à une seconde expérience pour voir si nous obtenons, dans des cas différents et avec plus de données, le même genre de résultats que ceux obtenus lors de notre première expérience.

10.3.2 Deuxième expérience

Puisque l'inadéquation entre le h relatif (qui ne variait que dans un ordre de grandeur dans l'expérience 1) et le ϵ semblait produire des résultats provenant de l'algorithme sans qu'il ait pu converger (voir B.1), une deuxième expérience avec un paramètre ϵ pris dans des ordres de grandeur plus restreints a été effectuée. Nous avons par contre fait varier le h relatif dans des ordres de grandeur plus larges.

La seconde expérience consistait à faire varier via une triple boucle for, le h relatif du maillage dans l'ensemble suivant : $\{0,05; 0,075; 0,1; 0,25; 0,5; 0,75\}$, le rapport des rayons dans l'ensemble :

$\{0,25; 0,5; 0,75\}$ et ϵ dans $\{0,025; 0,05; 0,075; 0,1; 0,75; 0,5; 0,75; 1\}$.

Nous allons commencer par essayer de corroborer la non convergence du code hxt lorsque nous travaillons avec des valeurs de ϵ trop éloignées du h relatif. Une première analyse a été effectuée via des boxplots de l'énergie par h relatif et par ϵ . Pour rappel voici la description de ce qu'est un boxplot (voir 10.15).

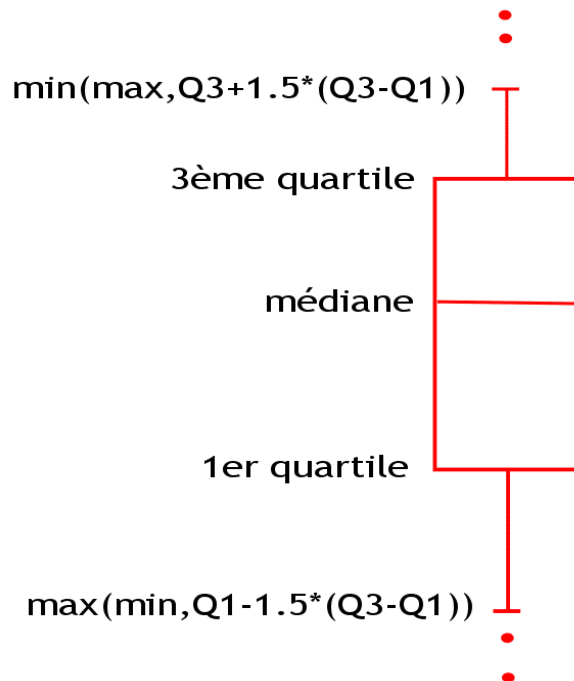


FIGURE 10.15 – Explication du boxplot ou boîte à moustaches provenant du site <https://www.stat4decision.com/fr/le-box-plot-ou-la-fameuse-boite-a-moustache/>.

Notre première analyse de l'énergie des données obtenues nous donne les résultats suivants (voir 10.16).

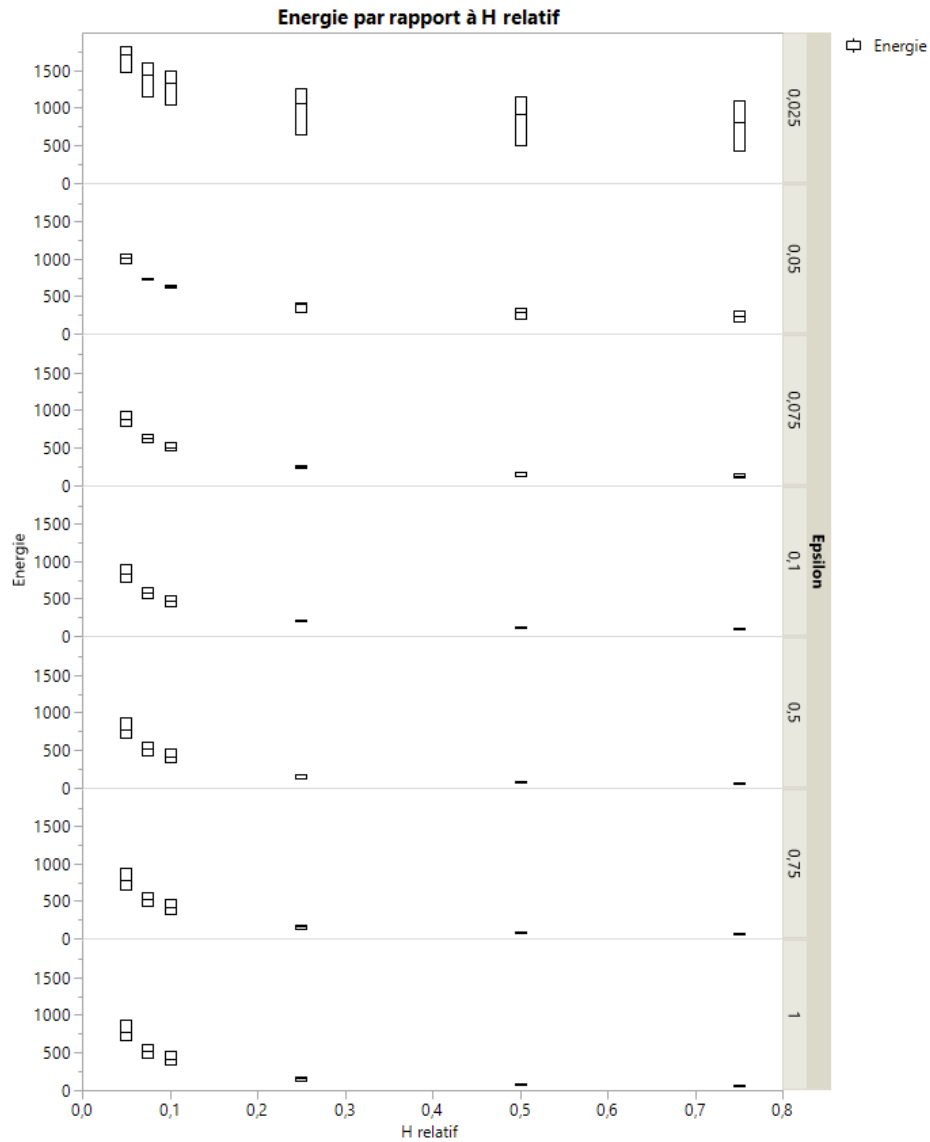


FIGURE 10.16 – Analyse de l'énergie en fonction du h relatif et de ϵ dans les données brutes.

Nous commençons par faire un constat inversé par rapport à précédemment. Nous n'obtenons plus aucune énergie supérieure à 2000 maintenant que ϵ et le h relatif sont pris suffisamment proches. Nous constatons par ailleurs, que plus ϵ est petit plus l'énergie augmente ce qui est mathématiquement évident lorsque nous observons la fonctionnelle énergétique de Ginzburg-Landau. Nous constatons aussi qu'un petit h relatif a aussi tendance à impliquer une plus grande énergie. Cela semble pouvoir s'expliquer par le fait que nous travaillons avec des points discrets qui possèdent chacun une certaine énergie, dès lors un h relatif plus petit implique plus d'éléments et donc plus de points dont nous prenons en compte l'énergie. Toutefois, comme nous le montre les graphes suivants (voir 10.17), nous avons toujours des données qui ne semblent pas provenir d'une relation linéaire entre l'énergie et le logarithme de ϵ comme ça devrait être le cas selon le théorème 8.5.1 et ce pour différentes tailles de mailles.

