



Raffinement local adaptatif et méthodes multiniveaux pour la simulation d'écoulements multiphasiques.

Sebastian Minjeaud

► To cite this version:

Sebastian Minjeaud. Raffinement local adaptatif et méthodes multiniveaux pour la simulation d'écoulements multiphasiques.. Mathématiques [math]. Aix-Marseille Université, 2010. Français. NNT : 2010AIX30051 . tel-00535892

HAL Id: tel-00535892

<https://theses.hal.science/tel-00535892>

Submitted on 13 Nov 2010

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

**THESE DE DOCTORAT DE
L'UNIVERSITE PAUL CEZANNE**

Spécialité :
Mathématiques

Présentée par
Sebastian Minjeaud

Pour obtenir le grade de
**DOCTEUR EN SCIENCES DE L'UNIVERSITÉ
PAUL CEZANNE**

Sujet de thèse
**Raffinement local adaptatif et méthodes multiniveaux
pour la simulation d'écoulements multiphasiques.**

Soutenue publiquement le **27 septembre 2010** devant le jury composé de

Mr Franck	BOYER	Université Paul Cézanne	Directeur
Mr Emmanuel	CREUSÉ	Université Lille 1	Examinateur
Mr Bruno	DESPRÉS	Université Paris 6	Examinateur
Mr Jean-Frédéric	GERBEAU	INRIA Rocquencourt	Rapporteur
Mr Jean-Luc	GUERMOND	Texas A&M University	Rapporteur
Mr Jean-Claude	LATCHÉ	IRSN Cadarache	Examinateur
Mr Bruno	PIAR	IRSN Cadarache	Encadrant

Remerciements

Mes premiers remerciements sont destinés à Franck Boyer et Bruno Piar.

Je remercie Franck pour sa grande disponibilité. Il m'a guidé dans mon travail de thèse. Son soutien et son accompagnement ont été exemplaires. Il a su me faire profiter de ses connaissances, tout en me laissant aborder les choses à ma manière. Nos relations ont souvent dépassé le cadre strictement professionnel, je l'en remercie.

Je remercie Bruno pour son encadrement au jour le jour. Nous avons partagé le même bureau pendant une grande partie de ces trois années de thèse. Ceci nous a permis d'échanger des discussions sur différentes problématiques scientifiques (directement liées à mon sujet de thèse ou non). Il a su me transmettre une partie de ses connaissances. Je le remercie en particulier pour ses explications sur la manière de développer rigoureusement un code informatique de calcul scientifique.

Je remercie Jean-Frédéric Gerbeau et Jean-Luc Guermond d'avoir accepté la tâche d'être rapporteur de mon manuscrit et de l'attention qu'ils ont portée à mon travail.

Je remercie Emmanuel Creusé et Bruno Deprès d'avoir accepté de participer à mon jury.

J'adresse un merci particulier à Jean-Claude Latché. Il a non seulement accepté de faire partie de mon jury mais il a également suivi mon travail tout au long de cette thèse. Il m'a en particulier aidé à encadrer le stage de master 2 de Fanny Dardahlon, en me conseillant et me proposant des sujets non directement liés à mon travail de thèse.

Je remercie Laurence Rigollet de m'avoir accueilli dans son laboratoire. Mes trois années de thèse passées au sein de ce laboratoire ont été très agréables. Je remercie en particulier Céline, Clément, Fabien, Fabrice, Jean, Lionel, Jean-Paul, Marc, Eric pour les différents échanges (scientifiques ou non) qui ont contribué à l'avancement du travail ainsi qu'à instaurer une bonne ambiance. Je remercie également les différents doctorants arrivés dans le laboratoire au cours de ma thèse : Walid, Fanny, Romain, Trung, Raphaël ainsi que ceux déjà présents à mon arrivée : Matthieu, Laura, Aurélien, Guillaume.

Un grand merci également au membre du LATP de m'avoir également accueilli dans leur laboratoire. Je remercie en particulier Florence, Pascal, Claudia, Pierre, Thierry, Raphaële, Guillemette, Emmanuel, Assia, Yannick, Catherine, Nicolas, Sylvie, Michel, Yves, Patricia, Jérôme, François, Marie. Merci, également à Gérard Henry, il m'a permis d'utiliser le matériel informatique du laboratoire dans de très bonnes conditions. Un petit clin d'oeil à tous les thésards : Federico, Sébastien, Jean-Christophe, Ismael, Mickaël, Jonathan, Clément.

Je remercie toute ma famille de m'avoir soutenu et encouragé tout au long de mes études. Enfin, je remercie Stella qui partage ma vie depuis plusieurs années maintenant. Elle a su me soutenir et me supporter pendant tous les moments difficiles de cette thèse. J'ai également une pensée pour sa famille.

Table des matières

Introduction

9

Partie 1

Algorithmes de raffinement local adaptatif et méthodes de préconditionnement multigrille

Chapitre I	Espace éléments finis mult niveau	31
I.1	Notations et définitions.....	31
I.2	Motif de raffinement.....	34
I.3	Espaces d'approximation éléments finis mult niveaux.....	36
I.3.1	Hierarchie d'espaces d'approximation conformes. Relation parents-enfants.	36
I.3.2	Bases mult niveaux et espaces d'approximation mult niveaux.....	41
Chapitre II	Procédure d'adaptation et preconditionneurs multigrilles.	43
II.1	Adaptation.....	43
II.1.1	Procédures de raffinement/déraffinement.....	43
II.1.2	Conservation de l'information.....	46
II.1.3	Procédure d'adaptation.....	48
II.1.4	Règles "au-plus-un-niveau-de-différence".....	49
II.1.5	Intégration numérique et opérateurs de transfert.....	51
II.2	Préconditionneurs mult niveaux.....	53
II.2.1	Méthodes des corrections successives.....	53
II.2.2	Méthodes des corrections successives liées à des sous-espaces.....	53
II.2.3	Méthodes multigrilles.....	56
II.2.4	Algorithme de coarsening d'un espace d'approximation mult niveau.....	58
II.2.5	Préconditionneurs mult niveaux.....	60
II.3	Validation sur un problème modèle stationnaire.....	61
II.3.1	Problème continu.....	61
II.3.2	Maillages initiaux et critère de raffinement.....	62
II.3.3	Raffinement local et preconditionneurs multigrilles.....	62
II.4	Conclusion.....	66
Chapitre III	Implémentation en parallèle dans la librairie PELICANS	67
III.1	Organisation du module de raffinement local.....	68
III.1.1	Éléments de référence, motifs de raffinement et indicateurs de (dé)raffinement.....	68

III.1.2	Champs discrets, cellules et fonctions de base	70
III.1.3	Algorithmes d'adaptation	72
III.2	Numérotation des inconnues dans la librairie PELICANS	79
III.2.1	Numérotation des degrés de liberté d'un champ : la classe <code>PDE_LinkDOF2Unknown</code>	79
III.2.2	Gestion de plusieurs champs : la classe <code>PDE_SystemNumbering</code>	80
III.3	Organisation du module de préconditionnement multineau	80
III.3.1	Algorithme de coarsening : la classe <code>PDE_AlgebraicCoarsener</code>	81
III.3.2	Préconditionneurs multineaux : la classe <code>PDE_GeometricMultilevel_PC</code>	83
III.4	Principe du parallélisme dans la librairie PELICANS	83
III.4.1	La bibliothèque MPI	84
III.4.2	Partitionnement du domaine	86
III.4.3	Numérotation globale des inconnues en parallèle	87
III.4.4	Algèbre linéaire distribuée	89
III.5	Raffinement local en parallèle	90
III.6	Multigrille en parallèle	92
III.7	Conclusion	93

Partie 2

Discretisation d'un modèle de type Cahn-Hilliard/Navier-Stokes (CH/NS)

Chapitre IV	Modèle triphasique consistant de type CH/NS	97
IV.1	Modèle de type Cahn-Hilliard/Navier-Stokes diphasique	97
IV.1.1	Modèle de Cahn-Hilliard diphasique	97
IV.1.2	Couplage aux équations de Navier-Stokes incompressibles	99
IV.2	Modèle de Cahn-Hilliard/Navier-Stokes triphasique	100
IV.2.1	Modèle de Cahn-Hilliard triphasique	100
IV.2.2	Couplage aux équations de Navier-Stokes incompressibles	111
Chapitre V	Discretisation du système de Cahn-Hilliard	115
V.1	Discretisation, existence et convergence des solutions approchées	115
V.1.1	Discretisation en temps	116
V.1.2	Discretisation en espace	116
V.1.3	Equivalence avec un système de deux équations couplées	118
V.1.4	Estimation d'énergie discrète	119
V.1.5	Théorème d'existence et de convergence	121
V.2	Différentes discrétisations pour les termes non linéaires	122
V.2.1	Remarques préliminaires	122
V.2.2	Discretisation implicite de la contribution de F_0	123
V.2.3	Discretisation convexe-concave de la contribution de F_0	125
V.2.4	Discretisation semi-implicite de la contribution de F_0	126
V.2.5	Discretisation semi-implicite de la contribution de P	127
V.2.6	Résumé des résultats	127
V.2.7	Schémas correspondants dans le cas diphasique	128
V.3	Illustrations numériques	130
V.3.1	Cas tests deux phases	130
V.3.2	Cas tests trois phases	136

V.4	Cas d'une mobilité dégénérée	141
V.5	Démonstrations des théorèmes d'existence et de convergence des solutions approchées	143
V.5.1	Démonstration du théorème V.9	144
V.5.2	Preuve du théorème V.10	146
V.6	Conclusion	154
Chapitre VI	Discretisation inconditionnellement stable du système CH/NS	155
VI.1	Discretisation du modèle CH/NS triphasique	155
VI.1.1	Discretisation en temps	155
VI.1.2	Discretisation en espace	160
VI.1.3	Equivalence du système de Cahn-Hilliard avec un système de deux équations	162
VI.2	Schéma correspondant dans le cas diphasique	163
VI.3	Stabilité inconditionnelle du schéma	165
VI.4	Existence de solutions au problème discret	166
VI.5	Convergence des solutions discrètes dans le cas homogène	171
VI.5.1	Bornes sur les solutions discrètes	173
VI.5.2	Argument de compacité, convergence des sous-suites	178
VI.5.3	Passage à la limite dans le schéma	181
VI.6	Conclusion	186
Chapitre VII	Méthode de projection incrémentale	189
VII.1	Éléments finis conformes	191
VII.1.1	Problème de Stokes	191
VII.1.2	Calcul d'un état d'équilibre : $\mathbf{f} = \nabla Q$	193
VII.1.3	Problème de Navier-Stokes à densité variable. Système couplé CH/NS	196
VII.2	Éléments finis non conformes	204
VII.2.1	Discretisation	205
VII.2.2	Opérateur elliptique discret pour la pression et conditions aux bords artificielles	206
VII.2.3	Formulation variationnelle et estimation d'erreur	206
VII.2.4	Tests numériques	207

Partie 3

Expérimentations numériques

Chapitre VIII	Etude de configurations diphasiques	215
VIII.1	Benchmark : une bulle immergée dans un liquide	215
VIII.1.1	Définition du cas test	216
VIII.1.2	Quantités du benchmark	216
VIII.1.3	Paramètres numériques et résultats	218
VIII.2	Forme d'une bulle montant dans un liquide	224
VIII.3	Calcul tridimensionnel	229
Chapitre IX	Etude de configurations triphasiques	233
IX.1	Bulle traversant une interface liquide-liquide	233
IX.1.1	Présentation du cas test	233
IX.1.2	Influence de la valeur de l'épaisseur d'interface ε	234
IX.1.3	Influence de la valeur du nombre de mailles dans l'interface	236
IX.1.4	Influence des schémas en temps sur le système Cahn-Hilliard	238
IX.1.5	Influence de la valeur du pas de temps	239
IX.1.6	Méthode de projection et Lagrangien augmenté	241

IX.1.7	Solveurs itératifs, raffinement local adaptatif et parallélisme	242
IX.2	Calcul tridimensionnel	250

Conclusions et perspectives	253
------------------------------------	------------

Bibliographie	257
----------------------	------------

Introduction

Nous nous intéressons dans ce manuscrit à la simulation numérique d'écoulements incompressibles triphasiques soumis aux tensions de surface à l'aide d'un modèle de type interfaces diffuses. Les différents aspects de la mécanique des fluides numériques sont abordés, de la modélisation aux expérimentations numériques en passant par la discrétisation et l'implémentation informatique. L'essentiel des contributions concerne néanmoins le développement de schémas numériques (discrétisation en espace et en temps) et leur analyse mathématique (stabilité, existence et convergence des solutions approchées).

Ce travail de thèse a été effectué au sein du laboratoire "étude de l'Incendie et développement de Méthodes pour la Simulation et les Incertitudes" (LIMSI) de l'Institut de Radioprotection et de Sécurité Nucléaire (IRSN). Le champ de compétences de l'IRSN couvre l'ensemble des risques liés aux rayonnements ionisants, utilisés dans l'industrie ou la médecine, ou encore les rayonnements naturels. Plus précisément, l'IRSN exerce des missions d'expertise et de recherche dans les domaines suivants :

- surveillance radiologique de l'environnement et intervention en situation d'urgence radiologique,
- radioprotection de l'homme,
- prévention des accidents majeurs dans les installations nucléaires,
- sûreté des réacteurs,
- sûreté des usines, des laboratoires, des transports et des déchets,
- expertise nucléaire de défense.

Le LIMSI est un laboratoire de la Direction de la Prévention des Accidents Majeurs (DPAM). Une part de son activité est constituée par l'étude des différentes situations auxquelles un réacteur nucléaire peut se trouver confronté depuis les conditions normales de fonctionnement jusqu'aux accidents graves qui sont le cadre général de cette thèse.

Une dégradation avancée d'un réacteur nucléaire à eau pressurisée lors d'un hypothétique accident majeur peut conduire, selon les scénarios envisagés, à la formation d'un bain de corium (mélange des matériaux fondus du cœur et de la cuve) dans le puits de cuve, composé de béton, qui constitue la dernière barrière de confinement. Le corium, encore chauffé par le dégagement de puissance résiduelle dû à la désintégration des produits de fission, interagit avec les structures en béton qui le contiennent, et le bain érode peu à peu le radier ainsi que les parois latérales. Cette interaction s'accompagne de relâchements importants de gaz : vaporisation de l'eau contenue dans le béton et formation de dioxyde de carbone principalement par décomposition du calcaire. Le bain est alors traversé par un flux de bulles.

Le corium liquide est un mélange complexe. Nous nous intéressons à une configuration probable du bain de corium dans laquelle deux phases principales, l'une majoritairement oxyde et l'autre majoritairement métallique, se séparent pour atteindre une géométrie stratifiée (dès que l'agitation engendrée par le flux gazeux tombe en deçà d'un certain seuil). Ce phénomène a un impact majeur sur le déroulement de l'accident : la couche métallique, beaucoup plus conductrice, constitue un pont thermique entre la couche oxyde, dans laquelle est générée l'essentiel de la puissance, et les parois ; la progression de l'érosion de la cavité en est fortement affectée ainsi que, en conséquence, les modes et temps de percée du puits de cuve (percée latérale ou verticale). De plus, le flux gazeux influence grandement les transferts entre les deux phases (modification des couches limites thermiques, changements topologiques de l'interface oxyde/métal avec entraînement éventuel du métal) pouvant accélérer l'ablation du béton dans une direction (horizontale ou verticale). La quantification de ces échanges thermiques et massiques reste un problème ouvert préjudiciable à la fiabilité des simulations d'accident actuelles [Cra07].

L'étude des échanges de masse et de chaleur entre deux phases liquides stratifiées lors du passage d'un flux de bulles fait l'objet, au LIMSI, depuis plusieurs années, d'une approche par simulation numérique directe. Un modèle mathématique a été élaboré et étudié au cours de la thèse de C. Lapuerta [BL06, Lap06].

Ce modèle et les différentes difficultés numériques identifiées dans [Lap06] constituent le point de départ du présent travail de thèse. Nous détaillons dans la suite de cette introduction les différentes contributions apportées au cours de ce travail :

- *modélisation* : le modèle comporte différents paramètres “non-objectifs” (*i.e.* qui ne correspondent pas à des propriétés physiques des fluides en présence). Ces paramètres peuvent avoir une influence importante sur les simulations et sont donc délicats à fixer. Nous apportons à travers des études numériques ou théoriques quelques éléments de compréhension nouveaux.
- *discrétisation en espace* : l'approximation des solutions du modèle nécessite une résolution fine (en espace) au voisinage des interfaces. Afin de limiter les coûts de calcul, nous avons mis au point et implémenté des algorithmes de raffinement local adaptatif pour des approximations de type éléments finis conformes de Lagrange.
- *préconditionnement des systèmes linéaires* : dans le cadre de la méthode de raffinement local évoquée ci-dessus nous avons proposé et implémenté un algorithme de “déaffinement” permettant de mettre en place de manière simple la structure nécessaire (*i.e.* une suite de sous-grilles emboîtées) au fonctionnement des preconditionneurs multigrilles géométriques.
- *discrétisation en temps* : différentes discrétisations en temps ont été proposées afin d'obtenir des schémas efficaces et robustes (notamment à grands pas de temps). Ces schémas ont été étudiés numériquement (comparaisons, courbes de convergence...) mais également d'un point de vue théorique (existence et convergence des solutions approchées...).
- *expérimentations numériques* : en parallèle aux travaux théoriques (évoqués ci-dessus) différentes expérimentations numériques ont été réalisées. En particulier, les développements informatiques (parallélisation du raffinement local) ont permis d'accéder aux premières simulations réalisées dans un cadre réellement tridimensionnel (*i.e.* en ne supposant aucune symétrie *a priori*).

Le plan que nous adoptons ci-dessous n'est pas rigoureusement identique à celui du manuscrit. Dans cette introduction, nous présentons tout d'abord le modèle afin de mieux appréhender les difficultés posées par son approximation numérique et ainsi justifier les choix de discrétisation que nous avons effectués ; alors que dans le corps du manuscrit nous avons choisi d'expliquer en première partie la méthode de raffinement local et les techniques de preconditionnement multigrilles de manière déconnectée de la problématique “écoulements incompressibles triphasiques” (exposée dans la seconde partie), puisque ces méthodes ont une portée bien plus générale et sont d'ailleurs maintenant utilisées dans d'autres contextes au sein du laboratoire.

1 Modèle de type Cahn-Hilliard/Navier-Stokes (chapitre IV)

Le modèle repose sur une représentation des interfaces par des zones d'épaisseur strictement positive, certes faible mais supérieure aux épaisseurs réelles. On parle de méthodes à interfaces diffuses (*cf.*, par exemple, [Ell89, BE92, AMW98, Jac99, Boy02, LS03] pour le cas diphasique et [EL91, EG97, GNS00, NGS05, KL05, BL06, GS06] pour le cas trois phases ou plus). Une phase i est décrite géométriquement par une fonction régulière c_i , appelée paramètre d'ordre (que nous prenons ici égale à la fraction volumique de la phase i dans le mélange), valant 1 dans la phase i , 0 en dehors, et variant continûment entre 0 et 1 dans les interfaces entre la phase i et les autres phases. Ainsi, la description de la position des phases nécessite l'introduction d'un paramètre d'ordre c_i par phase, ceux-ci étant néanmoins reliés par la relation $\sum_i c_i = 1$.

Le modèle mathématique est constitué d'un système d'équations aux dérivées partielles de type Cahn-Hilliard/Navier-Stokes. Les équations de Cahn-Hilliard permettent de modéliser la non-miscibilité des phases en maintenant l'épaisseur d'interface à une valeur prescrite ε et permettent également une représentation volumique naturelle des forces capillaires (dûes aux tensions de surface entre les différentes phases). L'hydrodynamique de l'écoulement est prise en compte par le couplage de ces équations au système de Navier-Stokes.

Dans la suite, Ω désigne un ouvert borné, connexe et régulier de $\mathbb{R}^d$ ($d = 2$ ou 3).

1.1 Modèle de Cahn-Hilliard triphasique

Le modèle de Cahn-Hilliard décrit l'évolution du système à travers la minimisation, sous la contrainte de conservation du volume, d'une énergie libre s'exprimant comme la somme de termes capillaires (mettant en jeu les tensions de surfaces) et d'un terme non linéaire F , le potentiel de Cahn-Hilliard :

$$\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(c_1, c_2, c_3) = \int_{\Omega} \frac{12}{\varepsilon} F(c_1, c_2, c_3) + \frac{3}{8} \varepsilon \Sigma_1 |\nabla c_1|^2 + \frac{3}{8} \varepsilon \Sigma_2 |\nabla c_2|^2 + \frac{3}{8} \varepsilon \Sigma_3 |\nabla c_3|^2 dx. \quad (1)$$

Le triplet de paramètres constants $\Sigma = (\Sigma_1, \Sigma_2, \Sigma_3)$ et la forme du potentiel F ont été déterminés de manière à ce que le modèle puisse prendre en compte correctement les valeurs des tensions de surface σ_{12} , σ_{13} et σ_{23} prescrites entre les différents couples de phases et soit "consistant" avec les situations deux-phases (ceci signifie que lorsque l'une des phases est absente, le modèle triphasique dégénère exactement en un modèle de Cahn-Hilliard diphasique, cf [AMW98, Boy02, Jac99, LS03]). Ceci a conduit aux expressions suivantes (cf [BL06, théorème 3.2]) :

$$\Sigma_i = \sigma_{ij} + \sigma_{ik} - \sigma_{jk}, \quad \forall i \in \{1, 2, 3\}, \quad (2)$$

$$F(\mathbf{c}) = \sigma_{12} c_1^2 c_2^2 + \sigma_{13} c_1^2 c_3^2 + \sigma_{23} c_2^2 c_3^2 + c_1 c_2 c_3 (\Sigma_1 c_1 + \Sigma_2 c_2 + \Sigma_3 c_3) + c_1^2 c_2^2 c_3^2 \Lambda(\mathbf{c}), \quad \forall \mathbf{c} \in \mathcal{S}, \quad (3)$$

où $\mathbf{c} = (c_1, c_2, c_3)$ et Λ est une fonction régulière.

Les coefficients $S_i = -\Sigma_i$ définis par (2) sont bien connus dans la littérature physique [RW82]. Le coefficient S_i est appelé coefficient d'étalement de la phase i à l'interface entre la phase j et k . Si S_i est positif (*i.e.* $\Sigma_i < 0$), alors l'étalement est dit total et si S_i est négatif, alors l'étalement est dit partiel (cf figure 1).

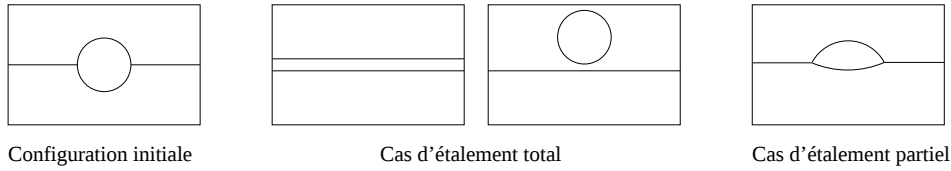


FIG. 1 – Situations d'étalement total et partiel

Nous ne supposons pas que les coefficients Σ_i sont positifs. Le modèle permet donc de prendre en compte certaines situations d'étalement total. Néanmoins il est nécessaire de supposer que la condition suivante est satisfaite :

$$\Sigma_1 \Sigma_2 + \Sigma_1 \Sigma_3 + \Sigma_2 \Sigma_3 > 0. \quad (4)$$

Cette condition garantit que les termes capillaires apportent une contribution positive à l'énergie $\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}$.

La minimisation, sous la contrainte de conservation du volume, de l'énergie libre $\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}$ conduit à l'écriture du système d'équations aux dérivées partielles suivant :

$$\begin{cases} \frac{\partial c_i}{\partial t} = \text{div} \left(\frac{M_0(\mathbf{c})}{\Sigma_i} \nabla \mu_i \right), & \text{pour } i = 1, 2, 3, \\ \mu_i = f_i^F(\mathbf{c}) - \frac{3}{4} \varepsilon \Sigma_i \Delta c_i, & \text{pour } i = 1, 2, 3, \end{cases} \quad (5)$$

où l'inconnue intermédiaire μ_i , appelée potentiel chimique, est la dérivée fonctionnelle de l'énergie $\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}$ par rapport à c_i , $M_0(\mathbf{c})$ est un coefficient de diffusion appelé *mobilité* qui peut éventuellement dépendre de $\mathbf{c}$ et

$$f_i^F(\mathbf{c}) = \frac{4\Sigma_T}{\varepsilon} \sum_{j \neq i} \left(\frac{1}{\Sigma_j} (\partial_i F(\mathbf{c}) - \partial_j F(\mathbf{c})) \right) \text{ avec } \Sigma_T \text{ défini par } \frac{3}{\Sigma_T} = \frac{1}{\Sigma_1} + \frac{1}{\Sigma_2} + \frac{1}{\Sigma_3}.$$

Ce choix de f_i^F , obtenu par l'utilisation d'un multiplicateur de Lagrange, impose que la condition $c_1 + c_2 + c_3 = 1$ est satisfaite à chaque instant.

Nous ajoutons à ce système les conditions aux bords et la condition initiale suivantes : pour $i = 1, 2$ et 3 ,

$$\nabla c_i \cdot \mathbf{n} = 0 \quad \text{et} \quad M_0 \nabla \mu_i \cdot \mathbf{n} = 0, \text{ sur } \partial\Omega, \quad (6)$$

$$\forall i \in \{1, 2, 3\}, \quad c_i(t=0) = c_i^0, \quad (7)$$

où $\mathbf{c}^0 = (c_1^0, c_2^0, c_3^0) \in (H^1(\Omega))^3$ (tel que $c_1^0 + c_2^0 + c_3^0 = 1$ presque partout) est donné.

Dans le cas où la mobilité M_0 est régulière et non dégénérée (*i.e.* $\inf_{\mathbf{c}} M_0(\mathbf{c}) > 0$) et sous la condition que le potentiel F soit positif et à croissance au plus polynomiale (ainsi que ces dérivées), un théorème d'existence de solutions faibles au problème (5) avec la condition initiale (7) et les conditions aux bords (6) a été prouvé dans [BL06]. Nous donnons une nouvelle démonstration de ce résultat pour des conditions aux bords plus générales (conditions mixtes de type de Dirichlet/Neumann pour les paramètres d'ordre) en passant à la limite dans le schéma numérique (*cf* section V.5.1).

Dans le cas d'étalement partiel, il est facile de montrer que les conditions requises sur le potentiel F (défini par (3)) sont satisfaites pour toute fonction Λ constante positive, et en particulier le choix $\Lambda = 0$ convient parfaitement. Cependant, la positivité du potentiel F devient non triviale dans des situations d'étalement total (puisque l'un des Σ_i est alors négatif). Dans [BL06], les auteurs ont démontré que le choix Λ constant suffisamment grand conduisait à un potentiel F positif et par conséquent à un système bien posé. Néanmoins, nous observons dans la pratique que l'influence du paramètre Λ sur le résultat des simulations est importante et sa valeur est par conséquent délicate à fixer. Nous adoptons alors la démarche de modélisation complémentaire suivante :

- nous justifions par une étude numérique que, dans le cas d'étalement partiel, le choix le plus simple $F = F_0$, *i.e.* $\Lambda = 0$, est celui qui convient. L'étude se base sur la simulation d'une lentille piégée entre deux phases stratifiées. Les solutions stationnaires numériques du système Cahn-Hilliard (5) obtenues pour différentes valeurs constantes de Λ sont comparées aux solutions “physiques” (les angles de contact entre les interfaces aux points triples sont donnés en fonction des tensions de surface par la loi de Young).
- nous montrons ensuite que $F_0(\mathbf{c})$ est positif lorsque $\mathbf{c} \in \mathbf{T} = \{\mathbf{c} \in \mathcal{S}, \forall i = 1, 2, 3, 0 \leq c_i \leq 1\}$ et ceci même en situation d'étalement total. Il n'y a donc *a priori* aucune raison que le terme $P = \Lambda c_1^2 c_2^2 c_3^2$ ait une influence dans ce domaine. Néanmoins, il reste indispensable en dehors du domaine $\mathbf{T}$ puisque la fonction F_0 peut y devenir négative.
- ces deux résultats combinés à celui de la proposition IV.11 font émerger l'idée d'utiliser un coefficient Λ dépendant de $\mathbf{c}$ comme fonction de “troncature” pour diminuer (ou supprimer) l'action du terme $c_1^2 c_2^2 c_3^2$ (non nécessaire) sur le domaine $\mathbf{T}$ sans la modifier en dehors ce qui permet de garantir la positivité du potentiel F correspondant.

Ce travail, encore en cours, est effectué en collaboration avec R. Bonhomme¹. Nous présentons les premiers résultats dans le chapitre IV.

Enfin nous terminons cette présentation du modèle de Cahn-Hilliard triphasique en énonçant les principales propriétés du système d'équations (5) :

- le système est indépendant de la numérotation attribuée (arbitrairement) aux phases. Cette propriété n'est pas vérifiée par tous les modèles de la littérature (*cf* par exemple [KL05]).
- la conservation du volume total $\int_{\Omega} c_i dx$ de la phase i au cours du temps est garantie.
- les solutions de ce système vérifient l'égalité d'énergie suivante :

$$\frac{\partial}{\partial t} [\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c})] + \sum_{i=1}^3 \int_{\Omega} \frac{M_0(\mathbf{c})}{\Sigma_i} |\nabla \mu_i|^2 dx = 0. \quad (8)$$

1.2 Couplage aux équations de Navier-Stokes incompressibles

Le couplage du système de Cahn-Hilliard avec les équations de Navier-Stokes incompressibles est réalisé (*cf* [Jac99, Boy02, KL05, LS03, Abe09a, Abe09b]) :

- en ajoutant un terme de transport $\mathbf{u} \cdot \nabla c_i$ dans l'équation d'évolution (première équation du système (5)) de chaque paramètre d'ordre c_i , $i \in \{1, 2, 3\}$,
- en définissant la densité et la viscosité comme des fonctions régulières des paramètres d'ordre $\mathbf{c}$ (constante et égale aux valeurs prescrites ρ_i , η_i dans chacune des phases),
- en ajoutant un terme de force capillaire $\sum_{i=1}^3 \mu_i \nabla c_i$ dans le second membre du bilan de quantité de mouvement (équations de Navier-Stokes).

De plus, nous adoptons une forme non standard des équations de Navier-Stokes. En effet, la densité, en tant que fonction du paramètre d'ordre, ne vérifie pas l'équation de conservation de la masse (un terme de diffusion supplémentaire est présent dans le second membre). Ainsi, il n'est pas possible de déduire le bilan d'énergie cinétique des formes conservatives ou non-conservatives des équations de Navier-Stokes.

¹doctorant encadré par B. Piar (IRSN) et J. Magnaudet (IMFT).

La forme des équations de Navier-Stokes ci-dessous, initialement proposée dans [GQ00], permet de démontrer le bilan d'énergie sans utiliser l'équation de conservation de la masse. Elle repose sur l'égalité suivante :

$$\frac{d}{dt} \int_{\Omega_t} \frac{1}{2} \varrho |\mathbf{u}|^2 dx = \int_{\Omega_t} \left[\sqrt{\varrho} \frac{\partial}{\partial t} (\sqrt{\varrho} \mathbf{u}) + (\varrho \mathbf{u} \cdot \nabla) \mathbf{u} + \frac{\mathbf{u}}{2} \operatorname{div} (\varrho \mathbf{u}) \right] \cdot \mathbf{u} dx,$$

le domaine Ω_t étant un domaine borné régulier arbitraire se déplaçant à la vitesse $\mathbf{u}$ du fluide [BF06].

Le modèle Cahn-Hilliard/Navier-Stokes diphasique considéré est alors constitué par les équations suivantes :

$$\begin{cases} \frac{\partial c_i}{\partial t} + \mathbf{u} \cdot \nabla c_i = \operatorname{div} \left(\frac{M_0(\mathbf{c})}{\Sigma_i} \nabla \mu_i \right), & \forall i = 1, 2, 3, \\ \mu_i = \frac{4\Sigma_T}{\varepsilon} \sum_{j \neq i} \left(\frac{1}{\Sigma_j} (\partial_i F(\mathbf{c}) - \partial_j F(\mathbf{c})) \right) - \frac{3}{4} \varepsilon \Sigma_i \Delta c_i, & \forall i = 1, 2, 3, \\ \sqrt{\varrho(\mathbf{c})} \frac{\partial}{\partial t} (\sqrt{\varrho(\mathbf{c})} \mathbf{u}) + (\varrho(\mathbf{c}) \mathbf{u} \cdot \nabla) \mathbf{u} + \frac{\mathbf{u}}{2} \operatorname{div} (\varrho(\mathbf{c}) \mathbf{u}) - \operatorname{div} (2\eta(\mathbf{c}) D(\mathbf{u})) + \nabla p = \sum_{i=1}^3 \mu_i \nabla c_i + \varrho(\mathbf{c}) \mathbf{g}, \\ \operatorname{div} \mathbf{u} = 0. \end{cases} \quad (9)$$

Nous imposons les conditions aux limites et les conditions initiales suivantes :

$$\nabla c_i \cdot \mathbf{n} = 0, \quad M_0 \nabla \mu_i \cdot \mathbf{n} = 0, \quad \text{et} \quad \mathbf{u} = 0 \text{ sur } \partial\Omega, \quad (10)$$

$$\forall i \in \{1, 2, 3\}, \quad c_i(t=0) = c_i^0, \quad \text{et} \quad \mathbf{u}(t=0) = \mathbf{u}^0, \quad \text{dans } \Omega, \quad (11)$$

où $\mathbf{c}^0 = (c_1^0, c_2^0, c_3^0) \in (H^1(\Omega))^3$ (tel que $c_1^0 + c_2^0 + c_3^0 = 1$ presque partout) et $\mathbf{u}^0 \in (H_0^1(\Omega))^d$ sont donnés.

Toute solution du système (9) vérifie les égalités de bilan suivantes :

– bilan du volume :

$$\frac{\partial}{\partial t} \left[\int_{\Omega} c_i dx \right] = 0.$$

– égalité d'énergie :

$$\frac{\partial}{\partial t} \left[\int_{\Omega} \frac{1}{2} \varrho(\mathbf{c}) |\mathbf{u}|^2 dx + \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}) \right] + \int_{\Omega} 2\eta(\mathbf{c}) |D\mathbf{u}|^2 dx + \sum_{i=1}^3 \int_{\Omega} \frac{M_0(\mathbf{c})}{\Sigma_i} |\nabla \mu_i|^2 dx = \int_{\Omega} \varrho(\mathbf{c}) \mathbf{g} \cdot \mathbf{u} dx.$$

Dans le cas où les trois fluides en présence ont la même densité, nous démontrons le résultat d'existence de solutions faibles suivant, en passant à la limite dans les schémas numériques :

Théorème 1 (Existence de solution faible dans le cas homogène)

Supposons que la mobilité M_0 est régulière et vérifie $\inf_{\mathbf{c}} M_0(\mathbf{c}) > 0$, et que le potentiel de Cahn-Hilliard F est positif et est au plus à croissance polynomiale (ainsi que ces dérivées). Supposons que les densités des trois fluides sont égales, c'est-à-dire $\varrho_1 = \varrho_2 = \varrho_3 = \varrho_0$, $\varrho_0 > 0$. Considérons le problème (9) avec la condition initiale (11) et les conditions aux bords (10). Alors, il existe une solution faible $(\mathbf{c}, \boldsymbol{\mu}, \mathbf{u}, p)$ sur $[0, t_f]$ telle que

$$\begin{aligned} \mathbf{c} &\in L^\infty(0, t_f; (H^1(\Omega))^3) \cap C^0([0, t_f]; (L^q(\Omega))^3), \text{ pour tout } q < 6, \\ \boldsymbol{\mu} &\in L^2(0, t_f; (H^1(\Omega))^3), \\ \mathbf{u} &\in L^\infty(0, t_f; (L^2(\Omega))^3) \cap L^2(0, t_f; (H^1(\Omega))^3), \\ \sum_{i=1}^3 c_i(t, x) &= 1, \text{ pour presque tout } (t, x) \in [0, t_f] \times \Omega. \end{aligned} \quad (12)$$

Le cas non homogène est encore un problème ouvert. Même si l'estimation d'énergie (et l'existence de solutions discrètes) restent vraies dans ce cas là, il est alors plus délicat d'obtenir les estimations donnant par compacité la convergence forte sur la vitesse qui est nécessaire pour passer à la limite dans les termes non linéaires. En effet, les équations de Navier-Stokes comportent alors un terme de la forme $\mathbf{u} \partial_t \varrho$ qui est difficile à contrôler.

Pour terminer la présentation du modèle dont nous souhaitons approcher numériquement les solutions, nous insistons sur les deux points suivants (cf également [Lap06, conclusion p.156]) qui constituent les motivations essentielles des travaux présentés dans la suite :

- Les paramètres d'ordre c_i (solution du système de Cahn-Hilliard) varient brutalement de 0 à 1 dans les interfaces (dont la taille caractéristique est donnée par le petit paramètre ε du modèle). Leur approximation au voisinage des interfaces requiert une finesse de maillage qui devient rédhibitoire si elle est appliquée à l'ensemble du domaine de calcul (notamment en géométrie tridimensionnelle, qui est la seule effectivement pertinente dans des études de simulation directe). Les techniques de raffinement local adaptatif deviennent alors des méthodes de choix, puisqu'elles permettent d'affiner dynamiquement la représentation discrète des inconnues en se focalisant sur les zones sensibles (choix de cellules de petites tailles au voisinage des interfaces), tout en limitant le nombre total de mailles.
- Des difficultés de convergence ont été observées dans la méthode de résolution du système de Cahn-Hilliard (méthode de linéarisation de Newton). Ces difficultés ont été partiellement résolues dans [Lap06] par l'utilisation d'une discrétisation adaptée à la forme des termes d'ordre 6 ($c_1^2 c_2^2 c_3^2$) du potentiel de Cahn-Hilliard. Cette discrétisation a été établie avec pour objectif de supprimer la contribution de ces termes au bilan d'énergie discret. Nous poursuivons ces travaux en nous intéressant :
 - aux discrétisations du terme d'ordre 4 du potentiel,
 - à la stabilité du découplage en temps des systèmes de Cahn-Hilliard et Navier-Stokes.

2 Raffinement local adaptatif et méthodes multigrilles (Partie 1)

2.1 Raffinement local adaptatif (section II.1)

Le problème discret en espace est formulé à l'aide de la méthode des éléments finis : les inconnues sont exprimées comme une combinaison d'un ensemble de fonctions locales, appelées fonctions de base, dont la résolution spatiale est directement déterminée par la taille des mailles qui leur sont associées. Ainsi, pour augmenter la précision dans une zone choisie, il suffit de diviser chaque maille de cette zone en quelques mailles de plus petit diamètre. Cependant, cela peut conduire à un placement non conforme (*cf* figure 2) qui rend délicate la formulation du problème discret. La méthode CHARMS (Conforming Hierarchical Adaptive Refinement MethodS) [KGS03] permet de prendre en compte implicitement les situations non conformes en adoptant le point de vue de "raffinement des fonctions de base".

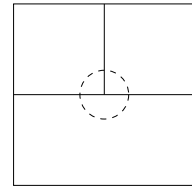


FIG. 2 – Position non conforme des mailles.

Cette approche repose sur la donnée d'une suite d'espaces d'approximation H^1 -conformes emboîtés $X_0 \subset \dots \subset X_J$, $J \geq 1$, engendrés par des ensembles B_j , $j \in \llbracket 0, J \rrbracket$ de fonctions de base de résolution spatiale d'autant plus fine que j est grand. Le raffinement local est alors réalisé en utilisant des espaces d'approximation composites (ou multiniveaux), c'est-à-dire des espaces engendrés par des fonctions de base sélectionnées dans chacun des ensembles B_j , $j \in \llbracket 0, J \rrbracket$ selon la résolution souhaitée dans chacune des parties du domaine. Dans ce cadre, Krysl, Grinspun et Schröder dans [KGS03] (*cf* aussi [GKS02, KTZ04, HK03]) ont proposé des procédures appelées CHARMS qui permettent de raffiner ou déraffiner les espaces d'approximation multiniveaux, c'est-à-dire ajouter ou retirer des fonctions de base des familles engendrant ces espaces tout en préservant leur indépendance linéaire.

La méthode repose sur la propriété fondamentale suivante : puisque $X_j \subset X_{j+1}$, toute fonction de base de B_j s'exprime comme une combinaison linéaire de certaines fonctions de base de B_{j+1} . Ces combinaisons linéaires établissent des relations parents/enfants entre les fonctions de base de deux niveaux consécutifs. Raffiner une fonction de base consiste à retirer la fonction de base et à ajouter ses enfants ; la déraffiner consiste à ajouter la fonction de base et à retirer ses enfants.

Nous donnons une construction précise (*cf* chapitre I) de la suite d'espaces emboîtés $X_0 \subset \dots \subset X_J$ pour des éléments finis conformes de type Lagrange (par exemple $\mathbb{P}_k, \mathbb{Q}_k$, $k \geq 1$). Cette construction repose sur la notion de motif de raffinement (*cf* définition I.11) que nous définissons comme la donnée d'un élément fini de référence (conforme de type Lagrange) et d'un maillage de son support géométrique (*cf* figure 3).

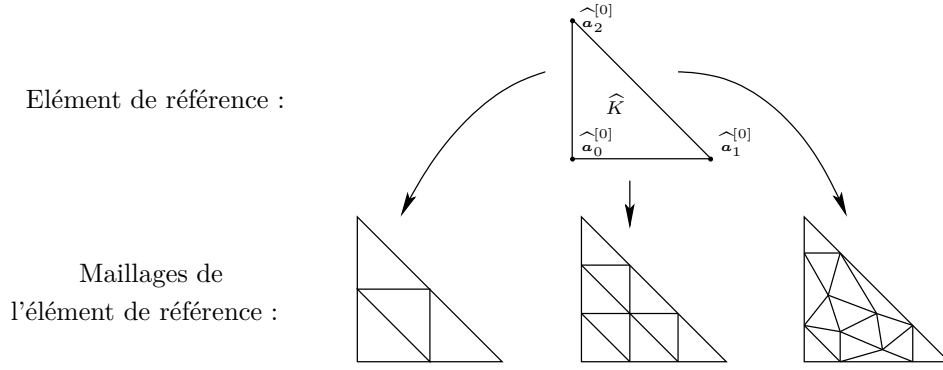


FIG. 3 – Trois exemples de motifs de raffinement associés à l'élément fini de référence Triangle- $\mathbb{P}_1$.

Sous des conditions de compatibilité du motif de raffinement, nous démontrons (*cf* proposition I.16) qu'en appliquant récursivement et uniformément ce dernier à une grille grossière $\mathcal{T}_0$ (non nécessairement structurée), nous construisons une suite de grilles emboîtées géométriquement conformes (*cf* figure 4). L'espace X_j est alors défini comme l'espace d'approximation éléments finis naturellement associé à $\mathcal{T}_j$.

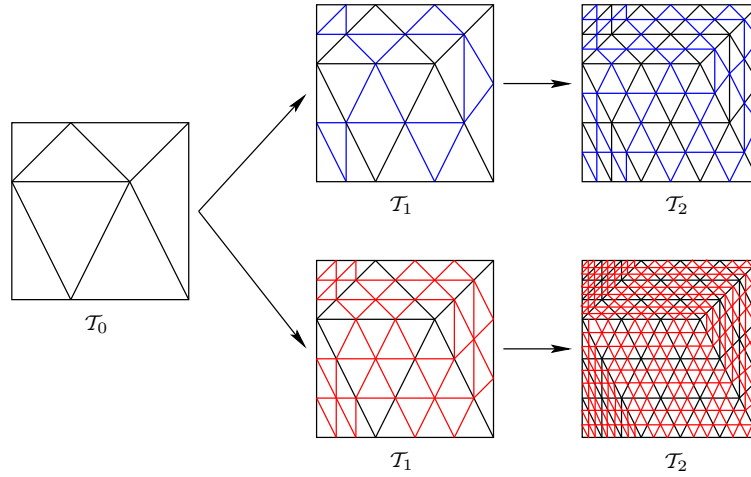


FIG. 4 – Exemples de suites de grilles emboîtées obtenues en appliquant les deux premiers motifs de raffinement présentés sur la figure 3.

Nous étudions ensuite une version modifiée des algorithmes d'adaptation quasi-hiérarchiques donnés dans [KGS03]. La différence essentielle provient de l'application d'une règle pratique "au-plus-un-niveau-de-différence" qui fait partie intégrante des algorithmes de [KGS03]. Nous avons externalisé cette règle, ce qui nous a en particulier permis d'en définir une variante garantissant que la largeur de bande des matrices assemblées reste bornée quelles que soient les étapes d'adaptation (*cf* section II.1.4).

Nous donnons des conditions précises suffisantes (sur les fonctions de bases constituant l'espace multiniveaux à l'initialisation du processus de raffinement) pour garantir que :

- l'algorithme d'adaptation que nous proposons préserve l'indépendance linéaire des fonctions de base sélectionnées sur les différents niveaux (*cf* proposition I.25).
- qu'aucune information n'est perdue lorsque nous raffinons une fonction de base (*cf* théorème II.8). Ceci signifie que, si $\mathcal{B}$ est obtenue à partir de $\mathcal{B}^*$ par le raffinement d'une fonction de base alors $\text{vect } \mathcal{B}^* \subset \text{vect } \mathcal{B}$, *i.e.* la base raffinée $\mathcal{B}$ autorise la représentation exacte de chaque fonction de la base originale $\mathcal{B}^*$.
- les espaces d'approximation obtenus en raffinant (resp. déraffinant) sont indépendants de l'ordre dans lequel les raffinements (resp. déraffinements) successifs sont réalisés (*cf* proposition II.12).

Quelques contre-exemples illustrent le fait que ces propriétés ne sont pas satisfaites dans le cas général (*cf* figures II.1 et II.2).

En outre, lors de la résolution de problèmes instationnaires, il est souvent nécessaire de calculer des intégrales faisant intervenir plusieurs champs discrets n'appartenant pas aux mêmes espaces d'approximation multigrilles.

Par exemple, lorsque le terme instationnaire est discrétisé par la méthode d'Euler, nous sommes amenés à calculer l'intégrale suivante :

$$\int_{\Omega} u_h^n \nu_h dx, \quad (13)$$

où u_h^n représente le champ explicite appartenant à l'espace d'approximation vect $\mathcal{B}^n$ à l'instant t^n et ν_h est une fonction test appartenant à l'espace d'approximation multigrilles vect $\mathcal{B}^{n+1}$ à l'instant t^{n+1} . A cause de l'adaptation de maillage, les deux espaces multigrilles vect $\mathcal{B}^n$ et vect $\mathcal{B}^{n+1}$ sont différents.

Nous montrons qu'il n'est pas nécessaire d'utiliser des opérateurs de transfert pour effectuer le calcul de telles intégrales. Ceci peut être réalisé en utilisant des règles de quadrature sur un maillage multigrille tenant compte des différences de supports des fonctions de base intervenant dans l'intégrale à calculer (*cf* section II.1.5).

2.2 Méthodes de préconditionnement multigrilles (section II.2)

Afin d'accélérer la résolution du problème discret, nous exploitons la structure multigrille créée par l'algorithme de raffinement local pour construire des préconditionneurs multigrilles [BZ00]. Ils permettent en effet d'obtenir un taux de convergence des solveurs linéaires indépendant du nombre d'inconnues et sont donc particulièrement attractifs.

Les méthodes itératives classiques (méthode de Richardson, méthode de Jacobi relaxée, méthode de Gauss-Seidel, *cf* [SVdV00]) sont peu performantes mais possèdent la qualité d'être de bons lisseurs. Ceci signifie qu'en quelques itérations elles permettent d'éliminer les hautes fréquences de l'erreur, la convergence des basses fréquences étant très lente.

Le principe des méthodes multigrilles (*cf* [Hac85, Wes92, TOS01]) est alors de conjuguer le pouvoir lissant de ces méthodes peu coûteuses à une correction effectuée sur une grille plus grossière (mais néanmoins suffisante pour corriger les basses fréquences de l'erreur). Ce principe peut être appliqué récursivement pour obtenir un algorithme utilisant plusieurs grilles.

Les méthodes multigrilles reposent sur les trois ingrédients suivants :

- opérateurs de transfert (projection et interpolation) entre les différentes grilles,
- algorithmes de lissage sur chacune des grilles,
- solveur (exact ou approché) sur la grille la plus grossière.

Le déroulement d'un cycle multigrille (plusieurs variantes existent nous parlons ici de *V-cycle*) est illustré par la figure 5.

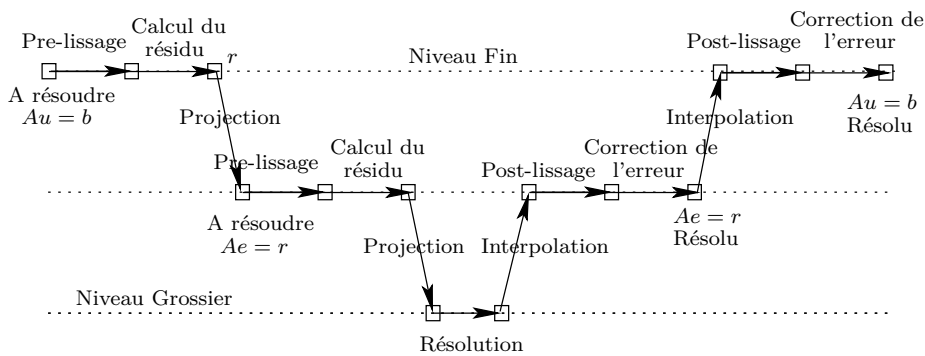


FIG. 5 – Fonctionnement du V-cycle

Nous définissons un algorithme (*cf* section II.2.4) permettant de reconstruire à partir d'un espace d'approximation éléments finis composite (contenant plusieurs niveaux de raffinement), obtenu grâce à la méthode de raffinement local décrite ci-dessus, une suite d'espaces emboîtés auxiliaires (*cf* figure 6). Ceci autorise alors à entrer dans le cadre abstrait multigrilles développé dans [BZ00]; les opérateurs de transfert entre les grilles étant déduits des relations parents-enfants de la méthode CHARMS.

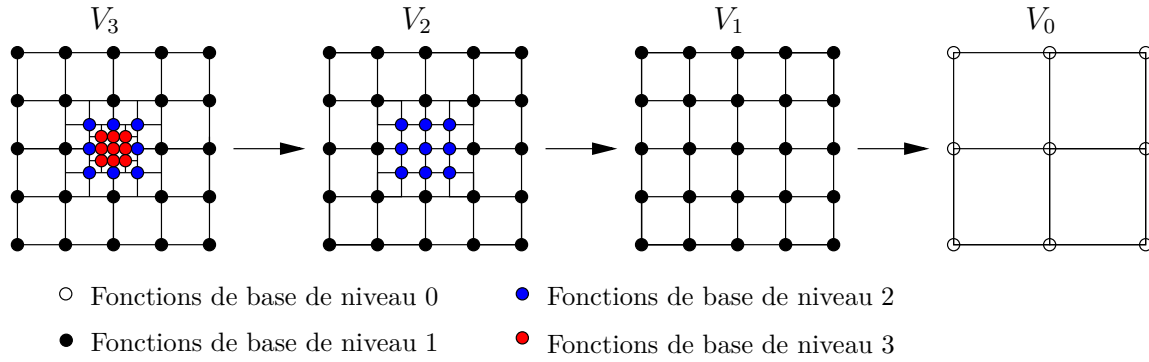


FIG. 6 – Exemple de construction de grille auxiliaire par déraffinement.

Nous utilisons les méthodes multigrilles obtenues comme préconditionneurs dans les méthodes itératives de Krylov (gradient conjugué, GMRES) pour la résolution du système de Cahn-Hilliard ainsi que pour la résolution de l'étape du calcul de l'incrément de pression de la méthode de projection incrémentale (pour la résolution du système de Navier-Stokes).

2.3 Implémentation dans la librairie PELICANS (chapitre III)

Les algorithmes de raffinement local et la méthode multigrille présentés ci-dessus ont été implémentés dans la plate-forme parallèle PELICANS (Plateforme Evolutive de LIBrairies de Composants pour l'Analyse Numérique et la Simulation). Cette librairie développée en C++ au sein du LIMSI fournit un ensemble de fonctionnalités pour faciliter le développement de logiciels de calcul scientifique. Elle est distribuée sous licence libre et est intégralement téléchargeable à l'adresse : <https://gforge.irsnn.fr/gf/project/pelicans>.

Les simulations en géométrie tridimensionnelle (non axisymétrique) sont particulièrement coûteuses en ressources informatiques (place mémoire et temps CPU). L'introduction des techniques de calcul parallèle dans les modules de raffinement local et préconditionneurs multigrilles a permis une exécution du code sur des systèmes à mémoire distribuée.

Le domaine de calcul est partitionné en plusieurs sous-domaines, chacun d'entre eux étant affecté à un processus. Chaque processus ne gère alors que les données relatives à la partie qui lui est associée. Les échanges nécessaires à la résolution globale du problème sont organisés grâce à l'utilisation de bibliothèques de communication par passage de message (MPI). La figure 7 montre un exemple de calcul effectué sur quatre processus, chacun ne sauvegardant que la partie du domaine qui lui est affectée.

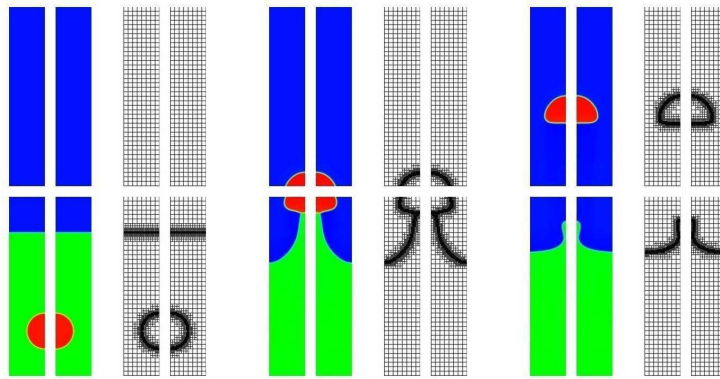


FIG. 7 – Exécution d'un calcul sur quatre processus

Le chapitre III de ce manuscrit décrit la structure actuelle des modules de raffinement local et préconditionneurs multiniveaux ainsi que leur fonctionnement en parallèle. L'objectif visé est double : à la fois insister sur cet aspect qui représente une partie importante de travail mais aussi faciliter les développements ultérieurs qui concerneraient ces modules.

3 Résolution numérique du système CH/NS (Partie 2)

Nous décrivons maintenant dans les trois sections qui suivent les travaux effectués sur la discrétisation en temps du modèle de Cahn-Hilliard/Navier-Stokes. Nous présentons de manière simplifiée les schémas en omettant les notations liées à la discrétisation en espace, mais il est important de mentionner que tous les résultats théoriques (estimations de stabilité, théorème d'existence et de convergence des solutions approchées) ont été établis sur le système complètement discret (en espace et en temps). La discrétisation en espace considérée est une discrétisation éléments finis conformes, nous supposons également que la condition inf-sup est satisfaite par les espaces d'approximation de la vitesse et de la pression. Par ailleurs, tous les schémas proposés sont utilisables dans le cas d'une adaptation du maillage (modifications des espaces d'approximation d'une itération en temps à la suivante) même si les résultats théoriques n'en tiennent pas compte (il est à noter qu'en pratique dans nos simulations le nombre de niveaux de raffinement reste constant au fil des itérations en temps).

3.1 Discrétisation en temps des équations de Cahn-Hilliard (chapitre V)

La principale difficulté provient de la non-convexité du potentiel de Cahn-Hilliard F . Nous proposons donc d'effectuer l'étude d'un schéma du type : pour $i = 1, 2, 3$,

$$\begin{cases} \frac{c_i^{n+1} - c_i^n}{\Delta t} = \operatorname{div} \left(\frac{M_0(\mathbf{c}^n)}{\Sigma_i} \nabla \mu_i^{n+1} \right), \\ \mu_i^{n+1} = D_i^F(\mathbf{c}^n, \mathbf{c}^{n+1}) - \frac{3}{4} \varepsilon \Sigma_i [(1 - \beta) \Delta c_i^n + \beta \Delta c_i^{n+1}], \end{cases} \quad (14)$$

où les fonctions (c_1^n, c_2^n, c_3^n) telles que $c_1^n + c_2^n + c_3^n = 1$ sont données, le réel β est compris entre 0.5 et 1 et

$$D_i^F(\mathbf{a}^n, \mathbf{a}^{n+1}) = \frac{4\Sigma_T}{\varepsilon} \sum_{j \neq i} \left(\frac{1}{\Sigma_j} (d_i^F(\mathbf{a}^n, \mathbf{a}^{n+1}) - d_j^F(\mathbf{a}^n, \mathbf{a}^{n+1})) \right), \quad \forall (\mathbf{a}^n, \mathbf{a}^{n+1}),$$

la discrétisation d_i^F de la i^{e} dérivée partielle $\partial_i F$ de la fonction F restant encore à préciser. Les solutions (éventuelles) de ce schéma vérifient l'estimation d'énergie suivante :

$$\begin{aligned} \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}^{n+1}) - \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}^n) + \Delta t \sum_{i=1}^3 \int_{\Omega} \frac{M_0(\mathbf{c}^n)}{\Sigma_i} |\nabla \mu_i^{n+1}|^2 dx \\ + \frac{3}{8} (2\beta - 1) \varepsilon \int_{\Omega} \sum_{i=1}^3 \Sigma_i |\nabla c_i^{n+1} - \nabla c_i^n|^2 dx \\ = \frac{12}{\varepsilon} \int_{\Omega} [F(\mathbf{c}^{n+1}) - F(\mathbf{c}^n) - \mathbf{d}^F(\mathbf{c}^n, \mathbf{c}^{n+1}) \cdot (\mathbf{c}^{n+1} - \mathbf{c}^n)] dx, \end{aligned} \quad (15)$$

où $\mathbf{d}^F(\cdot, \cdot)$ est le vecteur $(d_i^F(\cdot, \cdot))_{i=1,2,3}$.

Cette estimation est la contrepartie discrète de l'estimation (8). Cependant, nous constatons la présence de deux termes additionnels et, en conséquence, la validité de la décroissance de l'énergie (et des estimations *a priori* que nous pourrions déduire) au niveau discret dépend du signe de ces termes :

- Le dernier terme du membre de gauche de (15) est un terme standard de diffusion numérique dû à la discrétisation en temps de “ Δc_i ” de la seconde équation de (5). Ce terme a le “bon signe” puisque $\beta \geq 0.5$ et peut être supprimé en prenant $\beta = 0.5$.
- Le membre de droite de l'égalité (15) contient la discrétisation en temps $\mathbf{d}^F$ des termes non linéaires et, par suite, son signe dépend de la définition adoptée.

Ainsi, le choix de la discrétisation en temps $\mathbf{d}^F$ des termes non linéaires peut être guidé par une étude du membre de droite de l'égalité (15). Lorsque le membre de droite a le “bon signe”, *i.e.* est négatif, il est possible de l'éliminer pour obtenir une inégalité d'énergie. Les schémas sont donc écrits dans le but d'obtenir :

$$F(\mathbf{a}^{n+1}) - F(\mathbf{a}^n) - \mathbf{d}^F(\mathbf{a}^n, \mathbf{a}^{n+1}) \cdot (\mathbf{a}^{n+1} - \mathbf{a}^n) \leq 0, \quad \forall (\mathbf{a}^n, \mathbf{a}^{n+1}). \quad (16)$$

Puisque le potentiel de Cahn-Hilliard est non convexe, le schéma implicite classique ($\mathbf{d}^F(\mathbf{a}^n, \mathbf{a}^{n+1}) = \nabla F(\mathbf{a}^{n+1})$) ne satisfait pas une telle inégalité. Dans le cas d'étalement partiel, il est néanmoins possible de démontrer que la décroissance de l'énergie est satisfaite pour des petits pas de temps. L'équation (16) suggère d'explicitier la

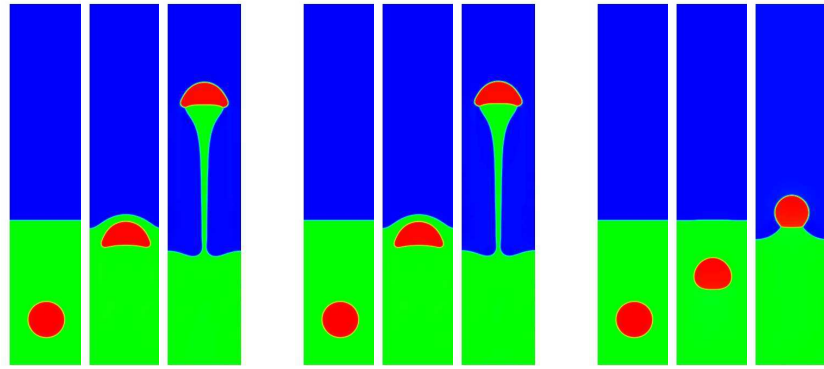
partie concave du potentiel (*cf* [Eyr93, Eyr98, KKL04b, KKL04a]). Ceci donne lieu à l'écriture d'un premier schéma (appelé schéma convexe-concave dans la suite) :

$$\mathbf{d}^F(\mathbf{a}^n, \mathbf{a}^{n+1}) = \nabla F^+(\mathbf{a}^{n+1}) + \nabla F^-(\mathbf{a}^n),$$

où F^+ et F^- sont respectivement des fonctions convexe et concave telle que $F = F^+ + F^-$. Pour ce schéma, l'estimation d'énergie est valable pour tout pas de temps mais les études numériques montrent qu'une forte erreur de troncature est introduite.

Nous proposons une discrétisation semi-implicite (*cf* section V.2.4), appelée "schéma semi-implicite" dans la suite, qui permet à la fois d'assurer la décroissance de l'énergie pour tout pas de temps et de limiter l'erreur de troncature. Cette discrétisation est basée sur des manipulations algébriques du polynôme F permettant d'obtenir (16).

Nous donnons une étude numérique détaillée (section V.3) permettant la comparaison de ces différents schémas. Le schéma semi-implicite représente un bon compromis entre précision et robustesse. Ceci est encore illustré dans la section IX.1 où des simulations (*cf* également figure 8) montrent cette fois l'influence que peut avoir la discrétisation du terme non-linéaire du système de Cahn-Hilliard sur les résultats obtenus pour la simulation d'une montée de bulle à l'aide du modèle couplé Cahn-Hilliard/Navier-Stokes. L'erreur de troncature commise dans l'approximation du système de Cahn-Hilliard par le schéma convexe-concave se manifeste par une telle sous-estimation de la vitesse de montée de bulle que ce dernier semble inutilisable pour ce type de simulations.



(a) Schéma implicite (b) Schéma semi-implicite (c) Schéma convexe-concave

FIG. 8 – Comparaison des différents schémas sur une simulation de montée de bulle aux mêmes instants physiques.

Il est également important de noter que l'estimation d'énergie permet de déduire des informations *a priori* sur les inconnues discrètes qui s'avèrent suffisantes pour démontrer leur existence (*cf* théorème V.9), et leur convergence (*cf* théorème V.10) vers la solution du modèle de Cahn-Hilliard. Les résultats théoriques sont présentés de manière synthétique dans le tableau 1. L'absence de résultats théoriques pour le schéma implicite dans le cas d'étalement total se confirme numériquement par une non convergence de la méthode de résolution du problème discret.

Schéma	Implicite	Convexe-concave	Semi-implicite
Etalement partiel	Décroissance énergie $\Delta t \leq \Delta t_0$	Décroissance énergie $\forall \Delta t$	Décroissance énergie $\forall \Delta t$
	Existence $\forall \Delta t$	Existence $\forall \Delta t$	Existence $\forall \Delta t$
	Convergence	Convergence	Convergence
Etalement total	Problèmes ouverts	Décroissance énergie $\forall \Delta t$	Décroissance énergie $\forall \Delta t$
		Existence $\forall \Delta t$	Existence $\forall \Delta t$
		Convergence	Convergence

TAB. 1 – Résultats théoriques obtenus pour les différents schémas.

Nous n'avons pas abordé d'un point de vue théorique les cas où la mobilité est dégénérée. Nous proposons néanmoins une discrétisation qui permet de contourner les difficultés numériques liées au fait que, dans ces cas là, la mobilité s'annule dans les phases pures. Dans [Lap06], l'auteur proposait d'ajouter une partie constante à la mobilité dégénérée :

$$M_0(\mathbf{c}) = M_{\text{cst}} + M_{\text{deg}} \prod_{i=1}^3 (1 - c_i)^2,$$

les coefficients M_{cst} et M_{deg} étant des constantes positives (éventuellement nulles). Le coefficient M_{cst} doit alors être inférieur, de quelques ordres de grandeur, au coefficient M_{deg} mais doit rester suffisamment grand pour que les difficultés numériques ne soient pas ressenties. Il peut alors être difficile d'ajuster la valeur de ce coefficient lorsque M_{deg} est par exemple de l'ordre de 10^{-7} .

Nous adaptons une idée trouvée dans la référence [BB99] au modèle triphasique considéré dans ce manuscrit : remplacer dans (14) le terme $\text{div} \left(M_0(\mathbf{c}^n) \nabla \mu_i^{n+1} \right)$ par le terme $\text{div} \left(|M_0|_\infty \nabla \mu_i^{n+1} + (M_0(\mathbf{c}^n) - |M_0|_\infty) \nabla \mu_i^n \right)$, où $|M_0|_\infty$ représente une constante supérieure à $\sup_{x \in \Omega} |M_0(\mathbf{c}^n(x))| > 0$.

Le point clé est que, d'un point de vue numérique, à chaque itération de la méthode de linéarisation (méthode de Newton), la matrice des systèmes linéaires à résoudre est exactement la même que celle obtenue lorsque la mobilité est constante de valeur $|M_0|_\infty$. Nous montrons que le schéma ainsi obtenu est encore stable (cf estimation (V.54)).

3.2 Discrétisation du système couplé Cahn-Hilliard/Navier-Stokes (chapitre VI)

Les derniers développements théoriques effectués au cours de cette thèse ont conduit à l'écriture d'un schéma numérique inconditionnellement stable pour la résolution du modèle Cahn-Hilliard/Navier-Stokes. Celui-ci autorise une résolution découplée des systèmes discrets de Cahn-Hilliard et Navier-Stokes. A notre connaissance, un tel résultat n'existe pas dans la littérature (y compris dans le cas diphasique). Le schéma comporte un terme de stabilisation d'ordre Δt et s'écrit de la manière suivante :

$$\begin{cases} \frac{c_i^{n+1} - c_i^n}{\Delta t} + \text{div} \left([c_i^n - \alpha_i] \left[\mathbf{u}^n - \frac{\Delta t}{\varrho^n} \sum_{j=1}^3 (c_j^n - \alpha_j) \nabla \mu_j^{n+1} \right] \right) = \text{div} \left(\frac{M_0(\mathbf{c}^n)}{\Sigma_i} \nabla \mu_i^{n+1} \right), \\ \mu_i^{n+1} = D_i^F(\mathbf{c}^n, \mathbf{c}^{n+1}) - \frac{3}{4} \Sigma_i \varepsilon \Delta c_i^{n+\beta}, \end{cases} \quad (17)$$

avec α_j une constante : $\alpha_j = \int_{\Omega} c_j^0 dx$, et

$$\begin{cases} \varrho^n \frac{\mathbf{u}^{n+1} - \mathbf{u}^n}{\Delta t} + \frac{1}{2} \frac{\varrho^{n+1} - \varrho^n}{\Delta t} \mathbf{u}^{n+1} \\ \quad + (\varrho^{n+1} \mathbf{u}^n \cdot \nabla) \mathbf{u}^{n+1} + \frac{\mathbf{u}^{n+1}}{2} \text{div} (\varrho^{n+1} \mathbf{u}^n) + \text{div} (2\eta^{n+1} D\mathbf{u}^{n+1}) \\ \quad + \nabla p^{n+1} = \varrho^{n+1} \mathbf{g} + \sum_{j=1}^3 (c_j^n - \alpha_j) \nabla \mu_j^{n+1}, \\ \text{div} (\mathbf{u}^{n+1}) = 0, \end{cases} \quad (18)$$

les conditions aux bords étant celles données pour le problème continu.

Ce schéma permet de conserver au niveau discret les principales propriétés du modèle continu (conservation du volume, somme des paramètres d'ordre égale à 1). En outre, nous démontrons l'inégalité d'énergie suivante :

$$\begin{aligned} & \left[\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}^{n+1}) + \frac{1}{2} \int_{\Omega} \varrho^{n+1} |\mathbf{u}^{n+1}|^2 dx \right] - \left[\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}^n) + \frac{1}{2} \int_{\Omega} \varrho^n |\mathbf{u}^n|^2 dx \right] \\ & + \Delta t \sum_{i=1}^3 \int_{\Omega} \frac{M_0(\mathbf{c}^n)}{\Sigma_i} |\nabla \mu_i^{n+1}|^2 dx + \Delta t \int_{\Omega} 2\eta^{n+1} |D\mathbf{u}^{n+1}|^2 dx \\ & + \frac{3}{8} (2\beta - 1) \varepsilon \int_{\Omega} \sum_{i=1}^3 \Sigma_i |\nabla c_i^{n+1} - \nabla c_i^n|^2 dx + \frac{1}{2} \int_{\Omega} \varrho^n [|\mathbf{u}^{n+1} - \mathbf{u}^*|^2 + |\mathbf{u}^* - \mathbf{u}^n|^2] dx \\ & = \Delta t \int_{\Omega} \varrho^{n+1} \mathbf{g} \cdot \mathbf{u}^{n+1} dx + \frac{12}{\varepsilon} \int_{\Omega} [F(\mathbf{c}^{n+1}) - F(\mathbf{c}^n) - \mathbf{d}^F(\mathbf{c}^n, \mathbf{c}^{n+1}) \cdot (\mathbf{c}^{n+1} - \mathbf{c}^n)] dx, \end{aligned} \quad (19)$$

où $\mathbf{d}^F(\cdot, \cdot)$ est le vecteur $(d_i^F(\cdot, \cdot))_{i=1,2,3}$ et $\mathbf{u}^* = \mathbf{u}^n - \frac{\Delta t}{\varrho^n} \sum_{j=1}^3 (c_j^n - \alpha_j) \nabla \mu_j^{n+1}$.

Nous démontrons que ce schéma admet des solutions et que, lorsque les trois fluides ont les mêmes densités, les solutions approchées convergent vers une solution faible du modèle Cahn-Hilliard/Navier-Stokes (cf [Fen06, KSW08] dans le cas deux phases).

Théorème 2 (Existence de solutions discrètes, théorème VI.15)

Etant donné $\mathbf{c}_h^n, \mathbf{u}_h^n$, supposons que

- *la mobilité M_0 est régulière et vérifie $\inf_{\mathbf{c}} M_0(\mathbf{c}) > 0$, et que le potentiel de Cahn-Hilliard F est positif et est au plus à croissance polynomiale (ainsi que ses dérivées).*
- *la discrétisation des termes non-linéaires $\mathbf{d}^F$ satisfait une inégalité du type (16).*

Alors, il existe au moins une solution $(\mathbf{c}_h^{n+1}, \mu_h^{n+1}, \mathbf{u}_h^{n+1}, p_h^{n+1})$ au problème discret en espace associé à (17)-(18) .

Pour tout $N \in \mathbb{N}$, nous pouvons introduire les fonctions du temps $t \in [0, t_f]$ suivantes :

$$\begin{aligned} c_{ih}^{\Delta t}(t, \cdot) &= \frac{t_{n+1} - t}{\Delta t} c_{ih}^n(\cdot) + \frac{t - t_n}{\Delta t} c_{ih}^{n+1}(\cdot), & \text{si } t \in]t_n, t_{n+1}[. \\ \mu_{ih}^{\Delta t}(t, \cdot) &= \mu_{ih}^{n+1}(\cdot), & \text{si } t \in]t_n, t_{n+1}[. \\ \mathbf{u}_h^{\Delta t}(t, \cdot) &= \frac{t_{n+1} - t}{\Delta t} \mathbf{u}_h^n(\cdot) + \frac{t - t_n}{\Delta t} \mathbf{u}_h^{n+1}(\cdot), & \text{si } t \in]t_n, t_{n+1}[. \end{aligned}$$

Théorème 3 (Théorème de convergence, théorème VI.18)

Nous supposons que les hypothèses du théorème 2 sont satisfaites, de manière qu'une solution $(\mathbf{c}_h^{\Delta t}, \mu_h^{\Delta t}, \mathbf{u}_h^{\Delta t}, p_h^{\Delta t})$ au problème discret en espace associé à (17)-(18) existe pour tout $\Delta t > 0$ et pour tout $h > 0$.

Alors, il existe une solution faible $(\mathbf{c}, \mu, \mathbf{u}, p)$ définie sur $[0, t_f]$ telle que (12). De plus, si la famille de maillages est régulière et quasi-uniforme, alors les suites $(\mathbf{c}_h^{\Delta t})$, $(\mu_h^{\Delta t})$ et $(\mathbf{u}_h^{\Delta t})$ satisfont, à sous-suite près, les convergences suivantes lorsque $h \rightarrow 0$, $\Delta t \rightarrow 0$:

$$\begin{aligned} \mathbf{c}_h^{\Delta t} &\rightarrow \mathbf{c} & \text{dans } C^0(0, t_f, (L^q)^3) \text{ fort, pour tout } q < 6, \\ \mathbf{u}_h^{\Delta t} &\rightarrow \mathbf{u} & \text{dans } L^2(0, t_f, (L^2)^d) \text{ fort,} \\ \mu_h^{\Delta t} &\rightharpoonup \mu & \text{dans } L^2(0, t_f, (H^1)^3) \text{ faible.} \end{aligned} \tag{20}$$

3.3 Méthode de projection incrémentale (Chapitre VII)

En pratique, plutôt que de résoudre directement le système discret (18) (par une méthode de type Lagrangien augmenté par exemple), nous utilisons la méthode de projection incrémentale [God79], moins coûteuse en temps de calcul. Il s'agit d'un découplage "prédiction-correction" du système. La première étape consiste à résoudre l'équation de bilan de quantité de mouvement en explicitant la pression et laissant temporairement la contrainte d'incompressibilité de côté. Dans une seconde étape, la vitesse prédite est projetée sur l'espace des fonctions à divergence nulle. Classiquement, l'étape de projection est effectuée en trois sous-étapes : le calcul de l'incrément de pression, puis les corrections de pression et vitesse. Sur le problème de Stokes, la méthode s'écrit de la manière suivante :

- **Etape 1** : Prédiction de vitesse

$$\frac{\tilde{\mathbf{u}}^{n+1} - \mathbf{u}^n}{\Delta t} - \Delta \tilde{\mathbf{u}}_h^{n+1} + \nabla p_h^n = \mathbf{f}^{n+1}. \tag{21}$$

- **Etape 2.1** : Calcul de l'incrément de pression

$$-\Delta \Phi^{n+1} = -\frac{1}{\Delta t} \operatorname{div}(\tilde{\mathbf{u}}^{n+1}). \tag{22}$$

– **Etape 2.2** : Correction de la pression

$$p_h^{n+1} = p_h^n + \Phi_h^{n+1}. \quad (23)$$

– **Etape 2.3** : Correction de la vitesse

$$\mathbf{u}^{n+1} = \tilde{\mathbf{u}}^{n+1} - \Delta t \nabla \Phi_h^{n+1}. \quad (24)$$

L'utilisation du raffinement local implique quelques modifications de ce schéma. Du fait de l'évolution des espaces d'approximation consécutive à l'adaptation de maillage, l'étape de correction de pression (23) n'a plus de sens clair puisque les deux pressions en jeu ne sont pas calculées dans les mêmes espaces d'approximation. Il est alors nécessaire d'introduire une sous-étape supplémentaire. Nous avons choisi de l'effectuer en début d'algorithme (étape de prédiction de pression, cf [GQ00]) afin de pouvoir remplacer la pression explicite par une de ses projections $\tilde{p}^{n+1}$ permettant de préserver les inégalités d'énergie et par conséquent la stabilité.

Par ailleurs, lors du couplage de ce schéma aux équations de Cahn-Hilliard, nous avons été confrontés à la problématique des courants parasites (vitesses de faible amplitude localisées au voisinage de l'interface, cf [SZ99, JTB02]). Celle-ci s'est avérée liée au fait que la méthode de projection ne permet pas de résoudre exactement le système de Navier-Stokes lorsque le second membre du bilan de quantité de mouvement s'écrit comme le gradient d'une fonction de l'espace d'approximation des pressions. Nous avons proposé une variante de la méthode permettant de corriger ce problème. Celle-ci consiste à tenir compte des variations du second membre dans l'étape de prédiction de pression (calcul de $\tilde{p}^{n+1}$) évoquée ci-dessus en l'écrivant de la manière suivante :

$$-\Delta \tilde{p}^{n+1} = -\Delta p^n + \operatorname{div}(\mathbf{f}^{n+1} - \mathbf{f}^n).$$

Ce principe est appliqué au modèle Cahn-Hilliard/Navier-Stokes dans la section VII.1.3.

En marge de ce travail, j'ai eu l'occasion au cours de ma thèse d'encadrer le stage de Master 2 de F. Dardhalon. Nous avons abordé, en collaboration avec J.C. Latché, l'étude de la méthode de projection dans un autre contexte discret : les éléments finis non conformes de bas degré. Ce type de discrétisation est utilisé dans le code de calcul ISIS développé au LMSI pour la simulation d'incendie. Les résultats obtenus sont résumés dans la section VII.2 (cf également [DLM10a, DLM10b]).

4 Expérimentations numériques (partie 3)

Nous présentons dans cette dernière partie quelques expérimentations numériques effectuées en parallèle aux travaux plus théoriques présentés ci-dessus.