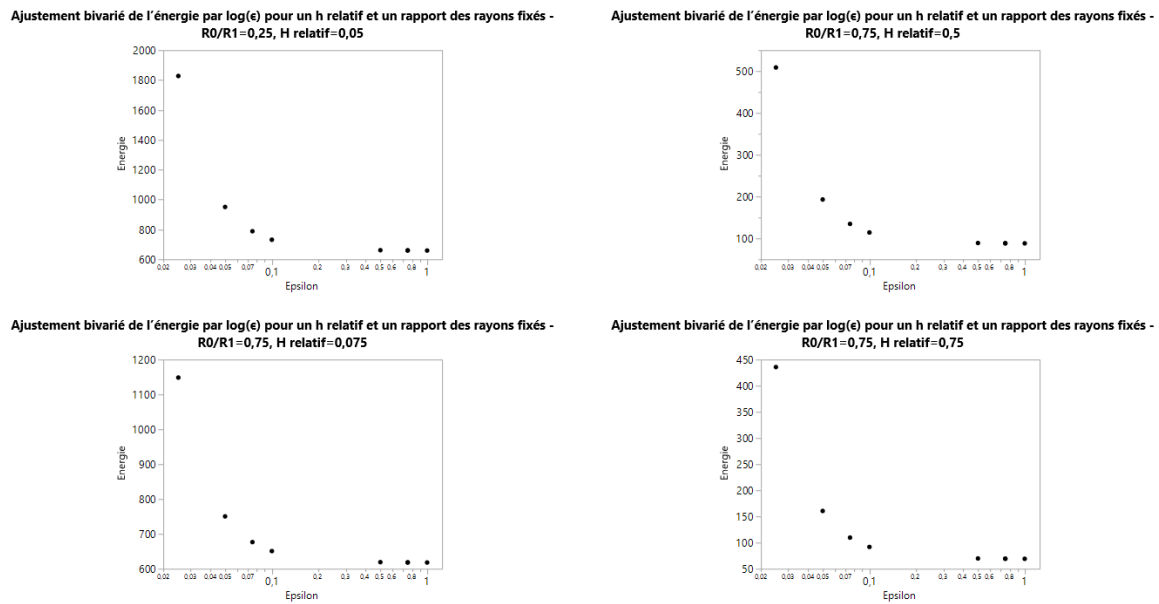


FIGURE 10.17 – Graphes des énergies en fonctions du logarithme de ϵ pour différentes paires de h relatif et de rapport des rayons.

Nous décidons à ce stade de ne garder que les données pour lesquelles le nombre de singularités est pair et l'énergie est inférieure à 1000 pour les mêmes raisons qu'à l'expérience 1 à savoir le fait qu'une énergie supérieure à 1000 est probablement le signe que le solveur numérique n'a pas convergé et que la caractéristique d'Euler-Poincaré implique que nous devons avoir un nombre pair de singularités. En regardant les graphes présentés ci-dessus (voir 10.17) nous pourrions décider de supprimer plus de données. Toutefois, la non-adéquation de chacun de ces graphes avec la forme voulue est déjà largement moins flagrante pour les données conservées et la méthodologie statistique nous incite à supprimer un minimum de données aberrantes (à nouveau, dans le cadre d'une expérience exploratoire, nous nous autorisons un choix subjectif).

Pour la deuxième fois, nous commençons par effectuer un modèle linéaire pour voir quels paramètres sont susceptibles d'expliquer l'apparition ou non de singularités. Nous obtenons les résultats suivants.

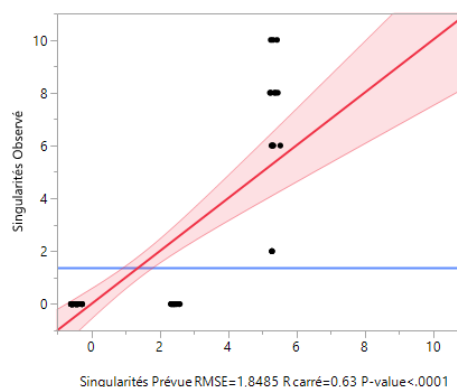


FIGURE 10.18 – Modèle linéaire censé expliquer le nombre de singularités via l'ensemble de nos paramètres.

Nous voyons que dans cette deuxième expérience, les données peuvent être un peu mieux expliquées via un modèle linéaire que précédemment. En effet, le R carré de 0,63 nous indique qu'à peu près 63 pourcent des données peuvent s'expliquer via un modèle linéaire (voir 10.5) alors que nous avons toujours le même nombre moyen de singularités (légèrement inférieur à 2).

R carré	0,630651
Racine de l'erreur quadratique moyenne	1,84848
Moyenne de la réponse	1,376623
Nombre d'observations	77

TABLE 10.5 – Résumé d'ajustement du modèle linéaire de la figure 10.18

Nous constatons via le tableau 10.6 qu'à nouveau et de façon encore plus marquée, seul le rapport des rayons peut expliquer avec une excellente certitude le nombre de singularités.

Source	LogWorth	P-value
R_0/R_1	14,825	0,00000
Energie	0,146	0,71456
H relatif	0,054	0,88380
Epsilon	0,018	0,96023

TABLE 10.6 – Résumé des effets du modèle linéaire de la figure 10.18.

Nous y voyons confirmation que le seul terme qui permet d'expliquer linéairement et de façon significative le nombre de singularités est le rapport R_0/R_1 .

Puisque la corrélation négative peut être directement constatée par le sens de la pente du modèle logistique (voir 10.6 pour l'expérience 1), nous allons directement produire ce genre de modèle sans plus nous attarder sur des modèles linéaires.

Modèle logarithmique

Nous allons utiliser le même type de modèles logistiques que lors de l'expérience 1. Puisque nous pourrions toujours soupçonner une explication non linéaire de nos données par des paramètres qui n'ont pas beaucoup de poids dans la régression linéaire faite précédemment, nous allons effectuer une régression logistique pour chacun des paramètres d'intérêt cette-fois afin toujours d'essayer d'expliquer le nombre de singularités.

Nous voyons via les tableaux et figures suivants que le modèle logistique ne permet pas de montrer plus d'influence sur le nombre de singularités des paramètres ϵ et h relatif que précédemment.

TABLE 10.7 – Estimations des coefficients du modèle logistique de la figure 10.19.

Coefficient	Estimation	Erreur standard	Inférieur à 95 %	Supérieur à 95 %
Taux de croissance	-1,881906	1,5842942	-4,987065	1,2232537
Point d'inflexion	-8,216053	7,000281	-21,93635	5,5042451
Asymptote	11447182	0	11447182	11447182

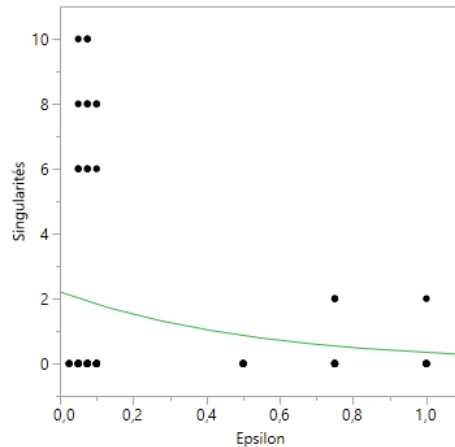


FIGURE 10.19 – Modèle tentant d’approximer le nombre de singularités de nos données via une régression logistique basée sur le paramètre ϵ .