Tout d'abord, nous reprenons un cas test proposé dans le benchmark [HTK⁺09]. Ceci nous permet d'étudier l'influence des paramètres non objectifs du modèle à savoir la mobilité et l'épaisseur d'interface. Nous constatons que l'épaisseur d'interface a dans ce cas là peu d'influence sur les résultats obtenus. La valeur du coefficient de mobilité en revanche peut considérablement influencer sur la vitesse de montée de bulle ainsi que sur la forme des bulles. Néanmoins nous observons qu'il existe une plage de valeurs pour lesquelles les résultats obtenus sont très similaires et proches des valeurs de référence fournies par les résultats du benchmark.

Nous montrons ensuite qu'il est possible de simuler de nombreux régimes d'écoulements diphasiques. Les résultats que nous obtenons (cf figure 9) sont à comparer avec les simulations effectuées dans [BM07] ou avec la classification expérimentale des formes de bulles effectuée dans [CGW78].

Enfin, nous montrons que les différents développements réalisés au cours de cette thèse ont rendu possibles des simulations de montées de bulle dans un cadre vraiment tridimensionnel sans supposer de symétrie *a priori* (cf figures 10 pour une simulation diphasique avec plusieurs bulles et 11 pour une simulation cas triphasique).

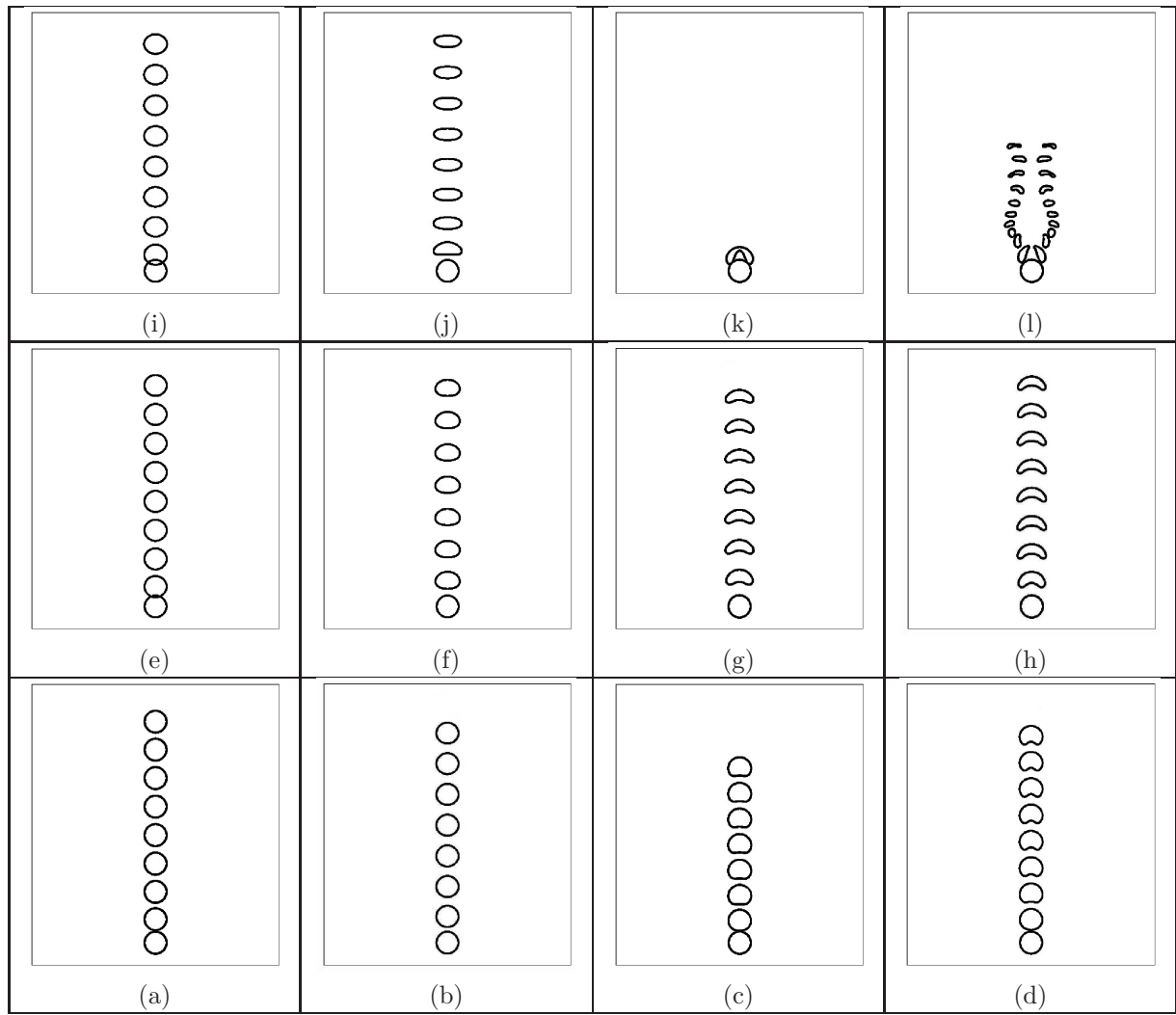


FIG. 9 – Ascension d'une bulle dans un liquide incompressible

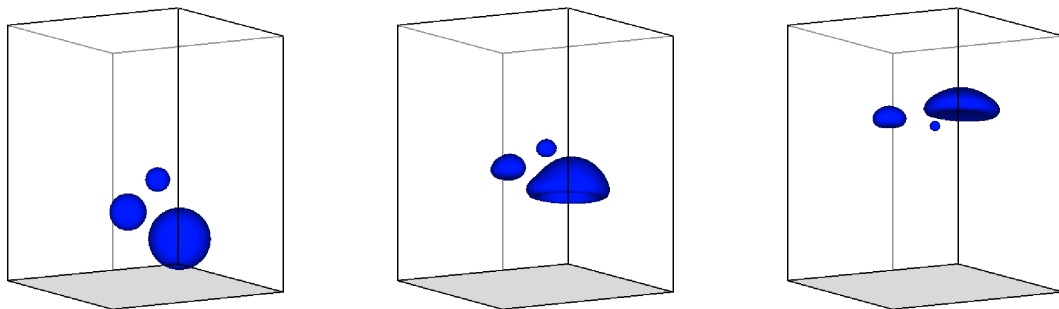


FIG. 10 – Calcul 3D disphasique

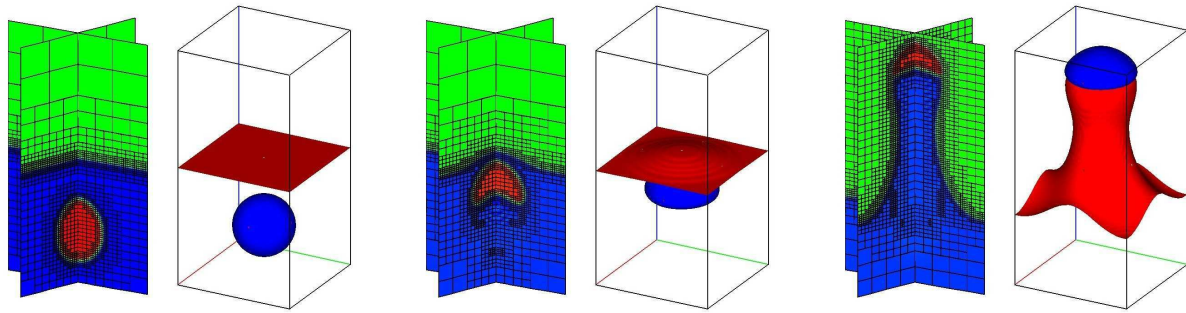


FIG. 11 – Calcul 3D triphasique

5 Publications

L'ensemble de ces travaux a donné lieu à l'écriture de trois publications dans des revues à comité de lecture et deux actes de conférence à comité de lecture. Par ailleurs, trois articles sont en cours de rédaction.

5.1 Articles dans des revues à comité de lecture

[BLMP09]

F. Boyer, C. Lapuerta, S. Minjeaud, B. Piar, *A local adaptive refinement method with multigrid preconditioning illustrated by multiphase flows simulations*, ESAIM Proceedings, Vol 27, pp. 15–53, 2009.

Article rédigé sur invitation des éditeurs suite au prix POSTER CANUM 2008.

Cet article décrit les algorithmes de raffinement local adaptatif et la procédure de “déaffinement” associée à la construction des preconditionneurs multiniveaux. Ces algorithmes sont illustrés par des simulations académiques (résolution du problème de Laplace (stationnaire) et du modèle de Cahn-Hilliard). Son contenu est essentiellement celui des chapitres I et II de ce manuscrit.

[BLM⁺]

F. Boyer, C. Lapuerta, S. Minjeaud, B. Piar, M. Quintard, *Cahn-Hilliard / Navier-Stokes model for the simulation of three-phase flows*, Transport in Porous Media, Vol 82(3), pp. 463–483, 2010.

Cet article donne une vue d'ensemble de l'établissement et des propriétés des équations du modèle de Cahn-Hilliard triphasique tel qu'il est présenté dans la thèse [Lap06] et dans le chapitre IV de cette thèse, du couplage aux équations de Navier-Stokes ainsi que les premières applications (notamment en trois dimensions) du raffinement local.

[BM10]

F. Boyer, S. Minjeaud, *Numerical Schemes for a three component Cahn-Hilliard model*, Mathematical modelling and numerical analysis, 2010, en révision, <http://hal.archives-ouvertes.fr/hal-00390065/fr/>

Cet article présente les discrétisations en temps du système de Cahn-Hilliard, leur étude théorique (existence et convergence des solutions approchées), ainsi que des tests numériques (courbes de convergence, comparaisons des schémas). Son contenu est essentiellement celui du chapitre V de ce manuscrit.

5.2 Articles en cours de rédaction

[DLM10a]

F. Dardhalon, J.C. Latché, S. Minjeaud, *Analysis of a projection method for low order nonconforming finite elements*, 2010.

Cet article contient une version détaillée de l'étude de la méthode de projection incrémentale dans le cas d'une discrétisation spatiale effectuée avec des éléments finis de bas degré non conformes (de type Rannacher-Turek) présentée dans la section VII.2 de ce manuscrit.

[MP10]

S. Minjeaud, B. Piar, *An adaptive pressure correction method without spurious velocities for diffuse-interface models of incompressible flows*, 2010.

Cet article reprend le contenu de la section VII.1 de ce manuscrit : variante de la méthode de projection incrémentale visant à réduire les vitesses parasites et une illustration par le phénomène des courants parasites lorsque le système de Navier-Stokes est couplé aux équations de Cahn-Hilliard.

[Min10]

S. Minjeaud, *An unconditionally stable uncoupled scheme for the approximation of a triphasic Cahn-Hilliard/Navier-Stokes system*, 2010.

Cet article reprend le contenu du chapitre VI de ce manuscrit : description et étude théorique du schéma découplé inconditionnellement stable pour le système couplé Cahn-Hilliard/Navier-Stokes.

5.3 Actes de conférences à comité de lecture

[BM]

F. Boyer, S. Minjeaud, *Fully discrete approximation of a three component Cahn-Hilliard model*, Proceedings of ALGORITMY the 18th Conference on Scientific Computing, (Vysoké Tatry - Podbanské, Slovaquie), 2009.

http://pc2.iam.fmph.uniba.sk/amuc/_contributed/algo2009/minjeaud.pdf

[DLM10b]

F. Dardhalon, J.C. Latché, S. Minjeaud, *On a projection method for piecewise-constant pressure nonconforming finite elements*,

Proceedings of MFD2010 the International congress in mathematical fluid dynamics and its applications (Rennes, France), 2010.

<http://hal.archives-ouvertes.fr/hal-00493589/fr/>

Partie 1

Algorithmes de raffinement local adaptatif et méthodes de préconditionnement multigrille

La mise en oeuvre de techniques de raffinement local adaptatif permet d'ajuster dynamiquement la résolution spatiale des méthodes d'approximation numériques en fonction de la position dans le domaine de calcul.

Pour accroître la résolution spatiale de certaines parties du domaine, une solution est d'utiliser des éléments de plus petit diamètre, *i.e.* de découper les cellules en cellules plus petites. La principale difficulté est alors de préserver à la fois la conformité des maillages et leur bonne qualité géométrique. La conformité interdit la présence de ce que l'on appelle les noeuds esclaves indésirables dans les applications. Les noeuds esclaves ne sont pas des degrés de liberté et sont quelquefois difficiles à prendre en compte, parce qu'ils modifient le stencil local de la matrice de rigidité. Lorsqu'ils existent, ils peuvent être gérés de diverses manières : élimination directe de ces "fausses inconnues", ajout de contraintes supplémentaires ou méthode de pénalisation et multiplicateur de Lagrange. A chaque fois, cela conduit à une modification des méthodes numériques et des schémas utilisés (donc du code de calcul).

Une autre méthode consiste à éliminer les non-conformités du maillage en divisant les cellules jusqu'à ce qu'il ne reste plus aucune non conformité. En deux dimensions, pour des maillages triangulaires, nous pouvons donner l'exemple bien connu du "red-green refinement" [BSW83]. Cette technique consiste à utiliser un découpage "régulier", appelé "red refinement", de chaque triangle, en quatre triangles semblables en connectant les milieux de ses arêtes. Ce raffinement préserve les propriétés géométriques des triangles mais crée des arêtes non conformes lorsque deux triangles, raffiné et non raffiné, sont adjacents. Un second type de découpage est alors utilisé, le "green refinement", connectant le milieu de l'arête non conforme au sommet opposé. Cela donne une bissection du triangle qui est "irrégulière" mais utilisée seulement aux endroits où des non conformités apparaissent. Bey [Bey95] et Zhang [Zha88] ont proposé la généralisation de cette méthode en trois dimensions. D'autres méthodes basées seulement sur la bissection ont été introduites par Rivara [Riv84, RI96] ou Mitchell [Mit91] en deux dimensions et Bänsch [Bän91] ou Maubach [Mau95] en trois dimensions. Toutes ces méthodes dépendent du choix de l'élément fini et de la dimension. De plus, leur implémentation peut devenir complexe surtout en trois dimensions.

Une alternative possible considérée dans cette partie est d'adopter le point de vue du raffinement des fonctions de base plutôt que celui des cellules. Cette approche repose sur la donnée d'une suite d'espaces d'approximation H^1 -conformes emboîtés $X_0 \subset \dots \subset X_J$, $J \geq 1$, engendrés par des ensembles B_j , $j \in \llbracket 0, J \rrbracket$ de fonctions de base de résolution spatiale d'autant plus fine que j est grand. Le raffinement local est alors réalisé en utilisant des espaces d'approximation composites (ou multiniveaux), c'est-à-dire des espaces engendrés par des fonctions de base sélectionnées dans chacun des ensembles B_j , $j \in \llbracket 0, J \rrbracket$ selon la résolution souhaitée dans chacune des parties du domaine. Dans ce cadre, Krysl, Grinspun et Schröder dans [KGS03] (*cf* aussi [GKS02, KTZ04, HK03]) ont proposé des procédures appelées CHARMS qui permettent de raffiner ou déraffiner les espaces d'approximation multiniveaux, c'est-à-dire ajouter ou retirer des fonctions de base des familles engendrant ces espaces tout en préservant leur indépendance linéaire. Nous donnons dans cette partie une construction précise de la séquence d'espaces emboîtés $X_0 \subset \dots \subset X_J$ pour des éléments finis conformes de type Lagrange (par exemple $\mathbb{P}_k$, $\mathbb{Q}_k$, $k \geq 1$). Nous montrons de manière détaillée qu'il est possible d'utiliser un seul motif de raffinement pour réaliser cette construction et nous étudions une version modifiée des algorithmes d'adaptation quasi-hiérarchique. La différence essentielle provient de l'application d'une règle pratique "au-plus-un-niveau-de-différence" qui fait partie intégrante des algorithmes de [KGS03]. Nous avons externalisé cette règle, ce qui nous a en particulier permis d'en définir une variante garantissant que la largeur de bande des matrices assemblées reste bornée quelles que soient les étapes d'adaptation. Finalement, nous incorporons à la méthode, des préconditionneurs multigrilles [Yse92, BPX90, Hac85, Zha88, BDY88]. Les avantages de la méthode sont les suivants :

- le problème discret (s'il est formulé de manière variationnelle) n'a pas besoin d'être modifié,
- les cellules sont divisées en cellules de même type en appliquant uniformément le même motif de raffinement,
- les non-conformités géométriques du maillage sont prises en compte de manière implicite,
- il n'y a pas de traitement particulier pour les différents éléments finis conformes de type Lagrange,
- la procédure ne dépend pas de la dimension en espace,
- les problèmes d'évolution en temps ne requièrent aucune construction d'opérateurs de transfert entre les différentes grilles à condition que l'on utilise une définition convenable des domaines d'intégration élémentaires lorsque le système est assemblé.

Cette partie est organisée comme suit. Dans le chapitre I, nous introduisons la notion de motif de raffinement pour des éléments finis conformes de type Lagrange et nous montrons comment une hiérarchie de maillage conforme peut être construite en appliquant récursivement un unique motif de raffinement à chaque cellule d'un maillage initial conforme quelconque. Nous étudions en détail les relations "parents-enfants" entre les fonctions de base de deux niveaux de raffinement successifs. En particulier, nous montrons que tous les coefficients dans cette relation linéaire peuvent être déduits facilement du motif de raffinement sur l'élément de référence. En

conséquence, la hiérarchie de maillage n'est jamais construite explicitement puisque toutes les informations sont disponibles sur l'élément de référence.

Dans le chapitre II, nous décrivons la procédure d'adaptation locale et nous établissons ses principales propriétés. En particulier, nous donnons des conditions précises suffisantes pour garantir que :

- l'algorithme d'adaptation que nous proposons préserve l'indépendance linéaire des fonctions de base sélectionnées sur les différents niveaux,
- qu'aucune information n'est perdue lorsque nous raffinons une fonction de base,
- les espaces d'approximation obtenus en raffinant (resp. déraffinant) sont indépendants de l'ordre dans lequel les raffinements (resp. déraffinements) successifs sont réalisés.

Quelques contre-exemples illustrent le fait que ces propriétés ne sont pas satisfaites dans le cas général.

Dans la seconde partie de ce chapitre II, nous montrons que la structure multiniveaux des espaces d'approximation construits peut être exploitée afin de dériver des préconditionneurs multigrilles efficaces. Tous les résultats de ce chapitre sont ensuite illustrés par des simulations numériques sur un problème modèle.

Le chapitre III est consacré à la description de l'implémentation informatique des méthodes de raffinement local et préconditionneur multigrilles. Nous donnons en particulier une description du fonctionnement de ces méthodes en parallèle.

Chapitre I

Espace éléments finis multiniveau

Dans ce chapitre, nous introduisons les principales définitions et notations associées aux concepts fondamentaux des procédures d'adaptation et de “coarsening” présentées dans le chapitre suivant. La notion de motif de raffinement que nous définissons comme la donnée d'un élément fini de référence (conforme de type Lagrange) et de maillage de son support géométrique permet de construire une suite de maillages conformes globalement raffinés auxquels sont associées les fonctions de base éléments finis correspondant à l'élément de référence choisi. Les fonctions associées à deux maillages consécutifs de cette suite sont reliées par une relation appelées relation parents-enfants qui est la base des algorithmes présentés dans le chapitre II.

I.1 Notations et définitions

Cette section rappelle, pour fixer les notations, quelques définitions et propriétés classiques, concernant les maillages et les éléments finis de type Lagrange. Ces propriétés seront très utilisées dans les sections I.2 et I.3.

Définition I.1 (Elément fini de Lagrange [RT88, §4.1])

Un élément fini de Lagrange est une triplet (K, Σ, P) où

- K est un sous-ensemble compact, connexe, lipschitzien de $\mathbb{R}^d$ ($d = 1, 2$ ou 3),
- $\Sigma = \{\mathbf{a}_k \in K; 1 \leq k \leq N\}$ est un ensemble de N points distincts de K , appelés noeuds de Lagrange,
- P est un espace vectoriel de fonctions $p : K \rightarrow \mathbb{R}$ tel que Σ est P -unisolvant [RT88, Def 4.1-1], i.e.

$$\forall (\alpha_1, \dots, \alpha_N) \in \mathbb{R}^N, \exists ! p \in P, \forall k \in \llbracket 1, N \rrbracket, p(\mathbf{a}_k) = \alpha_k.$$

L'élément (K, Σ, P) est appelé élément fini de Lagrange polygonal ssi K est un polygone.

Définition I.2 (Maillage [EG04, Def 1.49])

Soit $\omega \subset \mathbb{R}^d$ ($d = 1, 2$ ou 3) un domaine [EG04, Def 1.46]. Un maillage $\mathcal{T}$ de ω est un ensemble $\{K_e \subset \omega; 1 \leq e \leq N_e\}$ de sous-ensembles (appelés cellules) de ω compacts connexes lipschitziens et d'intérieur non vide tels que

- $\bar{\omega} = \bigcup_{e=1}^{N_e} K_e$ et,
- $\overset{\circ}{K}_e \cap \overset{\circ}{K}_f = \emptyset$ pour toute paire de cellules distinctes (K_e, K_f) .

Définition I.3 (Maillage éléments finis généré à partir d'un élément de référence [EG04, §1.3.2])

Soit $\mathcal{T} = \{K_e; 1 \leq e \leq N_e\}$ un maillage du domaine ω . Soit $(\hat{K}, \hat{\Sigma}, \hat{P})$ un élément fini de Lagrange polygonal, appelé ci-après élément de référence. Nous disons que le maillage $\mathcal{T}$ est généré à partir de l'élément de

référence $(\hat{K}, \hat{\Sigma}, \hat{P})$ ssi :

Pour tout $e \in \llbracket 1, N_e \rrbracket$, il existe un C^1 -difféomorphisme T_e de $\hat{K}$ vers K_e .

Lorsque les transformations T_e , $1 \leq e \leq N_e$, sont affines, le maillage est dit affine.

Pour tout $e \in \llbracket 1, N_e \rrbracket$, nous pouvons définir un élément fini de Lagrange (K_e, Σ_e, P_e) en posant :

- $\Sigma_e = T_e(\hat{\Sigma})$ et
- $P_e = \{p \circ T_e^{-1}; p \in \hat{P}\}$.

Remarque I.4

Nous supposons toujours que l'élément de référence est polygonal. Cependant, certaines applications T_e , peuvent conduire à des cellules K_e non polygonales [EG04, Figure 1.13].

Les définitions suivantes sont communément utilisées pour construire des espaces d'approximation X éléments finis $H^1(\omega)$ -conforme i.e. tels que $X \subset H^1(\omega)$ (cf proposition I.10).

Définition I.5 (Élément de classe C^0 [RT88, Def 5.1-2])

L'élément fini polygonal de Lagrange $(\hat{K}, \hat{\Sigma}, \hat{P})$ est un élément de classe C^0 ssi

- $\hat{P} \subset C^0(\hat{K})$ et
- pour toute face $\hat{F}$ de $\hat{K}$, $\hat{\Sigma} \cap \hat{F}$ est $\hat{P}|_{\hat{F}}$ -unisolvant où $\hat{P}|_{\hat{F}} = \{\hat{\varphi}|_{\hat{F}}; \hat{\varphi} \in \hat{P}\}$.

Définition I.6 (Conditions de compatibilité [RT88, Def 5.1-3])

Nous dirons que les conditions de compatibilité pour un élément de référence de Lagrange $(\hat{K}, \hat{\Sigma}, \hat{P})$ sont vérifiées ssi pour toutes faces $\hat{F}_1$ et $\hat{F}_2$ de $\hat{K}$, pour toute fonction affine inversible $\hat{A}$ telle que $\hat{F}_2 = \hat{A}(\hat{F}_1)$, nous avons

- $\hat{\Sigma} \cap \hat{F}_2 = \hat{A}(\hat{\Sigma} \cap \hat{F}_1)$ et
- $\{\hat{\varphi}|_{\hat{F}_1}; \hat{\varphi} \in \hat{P}_1\} = \{\hat{\varphi} \circ \hat{A}|_{\hat{F}_1}; \hat{\varphi} \in \hat{P}_2\}$.

Définition I.7 (Maillage géométriquement conforme [EG04, Def 1.55])

Un maillage $\mathcal{T} = \{K_e; 1 \leq e \leq N_e\}$ du domaine ω est dit géométriquement conforme ssi pour toutes cellules K_e et K_f ayant une intersection $(d-1)$ -dimensionnelle non vide, notée $F = K_e \cap K_f$, $T_e^{-1}(F)$ et $T_f^{-1}(F)$ sont des faces de $\hat{K}$, et il existe une transformation affine bijective $\hat{A} : T_e^{-1}(F) \rightarrow T_f^{-1}(F)$ telle que $\hat{A} \circ T_e^{-1} = T_f^{-1}$ sur F .

Remarque I.8

La définition I.7 implique, en particulier, que pour toute paire de cellules distinctes (K_e, K_f) , l'intersection $K_e \cap K_f$ est :

- ou vide, ou un sommet commun en dimension 1,
- ou vide, ou un sommet commun, ou une face commune en dimension 2,
- ou vide, ou un sommet commun, ou une arête commune, ou une face commune en dimension 3.

Un exemple de maillage géométriquement non conforme est présenté dans la figure I.1.

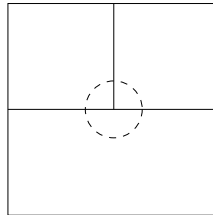


FIG. I.1 – Exemple de maillage géométriquement non conforme.

Remarque I.9

Dans le cas d'un maillage géométriquement conforme, la définition I.6 implique, en particulier, que les noeuds de Lagrange d'une face commune appartiennent à chacun des éléments partageant la face. Un exemple de position incompatible des noeuds est donné sur la figure I.2.

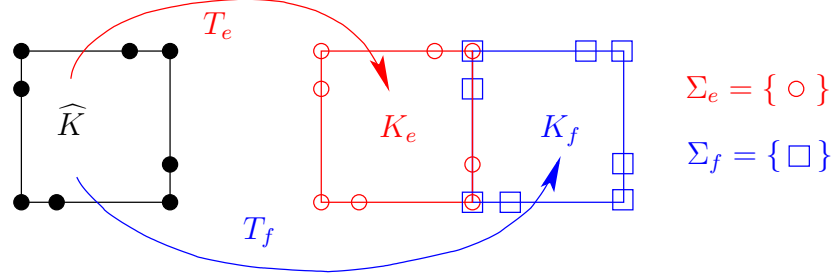


FIG. I.2 – Exemple de position incompatible des noeuds.

Soit $\mathcal{T} = \{K_e; 1 \leq e \leq N_e\}$ un maillage généré à partir de l'élément fini de référence de Lagrange $(\hat{K}, \hat{\Sigma}, \hat{P})$. D'après la définition I.3, nous pouvons associer à ce maillage :

- un ensemble de transformations géométriques $\{T_e : \hat{K} \rightarrow K_e; 1 \leq e \leq N_e\}$.
- un ensemble d'éléments finis de Lagrange $\{(K_e, \Sigma_e, P_e); 1 \leq e \leq N_e\}$.
- un ensemble de noeuds de Lagrange $\Sigma = \bigcup_{e=1}^{N_e} \Sigma_e$. Notons $N_{\text{ddl}} = \#\Sigma$ le nombre de noeuds de Lagrange; nous pouvons donc écrire $\Sigma = \{\mathbf{a}_i; 1 \leq i \leq N_{\text{ddl}}\}$.
- un ensemble de fonctions de base $\{\varphi_i; 1 \leq i \leq N_{\text{ddl}}\}$ défini comme suit. Pour $e \in \llbracket 1, N_e \rrbracket$ et $k \in \llbracket 1, \hat{N} \rrbracket$ où $\hat{N} = \#\hat{\Sigma}$, nous notons $I(e, k) \in \{1, \dots, N_{\text{ddl}}\}$ l'indice correspondant, dans la numérotation globale, au k^{e} noeud de Lagrange local dans la e^{e} cellule de $\mathcal{T}$. Cela signifie que :

$$\forall e \in \llbracket 1, N_e \rrbracket, \forall k \in \llbracket 1, \hat{N} \rrbracket, \quad \mathbf{a}_{I(e,k)} = T_e(\hat{\mathbf{a}}_k).$$

De plus, puisque $\hat{\Sigma}$ est $\hat{P}$ -unisolvant, nous avons :

$$\forall k \in \llbracket 1, \hat{N} \rrbracket, \exists \hat{\varphi}_k \in \hat{P}, \forall \ell \in \llbracket 1, \hat{N} \rrbracket, \quad \hat{\varphi}_k(\hat{\mathbf{a}}_\ell) = \delta_{k\ell}.$$

Nous pouvons associer à chaque noeud $\mathbf{a}_i$, $1 \leq i \leq N_{\text{ddl}}$, une fonction de base φ_i définie par morceaux sur chaque élément par

$$\varphi_i|_{K_e} = \begin{cases} \hat{\varphi}_k \circ T_e^{-1} & \text{s'il existe } k \in \llbracket 1, \hat{N} \rrbracket \text{ tel que } I(e, k) = i, \\ 0 & \text{sinon.} \end{cases}$$

Conformément à cette définition, nous notons $\text{supp}[\varphi_i]$ le sous-ensemble suivant de ω :

$$\text{supp}[\varphi_i] = \bigcup_{e \in \mathcal{E}} K_e \text{ où } \mathcal{E} = \{e \in \llbracket 1, N_e \rrbracket; \exists k \in \llbracket 1, \hat{N} \rrbracket, i = I(e, k)\}.$$

- un espace d'approximation $H^1(\omega)$ -conforme [EG04, Prop 1.74]

$$X = \{v \in \mathcal{C}^0(\omega); \forall e \in \llbracket 1, N_e \rrbracket, v|_{K_e} \in P_e\}.$$

L'avantage de cette construction classique est non seulement de produire un espace d'approximation $H^1(\omega)$ -conforme mais surtout de fournir une base explicite de cet espace. En effet, nous avons le résultat suivant :

Proposition I.10 ([EG04, Prop 1.78])

Soit $(\hat{K}, \hat{\Sigma}, \hat{P})$ un élément fini de référence de Lagrange, polygonal, de classe $\mathcal{C}^0$ et satisfaisant les conditions

de compatibilité I.6. Soit $\mathcal{T} = \{K_e; 1 \leq e \leq N_e\}$ un maillage géométriquement conforme de ω généré à partir de l'élément de référence $(\hat{K}, \hat{\Sigma}, \hat{P})$. Alors, nous avons $\varphi_k(\mathbf{a}_\ell) = \delta_{k\ell}$, et

$$\{\varphi_1, \dots, \varphi_{N_{ddl}}\} \text{ est une base de } X.$$

I.2 Motif de raffinement

Les définitions classiques rappelées dans la section I.1 permettent de définir la notion plus originale de motif de raffinement ainsi que les conditions de compatibilité associées.

Définition I.11 (Motif de raffinement)

Un motif de raffinement est un quadruplet $(\hat{K}, \hat{\Sigma}, \hat{P}, \hat{\mathcal{T}})$ où

- $(\hat{K}, \hat{\Sigma}, \hat{P})$ est un élément fini de référence de Lagrange, polygonal, de classe C^0 satisfaisant les conditions de compatibilité I.6,
- $\hat{\mathcal{T}} = \{\hat{K}_e^{[1]}; 1 \leq e \leq \hat{N}_e^{[1]}\}$ est un maillage affine géométriquement conforme de l'intérieur de $\hat{K}$ généré à partir de l'élément de référence $(\hat{K}, \hat{\Sigma}, \hat{P})$ lui-même.

En faisant référence à la section I.1, nous notons

- $\{\hat{T}_e : \hat{K} \rightarrow \hat{K}_e^{[1]}; 1 \leq e \leq \hat{N}_e^{[1]}\}$ l'ensemble des transformations géométriques,
- $\{(\hat{K}_e^{[1]}, \hat{\Sigma}_e^{[1]}, \hat{P}_e^{[1]}); 1 \leq e \leq \hat{N}_e^{[1]}\}$ l'ensemble des éléments finis de Lagrange,
- $\hat{\Sigma}^{[1]} = \{\hat{\mathbf{a}}_k^{[1]}; 1 \leq k \leq \hat{N}^{[1]}\}$ l'ensemble des noeuds de Lagrange,
- $\{\hat{\varphi}_k^{[1]}; 1 \leq k \leq \hat{N}^{[1]}\}$ l'ensemble des fonctions de base,

associés au maillage $\hat{\mathcal{T}}$.

Définition I.12 (Conditions de compatibilité)

Nous dirons que le motif de raffinement $(\hat{K}, \hat{\Sigma}, \hat{P}, \hat{\mathcal{T}})$ satisfait les conditions de compatibilité ssi

- pour tout $e \in \llbracket 1, \hat{N}_e^{[1]} \rrbracket$, $\{\varphi|_{\hat{K}_e^{[1]}}; \varphi \in \hat{P}\} \subset \hat{P}_e^{[1]}$,
- pour toutes faces $\hat{F}_1, \hat{F}_2$ de $\hat{K}$, pour toute fonction affine inversible $\hat{A}$ telle que $\hat{F}_2 = \hat{A}(\hat{F}_1)$, pour toutes faces $\hat{F}_1^{[1]} \subset \hat{F}_1$ d'un élément $\hat{K}_e^{[1]}$, $\hat{A}(\hat{F}_1^{[1]}) \subset \hat{F}_2$ est exactement la face d'un autre élément $\hat{K}_f^{[1]}$,
- chaque noeud de $\hat{\Sigma}$ est aussi un noeud de $\hat{\Sigma}^{[1]}$, i.e. $\hat{\Sigma} \subset \hat{\Sigma}^{[1]}$.

Remarque I.13

La définition I.12 est utilisée pour écarter les motifs de raffinement qui mèneront dans la suite à des maillages non conformes. Un exemple est donné dans la figure I.3.

Les conditions de compatibilité I.12 d'un motif de raffinement sont aussi utilisées pour établir l'équation de raffinement sur l'élément de référence qui est l'un des ingrédients fondamentaux de la méthode CHARMS (cf [KGS03]).

Proposition I.14 (Equation de raffinement sur le motif de raffinement)

Soit $(\hat{K}, \hat{\Sigma}, \hat{P}, \hat{\mathcal{T}})$ un motif de raffinement satisfaisant les conditions de compatibilité I.12. Nous avons la relation suivante :

$$\forall k \in \llbracket 1, \hat{N} \rrbracket, \quad \hat{\varphi}_k = \sum_{\ell=1}^{\hat{N}^{[1]}} \hat{\beta}_{k\ell} \hat{\varphi}_\ell^{[1]} \quad \text{où} \quad \hat{\beta}_{k\ell} = \hat{\varphi}_k(\hat{\mathbf{a}}_\ell^{[1]}).$$

Démonstration : D'après la proposition I.10, $\{\hat{\varphi}_1^{[1]}, \dots, \hat{\varphi}_{\hat{N}^{[1]}}^{[1]}\}$ est une base de l'espace

$$\left\{ v \in C^0(\hat{K}); \forall e \in \llbracket 1, \hat{N}_e^{[1]} \rrbracket, v|_{\hat{K}_e^{[1]}} \in \hat{P}_e^{[1]} \right\}.$$

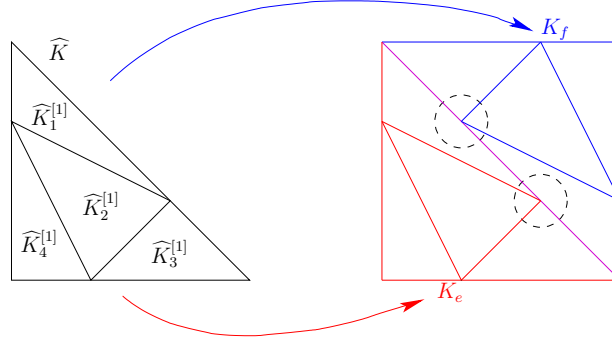


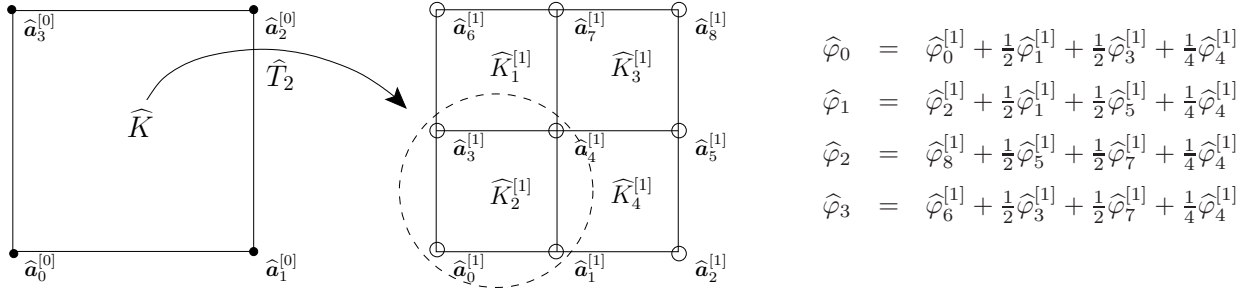
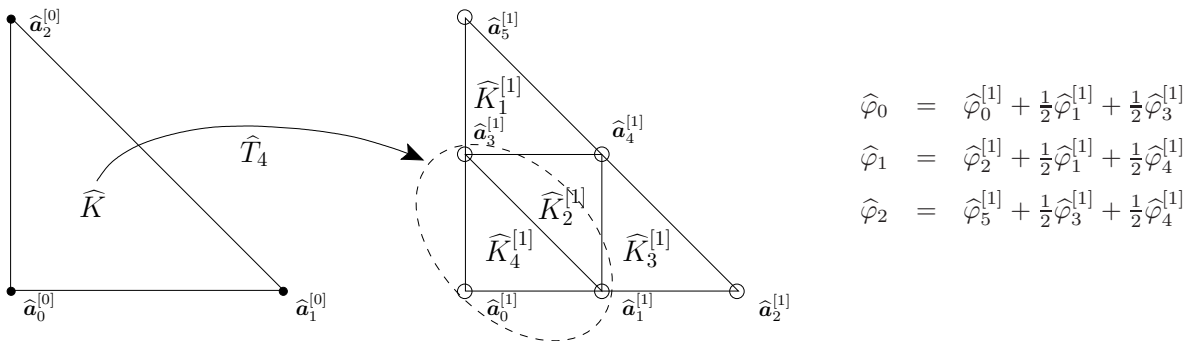
FIG. I.3 – Exemple de motif de raffinement incompatible.

Cependant, pour tout $k \in \{1, \dots, \hat{N}\}$, la fonction de base $\hat{\varphi}_k$ est dans $\mathcal{C}^0(\hat{K})$ et en utilisant les conditions de compatibilité I.12, nous avons $\forall e \in \llbracket 1, \hat{N}_e^{[1]} \rrbracket$, $\hat{\varphi}_k|_{\hat{K}_e^{[1]}} \in \hat{P}_e^{[1]}$. Ainsi, il existe des coefficients $\hat{\beta}_{k\ell}$ tels que

$$\hat{\varphi}_k = \sum_{\ell=1}^{\hat{N}^{[1]}} \hat{\beta}_{k\ell} \hat{\varphi}_\ell^{[1]}.$$

Finalement, les coefficients $\hat{\beta}_{kt}$ peuvent être obtenus grâce à la relation $\hat{\varphi}_\ell^{[1]}(\hat{\mathbf{a}}_t^{[1]}) = \delta_{\ell t}$. ■

Les notations introduites dans cette section sont illustrées par les figures I.4 et I.5 qui montrent les motifs de raffinement complets et toutes les équations de raffinement associées aux éléments carré- $\mathbb{Q}_1$ et triangle- $\mathbb{P}_1$.

FIG. I.4 – Motif et équations de raffinement associés à l'élément carré- $\mathbb{Q}_1$.FIG. I.5 – Motif et équations de raffinement associés à l'élément triangle- $\mathbb{P}_1$.

Puisque la configuration géométrique est plus compliquée pour l'élément carré- $\mathbb{Q}_2$, nous donnons seulement dans la figure I.6 les coefficients non nuls dans l'équation de raffinement associée à trois fonctions de base grossières (les autres relations pouvant être obtenues par symétrie). Plus précisément, sur chaque figure, la

fonction de base grossière, représentée par un point noir, est une combinaison linéaire des fonctions de base plus fines, représentées par des cercles, avec les coefficients mentionnés à côté.

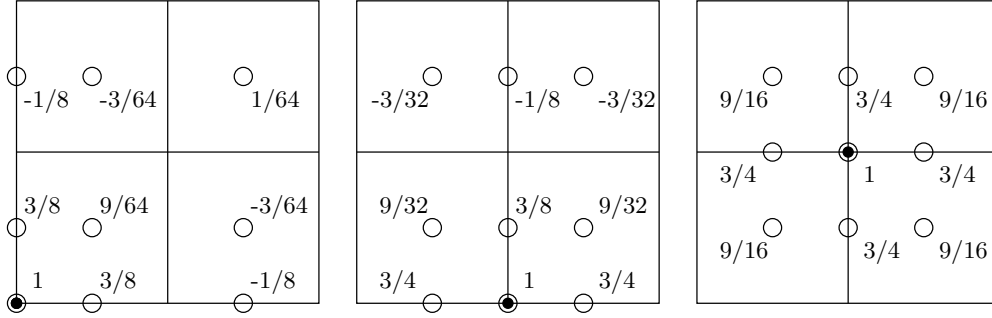


FIG. I.6 – Motif et équations de raffinement associés à l'élément carré- $\mathbb{Q}_2$.

I.3 Espaces d'approximation éléments finis multiniveaux

Soit Ω un domaine borné de $\mathbb{R}^d$, $d = 1, 2$ ou 3 . L'objectif de cette section est de donner un moyen automatique de construire des espaces d'approximation éléments finis multiniveaux $H^1(\Omega)$ -conformes à partir d'un maillage initial géométriquement conforme $\mathcal{T}_0$ et d'un motif de raffinement donnés. Soit $J \in \mathbb{N}^*$. Nous construisons une hiérarchie de maillages emboîtés $\mathcal{T}_0, \mathcal{T}_1, \dots, \mathcal{T}_J$ déduits à partir de $\mathcal{T}_0$ en appliquant uniformément et récursivement le motif de raffinement. Un espace d'approximation multiniveaux est ensuite obtenu en sélectionnant des fonctions de base associées à chacun des maillages $\mathcal{T}_j$, $0 \leq j \leq J$ suivant une méthode qui garantit l'indépendance linéaire des fonctions de base sélectionnées.

I.3.1 Hiérarchie d'espaces d'approximation conformes. Relation parents-enfants.

Soit $(\hat{K}, \hat{\Sigma}, \hat{P}, \hat{T})$ un motif de raffinement et $j \in \llbracket 0, J-1 \rrbracket$. Dans cette section, nous supposons qu'un maillage géométriquement conforme $\mathcal{T}_j = \{K_e^{[j]}; 1 \leq e \leq N_e^{[j]}\}$ de Ω est donné, que ce maillage $\mathcal{T}_j$ est généré à partir de l'élément de référence $(\hat{K}, \hat{\Sigma}, \hat{P})$, et nous expliquons ensuite comment le maillage $\mathcal{T}_{j+1}$ est construit.

Dans la suite, tous les objets mathématiques associés au maillage $\mathcal{T}_j$ seront marqués avec le signe j comme suit :

- $T_e^{[j]}$ est la transformation géométrique utilisée pour générer $K_e^{[j]}$, i.e. $K_e^{[j]} = T_e^{[j]}(\hat{K})$,
- $\{\mathbf{a}_1^{[j]}, \dots, \mathbf{a}_{N_{\text{ddl}}^{[j]}}^{[j]}\}$ est l'ensemble des noeuds de Lagrange du maillage $\mathcal{T}_j$, appelés noeuds de niveau j ,
- $B_j = \{\varphi_1^{[j]}, \dots, \varphi_{N_{\text{ddl}}^{[j]}}^{[j]}\}$ est l'ensemble des fonctions de base associé au maillage $\mathcal{T}_j$, appelées fonctions de base de niveau j ,
- $X_j = \{v \in \mathcal{C}^0(\Omega); \forall e \in \llbracket 1, N_e^{[j]} \rrbracket, v|_{K_e^{[j]}} \in P_e^{[j]}\}$ est l'espace d'approximation $H^1(\Omega)$ -conforme associé au maillage $\mathcal{T}_j$.

D'après la proposition I.10, nous avons le résultat suivant :

B_j est une base de l'espace d'approximation $H^1(\Omega)$ -conforme X_j .

Définition I.15

Nous définissons l'ensemble $\mathcal{T}_{j+1}$ comme suit :

$$\mathcal{T}_{j+1} = \left\{ K_{ef}^{[j+1]} = T_e^{[j]}(\hat{K}_f^{[1]}); \quad 1 \leq e \leq N_e^{[j]}, \quad 1 \leq f \leq \hat{N}_e^{[1]} \right\}.$$

Proposition I.16

L'ensemble $\mathcal{T}_{j+1}$ est un maillage géométriquement conforme de Ω généré à partir de l'élément de référence $\hat{K}$.

Démonstration :

1) Il est facile de voir que $\mathcal{T}_{j+1}$ est un maillage de Ω généré à partir de l'élément de référence $\hat{K}$.

- (i) Soit $x \in \bar{\Omega}$. Puisque l'ensemble $\mathcal{T}_j$ est un maillage de Ω , il existe $e \in \llbracket 1, N_e^{[j]} \rrbracket$ tel que $x \in K_e^{[j]} = T_e^{[j]}(\hat{K})$. Comme $\hat{\mathcal{T}}$ est un maillage de l'intérieur de $\hat{K}$, il existe $f \in \llbracket 1, \hat{N}_e^{[1]} \rrbracket$ tel quel $(T_e^{[j]})^{-1}(x) \in \hat{K}_f^{[1]}$ ou encore $x \in T_e^{[j]}(\hat{K}_f^{[1]})$. Ainsi, nous avons montré que

$$\bar{\Omega} \subset \bigcup K_{ef}^{[j+1]}, \quad 1 \leq e \leq N_e^{[j]}, \quad 1 \leq f \leq \hat{N}_e^{[1]}.$$

L'inclusion réciproque est évidente : $K_{ef}^{[j+1]} \subset K_e^{[j]} \subset \bar{\Omega}$, $\forall e \in \llbracket 1, N_e^{[j]} \rrbracket$, $\forall f \in \llbracket 1, \hat{N}_e^{[1]} \rrbracket$.

- (ii) Soit $1 \leq e, e' \leq N_e^{[j]}$, et $1 \leq f, f' \leq \hat{N}_e^{[1]}$. Supposons que $\overset{\circ}{K}_{ef}^{[j+1]} \cap \overset{\circ}{K}_{e'f'}^{[j+1]} \neq \emptyset$ et montrons que $e = e'$ et $f = f'$. Comme $\hat{K}_f^{[1]} \subset \hat{K}$, nous avons $K_{ef}^{[j+1]} \subset T_e^{[j]}(\hat{K}) = K_e^{[j]}$. Ainsi, $\overset{\circ}{K}_{ef}^{[j+1]} \cap \overset{\circ}{K}_{e'f'}^{[j+1]} \subset \overset{\circ}{K}_e^{[j]} \cap \overset{\circ}{K}_{e'}^{[j]}$. Cette dernière intersection est donc non vide et puisque $\mathcal{T}_j$ est un maillage de Ω ceci conduit à $e = e'$. Nous concluons maintenant grâce aux inclusions suivantes :

$$\begin{aligned} \overset{\circ}{K}_{ef}^{[j+1]} \cap \overset{\circ}{K}_{e'f'}^{[j+1]} &\subset \overset{\circ}{T}_e^{[j]}(\hat{K}_f^{[1]}) \cap \overset{\circ}{T}_{e'}^{[j]}(\hat{K}_{f'}^{[1]}) && \text{(par définition),} \\ &\subset \overset{\circ}{T}_e^{[j]}(\hat{K}_f^{[1]}) \cap \overset{\circ}{T}_{e'}^{[j]}(\hat{K}_{f'}^{[1]}) && \text{(puisque } T_e^{[j]} \text{ est un homéomorphisme),} \\ &\subset \overset{\circ}{T}_e^{[j]}(\hat{K}_f^{[1]} \cap \hat{K}_{f'}^{[1]}) && \text{(puisque } T_e^{[j]} \text{ est injective).} \end{aligned}$$

Ainsi, $\hat{K}_f^{[1]} \cap \hat{K}_{f'}^{[1]}$ est non vide, et par suite $f = f'$ puisque $\hat{\mathcal{T}}$ est un maillage de l'intérieur de $\hat{K}$.

- 2) Il reste à prouver que ce maillage est géométriquement conforme. Soit $K_{ef}^{[j+1]}$ et $K_{e'f'}^{[j+1]}$ deux cellules qui ont une intersection $(d-1)$ -dimensionnelle non vide, notée $F = K_{ef}^{[j+1]} \cap K_{e'f'}^{[j+1]}$. La démonstration est basée sur des arguments différents selon si $e = e'$ ou $e \neq e'$. Dans le premier cas, nous utilisons la conformité géométrique de $\hat{\mathcal{T}}$ alors que dans le second cas le résultat est déduit de la conformité géométrique de $\mathcal{T}_j$ et des conditions de compatibilité I.12.

- Supposons tout d'abord que $e = e'$. Nous avons $F = T_e^{[j]}(\hat{K}_f^{[1]} \cap \hat{K}_{f'}^{[1]})$ puisque $T_e^{[j]}$ est injective. Posons alors, $\hat{F}^{[1]} = \hat{K}_f^{[1]} \cap \hat{K}_{f'}^{[1]}$. Puisque $T_e^{[j]}$ est un $\mathcal{C}^1$ -difféomorphisme, $\hat{F}^{[1]}$ est l'intersection $(d-1)$ -dimensionnelle des deux cellules $\hat{K}_f^{[1]}$ et $\hat{K}_{f'}^{[1]}$ du maillage géométriquement conforme $\hat{\mathcal{T}}$ de l'intérieur de $\hat{K}$. Ainsi, $\hat{T}_f^{-1}(\hat{F}^{[1]})$ et $\hat{T}_{f'}^{-1}(\hat{F}^{[1]})$ sont des faces de $\hat{K}$ et il existe une application affine bijective $\hat{A} : \hat{T}_f^{-1}(\hat{F}^{[1]}) \rightarrow \hat{T}_{f'}^{-1}(\hat{F}^{[1]})$ telle que $\hat{A} \circ \hat{T}_f^{-1} = \hat{T}_{f'}^{-1}$ sur $\hat{F}^{[1]}$. On se rappelle maintenant que $\hat{F}^{[1]} = (T_e^{[j]})^{-1}(F)$ pour obtenir la conclusion.
- Supposons maintenant que $e \neq e'$. Puisque $\hat{K}_f^{[1]} \subset \hat{K}$ et $\hat{K}_{f'}^{[1]} \subset \hat{K}$, nous avons $F \subset K_e^{[j]} \cap K_{e'}^{[j]}$. Posons $G = K_e^{[j]} \cap K_{e'}^{[j]}$, nous avons donc $F \subset G$ et ainsi les cellules $K_e^{[j]}$ et $K_{e'}^{[j]}$ possèdent une intersection $d-1$ dimensionnelle non vide. Puisque, $\mathcal{T}_j$ est géométriquement conforme, $\hat{F}_1 = (T_e^{[j]})^{-1}(G)$ et $\hat{F}_2 = (T_{e'}^{[j]})^{-1}(G)$ sont des faces de $\hat{K}$ et il existe une application $\hat{A} : \hat{F}_1 \rightarrow \hat{F}_2$ affine bijective telle que $\hat{A} \circ (T_e^{[j]})^{-1} = (T_{e'}^{[j]})^{-1}$ sur G . Posons maintenant $\hat{F}_1^{[1]} = \hat{K}_f^{[1]} \cap \hat{F}_1$ et $\hat{F}_2^{[1]} = \hat{K}_{f'}^{[1]} \cap \hat{F}_2$. Il est facile de montrer que $F \subset T_e^{[j]}(\hat{F}_1^{[1]})$ et que $F \subset T_{e'}^{[j]}(\hat{F}_2^{[1]})$. Puisque F est $d-1$ dimensionnelle non vide, nous en déduisons que $\hat{F}_1^{[1]}$ et $\hat{F}_2^{[1]}$ sont des faces de $\hat{K}_f^{[1]}$ et $\hat{K}_{f'}^{[1]}$ respectivement (rappelons que ces polygones ou polyèdres sont entièrement inclus dans $\hat{K}$ dont $\hat{F}_1$ et $\hat{F}_2$ sont des faces). Les conditions de compatibilité I.12 sur le motif de raffinement impliquent alors que $\hat{A}(\hat{F}_1^{[1]})$ est une face d'un élément $\hat{K}_g^{[1]}$ du maillage

de l'élément de référence. Or nous avons :

$$\begin{aligned}
\widehat{A}(\widehat{F}_1^{[1]}) \cap \widehat{F}_2^{[1]} &= \widehat{A}(\widehat{K}_f^{[1]} \cap \widehat{F}_1) \cap \widehat{K}_{f'}^{[1]} \cap \widehat{F}_2 && \text{(par définition de } \widehat{F}_i^{[1]}), \\
&= \widehat{A} \circ (T_e^{[j]})^{-1} (K_{ef}^{[j+1]} \cap G) \cap (T_{e'}^{[j]})^{-1} (K_{ef'}^{[j+1]} \cap G) && \text{(par déf. de } \widehat{F}_i, K_{ef}^{[j+1]}, K_{ef'}^{[j+1]}), \\
&= (T_{e'}^{[j]})^{-1} (K_{ef}^{[j+1]} \cap K_{ef'}^{[j+1]} \cap G) && \text{(puisque } \widehat{A} \circ (T_e^{[j]})^{-1} = (T_{e'}^{[j]})^{-1} \text{ sur } G), \\
&= (T_{e'}^{[j]})^{-1} (F) && \text{(par définition de } F).
\end{aligned}$$

Ainsi, les faces $\widehat{A}(\widehat{F}_1^{[1]})$ et $\widehat{F}_2^{[1]}$ possèdent une intersection $d-1$ dimensionnelle non vide. Elles sont donc égales et il vient

$$\widehat{F}_1^{[1]} = (T_e^{[j]})^{-1}(F) \quad \text{et} \quad \widehat{F}_2^{[1]} = (T_{e'}^{[j]})^{-1}(F).$$

Ainsi, $(T_e^{[j]} \circ \widehat{T}_f)^{-1}(F)$ et $(T_{e'}^{[j]} \circ \widehat{T}_{f'})^{-1}(F)$ sont deux faces de $\widehat{K}$. On conclut en posant $A = (\widehat{T}_{f'})^{-1} \circ \widehat{A} \circ \widehat{T}_f$. L'application A est affine bijective de $(T_e^{[j]} \circ \widehat{T}_f)^{-1}(F)$ dans $(T_{e'}^{[j]} \circ \widehat{T}_{f'})^{-1}(F)$ et vérifie $A \circ (T_e^{[j]} \circ \widehat{T}_f)^{-1} = (T_{e'}^{[j]} \circ \widehat{T}_{f'})^{-1}$ sur F . ■

Notons également que le dernier point des conditions de compatibilité I.12 garantit de manière évidente que tout noeud de niveau j est aussi un noeud de niveau $j+1$:

$$\forall k \in \llbracket 1, N_{\text{ddl}}^{[j]} \rrbracket, \exists \ell \in \llbracket 1, N_{\text{ddl}}^{[j+1]} \rrbracket, \mathbf{a}_k^{[j]} = \mathbf{a}_\ell^{[j+1]}. \quad (\text{I.1})$$

Nous pouvons maintenant prouver que notre construction mène à des espaces d'approximation emboîtés :

Proposition I.17

Les espaces X_j et X_{j+1} sont emboîtés : $X_j \subset X_{j+1}$.

Démonstration : Soit $v \in X_j$, $e \in \{1, \dots, N_{\text{ddl}}^{[j]}\}$ et $f \in \{1, \dots, \widehat{N}_e^{[1]}\}$. Par définition, $v|_{K_e^{[j]}} \in P_e^{[j]}$. Ceci est équivalent à $v \circ T_e^{[j]} \in \widehat{P}$. Nous déduisons alors des conditions de compatibilité I.12 que $v \circ T_e^{[j]}|_{\widehat{K}_f^{[1]}} \in \widehat{P}_f^{[1]}$. Ainsi, nous obtenons $v \circ T_e^{[j]} \circ \widehat{T}_f \in \widehat{P}$, qui signifie exactement que $v|_{K_{ef}^{[j+1]}} \in P_{ef}^{[j+1]}$, et le résultat est prouvé. ■

Pour énoncer le résultat suivant, nous devons introduire une indexation pertinente des noeuds de niveau j et $j+1$. Les noeuds de niveau j et $j+1$ appartenant à $K_e^{[j]}$ sont par définition les images, par la transformation $T_e^{[j]}$, des noeuds de $\widehat{\Sigma}$ et $\widehat{\Sigma}^{[1]}$ respectivement. Ainsi nous notons par :

– $I^{[j]}(e, k)$ l'indice du noeud de niveau j image de $\widehat{\mathbf{a}}_k$ par la transformation $T_e^{[j]}$. Ceci signifie que nous avons :

$$\forall (e, k) \in \llbracket 1, N_e^{[j]} \rrbracket \times \llbracket 1, \widehat{N} \rrbracket, \quad \mathbf{a}_{I^{[j]}(e, k)}^{[j]} = T_e^{[j]}(\widehat{\mathbf{a}}_k).$$

– $I^{[j,1]}(e, \ell)$ l'indice du noeud de niveau $j+1$ image de $\widehat{\mathbf{a}}_\ell^{[1]}$ par la transformation $T_e^{[j]}$. Ceci signifie que nous avons :

$$\forall (e, \ell) \in \llbracket 1, N_e^{[j]} \rrbracket \times \llbracket 1, \widehat{N}^{[1]} \rrbracket, \quad \mathbf{a}_{I^{[j,1]}(e, \ell)}^{[j+1]} = T_e^{[j]}(\widehat{\mathbf{a}}_\ell^{[1]}).$$

Proposition I.18 (Equation de Raffinement)

Nous avons la relation suivante :

$$\forall i \in \llbracket 1, N_{\text{ddl}}^{[j]} \rrbracket, \quad \varphi_i^{[j]} = \sum_{t=1}^{N_{\text{ddl}}^{[j+1]}} \beta_{it}^{[j]} \varphi_t^{[j+1]}, \quad (\text{RE})$$

où les coefficients $\beta_{it}^{[j]}$ sont donnés par : $\forall (i, t) \in \llbracket 1, N_{\text{ddl}}^{[j]} \rrbracket \times \llbracket 1, N_{\text{ddl}}^{[j+1]} \rrbracket$,

$$\beta_{it}^{[j]} = \begin{cases} \widehat{\beta}_{k\ell} & \text{si } \exists (e, k, \ell) \in \llbracket 1, N_e^{[j]} \rrbracket \times \llbracket 1, \widehat{N} \rrbracket \times \llbracket 1, \widehat{N}^{[1]} \rrbracket \text{ t.q. } i = I^{[j]}(e, k) \text{ et } t = I^{[j,1]}(e, \ell), \\ 0 & \text{sinon.} \end{cases}$$

Remarque I.19

Remarquons qu'en pratique les coefficients $\widehat{\beta}_{k\ell}$ dépendent seulement du motif de raffinement. Ainsi, ces coefficients peuvent être calculés a priori. Ils sont en petit nombre et donc leur stockage ne nécessite qu'une très faible place mémoire. L'équation de raffinement ci-dessus peut donc être déduite sans aucun calcul de coefficients. En pratique (cf section III.1.3), les coefficients $\beta_{it}^{[j]}$ sont obtenus grâce à une boucle sur les cellules de niveau j incluses dans le support de $\text{supp}[\varphi_i^{[j]}]$ en posant : pour tout $e \in \llbracket 1, N_e^{[j]} \rrbracket$ tel que $K_e^{[j]} \subset \text{supp}[\varphi_i^{[j]}]$, pour tout $(k, \ell) \in \llbracket 1, \widehat{N} \rrbracket \times \llbracket 1, \widehat{N}^{[1]} \rrbracket$,

$$\beta_{I^{[j]}(e,k)I^{[j,1]}(e,\ell)}^{[j]} = \widehat{\beta}_{k\ell},$$

et les autres coefficients sont égaux à zéro. Remarquons qu'une telle boucle mène à parcourir plusieurs fois la même paire d'indices $(I^{[j]}(e, k), I^{[j,1]}(e, \ell))$ pour des indices e, k, ℓ distincts. La proposition I.18 garantit que les coefficients correspondants $\widehat{\beta}_{k\ell}$ sont les mêmes.

Démonstration de la proposition I.18 : Soit $i \in \llbracket 1, N_{\text{ddl}}^{[j]} \rrbracket$. La fonction de base $\varphi_i^{[j]}$ appartient à X_j . Puisque $X_j \subset X_{j+1}$ et B_{j+1} est une base de X_{j+1} , l'existence des coefficients $\beta_{it}^{[j]}$ est évidente.

Soit $(i, t) \in \llbracket 1, N_{\text{ddl}}^{[j]} \rrbracket \times \llbracket 1, N_{\text{ddl}}^{[j+1]} \rrbracket$. Nous avons :

$$\varphi_i^{[j]} = \sum_{t=1}^{N_{\text{ddl}}^{[j+1]}} \beta_{it}^{[j]} \varphi_t^{[j+1]}. \quad (\text{I.2})$$

Il reste à démontrer la formule donnant les coefficients $\beta_{it}^{[j]}$ en fonction des coefficients $\widehat{\beta}_{k\ell}$.

- Cas 1 : il existe $(e, k, \ell) \in \llbracket 1, N_e^{[j]} \rrbracket \times \llbracket 1, \widehat{N} \rrbracket \times \llbracket 1, \widehat{N}^{[1]} \rrbracket$ tels que $i = I^{[j]}(e, k)$ et $t = I^{[j,1]}(e, \ell)$.
La restriction de (I.2) à $K_e^{[j]}$ donne :

$$\widehat{\varphi}_k \circ \left(T_e^{[j]}\right)^{-1} = \sum_{l=1}^{\widehat{N}^{[1]}} \beta_{iI^{[j,1]}(e,l)}^{[j]} \varphi_{I^{[j,1]}(e,l)}^{[j+1]}.$$

D'après la proposition I.14, nous avons :

$$\sum_{l=1}^{\widehat{N}^{[1]}} \widehat{\beta}_{kl} \widehat{\varphi}_l^{[1]} \circ \left(T_e^{[j]}\right)^{-1} = \sum_{l=1}^{\widehat{N}^{[1]}} \beta_{iI^{[j,1]}(e,l)}^{[j]} \widehat{\varphi}_l^{[1]} \circ \left(T_e^{[j]}\right)^{-1}.$$

En évaluant cette égalité en $T_e^{[j]}(\widehat{\mathbf{a}}_\ell^{[1]})$, il vient $\beta_{iI^{[j,1]}(e,\ell)}^{[j]} = \widehat{\beta}_{k\ell}$. C'est exactement $\beta_{it}^{[j]} = \widehat{\beta}_{k\ell}$.

- Cas 2 : $\forall (e, k, \ell) \in \llbracket 1, N_e^{[j]} \rrbracket \times \llbracket 1, \widehat{N} \rrbracket \times \llbracket 1, \widehat{N}^{[1]} \rrbracket$, $i \neq I^{[j]}(e, k)$ ou $t \neq I^{[j,1]}(e, \ell)$.
Soient $e \in \llbracket 1, N_e^{[j]} \rrbracket$ et $\ell \in \llbracket 1, \widehat{N}^{[1]} \rrbracket$ tels que $t = I^{[j,1]}(e, \ell)$. Nous avons nécessairement, par hypothèse, $\forall k \in \llbracket 1, \widehat{N} \rrbracket$, $i \neq I^{[j]}(e, k)$. Donc, nous obtenons

$$0 = \sum_{s=1}^{N_{\text{ddl}}^{[j+1]}} \beta_{is}^{[j]} \varphi_s^{[j+1]}|_{K_e^{[j]}} = \sum_{v=1}^{\widehat{N}^{[1]}} \beta_{iI^{[j,1]}(e,v)}^{[j]} \widehat{\varphi}_v^{[1]} \circ \left(T_e^{[j]}\right)^{-1}.$$

L'évaluation de cette égalité en $T^{[j]}(\widehat{\mathbf{a}}_\ell^{[1]})$ donne $\beta_{iI^{[j,1]}(e,\ell)}^{[j]} = 0$. C'est exactement $\beta_{it}^{[j]} = 0$. ■

L'équation de raffinement (RE) introduit une relation entre les fonctions de base de niveau j et certaines fonctions de base de niveau $j+1$, appelées enfants.

Définition I.20 (Relation parents-enfants pour les fonctions de base)

Si le coefficient $\beta_{it}^{[j]}$ de l'équation (RE) est non nul, nous dirons que :

- la fonction de base $\varphi_i^{[j]}$ de niveau j est un parent de la fonction de base $\varphi_t^{[j+1]}$ de niveau $j+1$,
- la fonction de base $\varphi_t^{[j+1]}$ de niveau $j+1$ est un enfant de la fonction de base $\varphi_i^{[j]}$ de niveau j .

Pour cette raison, l'équation de raffinement (RE) est aussi appelée relation parents-enfants. De la même manière, nous pouvons définir une relation parents-enfants pour les cellules.

Définition I.21 (Relation parents-enfants pour les cellules)

- Soit $e \in \llbracket 1, N_e^{[j]} \rrbracket$.
- Pour tout $f \in \{1, \dots, \widehat{N}_e^{[1]}\}$, nous dirons que la cellule $K_{ef}^{[j+1]}$ de niveau $j+1$ est une cellule enfant de la cellule $K_e^{[j]}$ de niveau j .
 - A l'inverse, nous dirons que la cellule $K_e^{[j]}$ de niveau j est une cellule parent de chacune des cellules de niveau $j+1$, $K_{ef}^{[j+1]}$, pour $f \in \{1, \dots, \widehat{N}_e^{[1]}\}$.

Remarque I.22

Une cellule a au plus une cellule parent alors qu'une fonction de base peut avoir plusieurs parents. Néanmoins, la proposition suivante identifie certaines fonctions de base ayant un unique parent.

Proposition I.23 (Enfant privé)

- Soit $(k, \ell) \in \llbracket 1, N_{ddl}^{[j]} \rrbracket \times \llbracket 1, N_{ddl}^{[j+1]} \rrbracket$ tel que $\mathbf{a}_k^{[j]} = \mathbf{a}_\ell^{[j+1]}$ alors,
- $$\varphi_k^{[j]} \text{ est l'unique parent de } \varphi_\ell^{[j+1]}.$$

Démonstration : Pour $1 \leq i \leq N_{ddl}^{[j]}$, la relation parents-enfants donne :

$$\varphi_i^{[j]}(\mathbf{a}_k^{[j]}) = \sum_{t=1}^{N_{ddl}^{[j+1]}} \beta_{it}^{[j]} \varphi_t^{[j+1]}(\mathbf{a}_\ell^{[j+1]}).$$

C'est exactement :

$$\delta_{ik} = \beta_{i\ell}^{[j]}.$$

Ainsi, la fonction de base $\varphi_\ell^{[j+1]}$ a un unique parent qui est $\varphi_k^{[j]}$. ■

Résumé I.24

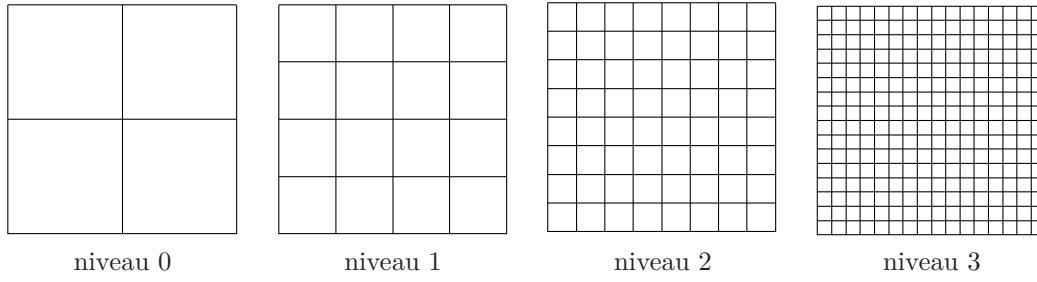
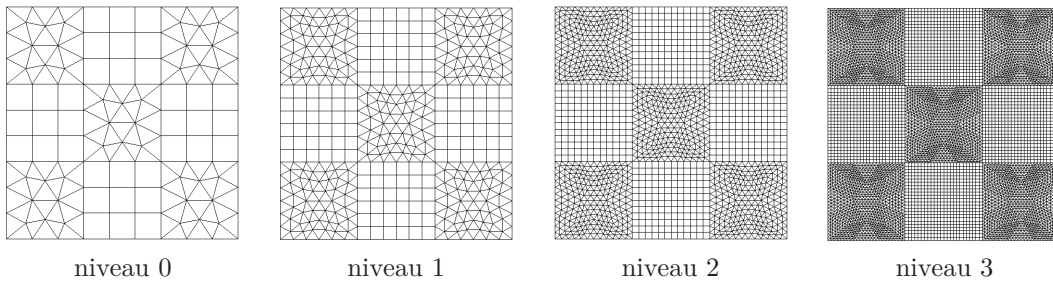
Soit $(\widehat{K}, \widehat{\Sigma}, \widehat{P}, \widehat{T})$ un motif de raffinement. Soit $\mathcal{T}_0$ un maillage géométriquement conforme de Ω généré à partir de l'élément de référence $(\widehat{K}, \widehat{\Sigma}, \widehat{P})$. En appliquant uniformément le motif de raffinement, nous pouvons construire :

- une hiérarchie de maillages emboîtés $\mathcal{T}_0, \mathcal{T}_1, \dots, \mathcal{T}_J$ (cf figures I.7, I.8, I.9),
- une hiérarchie d'espaces d'approximation éléments finis $H^1(\Omega)$ -conformes $X_0 \subset X_1 \subset \dots \subset X_J$,
- des ensembles de fonctions de base $B_0, B_1, \dots, B_J$, engendrant les espaces d'approximation précités, deux ensembles consécutifs étant reliés par les équations de raffinement.

La table I.1 donne un résumé des notations utilisées dans la suite.

	Maillages	Ensembles de fonctions de base	Espaces d'approximation
Niveau 0	$\mathcal{T}_0$	$B_0 = \{\varphi_k^{[0]}; k = 1, \dots, N_{ddl}^{[0]}\}$	$X_0 = \text{vect } B_0$
Niveau 1	$\mathcal{T}_1$	$B_1 = \{\varphi_k^{[1]}; k = 1, \dots, N_{ddl}^{[1]}\}$	$X_1 = \text{vect } B_1$
$\vdots$	$\vdots$	$\vdots$	$\vdots$
Niveau J	$\mathcal{T}_J$	$B_J = \{\varphi_k^{[J]}; k = 1, \dots, N_{ddl}^{[J]}\}$	$X_J = \text{vect } B_J$

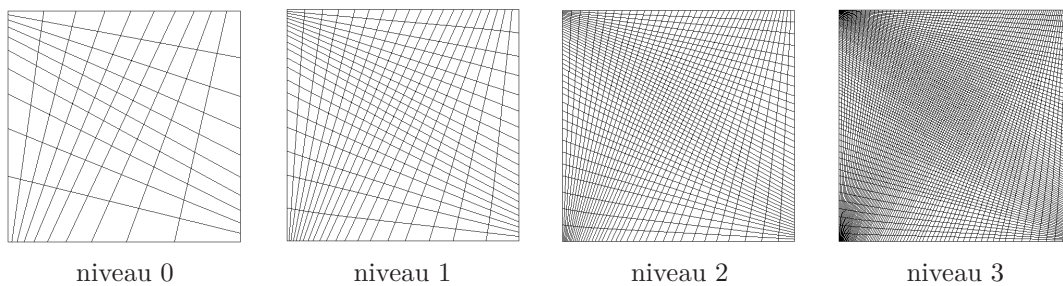
TAB. I.1 – Hiérarchie conceptuelle d'espaces éléments finis emboîtés.

FIG. I.7 – Carré- $\mathbb{Q}_1$. Maillages emboîtés $\mathcal{T}_j$.FIG. I.8 – Tri/Quadr-angle- $\mathbb{P}_1/\mathbb{Q}_1$. Maillages emboîtés $\mathcal{T}_j$.

Remarquons que la hiérarchie de maillages emboîtés n'est jamais construite explicitement. Cette structure conceptuelle est introduite pour expliquer la méthode de raffinement mais, en pratique (*cf* chapitre III), le motif de raffinement peut être appliqué localement aux endroits où les fonctions de base doivent être effectivement raffinées.

I.3.2 Bases multiniveaux et espaces d'approximation multiniveaux

Nous supposons dans cette section que la structure présentée dans le résumé I.24 est donnée et nous utilisons les mêmes notations (*cf* table I.1). L'objectif de cette section est d'expliquer comment nous pouvons sélectionner

FIG. I.9 – Quadrangle- $\mathbb{Q}_1$. Maillages emboîtés $\mathcal{T}_j$.

une famille libre de fonctions de base dans l'ensemble multiniveau $\bigcup_{j=0}^J B_j$.

Proposition I.25

Soit $\mathcal{B}$ un sous-ensemble de $\bigcup_{j=0}^J B_j$. Supposons que deux noeuds associés à deux fonctions de base distinctes de $\mathcal{B}$ n'aient jamais la même position géométrique, i.e.

$$\left(\begin{array}{l} \forall (j, j') \in \llbracket 0, J \rrbracket^2, \forall (k, k') \in \llbracket 1, N_{\text{ddl}}^{[j]} \rrbracket \times \llbracket 1, N_{\text{ddl}}^{[j']} \rrbracket \text{ tel que } \mathbf{a}_k^{[j]} = \mathbf{a}_{k'}^{[j']}, \\ (\varphi_k^{[j]} \in \mathcal{B}, \varphi_{k'}^{[j']} \in \mathcal{B}) \implies (j = j', k = k') \end{array} \right), \quad (\mathcal{P}_{\text{LI}})$$

alors $\mathcal{B}$ est une famille libre.

Remarque I.26

Remarquons que $(\mathbf{a}_k^{[j]} = \mathbf{a}_{k'}^{[j']}) \text{ et } j = j' \implies k = k'$.

Démonstration : La propriété $(\mathcal{P}_{\text{LI}})$ implique que

$$\left(\begin{array}{l} \forall (j, j') \in \llbracket 0, J \rrbracket^2, \forall (k, k') \in \llbracket 1, N_{\text{ddl}}^{[j]} \rrbracket \times \llbracket 1, N_{\text{ddl}}^{[j']} \rrbracket \text{ tel que } j' > j, \\ (\varphi_k^{[j]} \in \mathcal{B}, \varphi_{k'}^{[j']} \in \mathcal{B}) \implies \varphi_{k'}^{[j']}(\mathbf{a}_k^{[j]}) = 0 \end{array} \right), \quad (\text{I.3})$$

puisque, d'après (I.1), $\mathbf{a}_k^{[j]}$ est aussi un noeud de niveau j' et, par $(\mathcal{P}_{\text{LI}})$, ce noeud est nécessairement différent de $\mathbf{a}_{k'}^{[j']}$ (sinon $j' = j$). Considérons une combinaison linéaire de fonctions de base φ appartenant à $\mathcal{B}$ telle que

$$\sum_{\varphi \in \mathcal{B}} \lambda_{\varphi} \varphi = 0, \quad (\text{I.4})$$

et supposons que $\mathcal{E} = \{\varphi \in \mathcal{B}; \lambda_{\varphi} \neq 0\}$ est non vide. Nous pouvons alors définir

$$j_m = \min\{j \in \llbracket 0, J \rrbracket; \exists k \in \llbracket 1, N_{\text{ddl}}^{[j]} \rrbracket \text{ tel que } \varphi_k^{[j]} \in \mathcal{E}\},$$

et sélectionner $k_m \in \llbracket 1, N_{\text{ddl}}^{[j_m]} \rrbracket$ tel que $\varphi_{k_m}^{[j_m]} \in \mathcal{E}$. Soit $j \in \llbracket 0, J \rrbracket$ et $k \in \llbracket 1, N_{\text{ddl}}^{[j]} \rrbracket$ tel que $\varphi_k^{[j]} \in \mathcal{E}$ et $(j, k) \neq (j_m, k_m)$. Nous avons nécessairement $\varphi_k^{[j]}(\mathbf{a}_{k_m}^{[j_m]}) = 0$ (le cas $j = j_m$ est évident puisqu'alors $k \neq k_m$ et le cas $j > j_m$ se déduit de (I.3)). Donc, en évaluant la combinaison linéaire (I.4) au noeud $\mathbf{a}_{k_m}^{[j_m]}$ nous obtenons $\lambda_{\varphi_{k_m}^{[j_m]}} = 0$. C'est une contradiction et le résultat est prouvé. ■

Au vu de la proposition I.25, nous donnons la définition suivante d'une base multiniveau.

Définition I.27 (Base multiniveau et espace d'approximation multiniveau)

Nous appelons base multiniveau un sous-ensemble $\mathcal{B}$ de $\bigcup_{j=0}^J B_j$ satisfaisant la propriété $(\mathcal{P}_{\text{LI}})$. Par la proposition I.25, cet ensemble est bien libre. De plus, un espace engendré par une base multiniveau est appelé espace d'approximation multiniveau.

Remarque I.28

Soit $\mathcal{V} = \text{vect } \mathcal{B}$ un espace d'approximation multiniveau et $u = \sum_{\varphi \in \mathcal{B}} u_{\varphi} \varphi \in \mathcal{V}$. La coordonnée u_{φ} de u dans la base multiniveau $\mathcal{B}$ n'est pas nécessairement la valeur de u au noeud associé à la fonction de base φ puisque une fonction de base (d'un niveau plus "grossier") intervenant dans la décomposition de u peut avoir une contribution non nulle en ce noeud.

Chapitre II

Procédure d'adaptation et préconditionneurs multigrilles.

La procédure d'adaptation consiste à ajouter ou supprimer certaines fonctions de base d'une base multiniveau donnée $\mathcal{B}^*$ pour produire une nouvelle base multiniveau $\mathcal{B}$ ayant une résolution spatiale mieux adaptée au problème. Les algorithmes de raffinement et déraffinement doivent produire des familles de fonctions de base linéairement indépendantes, et il est également souhaitable qu'aucune information ne soit perdue au cours du processus de raffinement, ceci signifie que lorsque $\mathcal{B}$ est obtenue à partir de $\mathcal{B}^*$ par le raffinement d'une fonction de base alors toute fonction de base représentable dans la base $\mathcal{B}^*$ l'est (encore) dans la base $\mathcal{B}$. La section II.1 est dédiée aux démonstrations de telles propriétés. En outre, nous montrons que des ensembles de fonctions de base peuvent être raffinés (resp. déraffinés) et que le résultats obtenu est indépendant de l'ordre dans lequel les fonctions de base sont raffinées (resp. déraffinées). Enfin, nous nous intéressons à des problématique plus pratique comme la règle "au-plus-un-niveau-de-différence" qui permet de garantir que la largeur de bande des matrices assemblées n'augmentera pas à cause de la procédure de raffinement et la manière dont doivent être appliquées les règles de quadrature pour calculer des intégrales faisant intervenir des fonctions discrétisées dans des espaces d'approximation multiniveau différents. La section II.2 explique ensuite comment la structure multiniveau obtenue par l'algorithme d'adaptation peut être utilisée pour construire des preconditionneurs multigrilles.

II.1 Adaptation

II.1.1 Procédures de raffinement/déraffinement

Etant donnée une base multiniveau $\mathcal{B}$, nous introduisons la notion de fonctions de base $\mathcal{B}$ -raffinée. Cette notion est complètement indépendante d'un quelconque historique de la procédure d'adaptation (que nous n'avons d'ailleurs pas encore définie). Néanmoins nous verrons (cf remarque II.5) qu'ils sont complètement compatibles.

Définition II.1 (Fonction de base $\mathcal{B}$ -raffinée)

Soit $\mathcal{B}$ une base multiniveau. Soit $j \in \llbracket 0, J \rrbracket$ et $k \in \llbracket 1, N_{ddl}^{[j]} \rrbracket$. La fonction de base $\varphi_k^{[j]}$ est dite $\mathcal{B}$ -raffinée ssi :

$$\exists j' \in \llbracket j+1, J \rrbracket, \exists k' \in \llbracket 1, N_{ddl}^{[j']} \rrbracket, \text{ tels que } \varphi_{k'}^{[j']} \in \mathcal{B} \text{ et } \mathbf{a}_k^{[j]} = \mathbf{a}_{k'}^{[j']}.$$

De plus, si la condition ci-dessus est vraie avec $j' = j+1$ alors nous dirons que la fonction de base $\varphi_k^{[j]}$ est $\mathcal{B}$ -raffinée une seule fois.

Remarque II.2

La propriété $(\mathcal{P}_{LI})$ implique que :

- les indices j' et k' sont nécessairement uniques,
- une fonction de base $\mathcal{B}$ -raffinée ne peut appartenir à $\mathcal{B}$.

Notons également que, d'après la proposition I.23, si $\varphi_k^{[j]}$ est $\mathcal{B}$ -raffinée une seule fois alors $\varphi_k^{[j]}$ est l'unique parent de $\varphi_{k'}^{[j']}$.

Nous donnons maintenant un lemme qui sera utilisé dans les démonstrations à venir.

Lemme II.3

Soient $j \in \llbracket 0, J-1 \rrbracket$ et $j' \in \llbracket 0, J \rrbracket$ tels que $j' \leq j$. Soit $(k, \ell, \ell') \in \llbracket 1, N_{ddl}^{[j]} \rrbracket \times \llbracket 1, N_{ddl}^{[j+1]} \rrbracket \times \llbracket 1, N_{ddi}^{[j']} \rrbracket$. Si $\mathbf{a}_\ell^{[j+1]} = \mathbf{a}_{\ell'}^{[j']}$ et $\varphi_k^{[j]}$ est un parent de $\varphi_\ell^{[j+1]}$ alors le noeud $\mathbf{a}_k^{[j]}$ est nécessairement à la même position géométrique que les noeuds $\mathbf{a}_{\ell'}^{[j']}$ et $\mathbf{a}_\ell^{[j+1]}$.

Démonstration : Puisque $j' \leq j$, d'après (I.1) et une récurrence directe, il existe $t \in \llbracket 1, N_{ddl}^{[j]} \rrbracket$ tel que $\mathbf{a}_t^{[j]} = \mathbf{a}_\ell^{[j+1]}$. Nous pouvons appliquer la proposition I.23 qui prouve que $\varphi_t^{[j]}$ est l'unique parent de $\varphi_\ell^{[j+1]}$. Donc, $\varphi_k^{[j]} = \varphi_t^{[j]}$ et par suite, $k = t$. ■

Nous pouvons maintenant décrire les procédures de raffinement et de déraffinement.

Algorithme II.4 (Dé/Raffinement quasi-hiérarchique)

Soit $\mathcal{B}^*$ une base multiniveau.

- Raffinement : Soit $\varphi_k^{[j]}$ une fonction de base appartenant à $\mathcal{B}^*$.

Raffiner la fonction de base donnée $\varphi_k^{[j]} \in \mathcal{B}^*$ consiste à produire une nouvelle base multiniveau $\mathcal{B}$ à partir de $\mathcal{B}^*$ en

- enlevant cette fonction de base $\varphi_k^{[j]}$,
- ajoutant tous ses enfants $\varphi_\ell^{[j+1]}$ qui ne sont pas $\mathcal{B}^*$ -raffinés.

Ceci peut s'écrire de manière compacte comme suit :

$$\mathcal{B} = \mathcal{B}^* \setminus \{\varphi_k^{[j]}\} \cup \{\text{enfants de } \varphi_k^{[j]} \text{ non } \mathcal{B}^*\text{-raffinés}\}.$$

- Déraffinement : Soit $\varphi_k^{[j]}$ une fonction de base $\mathcal{B}^*$ -raffinée une seule fois et sans enfant $\mathcal{B}^*$ -raffiné.

Déraffiner la fonction de base donnée $\varphi_k^{[j]} \notin \mathcal{B}^*$ consiste à produire une nouvelle base multiniveau $\mathcal{B}$ à partir de $\mathcal{B}^*$ en

- ajoutant cette fonction de base $\varphi_k^{[j]}$,
- enlevant les enfants de $\varphi_k^{[j]}$ n'ayant pas d'autre parent $\mathcal{B}^*$ -raffiné.

Ceci peut s'écrire de manière compacte comme suit :

$$\mathcal{B} = \mathcal{B}^* \setminus \{\text{enfants de } \varphi_k^{[j]} \text{ sans autre parent } \mathcal{B}^*\text{-raffiné}\} \cup \{\varphi_k^{[j]}\}.$$

Démonstration : Nous devons prouver que les algorithmes de raffinement et déraffinement ci-dessus produisent réellement des bases multiniveaux. Nous allons en effet montrer que $\mathcal{B}$ satisfait la propriété $(\mathcal{P}_{LI})$. Soient $\varphi_k^{[j]}$, $\varphi_{k'}^{[j']}$ deux fonctions de base appartenant à $\mathcal{B}$ telles que $\mathbf{a}_k^{[j]} = \mathbf{a}_{k'}^{[j']}$. Nous devons montrer que $j = j'$.

- Raffinement : supposons que $\mathcal{B}$ est obtenue à partir de $\mathcal{B}^*$ par raffinement d'une fonction de base (appartenant à $\mathcal{B}^*$), notée $\varphi_{k_0}^{[j_0]} \in \mathcal{B}^*$. Nous avons donc

$$\mathcal{B} = \mathcal{B}^* \setminus \{\varphi_{k_0}^{[j_0]}\} \cup \{\text{enfants de } \varphi_{k_0}^{[j_0]} \text{ non } \mathcal{B}^*\text{-raffinés}\}.$$

- si $\varphi_k^{[j]}$ et $\varphi_{k'}^{[j']}$ appartiennent tous deux à $\mathcal{B}^*$ alors, puisque $\mathcal{B}^*$ satisfait la propriété $(\mathcal{P}_{LI})$, nous obtenons $j = j'$.
- sinon, supposons par exemple que $\varphi_k^{[j]}$ n'appartient pas à $\mathcal{B}^*$. Par définition de $\mathcal{B}$, $\varphi_k^{[j]}$ est alors un enfant de $\varphi_{k_0}^{[j_0]}$ qui n'est pas $\mathcal{B}^*$ -raffiné, disons $\varphi_{\ell_0}^{[j_0+1]}$, i.e. $k = \ell_0$, $j = j_0 + 1$.
- si $\varphi_{k'}^{[j']}$ $\in \mathcal{B}^*$ alors, puisque $\varphi_{\ell_0}^{[j_0+1]}$ n'est pas $\mathcal{B}^*$ -raffinée, nous avons $j_0 + 1 \geq j'$. Supposons que

- $j_0 + 1 > j'$. Le lemme II.3 donne $\mathbf{a}_{k_0}^{[j_0]} = \mathbf{a}_{k'}^{[j']}$. Puisque $\mathcal{B}^*$ satisfait la propriété $(\mathcal{P}_{LI})$, nous obtenons $j_0 = j'$ et $k_0 = k'$, mais $\varphi_{k_0}^{[j_0]} \notin \mathcal{B}$ et $\varphi_{k'}^{[j']} \in \mathcal{B}$. C'est une contradiction et par suite $j = j_0 + 1 = j'$.
- sinon, $\varphi_{k'}^{[j']}$ est un enfant de $\varphi_{k_0}^{[j_0]}$, disons $\varphi_{\ell_0'}^{[j_0+1]}$, i.e. $k' = \ell_0'$ et $j = j_0 + 1$. Donc, nous avons $j = j'$.
 - Déraffinement : Supposons que $\mathcal{B}$ est obtenue à partir de $\mathcal{B}^*$ par le déaffinement d'une fonction de base ($\mathcal{B}^*$ -raffinée une seule fois), notée $\varphi_{k_0}^{[j_0]}$. Nous avons donc

$$\mathcal{B} = \mathcal{B}^* \setminus \{\text{enfants de } \varphi_{k_0}^{[j_0]} \text{ sans autre parent } \mathcal{B}^*\text{-raffiné}\} \cup \{\varphi_{k_0}^{[j_0]}\}.$$

- si $\varphi_k^{[j]}$ et $\varphi_{k'}^{[j']}$ appartiennent toutes deux à $\mathcal{B}^*$ alors, puisque $\mathcal{B}^*$ satisfait la propriété $(\mathcal{P}_{LI})$, nous avons $j = j'$.
- sinon, supposons par exemple que $\varphi_k^{[j]}$ n'appartient pas à $\mathcal{B}^*$. Par définition de $\mathcal{B}$, nous avons $\varphi_k^{[j]} = \varphi_{k_0}^{[j_0]}$, i.e. $k = k_0$, $j = j_0$. En raisonnant par l'absurde, supposons que $j' \neq j$. Nous avons $j' \neq j_0$ et donc, $\varphi_{k'}^{[j']} \in \mathcal{B}^*$. Cependant, $\varphi_{k_0}^{[j_0]}$ est une fonction de base raffinée une seule fois, donc il existe $\ell_0 \in \llbracket 1, N_{\text{ddl}}^{[j_0]} \rrbracket$ tel que $\varphi_{\ell_0}^{[j_0+1]} \in \mathcal{B}^*$ et $\mathbf{a}_{\ell_0}^{[j_0+1]} = \mathbf{a}_{k_0}^{[j_0]}$. D'après la proposition I.23, $\mathbf{a}_{\ell_0}^{[j_0+1]} = \mathbf{a}_{k_0}^{[j_0]}$ implique que $\varphi_{k_0}^{[j_0]}$ est l'unique parent de $\varphi_{\ell_0}^{[j_0+1]}$. Ainsi, puisque $\varphi_{\ell_0}^{[j_0+1]}$ est un enfant de $\varphi_{k_0}^{[j_0]}$ sans autre parent $\mathcal{B}^*$ -raffiné, nous obtenons $\varphi_{\ell_0}^{[j_0+1]} \notin \mathcal{B}$. Cependant, puisque $\mathcal{B}^*$ satisfait la propriété $(\mathcal{P}_{LI})$ et puisque $\mathbf{a}_{\ell_0}^{[j_0+1]} = \mathbf{a}_{k'}^{[j']}$, nous avons $j' = j_0 + 1$ et $k' = \ell_0$. Finalement, nous obtenons $\varphi_{k'}^{[j']} \notin \mathcal{B}$. C'est une contradiction et par suite $j = j'$. ■

Remarque II.5

Cet algorithme est cohérent avec la définition II.1. En effet, d'après la proposition I.23,

- si $\mathcal{B}$ est obtenue à partir de $\mathcal{B}^*$ par raffinement d'une fonction de base $\varphi_k^{[j]}$, alors $\varphi_k^{[j]}$ est une fonction de base $\mathcal{B}$ -raffinée au sens de la définition II.1 ;
- si $\mathcal{B}$ est obtenue à partir de $\mathcal{B}^*$ par déaffinement de la fonction de base $\varphi_k^{[j]}$ $\mathcal{B}^*$ -raffinée une seule fois alors $\varphi_k^{[j]}$ n'est plus une fonction de base $\mathcal{B}$ -raffinée au sens de la définition II.1.

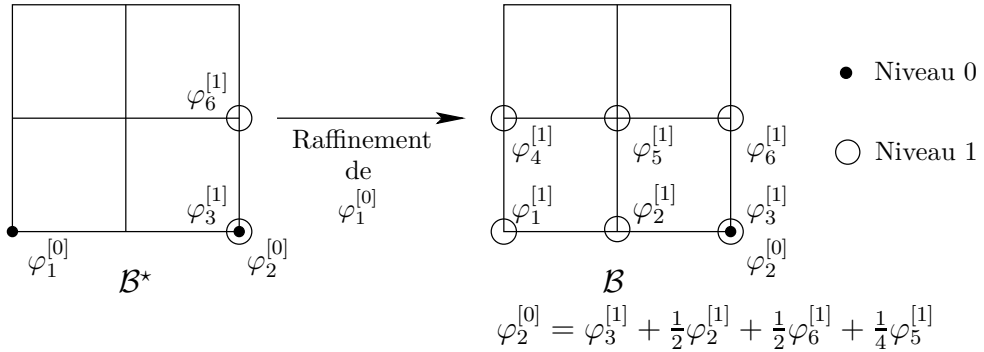


FIG. II.1 – Le raffinement ne préserve pas l'indépendance linéaire des ensembles multiniveaux ne satisfaisant pas $(\mathcal{P}_{LI})$.

Remarque II.6

Les procédures de raffinement et de déaffinement décrites dans l'algorithme II.4 ne préservent pas en général l'indépendance linéaire des ensembles de fonctions de base multiniveaux $\bigcup_{j=0}^J \tilde{\mathcal{B}}_j$ (avec $\tilde{\mathcal{B}}_j \subset \mathcal{B}_j$) ne satisfaisant pas la propriété $(\mathcal{P}_{LI})$. Un exemple est donné sur la figure II.1. La famille $\mathcal{B}^*$ représentée sur la figure de gauche est une famille libre (mais ne satisfait pas $(\mathcal{P}_{LI})$) alors que la famille $\mathcal{B}$, représentée sur la figure de droite, obtenue à partir de $\mathcal{B}^*$ par le raffinement de $\varphi_1^{[0]}$, n'est pas une famille libre car $\varphi_2^{[0]}$ s'exprime comme une combinaison linéaire des fonctions de base $\varphi_2^{[1]}$, $\varphi_3^{[1]}$, $\varphi_5^{[1]}$, $\varphi_6^{[1]}$.

II.1.2 Conservation de l'information

Une propriété souhaitable de la procédure de raffinement est qu'elle permette de conserver l'information. Cela signifie que, si $\mathcal{B}$ est obtenue à partir de $\mathcal{B}^*$ par le raffinement d'une fonction de base alors $\text{vect } \mathcal{B}^* \subset \text{vect } \mathcal{B}$, *i.e.* la base raffinée $\mathcal{B}$ autorise la représentation exacte de chaque fonction de la base originale $\mathcal{B}^*$. Cependant, la procédure de raffinement décrite dans l'algorithme II.4 n'est pas "sans perte". Un exemple est donné par la figure II.2. Néanmoins, nous pouvons prouver le résultat suivant :

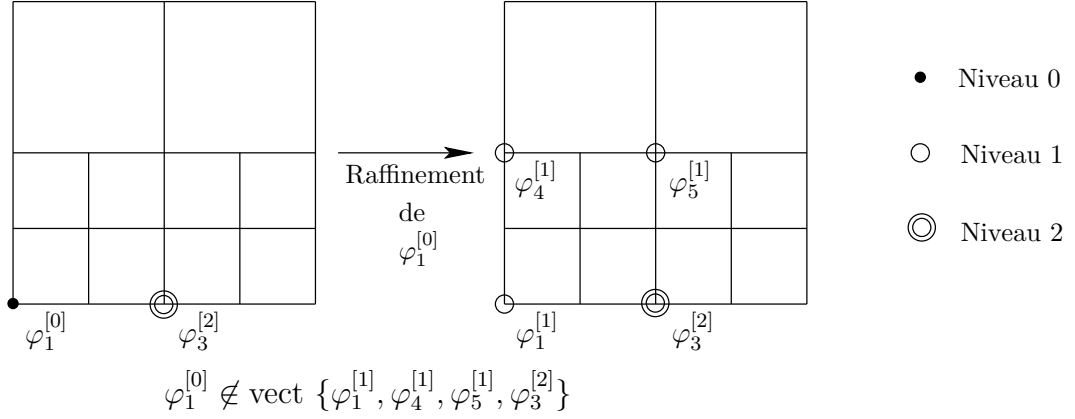


FIG. II.2 – Le raffinement d'une base multiniveau ne s'effectue pas nécessairement "sans perte".

Proposition II.7

Soit $\mathcal{B}$ une base multiniveau satisfaisant la propriété suivante :

Tout enfant d'une fonction de base $\mathcal{B}$ -raffinée est soit lui même $\mathcal{B}$ -raffiné, soit dans $\mathcal{B}$. ($\mathcal{P}_{LO}$)

Alors, toutes les fonctions de base $\mathcal{B}$ -raffinées appartiennent à $\text{vect } \mathcal{B}$.

Démonstration : Par une récurrence sur le niveau j des fonctions de base, nous prouvons l'énoncé (H_j) suivant pour tout $j \in \llbracket 0, J \rrbracket$:

Toute fonction de base $\mathcal{B}$ -raffinée de niveau j appartient à $\text{vect } \mathcal{B}$. (H_j)

Une fonction de base de niveau J ne peut pas être $\mathcal{B}$ -raffinée. Donc, l'énoncé (H_J) est vrai.

Soit $j \in \llbracket 0, J-1 \rrbracket$. Supposons que l'énoncé (H_{j+1}) est vrai et soit $\varphi_k^{[j]}$ une fonction de base raffinée de niveau j . D'après la proposition I.18, nous avons

$$\varphi_k^{[j]} = \sum_{\substack{\ell | \varphi_\ell^{[j+1]} \text{ est} \\ \text{un enfant de } \varphi_k^{[j]}}} \beta_{k\ell}^{[j]} \varphi_\ell^{[j+1]}.$$

De plus, la propriété $(\mathcal{P}_{LO})$ implique que toute fonction de base $\varphi_\ell^{[j+1]}$ intervenant dans la somme ci-dessus est soit dans $\mathcal{B}$ soit $\mathcal{B}$ -raffinée. Or lorsque $\varphi_\ell^{[j+1]}$ est $\mathcal{B}$ -raffinée, l'hypothèse de récurrence (H_{j+1}) mène à $\varphi_\ell^{[j+1]} \in \text{vect } \mathcal{B}$. Ainsi, $\varphi_k^{[j]} \in \text{vect } \mathcal{B}$ et la récurrence est établie. ■

Théorème II.8

Soit $\mathcal{B}$ une base multiniveau satisfaisant la propriété $(\mathcal{P}_{LO})$, et obtenue à partir d'une base multiniveau $\mathcal{B}^*$ par la procédure de raffinement (cf algorithme II.4). Alors, nous avons

$$\text{vect } \mathcal{B}^* \subset \text{vect } \mathcal{B}.$$

Démonstration : Supposons que $\mathcal{B}$ est obtenue à partir de $\mathcal{B}^*$ par le raffinement d'une fonction de base $\varphi_{k_0}^{[j_0]} \in \mathcal{B}^*$ et soit $\varphi_k^{[j]} \in \mathcal{B}^*$.

- si $\varphi_k^{[j]} \neq \varphi_{k_0}^{[j_0]}$ alors $\varphi_k^{[j]} \in \mathcal{B}$,
- sinon $\varphi_k^{[j]} = \varphi_{k_0}^{[j_0]}$ qui est $\mathcal{B}$ -raffinée. La proposition II.7 garantit que $\varphi_k^{[j]} \in \text{vect } \mathcal{B}$.

Alors $\text{vect } \mathcal{B}^* \subset \text{vect } \mathcal{B}$. ■

De plus, les procédures de raffinement et de déraffinement préservent la propriété $(\mathcal{P}_{\text{LO}})$. Commençons par démontrer le lemme suivant :

Lemme II.9

Soient $\mathcal{B}^$ et $\mathcal{B}$ deux bases multiniveaux. Soient $\varphi_{k^*}^{[j^*]}$ une fonction de base $\mathcal{B}^*$ -raffinée et $\varphi_k^{[j]}$ une fonction de base $\mathcal{B}$ -raffinée.*

- 1) *Si $\mathcal{B}$ est obtenue à partir de $\mathcal{B}^*$ par le raffinement d'une fonction de base $\varphi_{k_0}^{[j_0]} \in \mathcal{B}^*$ alors*
 - (i) $\varphi_{k^*}^{[j^*]}$ est aussi $\mathcal{B}$ -raffinée,
 - (ii) $\varphi_k^{[j]}$ est soit $\mathcal{B}^*$ -raffinée soit égale à $\varphi_{k_0}^{[j_0]}$.
- 2) *Si $\mathcal{B}$ est obtenue à partir de $\mathcal{B}^*$ par déraffinement d'une fonction de base $\mathcal{B}^*$ -raffinée une seule fois notée $\varphi_{k_0}^{[j_0]}$ alors*
 - (i) $\varphi_{k^*}^{[j^*]}$ est soit $\mathcal{B}$ -raffinée soit égale à $\varphi_{k_0}^{[j_0]}$,
 - (ii) $\varphi_k^{[j]}$ est aussi $\mathcal{B}^*$ -raffinée.
 - (iii) *De plus, si nous supposons que $\varphi_{k^*}^{[j^*]}$ est $\mathcal{B}^*$ -raffinée une seule fois alors $\varphi_{k^*}^{[j^*]}$ est soit également $\mathcal{B}$ -raffinée une seule fois soit égale à $\varphi_{k_0}^{[j_0]}$.*

Démonstration : Puisque $\varphi_{k^*}^{[j^*]}$ est $\mathcal{B}^*$ -raffinée, il existe $j^{*'} > j^*$ et $k^{*'} \in \llbracket 1, N_{\text{ddl}}^{[j^{*'}]} \rrbracket$ tels que $\varphi_{k^{*'}}^{[j^{*'}]} \in \mathcal{B}^*$ et $\mathbf{a}_{k^*}^{[j^*]} = \mathbf{a}_{k^{*'}}^{[j^{*'}]}$. De même, puisque $\varphi_k^{[j]}$ est $\mathcal{B}$ -raffinée, il existe $j' > j$ et $k' \in \llbracket 1, N_{\text{ddl}}^{[j']} \rrbracket$ tels que $\varphi_{k'}^{[j']} \in \mathcal{B}$ et $\mathbf{a}_k^{[j]} = \mathbf{a}_{k'}^{[j']}$.

- 1) Raffinement de $\varphi_{k_0}^{[j_0]}$.
 - (i) Si $\varphi_{k^{*'}}^{[j^{*'}]} \in \mathcal{B}$ alors $\varphi_{k^*}^{[j^*]}$ est $\mathcal{B}$ -raffinée; sinon $\varphi_{k^{*'}}^{[j^{*'}]} = \varphi_{k_0}^{[j_0]}$ et puisque $\varphi_{k_0}^{[j_0]}$ est $\mathcal{B}$ -raffinée, $\varphi_{k^*}^{[j^*]}$ est *a fortiori* $\mathcal{B}$ -raffinée.
 - (ii) Si $\varphi_{k'}^{[j']} \in \mathcal{B}^*$ alors $\varphi_k^{[j]}$ est $\mathcal{B}^*$ -raffinée et l'énoncé est prouvé. Sinon, $\varphi_{k'}^{[j']}$ est un enfant de $\varphi_{k_0}^{[j_0]}$, disons $\varphi_{\ell_0}^{[j_0+1]}$, i.e. $j' = j_0 + 1$, $k' = \ell_0$. Nous avons alors $\mathbf{a}_k^{[j]} = \mathbf{a}_{\ell_0}^{[j_0+1]}$ et $j_0 + 1 > j$, et donc d'après le lemme II.3, $\mathbf{a}_k^{[j]} = \mathbf{a}_{k_0}^{[j_0]}$. Si $j_0 = j$ alors $\varphi_k^{[j]} = \varphi_{k_0}^{[j_0]}$; sinon $j_0 > j$, puisque $\varphi_{k_0}^{[j_0]} \in \mathcal{B}^*$, $\varphi_k^{[j]}$ est $\mathcal{B}^*$ -raffinée.
- 2) Déraffinement de $\varphi_{k_0}^{[j_0]}$.
 - (i) Si $\varphi_{k^{*'}}^{[j^{*'}]} \in \mathcal{B}$ alors $\varphi_{k^*}^{[j^*]}$ est $\mathcal{B}$ -raffinée et l'énoncé est prouvé. Sinon, $\varphi_{k^{*'}}^{[j^{*'}]}$ est un enfant de $\varphi_{k_0}^{[j_0]}$, noté $\varphi_{\ell_0}^{[j_0+1]}$, i.e. $j^{*'} = j_0 + 1$, $k^{*'} = \ell_0$. Nous avons alors $\mathbf{a}_{k^*}^{[j^*]} = \mathbf{a}_{\ell_0}^{[j_0+1]}$ et $j_0 + 1 > j^*$. Ainsi, d'après le lemme II.3, $\mathbf{a}_{k^*}^{[j^*]} = \mathbf{a}_{k_0}^{[j_0]}$. Si $j_0 = j^*$ alors $\varphi_{k^*}^{[j^*]} = \varphi_{k_0}^{[j_0]}$; sinon $j_0 > j^*$, puisque $\varphi_{k_0}^{[j_0]} \in \mathcal{B}$, $\varphi_{k^*}^{[j^*]}$ est $\mathcal{B}$ -raffinée.
 - (ii) Si $\varphi_{k'}^{[j']} \in \mathcal{B}^*$ alors $\varphi_k^{[j]}$ est $\mathcal{B}^*$ -raffinée; sinon $\varphi_{k'}^{[j']} = \varphi_{k_0}^{[j_0]}$ et puisque $\varphi_{k_0}^{[j_0]}$ est $\mathcal{B}^*$ -raffinée, $\varphi_k^{[j]}$ est *a fortiori* $\mathcal{B}^*$ -raffinée.
 - (iii) Ici, nous supposons que $j^{*'} = j^* + 1$. Si $\varphi_{k^{*'}}^{[j^{*'}]} \in \mathcal{B}$ alors $\varphi_{k^*}^{[j^*]}$ est $\mathcal{B}$ -raffinée une seule fois et l'énoncé est prouvé. Sinon, $\varphi_{k^{*'}}^{[j^{*'}]}$ est un enfant de $\varphi_{k_0}^{[j_0]}$. Cependant, d'après la remarque II.2, $\varphi_{k^*}^{[j^*]}$ est l'unique parent de $\varphi_{k^{*'}}^{[j^{*'}]}$. Ainsi, nous obtenons $\varphi_{k^*}^{[j^*]} = \varphi_{k_0}^{[j_0]}$. ■

Proposition II.10

Les procédures de raffinement et déraffinement des bases multiniveaux (cf algorithme II.4) préservent la propriété $(\mathcal{P}_{\text{LO}})$.

Démonstration : Supposons que la base multiniveau $\mathcal{B}^*$ satisfasse la propriété $(\mathcal{P}_{\text{LO}})$. Soient $\varphi_k^{[j]}$ une fonction de base $\mathcal{B}$ -raffinée et $\varphi_{\ell}^{[j+1]}$ un enfant de $\varphi_k^{[j]}$. Nous devons montrer que : soit $\varphi_{\ell}^{[j+1]}$ appartient à $\mathcal{B}$ soit $\varphi_{\ell}^{[j+1]}$ est $\mathcal{B}$ -raffinée.

- Raffinement : Supposons que la base multiniveau $\mathcal{B}$ est obtenue à partir de $\mathcal{B}^*$ par le raffinement de la fonction de base $\varphi_{k_0}^{[j_0]} \in \mathcal{B}^*$. D'après le lemme II.9 propriété 1) (ii), soit $\varphi_k^{[j]}$ est $\mathcal{B}^*$ -raffinée soit $\varphi_k^{[j]}$ est égale à $\varphi_{k_0}^{[j_0]}$. Considérons les deux cas :

- Si $\varphi_k^{[j]}$ est $\mathcal{B}^*$ -raffinée. Puisque $\mathcal{B}^*$ satisfait la propriété $(\mathcal{P}_{LO})$, seulement deux cas sont possibles :
 - $\varphi_\ell^{[j+1]}$ est $\mathcal{B}^*$ -raffinée. D'après le lemme II.9 propriété 1)(i), cela implique que $\varphi_\ell^{[j+1]}$ est $\mathcal{B}$ -raffinée.
 - $\varphi_\ell^{[j+1]} \in \mathcal{B}^*$. Et alors, $\varphi_\ell^{[j+1]} \in \mathcal{B}$ ou $\varphi_\ell^{[j+1]} = \varphi_{k_0}^{[j_0]}$ qui est $\mathcal{B}$ -raffinée.
- Si $\varphi_k^{[j]} = \varphi_{k_0}^{[j_0]}$, tous ses enfants sont soit dans $\mathcal{B}$ soit $\mathcal{B}^*$ -raffinés. D'après le lemme II.9 propriété 1)(i), ils sont soit dans $\mathcal{B}$ soit $\mathcal{B}$ -raffinés.
- Déraffinement : Supposons que $\mathcal{B}$ est obtenue à partir de $\mathcal{B}^*$ par déraffinement d'une fonction de base raffinée une seule fois $\varphi_{k_0}^{[j_0]}$ de $\mathcal{B}^*$. D'après le lemme II.9 propriété 2) (ii), $\varphi_k^{[j]}$ est $\mathcal{B}^*$ -raffinée. Puisque $\mathcal{B}^*$ satisfait la propriété $(\mathcal{P}_{LO})$, seulement deux cas sont possibles :
 - $\varphi_\ell^{[j+1]}$ est $\mathcal{B}^*$ -raffinée. D'après le lemme II.9 propriété 2)(i), cela implique que $\varphi_\ell^{[j+1]}$ est $\mathcal{B}$ -raffinée ou $\varphi_\ell^{[j+1]} = \varphi_{k_0}^{[j_0]} \in \mathcal{B}$.
 - $\varphi_\ell^{[j+1]} \in \mathcal{B}^*$. Et donc, nous avons soit $\varphi_\ell^{[j+1]} \in \mathcal{B}$ soit $\varphi_\ell^{[j+1]}$ est un enfant de $\varphi_{k_0}^{[j_0]}$ sans autre parent $\mathcal{B}^*$ -raffiné. Le dernier cas est impossible. En effet, $\varphi_k^{[j]}$ est $\mathcal{B}^*$ -raffinée et alors, par unicité, nous devrions avoir $\varphi_k^{[j]} = \varphi_{k_0}^{[j_0]}$. Cependant, $\varphi_{k_0}^{[j_0]}$ n'est pas $\mathcal{B}$ -raffinée. C'est une contradiction et nous obtenons $\varphi_\ell^{[j+1]} \in \mathcal{B}$. ■

Remarque II.11

Notons que ces propriétés $(\mathcal{P}_{LI})$ ou $(\mathcal{P}_{LO})$ ne sont pas très restrictives puisqu'elles sont préservées par les procédures de (dé)raffinement et qu'elles sont de manière évidente vérifiées par la base grossière $\mathcal{B}_0$ qui est utilisée en pratique comme point de départ de l'algorithme.

II.1.3 Procédure d'adaptation

La procédure d'adaptation consiste à raffiner (resp. déraffiner) un ensemble de fonction de base. Ceci est tout à fait possible et le résultat est indépendant de l'ordre dans lequel les fonctions de base sont raffinées (resp. déraffinées).

Proposition II.12

Soit $\mathcal{B}^$ une base multiniveau.*

- 1) *Soit $\mathcal{E}^* \subset \mathcal{B}^*$. Il est possible de raffiner successivement toutes les fonctions de base appartenant à $\mathcal{E}^*$, produisant ainsi une base multiniveau $\mathcal{B}$ qui est indépendante de l'ordre dans lequel les fonctions de base de $\mathcal{E}^*$ ont été raffinées.*

Nous disons que $\mathcal{B}$ est obtenue à partir de $\mathcal{B}^$ par raffinement de l'ensemble des fonctions de base $\mathcal{E}^*$.*

- 2) *Soit $\mathcal{F}^*$ un ensemble de fonctions de base $\mathcal{B}^*$ -raffinées une seule fois et qui n'ont pas d'enfant $\mathcal{B}^*$ -raffiné. Il est possible de déraffiner successivement toutes les fonctions de base appartenant à $\mathcal{F}^*$, produisant ainsi une base multiniveau $\mathcal{B}$ qui est indépendante de l'ordre dans lequel les fonctions de base de $\mathcal{F}^*$ ont été déraffinées.*

Nous disons que $\mathcal{B}$ est obtenue à partir de $\mathcal{B}^$ par déraffinement de l'ensemble des fonctions de base $\mathcal{F}^*$.*

Démonstration : Dans les deux cas, nous devons d'abord montrer que des (dé)raffinements successifs sont possibles et ensuite que les bases multiniveaux $\mathcal{B}$ sont indépendantes de l'ordre dans lequel les fonctions de base sont (dé)raffinées.

- 1) – Soit $\varphi \in \mathcal{E}^*$. L'ensemble $\mathcal{E}^* \setminus \{\varphi\}$ est inclus dans la base multiniveau produite par le raffinement de φ puisque, dans cette procédure, seulement φ est supprimée de $\mathcal{B}^*$. Donc, toutes les fonctions de base de $\mathcal{E}^*$ peuvent être raffinées successivement.
- Il est suffisant de montrer que $\mathcal{B}$ est indépendante de l'ordre dans lequel les fonctions de base sont raffinées dans le cas où $\#\mathcal{E}^* = 2$, disons $\mathcal{E}^* = \{\varphi, \psi\}$. Nous notons $\overline{\mathcal{B}}$ la base multiniveau obtenue par le raffinement de φ dans $\mathcal{B}^*$ et $\mathcal{B}$ la base multiniveau obtenue par le raffinement de ψ dans $\overline{\mathcal{B}}$. Par définition, nous avons

$$\overline{\mathcal{B}} = \mathcal{B}^* \setminus \{\varphi\} \cup \{\text{enfants de } \varphi \text{ non } \mathcal{B}^* \text{-raffinés}\} \quad (\text{II.1})$$

et

$$\mathcal{B} = \overline{\mathcal{B}} \setminus \{\psi\} \cup \{\text{enfants de } \psi \text{ non } \overline{\mathcal{B}} \text{-raffinés}\}. \quad (\text{II.2})$$

En appliquant le lemme II.9 propriété 1) (i) et (ii), nous obtenons

$$\{\text{fonctions de base } \overline{\mathcal{B}} \text{-raffinées}\} = \{\text{fonctions de base } \mathcal{B}^* \text{-raffinées}\} \cup \{\varphi\}. \quad (\text{II.3})$$

Donc, en combinant (II.2) et (II.3), il vient

$$\mathcal{B} = \overline{\mathcal{B}} \setminus \{\psi\} \cup \{\text{enfants de } \psi \text{ non } \mathcal{B}^* \text{-raffinés et différents de } \varphi\}. \quad (\text{II.4})$$

Enfin, les égalités (II.1) et (II.4) donnent

$$\mathcal{B} = (\mathcal{B}^* \cup \{\text{enfants de } \varphi \text{ et } \psi \text{ qui ne sont pas } \mathcal{B}^* \text{-raffinés}\}) \setminus \{\varphi, \psi\}.$$

Cette expression montre que $\mathcal{B}$ ne dépend pas de l'ordre dans lequel les fonctions de base φ et ψ sont raffinées.

- 2) – Soit $\varphi \in \mathcal{F}^*$. Notons $\overline{\mathcal{B}}$ la base multiniveau obtenue par déraffinement de φ dans $\mathcal{B}^*$. Nous devons prouver que toute fonction de base de $\mathcal{F}^* \setminus \{\varphi\}$ peut être déraffinée dans $\overline{\mathcal{B}}$, *i.e.* que toute fonction de base de $\mathcal{F}^* \setminus \{\varphi\}$ est $\overline{\mathcal{B}}$ -raffinée une seule fois et n'a pas d'enfant $\overline{\mathcal{B}}$ -raffiné. Soit $\psi \in \mathcal{F}^* \setminus \{\varphi\}$. D'après le lemme II.9 propriété 2) (iii), puisque ψ est une fonction de base $\mathcal{B}^*$ -raffinée une seule fois et différente de φ , ψ est $\overline{\mathcal{B}}$ -raffinée une seule fois. En outre, si ψ a un enfant $\overline{\mathcal{B}}$ -raffiné alors en appliquant le lemme II.9 propriété 2) (i), cet enfant est aussi $\mathcal{B}^*$ -raffiné. C'est une contradiction et l'énoncé est prouvé.
- Il est suffisant de prouver que $\mathcal{B}$ est indépendante de l'ordre dans lequel les fonctions de base sont déraffinées dans le cas où $\#\mathcal{F}^* = 2$, disons $\mathcal{F}^* = \{\varphi, \psi\}$. Notons $\overline{\mathcal{B}}$ la base multiniveau obtenue à partir de $\mathcal{B}^*$ par déraffinement de φ et $\mathcal{B}$ la base multiniveau obtenue à partir de $\overline{\mathcal{B}}$ par déraffinement de ψ . Par définition, nous avons

$$\overline{\mathcal{B}} = \mathcal{B}^* \setminus \{\text{enfant de } \varphi \text{ qui n'ont pas d'autre parent } \mathcal{B}^* \text{-raffiné}\} \cup \{\varphi\} \quad (\text{II.5})$$

et

$$\mathcal{B} = \overline{\mathcal{B}} \setminus \{\text{enfant de } \psi \text{ qui n'ont pas d'autre parent } \overline{\mathcal{B}} \text{-raffiné}\} \cup \{\psi\}. \quad (\text{II.6})$$

En appliquant le lemme II.9 propriété 2) (i) et (ii), nous obtenons

$$\{\text{fonctions de base } \mathcal{B}^* \text{-raffinées}\} = \{\text{fonctions de base } \overline{\mathcal{B}} \text{-raffinées}\} \cup \{\varphi\}. \quad (\text{II.7})$$

Ainsi, en combinant (II.6) et (II.7), il vient

$$\mathcal{B} = \overline{\mathcal{B}} \setminus \{\text{enfants de } \psi \text{ qui n'ont pas de parent } \mathcal{B}^* \text{-raffiné à l'exception possible de } \varphi\} \cup \{\psi\}. \quad (\text{II.8})$$

Les égalités (II.5) et (II.8) donnent

$$\mathcal{B} = \mathcal{B}^* \setminus \{\text{enfants de } \varphi \text{ et } \psi \text{ qui n'ont pas d'autre parent } \mathcal{B}^* \text{-raffiné à l'exception possible de } \varphi \text{ et } \psi\} \cup \{\varphi, \psi\}.$$

Cette expression montre que $\mathcal{B}$ ne dépend pas de l'ordre dans lequel les fonctions de base φ et ψ ont été déraffinées. ■

Cette définition nous permet d'exprimer l'algorithme d'adaptation.

Algorithme II.13 (Procédure d'adaptation)

Soit $\mathcal{B}^$ une base multiniveau. Supposons que, grâce au critère de (dé)raffinement, les ensembles $\mathcal{E}^* \subset \mathcal{B}^*$ et $\mathcal{F}^*$ de fonctions de base à raffiner et à déraffiner soient donnés. La procédure d'adaptation est composée des deux étapes suivantes :*

- 1) *Raffinement de l'ensemble $\mathcal{E}^*$, produisant ainsi une nouvelle base multiniveau $\overline{\mathcal{B}}$.*
- 2) *Déraffinement de l'ensemble des fonctions de base de $\mathcal{F}^*$ qui sont $\overline{\mathcal{B}}$ -raffinées une seule fois et sans enfant $\overline{\mathcal{B}}$ -raffiné.*

Nous avons maintenant décrit complètement les algorithmes d'adaptation. Nous nous intéressons dans les sections suivantes à des problèmes plus pratiques tels que la taille de la bande de matrices obtenues après l'assemblage éléments finis et l'intégration numérique.

II.1.4 Règles “au-plus-un-niveau-de-différence”

Lorsque les espaces multiniveaux sont utilisés comme espaces d'approximation dans une méthode de Galerkin, il peut être intéressant d'imposer une condition sur le nombre de niveaux de raffinement séparant deux fonctions de base dont les supports s'intersectent. En effet, ceci permet de contrôler la largeur de bande de la matrice

de rigidité (*i.e.* nombre d'éléments non nuls par ligne) et ainsi garantir la structure creuse de cette matrice. Malheureusement, l'exemple donné sur la figure II.3 montre une situation simple où l'algorithme d'adaptation II.13 mène à la construction d'espaces d'approximation vect $\mathcal{B}_k$ ($k \in \mathbb{N}$) engendrés par $3k + 4$ fonctions de base dont les supports s'intersectent deux à deux. Le domaine est choisi carré, le maillage initial est formé d'une seule maille et la base multiniveau $\mathcal{B}_0$ constituée des quatre fonctions de base de niveau 0. La base multiniveau $\mathcal{B}_{k+1}$ est ensuite obtenue par raffinement de la fonction base (de niveau k) appartenant à $\mathcal{B}_k$ associée au noeud (indiqué par une flèche sur la figure) placé dans le coin en bas à gauche du domaine. Les matrices de rigidité associées à ces espaces d'approximation sont pleines (elles ne possèdent aucun élément nul *a priori*). Nous ajoutons alors des règles pratiques à l'algorithme d'adaptation II.13 pour éviter d'aboutir à ce type de configuration. Ces règles ont pour effet d'augmenter (resp. réduire) le nombre de fonctions de base effectivement raffinées (resp. déraffinées).

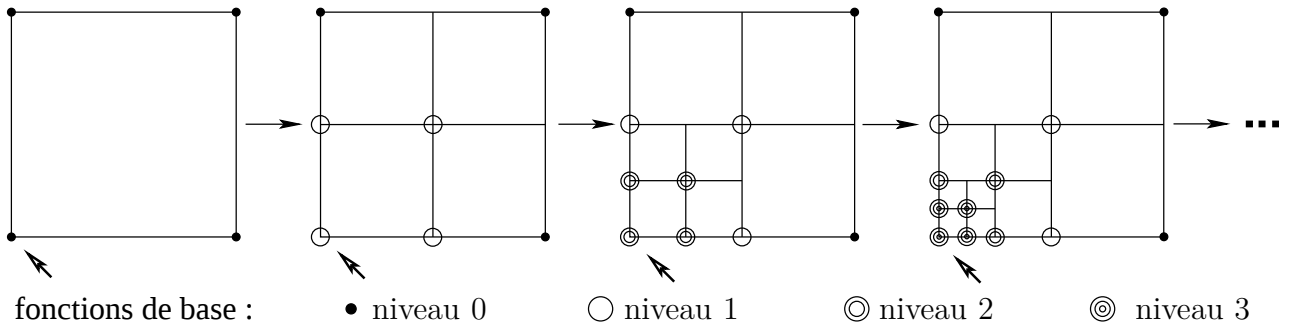


FIG. II.3 – Exemple d'espaces multiniveaux conduisant à assembler des matrices de rigidité pleines.

Pour cela nous introduisons la notion de descendants et d'ascendants d'une fonction de base $\varphi^{[j]}$. Au même titre que les enfants (resp. parents), les descendants (resp. ascendants) de la fonction de base $\varphi^{[j]}$ de niveau j sont des fonctions de base de niveau $j + 1$ (resp. $j - 1$). Nous les définissons comme suit :

Définition II.14 (Descendants et ascendants)

Soit $\varphi^{[j]}$ une fonction de base de niveau j . Nous appelons descendant (resp. ascendant) de la fonction de base $\varphi^{[j]}$ toute fonction de base de niveau $j + 1$ (resp. $j - 1$) dont le support intersecte celui de $\varphi^{[j]}$.

L'énoncé suivant précise alors le rôle de ces fonctions de base.

Algorithme II.15 (Règle “au-plus-un-niveau-de-différence”)

Reprenons les notations de l'algorithme II.13. Nous ajoutons à cet algorithme, les règles pratiques suivantes :

- 1) Lors du raffinement de $\varphi \in \mathcal{E}^*$, nous raffinons également tous ses ascendants. Ce principe s'applique récursivement.
- 2) Une fonction de base $\varphi \in \mathcal{F}^*$ peut être dérafinée seulement si aucun de ses descendants n'est $\overline{\mathcal{B}}$ -raffiné.

Les règles pratiques énoncées dans l'algorithme II.15 permettent de garantir que l'application de l'algorithme d'adaptation II.13 ne conduira pas à une augmentation de l'écart entre les niveaux des fonctions de base des espaces d'approximation produits. A titre d'exemple, reprenons la séquence de raffinement présentée sur la figure II.3 mais appliquons maintenant les règles pratiques énoncées dans l'algorithme II.15. Nous obtenons la séquence de raffinement présentée sur la figure II.4. Les règles II.15 forcent le raffinement de fonctions de base supplémentaires (marquées par des flèches en pointillés). Les bases multiniveaux contiennent plus de fonctions de base que celles présentées sur la figure II.3 mais les matrices assemblées contiennent moins de coefficients non nuls.

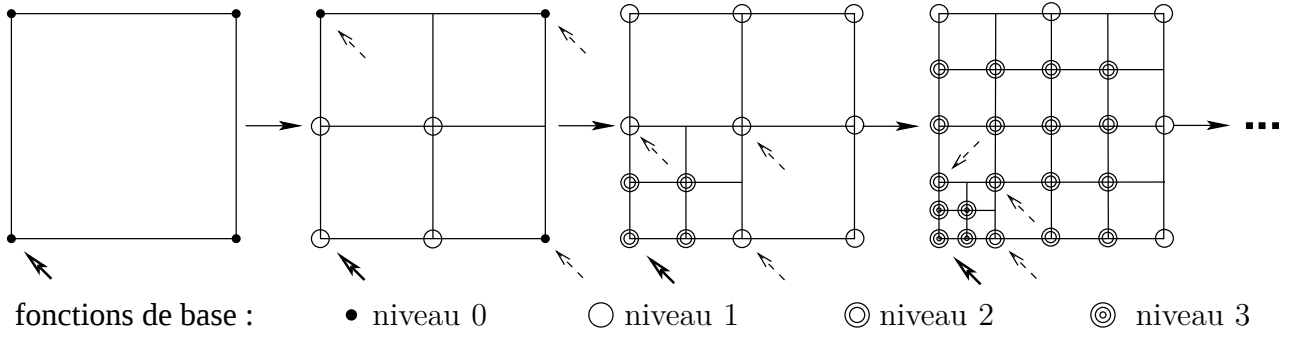


FIG. II.4 – Application de la règle “au-plus-un-niveau-de-différence” (à comparer avec la figure II.3)

Remarque II.16

Il est possible de choisir d'autres définitions des descendants et ascendants d'une fonction de base que celle proposée dans la définition II.14. Par exemple, nous pouvons retrouver les algorithmes d'adaptation présentés dans [KGS03] en définissant les ensembles de descendants (resp. ascendants) comme l'ensemble des enfants (resp. parents). Ces algorithmes ne permettent pas d'écarter des séquences de raffinement comme celles présentées sur la figure II.3 mais interdisent néanmoins certaines situations comme celles présentées sur la figure II.5 et donnent également de bons résultats en pratique.

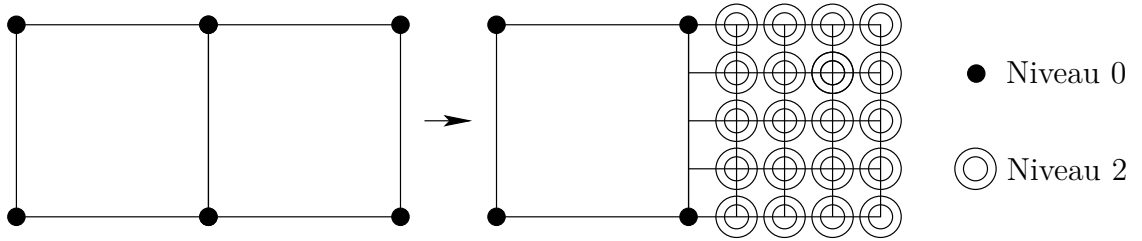


FIG. II.5 – Séquence de raffinement interdite par le critère “au-plus-un-niveau-de-différence”.

II.1.5 Intégration numérique et opérateurs de transfert

Lors de la résolution de problèmes instationnaires, il est souvent nécessaire de calculer des intégrales faisant intervenir plusieurs champs discrets n'appartenant pas aux mêmes espaces d'approximation multiniveaux.

Par exemple, lorsque le terme instationnaire $\partial_t u$ est discrétisé par la méthode d'Euler, nous sommes amenés à calculer l'intégrale suivante :

$$\int_{\Omega} u_h^n \nu_h \, dx, \quad (\text{II.9})$$

où u_h^n représente le champ explicite appartenant à l'espace d'approximation vect $\mathcal{B}^n$ à l'instant t^n et ν_h est une fonction test appartenant à l'espace d'approximation multiniveau vect $\mathcal{B}^{n+1}$ à l'instant t^{n+1} . A cause de l'adaptation de maillage, les deux espaces multiniveaux vect $\mathcal{B}^n$ et vect $\mathcal{B}^{n+1}$ sont différents.

De telles intégrales peuvent être calculées exactement en évitant tout transfert de champ grâce à la notion de maillages multiniveaux introduite ci-dessous. Dans la suite de ce paragraphe, $\mathcal{C}$ désigne l'union de toutes les bases multiniveaux générant les fonctions discrètes intervenant dans la formulation du problème discret. Dans l'exemple du calcul de l'intégrale ci-dessus, nous poserions $\mathcal{C} = \mathcal{B}^n \cup \mathcal{B}^{n+1}$. L'ensemble $\mathcal{C}$ permet de définir les degrés de liberté actifs : ce sont les degrés de liberté associés à chacune des fonctions de base intervenant dans la formulation du problème discret.

Définition II.17 (Degrés de liberté actifs)

Nous disons que $k \in \{1, \dots, N_{ddl}^{[j]}\}$ est un degré de liberté actif de niveau j , si et seulement si $\varphi_k^{[j]} \in \mathcal{C}$. Nous notons $\mathcal{A}_{ddl}^{[j]}$ l'ensemble des degrés de liberté actifs de niveau j .

Les deux définitions suivantes permettent d'introduire la notion de cellules actives qui constituent les domaines d'intégration élémentaires sur lesquels seront appliquées les règles de quadrature.

Définition II.18

Soit $j \in \llbracket 0, J \rrbracket$. Nous disons qu'une cellule $K^{[j]}$ de niveau j satisfait la propriété $(A_{K^{[j]}})$ si et seulement si

$$\forall j' \in \llbracket j+1, J \rrbracket, \forall k \in \mathcal{A}_{ddl}^{[j']}, \widehat{K_e^{[j]}} \cap \widehat{\text{supp}[\varphi_k^{[j']}] = \emptyset. \quad (A_{K^{[j]}})$$

Définition II.19 (Cellules actives)

Une cellule $K^{[j]}$ de niveau j de $\mathcal{T}_j$ est appelée cellule active si et seulement si :

- la cellule $K^{[j]}$ satisfait la propriété $(A_{K^{[j]}})$, et
- sa cellule parent $P(K^{[j]})$ (de niveau $j-1$) ne satisfait pas la propriété $A_{P(K^{[j]})}$, lorsque $j > 0$.

Remarque II.20

Cette définition garantit que toute cellule active est entièrement incluse (au sens large) dans une unique cellule du support de chacune des fonctions de base intervenant dans la formulation du problème discret (et ce malgré les différences potentielles de niveau entre ces différentes cellules). Ainsi, dans le cas d'éléments finis polynomiaux (par exemple $\mathbb{P}_k, \mathbb{Q}_k$) à condition que l'ordre des règles de quadrature soit suffisamment élevé, toutes les intégrales peuvent être calculées exactement.

Nous appelons alors maillage multiniveau, l'ensemble des cellules actives.

Proposition II.21 (Maillage multiniveau)

Pour $j \in \llbracket 0, J \rrbracket$, notons $\widetilde{\mathcal{T}}_j$ l'ensemble des cellules actives de $\mathcal{T}_j$. L'ensemble $\mathcal{T} = \bigcup_{j=0}^J \widetilde{\mathcal{T}}_j$ est un maillage de Ω appelé maillage multiniveau.

Démonstration : Soit $K_e^{[j]}$ et $K_{e'}^{[j']}$ deux cellules actives distinctes. Montrons que $\widehat{K_e^{[j]}} \cap \widehat{K_{e'}^{[j']}} = \emptyset$.

- Cas 1 : $j = j'$. Nous avons alors $e \neq e'$. Dans ce cas, $K_e^{[j]}$ et $K_{e'}^{[j']}$ sont deux cellules distinctes du maillage

$\mathcal{T}_j$. Alors, $\widehat{K_e^{[j]}} \cap \widehat{K_{e'}^{[j']}} = \emptyset$.

- Cas 2 : $j > j'$. En raisonnant par l'absurde, supposons que $\widehat{K_e^{[j]}} \cap \widehat{K_{e'}^{[j']}} \neq \emptyset$. Puisque $j > j'$, nous avons $K_e^{[j]} \subset K_{e'}^{[j']}$. Mais, la cellule $K_e^{[j]}$ est active et $\neg(A_{P(K_e^{[j]})})$ montre l'existence de $j_0 \geq j$ et $k_0 \in \mathcal{A}_{ddl}^{[j_0]}$

tels que $\widehat{P(K_e^{[j]})} \cap \widehat{\text{supp}[\varphi_{k_0}^{[j_0]}]} \neq \emptyset$. De plus, nous avons $j_0 > j'$ et la propriété $(A_{K_{e'}^{[j']}})$ montre que

$\widehat{K_{e'}^{[j']}} \cap \widehat{\text{supp}[\varphi_{k_0}^{[j_0]}]} = \emptyset$. Cependant, nous avons :

$$(j > j' \text{ et } K_e^{[j]} \subset K_{e'}^{[j']}) \implies \widehat{P(K_e^{[j]})} \subset \widehat{K_{e'}^{[j']}}.$$

Ainsi, $\emptyset \neq \left(\widehat{P(K_e^{[j]})} \cap \widehat{\text{supp}[\varphi_{k_0}^{[j_0]}]} \right) \subset \left(\widehat{K_{e'}^{[j']}} \cap \widehat{\text{supp}[\varphi_{k_0}^{[j_0]}]} \right) = \emptyset$. C'est une contradiction.

En conclusion, les intérieurs de deux cellules distinctes de $\mathcal{T}$ sont disjoints.

Soit maintenant $x \in \overline{\Omega}$. Puisque $\mathcal{T}_J$ est un maillage Ω , il existe une cellule $K_{e_J}^{[J]}$ de niveau J qui contient x . Alors, pour tout $j \in \llbracket 0, J-1 \rrbracket$, nous définissons $K_{e_j}^{[j]} = P(K_{e_{j+1}}^{[j+1]})$. Ainsi, pour tout $j \in \llbracket 0, J \rrbracket$, x appartient à la cellule $K_{e_j}^{[j]}$. Considérons l'ensemble $E = \left\{ j \in \llbracket 0, J \rrbracket, \forall \ell \geq j, K_{e_\ell}^{[\ell]} \text{ satisfait } (A_{K_{e_\ell}^{[\ell]}}) \right\}$. Nous avons $J \in E$, donc $E \neq \emptyset$. Soit $j_m = \min_{j \in E} j$, alors, par définition, $K_{e_{j_m}}^{[j_m]}$ est active et contient x . ■

Puisque $\mathcal{T}$ est une partition de Ω , l'intégrale (II.9) peut être décomposée en une somme sur toutes les cellules du maillage $\mathcal{T}$ sur lesquelles une règle de quadrature est appliquée menant à un calcul exact (cf remarque II.20).

Remarque II.22

Les maillages multiniveaux ne sont pas géométriquement conformes, mais les espaces d'approximation multiniveaux sont bien par construction $H^1(\Omega)$ -conforme puisque $\text{vect } \mathcal{B} \subset X_J \subset H^1(\Omega)$. La non-conformité des maillages multiniveaux n'est pas un problème puisqu'ils sont seulement utilisés comme domaines d'intégration élémentaires et pour la sauvegarde (donc la visualisation) des solutions discrètes.

II.2 Préconditionneurs multiniveaux

Dans cette section, nous définissons un algorithme (cf section II.2.4) permettant de reconstruire, à partir d'un espace d'approximation éléments finis multiniveaux, une suite d'espaces emboîtés auxiliaires. Ceci autorise alors à entrer dans le cadre abstrait multigrilles développé dans [BZ00] ; les opérateurs de transfert entre les grilles étant déduits des relations parents-enfants. Avant de décrire cet algorithme de "coarsening", nous rappelons le principe de fonctionnement des solveurs multiniveaux.

II.2.1 Méthodes des corrections successives

Supposons que nous ayons à résoudre le système linéaire $Au = b$. Nous considérons une classe de méthodes itératives (cf [Xu97]) toutes basées sur l'observation suivante : étant donnée une approximation u_{old} de la solution exacte u , le résidu $r = b - Au_{\text{old}}$ et l'erreur $e = u - u_{\text{old}}$ par rapport à cette approximation sont liés par la relation $Ae = r$.

Ainsi, il suffit de résoudre le système linéaire $Ae = r$ pour obtenir l'erreur e et par suite la solution exacte $u = u_{\text{old}} + e$. Bien sûr, la résolution du système $Ae = r$ est aussi difficile que celle du système de départ $Au = b$. Il est cependant possible de se contenter ici d'une résolution approchée en espérant que l'erreur sera corrigée correctement par une itération du processus. La structure des algorithmes présentés ci-après est la suivante : étant donnée une approximation u_{old} de la solution exacte u , une nouvelle approximation u_{new} est obtenue en trois étapes :

- (i) Calcul du résidu : $r = b - Au_{\text{old}}$.
- (ii) Calcul d'une solution approchée $\tilde{e}$ du système $Ae = r$.
- (iii) Correction par l'erreur approchée : $u_{\text{new}} = u_{\text{old}} + \tilde{e}$.

De manière plus compacte, l'algorithme précédent s'écrit :

$$u_{\text{new}} = u_{\text{old}} + B^{-1}(b - Au_{\text{old}}),$$

où B est un opérateur linéaire approchant A (celui qui permet le calcul de $\tilde{e} : B\tilde{e} = r$).

Remarque II.23

Nous donnons quelques exemples de choix possibles de la matrice B qui permettent de retrouver certaines méthodes itératives standard :

- Méthode de Richardson : $B = \omega \text{Id}$ où Id est la matrice identité et $0 < \omega < \frac{2}{\rho(A)}$, ($\rho(A)$ étant le rayon spectral de A).
- Méthode de Jacobi : $B = D$ où D est la partie diagonale de A .
- Méthode de Gauss-Seidel : $B = L$ où L est la partie triangulaire inférieure de A .

II.2.2 Méthodes des corrections successives liées à des sous-espaces

Revenons maintenant au contexte des éléments finis. Nous supposons que nous avons à résoudre le problème variationnel suivant : Trouver $u \in \mathcal{V}_h^F$ tel que

$$\forall v \in \mathcal{V}_h^F, \quad a(u, v) = l(v), \quad (\text{II.10})$$

où $a : \mathcal{V} \times \mathcal{V} \rightarrow \mathbb{R}$ est une forme bilinéaire continue et coercive sur l'espace de Hilbert $\mathcal{V}$, $b : \mathcal{V} \rightarrow \mathbb{R}$ est une forme linéaire continue et $\mathcal{V}_h^F$ est un espace d'approximation éléments finis de $\mathcal{V}$. Nous notons $\mathcal{B}^F$ une base de $\mathcal{V}_h^F$, A^F et b^F la matrice de rigidité et le second membre associé au problème (II.10), c'est-à-dire

$$A^F = [a(\varphi, \psi)]_{\varphi, \psi \in \mathcal{B}^F} \quad \text{et} \quad b^F = [l(\varphi)]_{\varphi \in \mathcal{B}^F}.$$

Nous avons à résoudre le système linéaire $A^F u^F = b^F$. L'idée est alors d'introduire un sous-espace $\mathcal{V}_h^C$ de $\mathcal{V}_h^F$ et de l'utiliser pour effectuer le calcul de l'erreur $\tilde{e}$ dans l'algorithme présenté dans la section précédente. Nous supposons qu'une base $\mathcal{B}^C$ du sous-espace $\mathcal{V}_h^C$ est donnée et nous notons $A^C = [a(\varphi, \psi)]_{\varphi, \psi \in \mathcal{B}^C}$ la matrice de rigidité associée. En outre, nous supposons que des opérateurs de transfert entre les espaces $\mathcal{V}_h^C$ et $\mathcal{V}_h^F$ sont connus par l'intermédiaire de leur représentations matricielles I_C^F et I_F^C (dans les bases $\mathcal{B}^C$ et $\mathcal{B}^F$). Classiquement, nous choisissons pour I_C^F la représentation matricielle de l'injection canonique de $\mathcal{V}_h^C$ dans $\mathcal{V}_h^F$ et pour I_F^C sa transposée. Ces notations sont résumées dans la table II.1.

Espaces	$\mathcal{V}_h^C \subset \mathcal{V}_h^F$	
Bases	$\mathcal{V}_h^C = \text{vect } \mathcal{B}^C$	$\mathcal{V}_h^F = \text{vect } \mathcal{B}^F$
Matrices	$A^C = [a(\varphi, \psi)]_{\varphi, \psi \in \mathcal{B}^C}$	$A^F = [a(\varphi, \psi)]_{\varphi, \psi \in \mathcal{B}^F}$
Seconds membres	$b^C = [l(\varphi)]_{\varphi \in \mathcal{B}^C}$	$b^F = [l(\varphi)]_{\varphi \in \mathcal{B}^F}$
Matrices de transfert	$\mathcal{V}_h^C \xrightarrow{I_C^F} \mathcal{V}_h^F$	
	$\mathcal{V}_h^F \xrightarrow{I_F^C} \mathcal{V}_h^C$	

TAB. II.1 – Résumé des notations

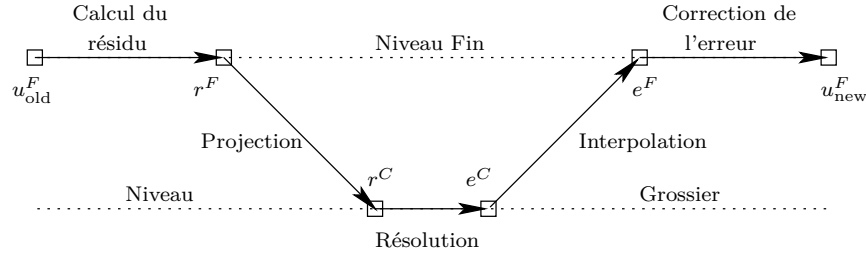


FIG. II.6 – Correction liée à un unique sous-espace

L'algorithme de correction (*cf* [Xu97, p.29-30]) s'écrit de la manière suivante :

$$\begin{array}{lll}
 u^F & \leftarrow & u_{\text{old}}^F \quad \text{donné.} \\
 r^F & \leftarrow & b^F - A^F u^F \quad \text{Calcul du résidu sur le "niveau fin".} \\
 r^C & \leftarrow & I_F^C r^F \quad \text{Projection du résidu sur le "niveau grossier".} \\
 e^C & \leftarrow & (A^C)^{-1} r^C \quad \text{Calcul (approché ou pas) de l'erreur sur le "niveau grossier".} \\
 e^F & \leftarrow & I_C^F e^C \quad \text{Interpolation de l'erreur sur le "niveau fin".} \\
 u_{\text{new}}^F & = & u^F + e^F \quad \text{Correction sur le "niveau fin".}
 \end{array}$$

Cet algorithme est illustré par la figure II.6 et peut encore s'écrire de manière plus compacte sous la forme :

$$u_{\text{new}}^F = u_{\text{old}}^F + I_C^F (A^C)^{-1} I_F^C (b^F - A^F u_{\text{old}}^F).$$

Certains algorithmes remplacent, pour le calcul de l'erreur, la matrice A^C par une matrice B^C l'approchant (*cf* remarque II.23) :

$$u_{\text{new}}^F = u_{\text{old}}^F + I_C^F (B^C)^{-1} I_F^C (b^F - A^F u_{\text{old}}^F).$$

L'algorithme présenté ci-dessus fait intervenir les corrections associées à un unique niveau. Lorsque nous disposons de plusieurs sous-espaces, nous pouvons effectuer des corrections par rapport à chacun de ces sous-espaces. Il existe plusieurs façons de combiner ces différentes corrections. Nous présentons quelques stratégies parmi les plus classiques : stratégies multiplicative et additive.

Nous adoptons maintenant les notations suivantes. Nous notons V_J l'espace d'approximation sur lequel est posé le système à résoudre : Trouver $u \in V_J$ tel que

$$\forall v \in V_J, a(u, v) = l(v).$$

Nous notons $\mathcal{B}_J$ une base de V_J , A_J et b_J la matrice de rigidité et le second membre associés. Nous supposons que nous disposons de J sous-espaces de V_J (pas nécessairement emboîtés) que nous notons V_j , $j \in \llbracket 0, J-1 \rrbracket$. Nous notons $\mathcal{B}_j$ une base de V_j , A_j et b_j la matrice de rigidité et le second membre associés, et enfin B_j désigne une matrice approchant la matrice A_j (cf remarque II.23). Nous supposons que des opérateurs de transfert entre les sous-espaces V_j , $j \in \llbracket 0, J-1 \rrbracket$ et l'espace V_J sont connus par l'intermédiaire de leur représentation matricielle I_j et ${}^t I_j$ (dans les bases $\mathcal{B}_j$ et $\mathcal{B}_J$). En pratique, nous utilisons pour I_j la représentation matricielle de l'injection canonique de V_j dans V_J . Les notations sont résumées dans la table II.2.

Espaces	$V_0, \dots, V_j, \dots, V_{J-1} \subset V_J$
Bases	$V_j = \text{vect } \mathcal{B}_j, j \in \llbracket 0, J-1 \rrbracket$
Matrices	$A_j = [a(\varphi, \psi)]_{\varphi, \psi \in \mathcal{B}_j}, j \in \llbracket 0, J \rrbracket$
Matrices approchées	$B_j, j \in \llbracket 0, J-1 \rrbracket$
Seconds membres	$b_j = [l(\varphi)]_{\varphi \in \mathcal{B}_j}, j \in \llbracket 0, J \rrbracket$
Matrices de transfert	$V_j \xrightarrow{I_j} V_J, j \in \llbracket 0, J-1 \rrbracket$
	$V_J \xrightarrow{{}^t I_j} V_j, j \in \llbracket 0, J-1 \rrbracket$

TAB. II.2 – Résumé des notations

Méthode multiplicative

La méthode multiplicative [Yse93, p.293-294] consiste à effectuer les corrections sur tous les niveaux successivement. Ainsi, le calcul de chaque correction prend en compte les corrections précédentes :

```

v = uold
Pour j = 0 ... J - 1
    v ← v + IjBj-1tIj(bJ - AJv)
Fin Pour
unew = v.

```

Cet algorithme est illustré par la figure II.7 dans le cas où $J = 2$.

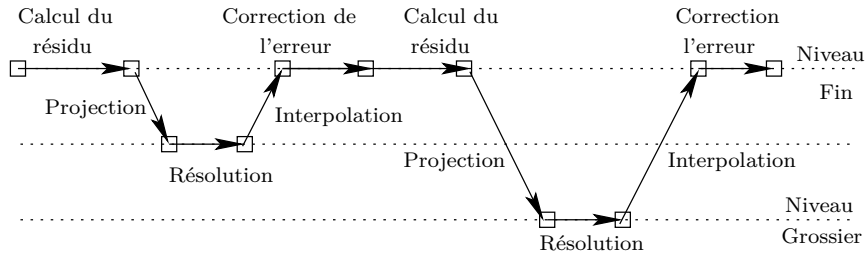


FIG. II.7 – Corrections liées à un sous-espace. Stratégie multiplicative

Méthode additive

La méthode additive [Yse93, p.297] consiste à calculer les corrections sur les différents niveaux en parallèle à partir de l'itérée courante. L'ensemble de ces corrections est ensuite apporté à l'itérée courante. Chaque

correction est donc calculée indépendamment des autres.

```

v = uold
Pour j = 0 ... J - 1
    v ← v + IjBj-1tj(bJ - AJuold)
Fin Pour
unew = v

```

Cet algorithme est illustré par la figure II.8 dans le cas où $J = 3$ et peut se réécrire de la manière suivante :

$$u_{\text{new}} = u_{\text{old}} + \left(\sum_{j=1}^J I_j B_j^{-1} t_j \right) (b - A u_{\text{old}}).$$

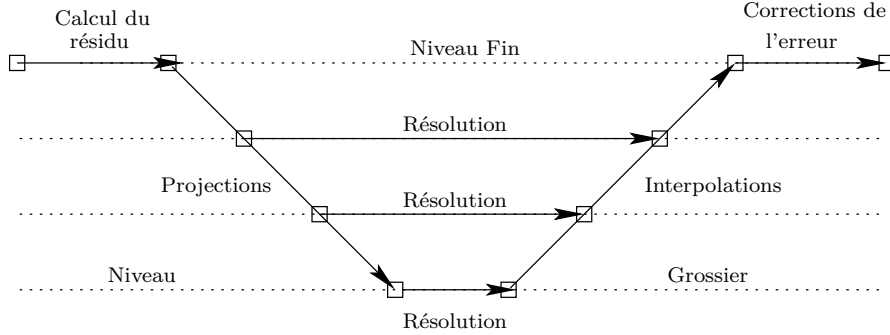


FIG. II.8 – Corrections liées à un sous-espace. Stratégie additive

Remarque II.24

Donnons un exemple d'application des algorithmes ci-dessus. Nous notons J la dimension de l'espace d'approximation V_J (ceci n'est pas gênant, dans les notations ci-dessus J désigne le nombre de sous-espaces utilisés pour effectuer les corrections et rien d'autre) ainsi nous pouvons noter $\mathcal{B}_J = \{\varphi_i^{[J]}\}_{i \in \llbracket 0, J-1 \rrbracket}$ et choisir, pour $j \in \{0 \dots J-1\}$, $V_j = \text{vect} \{\varphi_j^{[J]}\}$. L'espace V_j est donc de dimension 1 ; dans ce cas, I_j , la représentation matricielle de l'injection canonique de V_j dans V_J et A_j sont très simples : I_j est le i^e vecteur colonne de la base canonique de $\mathbb{R}^{N_J}$ et $A_j = [a(\Phi_j^j, \Phi_j^j)] = [(A_J)_{jj}]$. La correction associée à l'espace V_j est donc la suivante :

$$u_{\text{new}} = \begin{bmatrix} u_{\text{old},1} \\ \vdots \\ u_{\text{old},j-1} \\ u_{\text{old},j} + \frac{1}{(A_J)_{jj}} (b_J - A_J u_{\text{old}})_j \\ u_{\text{old},j+1} \\ \vdots \\ u_{\text{old},N_J} \end{bmatrix}$$

Selon la stratégie choisie, multiplicative ou additive, nous obtenons, pour cet exemple, respectivement la méthode de Gauss-Seidel ou celle de Jacobi.

II.2.3 Méthodes multigrilles

Les méthodes itératives classiques (méthode de Richardson, méthode de Jacobi relaxée, méthode de Gauss-Seidel, cf [SVdV00]) sont peu performantes mais possèdent la qualité d'être de bons lisseurs. Ceci signifie qu'en quelques itérations elles permettent d'éliminer les hautes fréquences de l'erreur, la convergence des basses fréquences étant très lente.

Le principe des méthodes multigrilles (cf [Hac85, Wes92, TOS01]) est alors de conjuguer le pouvoir lissant de ces méthodes peu coûteuses à la correction liée à un sous-espace associé à une grille grossière (mais néanmoins suffisante pour corriger les basses fréquences de l'erreur).

Méthode deux grilles

Nous reprenons les notations de la section précédente (*cf* table II.1) : les objets sur le niveau fin sont désignés à l'aide de la lettre F et ceux sur le niveau grossier à l'aide la lettre C . En outre, si ν est un entier, nous notons $S^\nu(u, A, b)$, le résultat obtenu après ν itérations (à partir de l'itéré u) d'une méthode lissante pour le système $Au = b$ et nous choisissons deux entiers positifs ν_1 et ν_2 .

L'algorithme deux grilles s'exprime alors de la manière suivante :

u^F	$\leftarrow u_{\text{old}}^F$	donné.
u^F	$\leftarrow S^{\nu_1}(u^F, A^F, b^F)$	Pré-lissage.
r^F	$\leftarrow b^F - A^F u^F$	Calcul du résidu sur le niveau fin $\mathcal{V}_h^F$.
r^C	$\leftarrow {}^t I_F^C r^F$	Projection du résidu sur le niveau grossier $\mathcal{V}_h^C$.
e^C	$\leftarrow (B^C)^{-1} r^C$	Calcul (approché ou pas) de l'erreur sur le niveau grossier $\mathcal{V}_h^C$.
e^F	$\leftarrow I_C^F e^C$	Interpolation de l'erreur sur le niveau fin $\mathcal{V}_h^F$.
u^F	$\leftarrow u^F + e^F$	Correction sur le niveau fin $\mathcal{V}_h^F$.
u_{new}^F	$= S^{\nu_2}(u^F, A^F, b^F)$	Post-lissage.

Cycles multigrilles

Nous supposons que nous disposons de plusieurs grilles emboîtées. Encore une fois nous reprenons les notations de la section précédente (*cf* table II.2) mais supposons en outre que $V_0 \subset V_1 \subset \dots \subset V_J$.

Le principe de l'algorithme est d'utiliser récursivement la méthode deux grilles. Pour résoudre le système linéaire $A_J u_J = b_J$ nous commençons par utiliser la méthode deux-grilles (*cf* paragraphe précédent) entre V_J et V_{J-1} . Il faut alors définir l'opérateur approchant A_{J-1} sur la grille V_{J-1} . Nous utilisons alors l'algorithme deux-grilles entre les niveaux V_{J-1} et V_{J-2} . Nous répétons ce processus récursivement jusqu'au niveau V_0 où nous choisissons $B_0 = A_0$.

Nous définissons donc récursivement l'algorithme multigrille $MG(j, u_{\text{old}}, f_j)$ effectuant une correction de l'erreur commise avec l'itérée u_{old} en vue de résoudre le système $A_j u_j = f_j$ de la manière suivante :

Si $j = 1$,	$MG(1, u_{\text{old}}, f_1) = A_1^{-1} f_1$	Résolution exacte sur le niveau le plus grossier.
Sinon		
u_j	$\leftarrow u_{\text{old}}$	
u_j	$\leftarrow S^{\nu_1}(u_j, A_j, f_j)$	Pré-lissage sur le niveau fin.
r_j	$\leftarrow f_j - A_j u_j$	Calcul du résidu sur le niveau fin.
r_{j-1}	$\leftarrow {}^t I_{j-1}^j r_j$	Projection du résidu du niveau fin sur le niveau grossier.
e_{j-1}	$\leftarrow MG(j-1, 0, r_{j-1})$	Calcul approché de l'erreur sur le niveau grossier V_{j-1}
e_j	$\leftarrow I_{j-1}^j e_{j-1}$	Interpolation de l'erreur du niveau grossier sur le niveau fin.
u_j	$\leftarrow u_j + e_j$	Correction sur le niveau fin.
u_j	$\leftarrow S^{\nu_2}(u_j, A_j, f_j)$	Post-lissage sur le niveau fin.
u_{new}	$\leftarrow u_j$	
$MG(j, u_{\text{old}}, f_j)$	$= u_{\text{new}}$	

La résolution du système $A_J u_J = b_J$ s'effectue alors de la manière suivante :

$u_J^0 = 0$	Itérée initiale donnée.
Pour $n \geq 0$, $u_J^{n+1} = MG(J, u_J^n, b_J)$	Itérations successives.

Ce cycle multigrille se nomme V-cycle, de nombreuses autres stratégies existent mais nous ne les présenterons pas dans ce manuscrit. Le fonctionnement de V-cycle est illustré dans le cas de trois niveaux par la figure II.9.

En pratique nous utilisons cette méthode comme le préconditionneur d'une méthode de Krylov (gradient conjugué ou GMRES, *cf* [SVdV00]) et non comme un solveur itératif (*cf* section II.2.5).