Bien que la figure 10.19 semble nous indiquer que nous avons plus de singularités pour des ϵ faibles (ce qui peut s’expliquer par une plus grande difficulté de convergence de l’algorithme qui pourrait alors s’arrêter sur des minima locaux et non globaux), l’adéquation entre notre régression logistique et nos données reste très grossière.

Regardons maintenant comment un modèle logistique basé sur le h relatif nous permet d’expliquer le nombre de singularités.

TABLE 10.8 – Estimations des coefficients du modèle logistique de la figure 10.20.

Coefficient	Estimation	Erreur standard	Inférieur à 95 %	Supérieur à 95 %
Taux de croissance	8,1928831	5,14e+104	-1e+105	1,01e+105
Point d’inflexion	-29,03327	1,82e+105	-3,6e+105	3,57e+105
Asymptote	1,3766234	0,6689827	0,0654413	2,6878055

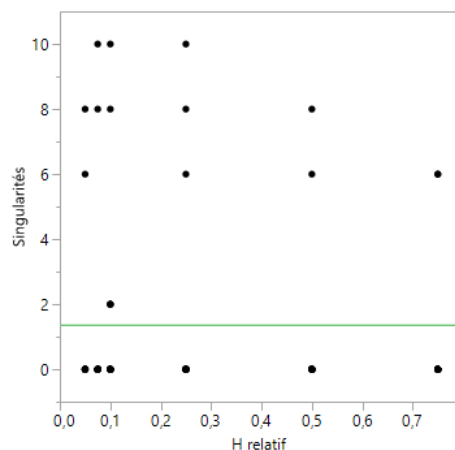


FIGURE 10.20 – Modèle tentant d’approximer le nombre de singularités de nos données via une régression logistique basée sur le h relatif.

Nous constatons sur la figure 10.20 que le h relatif ne semble pas du tout influencer le nombre de singularités.

Par contre, comme lors de l'expérience 1 et de façon encore plus marquée même, le rapport des rayons explique très bien le nombre de singularités avec un modèle logistique même si, comme précédemment, celui-ci est beaucoup plus performant pour expliquer des variables binaires.

TABLE 10.9 – Estimations des coefficients du modèle logistique de la figure 10.21.

Coefficient	Estimation	Erreur standard	Inférieur à 95 %	Supérieur à 95 %
Taux de croissance	-8,891827	1,8493296	-12,51645	-5,267207
Point d'inflexion	-1,437927	91715,029	-179759,6	179756,72
Asymptote	21570784	1,759e+13	-3,45e+13	3,448e+13

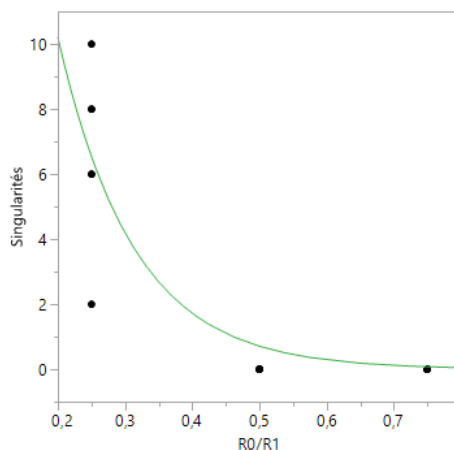


FIGURE 10.21 – Modèle tentant d'approximer le nombre de singularités de nos données via une régression logistique basée sur le rapport R_0/R_1 .

Nous y voyons à nouveau que c'est le rapport des rayons qui semble principalement expliquer le nombre de singularités même si, comme lors de l'expérience 1, nous n'avons pas encore utilisé la régression logistique dans son cadre habituel.

Passons à un modèle de régression logistique classique via l'utilisation de la même variable "en régime" que lors de l'expérience 1.

TABLE 10.10 – Estimations des coefficients du modèle logistique de la figure 10.22.

Coefficient	Estimation	Erreur standard	Inférieur à 95 %	Supérieur à 95 %
Taux de croissance	122,93481	9,3489518	104,6112	141,25841
Point d'inflexion	0,3877366	0,0085428	0,3709931	0,4044801
Asymptote	1,0000003	6,2338e-8	1,0000002	1,0000004

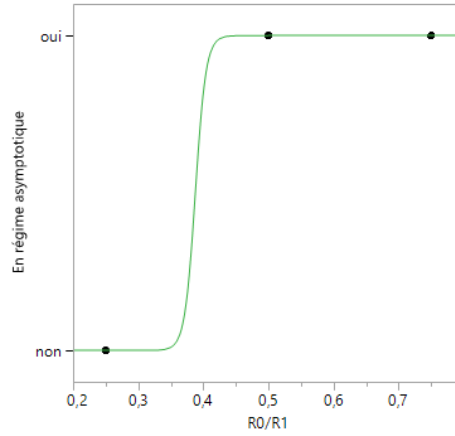


FIGURE 10.22 – Modèle tentant d’approximer une variable qualitative "en régime" issue de nos données via une régression logistique basée sur le rapport R_0/R_1 .

Notons que dans le graphe 10.22, les trois points visibles correspondent à un grand nombre de superposition de points (nous avons 77 observations qui rentrent en compte sur ce graphe).

Nous avons cette fois-ci une adéquation parfaite qui semble dire que le rapport des rayons peut à lui tout seul expliquer parfaitement le passage ou non en régime asymptotique !

Présentation de résultats

Nous allons montrer maintenant via des images associées aux données traitées, dans l’expérience 2 cette fois-ci, comment l’augmentation du rapport des rayons fait passer le champ de croix sur un disque à trou en régime asymptotique.

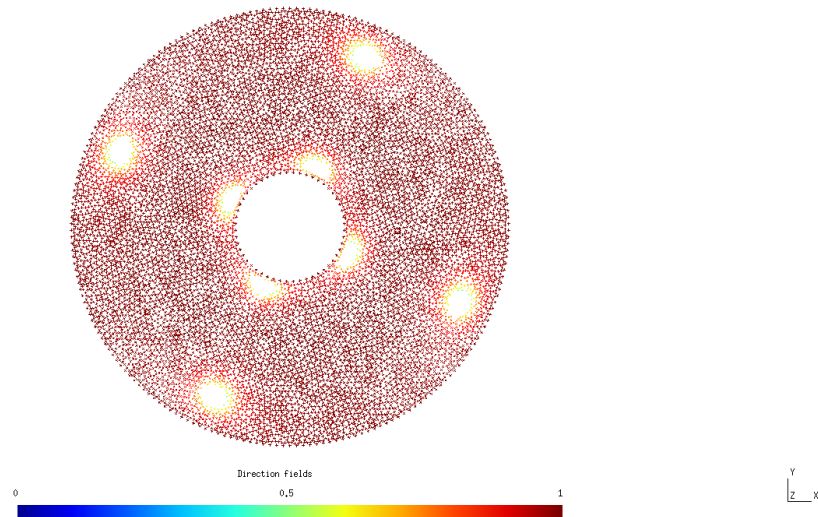


FIGURE 10.23 – Champ de croix avec 8 singularités et ayant une énergie valant 624,7643912 obtenu avec les paramètres suivants : h relatif = 0,1, $R_0/R_1 = 0,25$ et $\epsilon = 0,05$.