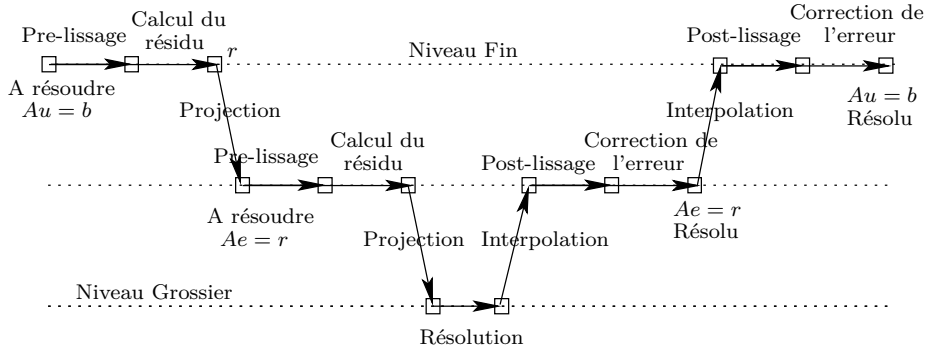


FIG. II.9 – Fonctionnement du V-cycle

Nous souhaitons maintenant utiliser ces algorithmes multiniveaux classiques dans le cas du raffinement local. Un des avantages cruciaux de la méthode de raffinement que nous avons présentée dans le début de ce chapitre est qu'elle va permettre de reconstruire facilement une sous-suite de "grilles" auxiliaires utilisées pour le solveur multigrille à travers un algorithme de coarsening des bases multiniveaux présenté dans la section suivante.

II.2.4 Algorithme de coarsening d'un espace d'approximation multiniveau

A partir d'une base multiniveau "fine" $\mathcal{B}^F$, l'algorithme suivant est utilisé pour construire une base multiniveau plus "grossière" $\mathcal{B}^C$.

Algorithme II.25 (Coarsening)

Soit $\mathcal{B}^F$ une base multiniveau. Soit $j_M = \max \{j \in \llbracket 0, J \rrbracket; \mathcal{B}^F \cap \mathcal{B}_j \neq \emptyset\}$ le plus haut niveau de raffinement dans $\mathcal{B}^F$. Une base multiniveau plus "grossière" $\mathcal{B}^C$, notée $\mathcal{B}^C = \text{coarsen}(\mathcal{B}^F)$, est obtenue à partir de $\mathcal{B}^F$ par déraffinement (cf algorithme II.4) de l'ensemble des fonctions de base $\mathcal{B}^F$ -raffinées de niveau $j_M - 1$.

La proposition suivante donne une formulation équivalente de l'algorithme ci-dessus.

Proposition II.26

Supposons que $\mathcal{B}^F$ soit une base multiniveau satisfaisant la propriété suivante :

Toute fonction de base de niveau j , $j \geq 1$, qui soit appartient à $\mathcal{B}$ soit est $\mathcal{B}$ -raffinée, a au moins un parent $\mathcal{B}$ -raffiné.

($\mathcal{P}_{\text{HI}}$)

Soit $j_M = \max \{j \in \llbracket 0, J \rrbracket; \mathcal{B}^F \cap \mathcal{B}_j \neq \emptyset\}$. La base multiniveau $\mathcal{B}^C = \text{coarsen}(\mathcal{B}^F)$ définie dans l'algorithme II.25 peut être obtenue par l'algorithme équivalent suivant :

- supprimer toutes les fonctions de base de niveau j_M de $\mathcal{B}^F$,
- ajouter toutes les fonctions de base $\mathcal{B}^F$ -raffinées de niveau $j_M - 1$.

Démonstration : Remarquons tout d'abord que chaque étape dans l'algorithme II.25 est le déraffinement d'une fonction de base de niveau $j_M - 1$, cela implique que les fonctions de base ajoutées sont de niveau $j_M - 1$ et que celles qui sont retirées sont de niveau j_M . Ainsi, les ensembles de fonctions de base ajoutées et supprimées sont disjoints. Dans l'algorithme II.25, une fonction de base qui est supprimée (resp. ajoutée) par un premier déraffinement ne peut être ajoutée (resp. supprimée) par une autre déraffinement ultérieur. De plus, la définition de j_M implique qu'il n'existe pas de fonction de base $\mathcal{B}^F$ -rafinée. Ainsi, puisqu'elles n'ont pas d'enfants $\mathcal{B}^F$ -raffinés, toutes les fonctions de base de niveau $j_M - 1$ peuvent être déraffinées. Le déraffinement d'une fonction de base implique que cette fonction est ajoutée et donc par suite que l'ensemble des fonctions de base ajoutées dans l'algorithme II.25 est exactement l'ensemble des fonctions de base $\mathcal{B}^F$ -raffinées. Il reste à montrer que toutes les fonctions de base de niveau j_M de $\mathcal{B}^F$ sont supprimées. En raisonnant par l'absurde, supposons qu'une fonction de base de niveau j_M appartienne à $\text{coarsen}(\mathcal{B}^F)$. Puisque la procédure de déraffinement préserve la propriété ($\mathcal{P}_{\text{HI}}$) (cf proposition II.27), cette fonction de base a au moins un parent $\text{coarsen}(\mathcal{B}^F)$ -raffiné. D'après le lemme II.9 propriété 2) (ii), ce parent est une fonction de base $\mathcal{B}^F$ -rafinée (de niveau $j_M - 1$). C'est une contradiction et le résultat est prouvé. ■

La propriété ($\mathcal{P}_{\text{HI}}$) garantit le fait que toutes les fonctions de base de niveaux j_M sont supprimées de $\mathcal{B}^F$.

De plus, cette propriété n'est pas très restrictive puisqu'elle est préservée par les procédures de raffinement et de déraffinement.

Proposition II.27

Les procédures de raffinement et de déraffinement décrites dans l'algorithme II.4 préservent la propriété $(\mathcal{P}_{\text{HI}})$.

Démonstration : Supposons que $\mathcal{B}^*$ soit une base multiniveau satisfaisant la propriété $(\mathcal{P}_{\text{HI}})$ et soit $j \in \llbracket 1, J \rrbracket$.

- Considérons tout d'abord une fonction de base $\varphi_k^{[j]}$ appartenant à $\mathcal{B}$.
- Raffinement : Supposons que la base multiniveau $\mathcal{B}$ est obtenue à partir de $\mathcal{B}^*$ par le raffinement d'une fonction de base $\varphi_{k_0}^{[j_0]} \in \mathcal{B}^*$.
 - Cas 1 : $\varphi_k^{[j]} \in \mathcal{B}^*$. Puisque $\mathcal{B}^*$ satisfait la propriété $(\mathcal{P}_{\text{HI}})$, $\varphi_k^{[j]}$ a au moins un parent $\mathcal{B}^*$ -raffiné. D'après le lemme II.9 propriété 1), ce parent est $\mathcal{B}$ -raffiné.
 - Cas 2 : $\varphi_k^{[j]}$ est un enfant de $\varphi_{k_0}^{[j_0]}$ et $\varphi_{k_0}^{[j_0]}$ est $\mathcal{B}$ -raffinée.
- Déraffinement : Supposons que $\mathcal{B}$ est obtenue à partir de $\mathcal{B}^*$ par le déraffinement d'une fonction de base raffinée une seule fois $\varphi_{k_0}^{[j_0]}$ de $\mathcal{B}^*$.
 - Cas 1 : $\varphi_k^{[j]} \in \mathcal{B}^* \setminus \{\text{enfant de } \varphi_k^{[j]} \text{ sans autre parent } \mathcal{B}^*\text{-raffiné}\}$.
Puisque $\mathcal{B}^*$ satisfait la propriété $(\mathcal{P}_{\text{HI}})$, $\varphi_k^{[j]}$ a au moins un parent $\mathcal{B}^*$ -raffiné. D'après le lemme II.9 propriété 2), soit ce parent est $\mathcal{B}$ -raffiné soit c'est $\varphi_{k_0}^{[j_0]}$.
Dans le premier cas, la preuve est finie, dans le second cas $\varphi_k^{[j]}$ est un enfant de $\varphi_{k_0}^{[j_0]}$ qui appartient à $\mathcal{B}^* \setminus \{\text{enfant de } \varphi_k^{[j]} \text{ sans autre parent } \mathcal{B}^*\text{-raffiné}\}$. Ainsi, $\varphi_k^{[j]}$ a un autre parent $\mathcal{B}^*$ -raffiné. Ce parent est par suite $\mathcal{B}$ -raffiné puisque différent de $\varphi_{k_0}^{[j_0]}$.
 - Cas 2 : $\varphi_k^{[j]} = \varphi_{k_0}^{[j_0]}$, $\varphi_{k_0}^{[j_0]}$ est $\mathcal{B}^*$ -raffinée. Puisque $\mathcal{B}^*$ satisfait la propriété $(\mathcal{P}_{\text{HI}})$, $\varphi_{k_0}^{[j_0]}$ a au moins un parent $\mathcal{B}^*$ raffiné. D'après le lemme II.9 propriété 2), ce parent est $\mathcal{B}$ -raffiné puisque différent de $\varphi_{k_0}^{[j_0]}$.
- Considérons maintenant une fonction de base $\varphi_k^{[j]}$ qui est $\mathcal{B}$ -raffinée.
 - Raffinement : Supposons que la base multiniveau $\mathcal{B}$ est obtenue à partir de $\mathcal{B}^*$ par raffinement d'une fonction de base $\varphi_{k_0}^{[j_0]} \in \mathcal{B}^*$. D'après le lemme II.9 propriété 1) (ii), $\varphi_k^{[j]}$ est soit $\mathcal{B}^*$ -raffinée soit égale à $\varphi_{k_0}^{[j_0]} \in \mathcal{B}^*$. Puisque $\mathcal{B}^*$ satisfait la propriété $(\mathcal{P}_{\text{HI}})$, dans les deux cas, $\varphi_k^{[j]}$ a au moins un parent $\mathcal{B}^*$ -raffiné. D'après le lemme II.9 propriété 1)(i), ce parent est aussi $\mathcal{B}$ -raffiné.
 - Déraffinement : Supposons que $\mathcal{B}$ soit obtenue à partir de $\mathcal{B}^*$ par déraffinement d'une fonction de base $\varphi_{k_0}^{[j_0]}$ de $\mathcal{B}^*$ raffinée une seule fois et sans enfant $\mathcal{B}^*$ -raffiné. D'après le lemme II.9 propriété 2) (ii), $\varphi_k^{[j]}$ est $\mathcal{B}^*$ -raffinée. Puisque $\mathcal{B}^*$ satisfait la propriété $(\mathcal{P}_{\text{HI}})$, $\varphi_k^{[j]}$ a au moins un parent $\mathcal{B}^*$ -raffiné. D'après le lemme II.9 propriété 2) (i), soit ce parent est $\mathcal{B}$ -raffiné soit il est égal à $\varphi_{k_0}^{[j_0]}$. Le dernier cas est impossible car $\varphi_{k_0}^{[j_0]}$ n'a pas d'enfant $\mathcal{B}^*$ -raffiné.

■

La dernière propriété importante de la procédure de coarsening est qu'elle produit des espaces emboîtés.

Proposition II.28

Soient $\mathcal{B}^F$ une base multiniveau satisfaisant la propriété $(\mathcal{P}_{\text{LO}})$ et $\mathcal{B}^C = \text{coarsen}(\mathcal{B}^F)$. Nous avons l'inclusion suivante :

$$\text{vect } \mathcal{B}^C \subset \text{vect } \mathcal{B}^F.$$

Démonstration : Soit $\varphi \in \mathcal{B}^C$. Si $\varphi \notin \mathcal{B}^F$ alors φ a été déraffinée et donc c'est une fonction de base $\mathcal{B}^F$ -raffinée de niveau $j_M - 1$. D'après la proposition II.7, $\varphi \in \text{vect } \mathcal{B}^F$. ■

Remarque II.29

Grâce à l'équation de raffinement, il est facile de construire la représentation matricielle de l'injection naturelle de $\text{vect } \mathcal{B}^C$ dans $\text{vect } \mathcal{B}^F$ (dans les bases multiniveaux associées $\mathcal{B}^C$ et $\mathcal{B}^F$).

En effet, si $\varphi \in \mathcal{B}^C$, nous avons

- soit $\varphi \in \mathcal{B}^F$,

- soit tous les enfants de φ appartiennent à $\mathcal{B}^F$ et donc, l'équation de raffinement donne une expression de φ comme combinaison linéaire des éléments de $\mathcal{B}^F$.

Remarque II.30

Les propriétés $(\mathcal{P}_{LO})$ et $(\mathcal{P}_{HI})$ utilisées dans cette section ne sont pas restrictives puisqu'elles sont préservées par les procédures de raffinement et de déraffinement. Il est très facile de vérifier que ces propriétés sont satisfaites par la base “grossière” $\mathcal{B}_0$ utilisée en pratique comme point de départ de l'algorithme d'adaptation (cf remarque II.11).

II.2.5 Préconditionneurs multiniveaux

A l'aide de l'algorithme présenté dans la section II.2.4, nous définissons récursivement la séquence $\{V_0, \dots, V_J\}$ d'espaces emboîtés déduite de $\mathcal{V}_h$ comme suit :

- nous posons tout d'abord $\mathcal{B}_J = \mathcal{B}$ et $V_J = \text{vect } \mathcal{B}_J = \mathcal{V}_h$,
- ensuite, pour $k = J, \dots, 1$, nous définissons une base multiniveau plus grossière $\mathcal{B}_{k-1}$ à partir de $\mathcal{B}_k$ par :

$$\mathcal{B}_{k-1} = \text{coarsen}(\mathcal{B}_k),$$

et l'espace d'approximation multiniveau correspondant :

$$V_{k-1} = \text{vect } \mathcal{B}_{k-1}.$$

D'après la proposition II.28, nous avons :

$$V_0 \subset V_1 \subset \dots \subset V_J.$$

Notons que la séquence auxiliaire $V_0 \subset \dots \subset V_J$ introduite ici ne reflète en aucun cas la procédure d'adaptation (dynamique) qui a conduit à $\mathcal{V}_h$. Un exemple simple comportant quatre niveaux de raffinement est donné sur la figure II.10.

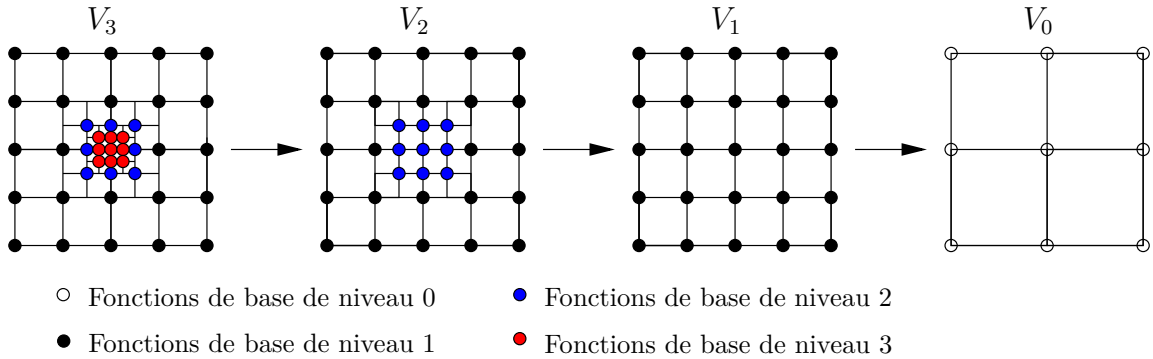


FIG. II.10 – Exemple de coarsening : de V_3 à V_0 .

Nous pouvons maintenant définir les trois composantes des préconditionneurs multigrilles (cf section II.2.3) : opérateurs “intergrilles” (permettant les transferts entre les différents espaces de la suite $\{V_0, \dots, V_J\}$), opérateurs approchés de l’opérateur A_J sur chacun des niveaux V_k et lisseurs sur chacun des niveaux V_k .

D’après la remarque II.29, il est facile de construire la représentation matricielle, notée I_{k-1}^k , de l’injection canonique de V_{k-1} dans V_k dans les bases $\mathcal{B}_{k-1}$ et $\mathcal{B}_k$. Nous définissons alors les opérateurs intergrilles de la façon suivante ; pour tout $k \in \llbracket 0, J \rrbracket$:

$$I_k = I_{J-1}^J I_{J-2}^{J-1} \dots I_k^{k+1}.$$

Nous pouvons définir les opérateurs approchés sur chaque espace V_k , pour tout $k \in \llbracket 0, J \rrbracket$ par :

$$A_k = I_k^t A_J I_k. \quad (\text{II.11})$$

Dans la suite nous utiliserons comme lisseurs, les méthodes de Jacobi et celle de Gauss-Seidel définies pour tout $k \in \llbracket 0, J \rrbracket$ de la manière suivante :

- Jacobi : $S_k = D_k$ où D_k est la partie diagonale de A_k .
- Gauss-Seidel : $S_k = T_k$ où T_k est la partie triangulaire supérieur de A_k .

Dans les deux cas, $S_k \in \mathfrak{M}_{\# \mathcal{B}_k, \# \mathcal{B}_k}(\mathbb{R})$.

Nous utilisons les deux préconditionneurs multiniveaux suivants :

- Version additive (cf [BPX90] et section II.2.2) :

$$P_a = \sum_{k=0}^J I_k S_k I_k^t, \quad (\text{II.12})$$

où S_k , $k \in \llbracket 0, J \rrbracket$, est la méthode de Jacobi.

- V-cycle (cf section II.2.3) : Dans cette partie S_k désigne la méthode de Gauss-Seidel. Nous définissons récursivement, pour tout $k \in \llbracket 0, J \rrbracket$, l'opérateur linéaire $MG_k : \mathbb{R}^{\# \mathcal{B}_k} \rightarrow \mathbb{R}^{\# \mathcal{B}_k}$. Nous posons tout d'abord $MG_0 = A_0^{-1}$ et pour tout $k \in \llbracket 1, J \rrbracket$, nous définissons $MG_k(f_k)$, $f_k \in \mathbb{R}^{\# \mathcal{B}_k}$ par les étapes suivantes :

- | | |
|---|------------------------------------|
| (0) $v_k \leftarrow 0$, | initialisation, |
| (1) $v_k \leftarrow v_k + S_k(f_k - A_k v_k)$, | étapes de préissage, |
| (2) $v_k \leftarrow v_k + I_{k-1} MG_{k-1}(I_{k-1}^t(f_k - A_k v_k))$, | correction de la grille grossière, |
| (3) $v_k \leftarrow v_k + S_k(f_k - A_k v_k)$, | étapes de postlissage, |
| (4) $MG_k(f_k) = v_k$. | |

Le préconditionneur multiplicatif est alors

$$P_m = MG_J. \quad (\text{II.13})$$

II.3 Validation sur un problème modèle stationnaire

Nous validons la méthode de raffinement local et les préconditionneurs multigrilles sur le problème modèle stationnaire suivant.

II.3.1 Problème continu

Soit $\Omega =]0, 1[^d$ ($d = 2$, ou 3). Considérons le problème de Laplace avec des conditions aux bords de type Dirichlet homogène :

$$\begin{cases} -\Delta u = f & \text{dans } \Omega, \\ u = 0 & \text{sur } \partial\Omega. \end{cases} \quad (\text{II.14})$$

Le terme source f est choisi de telle manière que la solution exacte u soit définie par :

$$\forall \mathbf{x} \in \mathbb{R}^d, u(\mathbf{x}) = H_\varepsilon(R - |\mathbf{x} - \mathbf{x}_C|),$$

où $R \in \mathbb{R}$, $\varepsilon \in \mathbb{R}_+^*$, $\mathbf{x}_C \in \mathbb{R}^d$ sont des paramètres réels et $H_\varepsilon : \mathbb{R} \rightarrow \mathbb{R}$ est définie par

$$\forall x \in \mathbb{R}, H_\varepsilon(x) = \begin{cases} 0 & \text{si } x < -\varepsilon, \\ \frac{1}{2} \left[1 + \frac{x}{\varepsilon} + \frac{1}{\pi} \sin\left(\pi \frac{x}{\varepsilon}\right) \right] & \text{si } |x| \leq \varepsilon, \\ 1 & \text{si } x > \varepsilon. \end{cases}$$

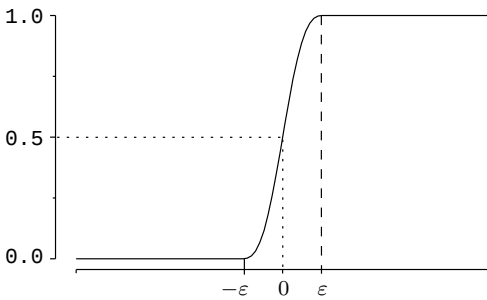


FIG. II.11 – Fonction H_ε .

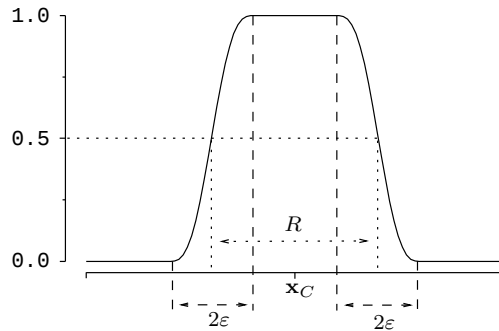


FIG. II.12 – Solution exacte du problème (II.14), $d = 1$.

La fonction H_ε est représentée sur la figure II.11 et l'interprétation des paramètres R , ε , $\mathbf{x}_C$ est expliquée sur la figure II.12 représentant la solution exacte u lorsque $d = 1$.

Pour les simulations numériques données dans la suite de ce chapitre nous avons posé :

$$\Omega =]0, 1[^d, \quad \mathbf{x}_C = (0.5, 0.5), \quad R = 0.3 \quad \text{et} \quad \varepsilon = 0.1.$$

II.3.2 Maillages initiaux et critère de raffinement

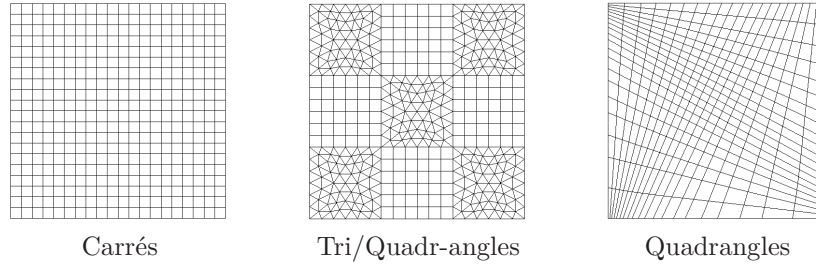


FIG. II.13 – Maillages initiaux.

Pour montrer les possibilités de la méthode de raffinement local, nous avons utilisé plusieurs types de maillages : carrés, triangles, quadrangles généraux, en deux dimensions et cube en trois dimensions. La figure II.13 montre les maillages initiaux, *i.e.* avant les étapes de raffinement, utilisés pour la validation en deux dimensions. En particulier, notons que la méthode permet de combiner plusieurs types d'éléments géométriques (par exemple triangle- $\mathbb{P}_1$ et carré- $\mathbb{Q}_1$) à condition que les éléments de référence et motifs de raffinement correspondants soient compatibles. En trois dimensions, le maillage initial est constitué de cubes réguliers obtenus en divisant chaque arête du domaine en 15 segments.

Dans cette section, la valeur de la solution exacte étant connue, nous utilisons un critère de raffinement géométrique, en choisissant *a priori* le nombre d'étapes de raffinement et la position des noeuds associés aux fonctions de base à raffiner. Nous utilisons le critère suivant :

Critère II.31 (Critère géométrique)

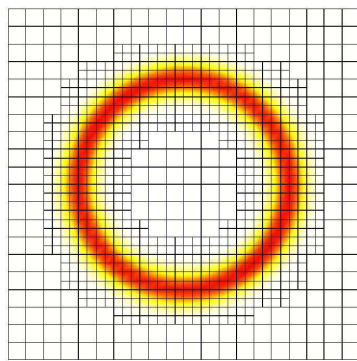
La fonction de base $\varphi_k^{[j]}$ est raffinée ssi

$$R - \varepsilon < \left| \mathbf{a}_k^{[j]} - \mathbf{x}_C \right| < R + \varepsilon.$$

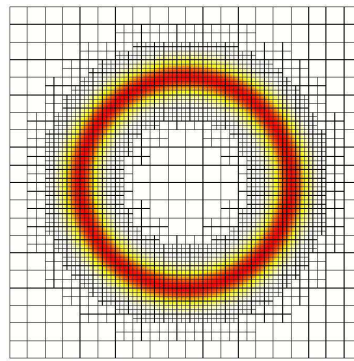
II.3.3 Raffinement local et préconditionneurs multigrilles

Pour chaque type de maillage initial, nous effectuons six calculs en augmentant le nombre d'étapes de raffinement (de un à six).

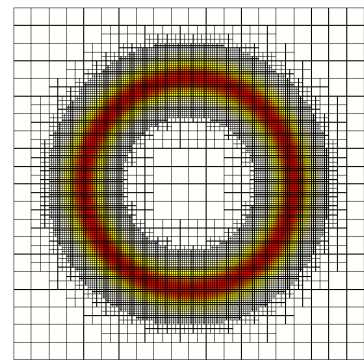
Sur les figures II.14, II.15, II.16 et II.17, nous représentons les maillages et la fonction $u_h(1 - u_h)$ où u_h est la solution approchée calculée lorsque nous effectuons une, deux ou trois étapes de raffinement. Ainsi, la zone colorée sur les figures représente "l'interface" calculée, *i.e.* la zone indiquée par le critère de raffinement II.31. Remarquons que les maillages dans notre méthode sont obtenus comme support des fonctions de base intervenant dans la formulation du problème (*cf* section II.1.5).



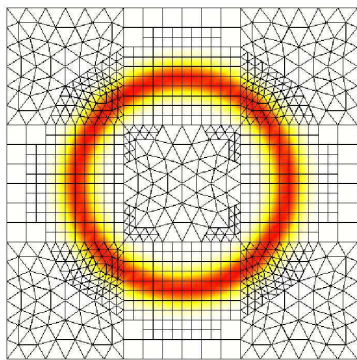
Une étape de raffinement



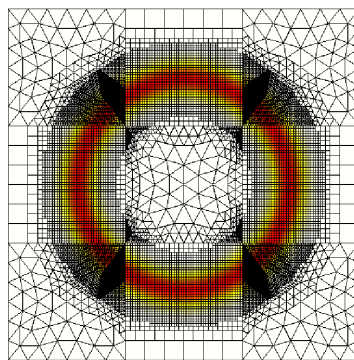
Deux étapes de raffinement



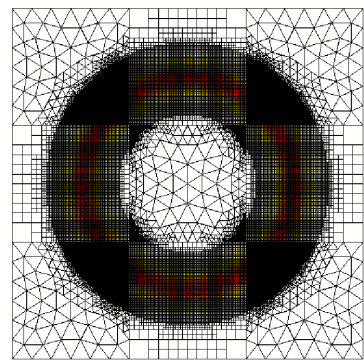
Trois étapes de raffinement

FIG. II.14 – Carrés- Q_1 . Maillages raffinés.

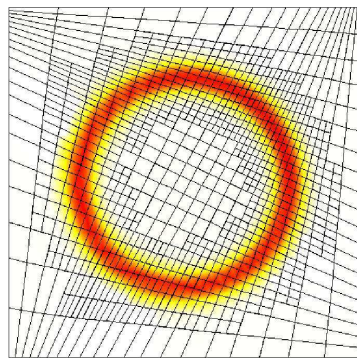
Une étape de raffinement



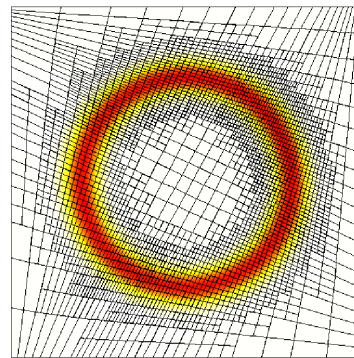
Deux étapes de raffinement



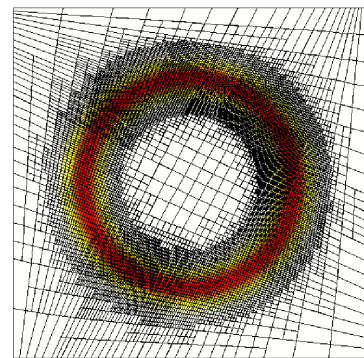
Trois étapes de raffinement

FIG. II.15 – Tri/Quadr-angles- P_1/Q_1 . Maillages raffinés.

Une étape de raffinement



Deux étapes de raffinement



Trois étapes de raffinement

FIG. II.16 – Quadrangles- Q_1 . Maillages raffinés.

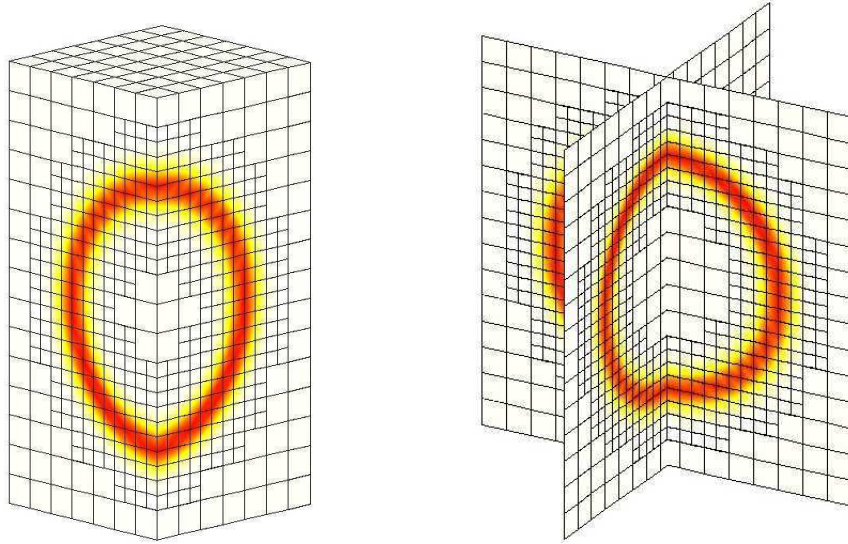


FIG. II.17 – Maillage raffiné en trois dimensions. Une étape de raffinement.

La table II.3 donne les ordres de convergence en norme L^2 , calculés par rapport à la taille des mailles du niveau de raffinement le plus fin utilisé dans le calcul. Nous obtenons des ordres de convergence en norme L^2 égaux à 2 pour les éléments finis d'ordre 1 ($\mathbb{P}_1$ et $\mathbb{Q}_1$) et égaux à 3 pour les éléments finis d'ordre 2 ($\mathbb{P}_2$ et $\mathbb{Q}_2$) comme attendu.

Dimensions	Maillages-Éléments de référence	Ordres de convergence
2D	Carrés- $\mathbb{Q}_1$	1.99
	Quadrangles- $\mathbb{Q}_1$	1.99
	Tri/Quadr-angles- $\mathbb{P}_1/\mathbb{Q}_1$	1.99
3D	Cubes- $\mathbb{Q}_1$	1.99
2D	Carrés- $\mathbb{Q}_2$	2.99
	Quadrangles- $\mathbb{Q}_2$	2.90
	Tri/Quadr-angles- $\mathbb{P}_2/\mathbb{Q}_2$	2.97

TAB. II.3 – Ordres de convergence en norme L^2 , calculés par rapport à la taille des mailles les plus fines utilisées dans le calcul.

Nous présentons maintenant les résultats obtenus avec les préconditionneurs multigrilles. Les tables II.4, II.5 et II.6 donnent le nombre d'itérations nécessaires à la méthode du gradient conjugué pour arriver à convergence (norme L^∞ relative du résidu inférieure à 10^{-10}) en fonction du nombre d'inconnues. La table II.4 donne ces résultats pour des éléments d'ordre 1 en deux dimensions, la table II.5 pour des éléments d'ordre 2 toujours en deux dimensions et enfin la table II.6 pour des éléments d'ordre 1 en trois dimensions. Différents préconditionneurs sont comparés : la classique factorisation LU incomplète (ILU0), et les versions additive (P_a) et multigrille (P_m) des préconditionneurs multiniveaux.

Le nombre total d'itérations dans la méthode du gradient conjugué est limité à 600, et dans les tables II.4, II.5 et II.6, nous notons par “—” les cas où la convergence n'est pas atteinte avant ce nombre maximal d'itérations.

Niveau de raffinement		1	2	3	4	5	6
Carrés- $\mathbb{Q}_1$	Nombre d'inconnues	893	3053	11021	41757	161233	633629
	Pas de préconditionneur	53	99	197	372	—	—
	ILU0	31	61	114	221	457	—
	P_a	18	23	27	29	33	34
	P_m	8	8	9	9	10	10
Quadrangles- $\mathbb{Q}_1$	Nombre d'inconnues	935	3217	11609	43593	168347	660710
	Pas de préconditionneur	176	357	—	—	—	—
	ILU0	45	84	183	388	—	—
	P_a	46	67	83	97	111	120
	P_m	15	16	20	21	23	24
Tri/Quadr-angles- $\mathbb{P}_1/\mathbb{Q}_1$	Nombre d'inconnues	869	2821	9893	37577	144969	569757
	Pas de preconditionneur	68	131	279	579	—	—
	ILU0	36	67	131	272	543	—
	P_a	21	29	35	40	43	47
	P_m	9	10	10	11	12	12

TAB. II.4 – Cas tests bidimensionnels. Nombre d'itérations dans la méthode du gradient conjugué en fonction du nombre d'inconnues.

Niveaux de raffinement		1	2	3	4
Carrés- $\mathbb{Q}_2$	Nombre d'inconnues	3653	12405	44485	166253
	Pas de préconditionneur	202	404	—	—
	ILU0	79	135	259	508
	P_a	36	56	69	77
	P_m	16	18	20	20
Quadrangles- $\mathbb{Q}_2$	Nombre d'inconnues	3827	13129	46881	175126
	Pas de preconditionneur	595	—	—	—
	ILU0	112	195	389	—
	P_a	83	127	158	178
	P_m	27	31	33	34
Tri/Quadr-angles- $\mathbb{P}_2/\mathbb{Q}_2$	Nombre d'inconnues	3569	11441	40293	151013
	Pas de préconditionneurs	180	361	—	—
	ILU0	80	152	309	—
	P_a	38	54	73	82
	P_m	17	18	20	21

TAB. II.5 – Cas tests bidimensionnels. Nombre d'itérations dans la méthode du gradient conjugué en fonction du nombre d'inconnues.

Niveau de raffinement		1	2	3
Cubes- $\mathbb{Q}_1$	Nombre d'inconnues	9942	63329	459063
	Pas de préconditionneurs	45	224	—
	ILU0	28	55	109
	P_a	21	31	35
	P_m	10	13	14

TAB. II.6 – Cas tests tridimensionnels. Nombre d'itérations dans le gradient conjugué en fonction du nombre d'inconnues.

Sans préconditionnement, le nombre d'itérations augmente rapidement avec le nombre d'inconnues. L'uti-

lisation des préconditionneurs permet de réduire le nombre d'itérations et en particulier, lorsque l'on utilise les préconditionneurs (P_m) et (P_a) le nombre d'itérations nécessaires pour arriver à convergence est (presque) indépendant de la taille du problème et significativement plus petit que lorsque ILU0 est utilisé.

II.4 Conclusion

Nous avons présenté dans ce chapitre, une méthode de raffinement local adaptatif. Celle-ci prend en compte implicitement les non-conformités possibles des maillages puisque les espaces d'approximation qu'elle permet de générer restent toujours H^1 -conformes. En outre, cette méthode permet facilement à partir d'un espace d'approximation donné, de construire une sous-suite d'espaces auxiliaires emboîtés qui peut être utilisée pour définir des préconditionneurs multiniveaux. Nous avons illustré, dans ce chapitre, les possibilités de la méthode sur des tests académiques; la partie 3 présente son application aux systèmes Cahn-Hilliard/Navier-Stokes (*cf* partie 2) pour la simulation d'écoulements incompressibles multiphasiques. Enfin, dans le chapitre suivant nous proposons une description précise des aspects liés à l'implémentation de ces méthodes.

Chapitre III

Implémentation en parallèle dans la bibliothèque PELICANS

Les méthodes décrites dans les chapitres I et II ont été implémentées, en collaboration avec B. Piar, dans la bibliothèque PELICANS (Plateforme Evolutive de Bibliothèques de Composants pour l'Analyse Numérique et la Simulation). Cette bibliothèque développée en C++, au sein du laboratoire DPAM/SEMIC/LIMSI de l'IRSN, fournit un ensemble de fonctionnalités pour faciliter le développement de logiciels de calcul scientifique. Elle est distribuée sous licence libre et est intégralement téléchargeable à l'adresse : <https://gforge.irsn.fr/gf/project/pelicans>.

Ce chapitre est dédié à la description de la version du code utilisée pour réaliser les simulations présentées dans la partie 3 de ce manuscrit. Nous nous concentrons sur les aspects directement reliés aux méthodes de raffinement local et préconditionneurs multigrilles, en insistant sur la structure du code source (*i.e.* les fonctionnalités apportées par chaque module et les relations entre ces différents modules) et les concepts qui ont régi son élaboration, la structure de données (liées aux choix particuliers de l'implémentation et à l'historique de développement de la plateforme PELICANS) étant décrite le plus succinctement possible, à titre d'exemple et pour donner un support concret à la discussion.

Nous exposons brièvement ci-après les principes généraux adoptés dans le développement de la plateforme PELICANS (un traité détaillé étant disponible dans la documentation PELICANS). Le code source de la bibliothèque est organisé en modules (structure de décomposition élémentaire regroupant des éléments définis par le programmeur et présentant une forte corrélation logique), les interactions entre ces différents modules étant gérées à l'aide des idées et techniques des programmations objet et par contrat.

La stratégie de conception employée est basée sur l'utilisation d'abstractions : un module abstrait (ou interface) représente un ensemble de comportements possibles, ceux-ci étant précisément spécifiés par contrat. En pratique, cela signifie que chaque méthode (ou fonction membre) du module satisfait des pré- et post-conditions : les pré-conditions sont les conditions exigées sur les arguments ou paramètres passés en entrée à la méthode ainsi que sur l'état de l'objet lui-même pour que l'exécution de celle-ci ait un sens, les post-conditions sont les conditions que la méthode garantit sur le résultat de son exécution (*i.e.* valeurs des paramètres et/ou état de l'objet). Ainsi, pré- et post-conditions constituent un véritable contrat que le module abstrait passe avec tous ses clients (ou utilisateurs), les pré-conditions étant la responsabilité du client, les post-conditions celle du module (abstrait).

Pour manipuler des abstractions, les langages objet fournissent le mécanisme fondamental d'héritage et les outils associés : le polymorphisme et le lien dynamique. L'héritage entre classes symbolise les relations hiérarchiques entre les concepts. Si la classe B hérite de la classe A (on dit également que la classe B dérive de la classe mère A) alors tout objet de type B est également de type A au sens où tous les services mis à disposition par A sont également mis à disposition par B, la classe B pouvant en définir de nouveaux. La classe B peut également redéfinir les méthodes de la classe A (on parle de surcharge des méthodes). La version de la méthode à exécuter sera alors déterminée par le type dynamique de l'objet utilisé pour l'appeler. Ceci signifie que la décision est prise au moment de l'exécution (et non de la compilation du code), on parle de lien dynamique. Ainsi, le contenu d'une telle méthode (déclarée virtuelle en C++) ne sera connu qu'au moment de son appel (on

parle de polymorphisme dynamique). C’est sur ce mécanisme subtil que repose la manipulation des abstractions.

Comme nous l’avons expliqué ci-dessus, la classe abstraite fixe de manière explicite les comportements de chacune des méthodes qu’elle définit (en les spécifiant par contrat). Chaque comportement particulier (pourvu qu’il respecte les spécifications de la classe mère) est alors implémenté dans une classe dérivée. Les clients n’ont pas besoin de connaître l’existence des classes dérivées, leur implémentation doit utiliser la classe abstraite en se basant, pour la correction du code, uniquement sur les spécifications qu’elle définit. La classe abstraite pourra *a posteriori* (i.e. à l’exécution) être remplacée par n’importe quelle classe dérivée sans changer le code source grâce au mécanisme de polymorphisme dynamique garantissant ainsi une extension possible des comportements (par l’ajout de classes dérivées) sans modification du code existant. On dit que le code est fermé pour les modifications et ouvert pour les extensions.

Les méthodes de raffinement local et multigrilles présentées dans les chapitres I et II se prêtent bien à ce type de programmation. La structure adoptée pour l’implémentation du module de raffinement local est exposée dans la section III.1. La section III.2 présente la stratégie utilisée pour numéroter les inconnues d’un problème discret (i.e. les différents degrés de liberté associés aux champs discrets inconnus), les objets effectuant ces tâches de numérotation étant intensivement utilisés dans le module de préconditionnement multiniveau dont la structure est exposée dans la section III.3. Et enfin, dans la section III.4, après une brève introduction au principe de fonctionnement de la librairie PELICANS en parallèle, sont présentés les développements qui ont été nécessaires à la bonne exécution en parallèle des modules décrits dans les sections III.1 et III.3.

Ce travail doit beaucoup à l’ensemble des développeurs PELICANS, en particulier à F. Babik et L. Chailan, puisque nous utilisons l’ensemble des fonctionnalités (parallèles) de la librairie.

Nous utilisons un pseudo-langage objet pour décrire les différents algorithmes. En particulier, nous employons la notation `objet->fonction` pour accéder à la fonction membre `fonction` de l’objet `objet`.

III.1 Organisation du module de raffinement local

La méthode de raffinement local présentée dans les chapitres I et II se déduit entièrement de la définition (abstraite) de motif de raffinement. Ainsi, nous utilisons le mécanisme d’abstraction (cf introduction de ce chapitre) pour implémenter la notion de motif de raffinement et tous les concepts sous-jacents : élément de référence, maillage du polyèdre de référence... La section III.1.1 décrit l’ensemble de ces classes abstraites. La section III.1.2 détaille ensuite les structures de données qui seront manipulées par les algorithmes d’adaptation présentés dans la section III.1.3. Ces algorithmes sont indépendants de l’élément de référence ou du motif de raffinement concret qui sera choisi par l’utilisateur, ils dépendent uniquement des classes abstraites. Ainsi, il est possible de permettre l’utilisation d’autres éléments finis de type Lagrange ou d’autres motifs de raffinement (que ceux actuellement disponibles) sans modifier aucune des classes présentées ci-dessous mais simplement en codant de nouvelles classes dérivées des classes abstraites (présentées dans la section III.1.1) implémentant concrètement ces éléments ou motifs de raffinement. Les fonctionnalités de raffinement local seront alors automatiquement disponibles.

III.1.1 Éléments de référence, motifs de raffinement et indicateurs de (dé)raffinement

Nous présentons dans cette section les classes abstraites de PELICANS implémentant les principaux concepts définis dans le chapitre I : la notion d’élément fini de référence de type Lagrange (cf définition I.1) et la notion de motif de raffinement (cf définition I.11). Les informations de type géométrique concernant ces objets sont stockées dans des structures abstraites séparées (les classes `GE_ReferencePolyhedron` et `GE_ReferencePolyhedronRefiner`) que nous ne décrivons que succinctement puisqu’elles interviennent très peu dans les algorithmes de raffinement local. Nous terminons cette section, par la description de la classe abstraite `PDE_AdaptationIndicator` qui implémente les fonctionnalités nécessaires à la sélection des fonctions à raffiner ou déraffiner.

Polyèdres de référence : la classe `GE_ReferencePolyhedron`

La classe `GE_ReferencePolyhedron` est une abstraction de la notion de support géométrique de référence (il s’agit du polyèdre noté $\hat{K}$ dans le chapitre I). Elle met à disposition l’ensemble des informations géométriques définissant le polyèdre de référence : dimension d’espace, nombre de sommets, coordonnées des sommets, aire ou volume, nombre de faces, connectivité faces/sommets, normales aux faces, centre de gravité...

EXEMPLES DE CLASSES CONCRÈTES DÉRIVÉES de la classe `GE_ReferencePolyhedron` :

`GE_ReferenceTriangle` (simplexe unité en 2D), `GE_ReferenceSquare` (carré unité),
`GE_ReferenceTetrahedron` (simplexe unité en 3D), `GE_ReferenceCube` (cube unité)...

Eléments de référence : la classe `PDE_ReferenceElement`

La classe `PDE_ReferenceElement` est une abstraction de la notion d'élément de référence de type Lagrange (*cf* définition I.1). Un objet `elt` de type `PDE_ReferenceElement` fournit les fonctionnalités suivantes :

- informations complètes sur le support géométrique de référence $\hat{K}$ en pointant vers un objet de la classe `GE_ReferencePolyhedron`,
- nombre N (`elt->nb_nodes()`), numérotation (locale) et coordonnées géométriques des noeuds de Lagrange $\hat{\mathbf{a}}_k$, $k = 1 \dots N$,
- valeurs des fonctions de base $\hat{\varphi}_k$, $k = 1 \dots N$, de leurs dérivées premières et secondes en chaque point du support géométrique.

EXEMPLES DE CLASSES CONCRÈTES DÉRIVÉES de la classe `PDE_ReferenceElement` :

`PDE_2D_P1_3nodes` (élément 2D triangle- $\mathbb{P}_1$), `PDE_2D_Q1_4nodes` (élément 2D carré- $\mathbb{Q}_1$),
`PDE_3D_Q1_8nodes` (élément 3D cube- $\mathbb{Q}_1$), `PDE_3D_Q2_27nodes` (élément 3D cube- $\mathbb{Q}_2$)...

Maillages du polyèdre de référence : la classe `GE_ReferencePolyhedronRefiner`

La classe `GE_ReferencePolyhedronRefiner` est une abstraction de la notion de maillage du support géométrique de l'élément de référence (notion faisant partie intégrante de la définition I.11 de motif de raffinement). Cette classe est associée à un élément de la classe `GE_ReferencePolyhedron`. Elle met les informations suivantes à disposition :

- nombre, numérotation et coordonnées des sommets du maillage,
- nombre et numérotation des sous-cellules,
- support géométrique des sous-cellules (supposé identique pour toutes les sous-cellules) en pointant sur un élément de la classe `GE_ReferencePolyhedron`,
- connectivité sous-cellules/sommets,
- informations diverses sur les faces (non décrites ici puisqu'elles ne sont pas directement liées aux méthodes présentées dans ce manuscrit).

EXEMPLES DE CLASSES CONCRÈTES DÉRIVÉES de la classe `GE_ReferencePolyhedronRefiner` :

`GE_ReferenceSquareWithSquares` (maillage du carré unité en quatre cellules),
`GE_ReferenceTriangleWithTriangles` (maillage du triangle unité en quatre cellules),
`GE_ReferenceCubeWithCubes` (maillage du cube unité en huit cellules)...

Motifs de raffinement : la classe `PDE_ReferenceElementRefiner`

La classe `PDE_ReferenceElementRefiner` est une abstraction de la notion de motif de raffinement (définition I.1). Nous la décrivons de manière plus précise que les précédentes puisqu'elle est le concept central de la méthode de raffinement local présentée dans les chapitres I et II. Conformément à la définition I.1, cette classe est associée à deux objets de classe respective `GE_ReferencePolyhedronRefiner` et `PDE_ReferenceElement`.

La classe `PDE_ReferenceElementRefiner` met à disposition des mécanismes de parcours des enfants, parents, descendants et ascendants de chaque noeud, ces relations sur les noeuds étant définies par analogie à celles existantes entre les fonctions de base qui leur sont associées (*cf* équation de raffinement I.14 et définitions I.20 et II.14). Les parcours sont formalisés dans la suite à l'aide d'itérateurs par quatre méthodes de la forme :

- `start>(*s)_iterator( ic, node )` : initialise l'itérateur associé à la sous-cellule `ic` et au noeud `node`,
- `(*)_is_valid()` : booléen indiquant si la position courante de l'itérateur est valide (vrai) ou non (faux), lorsque celle-ci ne l'est plus, le parcours est fini,
- `current(*)_node()` : numéro du noeud parcouru (*i.e.* celui désigné par la position courante de l'itérateur),
- `next(*)` : déplace l'itérateur vers la position suivante.

Le symbole `(*)` peut être remplacé par `child` (resp. `parent`, resp. `descendant`, resp. `ascendant`) pour un parcours des enfants (resp. parents, resp. descendants, resp. ascendants), la forme `(*s)` désignant le pluriel `children` (resp. `parents`, resp. `descendants`, resp. `ascendants`). Pour décrire plus précisément ces méthodes, nous adoptons les conventions de notations suivantes :

- `ic` désigne le numéro d'une sous-cellule du support de l'élément de référence,
- `rnode` désigne le numéro local d'un noeud de la sous-cellule de numéro `ic`,

- `cnode` désigne le numéro local d'un noeud de l'élément de référence.

De plus, lorsque ceci ne porte pas à confusion nous utiliserons le numéro de l'objet pour le désigner directement : ainsi, nous parlerons de la sous-cellule `ic` ou du noeud `cnode` à la place de la sous-cellule de numéro `ic` ou du noeud de numéro `cnode`. Revenons à la description des méthodes de parcours introduites ci-dessus. Sur chaque sous-cellule `ic` :

- étant donné le noeud `cnode` de la cellule de référence, la méthode `start_children_iterator( ic, cnode )` (resp. `start_descendants_iterator( ic, cnode )`) permet de commencer le parcours de tous les noeuds `rnnode` enfants (resp. descendants) du noeud `cnode` appartenant à la sous-cellule `ic`.
- étant donné un noeud `rnnode` de la sous-cellule `ic`, la méthode `start_parents_iterator( ic, rnnode )` (resp. `start_ascendants_iterator( ic, rnnode )`) permet de commencer le parcours de tous les noeuds `cnode` parents (resp. ascendants) du noeud `rnnode`.

Remarque III.1

A chaque élément de référence concret (i.e. à chaque classe dérivée de la classe abstraite `PDE_ReferenceElement`), nous devons pouvoir associer un motif de raffinement concret (i.e. une classe dérivée de la classe abstraite `PDE_ReferenceElementRefiner`). Le mécanisme de correspondance n'est pas décrit ici, nous noterons dans la suite `reference_element_refiner( elt )` le motif de raffinement associé à l'élément de référence `elt`.

EXEMPLES DE CLASSES CONCRÈTES DÉRIVÉES de la classe `PDE_ReferenceElementRefiner` :

`PDE_2D_P1_3nodesRefinerA` (associé à `PDE_2D_P1_3nodes` et `GE_ReferenceTriangleWithTriangles`, cf figure I.5),
`PDE_2D_Q1_4nodesRefinerA` (associé à `PDE_2D_Q1_4nodes` et `GE_ReferenceSquareWithSquares`, cf figure I.4),
`PDE_2D_Q2_9nodesRefinerA` (associé à `PDE_2D_Q2_9nodes` et `GE_ReferenceSquareWithSquares`, cf figure I.6),
`PDE_3D_Q1_8nodesRefinerA` (associé à `PDE_3D_Q1_8nodes` et `GE_ReferenceCubeWithCubes`)...

Indicateurs de (dé)raffinement : la classe `PDE_AdaptationIndicator`

La classe `PDE_AdaptationIndicator` est une abstraction de la notion d'indicateur de raffinement. Elle met à disposition les fonctionnalités suivantes :

- valeur d'un indicateur par cellule,
- décision de raffiner ou non une fonction de base.

La décision de raffiner ou non une fonction de base peut être basée (comme c'est souvent le cas) sur l'indicateur calculé sur chaque cellule (en faisant par exemple une moyenne sur les cellules du support de la fonction de base considérée).

Dans les listings présentés dans les sections suivantes, nous supposons qu'un objet appartenant à la classe `PDE_AdaptationIndicator`, noté `INDIC` est donné. Cet objet sera utilisé de la manière suivante : pour une fonction de base `bf` donnée, le booléen `INDIC->to_refined( bf )` (resp. `INDIC->to_unrefined( bf )`) est alors vrai lorsque l'indicateur a désigné `bf` comme une fonction de base à raffiner (resp. déraffiner).

EXEMPLES DE CLASSES CONCRÈTES DÉRIVÉES de la classe `PDE_AdaptationIndicator` :

`PDE_GeometricIndicator` (indicateur de raffinement géométrique, cf critère II.31),
`CH_InterfaceIndicator` (cf critère II.31)...

III.1.2 Champs discrets, cellules et fonctions de base

Cette section décrit les structures de données qui seront manipulées dans la section III.1.3 : il s'agit des champs discrets, des fonctions de base et des cellules du maillage.

Précisons tout d'abord la terminologie que nous allons utiliser dans la suite. Les algorithmes de raffinement local présentés dans la section III.1.3 présupposent que tous les champs inconnus (i.e. que nous devons calculer dans un pas de temps) possédant la même discrétisation sont exprimés dans la même base multiniveau, on parlera dans la suite de bases multiniveaux courantes. Ces bases multiniveaux courantes constituent l'ensemble des fonctions de base dites actives. Les algorithmes d'adaptation consistent donc à faire évoluer les bases multiniveaux courantes en activant ou désactivant des fonctions de base conformément aux procédures de l'algorithme II.4.

Champs discrets : la classe `PDE_DiscreteField`

La classe `PDE_DiscreteField` implémente la notion de champ discret. Cette classe est généralement utilisée pour représenter les inconnues du problème discret ou leurs représentations explicites (i.e. aux temps discrets

précédents) pour des problèmes instationnaires. Le rôle de cette classe est uniquement d'organiser l'ensemble des valeurs numériques (*i.e.* degrés de liberté) du champ discret sous forme d'une indexation formalisée. La numérotation est effectuée à l'aide des deux indices de noeuds et de composantes (pour les champs vectoriels). Les noeuds sont associés aux degrés de liberté de chacune des composantes. La numérotation des noeuds est propre à chaque objet de type `PDE_DiscreteField` et il est important de comprendre que chacun de ces objets contient un noeud associé à chacune des fonctions de base compatibles avec la discrétisation du champ sous-jacent. Le nombre de noeuds d'un champ est donc égal au nombre total de fonctions de base (actives ou non) compatibles avec sa discrétisation. Cependant, certaines fonctions de base, même dans le cas où elles seraient actives (resp. inactives), peuvent ne pas intervenir (resp. peuvent intervenir) dans la discrétisation du champ considéré. C'est par exemple le cas lors de la résolution d'un problème instationnaire : le problème discret fait à la fois intervenir des champs discrétisés dans les espaces d'approximation courants mais également des champs discrétisés dans les espaces d'approximation à des temps précédents (*cf* section II.1.5). C'est alors un maillage multiniveaux (*cf* définition II.21) qu'il faut utiliser pour effectuer les intégrations numériques. Sa définition est basée sur la notion de degré de liberté actif (*cf* définition II.17), c'est elle que nous implémentons en ajoutant la notion de noeud actif. Les noeuds actifs sont définis de manière indépendante pour chaque champ particulier. Le champ s'exprime alors comme combinaison linéaire des fonctions de base associées aux noeuds actifs. Tout noeud actif correspond à un degré de liberté actif : c'est de ce principe que sont déduites les cellules actives (*cf* définition II.19).

Un objet `ff` de la classe `PDE_DiscreteField` fournit les fonctionnalités suivantes :

- `ff->nb_nodes()` : nombre total de noeuds (actifs ou non),
- `ff->nb_components()` : nombre de composantes,
- diverses opérations sur les noeuds :
 - `ff->add_nodes( nb_new_nodes )` : crée `nb_new_nodes` nouveaux noeuds,
 - `ff->node_is_active( node )` : booléen indiquant si le noeud `node` est actif (vrai) ou non (faux),
 - `ff->set_node_active( node )`, `ff->set_node_inactive( node )` : active, désactive le noeud `node`,
 - `ff->start_nodes_iterator()`, `ff->node_is_valid()`, `ff->current_node()`, `ff->next_node()` : itérateur permettant de parcourir l'ensemble des noeuds (actifs ou non),
- diverses opérations (accès, modification...) sur la valeur du degré de liberté (correspondant à chaque couple formé d'un noeud et d'une composante).

Enfin, la classe `PDE_DiscreteField` permet de gérer les degrés de liberté imposés (par exemple, pour prendre en compte des conditions aux bords de type Dirichlet). Nous ne décrirons pas cet aspect dans ce manuscrit.

Remarque III.2

Remarquons que la classe `PDE_DiscreteField` en elle-même ne fournit aucune information relative au type de discrétisation utilisée, par exemple elle ne fournit pas de correspondance entre les valeurs de degrés de liberté stockées et l'approximation sous-jacente du champ continu. Ces correspondances très importantes pourront être retrouvées grâce aux objets `PDE_BasisFunctionCell` et `PDE_CellFE` décrits ci-dessous.

Fonctions de base : la classe `PDE_BasisFunctionCell`

La classe `PDE_BasisFunctionCell` définit la notion de fonction de base. Tout objet `bf` appartenant à la classe `PDE_BasisFunctionCell` possède les caractéristiques suivantes :

- `bf->refinement_level()` : entier indiquant le niveau de raffinement de `bf`,
- `bf->is_active()` : booléen indiquant si `bf` est active (vrai) ou non (faux) (*i.e.* si `bf` appartient ou non à la base multiniveau courante),
- `bf->is_refined()` : booléen indiquant si `bf` est raffinée (vrai) ou non (faux) par rapport à la base multiniveau courante, (*cf* définition II.1).

La présence du booléen `bf->is_refined()` nécessite quelques explications. Rappelons que les procédures de raffinement de l'algorithme II.4 font intervenir la définition II.1 de fonction de base raffinée (par rapport à la base multiniveau courante). Il faut donc pouvoir accéder à cette information dans le code de calcul. Nous avons choisi de la stocker dans chaque fonction de base. En effet, il est possible d'identifier les fonctions de base raffinées sans avoir recours à la définition II.1. L'exécution du code de calcul commence toujours par la création du maillage grossier T_0 et des fonctions de base associées, aucune des fonctions de base créées à l'initialisation n'est donc raffinée au sens de la définition II.1. Les fonctions de base raffinées au sens de cette définition seront donc celles que nous avons raffinées par la procédure de raffinement (*cf* remarque II.5) et pas encore déraffinées. L'information peut donc être mise à jour au cours de la procédure de (dé)raffinement.

Enfin, la classe `PDE_BasisFunctionCell` met à disposition les services suivants :

- modifications des caractéristiques :
 - `bf->set_refined()`, `bf->set_unrefined()` : modification de l'état "`bf->is_refined()`",
 - `bf->set_active()`, `bf->set_inactive()` : modification de l'état "`bf->is_active()`",
- `bf->start_cells_iterator()`, `bf->next_cell()`, `bf->cell_is_valid()` : itérateur permettant de parcourir les cellules du support de `bf`. A chaque position de l'itérateur, les informations suivantes sont disponibles :
 - `bf->current_cell()` : cellule du support (de type `PDE_MeshFE`) désignée par la position courante de l'itérateur,
 - `bf->reference_element_of_current_cell()` : élément de référence (de type `PDE_ReferenceElement`) associée à `bf` sur la cellule courante,
 - `bf->local_node_in_current_cell()` : numéro local du noeud associé à la fonction de base `bf` dans la cellule courante,
- `bf->start_fields_iterator()`, `bf->next_field()`, `bf->field_is_valid()` : itérateur sur les champs dont la discrétisation est compatible avec `bf`. A chaque position courante de cet itérateur, les informations disponibles sont :
 - `bf->current_field()` : champ (de type `PDE_DiscreteField`) désigné par la position courante de l'itérateur,
 - `bf->node_of_current_field()` : noeud du champ `bf->current_field()` associé à `bf`.
- `bf->node_of_DOF( ff )` : noeud du champ `ff` associé à `bf` (lorsque celui-ci existe).

Remarque III.3

Il ne faut pas confondre les noeuds renvoyés par les fonctions `bf->local_node_in_current_cell()` et `bf->node_of_current_field()` (ou `bf->node_of_DOF( ff )`). Le premier désigne un numéro de noeud local à une cellule (l'indexation d'un tel noeud est donc définie par et sur l'élément de référence) alors que le second désigne un noeud associé à un degré de liberté (par composante) d'un champ discret (l'indexation est réalisée par le champ discret lui-même).

Cellules : la classe `PDE_CellFE`

La classe `PDE_CellFE` implémente la notion de cellule élément fini. Tout objet `cell` de la classe `PDE_CellFE` met à disposition les informations suivantes :

- diverses informations d'ordre géométrique (nombre de sommets, de faces, connectivité face/sommets...), que nous ne détaillons pas ici,
- `cell->refinement_level()` : entier indiquant le niveau de raffinement,
- `cell->is_active()` : booléen indiquant si la cellule `cell` est active (vrai) ou non (faux),
- `cell->parent()` : cellule parent (elle est unique),
- `cell->nb_children()` : nombre de cellules enfants,
- `cell->child( ic )` : ic^e cellule enfant de `cell`. L'indexation est donnée (à la création des cellules enfants) par le mapping avec le maillage de l'élément de référence `GE_ReferencePolyhedronRefiner`. Ceci est important car c'est ce même indice `ic` qui est utilisé sur le motif de raffinement `PDE_ReferenceElementRefiner` pour retrouver la liste des enfants, parents, ascendants et descendants d'un noeud,
- `cell->reference_element( ff )` : élément de référence associé au champ `ff` sur la cellule `cell`,
- `cell->bf( node, elt )` : fonction de base associée au noeud local `node` de la cellule `cell` pour l'élément de référence `elt`.

Le booléen `cell->is_active()` est stocké dans chaque cellule. L'information n'est pas obtenue en appliquant *stricto sensu* la définition II.19 mais est mise à jour en fin de procédure d'adaptation (cf section III.1.3).

L'ensemble de toutes les cellules (ou objets de la classe `PDE_CellFE`) du domaine est accessible à travers un objet que nous noterons `GRID` dans les listings de la section suivante. Cet objet permet en particulier de parcourir toutes les cellules (quel que soit leur niveau de raffinement).

III.1.3 Algorithmes d'adaptation

Les fonctionnalités des classes décrites dans les sections précédentes permettent d'implémenter les algorithmes de raffinement local. Ceux-ci sont regroupés dans quatre classes :

- la classe `PDE_FamilyCHARMS` permet d'effectuer les parcours des enfants, parents, descendants ou ascendants d'une fonction de base. Son utilisation évite le stockage de ces informations dans chaque objet de type `PDE_BasisFunctionCell`,
- la classe `PDE_AdaptationRequest` permet d'appliquer les critères de raffinement ou déraffinement (de type `PDE_AdaptationIndicator`) ainsi que de faire respecter la règle "au-plus-un-niveau-de-différence" ,

- la classe `PDE_Activator` implémente les procédures de raffinement et de déraffinement de l'algorithme II.4, et permet de gérer l'activation et la désactivation des cellules en fin de procédure d'adaptation,
- la classe `PDE_AdapterCHARMS` est chargée de créer les nouveaux objets (fonctions de base, cellules, faces...), et d'implémenter le cycle complet de la procédure d'adaptation.

Enfants, parents, ascendants, descendants d'une fonction de base : la classe `PDE_FamilyCHARMS`

La classe `PDE_FamilyCHARMS` fournit les fonctionnalités nécessaires au parcours de tous les enfants, parents, descendants ou ascendants d'une fonction de base donnée `bf`. Ces parcours sont effectués en stockant préalablement les fonctions de base à parcourir (en faible nombre) grâce aux algorithmes donnés dans les listings III.1 et III.2. Le stockage des enfants (ou descendants) de `bf` s'effectue en parcourant chaque cellule du support de `bf`, sur chacune de ces cellules le motif de raffinement nous permet d'obtenir la liste des enfants (ou descendants) sous-cellule par sous-cellule. Le stockage des parents (ou ascendants) s'effectue également en parcourant les cellules du support de `bf`, puis en considérant, pour chacune d'elles, leur cellule parent sur laquelle le motif de raffinement permet de déduire la liste des parents (ou ascendants).

Listing III.1

Stockage des enfants existants d'une fonction de base `bf`
dans la liste `CHILDREN` en vue d'un parcours

```

1 // Parcours des cellules du support de bf : cell
2 bf->start_cells_iterator() ;
3 for( ; bf->cell_is_valid() ; bf->next_cell() )
4 {
5     cell = bf->current_cell() ;
6     // Sur la cellule cell,
7     // bf est la fonction de base associee
8     // - au noeud : node
9     // - a l'element de reference : elt
10    node = bf->local_node_in_current_cell() ;
11    elt = bf->reference_element_of_current_cell() ;
12    // Motif de raffinement : elr
13    // associe a l'element elt
14    elr = reference_element_refiner( elt ) ;
15
16    // Parcours des cellules enfants : rcell de cell
17    for( ic = 0 ; ic < cell->nb_children() ; ic++ )
18    {
19        rcell = cell->child( ic ) ;
20        // Parcours des enfants : rbf de la fonction
21        // de base bf sur la cellule rcell
22        // grace au motif de raffinement
23        elr->start_children_iterator( ic, node ) ;
24        for( ; elr->child_is_valid()
25            ; elr->next_child() )
26        {
27            rnode = elr->current_child_node() ;
28            rbf = rcell->bf( rnode, elt ) ;
29            // Attention : presuppone que
30            // l'element de reference associe
31            // a bf est le meme sur une cellule
32            // et sur toutes ses cellules enfant
33
34            // Si la fonction de base enfant existe,
35            // on l'ajoute a la liste
36            if( rbf != 0 ) CHILDREN->extend( rbf ) ;
37        }
38    }
39 }
```

Listing III.2

Stockage des parents existants d'une fonction de base `bf`
dans la liste `PARENTS` en vue d'un parcours

```

1 // Parcours des cellules du support de bf : rcell
2 bf->start_cells_iterator() ;
3 for( ; bf->cell_is_valid() ; bf->next_cell() )
4 {
5     cell = bf->current_cell() ;
6     // Sur la cellule cell,
7     // bf est la fonction de base associee
8     // - au noeud : node
9     // - a l'element de reference : elt
10    node = bf->local_node_in_current_cell() ;
11    elt = bf->reference_element_of_current_cell() ;
12    // Motif de raffinement : elr
13    // associe a l'element elt
14    elr = reference_element_refiner( elt ) ;
15
16    // la cellule parent : pcell de cell est unique
17    pcell = cell->parent() ;
18    // l'indice ic est tel que
19    // cell = pcell->child( ic )
20    ic = index_of_child_cell( pcell, cell ) ;
21
22    // Parcours des parents : pbf
23    // de la fonction de base bf
24    // sur la cellule cell
25    // grace au motif de raffinement
26    elr->start_parents_iterator( ic, node ) ;
27    for( ; elr->parent_is_valid()
28        ; elr->next_parent() )
29    {
30        pNode = elr->current_parent_node() ;
31        pbf = pcell->bf( pNode, elt ) ;
32        // Si la fonction de base parent existe,
33        // on l'ajoute a la liste
34        if( pbf != 0 ) PARENTS->extend( pbf ) ;
35    }
36 }
```

L'algorithme donnant la liste des descendants (resp. des ascendants) est similaire à celui donnant celle des enfants (resp. des parents). Il suffit de modifier les appels aux routines de parcours sur le motif de raffinement : l'algorithme de stockage des descendants est obtenu en remplaçant, dans le listing III.1, l'itérateur sur les enfants par un itérateur sur les descendants (par exemple, `elr->start_children_iterator` est remplacé par `elr->start_descendants_iterator`, `elr->child_is_valid` par `elr->descendant_is_valid`...); de même l'algorithme de stockage des ascendants est obtenu en remplaçant, dans le listing III.2, l'itérateur sur les parents par un itérateur sur les ascendants.

Remarque III.4

Dans le cas de la recherche des parents la boucle sur les cellules du support n'est pas nécessaire (contrairement au cas de la recherche des ascendants) : il suffit en fait de considérer une quelconque des cellules du support. En effet, toutes les cellules du support considérées conduiront à la même liste de parents. Notons K la cellule parent d'une cellule quelconque du support de `bf`. Le noeud associé à la fonction de base `bf` appartient à toutes les cellules du support de `bf` et par conséquent à leurs cellules parents (en particulier, à K). Considérons maintenant une fonction de base parent φ quelconque de `bf`. Ce parent φ est par définition non nul au noeud associé à `bf`. Ainsi, ce noeud est dans le support de φ , et par conséquent la cellule K fait partie du support de φ . Donc, ce parent φ sera trouvé à l'aide du motif de raffinement appliqué à K .

Dans la suite, nous noterons `FMS` un objet de la classe `PDE_FamilyCHARMS`, celui-ci permettant d'effectuer les parcours des enfants, parents, descendants ou ascendants d'une fonction de base `bf` à l'aide d'itérateurs :

- `start>(*s)_iterator( bf )` : initialise l'itérateur,
- `(*)_is_valid()` : booléen indiquant si la position courante de l'itérateur est valide (vrai) ou non (faux), lorsque celle-ci ne l'est plus, le parcours est fini,
- `current(*)()` : fonction de base parcourue (*i.e.* celle désignée par la position courante de l'itérateur),
- `next(*)` : déplace l'itérateur vers la position suivante.

Le symbole `(*)` peut être remplacé par `child` (resp. `parent`, resp. `desct`, resp. `asct`) pour un parcours des enfants (resp. parents, resp. descendants, resp. ascendants), la forme `(*s)` désignant le pluriel `children` (resp. `parents`, resp. `descts`, resp. `ascts`).

Sélection des fonctions de base à (dé)raffiner : la classe `PDE_AdaptationRequest`

Le rôle de la classe `PDE_AdaptationRequest` est de fournir les fonctionnalités nécessaires à la sélection des fonctions de base à raffiner et à déraffiner. Ceci est effectué en deux étapes :

- sélection des fonctions de base à l'aide des critères de raffinement et déraffinement (définis de manière abstraite par la classe `PDE_AdaptationIndicator`),
- application de critères plus spécifiques assurant le respect de la règle “au-plus-un-niveau-de-différence” (*cf* section II.1.4).

Ainsi, la première étape consiste à parcourir l'ensemble des fonctions de base et à sélectionner celles qui sont admissibles au sens défini par le critère de raffinement ou déraffinement. Les listings III.3 et III.4 présentent ces parcours, ils sont donnés à titre d'exemple puisque la façon de parcourir les fonctions de base est intimement liée à la structure de données dont nous disposons dans le code. Ils conduisent à la création de deux listes préliminaires `BFS_CRIT_R` et `BFS_CRIT_U` contenant les fonctions de base désignées par les critères de raffinement et déraffinement respectivement. Les listings III.5, III.6, III.7 et III.8 montrent ensuite comment la règle “au-plus-un-niveau-de-différence” est appliquée en pratique. Encore une fois, le point important est que toutes les informations sont déduites des relations données sur le motif de raffinement (via le parcours des descendants et ascendants *cf* section III.1.3). Les listings III.5 et III.6 présentent les routines de sélection des fonctions de base à proprement parler. Chacune de ces routines consiste en un parcours de la liste `BFS_CRIT_R` (resp. `BFS_CRIT_U`) fournie par la première étape, la décision d'ajouter ou non chaque fonction base parcourue à la liste définitive `BFS_R` (resp. `BFS_U`) étant déléguée aux fonctions auxiliaires présentées dans les listings III.7 et III.8. Conformément à l'algorithme II.15, nous ajoutons à la liste des fonctions de base désignées par le critère de raffinement tous ses ascendants (existants) et ce récursivement (*i.e.* les ascendants des ascendants ... jusqu'au niveau le plus grossier où les fonctions de base n'ont pas d'ascendant). Par contre, nous supprimons de la liste des fonctions désignées par le critère de déraffinement toutes celles ayant un parent ou un ascendant raffiné.

Listing III.3

Construction de la liste BFS_CRIT_R des fonctions de base désignées par le critère de raffinement

```

1 // Parcours de toutes les cellules : cell
2 // existantes (de tous les niveaux)
3 GRID->start_cells_iterator() ;
4 for( ; GRID->cell_is_valid() ; GRID->next_cell() )
5 {
6     cell = GRID->current_cell() ;
7
8     // limitation eventuelle sur
9     // le niveau le plus haut autorise
10
11     // Parcours des elements de reference : elt
12     // de la cellule cell
13     cell->start_elts_iterator() ;
14     for( ; cell->elt_is_valid() ; cell->next_elt() )
15     {
16         elt = cell->current_elt() ;
17         // Parcours des fonctions de base : bf
18         // de la cellule cell
19         elt->start_nodes_iterator() ;
20         for( ; elt->node_is_valid() ; elt->next_node() )
21         {
22             node = elt->current_node() ;
23             bf = cell->bf( node, elt ) ;
24             // Ajout des fonctions de base existantes,
25             // actives, non raffinees
26             // et designees par le critere
27             if( bf != 0 && bf->is_active() )
28             {
29                 if( !bf->is_refined() && to_refine( bf ) )
30                 {
31                     BFS_CRIT_R->extend( bf ) ;
32                 }
33             }
34         }
35     }
36 }

```

Listing III.4

Construction de la liste BFS_CRIT_U des fonctions de base désignées par le critère de déraffinement

```

1 // Parcours de toutes les cellules : cell
2 // existantes (de tous les niveaux)
3 GRID->start_cells_iterator() ;
4 for( ; GRID->cell_is_valid() ; GRID->next_cell() )
5 {
6     cell = GRID->current_cell() ;
7
8     // Parcours des elements de reference : elt
9     // de la cellule cell
10    cell->start_elts_iterator() ;
11    for( ; cell->elt_is_valid() ; cell->next_elt() )
12    {
13        elt = cell->current_elt() ;
14
15        // Parcours des fonctions de base : bf
16        // de la cellule cell
17        elt->start_nodes_iterator() ;
18        for( ; elt->node_is_valid() ; elt->next_node() )
19        {
20            node = elt->current_node() ;
21            bf = cell->bf( node, elt ) ;
22            // Ajout des fonctions de base existantes,
23            // raffinees et designees par le critere
24            if( bf != 0 )
25            {
26                if( bf->is_refined() && to_unrefine( bf ) )
27                {
28                    BFS_CRIT_U->extend( bf ) ;
29                }
30            }
31        }
32    }
33 }

```

Listing III.5

Construction de la liste finale BFS_R des fonctions de base à raffiner

```

1 // Parcours des fonctions de base
2 // pre-selectionnees au cours de la
3 // premiere etape
4 BFS_CRIT_R->start_items_iterator() ;
5 while( BFS_CRIT_R->item_is_valid() )
6 {
7     bf = BFS_CRIT_R->current_item() ;
8
9     // Ajout recursif de chaque fonction et
10    // de ses ascendants
11    extend_bf_with_ascendant( bf ) ;
12
13    BFS_CRIT_R->next_item() ;
14 }

```

Listing III.6

Construction de la liste finale BFS_U des fonctions de base à déraffiner

```

1 // Parcours des fonctions de base
2 // pre-selectionnees au cours de la
3 // premiere etape
4 BFS_CRIT_U->start_items_iterator() ;
5 while( BFS_CRIT_U->item_is_valid() )
6 {
7     bf = BFS_CRIT_U->current_item() ;
8
9     // Ajout conditionnel de la fonction
10    // de base
11    consider_unrefining_of( bf ) ;
12
13    BFS_CRIT_U->next_item() ;
14 }

```

Listing III.7

Fonctions auxiliaires : `extend_bf_with_ascendant( bf )`

```

1 //Ajout de la fonction de base
2 BFS_R->extend( bf ) ;
3 // Parcours de ses ascendants
4 FMS->start_ascts_iterator( bf ) ;
5 for( ; FMS->asct_is_valid() ; FMS->next_asct() )
6 {
7     pbf = FMS->current_asct() ;
8
9     // Ajout recursif de chaque ascendant actif
10    // et non raffine.
11    if( pbf->is_active() && !pbf->is_refined() )
12    {
13        extend_bf_with_ascendant( pbf ) ;
14    }
15 }

```

Listing III.8

Fonctions auxiliaires : `consider_unrefining_of( bf )`

```

1 // Est-ce que la fonction de base bf
2 // doit etre raffinee ?
3 may_be_unref = true ;
4
5 //Parcours des enfants de bf
6 FMS->start_children_iterator( bf ) ;
7 for( ; FMS->child_is_valid() ; FMS->next_child() )
8 {
9     rbf = FMS->current_child() ;
10    if( rbf->is_refined() ) may_be_unref= false ;
11 }
12 // La fonction de base bf ne sera pas deraffinee
13 // si un de ses enfants est raffinee
14
15 // Parcours des descendants de bf
16 FMS->start_descts_iterator( bf ) ;
17 for( ; FMS->descts_is_valid() ; FMS->next_desct() )
18 {
19     rbf = FMS->current_desct() ;
20     if( rbf->is_refined() ) may_be_unref = false ;
21 }
22 // La fonction de base bf ne sera pas deraffinee
23 // si un de ses descendants est raffinee.
24
25 // Ajout conditionnel de bf
26 if( may_be_unref ) BFS_U->extend( bf ) ;

```

Raffinement et déraffinement d'un ensemble de fonctions de base : la classe `PDE_Activator`

La classe `PDE_Activator` implémente l'algorithme II.13 de raffinement et de déraffinement d'un ensemble de fonctions de base. Rappelons que les ensembles des fonctions de base à raffiner (`BFS_R`) et à déraffiner (`BFS_U`) sont fournis par la classe `PDE_AdaptationRequest` présentée dans le paragraphe précédent.

Le raffinement de l'ensemble `BFS_R` (*cf* listing III.9) s'effectue par un parcours des fonctions de base `cbf` qu'il contient (l'ordre de parcours est indifférent, *cf* proposition II.12). Pour chacune de ces fonctions de base `cbf`, l'algorithme de raffinement II.4 est effectué :

- désactivation de la fonction de base `cbf`,
- activation de tous ses enfants `rbf` non raffinés.

De plus, la fonction de base `cbf` est maintenant raffinée au sens de la définition II.1 (*cf* remarque II.5). Ceci est gardé en mémoire à l'aide du booléen prévu à cet effet : `cbf->set_refined()`. Enfin, cet algorithme prend en compte la gestion des noeuds actifs de chaque champ. Pour cela la classe `PDE_Activator` définit une liste de champs exclus : le booléen `is_excluded( ff )` indique si le champ `ff` est exclu (vrai) ou non (faux). Cette liste doit être communiquée par l'utilisateur depuis le jeu de données, les champs exclus étant ceux dont la discrétisation n'est pas effectuée dans la base multiniveau courante. L'activation et la désactivation des noeuds ne les concernent pas.

Listing III.9

Raffinement d'un ensemble de fonctions de base

```

1 // Parcours de la liste BFS_R
2 // des fonctions de base a raffiner : cbf
3 BFS_R->start_items_iterator();
4 for( ; BFS_R->item_is_valid() ; BFS_R->next_item() )
5 {
6     cbf = BFS_R->current_item();
7
8     // desactivation de cbf
9     cbf->set_inactive();
10    // validation de l'indicateur de raffinement
11    cbf->set_refined();
12    // desactivation des noeuds des champs non exclus
13    // boucle sur les champs associes a cbf
14    cbf->start_fields_iterator();
15    for( ; cbf->field_is_valid() ; cbf->next_field() )
16    {
17        ff = cbf->current_field();
18        node = cbf->node_of_current_field();
19        if( !is_excluded( ff ) )
20        {
21            ff->set_node_inactive( node );
22        }
23    }
24    // Parcours des enfants de cbf : rbf
25    FMS->start_children_iterator( cbf );
26    for( ; FMS->child_is_valid() ; FMS->next_child() )
27    {
28        rbf = FMS->current_child();
29        if( !rbf->is_refined() )
30        {
31            // activation des enfants rbf (de cbf)
32            // non raffines
33            rbf->set_active();
34            // activation des noeuds des champs non exclus
35            // boucle sur les champs associes a cbf
36            rbf->start_fields_iterator();
37            for( ; rbf->field_is_valid() ; rbf->next_field() )
38            {
39                ff = rbf->current_field();
40                rnode = rbf->node_of_current_field();
41                if( !is_excluded( ff ) )
42                {
43                    ff->set_node_active( rnode );
44                }
45            }
46        }
47    }
48 }
49 }

```

Listing III.10

Déraffinement d'un ensemble de fonctions de base

```

1 // Parcours de la liste BFS_U
2 // des fonctions de base a deraffiner : cbf
3 BFS_U->start_items_iterator();
4 for( ; BFS_U->item_is_valid() ; BFS_U->next_item() )
5 {
6     cbf = BFS_U->current_item();
7     // activation de cbf
8     cbf->set_active();
9     // invalidation de l'indicateur de raffinement
10    cbf->set_unrefined();
11    // activation des noeuds des champs non exclus
12    // boucle sur les champs associes a cbf
13    cbf->start_fields_iterator();
14    for( ; cbf->field_is_valid() ; cbf->next_field() )
15    {
16        ff = cbf->current_field();
17        node = cbf->node_of_current_field();
18        if( !is_excluded( ff ) )
19            ff->set_node_active( node );
20    }
21    // Parcours des enfants de cbf : rbf
22    FMS->start_children_iterator( cbf );
23    for( ; FMS->child_is_valid() ; FMS->next_child() )
24    {
25        rbf = FMS->current_child();
26        // Est-ce que l'enfant rbf doit etre desactive ?
27        to_deactivate = true;
28        FMS->start_parents_iterator( rbf );
29        for( ; FMS->parent_is_valid() ; FMS->next_parent() )
30        {
31            pbf = FMS->current_parent();
32            if( pbf->is_refined() ) to_deactivate = false;
33        }
34        // desactivation des enfants rbf
35        // de cbf sans parent raffine
36        if( to_deactivate )
37        {
38            rbf->set_inactive();
39            // desactivation des noeuds des champs
40            // non exclus
41            // boucle sur les champs associes a rbf
42            rbf->start_fields_iterator();
43            for( ; rbf->field_is_valid() ; rbf->next_field() )
44            {
45                ff = rbf->current_field();
46                rnode = rbf->node_of_current_field();
47                if( !is_excluded( ff ) )
48                    ff->set_node_active( rnode );
49            }
50        }
51    }
52 }
53 }
54 }

```

Le déraffinement de l'ensemble **BFS_U** (cf listing III.10) s'effectue par un parcours des fonctions de base **cbf** qu'il contient (l'ordre de parcours n'a pas d'importance, cf proposition II.12). Pour chacune de ces fonctions de base **cbf**, l'algorithme de déraffinement II.4 est effectué :

- activation de la fonction de base **cbf**,
- désactivation de tous ses enfants **rbf** sans parent raffiné.

De plus, la fonction de base **cbf** n'est maintenant plus raffinée au sens de la définition II.1 (cf remarque II.5). Ceci est gardé en mémoire à l'aide du booléen prévu à cet effet : **cbf->set_unrefined()**. Cet algorithme prend également en compte la gestion des noeuds actifs de chaque champ (non exclu).

La classe **PDE_Activator** permet également de gérer l'activation et la désactivation (cf définition II.19) des cellules. Le listing III.11 détaille la procédure qui en est responsable. Cette procédure recherche les cellules qui

sont devenues actives après l'étape de déraffinement. Elle effectue un parcours niveau par niveau des cellules $K^{[j]}$ candidates à l'activation (ce sont celles qui constituent le support des fonctions de base de la liste BFS_U qui ont été déraffinées, seules cellules sur lesquelles des degrés de liberté actifs ont été supprimés). La cellule parent $P(K^{[j]})$ de la cellule $K^{[j]}$ ne vérifie nécessairement pas la propriété $A_{P(K^{[j]})}$ (cf section II.1.5) puisque nous avons activé une fonction de base de la cellule $K^{[j]}$ (au cours du déraffinement), et cette cellule est entièrement incluse dans $P(K^{[j]})$. Il ne reste donc qu'à déterminer si la cellule $K^{[j]}$ vérifie la propriété $A_{K^{[j]}}$. C'est la procédure `cell->all_basis_function_are_dropped()` (cf listing III.12) qui effectue cette vérification. Il faut parcourir récursivement toutes les cellules enfants de $K^{[j]}$ et vérifier qu'elles ne contiennent aucun noeud actif.

Listing III.11

Activation/Désactivation des cellules
après une étape de déraffinement

```

1 //Stockage des cellules candidates au deraffinement
2 //de niveau ll dans la liste cell_u[ll]
3 //Parcours de la liste BFS_U : bf
4 BFS_U->start_items_iterator() ;
5 for( ; BFS_U->item_is_valid() ; BFS_U->next_item() )
6 {
7     bf = BFS_U->current_item() ;
8     //parcours des cellules du support de bf : cell
9     bf->start_cells_iterator() ;
10    for( ; bf->cell_is_valid() ; bf->next_cell )
11    {
12        cell = bf->current_cell() ;
13        ll = cell->refinement_level() ;
14        //Ajout de cell
15        cell_u[ll].extend( cell ) ;
16    }
17 }
18
19 //Parcours de la liste cell_u
20 //niveau ll par niveau : cell
21 for( ll = 0 ; ll < cell_u.size() ; ll++ )
22 {
23     cell_u[ll].start_items_iterator() ;
24     for( ; cell_u[ll].item_is_valid()
25         ; cell_u[ll].next_item() )
26     {
27         cell = cell_u[ll].current_item() ;
28         //faut-il effectuer la (des)activation ?
29         do_it = true ;
30         //parcours des enfants de cell : rcell
31         cell->start_children_iterator() ;
32         for( ; cell->child_is_valid()
33             ; cell->next_child() )
34         {
35             rcell = cell->current_child() ;
36             //non s'il reste un noeud actif sur l'une
37             //des cellules enfants
38             do_it = do_it &&
39                 rcell->all_basis_function_are_dropped() ;
40         }
41         if( do_it )
42         {
43             //activation de cell
44             cell->set_active() ;
45             //desactivation de ses enfants
46             cell->start_children_iterator() ;
47             for( ; cell->child_is_valid()
48                 ; cell->next_child() )
49             {
50                 rcell = cell->current_child() ;
51                 rcell->set_inactive() ;
52             }
53         }
54     }
55 }
```

Listing III.12

`cell->all_basis_function_are_dropped()`

```

1 result = true ;
2
3 //verification recursive sur toutes les cellules
4 //enfants
5 cell->start_children_iterator() ;
6 for( ; cell->child_is_valid() ; cell->next_child() )
7 {
8     rcell = cell->current_child() ;
9     result = rcell->all_basis_functions_are_dropped() ;
10    if( !result ) break ;
11 }
12
13 //parcours des elements de reference de cell : elr
14 cell->start_elts_iterator() ;
15 for( ; cell->elt_is_valid() ; cell->next_elt() )
16 {
17     elt = cell->current_elt() ;
18     // Parcours des fonctions de base : bf
19     // de la cellule cell
20     for( ln = 0 ; ln < elt->nb_nodes() ; ln++ )
21     {
22         //
23         bf = cell->bf( elt, ln ) ;
24         if( bf != 0 )
25         {
26             //parcours des champs
27             bf->start_field_iterator() ;
28             for( ; bf->field_is_valid()
29                 ; bf->next_field() ) ;
30             {
31                 ff = bf->current_field() ;
32                 node = bf->node_of_DOF() ;
33                 result = !ff->node_is_active() ;
34             }
35         }
36         if( !result ) break ;
37     }
38 }
39 return( result ) ;
```

Adaptation : la classe `PDE_AdapterCHARMS`

La classe `PEL_AdapterCHARMS` implémente l'algorithme final adaptation. Cet algorithme reprend les briques de base présentées dans les paragraphes précédents. Les étapes de création des objets (étapes 2.ii. et 2.iii. de l'algorithme ci-dessous) ne sont pas décrites (puisqu'elles n'ont aucune portée générale).

Algorithme III.5 (Algorithme d'adaptation)

- | |
|--|
| <p>Etape 1 : Application des critères de raffinement et déraffinement :</p> <ul style="list-style-type: none"> i. Application du critère de raffinement : création de la liste <code>BFS_CRIT_R</code> (<i>cf</i> listing III.3), ii. Application du critère de déraffinement : création de la liste <code>BFS_CRIT_U</code> (<i>cf</i> listing III.4), <p>Etape 2 : Procédure de raffinement :</p> <ul style="list-style-type: none"> i. Sélection des fonctions de base à raffiner : création de la liste <code>BFS_R</code> (<i>cf</i> listing III.5), ii. Raffinement des cellules du support des fonctions de base de <code>BFS_R</code> : <ul style="list-style-type: none"> - désactivation de ces cellules, - activation de leurs cellules enfants (création si elles n'existent pas), iii. Création des noeuds et fonctions de base fines, iv. Raffinement de la liste <code>BFS_R</code> (<i>cf</i> listing III.9), <p>Etape 3 : Procédure de déraffinement :</p> <ul style="list-style-type: none"> i. Sélection des fonctions de base à déraffiner : création de la liste <code>BFS_U</code> (<i>cf</i> listing III.6), ii. Déraffinement de la liste <code>BFS_R</code> (<i>cf</i> listing III.10), iii. Déraffinement des cellules (<i>cf</i> listing III.11). |
|--|

Cet algorithme clos la présentation du module de raffinement local, nous exposons maintenant, en préliminaire à la description du module de préconditionnement multiniveau, le principe de numérotation des inconnues et les différentes options offertes par la librairie PELICANS.

III.2 Numérotation des inconnues dans la librairie PELICANS

Les inconnues dans la librairie PELICANS sont désignées par trois indices : un numéro désignant un champ discret (inconnu), puis deux numéros de noeud et de composante désignant un degré de liberté de ce champ. Il faut donc ordonner ces différents triplets pour construire *in fine* l'indexation des lignes et des colonnes de la matrice sous-jacente au problème à résoudre. Cette numérotation est effectuée en deux étapes :

- pour un champ fixé, la classe `PDE_LinkDOF2Unknown` permet de numéroté l'ensemble de ses degrés de liberté (associés par définition aux couples formés d'un noeud et d'une composante),
- la classe `PDE_SystemNumbering` permet ensuite de gérer la numérotation des degrés de liberté associés à différents champs.

Ces classes permettent de prendre en compte différentes possibilités de numérotation (*cf* figure III.1).

III.2.1 Numérotation des degrés de liberté d'un champ : la classe `PDE_LinkDOF2Unknown`

La classe `PDE_LinkDOF2Unknown` est dédiée à la numérotation des degrés de liberté associés à un champ discret donné (objet de type `PDE_DiscreteField`). Outre la gestion des degrés de liberté imposés (par exemple pour la prise en compte de conditions aux bords de type Dirichlet), cette classe permet d'écarter les noeuds non actifs du champ associé et offre deux options possibles pour l'organisation des degrés de liberté "inconnus" (*cf* figure III.1) :

- `"sequence_of_the_components"` : ordre lexicographique sur les couples (noeud, composante).
- `"sequence_of_the_nodes"` : ordre lexicographique sur les couples (composante, noeud).

Cette classe offre également la possibilité de considérer, en lieu et place des noeuds actifs, un sous-ensemble de noeuds prescrits par le client dans une liste de booléen `observed_nodes`. Cette fonctionnalité sera largement utilisée dans les algorithmes de coarsening présentés dans la section suivante.

Ainsi, la classe `PDE_LinkDOF2Unknown`, à son initialisation, effectue un parcours de noeuds et construit la numérotation en excluant : les noeuds inactifs (ou plus généralement ceux n'étant pas fournis dans une liste `observed_nodes`) et les degrés de liberté imposés.

Les fonctionnalités qu'un objet `link` de type `PDE_LinkDOF2Unknown` met à disposition sont :

- `link->field()` : champ associé (objet du type `PDE_DiscreteField`),
- `link->reset()` : ré-initialisation de la numérotation des degrés de liberté (exclusion des degrés de liberté imposés et des noeuds inactifs),

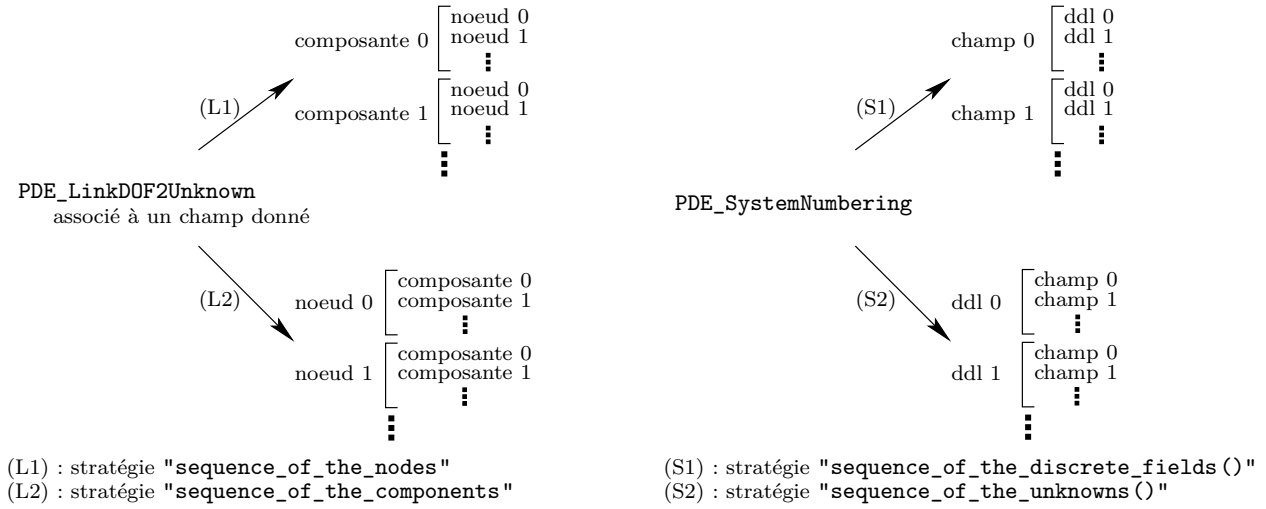


FIG. III.1 – Numérotation dans la librairie PELICANS

- `link->reset( observed_nodes )` : ré-initialisation de la numérotation des degrés de liberté (exclusion des degrés de liberté imposés et des noeuds `node` tels que `observed_nodes( node )=false`),
- `link->DOF_is_unknown( node, ic )` : booléen indiquant si le degré de liberté associé au couple (`node`, `ic`) est exclu de la numérotation (faux) ou non (vrai),
- `link->unknown_linked_to_DOF( node, ic )` : numéro associé au degré de liberté (`node`, `ic`),
- `link->unknown_vector_size()` : nombre total de degrés de liberté numérotés.

III.2.2 Gestion de plusieurs champs : la classe PDE_SystemNumbering

La classe `PDE_SystemNumbering` permet de gérer l'organisation des inconnues pour un système complet faisant intervenir plusieurs champs. Cette classe est donc associée à plusieurs objets de type `PDE_LinkDOF2Unknown`.

De nouveau, cette classe offre deux options de numérotation (cf figure III.1) :

- `"sequence_of_the_discrete_fields"` : ordre lexicographique sur les couples (champ, dof),
- `"sequence_of_the_unknowns"` ordre lexicographique sur les couples (dof, champ), cette dernière option n'est utilisable que lorsque les différents champs ont le même nombre de degrés de liberté.

Ces deux options peuvent être combinées avec toutes celles offertes par la classe `PDE_LinkDOF2Unknown` qui numérote les degrés de liberté.

Les fonctionnalités qu'un objet `nmb` de type `PDE_SystemNumbering` met à disposition sont :

- `nmb->nb_links()` : entier indiquant le nombre d'objet `PDE_LinkDOF2Unknown` associé à `nmb`,
- `nmb->reset()` : ré-initialisation des numérotations, effectue également la réinitialisation de tous les objets `PDE_LinkDOF2Unknown` associés (*i.e.* `nmb->link( i )->reset()`, pour tout `i`),
- `nmb->reset( observed_nodes_vv )` : ré-initialisation des numérotations, effectue également la réinitialisation de tous les objets `PDE_LinkDOF2Unknown` associés en tenant compte uniquement des noeuds indiqués par `observed_nodes_vv`, (*i.e.* `nmb->link( i )->reset( observed_nodes_vv( i )`), pour tout `i`),
- `nmb->nb_global_unknowns()` : nombre total d'inconnues,
- `nmb->global_unknown_for_DOF( node, ic, i )` : numéro de l'inconnue associée au degré de liberté défini par le couple (`node`, `ic`) du champ `nmb->link( i )->field()`.

Les fonctionnalités offertes par cette classe sont très utilisées par le module de préconditionnement multigrille puisqu'elles permettent de reconstruire automatiquement les numérotations des lignes et colonnes des matrices de passage entre les différents espaces d'approximation créés par l'algorithme de "coarsening" (cf algorithme II.25). Ceci est expliqué en détail dans la section III.3.

III.3 Organisation du module de préconditionnement multiniveau

Nous décrivons, dans cette section, l'organisation du module implémentant les méthodes de préconditionnement multiniveau. Rappelons que ces algorithmes reposent sur la définition d'une hiérarchie auxiliaire d'espaces d'approximation emboîtés, le plus grand étant celui sur lequel le problème à préconditionner est posé.

Nous avons présenté dans la section II.2 un algorithme (*cf* algorithme II.25) permettant de reconstruire, à partir de n'importe quelle base multiniveau, une suite de bases multiniveaux engendrant des espaces d'approximation multiniveaux emboîtés (décroissants) et de plus nous avons montré que les matrices des injections canoniques entre deux espaces successifs pouvaient facilement être déduites de la relation parents-enfants (*cf* remarque II.29).

Ainsi, il est possible d'appliquer l'algorithme de coarsening II.25 à chacune des bases multiniveaux courantes, la gestion de différents champs et de leurs multiples composantes (toutes discrétisées dans l'une des bases multiniveaux courantes) pouvant être déléguée à la classe `PDE_SystemNumbering`.

La section III.3.1 détaille l'algorithme de coarsening ainsi que celui qui permet de construire les matrices de passage (point central de la méthode). La classe abstraite `PDE_GeometricMultilevel_PC` implémentant les préconditionneurs multiniveaux est présentée dans la section III.3.2. Elle rassemble les différentes composantes des méthodes multigrilles (*cf* section II.2.5) à savoir :

- opérateurs de transfert (donnés par la classe `PDE_AlgebraicCoarsener`),
- opérateurs approchés (déduits de l'opérateur à préconditionner et des opérateurs de transfert par la formule II.11),
- lisseurs (à ce jour, seule la méthode de Gauss-Seidel est disponible, elle est implémentée dans la classe `LA_Smoothen`),
- solveur pour le problème approché dans l'espace d'approximation le plus grossier (classe abstraite `LA_Solver`, tous les solveurs de PELICANS sont à disposition ainsi que de nombreux couplages avec des bibliothèques extérieures).

III.3.1 Algorithme de coarsening : la classe `PDE_AlgebraicCoarsener`

La classe `PDE_AlgebraicCoarsener` s'organise autour de trois méthodes principales :

- la méthode `prepare_for_coarsening( nmb )` prend en argument un objet `nmb` de type `PDE_SystemNumbering` qui décrit l'ordre dans lequel sont organisées les inconnues du problème à préconditionner. Elle permet une initialisation de la structure interne de la classe (détaillée ci-dessous). En particulier, elle met en place une correspondance entre fonctions de base et noeuds d'un champ donné (ce pour tous les champs). Cette correspondance est stockée dans la structure `BFS` ; ainsi `BFS( ff, node )` est la fonction de base associée au noeud `node` du champ `ff`. Cette correspondance permet ensuite de travailler sur les noeuds plutôt que sur les fonctions de base en utilisant notamment les fonctionnalités offertes par l'objet `PDE_SystemNumbering`,
- la méthode `do_one_coarsening()` fait évoluer la structure interne de la classe pour effectuer un coarsening. Cette structure interne décrite ci-dessous contient en permanence les informations relatives à deux espaces d'approximation successifs (dans la hiérarchie que la classe est en train de construire). Dans la suite nous utiliserons la terminologie “niveau fin” et “niveau grossier” pour les désigner. La méthode `do_one_coarsening()` commence par écraser les structures concernant le niveau fin par celle du niveau grossier. Ainsi, l'ancien niveau fin disparaît et l'ancien niveau grossier devient le niveau fin. La suite de la méthode est consacrée à la construction du nouveau niveau grossier. Il est important de noter que cette méthode modifie uniquement la structure interne en particulier aucune fonction de base n'est activée ou désactivée pour effectuer le coarsening.
- la méthode `build_current_prolongation_matrix( mat )` permet de modifier la matrice `mat` passée en argument pour qu'elle devienne la matrice de passage du niveau grossier vers le niveau fin.

La structure interne de la classe `PDE_AlgebraicCoarsener` s'organise autour de six objets :

- `LEVEL_MAX` : entier désignant le plus grand des niveaux des fonctions de base actives (valeur initiale de l'entier j_M de la proposition II.26),
- `FINE_LEVEL` : entier désignant le niveau fin courant, cet entier est décrémenté à chaque coarsening (entier j_M de la proposition II.26),
- `NN_FINE[ff]` : noeuds du champ `ff` constituant le niveau fin courant (correspondant aux lignes de la matrice de passage),
- `NN_COAR[ff]` : noeuds du champ `ff` constituant le niveau grossier (correspondant aux colonnes de la matrice de passage),
- `ROW_NMB` : numérotation des inconnues pour les lignes de la matrice de passage,
- `COL_NMB` : numérotation des inconnues pour les colonnes de la matrice de passage.

Pour un champ `ff` donné, la correspondance noeud-fonction de base `BFS` permet en appliquant l'algorithme de coarsening II.25 de construire `NN_COAR[ff]` à partir de `NN_FINE[ff]`. L'objet `COL_NMB` est ensuite réinitialisé en utilisant les fonctionnalités de la classe `PDE_SystemNumbering` qui permettent de numérotter uniquement les degrés de liberté associés à `NN_COAR`. Une fois cette numérotation disponible, il est très facile de construire la

matrice de passage à l'aide de la relation parents-enfants.

Le listing III.13 présente l'initialisation de la structure interne (méthode `prepare_for_coarsening( nmb )`). Le listing III.14 présente la réalisation d'un coarsening (méthode `do_one_coarsening()`) et enfin, le listing III.15 présente la construction de la matrice de passage (méthode `build_current_prolongation_matrix( mat )`).

Listing III.13

Initialisation de la structure interne
`prepare_for_coarsening( nmb )`

```

1  /**Initialisation de ROW_NMB et COL_NMB
2  ROW_NMB = nmb
3  COL_NMB = nmb
4
5  /**Initialisation de BFS et NN_FINE
6  //Pour l'instant initialisation
7  //de toutes les valeurs a faux
8  NN_FINE.set( false );
9  //Parcours de toutes les cellules : cell
10 //existantes (de tous les niveaux)
11 GRID->start_cells_iterator();
12 for( ; GRID->cell_is_valid(); GRID->next_cell() )
13 {
14     cell = GRID->current_cell();
15     //Parcours de tous les champs associes a nmb
16     for( i = 0 ; i < nmb->nb_links(); i++ )
17     {
18         ff = nmb->link( i )->field();
19         //Element de reference associe au champs ff
20         //sur la cellule cell
21         elt = cell->reference_element( ff );
22         //Parcours des noeuds associes a l'element
23         //de reference elt
24         for( ln = 0 ; ln < elt->nb_nodes(); ln++ )
25         {
26             //fonction de base associe au
27             //ln-ieme noeud de l'element de
28             //reference elt sur la cellule cell
29             bf = cell->bf( elt, ln );
30             node = bf->node_of_DOF( ff );
31             //Association de la fonction
32             //de base bf au noeud node
33             BFS( ff, node ) = bf;
34             if( ff->node_is_active( node ) )
35             {
36                 NN_FINE( node ) = true;
37             }
38         }
39     }

```

Listing III.14
Réalisation d'un coarsening
`do_one_coarsening()`

```

1  //L'ancien niveau fin devient
2  //le nouveau niveau grossier
3  //Copie de COL_NMB dans ROW_NMB
4  ROW_NMB = COL_NMB;
5  //Copie de NN_COAR dans NN_FINE
6  NN_FINE = NN_COAR;
7  NN_COAR.set( false );
8
9  //Parcours de tous les champs
10 for( i = 0 ; i < COL_NMB->nb_links(); i++ )
11 {
12     ff = COL_NMB->link( i )->field();
13     //Parcours de tous les noeuds du champ ff
14     ff->start_nodes_iterator();
15     for( ; ff->node_is_valid(); ff->next_node() )
16     {
17         node = ff->current_node();
18         //Recuperation de la fonction de base associee
19         //au noeud node du champ ff
20         bf = BFS( ff, node );
21         //Si node est un noeud du niveau fin
22         if( NN_FINE( node ) )
23         {
24             //Si la fonction de base associe a node est
25             //du niveau le plus fin
26             if( bf->refinement_level() == FINE_LEVEL )
27             {
28                 //son parent est ajoute
29                 pbf = bf->leading_parent();
30                 if( pbf != 0 )
31                 {
32                     pnode = pbf->node_of_DOF( ff );
33                     NN_COAR( pnode ) = true;
34                 }
35             }
36             else
37             {
38                 //sinon, le noeud est conservee
39                 //au niveau grossier
40                 NN_COAR( node ) = true;
41             }
42         }
43     }
44 }
45 COL_NMB->reset( NN_COAR )

```

Listing III.15

Construction de la matrice de passage
`build_current_prolongation_matrix( mat )`

```

1  //Taille de la matrice
2  mat->re_initialize( ROW_NMB->nb_global_unknowns(),
3                    COL_NMB->nb_global_unknowns() );
4
5  for( i = 0 ; i < ROW_NMB->nb_links(); i++ )
6  {
7      rlink = ROW_NMB->link( i );
8      clink = COL_NMB->link( i );
9
10     ff = rlink->field();
11     ASSERT( ff == clink->field() );

```

```

12
13 //Parcours des noeuds associes au champ ff
14 ff->start_nodes_iterator() ;
15 for( ; ff->node_is_valid() ; ff->next_node() )
16 {
17     node = ff->current_node() ;
18     //Parcours des composantes du champ ff
19     for( ic = 0 ; ic < ff->nb_components() ; ic++ )
20     {
21         //Pour les inconnues grossiere
22         if( clink->DOF_is_unknown( node, ic ) )
23         {
24             i_col =
25                 COL_NMB->global_unknown_for_DOF( n, ic, clink ) ;
26             //si egalement inconnue fine : on met 1.
27             if( rlink->DOF_is_unknown( n, ic ) )
28             {
29                 i_row =
30                     ROW_NMB->global_unknown_for_DOF( n, ic, rlink )
31 ;
32                 mat->set_item( i_row, i_col, 1.0 ) ;
33             }
34             else
35             {
36                 bf = BFS( node, ff ) ;
37                 //Parcours des enfants de bf ;
38                 FMS->start_children_iterator( bf ) ;
39                 for( ; FMS->child_is_valid() ; FMS->next_child() )
40                 {
41                     cbf = FMS->current_child() ;
42                     xx = bf->current_refinement_coefficient() ;
43
44                     cnode = cbf->node_of_DOF( ff ) ;
45
46                     //on a necessairement (assertion)
47                     rlink->DOF_is_unknown( cnode, ic ) ;
48
49                     i_row =
50                         ROW_NMB->global_unknown_for_DOF( cnode,
51                                                             ic, rlink ) ;
52                     mat->set_item( i_row, i_col, xx ) ;
53                 }
54             }
55         }
56     }
57 }
58 }

```

III.3.2 Préconditionneurs multiniveaux : la classe PDE_GeometricMultilevel_PC

La classe permet de regrouper les différentes fonctionnalités nécessaires à la mise en oeuvre concrète de preconditionneurs multigrilles. Elle permet de récupérer l'ensemble des matrices de transfert fournies par la classe PDE_AlgebraicCoarsener, elle effectue le calcul des opérateurs approchés et fournit un lisseur par niveau (algorithme de Gauss-Seidel).

EXEMPLES DE CLASSES CONCRÈTES DÉRIVÉES de la classe PDE_GeometricMultilevel_PC :

- PDE_MG_PC (preconditionneur multigrille, version multiplicative, cf expression (II.13)),
- PDE_BPX_PC (preconditionneur multigrille, version additive, cf expression (II.12))...

III.4 Principe du parallélisme dans la librairie PELICANS

L'objectif que nous visons à travers les techniques de calcul parallèle est de permettre une exécution du code sur des systèmes à mémoire distribuée. Le modèle de programmation adopté dans la librairie PELICANS est un modèle par échange de message SPMD (Single Program Multiple Data). Cela signifie que le même programme est exécuté par plusieurs processus (pouvant eux-mêmes être exécutés sur différents processeurs physiques) pouvant échanger entre eux des données pour mener à bien la résolution du problème visé. Les échanges nécessaires sont

organisés grâce à l'utilisation de bibliothèques de communication par passage de messages (MPI) (*cf* section III.4.1).

Le domaine de calcul est partitionné en plusieurs sous-domaines, chacun d'entre eux étant affecté à un processus (*cf* section III.4.2). Chaque processus ne gère alors que les données relatives à la partie qui lui est associée. Certains objets, à la frontière entre deux sous-domaines, sont créés sur plusieurs processus. Une numérotation globale des inconnues, *i.e.* tenant compte de l'ensemble des sous-domaines, est alors constituée (*cf* section III.4.3). Ceci permet l'assemblage des matrices sous-jacentes au problème complet. Chaque processus se voit alors affecter d'une partie des inconnues : un processus ne stocke que certaines lignes des matrices à inverser (celles correspondant aux inconnues qui lui sont affectées, *cf* section III.4.4). Aucun processus n'est capable de reconstituer le système complet, celui-ci n'existe qu'à travers la numérotation globale, pour cette raison on parlera quelquefois de système linéaire logique.

III.4.1 La bibliothèque MPI

Il n'est pas question de décrire dans ce manuscrit tous les mécanismes de la bibliothèque MPI, cependant pour comprendre le fonctionnement de la librairie PELICANS en parallèle, il est nécessaire d'en expliquer les grands principes :

- toutes les variables du programme sont privées et résident dans la mémoire locale allouée à chaque processus,
- une donnée est échangée entre deux ou plusieurs processus via un appel, dans le programme, à des sous programmes particuliers.

La bibliothèque MPI fournit les fonctionnalités nécessaires pour mettre en oeuvre ces échanges (ou communications) par l'intermédiaire de communicateurs (ensemble de processus sur lesquels portent les opérations effectuées). Nous n'utiliserons que le communicateur par défaut, noté `COM` dans la suite, qui comprend l'ensemble de tous les processus. Deux grandes familles de communications se distinguent :

- les communications point à point : elles ont lieu entre deux processus : émetteur et récepteur. L'identification des processus récepteur et émetteur se fait par un entier, appelé rang, qui est attribué à chaque processus par la bibliothèque MPI lors de la phase d'initialisation. L'envoi (resp. la réception) d'une donnée `data` se fait par l'appel à une fonction du communicateur que nous noterons `send( rank, data )` (resp. `receive( rank, data )`), `rank` désignant le rang du processus récepteur (resp. émetteur). Notons que tout message envoyé doit être reçu (*i.e.* l'émetteur doit exécuter une commande `send` et le récepteur une commande `receive`) pour que le transfert soit effectif.
- les communications collectives : celles-ci concernent l'ensemble des processus (du communicateur), elles permettent de faire en une seule opération une série de communications point à point. Il en existe de nombreuses mais nous ne les détaillerons pas dans ce manuscrit.

Les figures III.2, III.3 et le listing III.16 illustrent un schéma classique de communication point à point entre les différents processus d'un communicateur. La figure III.2 montre comment les communications s'organisent dans le cas de cinq processus. Le processus de rang 0 envoie des données `data` au processus de rang 1 qui, après les avoir éventuellement modifiées par une procédure `algorithm`, les envoie lui-même au processus de rang 2, et ainsi de suite jusqu'au dernier processus qui, après avoir effectué ses propres modifications, les renvoie à tous les autres processus. Le programme effectuant cette succession de communication est détaillé par le listing III.16, les différentes actions réalisées par chacun des processus étant reportées, étape par étape, sur la figure III.3 dans le cas de trois processus. Sur la figure III.3, les flèches en pointillés signifient que le processus pointé est en attente de réception de donnée provenant du processus à l'origine de la flèche.

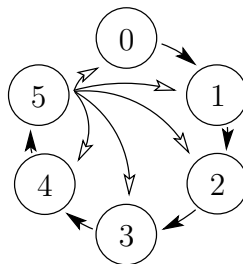


FIG. III.2 – Exemple de stratégie de communications point à point sur 5 processus

Listing III.16
Exemple de schéma de communications
Algorithme

```

1 nb_ranks = COM.nb_ranks() ;
2 rank = COM.rank() ;
3 last = nb_ranks-1 ;
4
5 if( rank > 0 )
6 {
7   //Chaque processus (hormis celui de rang 0) est
8   //mis en attente de reception des donnees : data
9   //provenant du processus dont le rang
10  //est immediatement inferieur
11  COM.receive( rank-1, data ) ; // (1)
12 }
13
14 //Algorithme modifiant eventuellement
15 //la valeur de data communiquee par
16 //le processeur precedent
17 algorithm( data ) ;
18
19 if( rank != last )
20 {
21   //Chaque processus (hormis celui de rang le plus
22   //grand) envoie les donnees : data qu'il vient
23   //de modifier au processus dont le rang est
24   //immediatement superieur et est mis en attente
25   //de reception des donnees : data modifiee par
26   //le processus de rang le plus grand
27   COM.send( rank+1, data ) ; // (1)
28   COM.receive( last, data ) ; // (2)
29 }
30 else
31 {
32   //Le processus de rang le plus grand envoie
33   //les donnees a tous les autres processus
34   for( i = 0 ; i < last ; i++ )
35   {
36     COM.send( i, data ) ; // (2)
37   }
38 }

```

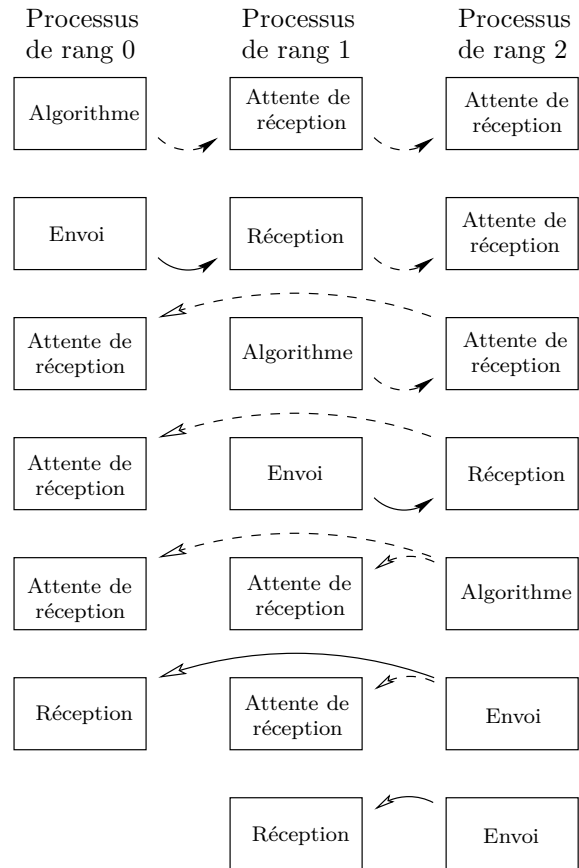


FIG. III.3 – Déroulement de l'algorithme sur trois processus étape par étape

Ce type de schéma de communication simple permet d'organiser les communications nécessaires à l'assemblage (cf section III.4.3) et à la résolution (cf section III.4.4) du système linéaire logique.

III.4.2 Partitionnement du domaine

Dans la librairie PELICANS, la répartition des tâches à effectuer entre les différents processus se fait par l'intermédiaire d'une décomposition du maillage $\mathcal{T}_0$ du domaine de calcul (*cf* figure III.4). Chacun des sous-maillages obtenus est affecté à un et un seul processus. Nous supposons que les sous-maillages sont sans recouvrement (la librairie PELICANS autorise l'utilisation de zones de recouvrement pour une gestion d'assemblages plus complexes (assemblages de type volumes finis, termes de saut...) mais celle-ci n'est pas compatible à ce jour avec le module de raffinement local que nous décrivons ici). Le partitionnement peut être réalisé par l'utilisateur (en prescrivant, pour chaque centre de maille du domaine, le rang du processus auquel la maille est affectée) ou peut être délégué à un logiciel extérieur (un couplage avec le logiciel METIS est fourni).

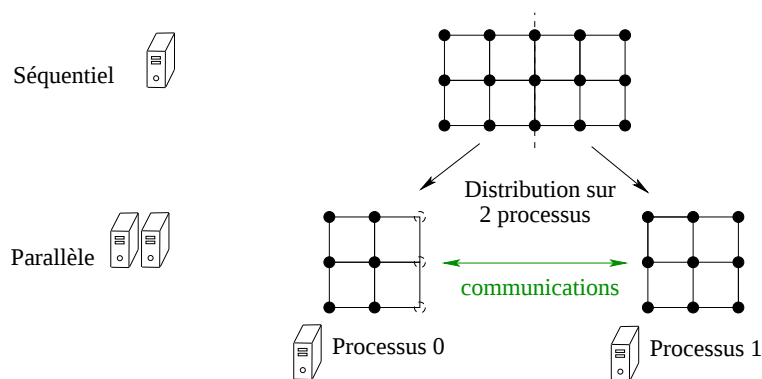


FIG. III.4 – Techniques de calcul parallèle – Partitionnement du domaine

Chaque processus gère le sous-maillage qui lui est affecté de manière similaire au mode de fonctionnement séquentiel à l'exception de la classe `PDE_SystemNumbering` qui met maintenant à disposition une correspondance entre les numéros d'inconnues (degrés de liberté de tous les champs) locaux à chaque processus et les numéros d'inconnues pour le système logique global, *i.e.* tenant compte de l'ensemble des sous-maillages. Une fois cette correspondance établie, la phase d'assemblage est effectuée en parallèle par tous les processus. Elle consiste, sur chaque processus, en un parcours de l'ensemble des mailles (rappelons qu'il n'y a pas de recouvrement entre les sous-maillages) et en l'ajout des contributions de chacune des mailles aux lignes et colonnes du système global logique fournies par la classe `PDE_SystemNumbering`.

L'obtention d'une numérotation pour le système global nécessite la mise en place de communications entre les différents processus. Cette responsabilité est déléguée aux classes `PDE_CrossProcessNodeNumbering` et `PDE_CrossProcessUnknownsNumbering` dont nous expliquons le fonctionnement dans la section suivante.

Remarque III.6

L'ensemble des objets utilisés en séquentiel (`PDE_MeshFE`, `PDE_BasisFunctionCell`, ...) est également créé en parallèle. Cependant chaque processus ne crée que les objets relatifs au sous-domaine qui lui est affecté. Notons que les objets à la frontière (par exemple certaines faces ou fonctions de base) entre deux sous-domaines sont créés par les deux processus auxquels sont affectés les sous-domaines de part et d'autre de la frontière. Cependant, ces objets sont affectés à un seul des deux processus par le jeu de la numérotation globale.

III.4.3 Numérotation globale des inconnues en parallèle

Dans cette section, nous expliquons comment est construite la numérotation globale des inconnues. Celle-ci est organisée processus par processus. Sur chacun des processus les différentes options de numérotation (*cf* section III.2) sont mises à disposition. Une première numérotation globale des noeuds associés à un champ est mise à disposition par la classe `PDE_CrossProcessNodeNumbering`. Celle-ci est ensuite utilisée pour la construction de la numérotation globale des degrés de liberté d'un champ par la classe `PDE_CrossProcessUnknownsNumbering`. Enfin, en utilisant ces deux numérotations globales, la classe `PDE_SystemNumbering` permet de reconstruire une numérotation globale de l'ensemble des inconnues.

Numérotation globales des noeuds : la classe `PDE_CrossProcessNodeNumbering`

La classe `PDE_CrossProcessNodeNumbering` est attachée à un champ discret `ff` (objet de type `PDE_DiscreteField`), elle permet de reconstruire une numérotation globale des noeuds de ce champ. Les informations que met à disposition un objet `cpn` de type `PDE_CrossProcessNodeNumbering` sont :

- `cpn->nb_global_nodes()` : nombre total de noeuds globaux associés au champ `ff`,
- `cpn->global_node_index( node )` : numéro global du noeud de numéro local `node` sur le processus,
- `cpn->local_node( glob_node )` : numéro local du noeud sur le processus de numéro global `glob_node`,
- `current_process_handles_node( node )` : booléen indiquant si le processus possède le noeud local `node`,
- `rank_of_process_handling( node )` : rang du processus possédant le noeud local `node`.

L'obtention de la numérotation globale des noeuds est réalisée de la manière suivante : tous les processus (dans l'ordre de leur rang) envoient, au processus de rang 0, les coordonnées géométriques associées aux noeuds du champ `ff`. Le processus 0 supprime les doublons, trie les noeuds suivant l'ordre lexicographique de leurs coordonnées en gardant les correspondances d'indices avec les listes envoyées. Le processus 0 communique ensuite cette correspondance à tous les autres processus qui en déduisent la numérotation globale. L'ordre de numérotation est indépendant de l'ordre de parcours des noeuds et du découpage du domaine. En outre, les noeuds du champ `ff` associés à un noeud géométrique placé à la frontière entre deux sous-domaines est affecté au processus de rang le plus petit. Dans la suite, nous utiliserons la terminologie “vue” pour un objet ou concept existant sur un processus et “possédé” pour un objet ou concept affecté au processus.

Numérotation des degrés de liberté `PDE_CrossProcessUnknownsNumbering`

La classe `PDE_CrossProcessUnknownsNumbering` est attachée à un objet de type `PDE_LinkDOF2Unknown` (*cf* section III.2), elle permet de reconstruire une numérotation globale des inconnues associées à l'objet `PDE_LinkDOF2Unknown`. En plus, des fonctionnalités apportées par l'objet `PDE_LinkDOF2Unknown` (gestion des noeuds actifs et degrés de liberté imposés), la classe `PDE_CrossProcessUnknownsNumbering` écarte les noeuds qui ne sont pas possédés. Les informations que met à disposition un objet `cpu` de type `PDE_CrossProcessUnknownsNumbering` sont :

- `cpu->nb_global_unknowns()` : nombre total d'inconnues associées au champ `ff`,
- `cpu->nb_unknowns_on_process( rank )` : nombre d'inconnues (associées au champ `ff`) possédées par le processus de rang `rank`,
- `cpu->global_unknown_index( i )` : numéro global de l'inconnue de numéro local au processus `i`,
- `cpu->rank_of_process_handling( i )` : rang du processus possédant l'inconnue de numéro local `i`.

Son fonctionnement reprend la stratégie simple de communication expliquée dans la section III.4.1. Les données `data` transférées entre chaque processus se résument à un entier `NB_UNKNOWNS` désignant le nombre d'inconnues globales numérotées à cet instant du parcours et l'algorithme effectué sur chacun des processus consiste à parcourir les degrés de liberté (dans l'ordre requis par l'option choisie) en écartant les noeuds non possédés (information fournie par la classe `PDE_CrossProcessNodeNumbering`) et les degrés de liberté écartés par la classe `PDE_LinkDOF2Unknown`.

Ainsi, la numérotation des inconnues est effectuée processus par processus : les inconnues vues sur le processus de rang 0 ont les numéros les plus petits. . .

Cet objet complète le rôle de la classe `PDE_LinkDOF2Unknown`. C'est donc à travers cette classe que nous y avons accès (`link->cross_process_numbering()`).

Numérotation globale des inconnues : la classe PDE_SystemNumbering

Les deux classes présentées précédemment servent exclusivement à faciliter la mise en place de la numérotation globale pour le système complet. Cette responsabilité est assumée par la classe `PDE_SystemNumbering`. De plus, aucune communication supplémentaire n'est nécessaire.

Comme en séquentiel, deux stratégies sont prises en compte (*cf* figure III.1). La stratégie (S2) du séquentiel reste facile à implémenter puisque nous avons connaissance d'un numéro d'inconnue global : pour obtenir le numéro du degré de liberté associé au noeud `node` de la composante `ic` du champ `link( i )->field()` il suffit d'appliquer la formule `unk_glob * nb_links() + i` où `unk_glob = cpu->global_unknown_index( unk_loc )`, `unk_loc = link( i )->unknown_linked_to_DOF( node, ic )` et `cpu = link( i )->cross_process_numbering()`. La stratégie (S1) est un peu plus complexe à implémenter. Il faut prendre en compte le fait que la numérotation s'effectue processus par processus. Considérons un champ de numéro `e` et un degré de liberté de ce champ de numéro global `ddl_glob`. Ce degré de liberté appartient à un processus, notons `r` son rang. Pour obtenir le numéro global de l'inconnue associée au degré de liberté `ddl_glob` du champ `e` il faut tenir compte de tous les degrés de liberté des autres champs sur les processus de rang strictement inférieur à `r` et de tous les degrés de liberté des champs dont le numéro est strictement plus petit que `e` sur le processus de rang `r`. Ceci induit donc un décalage entre la numérotation globale des degrés de liberté d'un champ donné et celle des inconnues. Pour un champ `e` donné, ce décalage est identique pour tous les degrés de liberté associés à un processus de rang `r`, il est calculé et stocké dans la variable `SHIFT(e,r)`. Le listing III.17 et les explications à droite montrent les détails du calcul de cette valeur. Le numéro d'inconnue du degré de liberté associé au noeud `node` de la composante `ic` du champ `link( i )->field()` est ensuite obtenu en posant : `SHIFT( i, r ) + unk_glob` où `unk_glob = cpu->global_unknown_index( unk_loc )`, `r = cpu->rank_of_process_handling( unk_loc )`, `unk_loc = link( i )->unknown_linked_to_DOF( node, ic )` et `cpu = link( i )->cross_process_numbering()`.

Listing III.17
Détails du calcul de la numérotation
(option S1)

```

1 for( r = 0 ; r < COM->nb_ranks() ; ++r )
2 {
3     temp( 0, r ) = 0 ;
4
5     for( size_t e=0 ; e<nb_links() ; ++e )
6     {
7         cpu = link( e )->cross_process_numbering() ;
8         temp( e + 1, r ) = temp( e, r )
9             + cpu->nb_unknowns_on_process( r ) ;
10    }
11    for( e = 0 ; e < nb_links() ; ++e )
12    {
13        if( r == 0 )
14        {
15            SHIFT( e, 0 ) = temp( e, 0 ) ;
16        }
17        else
18        {
19            SHIFT( e, r ) =
20                SHIFT( e, r - 1 ) + temp( e, r )
21                + temp( nb_links(), r - 1 )
22                - temp( e + 1, r - 1 ) ;
23        }
24    }
25 }

```

$$\begin{aligned}
 \text{SHIFT}(e, r) &= \sum_{\substack{p < r \\ f \neq e}} \text{nb_inc}_f(p) + \sum_{f < e} \text{nb_inc}_f(r) \\
 &= \sum_{\substack{p \leq r \\ f < e}} \text{nb_inc}_f(p) + \sum_{\substack{p < r \\ f > e}} \text{nb_inc}_f(p) \\
 \text{SHIFT}(e, 0) &= \sum_{f < e} \text{nb_inc}_f(0) \\
 \text{SHIFT}(e, r) &= \text{SHIFT}(e, r - 1) \\
 &\quad + \sum_{f < e} \text{nb_inc}_f(r) + \sum_{f > e} \text{nb_inc}_f(r - 1) \\
 \text{temp}(e, r) &= \sum_{f < e} \text{nb_inc}_f(r) \\
 \text{SHIFT}(e, 0) &= \text{temp}(e, 0) \\
 \text{SHIFT}(e, r) &= \text{SHIFT}(e, r - 1) + \text{temp}(e, r) \\
 &\quad + \text{temp}(\text{nb_chp}, r - 1) - \text{temp}(e + 1, r - 1)
 \end{aligned}$$

La figure III.5 illustre les deux stratégies les plus utilisées ; ce sont les combinaisons des options :

- "sequence_of_the_discrete_fields" et "sequence_of_the_nodes" (à gauche sur la figure),
- "sequence_of_the_unknowns" et "sequence_of_the_components" (à droite sur la figure).

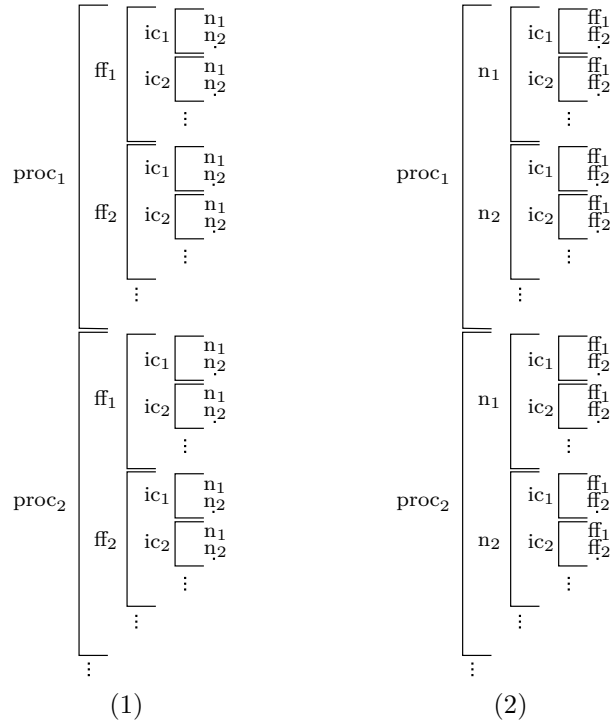


FIG. III.5 – Numérotation des inconnues en parallèle dans la librairie PELICANS

III.4.4 Algèbre linéaire distribuée

L'algèbre linéaire de la librairie PELICANS est également distribuée sur l'ensemble des processus : le stockage des matrices et vecteurs est réparti entre les différents processus et des solveurs linéaires itératifs préconditionnés parallèles sont fournis : la méthode gradient conjugué (classe `LA_CG_IS`), l'algorithme GMRES (classe `LA_GMRES_IS`), l'algorithme BICGSTAB (classe `LA_BiCGSTAB`).

Les matrices et vecteurs distribués

Chaque processus est responsable du stockage d'un bloc de lignes consécutives de chaque matrice (ou vecteur). De plus, ces blocs sont attribués aux processus de manière croissante : le processus de rang 0 possède le bloc constitué des n_0 premières lignes de la matrice, le processus de rang 1 possède le bloc constitué des lignes $n_0 + 1$ à $n_0 + n_1$... La taille des blocs n_0, n_1 ... est néanmoins *a priori* indépendante du partitionnement du domaine de calcul et de la numérotation globale des inconnues (*cf* section III.4.3).

Chaque processus stocke entièrement les lignes des matrices qui lui sont affectées et prévoit un espace mémoire tampon pour le stockage des éléments des autres lignes. Lorsqu'un élément d'une matrice est modifié (affectation ou addition), cette dernière est automatiquement placée dans un état dit de désynchronisation sur le processus où la modification a été effectuée. En effet, les autres processus ne peuvent avoir connaissance de la modification, en conséquence sa prise en compte n'est pas encore effective.

Deux états de désynchronisation sont possibles : `NotSync_add` et `NotSync_set`. Le premier `NotSync_add` correspond au cas où des opérations d'addition ont été effectuées sur les coefficients de la matrice, le second `NotSync_set` correspond à des opérations d'affectation de coefficients. Les deux états sont mutuellement exclusifs (*i.e.* aucune opération d'affectation ne sera permise sur une matrice en état de désynchronisation `NotSync_add` et réciproquement). En effet, au moment de la synchronisation, phase où les coefficients stockés dans les zones tampon sont transférés aux processus qui en sont responsables, il faut pouvoir décider du type d'opération à effectuer : remplacer les coefficients ou les additionner.

Ainsi, chaque processus peut mettre à jour n'importe quel élément de la matrice. Par contre, il ne peut accéder qu'à ceux des lignes dont il est responsable.

Ce type de stockage facilite les multiplications matrices-vecteurs puisqu'aucun transfert de coefficients de la matrice n'est nécessaire. Le produit matrice-vecteur nécessite néanmoins la connaissance du vecteur entier celui-ci étant rapatrié sur chaque processus en utilisant des objets de type `LA_Scatter`. Leur rôle est d'effectuer des changements d'indices entre un vecteur séquentiel et un vecteur de type abstrait (séquentiel ou parallèle) : cette classe réalise un simple changement d'indice lorsque les deux vecteurs sont séquentiels par contre lorsqu'ils sont parallèles elle gère entièrement toutes les communications nécessaires.

Les solveurs linéaires

Le parallélisme est essentiellement transparent pour la mise en oeuvre des solveurs linéaires puisque l'ingrédient principal sur lequel ils reposent est le produit matrice-vecteur.

Les préconditionneurs

Pour les préconditionneurs dont la généralisation n'est pas immédiate en parallèle (par exemple ILU0), la stratégie adoptée dans la librairie PELICANS est de les combiner avec une méthode de Jacobi par bloc.

Il faut garder en mémoire que les préconditionneurs utilisant cette stratégie ne sont pas les mêmes en séquentiel et en parallèle, ils diffèrent également lorsque le nombre de processus change.

Le préconditionneur ILU0 utilisant cette stratégie donne de bons résultats sur les problèmes que nous avons testés (système de Cahn-Hilliard, *cf* section IX.1.7). Pour les préconditionneurs nous avons adopté une autre stratégie qui consiste à modifier légèrement l'algorithme de Gauss-Seidel.

III.5 Raffinement local en parallèle

Les algorithmes de raffinement local peuvent être exécutés en parallèle sur chacun des processus. Nous devons cependant faire face à la difficulté où un objet serait créé unilatéralement à la frontière entre deux sous-domaines de la partition. Cette difficulté peut être aisément contournée en ajoutant une étape de synchronisation des objets à raffiner et à déraffiner entre les différents processus, garantissant ainsi que les objets aux frontières seront créés sur les processus de part et d'autre de la frontière.

Un objet BFN de type `PDE_CrossProcessBFNumbering` est chargé d'attribuer un identifiant global aux fonctions de base et de fournir la correspondance avec l'identifiant local (ou `id_number()`) :

- `BFN->global_bf_index( id_loc )` : identifiant global de la fonction de base dont l'identifiant local est `id_loc`,
- `BFN->local_bf_index( id_glob )` : identifiant local de la fonction de base dont l'identifiant global est `id_glob`.

L'identifiant global ainsi obtenu permet de communiquer les listes locales de fonctions de base à (dé)raffiner. La stratégie de communication utilisée est détaillée dans le listing III.18 (resp. III.19) pour la synchronisation des fonctions de base à raffiner (resp. déraffiner).

Une fonction de base est ajoutée dans la liste des fonctions de base à raffiner de tous les processus la voyant, dès que cette fonction de base est dans la liste d'un de ces processus.

Une fonction de base est supprimée de la liste des fonctions de base à déraffiner de chacun des processus la voyant, si un de ces processus ne la contient pas dans sa liste locale.

Listing III.18

Synchronisation des listes de fonctions de base à raffiner

```

1 //Creation d'une liste vide pour stocker les
2 //identifiants globaux des fonctions de base
3 //du processus courant
4 bfs_glob_idx( 0 ) ;
5
6 //Parcours de toutes les fonctions de base
7 //a raffiner sur le processus courant
8 BFS_R->start_items_iterator() ;
9 for( ; BFS_R->item_is_valid() ; BFS_R->next_item() )
10 {
11     bf = BFS_R->current_item() ;
12     //identifiant local de bf
13     id = bf->id_number() ;
14     //ajout de l'identifiant global de bf
15     //a la liste bfs_glob_idx
16     bfs_glob_idx.append( BFN->global_bf_index( id ) ) ;
17 }
18
19 //Pour tout r
20 for( r=0 ; r < COM->nb_ranks() ; ++r )
21 {
22     //si le processus courant est de rang r
23     //il copie sa liste bfs_global_idx
24     //dans la liste temporaire bfs_idx_dum
25     if( COM->rank() == r )
26     {
27         bfs_idx_dum = bfs_glob_idx ;
28     }
29     //le processus r envoie sa liste
30     //au processus courant
31     COM->broadcast( bfs_idx_dum, r ) ;
32
33     //le processus courant parcourt la liste envoyee
34     //par le processus r
35     for( i = 0 ; i < bfs_idx_dum.size() ; ++i )
36     {
37         //recuperation d'un identifiant local
38         loc_id = BFN->local_bf_index( bfs_idx_dum( i ) ) ;
39         //si la fonction de base est vue par
40         //le process courant
41         if( loc_id != PEL::bad_index() )
42         {
43             bf = BF_SET->item( loc_id ) ;
44             //on l'ajoute dans la liste du processus courant
45             if( !BFS_R->has( bf ) ) BFS_R->append( bf ) ;
46         }
47     }
48 }

```

Listing III.19

Synchronisation des listes de fonctions de base à déraffiner

```

1 //Creation d'une liste vide pour stocker les
2 //identifiants globaux des fonctions de base
3 //du processus courant
4 bfs_glob_idx( 0 ) ;
5
6 //Parcours de toutes les fonctions de base
7 //a deraffiner sur le processus courant
8 BFS_U->start_items_iterator() ;
9 for( ; BFS_U->item_is_valid() ; BFS_U->next_item() )
10 {
11     bf = BFS_U->current_item() ;
12     //identifiant local de bf
13     id = bf->id_number() ;
14     //ajout de l'identifiant global de bf
15     //a la liste bfs_glob_idx
16     bfs_glob_idx.append( BFN->global_bf_index( id ) ) ;
17 }
18
19 //Pour tout r
20 for( r=0 ; r < com->nb_ranks() ; ++r )
21 {
22     //si le processus courant est de rang r
23     //il copie sa liste bfs_global_idx
24     //dans la liste temporaire bfs_idx_dum
25     if( COM->rank() == r )
26     {
27         bfs_idx_dum = bfs_global_idx ;
28     }
29     //le processus r envoie sa liste
30     //au processus courant
31     COM->broadcast( bfs_idx_dum, r ) ;
32
33     //le processus courant parcourt la liste envoyee
34     //par le processus r
35     for( i = 0 ; i < bfs_idx_dum.size() ; ++i )
36     {
37         //recuperation d'un identifiant local
38         loc_id = BFN->local_bf_index( bfs_idx_dum( i ) ) ;
39         //Est-ce que la fonction de base est vue par
40         //le process courant ?
41         on_process = ( loc_id != PEL::bad_index() ) ;
42         //Est-ce que le processus courant possede
43         //cette fonction de base dans sa liste
44         //dans sa liste BFS_U ?
45         is_not_in_list = false ;
46         if( on_process )
47         {
48             bf = BF_SET->item( loc_id ) ;
49             is_not_in_list = !( BFS_U->has( bf ) ) ;
50         }
51         //si une des listes d'un des
52         //processus ne contient pas la fonction de base
53         to_remove = COM->boolean_or( is_not_in_list ) ;
54         if( to_remove && !is_not_in_list && on_process )
55         {
56             //on la supprime de la liste
57             //du processus courant
58             BFS_U->remove( bf ) ;
59         }
60     }
61 }

```

III.6 Multigrille en parallèle

L'algorithme de coarsening et la construction des matrices de passage présentés dans la section III.3 sont encore valides dans le cadre parallèle. Cependant, l'algorithme de Gauss-Seidel que nous utilisons comme lisseur ne s'étend pas naturellement au cas parallèle. En effet, cet algorithme est par définition séquentiel puisque la i^e étape nécessite la connaissance de la mise à jour des étapes précédentes.

Plaçons nous dans le cadre où A est une matrice symétrique définie positive, b le second membre. Une itération de l'algorithme de Gauss-Seidel peut s'exprimer ainsi

$$x_i \leftarrow x_i + \frac{1}{a_{ii}}(b - Ax)_i$$

ceci devant être exécuté une et une seule fois pour chacune des lignes i . Le calcul de la i^e composante du résidu $b - Ax$ au second membre tient compte des mises à jour du vecteur x effectuées par les étapes précédentes. Nous obtenons différentes variantes selon l'ordre de parcours des lignes que nous choisissons (l'algorithme standard parcourant les lignes dans l'ordre croissant).

Il est intéressant de remarquer que lorsque la répartition des lignes de la matrice correspond à celle des inconnues, la plupart des étapes ci-dessus peuvent se découpler. Notons $\text{first}(r)$ et $\text{last}(r)$ les première et dernière lignes de la matrice affectées au processus de rang r .

Nous définissons les ensembles suivants :

$$\text{Top}(r) = \{i \in \llbracket \text{first}(r), \text{last}(r) \rrbracket, \exists j < \text{first}(r), a_{ij} \neq 0 \text{ et } \forall j \geq \text{last}(r), a_{ij} = 0\},$$

$$\text{Bot}(r) = \{i \in \llbracket \text{first}(r), \text{last}(r) \rrbracket, \exists j \geq \text{last}(r), a_{ij} \neq 0 \text{ et } \forall j < \text{first}(r), a_{ij} = 0\},$$

$$\text{Int}(r) = \{i \in \llbracket \text{first}(r), \text{last}(r) \rrbracket, \forall j \text{ t.q. } j < \text{first}(r) \text{ ou } j \geq \text{last}(r), a_{ij} = 0\},$$

et

$$\text{Mid}(r) = \llbracket \text{first}(r), \text{last}(r) \rrbracket \setminus (\text{Top}(r) \cup \text{Bot}(r) \cup \text{Int}(r)).$$

Donnons nous maintenant deux indices de lignes i et j . Notons r (resp. p) le rang du processus auquel appartient la ligne i (resp. j). Ainsi, nous avons $i \in \llbracket \text{first}(r), \text{last}(r) \rrbracket$ et $j \in \llbracket \text{first}(p), \text{last}(p) \rrbracket$. Supposons arbitrairement que la ligne i soit mise à jour avant la ligne j . La valeur x_i intervient dans le calcul de x_j si et seulement si $a_{ji} \neq 0$ (ou de manière équivalente par symétrie $a_{ij} = 0$). Par ailleurs, puisqu'aucune communication entre les processus n'est nécessaire pour accéder à des valeurs locales (*i.e.* stockées sur le processus), nous pouvons directement supposer que $r \neq p$.

Supposons que $j \in \text{Int}(p)$. Par définition, nous avons $a_{ji} = 0$ (puisque nous avons supposé $r \neq p$). Ainsi, la valeur x_i n'intervient pas dans le calcul de x_j ou autrement dit : la modification d'une ligne de l'ensemble $\text{Int}(p)$ d'un processus ne requiert la connaissance d'aucune valeur non locale du vecteur x .

Supposons maintenant que $i \in \text{Int}(r)$. Par définition, nous trouvons $a_{ij} = 0$. Et de nouveau, la valeur x_i n'intervient pas dans le calcul de x_j et en conséquence : il n'est pas nécessaire d'effectuer de communication après la mise à jour d'une ligne de l'ensemble $\text{Int}(r)$.

Ainsi la mise à jour d'une ligne de l'ensemble $\text{Int}(s)$ sur le processus de rang s est complètement indépendante des calculs pouvant être effectués sur les autres processus quelles que soient les lignes qu'ils mettent à jour.

Nous envisageons maintenant le cas $i \in \text{Bot}(r)$. Supposons un instant que $a_{ij} \neq 0$. Alors nous avons nécessairement $j \geq \text{first}(r)$. Nous pouvons en déduire que $\text{first}(r) \leq i < \text{last}(r) < \text{first}(p) \leq j < \text{last}(p)$ et par suite que $r < p$. De plus, puisque par symétrie $a_{ji} \neq 0$ et $i < \text{first}(p)$, on en déduit $j \notin \text{Bot}(p)$ et $j \notin \text{Int}(p)$, autrement dit $j \in \text{Mid}(p) \cup \text{Top}(p)$.

La contraposée du résultat ci-dessus nous permet de déduire deux types d'informations :

- après avoir effectué la mise à jour d'une ligne de $\text{Bot}(r)$, il n'est pas nécessaire de faire de communications pour effectuer celle des lignes des processus de rang p strictement inférieur à r ou celle des lignes des ensembles $\text{Bot}(p)$ pour tout p . La mise à jour des lignes des ensembles $\text{Bot}(r)$ peut être effectuée en parallèle sur tous les processus,
- après avoir effectué la mise à jour de toutes les lignes des ensembles $\text{Bot}(s)$ pour tout s , pour permettre au processus de rang p de mettre à jour une ligne $j \in \llbracket \text{first}(p), \text{last}(p) \rrbracket$, il est suffisant que l'ensemble des processus lui communique les coefficients x_i des lignes $i < \text{first}(p)$ uniquement dans le cas où la ligne $j \in \text{Mid}(p) \cup \text{Top}(p)$ (aucun autre coefficient n'est nécessaire sinon).

Le même type de résultat peut être obtenu en ce qui concerne l'ensemble Top. Ainsi, nous avons les deux conséquences :

- après avoir effectué la mise à jour d'une ligne de $\text{Top}(r)$, il n'est pas nécessaire de faire de communication pour effectuer celle des lignes des processus de rang p strictement supérieur à r ou celle des lignes des ensembles $\text{Top}(p)$ pour tout p . La mise à jour des lignes des ensembles $\text{Top}(r)$ peut être effectuée en parallèle sur tous les processus,
- après avoir effectué la mise à jour de toutes les lignes des ensembles $\text{Top}(s)$ pour tout s , pour permettre au processus de rang p de mettre à jour une ligne $j \in \llbracket \text{first}(p), \text{last}(p) \rrbracket$, il est suffisant que l'ensemble des processus lui communique les coefficients x_i des lignes $i \geq \text{last}(p)$ uniquement dans le cas où la ligne $j \in \text{Mid}(p) \cup \text{Bot}(p)$ (aucun autre coefficient n'est nécessaire sinon).

Par contre, la mise à jour des lignes de l'ensemble Mid est effectuée de manière séquentielle (*i.e.* processus après processus). Ceci n'est pas pénalisant puisqu'elles sont en très petit nombre.

L'algorithme que nous avons implémenté est donc le suivant :

Algorithme III.7 (Algorithme de Gauss-Seidel parallèle)

- | |
|---|
| Etape 1 : Traitement des lignes Top |
| Etape 2 : Envois aux processus concernés |
| Etape 3 : Traitement de quelques lignes Int |
| Etape 4 : Réceptions des valeurs |
| Etape 5 : Traitement des lignes Mid avec communications |
| Etape 6 : Traitement des lignes Bot |
| Etape 7 : Envois aux processus concernés |
| Etape 8 : Traitement de la deuxième partie des lignes Int |
| Etape 9 : Réceptions des valeurs |

Un des intérêt de cet algorithme est qu'il permet, par l'utilisation des mécanismes de la bibliothèque MPI, de recouvrir le temps nécessaire à certaines communications, *i.e.* d'effectuer des opérations pendant que les données sont transférées à travers le réseau. C'est pour cette raison que les étapes 3 et 8 sont insérées entre les étapes d'envois et de reception de données. D'autres stratégies plus complexes sont envisageables (*cf* par exemple [Ada01]).

III.7 Conclusion

Le module de raffinement local est aujourd'hui utilisé dans le laboratoire pour différentes applications de mécanique des fluides. Les méthodes multigrilles, avec l'implémentation actuelle, n'ont pas fait leurs preuves en ce qui concerne les temps de calcul. Il est nécessaire d'envisager des développements supplémentaires pour augmenter l'efficacité (notamment en parallèle) de la construction des opérateurs approchés. Cette dernière est en effet effectuée par des produits matrice-matrice (*cf* formule (II.11)). Il serait sûrement moins coûteux de réaliser directement les assemblages de ces opérateurs.

Partie 2

Discrétisation d'un modèle de type Cahn-Hilliard/Navier-Stokes (CH/NS)

Chapitre IV

Modèle triphasique consistant de type Cahn-Hilliard/Navier-Stokes (CH/NS)

Ce chapitre est dédié à la présentation du système d'équations aux dérivées partielles dont nous souhaitons approcher numériquement les solutions. Ce système constitutif d'un modèle de type interfaces diffuses a été établi et étudié au cours de la thèse de Céline Lapuerta [Lap06] (voir également [BL06] et [BLM⁺]). Nous rappelons brièvement dans ce chapitre les principaux résultats obtenus au cours de la thèse précitée, tout en apportant quelques illustrations et éléments de compréhension nouveaux.

Le principe des modèles à interfaces diffuses est de supposer que les interfaces entre les phases du système ont une épaisseur ε faible mais non nulle. Les interfaces sont alors considérées comme des zones de mélange et la phase i peut être représentée par une indicatrice de phase régulière c_i appelée paramètre d'ordre (que nous prenons ici égale à la fraction volumique de la phase i dans le mélange). Ainsi, le système comporte autant d'inconnues c_i que de phases. Ces inconnues varient entre 0 et 1 (valeurs correspondant par convention aux phases pures) et sont reliées par la relation $\sum_i c_i = 1$.

Une dérivation complète de ce type de modèle pour des écoulements diphasiques est présentée dans les références [AMW98], [Boy02], [Jac99] et [LS03]. Sans mentionner tous les développements théoriques, nous décrivons succinctement dans la section IV.1 les équations obtenues ainsi que leur comportement et leurs principales propriétés.

La section IV.2 reprend ensuite le principe de généralisation à des écoulements triphasiques exposé dans la thèse [Lap06], l'idée directrice étant la construction d'un modèle permettant de retrouver exactement le modèle diphasique lorsque l'une des trois phases n'est pas présente.

Dans ce qui suit, nous nous plaçons sur un domaine Ω ouvert, borné, connexe et régulier de $\mathbb{R}^d$ ($d = 2$ ou $d = 3$).

IV.1 Modèle de type Cahn-Hilliard/Navier-Stokes diphasique

Le système de Cahn-Hilliard (cf section IV.1.1) permet de modéliser la non-miscibilité des phases en maintenant l'épaisseur de la zone de mélange (ou interface) à une valeur prescrite ε . Il autorise également une représentation volumique naturelle des forces capillaires (dûes aux tensions de surface entre les différentes phases). L'hydrodynamique de l'écoulement est prise en compte par le couplage de ces équations au système de Navier-Stokes (cf section IV.1.2).

IV.1.1 Modèle de Cahn-Hilliard diphasique

Lorsque seulement deux phases sont en présence, les deux inconnues du problème, c'est-à-dire les paramètres d'ordre c_1 et c_2 associés à chacune des deux phases, vérifient la relation $c_1 + c_2 = 1$. Le système peut donc être

décrit par un unique paramètre d'ordre que nous noterons $c = c_1 = 1 - c_2$.

Le modèle de Cahn-Hilliard diphasique repose sur un principe de minimisation, sous la contrainte de conservation du volume, d'une énergie, appelée énergie libre :

$$\mathcal{F}_{\sigma,\varepsilon}^{\text{diph}}(c) = \int_{\Omega} 12 \frac{\sigma}{\varepsilon} c^2 (1-c)^2 + \frac{3}{4} \sigma \varepsilon |\nabla c|^2 dx. \quad (\text{IV.1})$$

Cette énergie dépend de deux paramètres constants : la tension de surface σ entre les deux phases et l'épaisseur d'interface ε . Sa minimisation fait entrer en compétition les deux termes qui la composent :

- le premier, proportionnel à $\int_{\Omega} c^2 (1-c)^2 dx$, modélise la non-miscibilité des phases. En effet, la fonction $f(c) = c^2 (1-c)^2$ appelée potentiel de Cahn-Hilliard est en forme de double puits (cf figure IV.1 à gauche) et est minimale pour les valeurs $c = 0$ et $c = 1$ correspondant aux phases pures. Ainsi, ce premier terme est minimal pour des configurations où les interfaces sont infiniment fines.
- le second est proportionnel à $\int_{\Omega} |\nabla c|^2 dx$. Sa minimisation pénalise les fortes variations de c et donc tend à augmenter l'épaisseur ε de l'interface (moralelement proportionnelle à $|\nabla c|^{-1}$ dans l'interface).

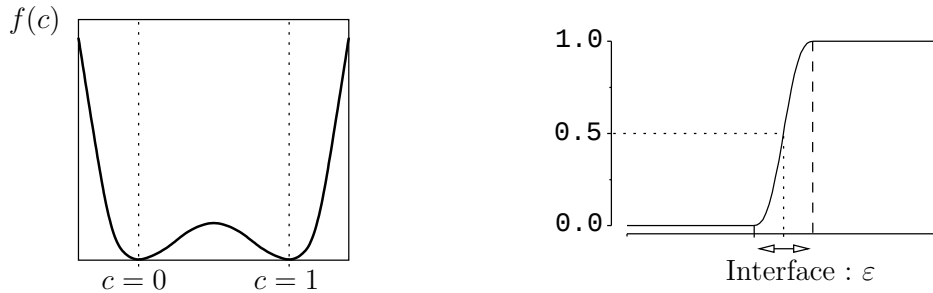


FIG. IV.1 – Potentiel de Cahn-Hilliard diphasique (à gauche) et profil d'équilibre 1D du paramètre d'ordre (à droite)

Les coefficients placés devant ces deux termes sont obtenus par la résolution du problème de minimisation en 1D sur un domaine infini. Le profil minimisant obtenu (cf figure IV.1 à droite) est

$$c_0(x) = 0.5 \left(1 + \tanh \left(\frac{2x}{\varepsilon} \right) \right), \quad (\text{IV.2})$$

ce qui permet d'identifier ε à l'épaisseur d'interface et son énergie est exactement

$$\mathcal{F}_{\sigma,\varepsilon}^{\text{diph}}(c_0) = \sigma, \quad (\text{IV.3})$$

que nous identifions donc à la tension de surface.

L'évolution en temps du paramètre d'ordre est alors décrite par le système d'équations aux dérivées partielles suivant :

$$\begin{cases} \frac{\partial c}{\partial t} = \operatorname{div} (M(c) \nabla \mu), \\ \mu = \frac{12}{\varepsilon} \sigma f'(c) - \frac{3}{2} \varepsilon \sigma \Delta c. \end{cases} \quad (\text{IV.4})$$

La première équation décrit l'évolution de c en fonction de l'inconnue intermédiaire μ appelée potentiel chimique qui est la dérivée fonctionnelle de l'énergie libre $\mathcal{F}_{\sigma,\varepsilon}^{\text{diph}}$ par rapport au paramètre d'ordre c . Le paramètre M est un coefficient de diffusion appelé mobilité pouvant dépendre de c .

Les propriétés importantes de ce système d'équations sont les suivantes :

- l'inconnue $1 - c$ associée à la seconde phase du système est formellement solution du même système d'équations.
- ce système d'équations garantit la conservation du volume total $\int_{\Omega} c dx$ de chaque phase au cours du temps si l'on impose la condition au bord $M(c) \nabla \mu \cdot \mathbf{n} = 0$ sur Γ la frontière de Ω . On peut facilement obtenir ce résultat en intégrant la première équation sur Ω .

– les solutions de ce système vérifient l'égalité d'énergie suivante :

$$\frac{\partial}{\partial t} [\mathcal{F}_{\sigma, \varepsilon}^{\text{diph}}(c)] + \int_{\Omega} M(c) |\nabla \mu|^2 dx = \int_{\Gamma} M(c) \mu \nabla \mu \cdot \mathbf{n} ds + \frac{3}{2} \sigma \varepsilon \int_{\Gamma} \frac{\partial c}{\partial t} \nabla c \cdot \mathbf{n} ds.$$

IV.1.2 Couplage aux équations de Navier-Stokes incompressibles

Le couplage du système de Cahn-Hilliard avec le système des équations de Navier-Stokes incompressibles est effectué de la manière suivante :

- un terme de transport $\mathbf{u} \cdot \nabla c$ du paramètre d'ordre c est ajouté dans la première équation du système de Cahn-Hilliard (IV.4).
- la densité ϱ et la viscosité η sont définies comme des fonctions régulières du paramètre d'ordre c .
- un terme de force capillaire $\mu \nabla c$ est ajouté dans le second membre du bilan de quantité de mouvement (équations de Navier-Stokes).

De plus, nous adoptons une forme non standard des équations de Navier-Stokes. En effet, la densité, en tant que fonction du paramètre d'ordre, ne vérifie pas l'équation de conservation de la masse. Un terme de diffusion supplémentaire est présent dans le second membre :

$$\begin{aligned} \frac{\partial \varrho}{\partial t} + \operatorname{div}(\varrho \mathbf{u}) &= \varrho'(c) \left[\frac{\partial c}{\partial t} + \operatorname{div}(c \mathbf{u}) \right] \\ &= \varrho'(c) \operatorname{div}(M \nabla \mu). \end{aligned}$$

Ainsi, les formes conservative ou non-conservative des équations de Navier-Stokes ne permettent pas de déduire le bilan d'énergie cinétique.