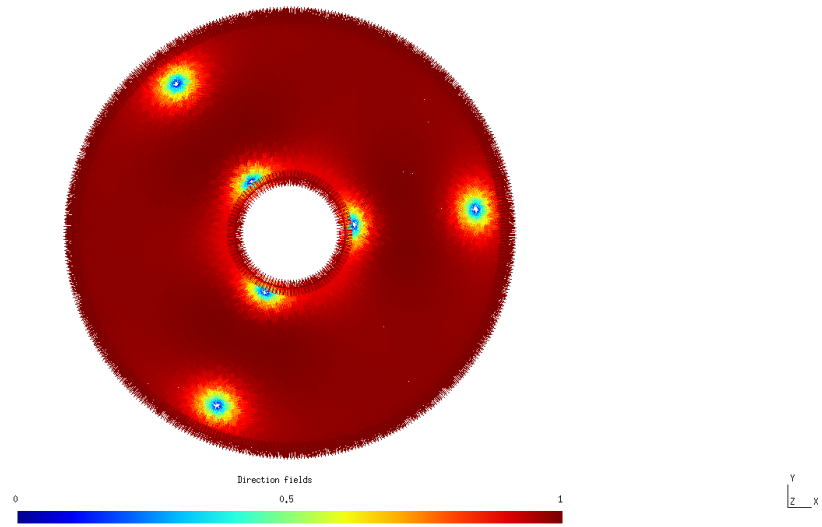


FIGURE 10.24 – Champ de croix avec 6 singularités et ayant une énergie valant 949,8375824 obtenu avec les paramètres suivants : h relatif = 0,05, $R_0/R_1 = 0,25$ et $\epsilon = 0,05$.

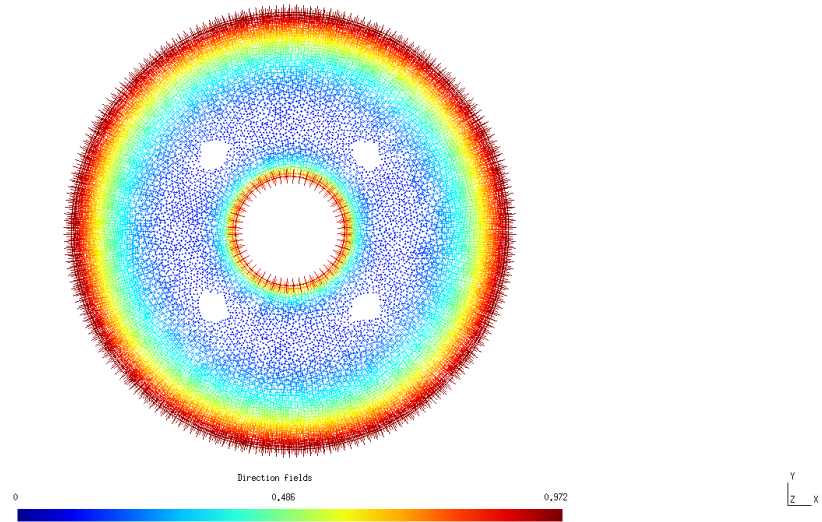


FIGURE 10.25 – Champ de croix avec 2 singularités et ayant une énergie valant 335,9914229 obtenu avec les paramètres suivants : h relatif = 0,1, $R_0/R_1 = 0,25$ et $\epsilon = 0,75$.

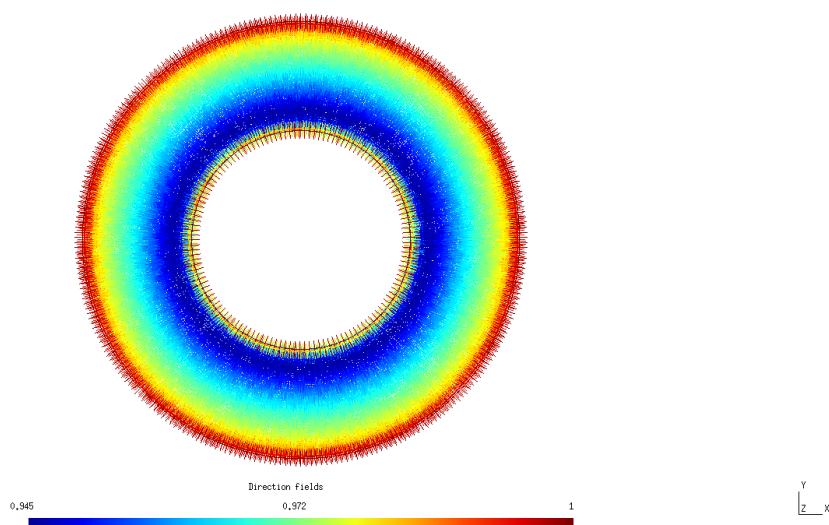


FIGURE 10.26 – Champ de croix avec 0 singularité et ayant une énergie valant 750,9883215 obtenu avec les paramètres suivants : h relatif = 0,075, $R_0/R_1 = 0,5$ et $\epsilon = 0,05$.

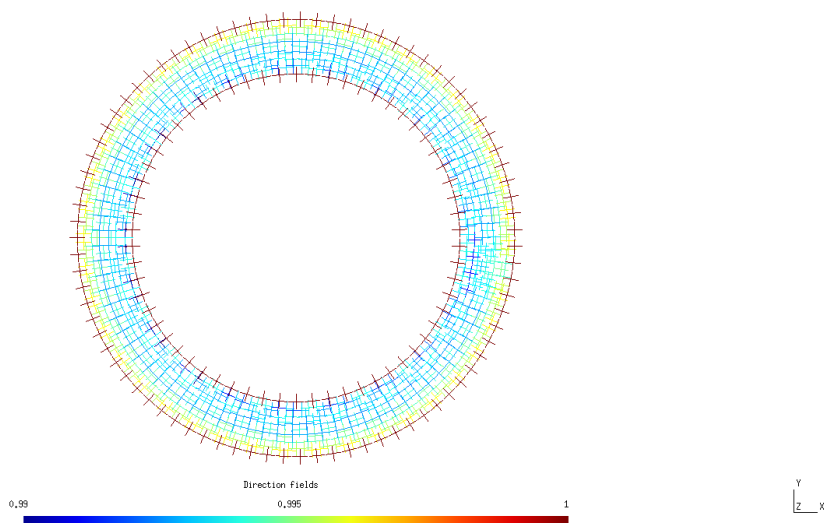


FIGURE 10.27 – Champ de croix avec 0 singularité et ayant une énergie valant 655,7821814 obtenu avec les paramètres suivants : h relatif = 0,25, $R_0/R_1 = 0,75$ et $\epsilon = 0,025$.

Nous terminons l'analyse de cette seconde expérience par un graphique qui semble montrer que les changements d'épsilon (pour ce qui concerne les données filtrées) ne changent pas qualitativement le résultat et que c'est bien le rapport des rayons qui permet de passer en régime (voir 10.28).