La forme des équations de Navier-Stokes ci-dessous, initialement proposée dans [GQ00], permet de montrer le bilan d'énergie sans utiliser l'équation de conservation de la masse. Elle repose sur l'égalité suivante :

$$\frac{d}{dt} \int_{\Omega_t} \frac{1}{2} \varrho |\mathbf{u}|^2 dx = \int_{\Omega_t} \left[\sqrt{\varrho} \frac{\partial}{\partial t} (\sqrt{\varrho} \mathbf{u}) + (\varrho \mathbf{u} \cdot \nabla) \mathbf{u} + \frac{\mathbf{u}}{2} \operatorname{div}(\varrho \mathbf{u}) \right] \cdot \mathbf{u} dx,$$

le domaine Ω_t étant un domaine borné régulier arbitraire se déplaçant à la vitesse $\mathbf{u}$ du fluide [BF06].

Le modèle Cahn-Hilliard/Navier-Stokes diphasique considéré est alors constitué par les équations suivantes :

$$\begin{cases} \frac{\partial c}{\partial t} + \mathbf{u} \cdot \nabla c = \operatorname{div}(M(c) \nabla \mu), \\ \mu = \frac{12}{\varepsilon} \sigma f'(c) - \frac{3}{2} \sigma \varepsilon \Delta c, \\ \sqrt{\varrho(c)} \frac{\partial}{\partial t} (\sqrt{\varrho(c)} \mathbf{u}) + (\varrho(c) \mathbf{u} \cdot \nabla) \mathbf{u} + \frac{\mathbf{u}}{2} \operatorname{div}(\varrho(c) \mathbf{u}) - \operatorname{div}(2\eta(c) D\mathbf{u}) + \nabla p = \mu \nabla c + \varrho(c) \mathbf{g}, \\ \operatorname{div} \mathbf{u} = 0, \end{cases} \quad (\text{IV.5})$$

où le vecteur $\mathbf{g}$ représente la gravité ; la densité et la viscosité sont définies, à partir d'une fonction de Heaviside régularisée ($\lambda = 0.5$)

$$h_{\lambda}(x) = \begin{cases} 0 & \text{si } x < -\lambda, \\ \frac{1}{2} \left(\frac{x}{\lambda} + \frac{1}{\pi} \sin \left(\pi \frac{x}{\lambda} \right) \right) & \text{si } -\lambda \leq x \leq \lambda, \\ 1 & \text{si } x > \lambda, \end{cases} \quad (\text{IV.6})$$

par la formule :

$$\varrho(c) = \frac{\varrho_1 h_{\lambda}(c - 0.5) + \varrho_2 h_{\lambda}(0.5 - c)}{h_{\lambda}(c - 0.5) + h_{\lambda}(0.5 - c)} \quad \text{et} \quad \eta(c) = \frac{\eta_1 h_{\lambda}(c - 0.5) + \eta_2 h_{\lambda}(0.5 - c)}{h_{\lambda}(c - 0.5) + h_{\lambda}(0.5 - c)},$$

ϱ_1 (resp. ϱ_2) et η_1 (resp. η_2) étant les valeurs supposées constantes de la densité et de la viscosité dans la phase 1 (resp. 2).

Les conditions aux limites et la condition initiale ne sont pas précisées pour l'instant, nous en discuterons ultérieurement.

Remarque IV.1

Les densités et viscosités sont définies de manière à prendre les valeurs prescrites dans chacune des phases. Notons que d'autres choix de régularisations sont possibles (cf [BM00]), en particulier d'autres valeurs de $\lambda < 0.5$; nous n'en discuterons pas dans ce manuscrit.

Remarque IV.2

La formulation des équations de Navier-Stokes donnée ci-dessus est reliée aux formes plus classiques (conservative et non-conservative) par les relations suivantes :

$$\begin{aligned} \sqrt{\varrho} \frac{\partial}{\partial t} (\sqrt{\varrho} \mathbf{u}) + (\varrho \mathbf{u} \cdot \nabla) \mathbf{u} + \frac{\mathbf{u}}{2} \operatorname{div} (\varrho \mathbf{u}) &= \frac{\partial}{\partial t} (\varrho \mathbf{u}) + \operatorname{div} (\varrho \mathbf{u} \otimes \mathbf{u}) - \frac{1}{2} \left[\frac{\partial \varrho}{\partial t} + \operatorname{div} (\varrho \mathbf{u}) \right] \mathbf{u} \\ &= \varrho \frac{\partial \mathbf{u}}{\partial t} + \varrho (\mathbf{u} \cdot \nabla) \mathbf{u} + \frac{1}{2} \left[\frac{\partial \varrho}{\partial t} + \operatorname{div} (\varrho \mathbf{u}) \right] \mathbf{u}. \end{aligned}$$

Ainsi, lorsque l'équation de conservation de la masse $\partial_t \varrho + \operatorname{div} (\varrho \mathbf{u}) = 0$ est vérifiée (cas limite où l'épaisseur d'interface est infiniment fine), ces formulations sont formellement équivalentes.

Formellement, les solutions du système ci-dessus vérifient l'estimation d'énergie suivante :

$$\begin{aligned} \frac{\partial}{\partial t} \left[\int_{\Omega} \frac{1}{2} \varrho(c) |\mathbf{u}|^2 dx + \mathcal{F}_{\sigma, \varepsilon}^{\text{diph}}(\mathbf{c}) \right] + \int_{\Gamma} \left(\frac{1}{2} \varrho(c) |\mathbf{u}|^2 \right) \mathbf{u} \cdot \mathbf{n} ds + \int_{\Omega} 2\eta(c) |D\mathbf{u}|^2 dx + \int_{\Omega} M(\mathbf{c}) |\nabla \mu|^2 dx \\ = \int_{\Omega} \varrho(c) \mathbf{g} \cdot \mathbf{u} dx + \int_{\Gamma} [2\eta(c) D\mathbf{u} \cdot \mathbf{n} - p\mathbf{n}] \cdot \mathbf{u} ds + \int_{\Gamma} M(c) \mu \nabla \mu \cdot \mathbf{n} ds + \frac{3}{2} \sigma \varepsilon \int_{\Gamma} \frac{\partial c}{\partial t} \nabla c \cdot \mathbf{n} ds. \end{aligned}$$

l'énergie créée par convection dans le système de Cahn-Hilliard se compensant exactement avec celle créée par capillarité dans les équations Navier-Stokes.

Enfin, l'équation de conservation de volume, obtenue en intégrant la première équation du système de Cahn-Hilliard, prend maintenant la forme suivante :

$$\frac{\partial}{\partial t} \left[\int_{\Omega} c dx \right] = \int_{\Gamma} [-c \mathbf{u} + M(c) \nabla \mu] \cdot \mathbf{n} ds.$$

IV.2 Modèle de Cahn-Hilliard/Navier-Stokes triphasique

Dans cette section, nous exposons l'extension du modèle de Cahn-Hilliard diphasique présenté dans la section IV.1.1 au cas triphasique puis son couplage aux équations de Navier-Stokes incompressibles. L'idée importante, introduite dans [BL06] et sur laquelle nous insistons ici, est la notion de consistance : le modèle triphasique doit reproduire exactement les situations diphasiques lorsque l'une des trois phases n'est pas présente.

IV.2.1 Modèle de Cahn-Hilliard triphasique

Le système comporte maintenant trois inconnues c_1 , c_2 et c_3 (représentant chacune des phases) liées par la relation :

$$c_1 + c_2 + c_3 = 1. \quad (\text{IV.7})$$

En d'autres termes, le vecteur $\mathbf{c} = (c_1, c_2, c_3)$ appartient à l'hyperplan $\mathcal{S} = \{(c_1, c_2, c_3) \in \mathbb{R}^3; c_1 + c_2 + c_3 = 1\}$ de $\mathbb{R}^3$.

Le modèle triphasique a été introduit dans [BL06] comme une généralisation du modèle deux phases. Les auteurs ont postulé que l'énergie libre trois phases pouvait s'écrire sous la forme suivante :

$$\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(c_1, c_2, c_3) = \int_{\Omega} \frac{12}{\varepsilon} F(c_1, c_2, c_3) + \frac{3}{8} \varepsilon \Sigma_1 |\nabla c_1|^2 + \frac{3}{8} \varepsilon \Sigma_2 |\nabla c_2|^2 + \frac{3}{8} \varepsilon \Sigma_3 |\nabla c_3|^2 dx. \quad (\text{IV.8})$$

Le triplet de paramètres constants $\Sigma = (\Sigma_1, \Sigma_2, \Sigma_3)$ et la forme du potentiel F ont été déterminés de manière à ce que le modèle puisse prendre en compte correctement les valeurs des tensions de surface σ_{12} , σ_{13} et σ_{23} prescrites entre les différents couples de phases et soit "consistant" avec les situations deux-phases (cf paragraphe suivant, expressions (IV.12) et (IV.13)).

Comme dans le cas diphasique, l'évolution du système est alors pilotée par la minimisation, sous la contrainte de conservation du volume, de l'énergie libre $\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}$ et l'évolution en temps des paramètres d'ordre $\mathbf{c} = (c_1, c_2, c_3)$ obéit au système d'équations suivant :

$$\begin{cases} \frac{\partial c_i}{\partial t} = \operatorname{div} \left(\frac{M_0(\mathbf{c})}{\Sigma_i} \nabla \mu_i \right), & \text{pour } i = 1, 2, 3, \\ \mu_i = f_i^F(\mathbf{c}) - \frac{3}{4} \varepsilon \Sigma_i \Delta c_i, & \text{pour } i = 1, 2, 3, \end{cases} \quad (\text{IV.9})$$

où $M_0(\mathbf{c})$ est une coefficient de diffusion appelé *mobilité* qui peut dépendre de $\mathbf{c}$ et

$$f_i^F(\mathbf{c}) = \frac{4\Sigma_T}{\varepsilon} \sum_{j \neq i} \left(\frac{1}{\Sigma_j} (\partial_i F(\mathbf{c}) - \partial_j F(\mathbf{c})) \right) \text{ avec } \Sigma_T \text{ défini par } \frac{3}{\Sigma_T} = \frac{1}{\Sigma_1} + \frac{1}{\Sigma_2} + \frac{1}{\Sigma_3}. \quad (\text{IV.10})$$

Ce choix de f_i^F , obtenu par l'utilisation d'un multiplicateur de Lagrange, impose que la condition (IV.7) soit satisfaite à chaque instant. Ainsi, l'une quelconque des inconnues peut être arbitrairement éliminée du système (IV.9). Nous montrerons que cette propriété est encore vraie au niveau discret (*cf* sections V.1.3 et VI.1.3).

Les propriétés importantes de ce système d'équations sont les suivantes :

- le système d'équation est indépendant de la numérotation attribuée (arbitrairement) aux phases. Cette propriété n'est pas vérifiée par tous les modèles de la littérature (*cf* par exemple [KL05]).
- la conservation du volume total $\int_{\Omega} c_i dx$ de la phase i au cours du temps est garantie si l'on impose la condition aux bords $M_0(\mathbf{c}) \nabla \mu_i \cdot \mathbf{n} = 0$ sur Γ la frontière de Ω .
- les solutions de ce système vérifient l'égalité d'énergie suivante :

$$\frac{\partial}{\partial t} [\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c})] + \sum_{i=1}^3 \int_{\Omega} \frac{M_0(\mathbf{c})}{\Sigma_i} |\nabla \mu_i|^2 dx = \sum_{i=1}^3 \int_{\Gamma} \frac{M_0(\mathbf{c})}{\Sigma_i} \mu_i \nabla \mu_i \cdot \mathbf{n} ds + \frac{3}{4} \sigma \varepsilon \sum_{i=1}^3 \int_{\Gamma} \Sigma_i \frac{\partial c_i}{\partial t} \nabla c_i \cdot \mathbf{n} ds. \quad (\text{IV.11})$$

Consistance avec le modèle diphasique

Pour spécifier complètement le modèle, il reste à fournir l'expression du triplet de paramètres constants Σ et celle du potentiel F . Ces paramètres ont été déterminés pour que le modèle triphasique (défini par (IV.8) et (IV.9)) coïncide exactement avec le modèle diphasique lorsque l'un des trois paramètres d'ordre est nul. Plus précisément, la consistance du modèle trois-phases avec le modèle deux-phases correspondant à la tension de surface σ_{12} (resp. σ_{13} , resp. σ_{23}) signifie que les propriétés suivantes sont vérifiées :

- lorsque le composant $i = 3$ (resp. $i = 2$, resp. $i = 1$) n'est pas présent dans le mélange, c'est-à-dire $c_i \equiv 0$, l'énergie libre $\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(c_1, c_2, c_3)$ du système triphasique est exactement égale à l'énergie libre $\mathcal{F}_{\sigma_{jk}, \varepsilon}^{\text{diph}}(c_j)$ du système diphasique correspondant aux deux autres phases $j = 1$ et $k = 2$ (resp. $j = 1$ et $k = 3$, resp. $j = 2$ et $k = 3$).
- lorsque le composant $i = 3$ (resp. $i = 2$, resp. $i = 1$) n'est pas présent dans le mélange à l'instant initial, le composant i ne doit pas apparaître au cours de l'évolution en temps.

Un résultat important démontré dans [BL06] est que le modèle défini par (IV.8) et (IV.9) est consistant avec le système diphasique associé à la tension de surface σ_{12} , (resp. σ_{13} , resp. σ_{23}) si et seulement si nous avons

$$\Sigma_i = \sigma_{ij} + \sigma_{ik} - \sigma_{jk}, \quad \forall i \in \{1, 2, 3\}, \quad (\text{IV.12})$$

et s'il existe une fonction régulière Λ telle que

$$F(\mathbf{c}) = \sigma_{12} c_1^2 c_2^2 + \sigma_{13} c_1^2 c_3^2 + \sigma_{23} c_2^2 c_3^2 + c_1 c_2 c_3 (\Sigma_1 c_1 + \Sigma_2 c_2 + \Sigma_3 c_3) + c_1^2 c_2^2 c_3^2 \Lambda(\mathbf{c}), \quad \forall \mathbf{c} \in \mathcal{S}. \quad (\text{IV.13})$$

Nous adoptons l'expression des coefficient Σ_i donnée par la relation (IV.12) dans la suite du manuscrit, l'expression du potentiel F choisie (définition de la fonction Λ) sera discutée ultérieurement.

Remarque IV.3

L'expression (IV.12) définissant les coefficients Σ_i implique que,

$$\forall i, j \in \{1, 2, 3\}, \quad \Sigma_i + \Sigma_j = 2\sigma_{ij} > 0.$$

En particulier, il existe au plus un coefficient Σ_i négatif.

Illustrons la propriété de consistance sur un exemple simple (*cf* figure IV.2). Considérons la configuration initiale 2D présentée sur les figures IV.2a et IV.2b : il s'agit de deux phases stratifiées (rouge et bleu sur la figure IV.2a, l'interface étant représentée en vert) ; les paramètres d'ordre (représentés en coupe verticale sur la figure IV.2b) ne dépendent que de la variable en ordonnée et sont initialisés avec la valeur du profil d'équilibre c_0 définie par (IV.2). Ainsi, à l'instant initial, nous avons $c_1(x, y) = c_0(y)$, $c_2(x, y) = 1 - c_0(y)$ et $c_3(x, y) = 0$ pour tout $(x, y) \in \Omega$. Les résultats obtenus après quelques itérations en temps du système de Cahn-Hilliard sont présentés sur les figures IV.2c (modèle consistant) et IV.2d (modèle non consistant). Nous observons l'apparition de la troisième phase lorsqu'un modèle non consistant ($F(\mathbf{c}) = \sigma_{12} c_1^2 c_2^2 + \sigma_{13} c_1^2 c_3^2 + \sigma_{23} c_2^2 c_3^2$) est utilisé alors

que le paramètre d'ordre c_3 reste nul lorsque le modèle est consistant ($F(\mathbf{c}) = \sigma_{12}c_1^2c_2^2 + \sigma_{13}c_1^2c_3^2 + \sigma_{23}c_2^2c_3^2 + c_1c_2c_3(\Sigma_1c_1 + \Sigma_2c_2 + \Sigma_3c_3)$).

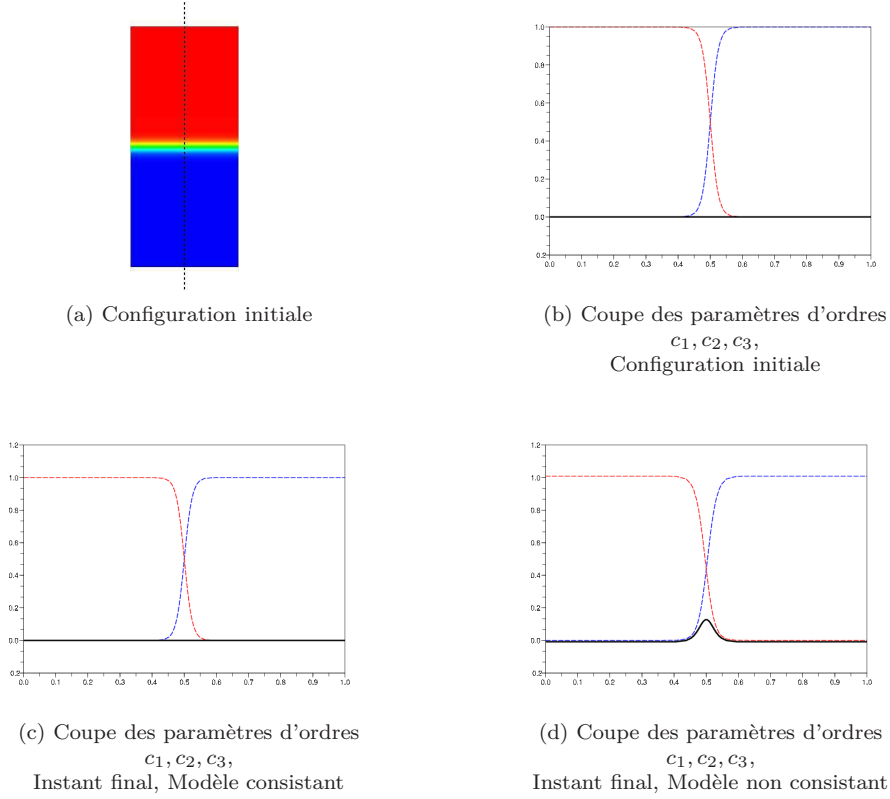


FIG. IV.2 – Illustration de la propriété de consistance du modèle

Remarque IV.4

Les coefficients $S_i = -\Sigma_i$ définis par (IV.12) sont bien connus dans la littérature physique [RW82]. Le coefficient S_i est appelé coefficient d'étalement de la phase i à l'interface entre la phase j et k . Si S_i est positif (i.e. $\Sigma_i < 0$), alors l'étalement est dit total et si S_i est négatif, alors l'étalement est dit partiel.

Notons que dans ce qui suit nous ne supposons pas que les coefficients Σ_i sont positifs, de sorte que le modèle présenté ci-dessus peut prendre en compte certaines situations d'étalement total. Cependant, comme il est prouvé dans [BL06], pour que le système soit bien posé, il est nécessaire de supposer que la condition suivante est vérifiée :

$$\Sigma_1\Sigma_2 + \Sigma_1\Sigma_3 + \Sigma_2\Sigma_3 > 0. \quad (\text{IV.14})$$

Cette condition est équivalente à la coercivité des termes capillaires et garantit que ces termes apportent une contribution positive à l'énergie libre. Ceci est détaillé dans la proposition suivante :

Proposition IV.5 ([BL06, Prop 2.1])

Soit $\Sigma = (\Sigma_1, \Sigma_2, \Sigma_3) \in \mathbb{R}^3$. Il existe $\underline{\Sigma} > 0$ tel que, pour tout $n \geq 1$, pour tout $(\xi_1, \xi_2, \xi_3) \in (\mathbb{R}^n)^3$ tel que $\xi_1 + \xi_2 + \xi_3 = 0$,

$$\Sigma_1|\xi_1|^2 + \Sigma_2|\xi_2|^2 + \Sigma_3|\xi_3|^2 \geq \underline{\Sigma} (|\xi_1|^2 + |\xi_2|^2 + |\xi_3|^2),$$

si et seulement si les deux conditions suivantes sont satisfaites

$$\Sigma_1\Sigma_2 + \Sigma_1\Sigma_3 + \Sigma_2\Sigma_3 > 0 \quad \text{et} \quad \Sigma_i + \Sigma_j > 0, \quad \forall i \neq j. \quad (\text{IV.15})$$

Cette proposition et le corollaire suivant (non triviaux essentiellement lorsque l'un des coefficients Σ_i est négatif) seront très utilisés dans la suite du manuscrit. En particulier, sous la condition (IV.15), la proposition IV.5 montre que la forme bilinéaire définie par $((\xi_1, \xi_2, \xi_3), (\eta_1, \eta_2, \eta_3)) \mapsto \sum_{i=1}^3 \Sigma_i \xi_i \cdot \eta_i$ est un produit

scalaire sur $\{(\xi_1, \xi_2, \xi_3) \in (\mathbb{R}^n)^3 \text{ tel que } \xi_1 + \xi_2 + \xi_3 = 0\}$. Le corollaire suivant est alors déduit en appliquant l'inégalité de Cauchy-Schwarz pour ce produit scalaire et l'inégalité de Young.

Corollaire IV.6

Soit $\Sigma = (\Sigma_1, \Sigma_2, \Sigma_3) \in \mathbb{R}^3$ satisfaisant la condition (IV.15). Alors, pour tout $(\xi_1, \xi_2, \xi_3) \in (\mathbb{R}^n)^3$, satisfaisant $\xi_1 + \xi_2 + \xi_3 = 0$, pour tout $(\eta_1, \eta_2, \eta_3) \in (\mathbb{R}^n)^3$, satisfaisant $\eta_1 + \eta_2 + \eta_3 = 0$, nous avons :

$$\left| \sum_{i=1}^3 \Sigma_i \xi_i \cdot \eta_i \right| \leq \frac{1}{2} \left(\sum_{i=1}^3 \Sigma_i |\xi_i|^2 + \sum_{i=1}^3 \Sigma_i |\eta_i|^2 \right).$$

Au vu de l'expression (IV.12) des coefficients Σ_i , la seconde partie de la condition (IV.15) est toujours vérifiée (cf remarque IV.3) et donc il est suffisant, pour appliquer les lemmes IV.5 et IV.6, de supposer la condition (IV.14).

Solution diphasique du système triphasique

La consistance peut également être interprétée de la manière suivante :

- (i) Supposons que (c, μ) soit une solution du système de Cahn-Hilliard diphasique (IV.4) associé à la tension de surface σ et la mobilité $M(c)$. Posons $\sigma_{12} = \sigma$, et choisissons σ_{23} et σ_{13} deux tensions de surface quelconques. Alors, il existe une solution particulière $((c_i, \mu_i))_{i \in \{1,2,3\}}$ du système triphasique (IV.9) associée aux tensions de surface σ_{12} , σ_{23} et σ_{13} et à la mobilité $M_0(c) = 2\sigma M(c_1)$ telle que $c_1 = c$, $c_2 = 1 - c$ et $c_3 = 0$.
- (ii) Supposons que $((c_i, \mu_i))_{i \in \{1,2,3\}}$ soit une solution diphasique (c'est-à-dire vérifiant $c_3 = \mu_3 = 0$) du système de Cahn-Hilliard triphasique (IV.9) associé aux tensions de surface σ_{12} , σ_{23} et σ_{13} et à la mobilité $M_0(c)$. Alors, il existe une solution (c, μ) du système de Cahn-Hilliard diphasique (IV.4) associé à la tension de surface $\sigma = \sigma_{12}$ et à la mobilité $M(c) = \frac{M_0(c, 1-c, 0)}{2\sigma}$ telle que $c = c_1$ et $1 - c = c_2$.

Ainsi, la correspondance s'effectue de la manière suivante, en posant :

$$c_1 = c, \quad c_2 = 1 - c, \quad \frac{\mu}{\sigma} = 2 \frac{\mu_1}{\Sigma_1} = -2 \frac{\mu_2}{\Sigma_2} \quad \text{et} \quad \sigma M = \frac{M_0}{2},$$

le paramètre ε est bien sûr fixé de manière commune aux modèles diphasique et triphasique.

L'équivalent de cette correspondance pour les schémas numériques est présenté dans la section V.1.3.

Existence de solutions faibles

Notons Γ la frontière du domaine Ω et supposons que celle-ci est divisée en deux parties distinctes $\Gamma = \Gamma_D^c \cup \Gamma_N^c$. Nous ajoutons au système précédent des conditions aux bords mixtes de type Dirichlet-Neumann pour chaque paramètre d'ordre c_i et des conditions aux bords de type Neumann pour chaque potentiel chimique μ_i , c'est-à-dire, pour $i = 1, 2$ et 3 ,

$$c_i = c_{iD} \quad \text{et} \quad M_0(c) \nabla \mu_i \cdot \mathbf{n} = 0, \quad \text{sur } \Gamma_D^c, \quad (\text{IV.16})$$

$$\nabla c_i \cdot \mathbf{n} = 0 \quad \text{et} \quad M_0(c) \nabla \mu_i \cdot \mathbf{n} = 0, \quad \text{sur } \Gamma_N^c, \quad (\text{IV.17})$$

où $\mathbf{c}_D = (c_{1D}, c_{2D}, c_{3D}) \in \left(H^{\frac{1}{2}}(\Gamma) \right)^3$ est donné tel que $\mathbf{c}_D(x) \in \mathcal{S}$ pour presque tout $x \in \Gamma$.

Remarque IV.7

Les conditions aux bords de type Neumann sur μ_i garantissent en particulier la conservation du volume de la phase i . En effet, nous avons

$$\frac{d}{dt} \left(\int_{\Omega} c_i dx \right) = \int_{\Gamma} \frac{1}{\Sigma_i} (-M_0(c) \nabla \mu_i) \cdot \mathbf{n} = 0.$$

Les conditions aux bords de type Neumann sur c_i imposent aux interfaces d'être normales aux frontières du domaine et les conditions de Dirichlet moins classiques sont utilisées sur les bords où l'écoulement est entrant pour simuler par exemple l'injection de la phase i lorsque le modèle de Cahn-Hilliard est couplé aux équations de Navier-Stokes (cf section IV.2.2).

Au vu des conditions aux bords (IV.16)-(IV.17), nous introduisons les espaces fonctionnels suivants :

$$\begin{aligned}\mathcal{V}^c &= \mathcal{V}^\mu = H^1(\Omega), \\ \mathcal{V}_D^{c_i} &= \{\nu^{c_i} \in H^1(\Omega); \nu^{c_i} = c_{iD} \text{ sur } \Gamma_D^c\}, \quad \text{pour } i = 1, 2 \text{ et } 3, \\ \mathcal{V}_{D,0}^c &= \{\nu^c \in H^1(\Omega); \nu^c = 0 \text{ sur } \Gamma_D^c\}, \\ \mathcal{V}_{D,\mathcal{S}}^c &= \{\mathbf{c} = (c_1, c_2, c_3) \in \mathcal{V}_D^{c_1} \times \mathcal{V}_D^{c_2} \times \mathcal{V}_D^{c_3}; \mathbf{c}(x) \in \mathcal{S} \text{ pour presque tout } x \in \Omega\}.\end{aligned}$$

Finalement, nous supposons qu'à l'instant initial, nous avons

$$c_i(t=0) = c_i^0, \quad (\text{IV.18})$$

où $\mathbf{c}^0 = (c_1^0, c_2^0, c_3^0) \in \mathcal{V}_{D,\mathcal{S}}^c$ est donné.

L'existence de solutions faibles au problème (IV.9) avec la condition initiale (IV.18) et des conditions aux bords de type Neumann (IV.17) ($\Gamma = \Gamma_N^c$) pour chacune des inconnues (c_i, μ_i) , a été prouvée dans [BL06] en 2D et 3D sous les hypothèses générales suivantes :

- La mobilité M_0 est une fonction bornée de classe $\mathcal{C}^1(\mathbb{R}^3)$ et il existe trois constantes positives M_1, M_2 et M_3 telles que :

$$\begin{aligned}\forall \mathbf{c} \in \mathcal{S}, \quad 0 < M_1 \leq M_0(\mathbf{c}) \leq M_2, \\ |DM_0(\mathbf{c})| \leq M_3.\end{aligned} \quad (\text{IV.19})$$

- Le potentiel de Cahn-Hilliard F est une fonction positive de classe $\mathcal{C}^2(\mathbb{R}^3)$ qui satisfait les hypothèses de croissance polynômiale suivantes : il existe $B_1 > 0$ et un réel p tel que $2 \leq p < +\infty$ si $d = 2$ ou $2 \leq p \leq 6$ si $d = 3$, et

$$\begin{aligned}\forall \mathbf{c} \in \mathcal{S}, \quad |F(\mathbf{c})| &\leq B_1 (1 + |\mathbf{c}|^p), \\ |DF(\mathbf{c})| &\leq B_1 (1 + |\mathbf{c}|^{p-1}), \\ |D^2F(\mathbf{c})| &\leq B_1 (1 + |\mathbf{c}|^{p-2}).\end{aligned} \quad (\text{IV.20})$$

Remarque IV.8

La définition de p implique que $H^1(\Omega) \subset L^p(\Omega)$. Rappelons qu'il existe une constante $C_{S,p}$ strictement positive telle que :

$$\forall u \in H^1(\Omega), \quad |u|_{L^p(\Omega)} \leq C_{S,p} |u|_{H^1(\Omega)}.$$

Théorème IV.9

Supposons que les coefficients $(\Sigma_1, \Sigma_2, \Sigma_3)$ satisfont (IV.14), que la mobilité M_0 satisfait (IV.19), et que le potentiel de Cahn-Hilliard F satisfait (IV.20). Considérons le problème (IV.9) avec la condition initiale (IV.18) et les conditions aux bords (IV.17) de type Neumann ($\Gamma = \Gamma_N^c$) pour chacune des inconnues (c_i, μ_i) . Alors, il existe une solution faible $(\mathbf{c}, \boldsymbol{\mu})$ sur $[0, +\infty[$ telle que

$$\begin{aligned}\mathbf{c} &\in L^\infty(0, +\infty; (H^1(\Omega))^3) \cap \mathcal{C}^0([0, +\infty[; (L^q(\Omega))^3), \text{ pour tout } q < 6, \\ \boldsymbol{\mu} &\in L^2(0, +\infty; (H^1(\Omega))^3), \\ \mathbf{c}(t, x) &\in \mathcal{S}, \text{ pour presque tout } (t, x) \in [0, +\infty[\times \Omega.\end{aligned}$$

Nous donnons dans ce manuscrit (cf sections V.1.5 et V.5) une preuve différente de ce résultat, pour les conditions aux bords plus générales (IV.16) et (IV.17) en passant à la limite dans le schéma numérique.

Remarque IV.10

Dans [BL06], un théorème d'unicité est aussi démontré sous des hypothèses supplémentaires sur la Hessienne du potentiel F . Notons que, en trois dimensions, la preuve requiert une mobilité constante et une légère modification de l'expression du terme de plus haut degré du potentiel F que nous n'envisagerons pas dans ce manuscrit.

Expression du potentiel de Cahn-Hilliard triphasique

Rappelons qu'une condition nécessaire (et suffisante puisque les coefficients Σ_i sont définis par (IV.12)) pour que le modèle triphasique soit consistant avec le modèle diphasique (cf section IV.2.1) est que le potentiel de Cahn-Hilliard soit de la forme suivante (cf équation (IV.13)) :

$$F(\mathbf{c}) = \underbrace{\sigma_{12}c_1^2c_2^2 + \sigma_{13}c_1^2c_3^2 + \sigma_{23}c_2^2c_3^2 + c_1c_2c_3(\Sigma_1c_1 + \Sigma_2c_2 + \Sigma_3c_3)}_{F_0(\mathbf{c})} + \underbrace{\Lambda(\mathbf{c})c_1^2c_2^2c_3^2}_{P(\mathbf{c})}, \quad \forall \mathbf{c} \in \mathbb{R}^3, \quad (\text{IV.21})$$

le coefficient Λ (à déterminer) pouvant dépendre éventuellement de $\mathbf{c}$.

Lorsque $\mathbf{c} \in \mathcal{S}$, il est possible d'obtenir une expression du potentiel F_0 en fonction du potentiel de Cahn-Hilliard diphasique f défini, rappelons-le, par

$$\forall c \in \mathbb{R}, \quad f(c) = c^2(1-c)^2. \quad (\text{IV.22})$$

En effet, par définition des coefficients Σ_i (cf (IV.12)), nous avons pour tout $\mathbf{c} \in \mathbb{R}^3$,

$$\begin{aligned} F_0(\mathbf{c}) &= \frac{\Sigma_1 + \Sigma_2}{2}c_1^2c_2^2 + \frac{\Sigma_1 + \Sigma_3}{2}c_1^2c_3^2 + \frac{\Sigma_2 + \Sigma_3}{2}c_2^2c_3^2 + c_1c_2c_3(\Sigma_1c_1 + \Sigma_2c_2 + \Sigma_3c_3) \\ &= \sum_{i=1}^3 \frac{\Sigma_i}{2}c_i^2(c_j + c_k)^2, \end{aligned}$$

où j et k désignent les deux indices appartenant à $\{1, 2, 3\}$ différents de i . Ainsi, sur $\mathcal{S}$, nous obtenons l'expression suivante de F_0 :

$$F_0(\mathbf{c}) = \sum_{i=1}^3 \frac{\Sigma_i}{2}f(c_i), \quad \forall \mathbf{c} \in \mathcal{S}. \quad (\text{IV.23})$$

Il est important de noter que dans les cas d'étalement partiel, *i.e.* $\Sigma_i > 0, \forall i = 1, 2, 3$, le potentiel F_0 satisfait les hypothèses (IV.20) et, en conséquence, le choix le plus simple $F = F_0$ est toujours acceptable. Cependant, dans les cas d'étalement total, *i.e.* l'un des Σ_i est négatif, le potentiel F_0 peut ne pas être minoré. Néanmoins, la proposition suivante, tirée de [BL06], garantit que $F = F_0 + P$ est une fonction positive satisfaisant (IV.20) à condition que le paramètre Λ (choisi constant, *i.e.* indépendant de $\mathbf{c}$) soit assez grand.

Proposition IV.11 ([BL06, Prop 3.7])

Si les coefficients $(\Sigma_1, \Sigma_2, \Sigma_3)$ satisfont la condition (IV.14), alors il existe $\Lambda_0 > 0$ tel que pour tout $\Lambda \in [\Lambda_0, +\infty[$, le potentiel F défini par (IV.21) est positif sur $\mathcal{S}$ et satisfait les propriétés (IV.20).

Ainsi, ces résultats permettent dans tous les cas (situations d'étalement partiel et total) d'obtenir une expression convenable du potentiel de Cahn-Hilliard (au sens où elle conduit à un modèle consistant bien posé). Cependant, en pratique, dans les cas d'étalement total, l'influence de la valeur du paramètre constant Λ sur les résultats des simulations est importante, son choix reste donc délicat.

Nous adoptons alors la démarche de modélisation suivante :

- nous justifions par une étude numérique que, dans le cas d'étalement partiel, le choix le plus simple $F = F_0$, *i.e.* $\Lambda = 0$, est celui qui convient. L'étude se base sur la simulation d'une lentille piégée entre deux phases stratifiées. Les solutions stationnaires numériques du système Cahn-Hilliard (IV.9) obtenues pour différentes valeurs constantes de Λ sont comparées aux solutions "physiques" (les angles de contact entre les interfaces aux points triples sont donnés en fonction des tensions de surface par la loi de Young).
- nous montrons ensuite que $F_0(\mathbf{c})$ est positif lorsque $\mathbf{c} \in \mathbf{T} = \{\mathbf{c} \in \mathcal{S}, \forall i = 1, 2, 3, 0 \leq c_i \leq 1\}$ et ceci même en situation d'étalement total. Il n'y a donc *a priori* aucune raison que le terme $P = \Lambda c_1^2c_2^2c_3^2$ ait une influence dans ce domaine. Néanmoins, il reste indispensable en dehors du domaine $\mathbf{T}$ puisque la fonction F_0 peut y devenir négative.
- ces deux résultats, combinés à celui de la proposition IV.11, font émerger l'idée d'utiliser un coefficient Λ dépendant de $\mathbf{c}$ comme fonction de "troncature", pour diminuer (ou supprimer) l'action du terme $c_1^2c_2^2c_3^2$ (non nécessaire) sur le domaine $\mathbf{T}$, sans la modifier en dehors, garantissant la positivité du potentiel F correspondant.

Ce travail, effectué en collaboration avec R. Bonhomme est encore en cours. Nous présentons ici les premiers résultats.

Remarque IV.12

Il peut être surprenant que nous considérions avec autant d'importance les valeurs de $\mathbf{c}$ en dehors du domaine $\mathbf{T}$ puisque celles-ci n'ont pas de sens physique (rappelons que c_i représente la fraction de phase i dans la

mélange). Cependant, aucun principe du maximum n'est satisfait par les équations de Cahn-Hilliard (d'ordre 4) et ainsi rien ne garantit qu'une solution (éventuelle) de l'équation reste dans le domaine $\mathbf{T}$ pour tout temps. Sans information sur le comportement du potentiel F à l'extérieur du domaine $\mathbf{T}$, il est donc impossible de prouver l'existence de solutions au problème continu, et du point de vue numérique une non convergence de l'algorithme de linéarisation (méthode de Newton) est observée systématiquement lorsque le terme P n'est pas utilisé dans les cas d'étalement total.

Expression du coefficient Λ – Etalement partiel

En anticipant sur les chapitres à venir du manuscrit où sont présentés en détail les schémas numériques, nous exposons dans cette section le résultat de simulations en géométrie axisymétrique tridimensionnelle d'une lentille initialement sphérique (de rayon $R = 0.08$) piégée entre deux phases stratifiées. La configuration initiale et la numérotation des phases sont données par la figure IV.3.

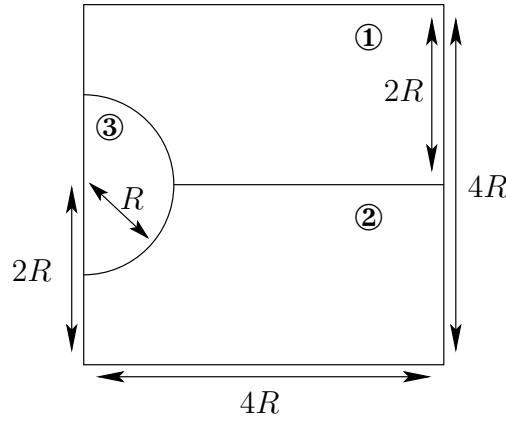


FIG. IV.3 – Configuration initiale, numérotation des phases

Les valeurs des tensions de surface choisies entre les trois fluides en présence :

$$\sigma_{12} = 0.07, \quad \sigma_{13} = 0.1912436 \quad \text{et} \quad \sigma_{23} = 0.2342246,$$

conduisent à des paramètres Σ_i tous positifs : $\Sigma_1 = 0.027019$, $\Sigma_2 = 0.112981$ et $\Sigma_3 = 0.3554682$.

Lorsqu'un équilibre est atteint, il est possible par un bilan de force (cf [Lap06, Annexe B]) d'obtenir les valeurs théoriques des angles de contact :

$$\cos \theta_i = \frac{\sigma_{jk}^2 - \sigma_{ij}^2 - \sigma_{ik}^2}{2\sigma_{ij}\sigma_{ik}},$$

où θ_i est l'angle de contact (au niveau du point triple) entre les plans tangents aux interfaces ij et ik (cf Figure IV.4), j et k désignant les deux indices différents de i .

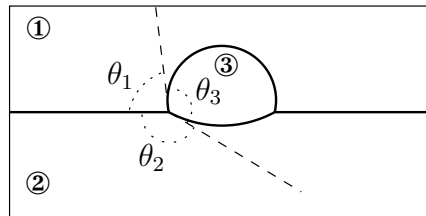
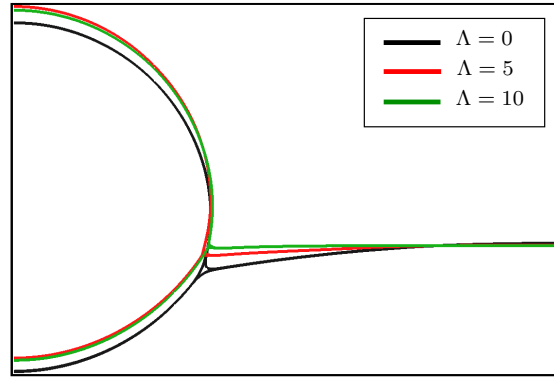


FIG. IV.4 – Définition des angles de contact

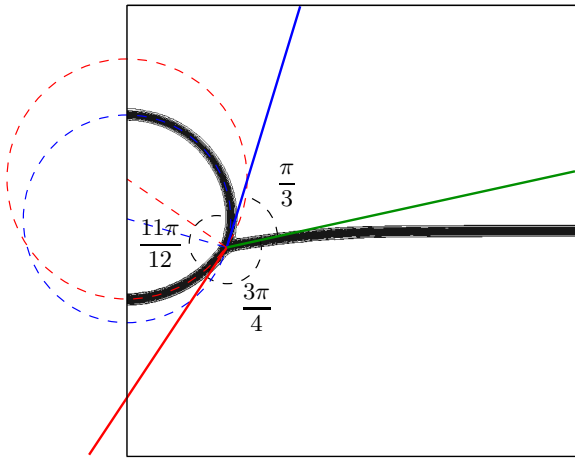
Avec les valeurs des tensions de surface données ci-dessus, nous obtenons les angles de contact suivants :

$$\theta_1 = \frac{11\pi}{12} (= 165^\circ), \quad \theta_2 = \frac{3\pi}{4} (= 135^\circ) \quad \text{et} \quad \theta_3 = \frac{\pi}{3} (= 60^\circ). \quad (\text{IV.24})$$

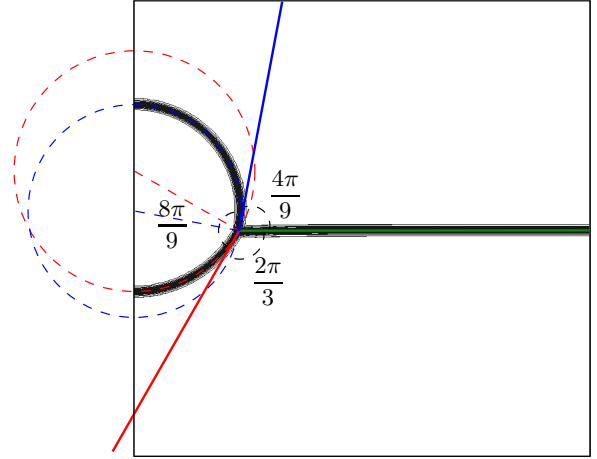
La figure IV.5 présente les résultats obtenus par la simulation en utilisant le potentiel F avec un coefficient Λ constant égal à 0, 5 ou 10. Sur la figure IV.5a, sont superposées les lignes de niveau ($c_i = 0.5$, $i = 1, 2, 3$) des paramètres d'ordre obtenus à l'état stationnaire pour chacune des trois simulations. Ceci permet en particulier de visualiser la position des interfaces entre les différentes phases. Nous observons que la valeur du paramètre Λ influe sur la forme de la lentille et sur les angles de contact obtenus à l'état stationnaire. Les figures IV.5b et IV.5c présentent alors le calcul de ces angles de contact pour les cas $\Lambda = 0$ et $\Lambda = 10$ respectivement (les cercles en pointillé ont été tracé afin de mieux approcher les tangentes aux interfaces). Nous constatons que les valeurs théoriques (IV.24) sont retrouvées seulement lorsque $\Lambda = 0$.



(a) Lignes de niveau ($c_1=0.5$, $c_2=0.5$, $c_3=0.5$) des trois paramètres d'ordre à l'état stationnaire pour différentes valeurs constantes de Λ .



(b) $\Lambda = 0$, Lignes de niveau ($0.05 \leq c_i \leq 0.95$) des trois paramètres d'ordre à l'état stationnaire (en noir) et angles de contact.



(c) $\Lambda = 10$, Lignes de niveau ($0.05 \leq c_i \leq 0.95$) des trois paramètres d'ordre à l'état stationnaire (en noir) et angles de contact.

FIG. IV.5 – Influence de la valeur du coefficient Λ constant dans une situation d'étalement partiel

Ainsi, le choix $F = F_0$, outre le fait qu'il soit mathématiquement convenable dans les situations d'étalement partiel, permet de retrouver les angles de contact aux points triples prédits par la théorie.

Expression du coefficient Λ – Etalement total

Nous montrons maintenant que même dans le cas d'étalement total, $F_0(\mathbf{c})$ est positif si $\mathbf{c} \in \mathcal{S}$ satisfait $0 \leq c_i \leq 1$.

Proposition IV.13

Si les coefficients $(\Sigma_1, \Sigma_2, \Sigma_3)$ satisfont (IV.14), alors nous avons

$$F_0(\mathbf{c}) \geq 0, \quad \forall \mathbf{c} \in \mathbf{T},$$

où $\mathbf{T} = \{\mathbf{c} \in \mathcal{S} / \forall i \in \{1, 2, 3\}, 0 \leq c_i \leq 1\}$.

Démonstration : Soit $\mathbf{c} \in \mathbf{T}$. Nous utilisons l'expression (IV.23) du potentiel F_0 . Le potentiel diphasique f (cf (IV.22)) étant positif, le résultat est trivial lorsque tous les Σ_i sont positifs. De plus, la valeur des coefficients Σ_i donnée par l'équation (IV.12) implique qu'au plus un de ces coefficients est négatif ($\Sigma_i + \Sigma_j = 2\sigma_{ij} \geq 0$, cf remarque IV.3). Supposons, sans perte de généralité, que le coefficient Σ_3 est négatif. Les deux coefficients Σ_1 et Σ_2 sont donc positifs. De plus, la condition (IV.14) implique que

$$|\Sigma_3| < \frac{\Sigma_1 \Sigma_2}{\Sigma_1 + \Sigma_2}.$$

Ainsi, puisque f est positive, nous obtenons,

$$F_0(\mathbf{c}) \geq \Sigma_1 f(c_1) + \Sigma_2 f(c_2) - \frac{\Sigma_1 \Sigma_2}{\Sigma_1 + \Sigma_2} f(c_3).$$

En multipliant par $(\Sigma_1 + \Sigma_2) \geq 0$, il vient :

$$\begin{aligned} (\Sigma_1 + \Sigma_2) F_0(\mathbf{c}) &\geq \Sigma_1 (\Sigma_1 + \Sigma_2) f(c_1) + \Sigma_2 (\Sigma_1 + \Sigma_2) f(c_2) - \Sigma_1 \Sigma_2 f(c_3), \\ &\geq \Sigma_1^2 f(c_1) + \Sigma_2^2 f(c_2) + \Sigma_1 \Sigma_2 (f(c_1) + f(c_2) - f(c_3)). \end{aligned}$$

Or nous avons, par l'inégalité de Young,

$$\Sigma_1^2 f(c_1) + \Sigma_2^2 f(c_2) \geq 2 \Sigma_1 \Sigma_2 c_1 (1 - c_1) c_2 (1 - c_2),$$

puisque $f(c) = c^2(1 - c)^2$. Ainsi, il vient

$$\begin{aligned} (\Sigma_1 + \Sigma_2) F_0(\mathbf{c}) &\geq \Sigma_1 \Sigma_2 \left[2c_1(1 - c_1)c_2(1 - c_2) + f(c_1) + f(c_2) - f(c_3) \right] \\ &\geq \Sigma_1 \Sigma_2 \left[\left[c_1(1 - c_1) + c_2(1 - c_2) \right]^2 - \left[c_3(1 - c_3) \right]^2 \right]. \end{aligned}$$

L'identité remarquable $a^2 - b^2 = (a - b)(a + b)$ conduit alors à

$$(\Sigma_1 + \Sigma_2) F_0(\mathbf{c}) \geq \Sigma_1 \Sigma_2 [c_1(1 - c_1) + c_2(1 - c_2) + c_3(1 - c_3)] [c_1(1 - c_1) + c_2(1 - c_2) - c_3(1 - c_3)]. \quad (\text{IV.25})$$

De plus, nous avons $c_3 = 1 - c_1 - c_2$ (puisque $\mathbf{c} \in \mathcal{S}$) et donc

$$c_1(1 - c_1) + c_2(1 - c_2) - c_3(1 - c_3) = 2c_1c_2.$$

Comme $\mathbf{c} \in \mathbf{T}$, tous les termes c_i et $1 - c_i$ sont positifs. Par suite, le second membre de (IV.25) est positif et nous obtenons la conclusion. ■

La proposition IV.11 nous autorise à considérer un coefficient Λ variable s'annulant à l'intérieur de $\mathbf{T}$. Pour cela, nous posons $\mathbf{T}_\eta = \{\mathbf{c} \in \mathcal{S} / \forall i \in \{1, 2, 3\}, c_i \geq \eta\}$, $\eta \in]0, \frac{1}{2}]$. Le domaine $\mathbf{T}_\eta$ est inclus dans $\mathbf{T}$ (cf figure IV.6).

Nous choisissons une fonction $\varphi_\eta : \mathbb{R}^2 \rightarrow \mathbb{R}$ de classe $\mathcal{C}^\infty$ telle que

$$0 \leq \varphi_\eta(c_1, c_2) \leq 1, \text{ si } (c_1, c_2, 1 - c_1 - c_2) \in \mathbf{T} \setminus \mathbf{T}_\eta \quad \text{et} \quad \varphi_\eta(c_1, c_2) = \begin{cases} 0 & \text{si } (c_1, c_2, 1 - c_1 - c_2) \in \mathbf{T}_\eta, \\ 1 & \text{si } (c_1, c_2, 1 - c_1 - c_2) \notin \mathbf{T}. \end{cases}$$

Nous définissons ensuite $\varphi_\eta : \mathbb{R}^3 \rightarrow \mathbb{R}$ par $\varphi_\eta(\mathbf{c}) = \varphi_\eta(c_1, c_2)$ (cf figure IV.6). La fonction φ_η est de classe $\mathcal{C}^\infty$ et satisfait :

$$0 \leq \varphi_\eta(\mathbf{c}) \leq 1, \text{ si } \mathbf{c} \in \mathbf{T} \setminus \mathbf{T}_\eta \quad \text{et} \quad \varphi_\eta(\mathbf{c}) = \begin{cases} 0 & \text{si } \mathbf{c} \in \mathbf{T}_\eta, \\ 1 & \text{si } \mathbf{c} \notin \mathbf{T}. \end{cases}$$

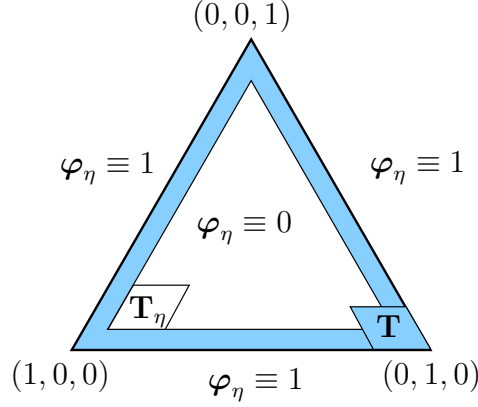


FIG. IV.6 – Définition de la fonction de troncature φ_η sur l'hyperplan $\mathcal{S}$ (représenté en coordonnées barycentriques (c_1, c_2, c_3))

Nous avons le résultat suivant :

Proposition IV.14

Si les coefficients $(\Sigma_1, \Sigma_2, \Sigma_3)$ vérifient la condition (IV.14), alors il existe $\Lambda_0 > 0$ tel que pour tout $\lambda \in [\Lambda_0, +\infty[$ et pour tout $\eta \in]0, \frac{1}{3}[$, le potentiel F défini par (IV.21) avec $\Lambda(\mathbf{c}) = \lambda \varphi_\eta(\mathbf{c})$ est positif et satisfait les propriétés (IV.20). De plus, nous avons

$$\forall \mathbf{c} \in \mathbf{T}_\eta, \quad F(\mathbf{c}) = F_0(\mathbf{c}), \quad (\text{IV.26})$$

et

$$\forall \mathbf{c} \in \mathbf{T} \setminus \mathbf{T}_\eta, \quad |F(\mathbf{c}) - F_0(\mathbf{c})| \leq \frac{3}{16} \lambda \eta^2. \quad (\text{IV.27})$$

Démonstration : La positivité de F se déduit des propositions IV.11 et IV.13. En effet, la proposition IV.11 nous donne $\Lambda_0 > 0$ tel que $\forall \lambda \geq \Lambda_0, \forall \mathbf{c} \in \mathcal{S}, \quad F_0(\mathbf{c}) + \lambda c_1^2 c_2^2 c_3^2 \geq 0$. Nous distinguons alors les deux cas suivants :

- si $\mathbf{c} \in \mathbf{T}$, nous avons $F(\mathbf{c}) \geq F_0(\mathbf{c}) \geq 0$ d'après la proposition IV.13.
- sinon, $\mathbf{c} \notin \mathbf{T}$, et nous avons $\varphi_\eta(\mathbf{c}) = 1$. Ainsi, puisque $F(\mathbf{c}) = F_0(\mathbf{c}) + \lambda c_1^2 c_2^2 c_3^2$, nous avons pour tout $\lambda \geq \Lambda_0, F(\mathbf{c}) \geq 0$.

En conclusion, la fonction F définie par (IV.21) avec $\Lambda(\mathbf{c}) = \lambda \varphi_\eta(\mathbf{c})$ est positive sur $\mathcal{S}$ pour tout $\lambda > \Lambda_0$, pour tout $\eta \in]0, \frac{1}{3}[$. La relation (IV.26) se déduit facilement de la définition de la fonction φ_η . La relation (IV.27) est obtenue en remarquant que $\forall \mathbf{c} \in \mathbf{T} \setminus \mathbf{T}_\eta, \exists i \in \{1, 2, 3\}$ tel que $0 \leq c_i \leq \eta$. En effet, puisque $0 \leq \varphi_\eta(\mathbf{c}) \leq 1$, nous avons, pour $\mathbf{c} \in \mathbf{T} \setminus \mathbf{T}_\eta$,

$$\begin{aligned} |F(\mathbf{c}) - F_0(\mathbf{c})| &\leq \lambda c_i^2 c_j^2 c_k^2 \\ &\leq \frac{3}{16} \lambda \eta^2. \end{aligned}$$

Enfin, les propriétés (IV.20) s'obtiennent facilement. Sur $\mathbf{T}$, la fonction F est de classe $\mathcal{C}^2$, donc elle-même, ses dérivées premières et secondes sont bornées. À l'extérieur de $\mathbf{T}$, la fonction φ_η est constante égale à 1 et donc ses dérivées premières et secondes sont nulles. La proposition IV.11 garantit alors que les majorations sont également satisfaites en dehors de $\mathbf{T}$. ■

Cette dernière proposition montre que d'un point de vue mathématique, le choix $\Lambda(\mathbf{c}) = \lambda \varphi_\eta(\mathbf{c})$ est acceptable (au même titre que Λ constant). De plus, les deux (in)égalités (IV.27) et (IV.26) montrent que l'influence de ce terme $\Lambda(\mathbf{c})$ pour des valeurs de $\mathbf{c}$ dans le domaine $\mathbf{T}$ est contrôlée par le paramètre η de régularisation de la fonction de troncature φ_η .

Illustrons maintenant ces résultats théoriques par quelques simulations. En pratique, nous choisissons l'expression suivante de la fonction φ_η :

$$\varphi_\eta(\mathbf{c}) = 1 - g_\eta(c_1)g_\eta(c_2)g_\eta(c_3),$$

où $g_\eta(x) = \frac{1}{\alpha^5}x^3(6x^2 - 15\alpha x + 10\alpha^2)$ et $\eta = 0.2$. Nous reprenons tout d'abord le cas test en situation d'étalement partiel présenté dans le paragraphe précédent mais cette fois en utilisant la définition $\Lambda(\mathbf{c}) = \lambda \varphi_\eta(\mathbf{c})$. Les résultats obtenus pour les valeurs $\lambda = 0$, $\lambda = 5$ et $\lambda = 10$ sont superposés sur la figure IV.7. Nous observons que l'utilisation du nouveau potentiel diminue fortement l'influence du coefficient $\lambda = \sup_{\mathbf{c}} \varphi_\eta(\mathbf{c})$.

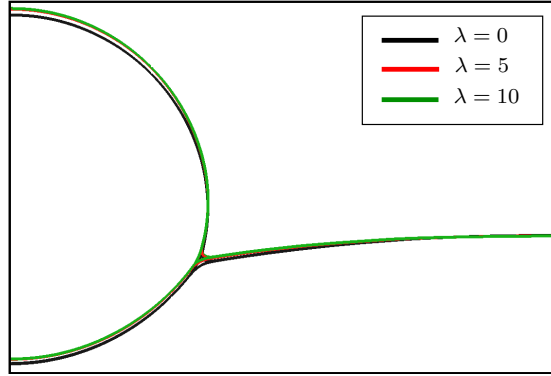


FIG. IV.7 – Situation d'étalement partiel, lignes de niveau ($c_1=0.5$, $c_2=0.5$, $c_3=0.5$) des trois paramètres d'ordre à l'état stationnaire pour différentes valeurs de λ avec $\Lambda(\mathbf{c}) = \lambda \varphi_\eta(\mathbf{c})$.

Nous proposons, ensuite, un cas test en situation d'étalement total. Les valeurs des tensions de surfaces entre les trois fluides en présence sont les suivantes : $\sigma_{12} = 0.01$, $\sigma_{13} = 0.03$, $\sigma_{23} = 0.01$, l'indice 1 (resp. 2) étant celui qui désigne la phase placée au dessus (resp. au dessous) de la lentille, et l'indice 3 celui qui désigne la phase contenue dans la lentille. Le coefficient Σ_2 est donc négatif : à l'état stationnaire les interfaces entre la phase 1 et la phase 3 ne devraient plus exister, c'est-à-dire que la lentille devrait être complètement incluse dans la phase 2 (celle du dessous). Numériquement, il est difficile d'atteindre l'état stationnaire, néanmoins nous pouvons comparer les résultats obtenus à différents instants pour les différents potentiels : $\Lambda = \text{cste}$ et $\Lambda(\mathbf{c}) = \lambda \varphi_\eta(\mathbf{c})$. Les résultats sont montrés sur la figure IV.8 : les deux rangées de graphiques du haut et du milieu sont réalisées avec $\Lambda = \text{cste}$ (0.1 et 2 respectivement), celle du bas est réalisée avec $\Lambda(\mathbf{c}) = \lambda \varphi_\eta(\mathbf{c})$ ($\lambda = 2$). Nous observons que la bulle s'extraît complètement de la phase du haut pour des valeurs faibles de $\Lambda = \text{cste}$ (par exemple $\Lambda = 0.1$). Par contre, dès que la valeur de Λ augmente (par exemple $\Lambda = 2$) les résultats obtenus sont différents et pour $t = 1.96$ la bulle n'est pas encore extraite. L'utilisation du potentiel $\Lambda(\mathbf{c}) = \lambda \varphi_\eta(\mathbf{c})$ permet de fortement réduire l'influence du coefficient puisqu'avec la même valeur $\lambda = 2$ nous retrouvons des résultats similaires au cas $\Lambda = 0.1$.

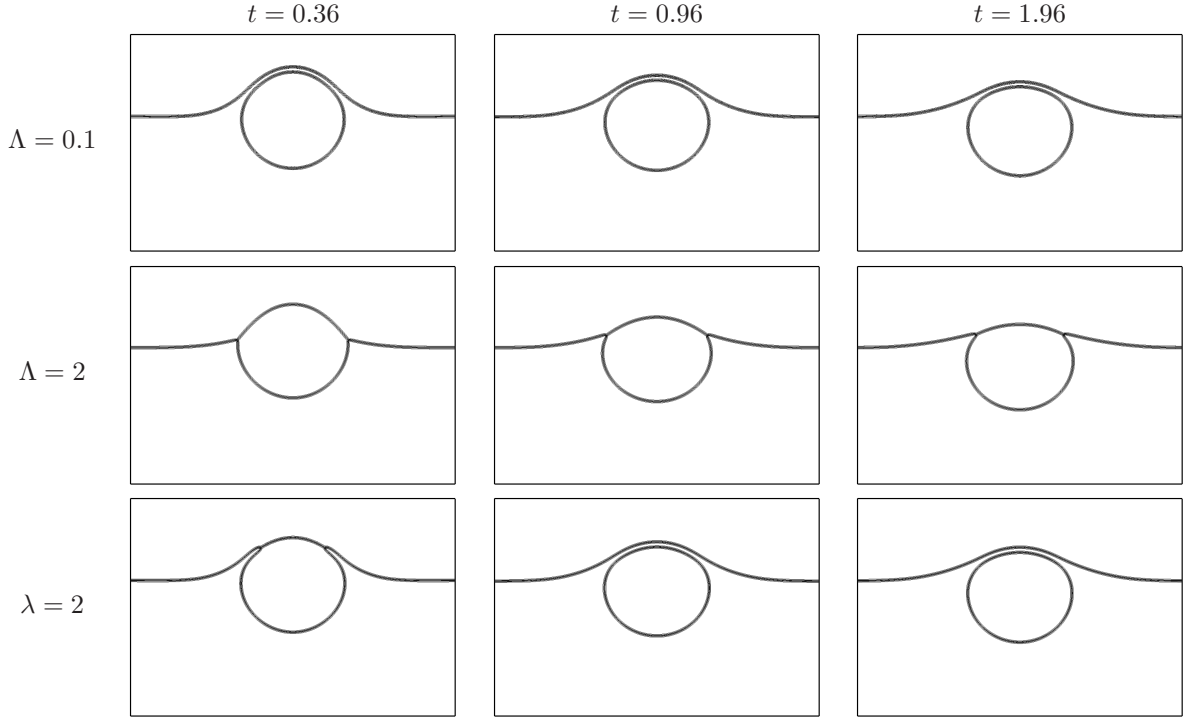


FIG. IV.8 – Situation d'étalement total, comparaison entre les potentiels F pour $\Lambda(\mathbf{c}) = \text{cste}$ et $\Lambda(\mathbf{c}) = \lambda\varphi_\eta(\mathbf{c})$.

IV.2.2 Couplage aux équations de Navier-Stokes incompressibles

De manière similaire au cas diphasique, le couplage avec les équations de Navier-Stokes incompressibles se fait :

- en ajoutant un terme de transport $\mathbf{u} \cdot \nabla c_i$ dans l'équation d'évolution (première équation du système (IV.9)) de chaque paramètre d'ordre c_i , $i \in \{1, 2, 3\}$.
- en définissant la densité et la viscosité comme des fonctions régulières des paramètres d'ordre $\mathbf{c}$.
- en ajoutant un terme de force capillaire $\sum_{i=1}^3 \mu_i \nabla c_i$ dans le second membre du bilan de quantité de mouvement (équations de Navier-Stokes).

Nous utilisons la même formulation des équations de Navier-Stokes que celle présentée dans le cas diphasique (cf section IV.1.2).

Un modèle de type Cahn-Hilliard/Navier-Stokes triphasique

Le modèle Cahn-Hilliard/Navier-Stokes que nous étudions est donc le suivant :

$$\left\{ \begin{array}{l} \frac{\partial c_i}{\partial t} + \mathbf{u} \cdot \nabla c_i = \text{div} \left(\frac{M_0}{\Sigma_i} \nabla \mu_i \right), \quad \forall i = 1, 2, 3, \\ \mu_i = \frac{4\Sigma_T}{\varepsilon} \sum_{j \neq i} \left(\frac{1}{\Sigma_j} (\partial_i F(\mathbf{c}) - \partial_j F(\mathbf{c})) \right) - \frac{3}{4} \varepsilon \Sigma_i \Delta c_i, \quad \forall i = 1, 2, 3, \\ \sqrt{\varrho(\mathbf{c})} \frac{\partial}{\partial t} (\sqrt{\varrho(\mathbf{c})} \mathbf{u}) + (\varrho(\mathbf{c}) \mathbf{u} \cdot \nabla) \mathbf{u} + \frac{\mathbf{u}}{2} \text{div} (\varrho(\mathbf{c}) \mathbf{u}) - \text{div} (2\eta(\mathbf{c}) D(\mathbf{u})) + \nabla p = \sum_{i=1}^3 \mu_i \nabla c_i + \varrho(\mathbf{c}) \mathbf{g}, \\ \text{div} \mathbf{u} = 0, \end{array} \right. \quad (\text{IV.28})$$

où le vecteur $\mathbf{g}$ représente la gravité; la densité et la viscosité sont définies par :

$$\varrho(\mathbf{c}) = \frac{\sum_{i=1}^3 \varrho_i h_\lambda(c_i - 0.5)}{\sum_{i=1}^3 h_\lambda(c_i - 0.5)} \quad \text{et} \quad \eta(\mathbf{c}) = \frac{\sum_{i=1}^3 \eta_i h_\lambda(c_i - 0.5)}{\sum_{i=1}^3 h_\lambda(c_i - 0.5)}, \quad (\text{IV.29})$$

avec ϱ_1 (resp. ϱ_2 , resp. ϱ_3) et η_1 (resp. η_2 , resp. η_3) les valeurs supposées constantes dans la phase 1 (resp. 2, resp. 3) et la fonction h_λ ($\lambda = 0.5$) définie par IV.6.

Toute solution du système (IV.28) vérifie les égalités de bilan suivantes :

– bilan du volume :

$$\frac{\partial}{\partial t} \left[\int_{\Omega} c_i dx \right] = \int_{\Gamma} \left[-c_i \mathbf{u} + \frac{M_0(\mathbf{c})}{\Sigma_i} \nabla \mu_i \right] \cdot \mathbf{n} ds.$$

– égalité d'énergie :

$$\begin{aligned} \frac{\partial}{\partial t} \left[\int_{\Omega} \frac{1}{2} \varrho(\mathbf{c}) |\mathbf{u}|^2 dx + \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}) \right] + \int_{\Gamma} \left(\frac{1}{2} \varrho(\mathbf{c}) |\mathbf{u}|^2 \right) \mathbf{u} \cdot \mathbf{n} ds + \int_{\Omega} 2\eta(\mathbf{c}) |D\mathbf{u}|^2 dx + \sum_{i=1}^3 \int_{\Omega} \frac{M_0(\mathbf{c})}{\Sigma_i} |\nabla \mu_i|^2 dx = \\ \int_{\Omega} \varrho(\mathbf{c}) \mathbf{g} \cdot \mathbf{u} dx + \int_{\Gamma} [2\eta(\mathbf{c}) D\mathbf{u} \cdot \mathbf{n} - p\mathbf{n}] \cdot \mathbf{u} ds + \sum_{i=1}^3 \int_{\Gamma} \frac{M_0(\mathbf{c})}{\Sigma_i} \mu_i \nabla \mu_i \cdot \mathbf{n} ds + \frac{3}{4} \sigma \varepsilon \sum_{i=1}^3 \int_{\Gamma} \Sigma_i \frac{\partial c_i}{\partial t} \nabla c_i \cdot \mathbf{n} ds. \end{aligned}$$

Existence de solutions faibles

Nous ajoutons au système précédent des conditions aux bords de type Neumann pour chaque paramètre d'ordre c_i et pour chaque potentiel chimique μ_i , c'est-à-dire, pour $i = 1, 2$ et 3 ,

$$\nabla c_i \cdot \mathbf{n} = 0 \text{ et } M_0 \nabla \mu_i \cdot \mathbf{n} = 0, \text{ sur } \Gamma, \quad (\text{IV.30})$$

ainsi que des conditions aux bords de type Dirichlet homogène pour la vitesse, c'est-à-dire

$$\mathbf{u} = 0 \text{ sur } \Gamma. \quad (\text{IV.31})$$

Remarque IV.15

Il est possible de considérer d'autres types de conditions aux limites. Dans les applications numériques présentées dans le chapitre 3, nous considérons des conditions aux bords de type glissement :

$$\mathbf{u} \cdot \mathbf{n} = 0 \text{ et } D\mathbf{u} \cdot \mathbf{n} \cdot \mathbf{t} = 0,$$

où $\mathbf{t}$ est le vecteur tangentiel au bord du domaine. Il est également possible, pour simuler l'injection de bulles, d'utiliser des conditions de type Dirichlet non homogène sur la vitesse combinée à des conditions aux bords mixtes du type (IV.16)-(IV.17) pour les paramètres d'ordre c_i .

Au vu des conditions aux bords (IV.30) et (IV.31), nous introduisons les espaces fonctionnels suivants :

$$\begin{aligned} \mathcal{V}^c &= \mathcal{V}^\mu = H^1(\Omega), \\ \mathcal{V}_S^c &= \{\mathbf{c} = (c_1, c_2, c_3) \in (H^1(\Omega))^3; \mathbf{c}(x) \in \mathcal{S} \text{ pour presque tout } x \in \Omega\}, \\ \mathcal{V}^u &= (H^1(\Omega))^d, \\ \mathcal{V}_0^u &= (H_0^1(\Omega))^d, \\ \mathcal{V}^p &= \{p \in L^2(\Omega), \int_{\Omega} p dx = 0\}. \end{aligned}$$

Finalement, nous supposons qu'à l'instant initial, nous avons

$$c_i(t=0) = c_i^0, \quad \text{et} \quad \mathbf{u}(t=0) = \mathbf{u}^0, \quad (\text{IV.32})$$

où $\mathbf{c}^0 = (c_1^0, c_2^0, c_3^0) \in \mathcal{V}_S^c$ et $\mathbf{u}^0 \in \mathcal{V}_0^u$ sont donnés.

Nous montrerons le théorème d'existence ci-dessous dans le chapitre VI en effectuant un passage à la limite dans les schémas numériques.

Théorème IV.16 (Existence de solution faible dans le cas homogène)

Supposons que les coefficients $(\Sigma_1, \Sigma_2, \Sigma_3)$ vérifient la condition (IV.14), la mobilité satisfait (IV.19), et que le potentiel de Cahn-Hilliard F satisfait la condition (IV.20). Supposons que les densités des trois fluides soient égales, c'est-à-dire $\varrho_1 = \varrho_2 = \varrho_3 = \varrho_0$, $\varrho_0 \in \mathbb{R}$. Considérons le problème (IV.28) avec la condition initiale (IV.32) et les conditions aux bords (IV.30)-(IV.31). Alors, il existe une solution faible $(\mathbf{c}, \boldsymbol{\mu}, \mathbf{u}, p)$ sur $[0, t_f[$ telle que

$$\mathbf{c} \in L^\infty(0, t_f; (H^1(\Omega))^3) \cap C^0([0, t_f]; (L^q(\Omega))^3), \text{ pour tout } q < 6,$$

$$\boldsymbol{\mu} \in L^2(0, t_f; (H^1(\Omega))^3),$$

$$\mathbf{u} \in L^\infty(0, t_f; (L^2(\Omega))^3) \cap L^2(0, t_f; (H^1(\Omega))^3),$$

$$\mathbf{c}(t, x) \in \mathcal{S}, \text{ pour presque tout } (t, x) \in [0, t_f] \times \Omega.$$

Chapitre V

Discrétisation du système de Cahn-Hilliard

Ce chapitre est dédié à l'étude de schémas numériques pour le système de Cahn-Hilliard (IV.9) avec les conditions aux bords (IV.16)-(IV.17) et la condition initiale (IV.18). La difficulté provient essentiellement de la non-convexité du potentiel de Cahn-Hilliard. Il est en effet souhaitable que le choix de la discrétisation en temps conduise à une estimation d'énergie, gage de stabilité. Cependant, le schéma implicite couramment utilisé ne permet pas de garantir la décroissance de l'énergie discrète, nous amenant ainsi à considérer d'autres discrétisations mieux adaptées à la forme des équations de Cahn-Hilliard.

Dans la section V.1, nous donnons les schémas numériques que nous utilisons pour approcher les solutions du système (IV.9). La discrétisation des termes non linéaires est exprimée de manière abstraite, et des conditions suffisantes pour garantir l'existence et la convergence des solutions discrètes sont fournies. Dans la section V.2, nous précisons différents choix possibles pour ces discrétisations et montrons leurs principales propriétés. Nous donnons dans la section V.5 la preuve des théorèmes d'existence et de convergence énoncés dans la section V.1. Notons que nous ne supposons pas l'existence d'une solution du problème continu : nous l'obtenons comme conséquence de la convergence du schéma numérique. Ainsi, nous donnons une nouvelle preuve du théorème IV.9 pour les conditions aux bords plus générales (IV.16)-(IV.17). La section V.3 illustre par des tests numériques les résultats obtenus, en particulier avec la simulation d'étalement d'une lentille piégée entre deux phases. Ces simulations nous permettent de conclure que la discrétisation semi-implicite en temps que nous proposons est un bon compromis entre précision et robustesse.

Il est important de noter que tous les résultats théoriques sont établis en supposant que la mobilité est non dégénérée (*i.e.* sa borne inférieure est strictement positive). Dans la section V.4, nous fournissons néanmoins une discrétisation stable permettant d'éviter certains problèmes numériques lorsque la mobilité est dégénérée (*i.e.* elle s'annule dans les phases pures, ou autrement dit aux triplets $(1, 0, 0)$, $(0, 1, 0)$ et $(0, 0, 1)$).

V.1 Discrétisation, existence et convergence des solutions approchées

Nous décrivons tout d'abord dans la section V.1.1 une semi-discrétisation en temps. La discrétisation en temps des termes non linéaires est formulée de manière très générale ; plusieurs choix particuliers sont donnés dans la section V.2. Dans la section V.1.2, nous donnons la discrétisation en espace qui est réalisée par une approximation de Galerkin et la méthode des éléments finis. Le problème discret est formulé dans un premier temps en utilisant les trois couples d'inconnues (c_{ih}^n, μ_{ih}^n) , $i = 1, 2, 3$. Nous montrons ensuite dans la section V.1.3 que ce problème peut être formulé de manière équivalente en utilisant seulement deux couples d'inconnues de notre choix, le troisième étant déduit *a posteriori*. Finalement, le reste de cette section est consacré à l'analyse du problème discret en temps et en espace.

Nous adoptons l'approche suivante :

- des estimations *a priori* sont obtenues à partir de l'égalité d'énergie donnée dans la section V.1.4.
- le problème discret non-linéaire est relié par homotopie à un problème linéaire. L'existence d'une solution approchée est alors déduite des estimations *a priori* mentionnées précédemment et de l'existence d'une solution au problème linéaire (en appliquant la théorie du degré topologique).
- la convergence des solutions approchées est obtenue à partir des estimations *a priori* en utilisant des résultats de compacité.

Les théorèmes d'existence et de convergence sont énoncés dans la section V.1.5 mais leur démonstration est reportée à la section V.5.

V.1.1 Discrétisation en temps

Soit $N \in \mathbb{N}^*$ et $t_f \in]0, +\infty[$. L'intervalle de temps $[0, t_f]$ est uniformément discrétisé avec un pas de temps fixe $\Delta t = \frac{t_f}{N}$. Pour $n \in \llbracket 0, N \rrbracket$, nous définissons $t_n = n\Delta t$.

Soit $n \in \mathbb{N}$. Nous supposons que les fonctions $(c_1^n, c_2^n, c_3^n) \in \mathcal{V}_{D,S}^c$ sont données. Nous utilisons une discrétisation semi-implicite en portant une attention particulière aux termes non linéaires.

Le schéma est écrit sous une forme très générale : pour $i = 1, 2, 3$,

$$\begin{cases} \frac{c_i^{n+1} - c_i^n}{\Delta t} = \operatorname{div} \left(\frac{M_0^{n+\alpha}}{\Sigma_i} \nabla \mu_i^{n+1} \right), \\ \mu_i^{n+1} = D_i^F(\mathbf{c}^n, \mathbf{c}^{n+1}) - \frac{3}{4} \varepsilon \Sigma_i \Delta c_i^{n+\beta}, \end{cases} \quad (\text{V.1})$$

où • $M_0^{n+\alpha} = M_0((1-\alpha)\mathbf{c}^n + \alpha\mathbf{c}^{n+1})$ avec $\alpha \in [0, 1]$,

• $c_i^{n+\beta} = (1-\beta)c_i^n + \beta c_i^{n+1}$ avec $\beta \in \left[\frac{1}{2}, 1\right]$,

• $D_i^F(\mathbf{a}^n, \mathbf{a}^{n+1}) = \frac{4\Sigma_T}{\varepsilon} \sum_{j \neq i} \left(\frac{1}{\Sigma_j} (d_i^F(\mathbf{a}^n, \mathbf{a}^{n+1}) - d_j^F(\mathbf{a}^n, \mathbf{a}^{n+1})) \right), \quad \forall (\mathbf{a}^n, \mathbf{a}^{n+1}) \in \mathcal{S}^2. \quad (\text{V.2})$

Les fonctions d_i^F représentent une discrétisation semi-implicite des dérivées partielles $\partial_{c_i} F$ de F .

Pour le moment, nous supposons seulement que

$$\forall \mathbf{c} \in \mathcal{S}, \quad D_i^F(\mathbf{c}, \mathbf{c}) = f_i^F(\mathbf{c}), \quad (\text{V.3})$$

pour assurer la consistance. Plusieurs choix possibles de discrétisation seront proposés et étudiés dans la section V.2.

Au vu de (IV.16) et (IV.17), les conditions aux bords discrètes sont, pour $i = 1, 2, 3$,

$$\begin{aligned} c_i^{n+1} = c_{iD} \quad & \text{et} \quad M_0 \nabla \mu_i^{n+1} \cdot \mathbf{n} = 0, \quad \text{sur } \Gamma_D^c, \\ \nabla c_i^{n+1} \cdot \mathbf{n} = 0 \quad & \text{et} \quad M_0 \nabla \mu_i^{n+1} \cdot \mathbf{n} = 0, \quad \text{sur } \Gamma_N^c. \end{aligned}$$

V.1.2 Discrétisation en espace

Pour la discrétisation en espace, nous utilisons une approximation de Galerkin et la méthode des éléments finis.

Soient $\mathcal{V}_h^c$ et $\mathcal{V}_h^\mu$ deux espaces d'approximation éléments finis de $\mathcal{V}^c$ et $\mathcal{V}^\mu$ respectivement. Puisque les paramètres d'ordre vérifient des conditions aux bords de Dirichlet non homogènes sur la frontière Γ_D^c , nous utilisons c_{ih}^0 comme un relèvement de c_{iD} dans $\mathcal{V}^c$ et nous supposons que les fonctions $c_{ih}^0 \in \mathcal{V}_h^c$ sont données pour tout $i \in \{1, 2, 3\}$, pour tout $h > 0$ de manière que

$$\mathbf{c}_h^0(x) \in \mathcal{S}, \quad \forall h > 0, \quad \text{pour presque tout } x \in \Omega \quad \text{et} \quad \|\mathbf{c}_h^0 - \mathbf{c}^0\|_{(\mathbf{H}^1(\Omega))^3} \xrightarrow{h \rightarrow 0} 0.$$

Ces fonctions c_{ih}^0 peuvent être obtenues à partir de c_i^0 par projection $\mathbf{H}^1(\Omega)$, ou comme c'est le cas en pratique, par interpolation éléments finis pourvu que c_i^0 soit assez régulière.

Nous pouvons maintenant définir les espaces d'approximation suivants :

$$\begin{aligned}\mathcal{V}_{Dh,0}^c &= \{\nu_h^c \in \mathcal{V}_h^c; \nu_h^c = 0 \text{ sur } \Gamma_D^c\}, \\ \mathcal{V}_{Dh}^{c_i} &= c_{ih}^0 + \mathcal{V}_{Dh,0}^c, \\ \mathcal{V}_{Dh,S}^c &= \{\mathbf{c}_h = (c_{1h}, c_{2h}, c_{3h}) \in \mathcal{V}_{Dh}^{c_1} \times \mathcal{V}_{Dh}^{c_2} \times \mathcal{V}_{Dh}^{c_3}; \mathbf{c}_h(x) \in \mathcal{S} \text{ pour presque tout } x \in \Omega\}.\end{aligned}$$

Les hypothèses générales requises sur les espaces d'approximation sont les suivantes :

$$\bullet 1 \in \mathcal{V}_h^c \quad \text{et} \quad 1 \in \mathcal{V}_h^\mu, \quad (\text{V.4})$$

$$\bullet \forall \nu^\mu \in \mathcal{V}^\mu, \quad \inf_{\nu_h^\mu \in \mathcal{V}_h^\mu} |\nu^\mu - \nu_h^\mu|_{H^1(\Omega)} \xrightarrow{h \rightarrow 0} 0 \quad \text{et} \quad \forall \nu^c \in \mathcal{V}_{D,0}^c, \quad \inf_{\nu_h^c \in \mathcal{V}_{Dh,0}^c} |\nu^c - \nu_h^c|_{H^1(\Omega)} \xrightarrow{h \rightarrow 0} 0, \quad (\text{V.5})$$

• il existe une constante strictement positive C indépendante de h telle que

$$\forall \nu^c \in \mathcal{V}^c, \quad \left| \Pi_0^{\mathcal{V}^c}(\nu^c) \right|_{H^1(\Omega)} \leq |\nu^c|_{H^1(\Omega)} \quad \text{et} \quad \forall \nu^\mu \in \mathcal{V}^\mu, \quad \left| \Pi_0^{\mathcal{V}^\mu}(\nu^\mu) \right|_{H^1(\Omega)} \leq C |\nu^\mu|_{H^1(\Omega)}, \quad (\text{V.6})$$

où $\Pi_0^{\mathcal{V}^\mu}$ est la projection $L^2(\Omega)$ sur $\mathcal{V}_h^\mu$,

$$\bullet \mathcal{V}_h^c \subset \mathcal{V}_h^\mu. \quad (\text{V.7})$$

Remarque V.1

L'hypothèse (V.6) est vraie, par exemple, lorsque nous considérons une famille de triangulation quasi-uniforme et les espaces d'approximation associés à des éléments finis conformes de Lagrange correspondants [EG04, p.72 (1.117)].

Supposons que $\mathbf{c}_h^n \in \mathcal{V}_{Dh,S}^c$ est donné, l'approximation de Galerkin du problème (V.1) au temps t_{n+1} s'écrit de la manière suivante :

Problème V.2 (Formulation avec trois paramètres d'ordre)

$$\begin{aligned} & \text{Trouver } (\mathbf{c}_h^{n+1}, \boldsymbol{\mu}_h^{n+1}) \in \mathcal{V}_{Dh}^{c_1} \times \mathcal{V}_{Dh}^{c_2} \times \mathcal{V}_{Dh}^{c_3} \times (\mathcal{V}_h^\mu)^3 \text{ tels que } \forall \nu_h^c \in \mathcal{V}_{Dh,0}^c, \forall \nu_h^\mu \in \mathcal{V}_h^\mu, \text{ nous avons, pour} \\ & i = 1, 2, 3, \\ & \left\{ \begin{aligned} \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx &= - \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} \nabla \mu_{ih}^{n+1} \cdot \nabla \nu_h^\mu dx, \\ \int_{\Omega} \mu_{ih}^{n+1} \nu_h^c dx &= \int_{\Omega} D_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \nu_h^c dx + \int_{\Omega} \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^{n+\beta} \cdot \nabla \nu_h^c dx, \end{aligned} \right. \quad (\text{V.8}) \\ & \text{où } M_{0h}^{n+\alpha} = M_0((1-\alpha)\mathbf{c}_h^n + \alpha\mathbf{c}_h^{n+1}) \text{ et } c_{ih}^{n+\beta} = (1-\beta)c_{ih}^n + \beta c_{ih}^{n+1}. \end{aligned}$$

Notons que nous ne cherchons pas $\mathbf{c}_h^{n+1}$ dans $\mathcal{V}_{Dh,S}^c$. La contrainte $c_{1h}^{n+1} + c_{2h}^{n+1} + c_{3h}^{n+1} = 1$ est naturellement imposée par la forme particulière des D_i^F dans le modèle (cf Théorème V.6).

Remarque V.3

L'hypothèse (V.4) autorise à prendre $\nu_h^\mu \equiv 1$ dans la première équation de (V.8). Cela permet d'obtenir, au niveau discret, l'exacte conservation du volume de chaque phase :

$$\int_{\Omega} c_{ih}^{n+1} dx = \int_{\Omega} c_{ih}^n dx, \quad \forall i \in \{1, 2, 3\}, \quad \forall n \in \llbracket 0, N-1 \rrbracket. \quad (\text{V.9})$$

V.1.3 Equivalence avec un système de deux équations couplées

En pratique, nous résolvons seulement les équations satisfaites par (c_1, c_2, μ_1, μ_2) . En effet, supposons que $\mathbf{c}_h^n \in \mathcal{V}_{Dh,S}^c$ est donné, alors le problème V.2 est équivalent au suivant :

Problème V.4 (Formulation avec deux paramètres d'ordre)

$$\begin{aligned} & \text{Trouver } (c_{1h}^{n+1}, c_{2h}^{n+1}, \mu_{1h}^{n+1}, \mu_{2h}^{n+1}) \in \mathcal{V}_{Dh}^{c_1} \times \mathcal{V}_{Dh}^{c_2} \times (\mathcal{V}_h^\mu)^2 \text{ tel que } \forall \nu_h^c \in \mathcal{V}_{Dh,0}^c, \forall \nu_h^\mu \in \mathcal{V}_h^\mu, \text{ nous avons, pour } \\ & i = 1 \text{ et } 2, \\ & \left\{ \begin{aligned} \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx &= - \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} \nabla \mu_{ih}^{n+1} \cdot \nabla \nu_h^\mu dx, \\ \int_{\Omega} \mu_{ih}^{n+1} \nu_h^c dx &= \int_{\Omega} D_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \nu_h^c dx + \int_{\Omega} \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^{n+\beta} \cdot \nabla \nu_h^c dx, \end{aligned} \right. \quad (\text{V.10}) \\ & \text{avec } \mathbf{c}_h^{n+1} = (c_{1h}^{n+1}, c_{2h}^{n+1}, 1 - c_{1h}^{n+1} - c_{2h}^{n+1}). \end{aligned}$$

Il reste ensuite à définir

$$c_{3h}^{n+1} = 1 - c_{1h}^{n+1} - c_{2h}^{n+1} \quad \text{et} \quad \mu_{3h}^{n+1} = - \left(\frac{\Sigma_3}{\Sigma_1} \mu_{1h}^{n+1} + \frac{\Sigma_3}{\Sigma_2} \mu_{2h}^{n+1} \right). \quad (\text{V.11})$$

Remarque V.5

Notons que dans la suite, dans les systèmes où seulement les inconnues $(c_{1h}^{n+1}, \mu_{1h}^{n+1}, c_{2h}^{n+1}, \mu_{2h}^{n+1})$ sont présentes, la notation $\mathbf{c}_h^{n+1}$ représente le vecteur $(c_{1h}^{n+1}, c_{2h}^{n+1}, 1 - c_{1h}^{n+1} - c_{2h}^{n+1})$.

Théorème V.6

Le problème (V.8) est équivalent au problème (V.10)-(V.11). En particulier, toute solution $(\mathbf{c}_h^{n+1}, \mu_h^{n+1})$ du problème V.2 satisfait

$$\sum_{i=1}^3 c_{ih}^{n+1} = 1 \quad \text{et} \quad \sum_{i=1}^3 \frac{\mu_{ih}^{n+1}}{\Sigma_i} = 0. \quad (\text{V.12})$$

Démonstration : Tout d'abord, en utilisant la définition de D_i^F donnée par (V.2), après avoir ré-ordonné les termes, nous trouvons (j et k les deux indices différents de i) :

$$\sum_{i=1}^3 \frac{1}{\Sigma_i} D_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) = \frac{4\Sigma_T}{\varepsilon} \sum_{i=1}^3 \left(\frac{1}{\Sigma_i} \left(\frac{1}{\Sigma_j} + \frac{1}{\Sigma_k} \right) - \frac{1}{\Sigma_i \Sigma_j} - \frac{1}{\Sigma_i \Sigma_k} \right) d_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) = 0. \quad (\text{V.13})$$

Supposons maintenant que le problème (V.10)-(V.11) est satisfait. Alors, en ajoutant les équations du système (V.10) pour $i = 1, 2$ et en utilisant (V.11) et (V.13), nous obtenons

$$\left\{ \begin{aligned} \int_{\Omega} \frac{(1 - c_{3h}^{n+1}) - (1 - c_{3h}^n)}{\Delta t} \nu_h^\mu dx &= - \int_{\Omega} M_{0h}^{n+\alpha} \nabla \left(-\frac{\mu_{3h}^{n+1}}{\Sigma_3} \right) \cdot \nabla \nu_h^\mu dx, \\ \int_{\Omega} \left(-\frac{\mu_{3h}^{n+1}}{\Sigma_3} \right) \nu_h^c dx &= \int_{\Omega} \left(-\frac{1}{\Sigma_3} D_3^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \right) \nu_h^c dx + \frac{3}{4} \varepsilon \int_{\Omega} \nabla (1 - c_{3h}^{n+\beta}) \cdot \nabla \nu_h^c dx. \end{aligned} \right.$$

Cela prouve que c_{3h}^{n+1} satisfait (V.8) pour $i = 3$.

Réciproquement, si nous supposons que (V.8) est satisfait, alors en ajoutant les équation pour $i = 1, 2, 3$, grâce à l'égalité (V.13), nous obtenons

$$\left\{ \begin{aligned} \int_{\Omega} \frac{S_h^{n+1} - S_h^n}{\Delta t} \nu_h^\mu dx &= - \int_{\Omega} M_{0h}^{n+\alpha} \nabla \Theta_h^{n+1} \cdot \nabla \nu_h^\mu dx \\ \int_{\Omega} \Theta_h^{n+1} \nu_h^c dx &= \frac{3}{4} \varepsilon \int_{\Omega} [(1 - \beta) \nabla S_h^n + \beta \nabla S_h^{n+1}] \cdot \nabla \nu_h^c dx, \end{aligned} \right. \quad (\text{V.14})$$

où $S_h^\ell = \sum_{i=1}^3 c_{ih}^\ell$ et $\Theta_h^\ell = \sum_{i=1}^3 \frac{\mu_{ih}^\ell}{\Sigma_i}$ pour $\ell = n$ et $\ell = n + 1$. Ces équations sont satisfaites pour tout $\nu_h^\mu \in \mathcal{V}_h^\mu$ et

pour tout $\nu_h^c \in \mathcal{V}_{Dh,0}^c$. En particulier, nous prenons $\nu_h^\mu = \Theta_h^{n+1}$ et $\nu_h^c = \frac{S_h^{n+1} - S_h^n}{\Delta t} \in \mathcal{V}_{Dh,0}^c$, de manière à ce

que le membre de gauche des deux équations (V.14) soit le même. En notant que

$$[(1 - \beta)\nabla S_h^n + \beta\nabla S_h^{n+1}] \cdot \nabla(S_h^{n+1} - S_h^n) = \frac{1}{2} \left(|\nabla S_h^{n+1}|^2 - |\nabla S_h^n|^2 + (2\beta - 1)|\nabla S_h^{n+1} - \nabla S_h^n|^2 \right),$$

nous obtenons finalement l'égalité

$$\frac{3}{8}\varepsilon \int_{\Omega} \left(|\nabla S_h^{n+1}|^2 - |\nabla S_h^n|^2 + (2\beta - 1)|\nabla S_h^{n+1} - \nabla S_h^n|^2 \right) dx + \Delta t \int_{\Omega} M_{0h}^{n+\alpha} |\nabla \Theta_h^{n+1}|^2 dx = 0. \quad (\text{V.15})$$

Puisque $S_h^n \equiv 1$ ($\mathbf{c}_h^n \in \mathcal{V}_{Dh,S}$), M_0 est positive et $\beta \geq \frac{1}{2}$, le membre de gauche de (V.15) est une somme de termes négatifs. En particulier, $\nabla S_h^{n+1} \equiv 0$ et $\nabla \Theta_h^{n+1} \equiv 0$. Donc, les fonctions S_h^{n+1} et Θ_h^{n+1} sont constantes. En ré-injectant ces constantes dans les équations du système (V.14), nous obtenons $S_h^{n+1} \equiv 1$ et $\Theta_h^{n+1} \equiv 0$. Ainsi, le couple $(\mathbf{c}_h^{n+1}, \boldsymbol{\mu}_h^{n+1})$ satisfait (V.12) et en conséquence le système (V.10)-(V.11). ■

V.1.4 Estimation d'énergie discrète

L'estimation d'énergie pour notre problème est obtenue par un calcul similaire à celui utilisé pour prouver l'équivalence entre les problèmes V.2 et V.4 dans la démonstration du théorème V.6.

Proposition V.7 (Egalité d'énergie discrète)

Soit $\mathbf{c}_h^n \in \mathcal{V}_{Dh,S}^c$. Supposons qu'il existe une solution $(\mathbf{c}_h^{n+1}, \boldsymbol{\mu}_h^{n+1})$ au problème V.2. Alors, l'égalité suivante est vérifiée :

$$\begin{aligned} \mathcal{F}_{\Sigma,\varepsilon}^{\text{triph}}(\mathbf{c}_h^{n+1}) - \mathcal{F}_{\Sigma,\varepsilon}^{\text{triph}}(\mathbf{c}_h^n) + \Delta t \sum_{i=1}^3 \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx \\ + \frac{3}{8}(2\beta - 1)\varepsilon \int_{\Omega} \sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 dx \\ = \frac{12}{\varepsilon} \int_{\Omega} [F(\mathbf{c}_h^{n+1}) - F(\mathbf{c}_h^n) - \mathbf{d}^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \cdot (\mathbf{c}_h^{n+1} - \mathbf{c}_h^n)] dx, \end{aligned} \quad (\text{V.16})$$

où $\mathbf{d}^F(\cdot, \cdot)$ est le vecteur $(d_i^F(\cdot, \cdot))_{i=1,2,3}$.

Démonstration : Tout d'abord, la définition de l'énergie libre triphasique $\mathcal{F}_{\Sigma,\varepsilon}^{\text{triph}}$ donnée par (IV.8) conduit à l'égalité

$$\mathcal{F}_{\Sigma,\varepsilon}^{\text{triph}}(\mathbf{c}_h^{n+1}) - \mathcal{F}_{\Sigma,\varepsilon}^{\text{triph}}(\mathbf{c}_h^n) = \int_{\Omega} \frac{12}{\varepsilon} (F(\mathbf{c}_h^{n+1}) - F(\mathbf{c}_h^n)) dx + \int_{\Omega} \sum_{i=1}^3 \frac{3}{8} \Sigma_i \varepsilon \left(|\nabla c_{ih}^{n+1}|^2 - |\nabla c_{ih}^n|^2 \right) dx. \quad (\text{V.17})$$

Par ailleurs, choisir $\nu_h^\mu = \mu_{ih}^{n+1}$ et $\nu_h^c = \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t}$ dans le système (V.8), donne pour $i = 1, 2, 3$,

$$\left\{ \begin{aligned} \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \mu_{ih}^{n+1} dx &= - \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx, \\ \int_{\Omega} \mu_{ih}^{n+1} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} dx &= \int_{\Omega} \frac{4\Sigma_T}{\varepsilon} \sum_{j \neq i} \left(\frac{1}{\Sigma_j} (d_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) - d_j^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1})) \right) \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} dx \\ &\quad + \int_{\Omega} \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^{n+\beta} \cdot \nabla \left(\frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \right) dx. \end{aligned} \right. \quad (\text{V.18})$$

Rappelons que $c_{ih}^{n+\beta} = (1 - \beta)c_{ih}^n + \beta c_{ih}^{n+1}$. L'égalité suivante est donc satisfaite

$$\nabla c_{ih}^{n+\beta} \cdot \nabla (c_{ih}^{n+1} - c_{ih}^n) = \frac{1}{2} \left(|\nabla c_{ih}^{n+1}|^2 - |\nabla c_{ih}^n|^2 + (2\beta - 1)|\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 \right).$$

En ré-ordonnant les termes et en utilisant $\sum_{i=1}^3 (c_{ih}^{n+1} - c_{ih}^n) = 0$ (théorème V.6), nous obtenons également

$$\sum_{i=1}^3 \sum_{j \neq i} \left(\frac{1}{\Sigma_j} (d_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) - d_j^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1})) \right) (c_{ih}^{n+1} - c_{ih}^n) = \frac{3}{\Sigma_T} \sum_{i=1}^3 (c_{ih}^{n+1} - c_{ih}^n) d_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}).$$

Ainsi, nous déduisons de (V.18) que

$$\begin{aligned} \Delta t \sum_{i=1}^3 \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx &= - \frac{12}{\varepsilon} \int_{\Omega} \sum_{i=1}^3 d_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) (c_{ih}^{n+1} - c_{ih}^n) dx \\ &\quad - \frac{3}{8} \varepsilon \int_{\Omega} \sum_{i=1}^3 \Sigma_i (|\nabla c_{ih}^{n+1}|^2 - |\nabla c_{ih}^n|^2) dx \\ &\quad - \frac{3}{8} (2\beta - 1) \varepsilon \int_{\Omega} \sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 dx. \end{aligned} \quad (\text{V.19})$$

La conclusion est obtenue en ajoutant (V.17) et (V.19). ■

Remarque V.8

Bien que les coefficients Σ_i ne soient pas nécessairement positifs, les deux termes $\sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2$ et $\sum_{i=1}^3 \frac{|\nabla \mu_{ih}^{n+1}|^2}{\Sigma_i}$, présents dans le second membre de l'équation (V.16), sont positifs lorsque la condition (IV.14) est vérifiée par les coefficients $(\Sigma_1, \Sigma_2, \Sigma_3)$. En effet, dans ce cas, la proposition IV.5 montre que

$$\sum_{i=1}^3 \frac{|\nabla \mu_{ih}^{n+1}|^2}{\Sigma_i} = \sum_{i=1}^3 \Sigma_i \frac{|\nabla \mu_{ih}^{n+1}|^2}{\Sigma_i^2} \geq \underline{\Sigma} \sum_{i=1}^3 \frac{|\nabla \mu_{ih}^{n+1}|^2}{\Sigma_i^2} \geq 0, \quad \text{puisque} \quad \sum_{i=1}^3 \frac{\nabla \mu_{ih}^{n+1}}{\Sigma_i} = 0 \quad (\text{théorème V.6}),$$

et

$$\sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 \geq \underline{\Sigma} \sum_{i=1}^3 |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 \geq 0, \quad \text{puisque} \quad \sum_{i=1}^3 \nabla (c_{ih}^{n+1} - c_{ih}^n) = 0 \quad (\text{théorème V.6}).$$

L'égalité (V.16) est une version discrète de l'égalité d'énergie (IV.11) satisfaite par les solutions $(\mathbf{c}, \boldsymbol{\mu})$ du modèle de Cahn-Hilliard (IV.9) :

$$\frac{d}{dt} [\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c})] = - \int_{\Omega} \sum_{i=1}^3 \frac{M_0(\mathbf{c})}{\Sigma_i} |\nabla \mu_i|^2 dx.$$

Cette égalité montre en particulier que l'énergie associée aux solutions du système (IV.9) décroît avec le temps.

Au niveau discret, l'égalité d'énergie (V.16) devrait fournir non seulement la décroissance de l'énergie discrète mais aussi les premières estimations *a priori* très utiles pour prouver l'existence des solutions approchées et leur convergence vers une solution faible du problème (IV.9). Cependant, deux termes additionnels apparaissent dans la contre-partie discrète de (IV.11) et, en conséquence, la validité de la décroissance de l'énergie et des estimations *a priori* dépend du signe de ces termes :

- le dernier terme du membre de gauche de (V.16) est un terme standard de diffusion numérique dû à la discrétisation en temps de “ Δc_i ” de la seconde équation de (IV.9). Ce terme a le “bon signe” puisque $\beta \geq 0.5$ (cf remarque V.8) et peut être supprimé en prenant $\beta = 0.5$.
- le membre de droite de l'égalité (V.16) contient la discrétisation en temps $\mathbf{d}^F$ des termes non linéaires et, en conséquence, son signe dépend des choix particuliers de $\mathbf{d}^F$.

Ainsi, le choix de la discrétisation en temps $\mathbf{d}^F$ des termes non linéaires peut être guidé par une étude du membre de droite de l'égalité (V.16). La situation la plus simple est celle où $\mathbf{d}^F$ est tel que ce terme soit nul. Dans ce

cas, l'égalité d'énergie discrète est similaire à celle obtenue au niveau continu. Lorsque le membre de droite a le "bon signe", *i.e.* est négatif, il est encore possible de l'éliminer pour obtenir une inégalité d'énergie. Plus généralement, il est suffisant de pouvoir contrôler le membre de droite de (V.16) pour obtenir les estimations *a priori* désirées (*cf* section V.2.2). C'est la raison pour laquelle, dans la section suivante, les hypothèses sur la discrétisation des termes non linéaires sont données sous la forme d'estimations faisant intervenir les termes de l'égalité d'énergie (V.16). Dans les deux théorèmes V.9 et V.10, ces hypothèses sont utilisées pour obtenir des estimations sur les solutions approchées $(\mathbf{c}_h^{n+1}, \boldsymbol{\mu}_h^{n+1})$ (dans des normes appropriées). Le point clé est que dans le théorème d'existence, ces bornes peuvent dépendre de la solution approchée $\mathbf{c}_h^n$ au temps précédent, du pas de temps Δt ou du pas du maillage h (toutes ces quantités étant fixées) alors que, dans le théorème de convergence, il est crucial que ces estimations *a priori* soient indépendantes du pas de temps Δt et du pas de maillage h . Les différentes hypothèses seront validées pour tous les schémas présentés dans la section V.2.

V.1.5 Théorème d'existence et de convergence

Cette section est dédiée aux énoncés des théorèmes d'existence et de convergence des solutions approchées, dont les démonstrations seront données dans la section V.5. Tout d'abord, nous donnons les hypothèses générales sur la discrétisation des termes non-linéaire $\mathbf{d}^F : \mathbb{R}^3 \times \mathbb{R}^3 \rightarrow \mathbb{R}^3$. La fonction $\mathbf{d}^F$ est de classe $\mathcal{C}^1(\mathbb{R}^3 \times \mathbb{R}^3)$ et satisfait une hypothèse de croissance polynomiale : il existe une constante $B_1 \geq 0$ et un réel p tels que $2 \leq p < +\infty$ si $d = 2$ ou $p = 6$ si $d = 3$ et

$$\begin{aligned} \forall i \in \{1, 2, 3\}, \forall (\mathbf{a}^n, \mathbf{a}^{n+1}) \in \mathcal{S}^2, \quad & |d_i^F(\mathbf{a}^n, \mathbf{a}^{n+1})| \leq B_1 \left(1 + |\mathbf{a}^n|^{p-1} + |\mathbf{a}^{n+1}|^{p-1}\right), \\ & |D(d_i^F(\mathbf{a}^n, \cdot))(\mathbf{a}^{n+1})| \leq B_1 \left(1 + |\mathbf{a}^n|^{p-2} + |\mathbf{a}^{n+1}|^{p-2}\right). \end{aligned} \quad (\text{V.20})$$

Théorème V.9 (Existence de solutions discrètes)

Soit $\mathbf{c}_h^n \in \mathcal{V}_{Dh, \mathcal{S}}^c$. Nous supposons que :

- les coefficients $(\Sigma_1, \Sigma_2, \Sigma_3)$ satisfont (IV.14), la mobilité satisfait (IV.19), et que le potentiel de Cahn-Hilliard F satisfait (IV.20),
- la discrétisation des termes non-linéaires $\mathbf{d}^F$ satisfait (V.20) et la propriété suivante : il existe $K_1^{\mathbf{c}_h^n} > 0$ (pouvant dépendre de $\mathbf{c}_h^n$) tel que :

$$\int_{\Omega} [F(\mathbf{a}_h^{n+1}) - F(\mathbf{c}_h^n) - \mathbf{d}^F(\mathbf{c}_h^n, \mathbf{a}_h^{n+1}) \cdot (\mathbf{a}_h^{n+1} - \mathbf{c}_h^n)] dx \leq K_1^{\mathbf{c}_h^n}, \quad \forall \mathbf{a}_h^{n+1} \in \mathcal{V}_{Dh, \mathcal{S}}^c. \quad (\text{V.21})$$

Alors, il existe au moins une solution $(\mathbf{c}_h^{n+1}, \boldsymbol{\mu}_h^{n+1})$ au problème V.2.

Pour tout $N \in \mathbb{N}$, nous pouvons maintenant introduire les fonctions du temps $t \in [0, t_f]$ suivantes :

$$\underline{c}_{ih}^N(t, \cdot) = c_{ih}^n(\cdot), \quad \text{si } t \in]t_n, t_{n+1}[, \quad (\text{V.22})$$

$$\bar{c}_{ih}^N(t, \cdot) = c_{ih}^{n+1}(\cdot), \quad \text{si } t \in]t_n, t_{n+1}[, \quad (\text{V.23})$$

$$c_{ih}^N(t, \cdot) = \frac{t_{n+1} - t}{\Delta t} c_{ih}^n(\cdot) + \frac{t - t_n}{\Delta t} c_{ih}^{n+1}(\cdot), \quad \text{si } t \in]t_n, t_{n+1}[. \quad (\text{V.24})$$

Pour les potentiels chimiques, nous introduisons les fonctions constantes par morceaux en temps suivantes : pour tout $N \in \mathbb{N}$, soit :

$$\mu_{ih}^N(t, \cdot) = \mu_{ih}^{n+1}(\cdot), \quad \text{si } t \in]t_n, t_{n+1}[. \quad (\text{V.25})$$

Théorème V.10 (Théorème de convergence)

Supposons que les hypothèses du théorème V.9 soient satisfaites, de manière qu'une solution $(\mathbf{c}_h^N, \boldsymbol{\mu}_h^N)$ au problème V.2 existe pour tout $N \in \mathbb{N}^*$ et pour tout $h > 0$. Supposons que $\beta \in]\frac{1}{2}, 1]$, que la propriété de consistance (V.3) soit vérifiée et qu'il existe deux constantes $C > 0$ et $\Delta t_0 > 0$ telles que pour tout $\Delta t \leq \Delta t_0$ et pour tout $n \in \llbracket 0, N-1 \rrbracket$,

$$\begin{aligned} & \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^{n+1}) - \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^n) \\ & + C \left[\Delta t \sum_{i=1}^3 \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx + \frac{3}{8} (2\beta - 1) \varepsilon \int_{\Omega} \sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 dx \right] \leq 0. \end{aligned} \quad (\text{V.26})$$

Considérons le problème (IV.9) avec la condition initiale (IV.18) et les conditions aux bords (IV.17). Alors, il existe une solution faible $(\mathbf{c}, \boldsymbol{\mu})$ définie sur $[0, t_f]$ telle que

$$\mathbf{c} \in L^\infty(0, t_f; (H^1(\Omega))^3) \cap C^0([0, t_f]; (L^q(\Omega))^3), \text{ pour tout } q < 6,$$

$$\boldsymbol{\mu} \in L^2(0, t_f; (H^1(\Omega))^3),$$

$$\mathbf{c}(t, x) \in \mathcal{S}, \text{ pour presque tout } (t, x) \in [0, t_f] \times \Omega,$$

et pour toute suite $(h_K)_{K \in \mathbb{N}^*}$ telle que $h_K \xrightarrow{K \rightarrow +\infty} 0$, les suites $(\mathbf{c}_{h_K}^N)_{(N, K) \in (\mathbb{N}^*)^2}$ et $(\boldsymbol{\mu}_{h_K}^N)_{(N, K) \in (\mathbb{N}^*)^2}$, définie par (V.8), satisfont, à sous-suites près, les convergences suivantes lorsque $\min(N, K) \rightarrow +\infty$:

$$\mathbf{c}_{h_K}^N \rightarrow \mathbf{c} \quad \text{dans } C^0(0, t_f, (L^q)^3) \text{ fort}, \quad \text{pour tout } q < 6, \quad (\text{V.27})$$

$$\boldsymbol{\mu}_{h_K}^N \rightharpoonup \boldsymbol{\mu} \quad \text{dans } L^2(0, t_f, (H^1)^3) \text{ faible}. \quad (\text{V.28})$$

Remarque V.11

Dans le théorème V.10, nous supposons que $\frac{1}{2} < \beta \leq 1$. En effet, le dernier terme dans le membre de gauche de l'égalité (V.26) (qui disparaît dans le cas où β est égal à $\frac{1}{2}$) est crucial dans les estimations des résidus (cf section V.5.2) et dans la démonstration de l'estimation d'énergie pour le schéma implicite (cf section V.2.2).

Remarque V.12

Sous des hypothèses supplémentaires sur la Hessienne du potentiel de Cahn-Hilliard F , il est prouvé dans [BL06] que le modèle (IV.9) a une unique solution faible. Dans ce cas, nous pouvons conclure que les convergences énoncées dans le théorème V.10 pour des sous-suites sont vérifiées par les suites entières $(\mathbf{c}_{h_K}^N, \boldsymbol{\mu}_{h_K}^N)$.

V.2 Différentes discrétisations pour les termes non linéaires

Dans cette section, nous présentons différents choix possibles pour la discrétisation des termes non-linéaires $\mathbf{d}^F$. Puisque l'expression (IV.21) du potentiel de Cahn-Hilliard triphasique F fournit une décomposition naturelle : $F = F_0 + P$, nous allons choisir des discrétisations de la forme $d_i^F = d_i^{F_0} + d_i^P$ où $d_i^{F_0}$ et d_i^P représentent une discrétisation de $\partial_{c_i} F_0$ et de $\partial_{c_i} P$ respectivement. Nous donnons trois choix possibles de discrétisation de la contribution de F_0 dans les sections V.2.2, V.2.3 et V.2.4, et une discrétisation semi-implicite de la contribution de P dans la section V.2.5. Dans chacune de ces sections, les estimations du membre de droite de (V.16) sont prouvées. Les résultats sont rassemblés dans la section V.2.6 pour obtenir l'existence des solutions approchées et leur convergence vers une solution faible de (IV.9). Finalement dans la section V.2.7, nous montrons que la propriété de consistance algébrique (cf section IV.2.1) a un équivalent discret en identifiant les schémas qui sont obtenus lorsque seulement deux phases sont présentes.

V.2.1 Remarques préliminaires

La relation $c_1 + c_2 + c_3 = 1$ permet d'obtenir une expression équivalente de F_0 sur l'hyperplan $\mathcal{S}$. Cette expression fait intervenir le potentiel de Cahn-Hilliard diphasique f dont nous rappelons la définition :

$$f(x) = x^2(1-x)^2, \quad \forall x \in \mathbb{R}. \quad (\text{V.29})$$

En effet, la fonction définie par :

$$\hat{F}_0(\mathbf{c}) = \sum_{i=1}^3 \frac{\Sigma_i}{2} f(c_i), \quad \forall \mathbf{c} \in \mathbb{R}^3, \quad (\text{V.30})$$

est égale à F_0 sur l'hyperplan $\mathcal{S}$:

$$\hat{F}_0(\mathbf{c}) = F_0(\mathbf{c}), \quad \forall \mathbf{c} \in \mathcal{S}.$$

Ces deux expressions différentes peuvent être utilisées de manière équivalente puisque nous pouvons facilement prouver que :

$$\nabla F_0(\mathbf{c}) \cdot \boldsymbol{\xi} = \nabla \hat{F}_0(\mathbf{c}) \cdot \boldsymbol{\xi}, \quad \forall (\mathbf{c}, \boldsymbol{\xi}) \in \mathcal{S}^2, \quad (\text{V.31})$$

et en conséquence, au vu de l'expression (IV.10), nous avons :

$$f_i^{\hat{F}_0}(\mathbf{c}) = f_i^{F_0}(\mathbf{c}), \quad \forall \mathbf{c} \in \mathcal{S}.$$

V.2.2 Discrétisation implicite de la contribution de F_0

La discrétisation implicite correspond à la définition suivante :

$$\mathbf{d}^{F_0}(\mathbf{a}^n, \mathbf{a}^{n+1}) = \nabla F_0(\mathbf{a}^{n+1}), \quad \forall (\mathbf{a}^n, \mathbf{a}^{n+1}) \in \mathcal{S}^2. \quad (\text{V.32})$$

Nous prouvons respectivement que la contribution de F_0 à l'estimation d'énergie satisfait l'estimation (V.21) du théorème V.9 et l'estimation (V.26) du théorème V.10 lorsque nous utilisons le schéma implicite (V.32). Notons que nous supposons ici que $\Sigma_i > 0$, $\forall i \in \{1, 2, 3\}$, c'est-à-dire que nous nous plaçons dans le cas d'une situation d'étalement partiel.

Dans les situations d'étalement total (*i.e.* lorsque l'un des Σ_i est négatif), la démonstration des théorèmes d'existence et de convergence, lorsque le schéma implicite (V.32) est utilisé, est encore un problème ouvert. Nous observons, dans les expérimentations numériques, que dans ce cas la méthode de linéarisation de Newton (utilisée pour la résolution du problème V.4) peut ne pas converger (*cf* table V.4 de la section V.3).

Existence de solutions discrètes

Nous commençons par montrer que, dans le cas où tous les Σ_i sont positifs, l'hypothèse (V.21) du théorème V.9 est satisfaite. La démonstration utilise l'expression (V.30) de F_0 (valable sur l'hyperplan $\mathcal{S}$) et la remarque préliminaire de la section V.2.1.

Proposition V.13

Soit $\mathbf{c}_h^n \in \mathcal{V}_{Dh, \mathcal{S}}^c$. Supposons que $\forall i \in \{1, 2, 3\}$, $\Sigma_i > 0$. Alors, il existe $K_1^{\mathbf{c}_h^n} > 0$ (dépendant éventuellement de $\mathbf{c}_h^n$) tel que :

$$\int_{\Omega} [F_0(\mathbf{a}_h^{n+1}) - F_0(\mathbf{c}_h^n) - \nabla F_0(\mathbf{a}_h^{n+1}) \cdot (\mathbf{a}_h^{n+1} - \mathbf{c}_h^n)] dx \leq K_1^{\mathbf{c}_h^n}, \quad \forall \mathbf{a}_h^{n+1} \in \mathcal{V}_{Dh, \mathcal{S}}^c. \quad (\text{V.33})$$

Démonstration : Commençons par une inégalité élémentaire. Rappelons que la fonction f est définie par (V.29), et soit $y \in \mathbb{R}$ fixé. Nous définissons une fonction auxiliaire g par

$$g(x) = f(x) - f(y) - f'(x)(x - y).$$

Cette fonction est polynomiale d'ordre 4 et a un coefficient dominant négatif. Elle admet donc un maximum que nous noterons x_0 . Ce maximum x_0 dépend *a priori* de y mais satisfait $g'(x_0) = 0$ *i.e.* $-f''(x_0)(x_0 - y) = 0$. En conséquence, nous avons seulement deux cas possibles, ou bien $x_0 = y$ ou bien x_0 est solution de l'équation du second degré : $f''(x) = 0$. Dans le cas où $x_0 = y$, nous avons, pour tout $x \in \mathbb{R}$, $g(x) \leq g(y) = 0$. Dans le second cas, x_0 est indépendant de y et nous obtenons $g(x) \leq f(x_0) - f(y) - f'(x_0)(x_0 - y)$. Ainsi, dans tous les cas, en posant $C_1 = |f(x_0) - f'(x_0)x_0|$ et $C_2 = |f'(x_0)|$, nous avons

$$f(x) - f(y) - f'(x)(x - y) \leq C_1 + C_2|y| + |f(y)|, \quad \forall x \in \mathbb{R}, \quad (\text{V.34})$$

où C_1 et C_2 sont deux constantes indépendantes de x et y .

En combinant (V.30) et (V.34), puisque tous les Σ_i sont positifs, nous avons, pour tout $\mathbf{a}_h^{n+1} \in \mathcal{V}_{Dh, \mathcal{S}}^c$,

$$\begin{aligned} & \int_{\Omega} \hat{F}_0(\mathbf{a}_h^{n+1}) - \hat{F}_0(\mathbf{c}_h^n) - \nabla \hat{F}_0(\mathbf{a}_h^{n+1}) \cdot (\mathbf{a}_h^{n+1} - \mathbf{c}_h^n) dx \\ & \leq C_1 |\Omega| \sum_{i=1}^3 \frac{\Sigma_i}{2} + C_2 \sum_{i=1}^3 \frac{\Sigma_i}{2} \int_{\Omega} |c_{ih}^n| dx + \sum_{i=1}^3 \frac{\Sigma_i}{2} \int_{\Omega} |f(c_{ih}^n)| dx := K_1^{\mathbf{c}_h^n}. \end{aligned}$$

La conclusion est alors obtenue grâce à l'égalité (V.31). Le membre de droite $K_1^{\mathbf{c}_h^n}$ dépend uniquement de Σ et $\mathbf{c}_h^n$. Nous terminons en remarquant que $c_{ih}^n \in \mathcal{V}^c \subset H^1(\Omega) \subset L^q(\Omega)$ pour tout $q \leq 6$ et que f est un polynôme d'ordre 4. Ainsi, il existe une constante C indépendante de $\mathbf{c}_h^{n+1}$ et $\mathbf{c}_h^n$ telle que

$$K_1^{\mathbf{c}_h^n} \leq C \left(1 + |\mathbf{c}_h^n|_{H^1(\Omega)}^4 \right).$$

Ceci conclut la démonstration. ■

D'après le théorème V.9, l'existence de solutions approchées est obtenue comme un corollaire immédiat de la proposition précédente.

Corollaire V.14

Soit $\mathbf{c}_h^n \in \mathcal{V}_{Dh,S}^c$. Supposons que $\forall i \in \{1, 2, 3\}$, $\Sigma_i > 0$ et que la mobilité vérifie l'hypothèse (IV.19). Alors, il existe une solution au problème V.2 où $\mathbf{d}^F$ est défini par (V.32).

Convergence des solutions approchées

L'estimation établie dans la proposition V.13 est valable pour tout pas de temps mais n'est pas suffisante pour montrer le théorème de convergence. Dans cette section, nous donnons une autre estimation (qui correspond à l'hypothèse (V.26) du théorème V.10) valide seulement pour des pas de temps suffisamment petits.

Proposition V.15

Supposons que $\forall i \in \{1, 2, 3\}$, $\Sigma_i > 0$ et que la mobilité vérifie (IV.19). Alors pour tout pas de temps Δt tel que $\Delta t \leq \Delta t_0 = \frac{(2\beta - 1)\varepsilon^3}{24M_2}$, nous obtenons

$$\begin{aligned} \mathcal{F}_{\Sigma,\varepsilon}^{\text{triph}}(\mathbf{c}_h^{n+1}) - \mathcal{F}_{\Sigma,\varepsilon}^{\text{triph}}(\mathbf{c}_h^n) + \frac{\Delta t}{2} \sum_{i=1}^3 \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx \\ + \frac{3}{16} \varepsilon (2\beta - 1) \int_{\Omega} \sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 dx \leq 0. \end{aligned} \quad (\text{V.35})$$

Démonstration : Considérons la fonction f définie par (V.29). Puisque $\inf_{\mathbb{R}} f'' = -1$, nous obtenons

$$f(x) - f(y) - f'(x)(x - y) \leq \frac{(x - y)^2}{2}, \quad \forall x \in \mathbb{R}, \forall y \in \mathbb{R}.$$

Puisque tous les Σ_i sont positifs, nous déduisons de l'inégalité ci-dessus que

$$\begin{aligned} \hat{F}_0(\mathbf{c}_h^{n+1}) - \hat{F}_0(\mathbf{c}_h^n) - \nabla \hat{F}_0(\mathbf{c}_h^{n+1}) \cdot (\mathbf{c}_h^{n+1} - \mathbf{c}_h^n) &= \sum_{i=1}^3 \frac{\Sigma_i}{2} (f(c_{ih}^{n+1}) - f(c_{ih}^n) - f'(c_{ih}^{n+1})(c_{ih}^{n+1} - c_{ih}^n)) \\ &\leq \sum_{i=1}^3 \frac{\Sigma_i}{4} |c_{ih}^{n+1} - c_{ih}^n|^2. \end{aligned}$$

D'après l'égalité (V.31), le membre de gauche de l'inégalité ci-dessus est exactement l'intégrande du membre de droite de l'inégalité d'énergie (V.16). Nous obtenons alors l'estimation

$$\begin{aligned} \mathcal{F}_{\Sigma,\varepsilon}^{\text{triph}}(\mathbf{c}_h^{n+1}) - \mathcal{F}_{\Sigma,\varepsilon}^{\text{triph}}(\mathbf{c}_h^n) + \Delta t \sum_{i=1}^3 \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx \\ \leq \frac{3}{\varepsilon} \int_{\Omega} \sum_{i=1}^3 \Sigma_i |c_{ih}^{n+1} - c_{ih}^n|^2 dx - \frac{3}{8} \varepsilon (2\beta - 1) \int_{\Omega} \sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 dx. \end{aligned} \quad (\text{V.36})$$

Il reste maintenant à estimer le terme $\sum_{i=1}^3 \int_{\Omega} \Sigma_i |c_{ih}^{n+1} - c_{ih}^n|^2 dx$. Nous prenons $\nu_h^\mu = \Sigma_i (c_{ih}^{n+1} - c_{ih}^n)$ comme fonction test dans la première équation de (V.8) (remarquons que $\nu_h^\mu \in \mathcal{V}_h^\mu$ puisque $\mathcal{V}_{Dh,0}^c \subset \mathcal{V}_h^\mu$ (hypothèse (V.7))). Donc, nous obtenons

$$\int_{\Omega} \Sigma_i \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} (c_{ih}^{n+1} - c_{ih}^n) dx = - \int_{\Omega} \Sigma_i \frac{M_{0h}^{n+\alpha}}{\Sigma_i} \nabla \mu_{ih}^{n+1} \cdot \nabla (c_{ih}^{n+1} - c_{ih}^n) dx.$$

Nous ajoutons ensuite ces équations pour $i = 1, 2, 3$, et appliquons le corollaire IV.6, pour obtenir

$$\int_{\Omega} \sum_{i=1}^3 \Sigma_i |c_{ih}^{n+1} - c_{ih}^n|^2 dx \leq \frac{\Delta t}{2} \int_{\Omega} M_{0h}^{n+\alpha} \left[\frac{\varepsilon}{3} \sum_{i=1}^3 \frac{|\nabla \mu_{ih}^{n+1}|^2}{\Sigma_i} + \frac{3}{\varepsilon} \sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 \right] dx.$$

Puisque la mobilité est majorée (cf équation (IV.19)), nous avons

$$\frac{3}{\varepsilon} \int_{\Omega} \sum_{i=1}^3 \Sigma_i |c_{ih}^{n+1} - c_{ih}^n|^2 dx \leq \frac{\Delta t}{2} \sum_{i=1}^3 \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx + \frac{9M_2\Delta t}{2\varepsilon^2} \sum_{i=1}^3 \int_{\Omega} \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 dx. \quad (\text{V.37})$$

Ainsi, en combinant l'inégalité (V.36) et (V.37), nous avons finalement

$$\begin{aligned} \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^{n+1}) - \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^n) + \Delta t \sum_{i=1}^3 \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx &\leq \frac{\Delta t}{2} \sum_{i=1}^3 \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx \\ &\quad + \frac{9M_2\Delta t}{2\varepsilon^2} \sum_{i=1}^3 \int_{\Omega} \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 dx \\ &\quad - \frac{3}{8}\varepsilon(2\beta - 1) \int_{\Omega} \sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 dx. \end{aligned}$$

La conclusion est obtenue en utilisant le fait que $\Delta t \leq \frac{(2\beta - 1)\varepsilon^3}{24M_2}$. ■

Corollaire V.16

Sous les hypothèses du corollaire V.14, les conclusions du théorème V.10 sont satisfaites lorsque $\mathbf{d}^F$ est défini par le schéma implicite (V.32).

V.2.3 Discrétisation convexe-concave de la contribution de F_0

Dans la section V.2.2, nous avons remarqué que le schéma implicite (V.32) ne garantissait la décroissance de l'énergie que pour des pas de temps suffisamment petits. De plus, les résultats n'étaient valables que dans le cas de situations d'étalement partiel. Pour s'affranchir de ces problèmes, dans cette section, nous cherchons une discrétisation $\mathbf{d}^{F_0}$ telle que :

$$F_0(\mathbf{a}^{n+1}) - F_0(\mathbf{a}^n) - \mathbf{d}^{F_0}(\mathbf{a}^n, \mathbf{a}^{n+1}) \cdot (\mathbf{a}^{n+1} - \mathbf{a}^n) \leq 0, \quad \forall (\mathbf{a}^n, \mathbf{a}^{n+1}) \in \mathcal{S}. \quad (\text{V.38})$$

Supposons un instant que $\mathbf{d}^{F_0}$ soit défini par la discrétisation implicite (V.32) ; l'inégalité (V.38) serait vérifiée si la fonction F_0 était convexe sur l'hyperplan $\mathcal{S}$. De la même manière, si l'on utilisait une discrétisation explicite l'inégalité (V.38) serait vérifiée si la fonction F_0 était concave sur l'hyperplan $\mathcal{S}$. Malheureusement, le potentiel F_0 n'est ni convexe ni concave, néanmoins ces remarques fournissent un moyen naturel (cf [Lap06]) d'obtenir une discrétisation $\mathbf{d}^{F_0}$ qui satisfait (V.38) si la fonction F_0 se décompose comme la somme d'une fonction convexe et d'une fonction concave. En effet, si $F_0 = F_0^+ + F_0^-$ avec F_0^+ convexe et F_0^- concave alors nous pouvons définir

$$\mathbf{d}^{F_0}(\mathbf{a}^n, \mathbf{a}^{n+1}) = \nabla F_0^+(\mathbf{a}^{n+1}) + \nabla F_0^-(\mathbf{a}^n). \quad (\text{V.39})$$

Le potentiel de Cahn-Hilliard diphasique a une décomposition convexe-concave naturelle (cf [Eyr98]) :

$$f(x) = \underbrace{\left(x - \frac{1}{2}\right)^4}_{:=f^+(x)} + \underbrace{\frac{1}{16}(1 - 2(2x - 1)^2)}_{:=f^-(x)}. \quad (\text{V.40})$$

Cette décomposition mène facilement à une décomposition convexe-concave de $\hat{F}_0$ en utilisant les définitions suivantes :

$$\begin{aligned} F_0^+(\mathbf{c}) &= \sum_{i=1}^3 \frac{\Sigma_i^+}{2} f^+(c_i) - \sum_{i=1}^3 \frac{\Sigma_i^-}{2} f^-(c_i) \\ F_0^-(\mathbf{c}) &= \sum_{i=1}^3 \frac{\Sigma_i^+}{2} f^-(c_i) - \sum_{i=1}^3 \frac{\Sigma_i^-}{2} f^+(c_i), \end{aligned}$$

où $\Sigma_i^+ = \max(\Sigma_i, 0)$ et $\Sigma_i^- = -\min(\Sigma_i, 0)$.

Puisque F_0 et $\hat{F}_0$ coïncident sur l'hyperplan $\mathcal{S}$ (cf (V.31)), l'inégalité (V.38) est satisfaite et par suite les hypothèses (V.21) du théorème V.9 et (V.26) du théorème V.10 sont satisfaites pour la contribution de F_0 lorsque l'on utilise la discrétisation convexe-concave (V.39) (cf section V.2.6 pour plus de détails). Ces hypothèses sont satisfaites pour tout pas de temps Δt et même dans le cas d'étalement total.

V.2.4 Discrétisation semi-implicite de la contribution de F_0

Le schéma convexe-concave présenté dans la section V.2.3 garantit la décroissance de l'énergie pour tout pas de temps et même dans le cas de situations d'étalement total. Cependant, il souffre d'un important manque de précision (cf figure V.3, V.7 et V.15 dans la section V.3). Ceci est certainement dû au fait que la discrétisation convexe-concave partage inégalement les deux parties du potentiel de Cahn-Hilliard qui devraient agir ensemble ou plutôt entrer en compétition. Nous proposons dans cette section une discrétisation semi-implicite plus spécifique construite dans le but d'obtenir

$$F_0(\mathbf{a}^{n+1}) - F_0(\mathbf{a}^n) - \mathbf{d}^{F_0}(\mathbf{a}^n, \mathbf{a}^{n+1}) \cdot (\mathbf{a}^{n+1} - \mathbf{a}^n) = 0, \quad \forall (\mathbf{a}^n, \mathbf{a}^{n+1}) \in \mathcal{S}^2. \quad (\text{V.41})$$

Pour le modèle de Cahn-Hilliard diphasique, dans [KKL04b] et [KKL04a], les auteurs donnent d'autres discrétisations de ce type obtenues grâce à des développements de Taylor du potentiel de Cahn-Hilliard (cf remarque V.21).

Pour simplifier les notations, nous notons $\mathbf{a} := \mathbf{a}^n$ et $\mathbf{b} := \mathbf{a}^{n+1}$ dans les calculs suivants. Nous écrivons $F_0(\mathbf{b}) - F_0(\mathbf{a})$ comme la somme de termes contenant δ_1 , δ_2 ou δ_3 en facteur, avec $\delta_i = b_i - a_i$ pour $i = 1, 2, 3$. Puisque $F_0(c_1, c_2, c_3) = \sigma_{12}c_1^2c_2^2 + \sigma_{13}c_1^2c_3^2 + \sigma_{23}c_2^2c_3^2 + c_1c_2c_3(\Sigma_1c_1 + \Sigma_2c_2 + \Sigma_3c_3)$, il est suffisant de considérer séparément les termes de la forme $b_i^2b_jb_k - a_i^2a_ja_k$ avec $(i, j, k) \in \{1, 2, 3\}^3$. Nous utilisons les identités $a_i^2 = b_i^2 - (a_i + b_i)\delta_i$ et $a_j = b_j - \delta_j$ pour introduire δ_i , δ_j et δ_k dans la formule :

$$\begin{aligned} b_i^2b_jb_k - a_i^2a_ja_k &= b_i^2(b_jb_k - a_ja_k) + (a_i + b_i)a_ja_k\delta_i \\ &= b_i^2(b_j\delta_k + a_k\delta_j) + (a_i + b_i)a_ja_k\delta_i \\ &= (a_i + b_i)a_ja_k\delta_i + b_i^2a_k\delta_j + b_i^2b_j\delta_k. \end{aligned}$$

Nous utilisons maintenant cette expression pour construire une formule symétrique ayant pour objectif d'obtenir au moins formellement une discrétisation d'ordre 2. En intervertissant les rôles de j et k , nous pouvons obtenir

$$b_i^2b_jb_k - a_i^2a_ja_k = (a_i + b_i)a_ja_k\delta_i + \frac{1}{2}b_i^2(a_k + b_k)\delta_j + \frac{1}{2}b_i^2(a_j + b_j)\delta_k,$$

et finalement, en intervertissant les rôles de $\mathbf{a}$ et $\mathbf{b}$, nous obtenons

$$b_i^2b_jb_k - a_i^2a_ja_k = \frac{1}{2}(a_i + b_i)(a_ja_k + b_jb_k)\delta_i + \frac{1}{4}(a_i^2 + b_i^2)(a_k + b_k)\delta_j + \frac{1}{4}(a_i^2 + b_i^2)(a_j + b_j)\delta_k. \quad (\text{V.42})$$

Nous obtenons une formule pour les termes de la forme $b_i^2b_j^2 - a_i^2a_j^2$ en prenant $k = j$ dans (V.42) :

$$b_i^2b_j^2 - a_i^2a_j^2 = \frac{1}{2}(a_i + b_i)(a_j^2 + b_j^2)\delta_i + \frac{1}{2}(a_i^2 + b_i^2)(a_j + b_j)\delta_j. \quad (\text{V.43})$$

Ainsi, nous proposons de définir, pour tout $i \in \{1, 2, 3\}$, l'approximation consistante suivante des termes non-linéaires :

$$\begin{aligned} d_i^{F_0}(\mathbf{a}^n, \mathbf{a}^{n+1}) &= \frac{\Sigma_i}{4} [a_i^{n+1} + a_i^n] [(a_j^{n+1} + a_k^{n+1})^2 + (a_j^n + a_k^n)^2] \\ &\quad + \frac{\Sigma_j}{4} [(a_j^{n+1})^2 + (a_j^n)^2] [a_i^{n+1} + a_k^{n+1} + a_i^n + a_k^n] \\ &\quad + \frac{\Sigma_k}{4} [(a_k^{n+1})^2 + (a_k^n)^2] [a_i^{n+1} + a_j^{n+1} + a_i^n + a_j^n]. \end{aligned} \quad (\text{V.44})$$

Nous pouvons déduire de la définition de F_0 et des formules (V.42), (V.43) et (V.44) que, pour tout $\mathbf{a}^n \in \mathcal{S}$ et $\mathbf{a}^{n+1} \in \mathcal{S}$

$$F_0(\mathbf{a}^{n+1}) - F_0(\mathbf{a}^n) = \sum_{i=1}^3 d_i^{F_0}(\mathbf{a}^n, \mathbf{a}^{n+1})(a_i^{n+1} - a_i^n), \quad \forall (\mathbf{a}^n, \mathbf{a}^{n+1}) \in \mathcal{S}^2,$$

et que pour tout $\mathbf{c} \in \mathcal{S}$,

$$d_i^{F_0}(\mathbf{c}, \mathbf{c}) = \frac{\partial F_0}{\partial c_i}(\mathbf{c}).$$

Ainsi, à partir de l'égalité (V.41), nous pouvons déduire que les hypothèses (V.21) du théorème V.9 et (V.26) du théorème V.10 sont satisfaites pour la contribution de F_0 lorsque nous utilisons la discrétisation semi-implicite (V.44) (cf section V.2.6 pour plus de détails). Ces hypothèses sont satisfaites pour tout pas de temps Δt et même dans les situations d'étalement total.

V.2.5 Discrétisation semi-implicite de la contribution de P

Rappelons la définition de P :

$$P(\mathbf{c}) = \Lambda c_1^2 c_2^2 c_3^2.$$

Nous considérons seulement la discrétisation semi-implicite de la contribution de P proposée dans [Lap06]. Les simulations numériques réalisées dans [Lap06] montrent les difficultés à utiliser une discrétisation implicite pour ce terme (non convergence de la méthode de linéarisation de Newton pour la résolution du problème V.4). De plus, nous n'avons pas de décomposition convexe-concave naturelle de P .

Pour obtenir une estimation d'énergie discrète, nous cherchons des fonctions d_1^P , d_2^P et d_3^P telles que $d_i^P(\mathbf{c}, \mathbf{c}) = \frac{\partial P}{\partial c_i}(\mathbf{c})$, $\forall \mathbf{c} \in \mathcal{S}$ et

$$P(\mathbf{a}^{n+1}) - P(\mathbf{a}^n) - \mathbf{d}^P(\mathbf{a}^n, \mathbf{a}^{n+1}) \cdot (\mathbf{a}^{n+1} - \mathbf{a}^n) = 0, \quad \forall (\mathbf{a}^n, \mathbf{a}^{n+1}) \in \mathcal{S}^2. \quad (\text{V.45})$$

Nous définissons pour $i \in \{1, 2, 3\}$, $\delta_i = b_i - a_i$ et nous utilisons l'identité $a_i^2 = b_i^2 - (a_i + b_i)\delta_i$ et ainsi l'égalité (V.43) qui nous permet d'introduire δ_i , δ_j et δ_k dans les termes $b_i^2 b_j^2 b_k^2$, $(i, j, k) \in \{1, 2, 3\}^3$:

$$\begin{aligned} b_i^2 b_j^2 b_k^2 - a_i^2 a_j^2 a_k^2 &= b_i^2 (b_j^2 b_k^2 - a_j^2 a_k^2) + (a_i + b_i) a_j^2 a_k^2 \delta_i \\ &= (a_i + b_i) a_j^2 a_k^2 \delta_i + \frac{1}{2} b_i^2 (a_j + b_j) (a_k^2 + b_k^2) \delta_j + \frac{1}{2} b_i^2 (a_j^2 + b_j^2) (a_k + b_k) \delta_k. \end{aligned} \quad (\text{V.46})$$

En ajoutant les trois formules données par (V.46) avec $(i, j, k) = (1, 2, 3)$, $(2, 1, 3)$ et $(3, 1, 2)$, nous obtenons

$$\begin{aligned} b_1^2 b_2^2 b_3^2 - a_1^2 a_2^2 a_3^2 &= \frac{1}{3} \left[a_2^2 a_3^2 + \frac{1}{2} b_2^2 a_3^2 + \frac{1}{2} a_2^2 b_3^2 + b_2^2 b_3^2 \right] (a_1 + b_1) \delta_1 \\ &\quad + \frac{1}{3} \left[a_1^2 a_3^2 + \frac{1}{2} b_1^2 a_3^2 + \frac{1}{2} a_1^2 b_3^2 + b_1^2 b_3^2 \right] (a_2 + b_2) \delta_2 \\ &\quad + \frac{1}{3} \left[a_1^2 a_2^2 + \frac{1}{2} b_1^2 a_2^2 + \frac{1}{2} a_1^2 b_2^2 + b_1^2 b_2^2 \right] (a_3 + b_3) \delta_3. \end{aligned}$$

Ainsi, en définissant

$$d_i^P(\mathbf{a}^n, \mathbf{a}^{n+1}) = \frac{\Lambda}{3} (a_i^n + a_i^{n+1}) \left[(a_j^n)^2 (a_k^n)^2 + \frac{1}{2} (a_j^{n+1})^2 (a_k^n)^2 + \frac{1}{2} (a_j^n)^2 (a_k^{n+1})^2 + (a_j^{n+1})^2 (a_k^{n+1})^2 \right], \quad (\text{V.47})$$

nous obtenons la propriété (V.45) et pour tout $\mathbf{c} \in \mathcal{S}$,

$$d_i^P(\mathbf{c}, \mathbf{c}) = \frac{\partial P}{\partial c_i}(\mathbf{c}).$$

Ainsi, comme dans la section précédente, à partir de l'inégalité (V.45), nous pouvons déduire que les hypothèses (V.21) du théorème V.9 et (V.26) du théorème V.10 sont satisfaites pour la contribution de P lorsque nous utilisons la discrétisation semi-implicite (V.47) (cf section V.2.6 pour plus de détails).

V.2.6 Résumé des résultats

Dans les sections V.2.2, V.2.3, V.2.4 et V.2.5, nous avons présenté séparément plusieurs discrétisations $\mathbf{d}^{F_0}$ pour la contribution de F_0 et une discrétisation $\mathbf{d}^P$ pour la contribution de P (rappelons que le potentiel de Cahn-Hilliard F est défini par $F = F_0 + P$). Pour la contribution de P , nous n'avons considéré qu'un seul schéma semi-implicite défini par (V.47). Ce schéma peut ensuite être combiné avec trois discrétisations possibles pour

la contribution de F_0 : lorsque nous choisissons la discrétisation (V.32), resp. (V.39), resp. (V.44), nous faisons référence au schéma obtenu par implicite, resp. convexe-concave, resp. semi-implicite.

Nous pouvons maintenant énoncer les théorèmes d'existence et de convergence grâce aux théorèmes généraux V.9 et V.10 et aux estimations (V.33), (V.35), (V.38) et (V.41) valides pour chaque discrétisation particulière. Avant tout, rappelons que le potentiel de Cahn-Hilliard $F = F_0 + P$ satisfait l'hypothèse (IV.20) (en fait seul l'hypothèse de positivité est non triviale, cf Proposition IV.11) à condition que :

- $\Lambda \geq 0$ lorsque $\Sigma_i > 0$ pour tout $i \in \{1, 2, 3\}$,
- $\Lambda \geq \Lambda_0$ lorsque la condition (IV.14) sur les coefficients $(\Sigma_1, \Sigma_2, \Sigma_3)$ est satisfaite (cela autorise l'existence d'au plus un coefficient Σ_i négatif).

Dans le premier cas, les théorèmes d'existence et de convergence sont prouvés pour les trois schémas (implicite, convexe-concave et semi-implicite) alors que lorsqu'un coefficient Σ_i est négatif, les théorèmes d'existence et de convergence ne sont prouvés que pour les discrétisations convexe-concave et semi-implicite, l'existence de solution pour le schéma implicite étant encore un problème ouvert. Notons que dans ce dernier cas, nous observons dans plusieurs expérimentations numériques (cf section V.3), une non convergence de la méthode de linéarisation de Newton dans la résolution du problème V.4. Dans le cas où tous les coefficients Σ_i sont positifs, nous pouvons aussi remarquer que le schéma implicite garantit la décroissance de l'énergie seulement pour de petits pas de temps (cf Proposition V.15) alors que les schémas convexe-concave et semi-implicite la garantissent pour tout pas de temps. Tous ces résultats sont énoncés dans les propositions V.17 et V.18 et résumés dans la table V.1.

Proposition V.17 (Etalement partiel)

Nous supposons que $\forall i \in \{1, 2, 3\}$, $\Sigma_i > 0$, que $F = F_0 + P$ avec $\Lambda \geq 0$ et que la mobilité satisfait (IV.19). Alors, il existe au moins une solution au problème V.2 où $\mathbf{d}^F$ correspond au schéma implicite, convexe-concave ou semi-implicite. De plus, si $\frac{1}{2} < \beta \leq 1$ alors les conclusions du théorème V.10 sont satisfaites.

Proposition V.18 (Etalement total)

Nous supposons que le triplet de coefficients Σ satisfait (IV.14), que $F = F_0 + P$ avec $\Lambda \geq \Lambda_0$ (cf proposition IV.11) et que la mobilité satisfait (IV.19). Alors, il existe au moins une solution au problème V.2 où $\mathbf{d}^F$ correspond au schéma convexe-concave ou semi-implicite. De plus, si $\frac{1}{2} < \beta \leq 1$ alors les conclusions du théorème V.10 sont satisfaites.

Schémas	Déf.	Implicite	Convexe-concave	Semi-implicite
		Semi-impl. (V.47)		
$\mathbf{d}^P$		Impl. (V.32)	Conv.-conc. (V.39)	Semi-impl. (V.44)
$\mathbf{d}^{F_0}$				
$\forall i, \Sigma_i > 0$	$\mathbf{d}^F = \mathbf{d}^{F_0} + \mathbf{d}^P$ $\Lambda \geq 0$	Décroiss. de l'énergie $\Delta t \leq \Delta t_0$ Existence $\forall \Delta t$ Convergence ($\beta > 1/2$)	Décroiss. de l'énergie $\forall \Delta t$ Existence $\forall \Delta t$ Convergence ($\beta > 1/2$)	
$\exists i, \Sigma_i < 0$ t.q. (IV.14)	$\mathbf{d}^F = \mathbf{d}^{F_0} + \mathbf{d}^P$ $\Lambda \geq \Lambda_0$	Problème ouvert	Décroiss. de l'énergie $\forall \Delta t$ Existence $\forall \Delta t$ Convergence ($\beta > 1/2$)	

TAB. V.1 – Résumé des résultats théoriques

V.2.7 Schémas correspondants dans le cas diphasique.

Considérons un système avec deux composants (noté ci-dessous avec les indices 1 et 2 respectivement) et supposons que l'évolution des paramètres d'ordre c_i , ($i = 1, 2$) et des potentiels chimiques $\tilde{\mu}_i$, ($i = 1, 2$) associés à ces deux phases soit gouvernée par le modèle de Cahn-Hilliard diphasique :

$$\begin{cases} \frac{\partial c_i}{\partial t} = \operatorname{div} (M(c_1, c_2) \nabla \tilde{\mu}_i), & \text{pour } i = 1, 2, \\ \tilde{\mu}_i = \frac{12}{\varepsilon} \sigma_{12} f'(c_i) - \frac{3}{2} \varepsilon \sigma_{12} \Delta c_i & \text{pour } i = 1, 2, \end{cases} \quad (\text{V.48})$$

où ε est l'épaisseur d'interface, $M(c_1, c_2)$ la mobilité et σ_{12} la tension de surface entre les deux composants. Les inconnues sont liées par les relations suivantes : $c_1 + c_2 = 1$ et $\tilde{\mu}_1 + \tilde{\mu}_2 = 0$.

La consistance algébrique (cf section IV.2.1) garantit que le triplet $(c_1, c_2 = 1 - c_1, c_3 = 0)$ est une solution particulière du modèle de Cahn-Hilliard triphasique (IV.9) (avec $M_0(\mathbf{c}) = 2\sigma_{12}M(c_1, c_2)$) et pour tout choix des valeurs des tensions de surface σ_{13} et σ_{23} impliquant le troisième composant. Dans ce cas, les potentiels chimiques triphasiques sont donnés par $\mu_i = \frac{\Sigma_i}{2\sigma_{12}}\tilde{\mu}_i$ pour $i = 1, 2$ et $\mu_3 = 0$.

Des résultats équivalents peuvent être obtenus pour le système complètement discret et nous pouvons identifier les schémas correspondants au modèle diphasique (V.48) : Etant donné $(c_{ih}^n, \mu_{ih}^n) \in \mathcal{V}_{Dh}^{c_i} \times \mathcal{V}_h^\mu$,

- Schéma implicite dans le cas diphasique : pour $i = 1, 2$,

Trouver $(c_{ih}^{n+1}, \mu_{ih}^{n+1}) \in \mathcal{V}_{Dh}^{c_i} \times \mathcal{V}_h^\mu$ tel que

$$\begin{cases} \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx = - \int_{\Omega} M(c_{1h}^{n+\alpha}, c_{2h}^{n+\alpha}) \nabla \tilde{\mu}_{ih}^{n+1} \nabla \nu_h^\mu dx, & \forall \nu_h^\mu \in \mathcal{V}_h^\mu, \\ \int_{\Omega} \tilde{\mu}_{ih}^{n+1} \nu_h^c dx = \frac{12}{\varepsilon} \sigma_{12} \int_{\Omega} f'(c_{ih}^{n+1}) \nu_h^c dx + \frac{3}{2} \varepsilon \sigma_{12} \int_{\Omega} \nabla c_{ih}^{n+\beta} \nabla \nu_h^c dx, & \forall \nu_h^c \in \mathcal{V}_{Dh,0}^c. \end{cases} \quad (\text{V.49})$$

- Schéma convexe-concave dans le cas diphasique : pour $i = 1, 2$,

Trouver $(c_{ih}^{n+1}, \mu_{ih}^{n+1}) \in \mathcal{V}_{Dh}^{c_i} \times \mathcal{V}_h^\mu$ tel que

$$\begin{cases} \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx = - \int_{\Omega} M(c_{1h}^{n+\alpha}, c_{2h}^{n+\alpha}) \nabla \tilde{\mu}_{ih}^{n+1} \nabla \nu_h^\mu dx, & \forall \nu_h^\mu \in \mathcal{V}_h^\mu, \\ \int_{\Omega} \tilde{\mu}_{ih}^{n+1} \nu_h^c dx = \frac{12}{\varepsilon} \sigma_{12} \int_{\Omega} [(f^+)'(c_{ih}^{n+1}) + (f^-)'(c_{ih}^n)] \nu_h^c dx \\ \quad + \frac{3}{2} \varepsilon \sigma_{12} \int_{\Omega} \nabla c_{ih}^{n+\beta} \nabla \nu_h^c dx, & \forall \nu_h^c \in \mathcal{V}_{Dh,0}^c, \end{cases} \quad (\text{V.50})$$

où $f = f^+ + f^-$ est la décomposition convexe-concave de f donnée dans (V.40).

- Schéma semi-implicite dans le cas diphasique : pour $i = 1, 2$,

Trouver $(c_{ih}^{n+1}, \mu_{ih}^{n+1}) \in \mathcal{V}_{Dh}^{c_i} \times \mathcal{V}_h^\mu$ tel que

$$\begin{cases} \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx = - \int_{\Omega} M(c_{1h}^{n+\alpha}, c_{2h}^{n+\alpha}) \nabla \tilde{\mu}_{ih}^{n+1} \nabla \nu_h^\mu dx, & \forall \nu_h^\mu \in \mathcal{V}_h^\mu, \\ \int_{\Omega} \tilde{\mu}_{ih}^{n+1} \nu_h^c dx = \frac{12}{\varepsilon} \sigma_{12} \int_{\Omega} \left[f' \left(\frac{c_{ih}^n + c_{ih}^{n+1}}{2} \right) - \frac{1}{2} (1 - c_{ih}^n - c_{ih}^{n+1}) (c_{ih}^{n+1} - c_{ih}^n)^2 \right] \nu_h^c dx \\ \quad + \frac{3}{2} \varepsilon \sigma_{12} \int_{\Omega} \nabla c_{ih}^{n+\beta} \nabla \nu_h^c dx, & \forall \nu_h^c \in \mathcal{V}_{Dh,0}^c. \end{cases} \quad (\text{V.51})$$

Proposition V.19

Les schémas diphasiques ci-dessus (V.49), (V.50) et (V.51) ont au moins une solution. De plus, en définissant $M_0 = 2\sigma_{12}M$, $\mu_{ih}^{n+1} = \frac{\Sigma_i}{2\sigma_{12}}\tilde{\mu}_{ih}^{n+1}$ pour $i = 1, 2$ et $\mu_{3h}^{n+1} = 0$, nous avons le résultat suivant : si $((c_{1h}^{n+1}, \tilde{\mu}_{1h}^{n+1}), (c_{2h}^{n+1}, \tilde{\mu}_{2h}^{n+1}))$ est une solution de (V.49) (resp. (V.50), resp. (V.51)) alors $((c_{1h}^{n+1}, \mu_{1h}^{n+1}), (c_{2h}^{n+1}, \mu_{2h}^{n+1}), (0, 0))$ est une solution du problème triphasique (V.2) où $\mathbf{d}^F$ est donné par (V.32) (resp. (V.39), resp. (V.44)).

Remarque V.20

L'expression des potentiels chimiques triphasiques diffère de celle des potentiels chimiques diphasiques $\tilde{\mu}_i$ mais les quantités d'intérêts dans nos applications sont les paramètres d'ordre qui indiquent la position des phases et les forces capillaires f_{ca} qui sont utilisées pour le couplage avec les équations de Navier-Stokes. Dans le modèle triphasique, nous utilisons l'expression $f_{ca} = \sum_{i=1}^3 \mu_i \nabla c_i$. Le point clé est que dans le cas où $c_3 = 0$, nous avons

$$\begin{aligned} f_{ca} &= \mu_1 \nabla c_1 + \mu_2 \nabla c_2 \\ &= \frac{\Sigma_1}{2\sigma_{12}} \tilde{\mu}_1 \nabla c_1 + \frac{\Sigma_2}{2\sigma_{12}} (-\tilde{\mu}_1) \nabla (1 - c_1) \\ &= \tilde{\mu}_1 \nabla c_1, \end{aligned}$$

qui est l'expression des forces capillaires dans le cas diphasique.

Remarque V.21

Dans [KKL04a], les auteurs proposent un schéma semi-implicite pour la discrétisation du modèle de Cahn-Hilliard deux-phases. Ce schéma est obtenu à partir de développements de Taylor. Il est intéressant de remarquer que, lorsque seulement deux phases sont présentes, le schéma semi-implicite que nous présentons dans ce manuscrit est très proche (mais différent) du schéma donné dans [KKL04a]. En effet, nous avons la relation suivante :

$$\left[f' \left(\frac{c_h^n + c_h^{n+1}}{2} \right) - \frac{1}{2} (1 - c_h^n - c_h^{n+1}) (c_h^{n+1} - c_h^n)^2 \right] - \left[f'(c_h^{n+1}) - \frac{f''(c_h^{n+1})}{2} (c_h^{n+1} - c_h^n) + \frac{f'''(c_h^{n+1})}{6} (c_h^{n+1} - c_h^n)^2 \right] = (c_h^{n+1} - c_h^n)^3,$$

où le premier terme du membre de gauche est la discrétisation de $f'(c)$ proposée dans ce manuscrit et le second est la discrétisation proposée dans [KKL04a].

V.3 Illustrations numériques

Dans cette section, nous présentons quelques illustrations numériques en une dimension et deux dimensions permettant de comparer les différentes discrétisations en temps des termes non linéaires, présentées dans la section V.2.

Nous utilisons la notation suivante pour désigner les différents schémas :

- Impl. désigne la discrétisation implicite (V.32) pour la contribution de F_0 combinée à la discrétisation semi-implicite (V.47) pour la contribution de P et $\beta = 1$,
- CC. désigne la discrétisation convexe-concave (V.39) pour la contribution de F_0 combinée à la discrétisation semi-implicite (V.47) pour la contribution de P et $\beta = 1$,
- SImpl(β). désigne la discrétisation semi-implicite (V.44) pour la contribution de F_0 combinée à la discrétisation semi-implicite (V.47) pour la contribution de P et la valeur donnée de β ,
- SImpl. désigne le schéma SImpl(1).

En 1D, la discrétisation spatiale est réalisée sur une grille uniforme par éléments finis linéaires. En 2D, pour limiter les temps de calcul, nous utilisons des éléments finis de Lagrange $\mathbb{Q}_1$ sur des maillages adaptatifs localement raffinés (cf partie 1). Le critère de raffinement impose la valeur du plus petit diamètre $h_{\text{interface}}$ d'une cellule et assure que les aires raffinées sont aux voisinages des interfaces (*i.e.* où aucun paramètre d'ordre n'est égal à 1). En pratique, nous prenons $h_{\text{interface}} = \frac{\varepsilon}{2}$ pour garantir la présence d'au moins deux cellules dans l'interface. Nous renvoyons à l'introduction de la partie 3 pour une définition plus précise du critère de raffinement. Dans tous les tests bidimensionnels, les solutions approchées sont visualisées grâce aux lignes de niveau de la fonction :

$$(c_1, c_2, c_3) \mapsto (1 - c_1)(1 - c_2)(1 - c_3), \quad (\text{V.52})$$

qui n'est pas nulle seulement dans l'interface. Les figures présentant les solutions approchées montrent également les maillages localement raffinés utilisés pour les calculs correspondants.

Pour les études de convergence, pour chaque schéma, les différentes solutions approchées $\mathbf{c}^{\Delta t_j}$ sont calculées en utilisant les pas de temps Δt_j . Puisque, aucune solution analytique du système de Cahn-Hilliard (IV.9) n'est connue, nous utilisons une solution approchée $\mathbf{c}^{\Delta t_{\text{ref}}}$ obtenue avec un pas de temps de référence Δt_{ref} comme solution de référence. Evidemment, le pas de temps de référence Δt_{ref} est supposé assez petit devant Δt_j . Bien que le critère de raffinement soit le même pour tous les calculs, les grilles raffinées peuvent différer légèrement d'un calcul à l'autre puisque les pas de temps sont différents. Cependant, la norme L^2 de l'erreur

$$e_j(t) = \left| \mathbf{c}^{\Delta t_j}(t, \cdot) - \mathbf{c}^{\Delta t_{\text{ref}}}(t, \cdot) \right|_{(L^2(\Omega))^3},$$

à un instant fixé t , est exactement calculée sur la grille uniforme de taille $h_{\min}$ pendant une étape de post-traitement.

V.3.1 Cas tests deux phases

Dans cette section, les schémas sont comparés sur des cas tests ne faisant intervenir que deux phases. Nous résolvons numériquement le problème trois phases V.4 mais le troisième paramètre d'ordre c_3 est initialisé à

zéro sur tout le domaine de manière à ce que les deux phases en présence soient décrites par les paramètres d'ordre c_1 et $c_2 = 1 - c_1$. La propriété de consistance (*cf* sections IV.2.1 et V.2.7) garantit que le paramètre d'ordre c_3 restera nul durant toute la simulation et en conséquence, les schémas que nous comparons sont en fait ceux présentés dans la section V.2.7.

Les cas tests donnés illustrent les deux comportements différents du système de Cahn-Hilliard : le premier est la stabilité de l'épaisseur d'interface observée proche de ε et le second est le déplacement de l'interface sous l'influence des tensions de surface.