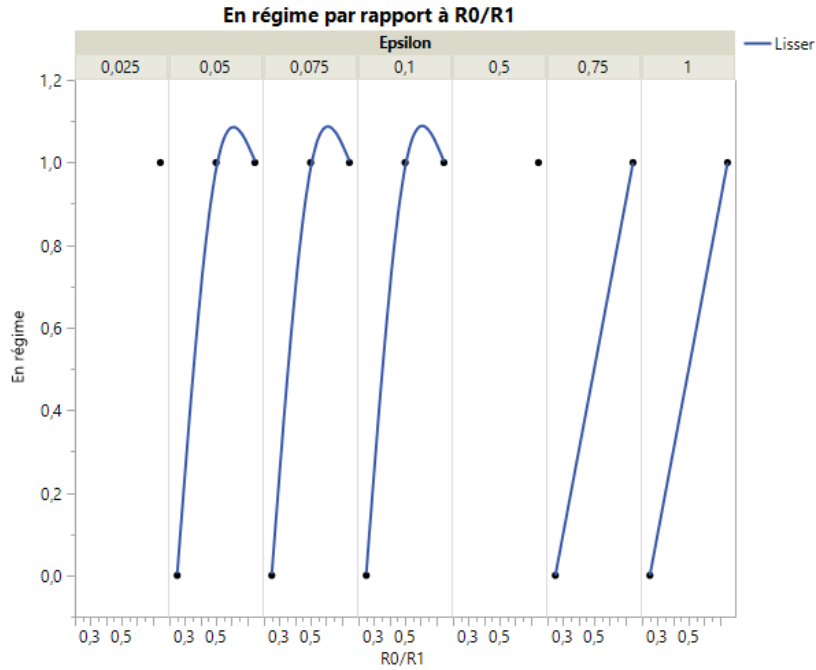


FIGURE 10.28 – Passage en régime via le rapport de rayons par valeurs de epsilon pour l’expérience 2.

Ce graphique appuie de façon plus complète l’observation faite en fin d’expérience 1 à savoir que les variations d’ ϵ ne changent pas qualitativement nos résultats.

10.4 Conclusion des expériences

Nous allons commencer par résumer nos observations puis nous tenterons de les expliquer.

Nous avons d’abord montrer que des valeurs trop différentes de h relatif et de ϵ semblaient compliquer la convergence du schéma numérique. Nous avons ensuite constaté que, dans nos expériences et via une régression logistique, le rapport des rayons est le paramètre qui permet le mieux d’expliquer ou non le passage en régime asymptotique. Nous pouvons aussi noter en observant les différents disques obtenus que nous avons une force de répulsion des bords et que cette répulsion est d’autant plus grande que la courbure du bord est faible.

Commençons par tenter d’expliquer notre première observation. Voici une explication possible. Lorsque notre taille de mailles est trop grande par rapport à ϵ , le poids du terme non linéaire est trop grand et les tailles de maille sont trop petites pour que l’algorithme réussisse à minimiser suffisamment cette partie non linéaire suffisamment rapidement dès lors l’algorithme est incapable de converger avec un nombre d’itérations inférieur au nombre maximal d’itérations. La minimisation à effectuer est trop intense au vue du nombre de degrés de liberté (de mailles) et du nombre d’itérations possibles.

Tentons maintenant de donner une explication possible pour notre deuxième observation. Pour avoir une chance d’atteindre le régime asymptotique, nous avons besoin qu’ ϵ (qui correspond à une taille caractéristique du problème continu qui est approximé) soit inférieur à certaines mesures caractéristiques du problème numérique discret. Or, le rapport des rayons lorsqu’il tend vers 1 permet aux éléments d’avoir une rotation plus limitée sur les bords qu’ils doivent épouser. Dès lors, plus ce rapport des rayons tend vers 1 et plus la taille des éléments peut être grande tout en les laissant ajustés aux bords de notre domaine. Donc le rapport des rayons semble en ce

sens garant d'une certaine grandeur caractéristique du problème discret. Ainsi, ϵ va rapidement être suffisamment petit par rapport à une taille caractéristique de notre problème et donc nous allons pouvoir atteindre le régime asymptotique.

Pour notre troisième observation, il est plus facile de proposer une explication possible. En effet, nous imposons fortement à nos croix d'épouser le bord avec une norme 1, dès lors, il est impossible pour une singularité de venir coller le bord et ce d'autant plus que le bord est plat car alors, la singularité subit dans un même voisinage l'influence de plus de croix fixes du bord avec des orientations proches que si le bord à une forte courbure (courbure qui impliquera que les croix du bord auront rapidement des orientations très différentes).

10.5 Remarques supplémentaires

Il est important de noter que nous avons travaillé précédemment dans un cas où l'utilisation d'un gradient covariant ne changeait rien puisque nous travaillions alors sur une surface de courbure nulle. Toutefois, l'introduction d'un gradient covariant pourrait potentiellement améliorer la méthode numérique puisque, dans le cas d'un tore par exemple, elle ne permet, sans l'utilisation d'un gradient covariant, et pour les tores sur lesquels cela a été essayé, pas encore de construire un champ de croix sans singularité. Or, mathématiquement, les minimiseurs asymptotiques ne devraient jamais contenir de singularités sur un tore. Il serait intéressant d'étudier l'influence de l'utilisation d'un gradient covariant pour construire un champ de croix sur un tore ou sur d'autres surfaces de courbure non nulle mais de caractéristique d'Euler-Poincaré égale à 0. Notons que l'introduction d'un gradient covariant dans la méthode numérique est une tâche ardue qui n'a pas été réalisée dans le cadre de ce mémoire.

Notons aussi que les conclusions effectuées le sont dans le cadre des valeurs de ϵ utilisées. En effet, même en supposant que la fusion de singularités opposées ait un coût fixe, ce coût fixe serait alors dépendant du domaine et potentiellement très élevé. Dès lors, le passage en régime asymptotique qui dépendant d'une comparaison entre le logarithme de ϵ et ce coût fixe peut nécessiter des ϵ très petits (qui ne sont potentiellement pas atteignables avec les méthodes numériques utilisées dans ce chapitre). Ainsi, si nous cherchons à obtenir numériquement des informations sur le problème mathématique de la forme et de l'existence de minimiseurs pour notre fonctionnelle de Ginzburg-Landau sur différentes surfaces, il faut probablement adapter l'algorithme numérique.

Chapitre 11

Conclusion

Ce mémoire prenait place dans le cadre de l'étude de méthodes destinées à produire des maillages quadrangulaires réguliers. Pour cela, il s'est basé sur une fonctionnelle étant censée fournir des informations sur la faisabilité de tels maillages et sur le placement de points singuliers si nécessaire. Ce sont les minimiseurs de cette fonctionnelle (ou de fonctionnelles analogues) prise avec un paramètre ϵ tendant vers 0 qui nous permettent de trouver, lorsqu'il doit y en avoir, les points supposés optimaux pour placer, dans un maillage quadrangulaire le plus régulier possible de la surface considérée, des nœuds singuliers. Lorsque nous faisons tendre le paramètre ϵ vers 0, nous passons mathématiquement en régime asymptotique et ce sont les minimiseurs du problème asymptotique qui nous permettent de donner des localisations optimales aux éventuels nœuds singuliers d'un maillage quadrangulaire régulier sur la même surface.

Après avoir motivé cette étude dans le premier chapitre, nous avons démontré un argument topologique pouvant être utilisé pour montrer que toute surface fermée ne peut pas être maillée uniquement avec des éléments quadrangulaires.

La fonctionnelle introduite au début a ensuite été étudiée. Nous avons notamment montré l'existence de solutions parmi un ensemble de fonctions dites dans certains contextes d'énergie finie pour des surfaces suffisamment régulières. Nous avons aussi montré que, dans le cas particulier de surfaces de type tore T^2 , les minimiseurs de notre fonctionnelle appartiennent à un ensemble de fonctions qui sont des analogues non-linéaires et covariants d'applications harmoniques.

Nous avons ensuite étudié l'apparition de défauts dans les minimiseurs d'une fonctionnelle du même type que celle étudiée sur certaines surfaces lorsqu'un paramètre ϵ tend vers 0. Ces défauts apparaissent lorsque nous essayons de mettre en rotation des vecteurs de norme 1 autour d'un point. Nous sommes alors obligé de faire tendre la norme de nos vecteurs tangents vers 0 en atteignant ce point. Vu leur nature ces défauts ont donc été nommés "tourbillons". La localisation de ces tourbillons peut être mise en correspondance avec la localisation des points singuliers de maillages quadrangulaires. Ces points sont ceux ayant plus de voisins que ne le nécessite la production d'un maillage quadrangulaire parfaitement régulier.

Nous avons ensuite tenté de formaliser des objets mathématiques qui peuvent représenter des champs de croix. Nous avons vu que ceux-ci peuvent être définis de différentes manières et notamment via des polynômes homogènes harmoniques de degré 4 à 2 variables ou via une paire d'un cosinus et d'un sinus d'un angle multiplié par 4.

Finalement nous avons étudié, pour une fonctionnelle du même type que celle initialement introduite, l'usage numérique de ce genre de méthode pour la production de maillages quadrangulaires. Le but était de déterminer les paramètres numériques permettant le passage en régime asymptotique pour des surfaces étant des disques à trou. Mathématiquement, les minimiseurs du régime asymptotique minimisent le nombre de singularité et optimisent leur localisation. Dès lors, numériquement, c'est sur base des minimiseurs du problème asymptotique que nous cherchons à placer les points singuliers de maillages quadrangulaires se voulant les plus réguliers possibles.

Toutefois, il est numériquement impossible de faire tendre ϵ vers 0 et difficile d'obtenir des solutions correspondant au régime asymptotique.

L'étude effectuée nous a permis de constater que, pour des disques à un trou, les limitations numériques, notamment l'incapacité de produire des nombres suffisamment proches de zéro, impliquaient que ce n'était pas le paramètre ϵ , mais le rapport des rayons qui permettait d'expliquer le passage numérique en régime asymptotique.

Dans le cadre de surfaces de type disque à un trou central, le passage en régime asymptotique semble donc régi par le rapport des rayons (du moins pour les ϵ considérés). Plus les rayons sont proches l'un de l'autre et plus nous sommes susceptibles de passer en régime asymptotique numériquement. C'est sans doute encore à creuser, mais nous dégageons deux pistes d'explications pour ce phénomène. Nous avons d'une part un phénomène de répulsion par le bord (ce d'autant plus que la courbure du bord est faible) et d'autre part le rapport des rayons semble être garant d'une certaine grandeur caractéristique du problème discret de maillage en ce sens qu'il est plus simple d'ajuster des quadrangles entre deux cercles proches qu'entre deux cercles concentriques de tailles très différentes. Par ailleurs, des paramètres d'ordres de grandeur trop différents semblent complexifier le passage numérique en régime asymptotique ce qui pourrait s'expliquer par des difficultés de convergence de l'algorithme dans ces cas.

En conclusion, l'étude mathématique de fonctionnelles de Ginzburg-Landau est intéressante dans le cadre de la production de maillages quadrangulaires réguliers puisque les singularités des minimiseurs de celles-ci semblent être des endroits de choix pour placer numériquement des nœuds singuliers dans nos maillages.

Sans changer le code actuelle, de nouvelles expériences pourraient être intéressantes. En effet, par exemple, nous nous attendons de façon intuitive à ce qu'il soit plus facile d'avoir une distorsion moyenne de nos éléments plus faible dans un maillage fin que dans un maillage grossier pour une même configuration de singularités. Il serait dès lors intéressant de mesurer l'effet de la finesse du maillage sur l'apparition de singularités et sur la distorsion moyenne des croix produites.

Par ailleurs, de façon à se mettre plus en adéquation avec de nombreuses surfaces d'intérêt, il serait intéressant de prendre en compte l'aspect non-euclidien de certaines d'entre elles. Ainsi, l'introduction dans les outils numériques du gradient covariant, par exemple via l'utilisation de la fonctionnelle (1.1), peut s'avérer intéressante. L'étude plus large de l'intégration numérique de géométrie non euclidienne peut aussi être profitable.

Toutefois, il y a, comme nous l'avons vu, un autre aspect puisque les fonctionnelles utilisées dépendent de paramètres dont il est impossible (du moins pour le moment) de reproduire le comportement mathématiquement voulu de façon numérique. Dès lors, un deuxième volet d'investigations serait intéressant pour tenter de comprendre l'interaction numérique entre ces paramètres et le comportement des minimiseurs des fonctionnelles utilisées. Une façon de procéder pourrait notamment être dans un premier temps d'effectuer des études statistiques pour essayer de cerner les caractéristiques et paramètres majeurs de problèmes numériques emblématiques.

Il serait aussi intéressant de réfléchir aux types de maillages que nous souhaitons construire (comme ce fut mentionné dans le chapitre sur les motivations de ce travail de fin d'études, il n'y a pas toujours de réponse claire à cette interrogation).

Finalement, d'une part, pour tenter d'obtenir numériquement des informations supplémentaires sur le problème mathématique, il pourrait être intéressant d'essayer d'adapter le schéma numérique utilisé. D'autre part, pour valider les résultats numériques de nouvelles méthodes, continuer l'étude des fonctionnelles mathématiques associées sur des surfaces intéressantes est essentiel.

Annexes A

Le degré de Brouwer

L'essentiel de l'introduction suivante au degré de Brouwer peut être trouvée à l'url suivante : <https://math.unice.fr/~hoering/students/cartau.pdf>.

Dans cette partie, nous allons développer une notion qui permet d'associer un entier à une application lisse entre variétés. Il s'agit d'une généralisation de l'indice pour les courbes dans le plan. Conformément à l'intuition, un tel entier ne changera pas si l'on "déforme continûment" l'application en question (formellement, si nous regardons une application qui lui est homotope). Cet invariant puissant est un bon outil de classification, et a des applications surprenantes, en géométrie différentielle par exemple.

A.1 Rappels sur l'orientation des espaces vectoriels

Considérons E , un $\mathbb{R}$ -espace vectoriel de dimension finie. Nous définissons une relation d'équivalence sur l'ensemble des bases de E comme suit : si $(e_i)_i$ et $(e'_i)_i$ avec $i = 1, \dots, n$ sont deux bases de E , alors elles sont équivalentes si et seulement si la matrice de passage entre ces deux bases a un déterminant positif. Cette relation partitionne l'ensemble des bases de E en deux classes d'équivalence.

On appelle orientation sur un espace vectoriel (réel) le choix d'une de ces deux classes, et on appelle espace vectoriel orienté un couple (E, e) où E est un espace vectoriel, et $e \equiv [e]$ est une classe d'équivalence de bases. Nous disons que $[e]$ est l'orientation de E . La base canonique dans $\mathbb{R}^n$, par exemple, induit une orientation qu'on appellera orientation canonique sur $\mathbb{R}^n$.

Remarque A.1.1. L'ordre des vecteurs dans la base considérée joue un rôle crucial dans l'orientation déterminée par ladite base.

A.2 Orientation des variétés

La question d'une orientation sur les variétés se pose directement par le biais des espaces tangents qui ont une structure naturelle d'espace vectoriel (de manière informelle, nous voulons faire un choix "continu" d'orientation pour chaque espace tangent de la variété).

Définition A.2.1. Soient M et N des n -variétés lisses, sans bord, et orientables. Nous supposons M compacte et N connexe. Soit également $f : M \rightarrow N$ une application lisse. Soit $x \in M$ un point régulier tel que $df_x : TM_x \rightarrow TN_{f(x)}$ soit un isomorphisme entre espaces vectoriels de dimension n . Nous définissons le signe sgn de df_x comme étant égal à 1 si df_x préserve l'orientation et -1 sinon. Alors, pour toute valeur régulière $y \in N$, nous posons :

$$\deg(f; y) = \sum_{x \in f^{-1}(y)} sgn(df_x)$$

qui est le degré de Brouwer de l'application f au point y . Comme $\#f^{-1}$ est localement constant sur les valeurs régulières, alors $\deg(f; y)$ l'est aussi.

Le résultat principal qui rend cet outil aussi intéressant est le suivant : $\deg(f; y)$ ne dépend pas de la valeur régulière choisie. Nous parlerons alors simplement de $\deg(f)$. De plus, le degré est invariant par homotopie lisse.

Lemme A.2.1. *Supposons que M soit le bord d'une variété compacte et orientable X , et que M soit orientée en tant que bord de X comme nous l'avons vu précédemment. Soit $f : M \rightarrow N$ lisse. Si il existe $F : X \rightarrow N$ lisse telle que $F|_M = f$, alors $\deg(f; y) = 0$ pour toute valeur régulière $y \in N$.*

Théorème A.2.1. *Si $f \sim g$, alors $\deg(f) = \deg(g)$*

Proposition A.2.1. $\deg(g \circ f) = \deg(g) \deg(f)$

Remarque A.2.1. Un difféomorphisme $f : M \rightarrow N$ aura toujours un degré $+1$ ou -1 (car c'est une bijection), en fonction de s'il préserve ou change l'orientation. En particulier, ceci implique qu'un difféomorphisme d'une variété lisse et sans bord qui change l'orientation ne peut pas être lissement homotope à l'identité (car le degré est invariant par homotopie, et $+1 \neq -1$).

Théorème A.2.2. *Soit $f : S^n \rightarrow S^n$ une application lisse. Si $\deg(f) \neq (-1)^{n+1}$, alors f admet un point fixe.*

Proposition A.2.2. *Si n est un entier pair, alors l'application antipodale $a : S^n \rightarrow S^n$ n'est pas homotope à l'identité.*

Annexes B

Analyse statistique

B.1 Graphes concernant le rejet de certaines données dans la première expérience

Parmi les données obtenues lors de l'expérience 1, certaines sont associées à une énergie particulièrement élevée. Cela nous fait penser que ces données proviennent de résolutions numériques qui ont atteint le nombre maximal d'itérations du schéma numérique. Via les données brutes, il a été observé que tant que ϵ est plus ou moins proportionnel au h relatif, nous n'avons pas d'énergies supérieures à 10000. Après, pour chaque diminution de l'ordre de grandeur de ϵ , nous gagnons au moins un ordre de grandeur en énergie.

Pour illustrer cela, nous avons dans la figure B.2 : une ligne par valeur de ϵ , un axe des x qui correspond au h relatif et au croisement d'un ϵ et d'un h relatif nous avons un boxplot. Pour rappel voici la description de ce qu'est un boxplot (voir 10.15).

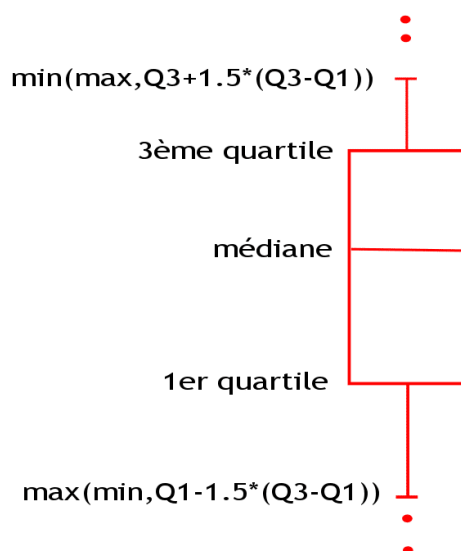


FIGURE B.1 – Explication du boxplot ou boîte à moustaches provenant du site <https://www.stat4decision.com/fr/le-box-plot-ou-la-fameuse-boite-a-moustache/>.

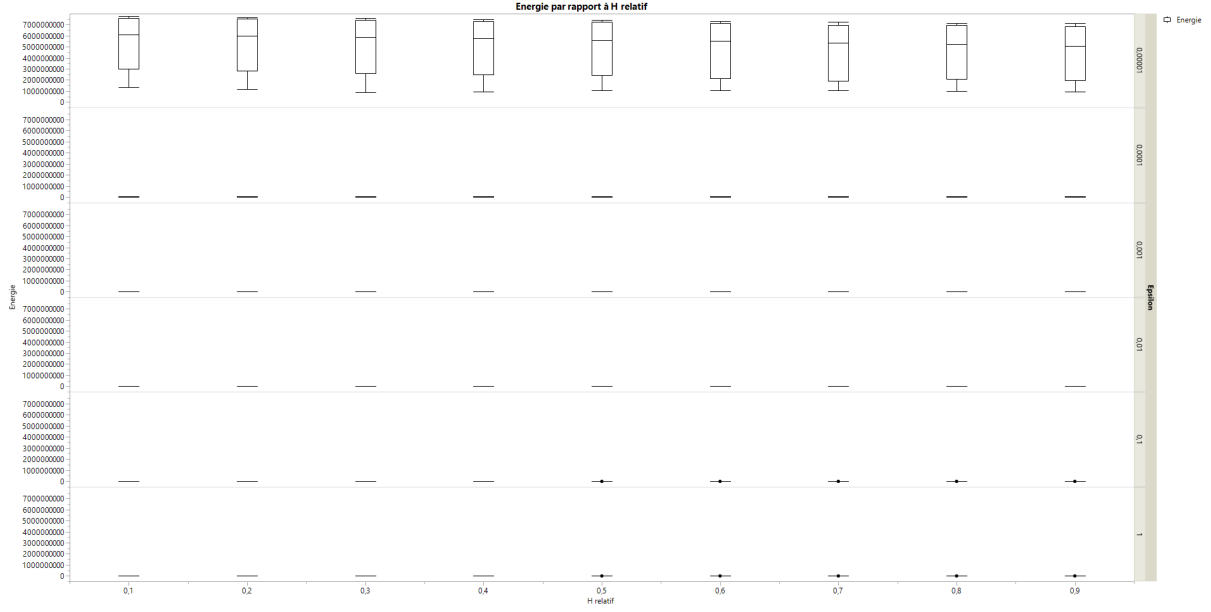


FIGURE B.2 – Analyse de l'énergie en fonction du h relatif et de ϵ dans les données brutes.

Nous constatons sur la figure B.2 que, lorsque ϵ est trop petit par rapport au h relatif, les ordres de grandeurs des énergies sont énormes et puisque nous gardons la même échelle pour tous les boxplots de la figure, ceux pour lesquels aucune valeur de l'énergie n'est très élevée se retrouvent complètement écrasés puisqu'ils sont dessinés dans une échelle qui ne leur convient pas. Ainsi, nous supposons que ces énergies ne correspondent pas à des minima de schémas de Ginzburg-Landau mais qu'ils correspondent plutôt à des énergies obtenues après non convergence de la résolution numérique de la fonctionnelle énergétique de Ginzburg-Landau sur les disques à trous correspondant. Nous supposons donc que ces données proviennent de résolutions numériques qui ont atteint le nombre maximal d'itérations du schéma numérique.

B.2 Données supplémentaires concernant l'expérience 1

100.0%	maximum	556,3435123
99.5%		556,3435123
97.5%		499,21945689
90.0%		296,15746794
75.0%	quartile	216,983046
50.0%	médiane	153,6781042
25.0%	quartile	112,3516306
10.0%		89,286485824
2.5%		69,005731925
0.5%		46,31398333
0.0%	minimum	46,31398333

TABLE B.1 – Quantiles de l'énergie pour l'expérience 1.

Moyenne	184,25801
Écart-type	109,2676
Erreur standard de la moyenne	9,7343318
Limite supérieure de l'intervalle de confiance de la moyenne pour 95 %	203,52346
Limite inférieure de l'intervalle de confiance de la moyenne pour 95 %	164,99256
Nombre d'observations	126

TABLE B.2 – Statistiques de résumé de la distribution de l'énergie pour l'expérience 1.

B.3 Un des scripts bash utilisé pour l'analyse statistique

```
#!/bin/bash
echo "H relatif; R0; Epsilon; Singularities; Energy" > sorties.txt
for r0fact in {1..9}
do
  varR0=$(echo "scale=5; 0.1*$r0fact" | bc)
  sed -n 1p diskHole.geo > diskHoleTmp.geo
  sed -n 2p diskHole.geo >> diskHoleTmp.geo
  sed -n 3p diskHole.geo >> diskHoleTmp.geo
  sed -n 4p diskHole.geo >> diskHoleTmp.geo
  echo "Disk(2) = {0, 0, 0, 0"$varR0", 0"$varR0"};" >> diskHoleTmp.geo
  sed -n 6p diskHole.geo >> diskHoleTmp.geo
  sed -n 7p diskHole.geo >> diskHoleTmp.geo
  sed -n 8p diskHole.geo >> diskHoleTmp.geo
  for expo in {0..5}
  do
    varEpsilon=$(echo "scale=5;1*10^(-$expo)" | bc)
    for scale in {1..9}
    do
      varScale=$(echo "scale=5; 0.1*$scale" | bc)
      echo -e "r0 = "$varR0", epsilon = "$varEpsilon", hscale = "$varScale" \n \n"
      varName="r0_"$varR0"_epsilon_"$varEpsilon"_hrelat_"$varScale
      gmsh -2 -clscale $varScale diskHoleTmp.geo -o maillage_$varName.msh >
      infoMaillage.txt
      ../bin/hxtGinzburg maillage_$varName.msh --epsilon $varEpsilon > temp.txt
      mv directions.pos directions_$varName.pos
      tail -n 2 temp.txt > temp1.txt
      cat temp.txt | awk '{ print $1 }' > temp2.txt
      varEnergy=$(head -1 temp2.txt)
      varSingularities=$(tail -1 temp2.txt)
      echo $varScale";" $varR0";" $varEpsilon";" $varSingularities $varEnergy >>
      sorties.txt
    done
  done
done
done
```

diskHole.sh

Bibliographie

- [1] Grégoire Allaire. *Analyse numérique et optimisation*. Les éditions de l'école polytechnique, 2007.
- [2] S. Antontsev, S. Shmarev, A. Braides, M. del Pino, M. Musso, J. Hernandez, F.J. Mancebo, S. Kichenassamy, A. Lyaghfour, and L.A. Peletier. *Handbook of Differential Equations : Stationary Partial Differential Equations, Volume 3*. Elsevier, 2006.
- [3] K. J. Bathe. *Numerical methods in finite element analysis*. Prentice-Hall, 1976.
- [4] P.-A. Beaufort, C. Geuzaine, and J.-F. Remacle. Understanding directional fields over two-manifolds from Ginzburg-Landau functional. *submitted to Journal of Computational Physics*, 2018.
- [5] P.-A. Beaufort, J. Lambrechts, F. Henrotte, C. Geuzaine, and J.-F. Remacle. Computing cross fields - A PDE approach based on the Ginzburg-Landau theory. *Procedia Engineering*, 203 :219–231, September 2017.
- [6] M. Berger and R. Gostiaux. *Géométrie différentielle : variétés, courbes et surfaces*. P.U.F., 1992.
- [7] F. Bethuel. Variational methods for Ginzburg-Landau equations. *Calculus of Variations and Geometric Evolutions Problem*, pages 1–43, 1999.
- [8] F. Bethuel, H. Brezis, and F. Hélein. Asymptotics for the minimization of a Ginzburg-Landau fonctional. *Calculus of Variations and Partial Differential Equations*, 1(no.2), 1993.
- [9] F. Bethuel, H. Brezis, and F. Hélein. Ginzburg-Landau vortices. *Progress in Nonlinear Differential Equations and their Applications*, vol. 13, 1994.
- [10] F. Boyer and P. Fabrie. *Éléments d'analyse pour l'étude de quelques modèles d'écoulements de fluides visqueux incompressibles*. Springer, 2006.
- [11] H. Brezis. *Analyse fonctionnelle*. Dunod, 2005.
- [12] Claude Le Bris. *Systèmes multi-échelles*. Springer, 2005.
- [13] D.Eppstein. *Twenty proofs of Euler's formula : $V-E+F=2$* . The Geometry Junkyard, 2014. <http://www.ics.uci.edu/~eppstein/junkyard/euler>.
- [14] C. Geuzaine and J.-F. Remacle. Gmsh : a three-dimensional finite element mesh generator with built-in pre- and post- processing facilities. *International Journal for Numerical Methods in Engineering*, 79 (11) :pp. 1309–1331, 2009.
- [15] Roger Godement. *Analyse mathématique III : Fonctions analytiques, différentielles et variétés, surfaces de Riemann*. Springer, 2001.
- [16] Luc Haine. LMAT 2110 - Éléments de géométrie différentielle. Université catholique de Louvain, Chemin du Cyclotron 2, 1348 Louvain-la-Neuve, 2014-2015.
- [17] W. F. Harris. Disclinations. *Scientific American, Inc*, 1977.
- [18] Allen Hatcher. *Algebraic Topology*. Cambridge University Press, 2002.
- [19] Jean Hladik and Pierre-Emmanuel Hladik. *Le calcul tensoriel en physique - Cours et exercices corrigés*. Dunod, 1999.

- [20] Yvette Kosmann-Schwarzbach. *Groups and Symmetries*. Springer, 2010.
- [21] François Laudenbach. *Calcul différentiel et intégral*. Les éditions de l'école polytechnique, 2005.
- [22] Henri Poincaré. Sur la généralisation d'un théorème d'Euler relatif aux polyèdres. *Comptes rendus hebdomadaires des séances de l'Académie des sciences*, 1893.
- [23] Claude Semay and Bernard Silvestre-Brac. *Introduction au calcul tensoriel, Applications à la physique*. Dunod, 2007.
- [24] M. Struwe. On the asymptotic behaviour of minimizers of the Ginzburg-Landau model in 2 dimensions. *Differential Integral Equations*, 8(no. 1) :224, 1995.
- [25] Endre Süli and David Mayers. *An introduction to Numerical Analysis*. Cambridge University Press, 2014.