Dynamique de l'interface

Le premier cas test est réalisé sur le domaine $] -1, 1[$ avec les paramètres suivants : l'épaisseur d'interface $\varepsilon = 0.5$, une mobilité constante $M_0 = 8$ et une tension de surface entre les deux phases présentes $\sigma = 1$. Nous imposons des conditions aux bords de type Neumann pour les paramètres d'ordre et les potentiels chimiques. La donnée initiale est définie par :

$$c_1^0(x) = \frac{1}{2} \left(1 + \tanh \left(\frac{2x}{10\varepsilon} \right) \right), \quad \text{et} \quad c_2^0(x) = 0, \quad \forall x \in [-1, 1].$$

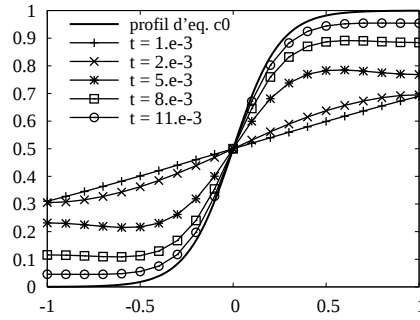


FIG. V.1 – Evolution de paramètre d'ordre c_1 en utilisant le schéma implicite avec $\Delta t = 10^{-3}$

La figure V.1 montre l'évolution du paramètre d'ordre c_1 vers la forme d'équilibre. Nous avons également représenté sur cette figure l'approximation de la solution stationnaire :

$$c_0(x) = \frac{1}{2} \left(1 + \tanh \left(\frac{2x}{\varepsilon} \right) \right), \quad \forall x \in \mathbb{R},$$

obtenue en résolvant exactement le problème suivant, posé sur un domaine infini :

$$\begin{cases} -\frac{3}{2}\sigma\varepsilon c_0''(x) + 12\frac{\sigma}{\varepsilon}f'(c_0(x)) = 0, & \forall x \in \mathbb{R}, \\ \lim_{+\infty} c_0 = 1, \quad \lim_{-\infty} c_0 = 0, \quad c_0(0) = \frac{1}{2}. \end{cases}$$

La figure V.2 présente l'étude de convergence. La solution de référence est calculée avec le schéma SImpl.(0.5) et $\Delta t_{\text{ref}} = 10^{-8}$. Nous avons effectué plusieurs calculs en utilisant les différents schémas et pour chacun des pas de temps suivants : $\Delta t_j : 2.10^{-4}, 5.10^{-4}, 10^{-4}, 2.10^{-5}, 5.10^{-5}, 10^{-5}, 10^{-6}$. Les normes L^2 des erreurs correspondantes $e_j(t)$ au temps $t = 0.01$ sont représentées sur la figure de gauche et les ordres de convergence de chacun des schémas sont donnés dans le tableau de droite.

Nous observons une convergence d'ordre 1 pour les schémas Impl., CC. et SImpl. et une convergence d'ordre (environ) 2 pour le schéma SImpl.(0.5). Notons également que le schéma CC. est moins précis que les trois autres, alors que le schéma SImpl. permet d'atteindre la même précision que le schéma Impl.

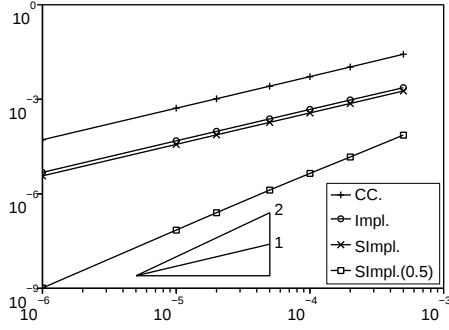


Schéma	Ordre de convergence
CC.	1.0
SImpl.	1.0
Impl.	1.0
SImpl.(0.5)	1.8

FIG. V.2 – Erreurs $e_j(t) = |\mathbf{c}^{\Delta t_j}(t, \cdot) - \mathbf{c}^{\Delta t_{\text{ref}}}(t, \cdot)|_{(L^2(\Omega))^3}$ au temps $t = 0.01$ en fonction du pas de temps Δt_j (à gauche) et ordres de convergence (à droite) obtenus avec les différents schémas.

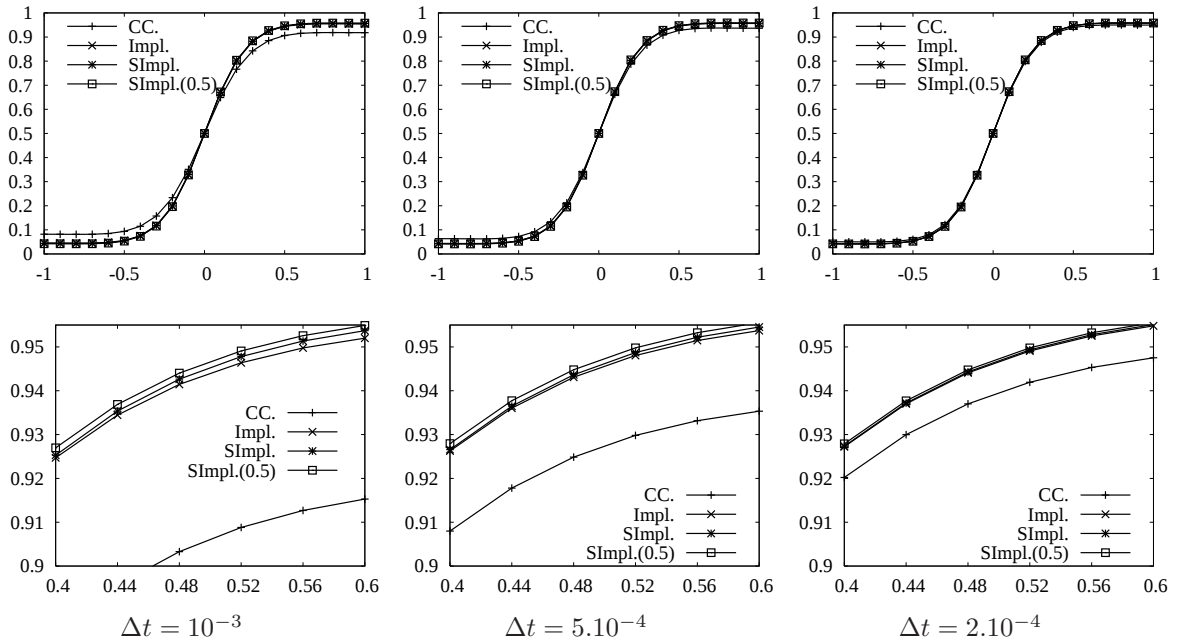


FIG. V.3 – Paramètre d'ordre c_1 en fonction de la variable d'espace x au temps $t = 0.01$. En haut : $x \in [-1, 1]$, En bas : $x \in [0.4, 0.6]$ (zoom)

L'influence des différents schémas sur l'allure de la solution est illustrée par la figure V.3. Nous représentons, pour différents pas de temps, le paramètre d'ordre c_1 en fonction de la variable de l'espace sur le domaine entier (en haut) et sur une partie du domaine en zoomant (en bas). Les schémas Impl., SImpl, SImpl.(0.5) donnent des résultats très proches alors que le schéma CC. donne un profil différent.

Bulle éllipsoïdale - Conditions aux bords de type Neumann

Ce test est réalisé sur le domaine $] -0.2, 0.2[^2$ avec les paramètres suivants : l'épaisseur d'interface $\varepsilon = 0.01$, une mobilité constante $M_0 = 10^{-4}$ et une tension de surface entre les deux phases $\sigma = 1$. Nous imposons une condition aux bords de type Neumann pour les paramètres d'ordre ainsi que pour les potentiels chimiques. La donnée initiale est définie par :

$$c_1^0(x, y) = \frac{1}{2} \left[1 + \tanh \left(\frac{2}{\varepsilon} \left[\left(\frac{x^2}{a^2} + a^2 y^2 \right)^{\frac{1}{2}} - 0.1 \right] \right) \right], \quad c_2^0(x, y) = 0, \quad \forall (x, y) \in [-0.2, 0.2]^2,$$

où $a = 1.5$.

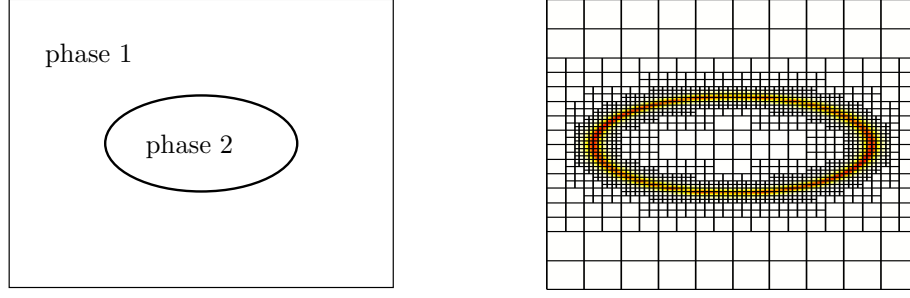
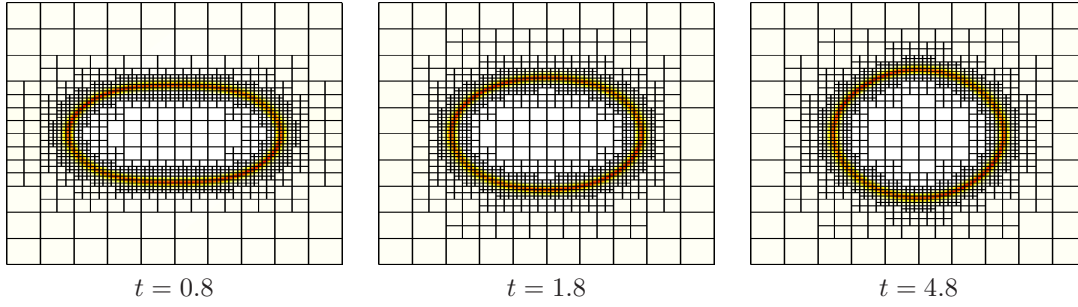
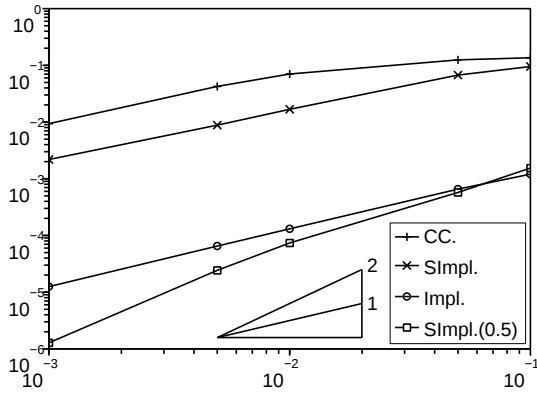


FIG. V.4 – Configuration du cas test (à gauche) et position initiale (à droite)

La figure V.4 montre la configuration du cas test (à gauche) et la position des interfaces ainsi que le maillage à l'instant initial (à droite). Rappelons que la représentation des interfaces est réalisée grâce aux lignes de niveau de la fonction définie par (V.52).

FIG. V.5 – Evolution de la position de l'interface en utilisant le schéma Impl. avec $\Delta t = 5.10^{-4}$

La figure V.5 montre l'évolution de la position des interfaces. Le système tend vers la position qui minimise la longueur des interfaces tout en conservant le volume de chaque phase, *i.e.* une interface circulaire. Notons que l'état stationnaire n'est pas encore atteint à la fin de notre calcul ($t = 4.8$).



j	Δt_j	$\frac{\ln(e_{j+1}/e_j)}{\ln(\Delta t_{j+1}/\Delta t_j)}$			
		CC.	SImpl.	SImpl.(0.5)	Impl.
1	10^{-1}	0.1	0.5	1.4	0.9
2	5.10^{-2}	0.4	0.9	1.3	1.0
3	10^{-2}	0.7	0.9	1.6	1.0
4	5.10^{-3}	0.9	0.9	1.8	1.0
5	10^{-3}	-	-	-	-

FIG. V.6 – Erreurs $e_j(t) = |\mathbf{c}^{\Delta t_j}(t, \cdot) - \mathbf{c}^{\Delta t_{\text{ref}}}(t, \cdot)|_{(L^2(\Omega))^3}$ au temps $t = 3.8$ en fonction du pas de temps Δt_j (à gauche) et ordres de convergence (à droite) obtenus en utilisant les différents schémas

La figure V.6 présente une étude de convergence. La solution de référence est calculée en utilisant le schéma SImpl.(0.5) avec $\Delta t_{\text{ref}} = 5.10^{-4}$. Nous avons effectué plusieurs calculs avec les différents schémas et pour chacun des pas de temps suivants Δt_j : 10^{-1} , 5.10^{-2} , 10^{-2} , 5.10^{-3} , 10^{-3} . Les normes L^2 des erreurs correspondantes $e_j(t)$ au temps $t = 3.8$ sont représentées sur la figure de gauche et les ordres de convergence obtenus sont donnés dans le tableau de droite. Nous obtenons essentiellement les mêmes résultats qu'en dimension un, *i.e.*

une convergence d'ordre 1 pour les schémas Impl., CC. et SImpl. et une convergence d'ordre 2 pour le schéma SImpl.(0.5). On note encore que le schéma CC. est moins précis que les trois autres. L'influence des différents schémas sur l'allure de la solution est illustrée par la figure V.7. Nous montrons la position de l'interface obtenue avec les différents schémas et différents pas de temps. Les schémas Impl., SImpl.(0.5) donnent des résultats très similaires pour tous les pas de temps testés. Nous observons des différences sur les formes de bulle lorsque le schéma produit une erreur plus grande que 10^{-2} , par exemple avec le schéma CC.

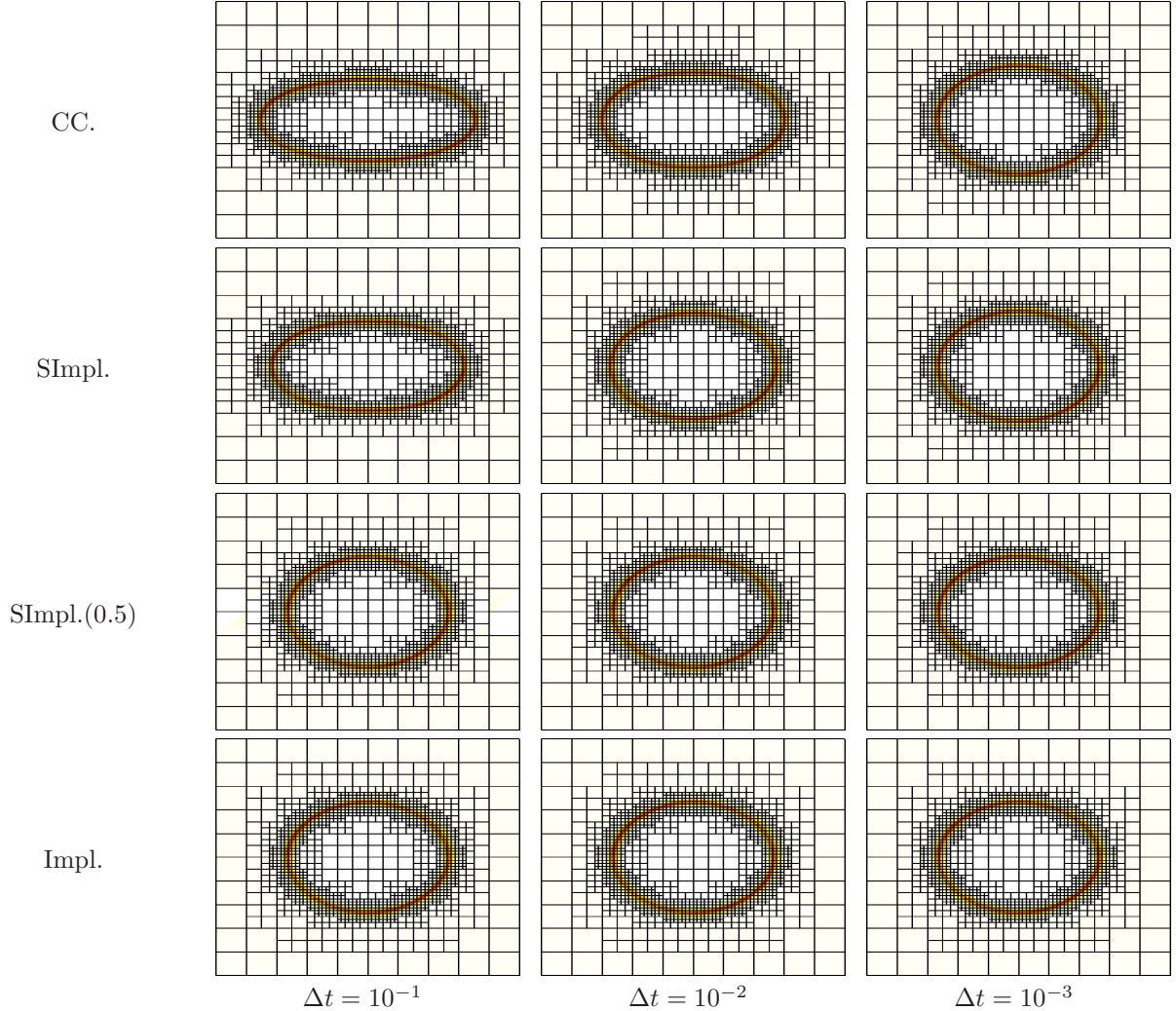


FIG. V.7 – Influence des schémas sur la forme de la bulle à $t = 3$.

Bulle ellipsoïdale - Conditions aux bords de type Dirichlet

Ce cas test est réalisé sur le domaine $] -0.1, 0.1[\times] 0, 0.2[$ avec les paramètres suivants : une épaisseur d'interface $\varepsilon = 6.10^{-3}$, une mobilité constante $M_0 = 10^{-4}$ et une tension de surface entre les deux phases présentes $\sigma = 1$. La donnée initiale est définie par :

$$c_1^0(x, y) = \frac{1}{2} \left[1 + \tanh \left(\frac{2}{\varepsilon} \left(4x^2 + \frac{y^2}{12.25} \right)^{\frac{1}{2}} - 0.05 \right) \right], \quad c_2^0(x, y) = 0, \quad \text{pour tout } (x, y) \in \Omega.$$

La figure V.8 montre la configuration initiale (à gauche) ainsi que la position des interfaces et le maillage à l'instant initial (à droite). Nous imposons des conditions de Neumann pour les paramètres d'ordre et les potentiels chimiques, à l'exception du bord situé au bas du domaine, *i.e.* $[-0.1, 0.1] \times \{0\}$, où des conditions aux bords de type Dirichlet sont imposées aux paramètres d'ordre. Rappelons que la représentation des interfaces est obtenue par le tracé des lignes de niveau de la fonction définie par (V.52).

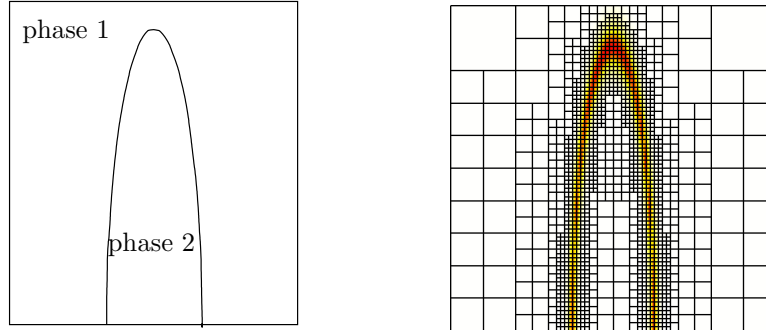
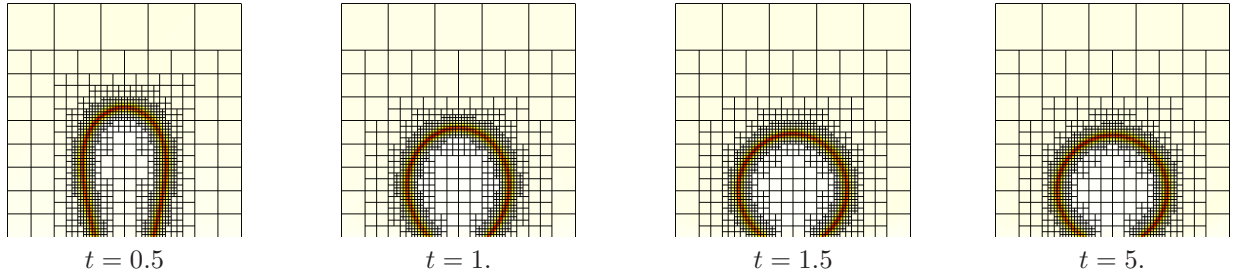
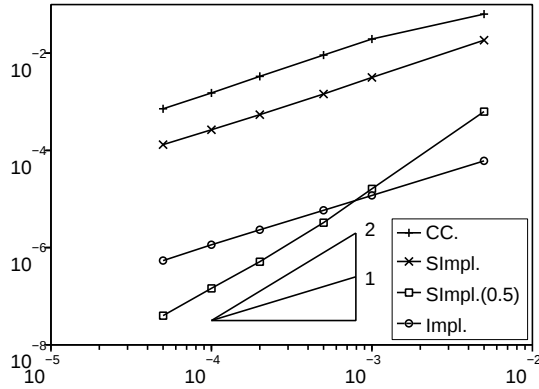


FIG. V.8 – Configuration du cas test (à gauche) et position initiale de l'interface (à droite)

FIG. V.9 – Evolution de la position des interfaces en utilisant le schéma Impl. avec $\Delta t = 5.10^{-5}$

La figure V.9 montre l'évolution de la position des interfaces. Le système tend vers une position qui minimise la longueur des interfaces tout en conservant le volume de chacune des phases, les interfaces décrivent un arc de cercle puisque la valeur des paramètres d'ordre est imposée sur la partie basse du domaine.



j	Δt_j	$\frac{\ln(e_{j+1}/e_j)}{\ln(\Delta t_{j+1}/\Delta t_j)}$			
		CC.	SImpl.	SImpl.(0.5)	Impl.
1	5.10^{-3}	0.7	1.1	2.3	1.0
2	10^{-3}	1.1	1.1	2.3	1.0
3	5.10^{-4}	1.1	1.1	2.0	1.0
4	2.10^{-4}	1.1	1.0	1.8	1.0
5	10^{-4}	1.1	1.0	1.9	1.1
6	5.10^{-5}	-	-	-	-

FIG. V.10 – Erreurs $e_j(t) = |\mathbf{c}^{\Delta t_j}(t, \cdot) - \mathbf{c}^{\Delta t_{\text{ref}}}(t, \cdot)|_{(L^2(\Omega))_3}$ au temps $t = 1.5$ en fonction du pas de temps Δt_j (à gauche) et ordres de convergence (à droite) obtenus avec les différents schémas.

La figure V.10 présente l'étude de convergence. La solution de référence est calculée avec le schéma SImpl.(0.5) et $\Delta t_{\text{ref}} = 10^{-5}$. Plusieurs calculs ont été réalisés en utilisant les différents schémas et pour chacun des pas de temps Δt_j suivants : 5.10^{-3} , 10^{-3} , 5.10^{-4} , 2.10^{-4} , 10^{-4} , 5.10^{-5} . Les normes L^2 des erreurs correspondantes $e_j(t)$ au temps $t = 1.5$ sont représentées sur la figure de gauche et les ordres de convergence obtenus sont donnés dans le tableau de droite. Nous obtenons une convergence d'ordre 1 pour les schémas CC., Impl. et SImpl. et une convergence d'ordre 2 pour le schéma SImpl.(0.5). Remarquons que les schémas SImpl.(0.5) et Impl. donne des résultats significativement plus précis que ceux obtenus avec le schéma CC.

V.3.2 Cas tests trois phases

Dans cette section, nous illustrons les propriétés des différents schémas en dimension 2 par l'étalement d'une lentille liquide piégée entre deux autres phases stratifiées. La solution initiale étant moins régulière que dans les tests présentés dans la section précédente, nous évitons de donner la valeur 0.5 au paramètre β , puisque celle-ci correspond à la limite de stabilité inconditionnelle du schéma de Crank-Nicholson. De plus, pour la même raison, nous utilisons la valeur $\beta = 1$ (*i.e.* une discrétisation implicite du terme de diffusion) pour la première itération en temps (pour le schéma SImpl.(\(\beta\))).

Situation d'étalement partiel

Les valeurs des paramètres sont données dans la table V.2. Notons que dans ce cas, tous les Σ_i , $i = 1, 2, 3$, sont positifs. Ainsi, nous prenons $\Lambda = 0$ (*cf* section V.2.6), de manière que le potentiel de Cahn-Hilliard est $F = F_0$.

Ω	ε	M_0	σ_{12}	σ_{13}	σ_{23}	Σ_1	Σ_2	Σ_3	Λ
$] - 0.3; 0.3[\times] - 0.15; 0.15[$	10^{-2}	10^{-4}	1	0.8	1.4	0.4	1.6	1.2	0

TAB. V.2 – Valeurs des paramètres pour le cas test trois phases en situation d'étalement partiel

La donnée initiale $\mathbf{c}^0$ est définie par

$$\begin{aligned} c_1^0(\mathbf{x}) &= \frac{1}{2} \left[1 + \tanh \left(\frac{2}{\varepsilon} \min(|\mathbf{x}| - 0.1, y) \right) \right], \\ c_2^0(\mathbf{x}) &= \frac{1}{2} \left[1 - \tanh \left(\frac{2}{\varepsilon} \max(-|\mathbf{x}| + 0.1, y) \right) \right], \\ c_3^0(\mathbf{x}) &= 1 - c_1(\mathbf{x}) - c_2(\mathbf{x}), \end{aligned}$$

où $\mathbf{x} = (x, y) \in \Omega$. Celle-ci correspond (*cf* figure V.11) à une bulle sphérique de phase 3 captive entre les deux phases stratifiées 1 et 2. Rappelons que la représentation des interfaces est réalisée grâce à la fonction définie par (V.52).

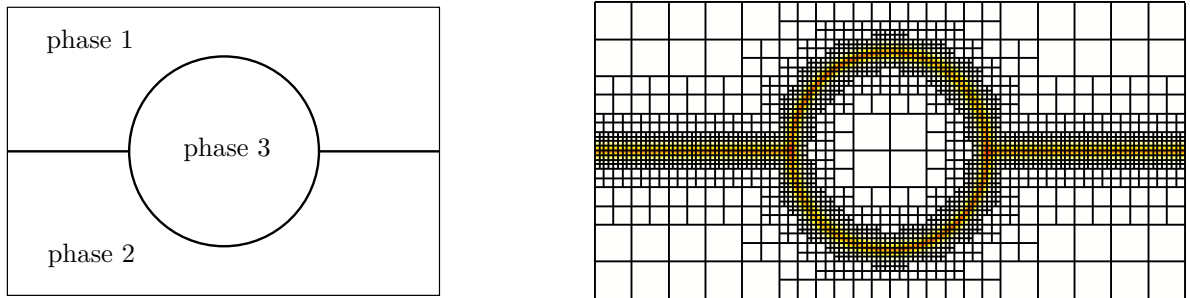


FIG. V.11 – Configuration du cas test (à gauche) et position initiale des interfaces (à droite)

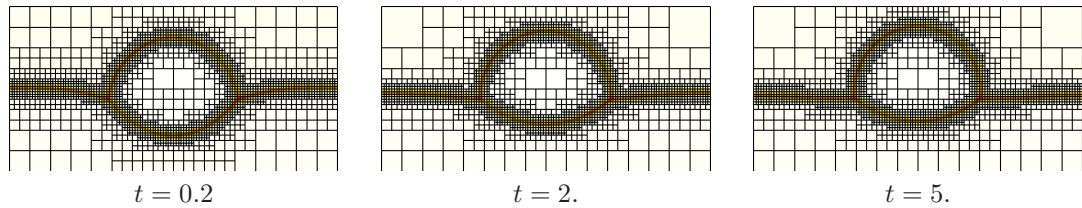
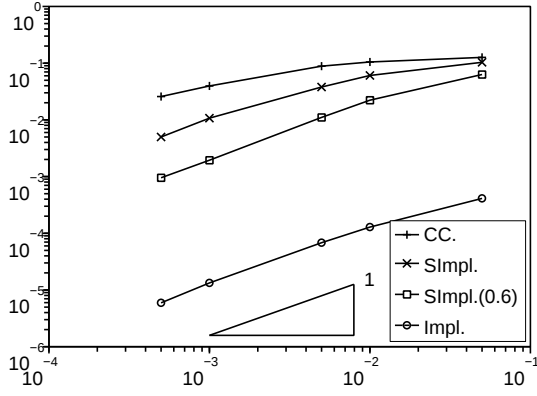


FIG. V.12 – Evolution de la position de l'interface en utilisant le schéma Impl. avec $\Delta t = 10^{-4}$

La figure V.12 montre l'évolution de la position des interfaces. A l'équilibre, la forme attendue de la lentille est l'intersection de deux arcs de cercle (les angles de contact dépendant des trois tensions de surface à travers la loi de Young).



j	Δt_j	$\frac{\ln(e_{j+1}/e_j)}{\ln(\Delta t_{j+1}/\Delta t_j)}$			
		CC.	SImpl.	SImpl.(0.6)	Impl.
1	$5 \cdot 10^{-2}$	0.1	0.3	0.6	0.7
2	10^{-2}	0.2	0.7	1.0	0.9
3	$5 \cdot 10^{-3}$	0.5	0.8	1.1	1.0
4	10^{-3}	0.6	1.1	1.0	1.2
5	$5 \cdot 10^{-4}$	-	-	-	-

FIG. V.13 – Erreurs $e_j(t) = |\mathbf{c}^{\Delta t_j}(t, \cdot) - \mathbf{c}^{\Delta t_{\text{ref}}}(t, \cdot)|_{(L^2(\Omega))^3}$ au temps $t = 2$. en fonction du pas de temps Δt_j (à gauche) et ordres de convergence (à droite) obtenus avec les différents schémas

La figure V.13 présente l'étude de convergence. La solution de référence est calculée en utilisant le schéma Impl. avec $\Delta t_{\text{ref}} = 10^{-4}$. Nous avons réalisé plusieurs calculs avec les différents schémas et les pas de temps Δt_j suivants : $5 \cdot 10^{-2}$, 10^{-2} , $5 \cdot 10^{-3}$, 10^{-3} , $5 \cdot 10^{-4}$. Les normes L^2 de chaque erreur correspondante $e_j(t)$ au temps $t = 2$ sont présentées sur la figure de gauche et les ordres de convergence obtenus sont donnés dans le tableau de droite. Comme attendu, nous obtenons une convergence d'ordre 1 pour les quatre schémas. Néanmoins, le schéma Impl. est le plus précis. Nous observons en particulier trois ordres de grandeur de différence avec le schéma CC.

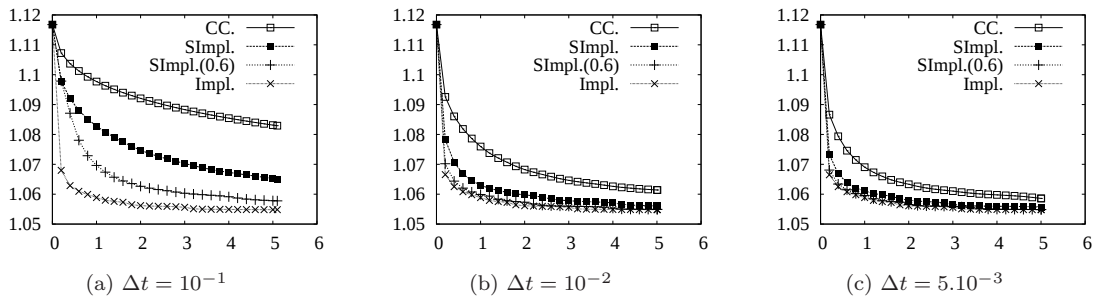
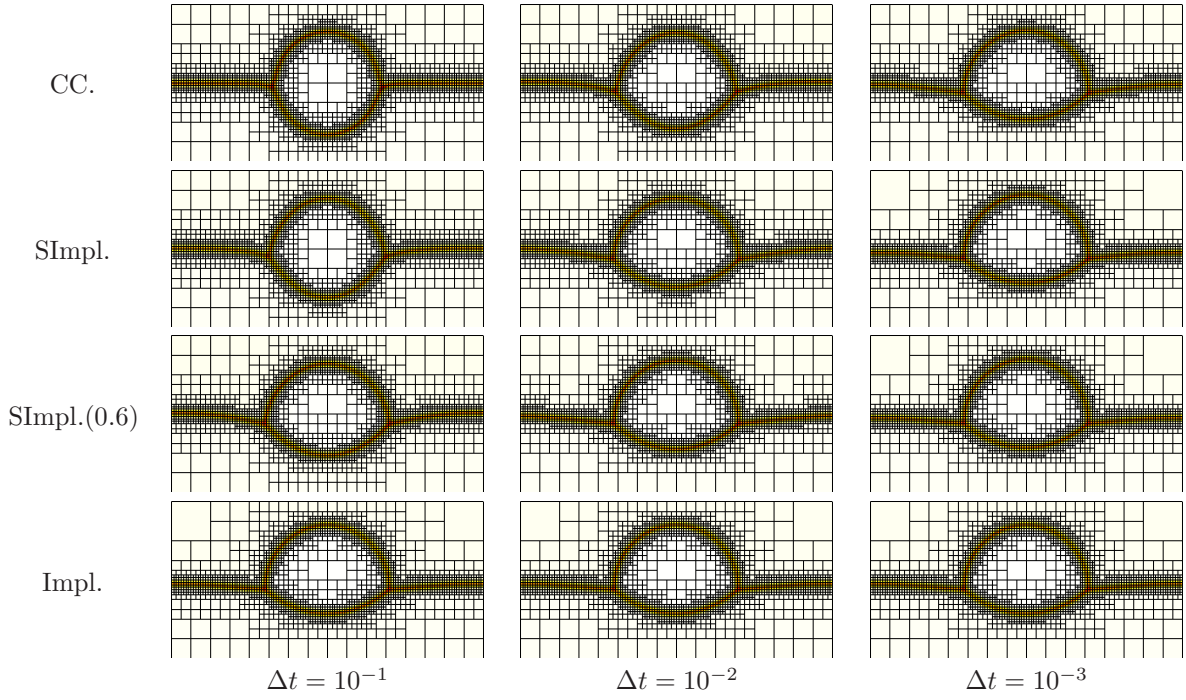


FIG. V.14 – Evolution de l'énergie en fonction du temps en situation d'étalement partiel

La figure V.14 présente les courbes d'énergie discrète $\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^n)$ en fonction du temps $t_n \in [0, t_f]$. Pour chacun des quatre schémas, nous avons réalisé trois simulations avec $\Delta t = 10^{-1}$, 10^{-2} et $5 \cdot 10^{-3}$.

FIG. V.15 – Influence des schémas sur la forme de la bulle à $t = 2$.

La figure V.15 montre l'influence des schémas sur la forme de la bulle à l'instant $t = 2$. Les résultats obtenus à l'aide du schéma Impl. sont très similaires pour les différents pas de temps testés. Pour les grands pas de temps, le schéma CC. ne donne pas la forme de bulle attendue. Ce phénomène est considérablement réduit par l'utilisation des schémas SImpl. ou SImpl.(0.6).

Situation d'étalement total

Les valeurs des paramètres utilisés pour ce cas test sont présentées dans la table V.3.

Ω	ε	M_0	σ_{12}	σ_{13}	σ_{23}	Σ_1	Σ_2	Σ_3	Λ
$] -0.3; 0.3[\times] -0.3; 0.2[$	10^{-2}	10^{-4}	1	1	3	-1	3	3	$7/3$

TAB. V.3 – Valeurs des paramètres pour le cas test trois phases en situation d'étalement total

La donnée initiale $\mathbf{c}^0$ est définie par

$$\begin{aligned}
 c_1^0(\mathbf{x}) &= \frac{1}{2} \left[1 + \tanh \left(\frac{2}{\varepsilon} \min(\sqrt{x^2 + y^2} - 0.1, y) \right) \right], \\
 c_2^0(\mathbf{x}) &= \frac{1}{2} \left[1 - \tanh \left(\frac{2}{\varepsilon} y \right) \right], \\
 c_3^0(\mathbf{x}) &= 1 - c_1(\mathbf{x}) - c_2(\mathbf{x}),
 \end{aligned}$$

où $\mathbf{x} = (x, y) \in \Omega$.

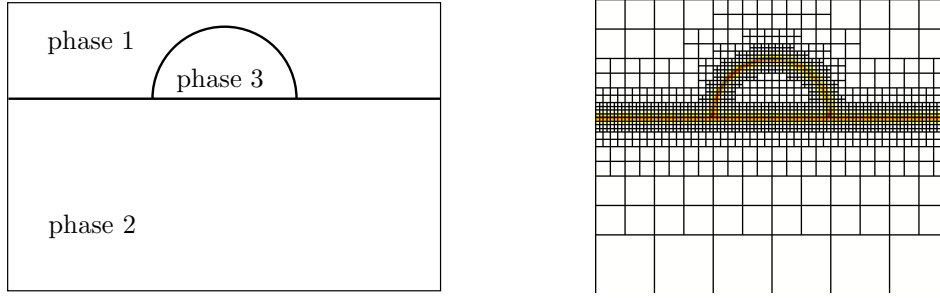
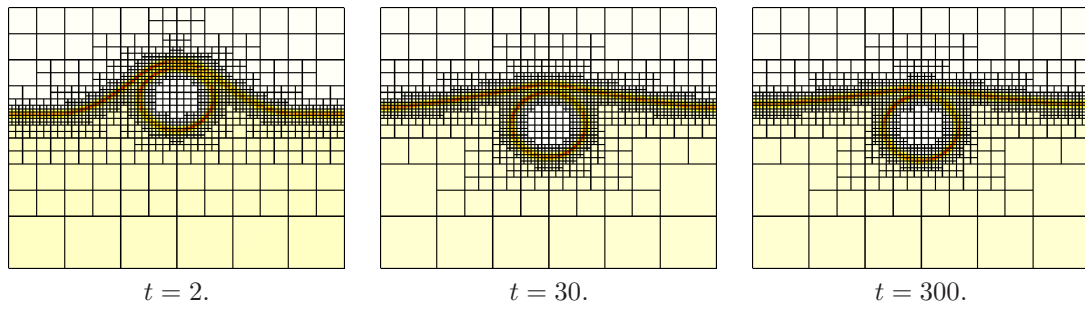
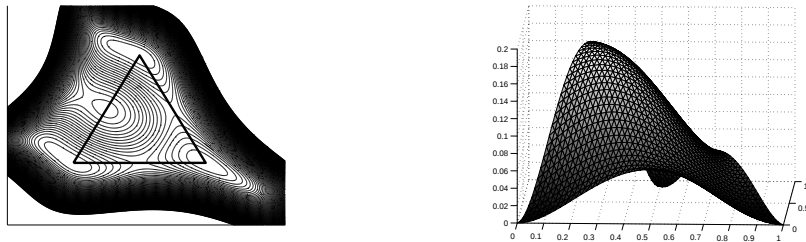


FIG. V.16 – Configuration du cas test (à gauche) et position initiale des interfaces (à droite)

Ceci correspond (*cf* figure V.16) à une bulle de phase 3 initialement posée sur l'interface entre les deux phases stratifiées 1 et 2.

FIG. V.17 – Evolution de la position de l'interface en utilisant le schéma SImpl. avec $\Delta t = 10^{-3}$

Dans ce cas, Σ_1 est négatif mais la condition (IV.14) est vérifiée. Cela correspond au phénomène d'extraction de la bulle (*cf* figure V.17) : à l'état stationnaire la bulle est entièrement dans l'une des deux autres phases. Nous avons choisi Λ constant suffisamment grand pour garantir la positivité du potentiel de Cahn-Hilliard F (*cf* section V.2.6). Nous prenons ici $\Lambda = 7$.

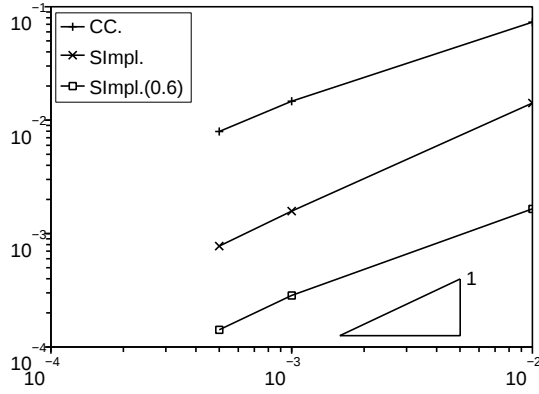
FIG. V.18 – Potentiel de Cahn-Hilliard F en utilisant des coordonnées barycentriques

La figure V.18 montre que le potentiel de Cahn-Hilliard F correspondant à la forme attendue : F n'est pas négatif et a exactement trois minima qui correspondent aux phases pures. Le potentiel F est représenté sur l'hyperplan $\mathcal{S}$ en utilisant des coordonnées barycentriques.

Nous avons réalisé plusieurs simulations en utilisant les différents schémas avec les pas de temps Δt suivants : 10^{-1} , $5 \cdot 10^{-2}$, 10^{-2} , $5 \cdot 10^{-3}$, 10^{-3} , $5 \cdot 10^{-4}$, 10^{-4} . Nous observons que la méthode de linéarisation de Newton ne converge pas lorsque nous utilisons le schéma Impl. à moins que le pas temps soit plus petit que 10^{-4} . La table V.4 donne le nombre maximum d'itérations de la méthode de Newton nécessaires pour arriver à convergence pour tous les pas de temps de la simulation. Le schéma CC. est le plus robuste puisque les calculs se passent bien pour toutes les valeurs des pas temps testées. Les schémas SImpl. et SImpl.(0.6) fonctionnent pour une large gamme de pas de temps.

Schéma \ Δt	10^{-1}	5.10^{-2}	10^{-2}	5.10^{-3}	10^{-3}	5.10^{-4}	10^{-4}
CC.	5	5	5	5	5	5	4
SImpl.	-	-	9	9	6	6	5
SImpl.(0.6)	-	-	29	-	7	6	5
Impl.	-	-	-	-	-	-	7
CPU time	5min	9min	40min	1h10	5h45	11h	53h

TAB. V.4 – Nombre d’itérations de la méthode de Newton. Le symbole “–” signifie que la méthode n’a pas convergé



j	Δt_j	$\frac{\ln(e_{j+1}/e_j)}{\ln(\Delta t_{j+1}/\Delta t_j)}$		
		CC.	SImpl.	SImpl.(0.6)
1	10^{-2}	0.7	1.0	0.8
2	10^{-3}	0.9	1.0	1.0
3	5.10^{-4}	-	-	-

FIG. V.19 – Erreurs $e_j(t) = |\mathbf{c}^{\Delta t_j}(t, \cdot) - \mathbf{c}^{\Delta t_{\text{ref}}}(t, \cdot)|_{(L^2(\Omega))^3}$ au temps $t = 3.8$ en fonction des pas de temps Δt_j (à gauche) et ordres de convergence obtenus (à droite) avec les différents schémas.

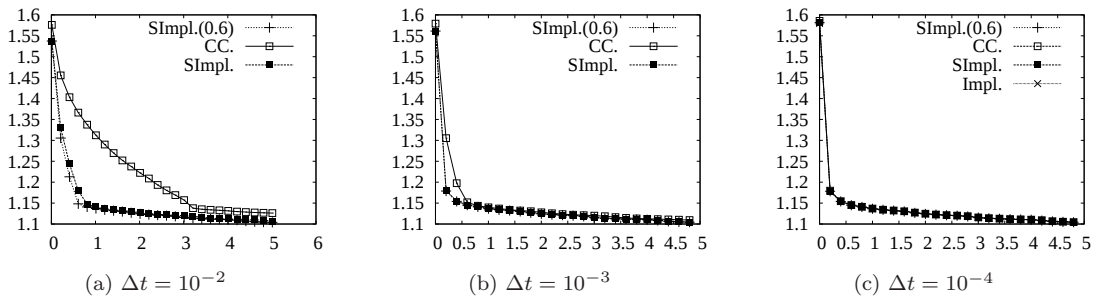


FIG. V.20 – Evolution de l’énergie discrète en fonction du temps pour une situation d’étalement total

Ces résultats doivent être confrontés aux ordres de convergence présentés sur la figure V.19. En effet, le schéma CC. est moins précis que les schémas SImpl. et SImpl.(0.6), même si les trois sont d’ordre 1. Nous pouvons également visualiser les différences entre les schémas grâce à la figure V.20 qui montrent comment l’énergie discrète évolue au cours du temps lorsque l’on utilise les différents schémas. Nous avons réalisé des calculs avec $\Delta t = 10^{-2}$, 10^{-3} , 10^{-4} et nous observons que les schémas SImpl. et SImpl.(0.6) donnent des résultats plus précis que ceux obtenus avec le schéma CC.

V.4 Cas d'une mobilité dégénérée

Dans les sections précédentes de ce chapitre, nous avons supposé que la mobilité M_0 satisfait (IV.19) :

$$\forall \mathbf{c} \in \mathcal{S}, \quad 0 < M_1 \leq M_0(\mathbf{c}) \leq M_2.$$

Il est également intéressant d'utiliser une mobilité qui s'annule dans les phases pures (*i.e.* en 0 et 1). Aussi bien d'un point de vue numérique que théorique, l'étude du système de Cahn-Hilliard devient alors plus complexe. Nous ne réalisons pas cette étude dans ce manuscrit, nous renvoyons aux références [BBG01] et [BBG99]. Nous donnons néanmoins dans cette section, le schéma que nous utilisons dans le cas d'une mobilité dégénérée et montrons qu'il satisfait encore une estimation d'énergie.

En pratique, nous choisissons une mobilité de la forme :

$$M_0(\mathbf{c}) = M_{\text{deg}} \prod_{i=1}^3 (1 - c_i)^2,$$

le coefficient M_{deg} étant une constante strictement positive.

Nous adaptons une idée trouvée dans la référence [BB99] au modèle triphasique considéré dans ce manuscrit :

$$\begin{aligned} &\text{remplacer dans (V.1) le terme } \operatorname{div} \left(M_0(\mathbf{c}^{n+\alpha}) \nabla \mu_i^{n+1} \right) \text{ par le terme} \\ &\operatorname{div} \left(|M_0|_{\infty} \nabla \mu_i^{n+1} + (M_0(\mathbf{c}^{n+\alpha}) - |M_0|_{\infty}) \nabla \mu_i^n \right), \end{aligned}$$

où $|M_0|_{\infty}$ représente une constante supérieure à $\sup_{x \in \Omega} |M_0(\mathbf{c}^{n+\alpha}(x))| > 0$.

Ce terme constitue une discrétisation à l'ordre 1 de $\operatorname{div} \left(M_0(\mathbf{c}) \nabla \mu_i \right)$. Nous pouvons encore l'écrire de la manière suivante :

$$\underbrace{\operatorname{div} \left(M_0(\mathbf{c}^{n+\alpha}) \nabla \mu_{ih}^{n+1} \right)}_{\text{implicite}} + \underbrace{\operatorname{div} \left[(|M_0|_{\infty} - M_0(\mathbf{c}^{n+\alpha})) (\nabla \mu_{ih}^{n+1} - \nabla \mu_{ih}^n) \right]}_{\text{formellement } \leq \Delta t}.$$

Le point clé est que, d'un point de vue numérique, à chaque itération de la méthode de linéarisation (méthode de Newton), la matrice des systèmes linéaires est exactement la même que celle obtenue lorsque la mobilité est constante de valeur $|M_0|_{\infty}$.

Lorsque la mobilité est dégénérée nous utilisons donc le schéma suivant :

Problème V.22 (Formulation à mobilité dégénérée)

Trouver $(\mathbf{c}_h^{n+1}, \mu_h^{n+1}) \in \mathcal{V}_{Dh}^{c_1} \times \mathcal{V}_{Dh}^{c_2} \times \mathcal{V}_{Dh}^{c_3} \times (\mathcal{V}_h^{\mu})^3$ tels que $\forall \nu_h^c \in \mathcal{V}_{Dh,0}^c$, $\forall \nu_h^{\mu} \in \mathcal{V}_h^{\mu}$, nous avons, pour $i = 1, 2, 3$,

$$\begin{cases} \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^{\mu} dx = - \int_{\Omega} \frac{|M_0|_{\infty}}{\Sigma_i} \nabla \mu_{ih}^{n+1} \cdot \nabla \nu_h^{\mu} dx - \int_{\Omega} \frac{M_{0h}^{n+\alpha} - |M_0|_{\infty}}{\Sigma_i} \nabla \mu_{ih}^n \cdot \nabla \nu_h^{\mu} dx, \\ \int_{\Omega} \mu_{ih}^{n+1} \nu_h^c dx = \int_{\Omega} D_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \nu_h^c dx + \int_{\Omega} \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^{n+\beta} \cdot \nabla \nu_h^c dx, \end{cases} \quad (\text{V.53})$$

où $M_{0h}^{n+\alpha} = M_0((1-\alpha)\mathbf{c}_h^n + \alpha\mathbf{c}_h^{n+1})$, $c_{ih}^{n+\beta} = (1-\beta)c_{ih}^n + \beta c_{ih}^{n+1}$ et $|M_0|_{\infty}$ représente une constante supérieure à $\sup_{x \in \Omega} |M_{0h}^{n+\alpha}(x)| > 0$.

Pour mettre en pratique ce schéma, nous devons connaître la valeur de $|M_0|_{\infty}$. Plutôt que de calculer cette valeur à chaque itération en temps, nous utilisons le fait que les valeurs numériques des paramètres d'ordre restent très proches de l'intervalle $[0, 1]$ (même si *a priori* nous ne connaissons pas de bornes L^{∞}), dans le cas contraire, ils perdraient tout sens physique. De plus, la somme des trois paramètres d'ordre est toujours égale à 1. En conséquence, nous pouvons supposer que

$$\prod_{i=1}^3 (1 - c_{ih}^n) \leq 0.1,$$

et nous posons

$$|M_0|_\infty = 0.1 M_{\text{deg}}.$$

L'estimation d'énergie s'obtient en prenant $\nu_h^c = \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t}$ et $\nu_h^\mu = \mu_{ih}^{n+1}$ comme fonction test dans (V.53) (cf démonstration de la proposition V.7). Les termes associés au membre de droite de la première équation du système (V.53) s'écrivent alors :

$$- \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx - \int_{\Omega} \frac{|M_0|_\infty - M_{0h}^{n+\alpha}}{\Sigma_i} \nabla \mu_{ih}^{n+1} \cdot \nabla (\mu_{ih}^{n+1} - \mu_{ih}^n) dx.$$

En utilisant la formule $s(s-t) = \frac{1}{2} [s^2 - t^2 + (s-t)^2]$, ces termes peuvent s'écrire

$$- \int_{\Omega} \frac{|M_0|_\infty + M_{0h}^{n+\alpha}}{2\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx + \int_{\Omega} \frac{|M_0|_\infty - M_{0h}^{n+\alpha}}{2\Sigma_i} |\nabla \mu_{ih}^n|^2 dx - \int_{\Omega} \frac{|M_0|_\infty - M_{0h}^{n+\alpha}}{2\Sigma_i} |\nabla \mu_{ih}^{n+1} - \nabla \mu_{ih}^n|^2 dx.$$

Nous obtenons alors l'estimation d'énergie suivante :

$$\begin{aligned} \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^{n+1}) - \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^n) &+ \frac{\Delta t}{2} |M_0|_\infty \sum_{i=1}^3 \Sigma_i \left[\left| \frac{\nabla \mu_{ih}^{n+1}}{\Sigma_i} \right|_0^2 - \left| \frac{\nabla \mu_{ih}^n}{\Sigma_i} \right|_0^2 \right] \\ &+ \frac{\Delta t}{2} \int_{\Omega} M_{0h}^{n+\alpha} \sum_{i=1}^3 \Sigma_i \left[\left| \frac{\nabla \mu_{ih}^{n+1}}{\Sigma_i} \right|^2 + \left| \frac{\nabla \mu_{ih}^n}{\Sigma_i} \right|^2 \right] dx \\ &+ \frac{\Delta t}{2} \int_{\Omega} [|M_0|_\infty - M_{0h}^{n+\alpha}] \sum_{i=1}^3 \Sigma_i \left[\left| \frac{\nabla \mu_{ih}^{n+1}}{\Sigma_i} - \frac{\nabla \mu_{ih}^n}{\Sigma_i} \right|^2 \right] dx \\ &+ \frac{3}{8} (2\beta - 1) \varepsilon \int_{\Omega} \sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 dx \\ &= \frac{12}{\varepsilon} \int_{\Omega} [F(\mathbf{c}_h^{n+1}) - F(\mathbf{c}_h^n) - \mathbf{d}^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \cdot (\mathbf{c}_h^{n+1} - \mathbf{c}_h^n)] dx. \end{aligned} \quad (\text{V.54})$$

Nous effectuons quelques tests numériques diphasiques académiques pour illustrer l'utilité de ce schéma. L'objectif est de comparer les résultats obtenus à ceux donnés par une autre stratégie possible (proposée dans [Lap06]) qui consiste à conserver le schéma plus standard décrit par le problème V.2 mais à ajouter une constante à la mobilité (qui n'est alors plus dégénérée) :

$$M_0(\mathbf{c}) = M_{\text{cst}} + M_{\text{deg}} \prod_{i=1}^3 (1 - c_i)^2.$$

Nous effectuons des simulations pour plusieurs valeurs du coefficient M_{cst} . Dans toutes les simulations, nous utilisons la discrétisation semi-implicite (V.44) des termes non linéaires et le paramètre β est choisi égal à 0.6 (nous effectuons les 5 premières itérations en temps avec $\beta = 1$). Nous noterons le schéma standard SImpl.(0.6) et le schéma décrit dans cette section SImpl(0.6,m).

La donnée initiale est représentée sur la figure V.21.

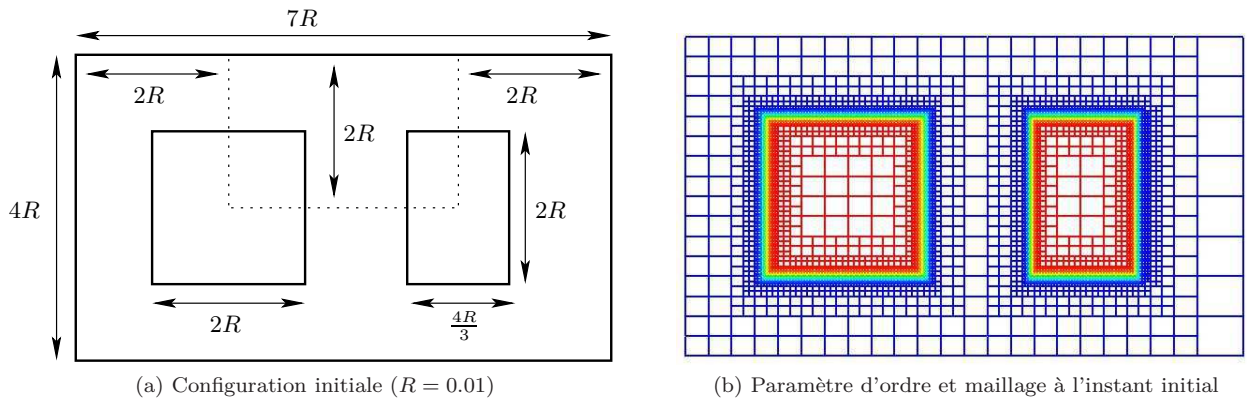


FIG. V.21 – Configuration et maillage initiaux

Le pas de temps Δt est fixé à 10^{-5} , l'épaisseur d'interface ε à 10^{-3} et la tension de surface σ entre les deux fluides à 1. Le maillage initial avant raffinement est rectangle structuré constitué de 12×8 cellules. Nous choisissons le même critère de raffinement que dans la section V.3 avec $h_{\min} = \frac{\varepsilon}{2}$.

Les systèmes linéaires à chaque itération de la méthode de Newton sont résolus par un solveur itératif : GMRES (5000 itérations maximum).

Des tests sont effectués sur le comportement du système lorsque la mobilité est dégénérée ou non. Les résultats sont présentés dans la figure V.22 : à mobilité constante, la configuration des phases à la fin de la simulation est constituée d'un seul disque (à gauche) alors qu'à mobilité dégénérée (à droite) elle est constituée de deux disques distincts. Le fait que la mobilité soit dégénérée diminue fortement la diffusion dans les phases pures. La figure V.22 indique quel type de comportement correspond aux différentes valeurs de mobilité testées.

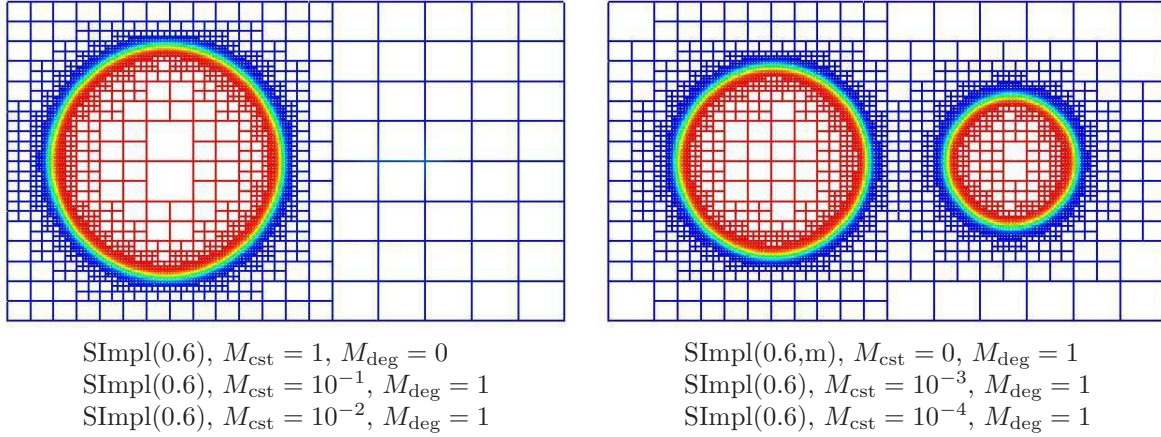


FIG. V.22 – Paramètre d'ordre après 100 pas de temps

Il est important de noter que pour un certain nombre de valeurs de la mobilité, lorsque nous utilisons le schéma SImpl(0.6), le solveur itératif ne converge pas. Ces valeurs sont indiquées ci-dessous :

- SImpl(0.6), $M_{\text{cst}} = 0$, $M_{\text{deg}} = 1$,
- SImpl(0.6), $M_{\text{cst}} = 10^{-5}$, $M_{\text{deg}} = 1$,
- SImpl(0.6), $M_{\text{cst}} = 10^{-6}$, $M_{\text{deg}} = 1$.

Le schéma décrit dans cette section permet donc d'obtenir des résultats qualitativement corrects sans ajouter un paramètre supplémentaire M_{cst} dont le choix peut se révéler délicat : lorsque ce paramètre est trop faible nous avons les mêmes difficultés numériques que dans le cas où la mobilité est dégénérée, lorsqu'il est trop grand, le système se comporte comme si la mobilité était constante.

V.5 Démonstrations des théorèmes d'existence et de convergence des solutions approchées

Rappelons tout d'abord le résultat suivant qui sera très utile dans la suite :

Lemme V.23 (Inégalité de Poincaré)

Soit θ une fonction de $H^1(\Omega)$ donnée telle que $m(\theta) \neq 0$. Il existe une constante $C_{p,\theta} > 0$ telle que

$$\forall \nu \in H^1(\Omega), \quad |\nu|_{H^1(\Omega)} \leq C_{p,\theta} \left[|\nabla \nu|_{L^2(\Omega)} + |m(\nu\theta)| \right]. \quad (\text{V.55})$$

V.5.1 Démonstration du théorème V.9

Nous allons prouver l'existence de solutions au problème (V.8). Les points clés sont les estimations *a priori* données par l'estimation d'énergie discrète et le lemme suivant issu de la théorie du degré topologique [Dei85].

Lemme V.24 (Degré topologique)

Soit W un espace vectoriel de dimension finie et G une fonction continue de W dans W . Supposons qu'il existe une fonction continue H de $W \times [0; 1]$ dans W satisfaisant :

- (i) $H(\cdot, 1) = G$ et $H(\cdot, 0)$ est affine,
 - (ii) $\exists R > 0$ tel que $\forall (w, \delta) \in W \times [0; 1]$, si $H(w, \delta) = 0$ alors $|w|_W \neq R$,
 - (iii) l'équation $H(w, 0) = 0$ a une solution $w \in W$ telle que $|w|_W < R$,
- Alors il existe au moins une solution $w \in W$ telle que $G(w) = 0$ et $|w|_W < R$.*

Reformulation du problème

Soit W l'espace vectoriel de dimension finie $\left(\mathcal{V}_{Dh,0}^c\right)^2 \times \left(\mathcal{V}_h^\mu\right)^2$. Nous définissons la norme suivante sur W ,

$$|w|_W^2 = |\tilde{c}_{1h}|_{H^1(\Omega)}^2 + |\tilde{c}_{2h}|_{H^1(\Omega)}^2 + |\mu_{1h}|_{H^1(\Omega)}^2 + |\mu_{2h}|_{H^1(\Omega)}^2, \quad \forall w = (\tilde{c}_{1h}, \tilde{c}_{2h}, \mu_{1h}, \mu_{2h}) \in W,$$

et nous introduisons la fonction H telle que

$$H : W \times [0; 1] \rightarrow W$$

$$(w^{n+1}, \delta) = (\tilde{c}_{1h}^{n+1}, \tilde{c}_{2h}^{n+1}, \mu_{1h}^{n+1}, \mu_{2h}^{n+1}, \delta) \mapsto (\mathcal{R}_\delta^{\mu_1}, \mathcal{R}_\delta^{c_1}, \mathcal{R}_\delta^{\mu_2}, \mathcal{R}_\delta^{c_2})$$

où $\mathcal{R}_\delta^{\mu_1}$, $\mathcal{R}_\delta^{c_1}$, $\mathcal{R}_\delta^{\mu_2}$ et $\mathcal{R}_\delta^{c_2}$ sont définis par leurs coordonnées dans la base éléments finis $(\nu_I^c)_{I \in \llbracket 1, N^c \rrbracket}$ (resp. $(\nu_I^\mu)_{I \in \llbracket 1, N^\mu \rrbracket}$) de $\mathcal{V}_{Dh,0}^c$ (resp. $\mathcal{V}_h^\mu$) :

$$\text{pour } I \in \llbracket 1, N^\mu \rrbracket, \quad (\mathcal{R}_\delta^{\mu_i})_I = \int_\Omega \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_I^\mu dx + \int_\Omega \frac{M_{0h\delta}^{n+\alpha}}{\Sigma_i} \nabla \mu_{ih}^{n+1} \cdot \nabla \nu_I^\mu dx,$$

$$\text{pour } I \in \llbracket 1, N^c \rrbracket, \quad (\mathcal{R}_\delta^{c_i})_I = \int_\Omega \mu_{ih}^{n+1} \nu_I^c dx - \int_\Omega \delta D_i(c_h^n, c_h^{n+1}) \nu_I^c dx - \int_\Omega \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^{n+\beta} \cdot \nabla \nu_I^c dx,$$

avec $c_{ih}^{n+1} = \tilde{c}_i^{n+1} + c_{iDh}$, $c_{ih}^n = \tilde{c}_i^n + c_{iDh}$ et $M_{0h\delta}^{n+\alpha} = M_0((1 - \delta\alpha)\mathbf{c}_h^n + \delta\alpha\mathbf{c}_h^{n+1})$. La fonction G est définie

$$G : W \rightarrow W$$

$$w \mapsto H(w, 1)$$

Le problème "Trouver w^{n+1} tel que $G(w^{n+1}) = 0$ " est équivalent au problème (V.8). Pour démontrer le théorème, nous allons montrer que les fonctions H et G satisfont les hypothèses du lemme V.24. La continuité de la fonction H est obtenue en utilisant (V.20) et le théorème de Lebesgue. La fonction $H(\cdot, 0)$ est clairement affine par construction.

Validation de l'hypothèse (ii) du lemme V.24

Soit $(w^{n+1}, \delta) \in W \times [0; 1]$ tel que $H(w^{n+1}, \delta) = 0$. Remarquons que $H(w^{n+1}, \delta) = 0$ revient à dire que $w^{n+1} = (\tilde{c}_{1h}^{n+1}, \tilde{c}_{2h}^{n+1}, \mu_{1h}^{n+1}, \mu_{2h}^{n+1})$ est solution d'un problème similaire à (V.10) avec δF à la place de F , $\delta \mathbf{d}^F(\mathbf{c}^n, \mathbf{c}^{n+1})$ comme choix de discrétisation des termes non linéaires et une mobilité un peu modifiée. Il est possible d'appliquer le théorème V.7 (la modification de la mobilité M_{0h} ne changeant pas les calculs). Nous obtenons l'égalité suivante

$$\mathcal{F}_{\Sigma, \varepsilon, \delta}^{\text{triph}}(\mathbf{c}_h^{n+1}) - \mathcal{F}_{\Sigma, \varepsilon, \delta}^{\text{triph}}(\mathbf{c}_h^n) + \Delta t \sum_{i=1}^3 \int_\Omega \frac{M_{0h\delta}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx + \frac{3}{8} (2\beta - 1) \varepsilon \int_\Omega \sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 dx =$$

$$\frac{12}{\varepsilon} \delta \int_\Omega [F(\mathbf{c}_h^{n+1}) - F(\mathbf{c}_h^n) - \mathbf{d}^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \cdot (\mathbf{c}_h^{n+1} - \mathbf{c}_h^n)] dx,$$

avec $\mathcal{F}_{\Sigma, \varepsilon, \delta}^{\text{triph}}(\mathbf{c}_h^k) = \int_{\Omega} \delta \frac{12}{\varepsilon} F(\mathbf{c}_h^k) + \sum_{i=1}^3 \frac{3}{8} \varepsilon \Sigma_i |\nabla c_{ih}^k|^2 dx$. En utilisant l'hypothèse (V.21) et la remarque V.8, nous obtenons

$$\mathcal{F}_{\Sigma, \varepsilon, \delta}^{\text{triph}}(\mathbf{c}_h^{n+1}) + \Delta t \int_{\Omega} M_{0h\delta}^{n+\alpha} \sum_{i=1}^3 \frac{|\nabla \mu_{ih}^{n+1}|^2}{\Sigma_i} dx \leq \mathcal{F}_{\Sigma, \varepsilon, \delta}^{\text{triph}}(\mathbf{c}_h^n) + \delta \frac{12}{\varepsilon} K_1^{\mathbf{c}_h^n}. \quad (\text{V.56})$$

Puisque la mobilité est minorée (hypothèse (IV.19)) et d'après la remarque V.8, le second terme du membre de gauche de (V.56) est minoré :

$$\int_{\Omega} M_1 \underline{\Sigma} \sum_{i=1}^3 \frac{|\nabla \mu_{ih}^{n+1}|^2}{\Sigma_i^2} dx \leq \int_{\Omega} M_{0h\delta}^{n+\alpha} \sum_{i=1}^3 \frac{|\nabla \mu_{ih}^{n+1}|^2}{\Sigma_i} dx. \quad (\text{V.57})$$

De plus, puisque $F \geq 0$ et $\delta \leq 1$, nous avons

$$\mathcal{F}_{\Sigma, \varepsilon, \delta}^{\text{triph}}(\mathbf{c}_h^k) \leq \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^k), \quad (\text{V.58})$$

et d'après (V.56), (V.57), (V.58) et la proposition IV.5, il existe une constante $K_2^{\mathbf{c}_h^n} = \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^n) + \frac{12}{\varepsilon} K_1^{\mathbf{c}_h^n} > 0$ indépendante de δ et $\mathbf{c}_h^{n+1}$ tel que

$$\int_{\Omega} \delta \frac{12}{\varepsilon} F(\mathbf{c}_h^{n+1}) + \frac{3}{8} \varepsilon \underline{\Sigma} \sum_{i=1}^3 |\nabla c_{ih}^{n+1}|^2 dx + \Delta t \int_{\Omega} M_1 \underline{\Sigma} \sum_{i=1}^3 \frac{|\nabla \mu_{ih}^{n+1}|^2}{\Sigma_i^2} dx \leq K_2^{\mathbf{c}_h^n}. \quad (\text{V.59})$$

Puisque F est positif et $\delta \geq 0$, nous obtenons la borne suivante pour les second et troisième termes du membre de gauche de (V.59) : pour $i = 1, 2, 3$,

$$|\nabla c_{ih}^{n+1}|_{L^2} \leq \sqrt{\frac{8}{3} \frac{K_2^{\mathbf{c}_h^n}}{\varepsilon \underline{\Sigma}}} := K_3^{\mathbf{c}_h^n} \quad \text{et} \quad |\nabla \mu_{ih}^{n+1}|_{L^2} \leq \max_{i=1,2,3} (|\Sigma_i|) \sqrt{\frac{K_2^{\mathbf{c}_h^n}}{M_1 \underline{\Sigma} \Delta t}} := K_4^{\mathbf{c}_h^n}.$$

Nous utilisons maintenant la forme discrète de la conservation du volume (V.9) : $m(c_{ih}^{n+1}) = m(c_{ih}^n)$. Ainsi, grâce à l'inégalité de Poincaré (V.55) (avec $\theta \equiv 1$), il existe une constante positive C_P telle que

$$|c_{ih}^{n+1}|_{H^1(\Omega)} \leq C_P (|\nabla c_{ih}^{n+1}|_{L^2} + m(c_{ih}^{n+1})) = C_P (|\nabla c_{ih}^{n+1}|_{L^2} + m(c_{ih}^n)),$$

et il existe une constante positive $K_5^{\mathbf{c}_h^n} = C_P (K_3^{\mathbf{c}_h^n} + m(c_{ih}^n))$ indépendante de δ et $\mathbf{c}_h^{n+1}$ telle que

$$|c_{ih}^{n+1}|_{H^1(\Omega)} \leq K_5^{\mathbf{c}_h^n}. \quad (\text{V.60})$$

Il reste à borner la moyenne $m(\mu_{ih}^{n+1})$. A cause des conditions au bord de Dirichlet imposées à $\mathbf{c}$, les constantes n'appartiennent pas à l'espace d'approximation $\mathcal{V}_{Dh,0}^c$. Donc, nous prenons une fonction fixée θ_h de $\mathcal{V}_{Dh,0}^c$ telle que $m(\theta_h) \neq 0$. Puisque $\mathcal{R}_{\delta}^{c_i} = 0$, nous avons

$$m(\mu_{ih}^{n+1} \theta_h) = \int_{\Omega} \delta D_i^F(\mathbf{c}_h^{n+1}, \mathbf{c}_h^n) \theta_h dx + \int_{\Omega} \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^{n+\beta} \cdot \nabla \theta_h dx.$$

Ceci peut être contrôlé par $|\mathbf{c}_h^{n+1}|_{H^1(\Omega)}$ et $|\mathbf{c}_h^n|_{H^1(\Omega)}$ sous l'hypothèse (V.20). En effet, la croissance polynomiale (V.20) de d_i^F implique qu'il existe une constante positive $C_1 = \frac{16 \Sigma_T}{3 \Sigma_m} B_1$ telle que

$$|D_i^F(\mathbf{c}_h^{n+1}, \mathbf{c}_h^n)| \leq C_1 (1 + |\mathbf{c}_h^{n+1}|^{p-1} + |\mathbf{c}_h^n|^{p-1}).$$

Ainsi, puisque $\delta \leq 1$, et en utilisant (V.60), nous obtenons

$$\begin{aligned} m(\mu_{ih}^{n+1} \theta_h) &\leq C_1 |\theta_h|_{L^\infty(\Omega)} \left(|\Omega| + |\mathbf{c}_h^{n+1}|_{L^{p-1}}^{p-1} + |\mathbf{c}_h^n|_{L^{p-1}}^{p-1} \right) + \frac{3}{4} \Sigma_M \varepsilon (|\nabla c_{ih}^n|_{L^2} + |\nabla c_{ih}^{n+1}|_{L^2}) |\nabla \theta_h|_{L^2} \\ &\leq C_1 |\theta_h|_{L^\infty(\Omega)} \left(|\Omega| + (K_5^{\mathbf{c}_h^n})^{p-1} + |\mathbf{c}_h^n|_{H^1}^{p-1} \right) + \frac{3}{4} \Sigma_M \varepsilon (|c_{ih}^n|_{H^1} + K_5^{\mathbf{c}_h^n}) |\theta_h|_{H^1} := K_6^{\mathbf{c}_h^n}. \end{aligned}$$

Grâce à l'inégalité de Poincaré (V.55), il existe une constante C_{p,θ_h} telle que

$$|\mu_{ih}^{n+1}|_{H^1(\Omega)} \leq C_{p,\theta_h} (|\nabla \mu_{ih}^{n+1}|_{L^2} + m(\mu_{ih}^{n+1} \theta_h)) \leq C_{p,\theta_h} (K_4^{c_h^n} + K_6^{h,c_h^n}). \quad (V.61)$$

Ainsi, en combinant (V.60) et (V.61), nous obtenons une constante positive $K^{c_h^n}$ indépendante de δ et $\mathbf{c}_h^{n+1}$ telle que

$$|w^{n+1}|_W \leq K^{c_h^n}.$$

Donc, prendre $R > K^{c_h^n} \geq 0$ garantit que pour tout $(w, \delta) \in W \times [0; 1]$, $H(w, \delta) = 0 \implies |w|_W \neq R$.

Validation de l'hypothèse (iii) du lemme V.24

Nous devons montrer l'existence d'une solution au problème linéaire $H(w^{n+1}, 0) = 0$. Ce problème peut être écrit sous la forme variationnelle :

Trouver $(\tilde{c}_h^{n+1}, \mu_h^{n+1}) \in (\mathcal{V}_{Dh,0}^c)^3 \times (\mathcal{V}_h^\mu)^3$ telle que $\forall i = 1, 2, 3, \forall \nu_h^\mu \in \mathcal{V}_h^\mu, \forall \nu_h^c \in \mathcal{V}_{Dh,0}^c$,

$$a_i((\tilde{c}_h^{n+1}, \mu_h^{n+1}), (\nu_h^c, \nu_h^\mu)) = \int_{\Omega} \tilde{c}_{ih}^n \nu_h^\mu dx - \int_{\Omega} \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^n \cdot \nabla \nu_h^c dx,$$

où

$$a_i((\tilde{c}_h^{n+1}, \mu_h^{n+1}), (\nu_h^c, \nu_h^\mu)) = \int_{\Omega} \left[\tilde{c}_{ih}^{n+1} \nu_h^\mu + \frac{M_{0h}^n}{\Sigma_i} \Delta t \nabla \mu_{ih}^{n+1} \cdot \nabla \nu_h^\mu \right] dx + \int_{\Omega} \left[\frac{3}{4} \Sigma_i \varepsilon \beta \nabla \tilde{c}_{ih}^{n+1} \cdot \nabla \nu_h^c - \mu_{ih}^{n+1} \nu_h^c \right] dx,$$

avec

$$M_{0h}^n = M_0(\mathbf{c}_h^n) \quad \text{et} \quad c_{ih}^n = \beta c_{iDh} + (1 - \beta) c_{ih}^n.$$

Puisque ce problème linéaire est posé en dimension finie, il est suffisant de montrer que, pour tout $(\tilde{c}_{ih}^{n+1}, \mu_{ih}^{n+1}) \in (\mathcal{V}_{Dh,0}^c)^3 \times (\mathcal{V}_h^\mu)^3$:

$$(a_i((\tilde{c}_{ih}^{n+1}, \mu_{ih}^{n+1}), (\nu_h^c, \nu_h^\mu)) = 0, \quad \forall (\nu_h^c, \nu_h^\mu) \in (\mathcal{V}_{Dh,0}^c)^3 \times (\mathcal{V}_h^\mu)^3) \implies (\tilde{c}_{ih}^{n+1}, \mu_{ih}^{n+1}) = (0, 0).$$

Soit $(\tilde{c}_{ih}^{n+1}, \mu_{ih}^{n+1}) \in (\mathcal{V}_{Dh,0}^c)^3 \times (\mathcal{V}_h^\mu)^3$ tel que

$$(a_i((\tilde{c}_{ih}^{n+1}, \mu_{ih}^{n+1}), (\nu_h^c, \nu_h^\mu)) = 0, \quad \forall (\nu_h^c, \nu_h^\mu) \in (\mathcal{V}_{Dh,0}^c)^3 \times (\mathcal{V}_h^\mu)^3), \quad (V.62)$$

Prenons $(\nu_h^c, \nu_h^\mu) = (\tilde{c}_{ih}^{n+1}, \mu_{ih}^{n+1})$ dans (V.62), nous obtenons :

$$\int_{\Omega} \tilde{c}_{ih}^{n+1} \mu_{ih}^{n+1} dx + \int_{\Omega} \frac{M_{0h}^n}{\Sigma_i} \Delta t |\nabla \mu_{ih}^{n+1}|^2 dx + \frac{3}{4} \Sigma_i \varepsilon \beta \int_{\Omega} |\nabla \tilde{c}_{ih}^{n+1}|^2 dx - \int_{\Omega} \mu_{ih}^{n+1} \tilde{c}_{ih}^{n+1} dx = 0.$$

C'est équivalent à

$$\int_{\Omega} M_{0h}^n \Delta t |\nabla \mu_{ih}^{n+1}|^2 dx + \frac{3}{4} \Sigma_i^2 \varepsilon \beta \int_{\Omega} |\nabla \tilde{c}_{ih}^{n+1}|^2 dx = 0.$$

Puisque la mobilité satisfait (IV.19), nous obtenons : $\nabla \mu_{ih}^{n+1} = \nabla \tilde{c}_{ih}^{n+1} = 0$. Donc, $\tilde{c}_{ih}^{n+1}$ et μ_{ih}^{n+1} sont constants. En ré-injectant ces constantes dans (V.62), nous obtenons

$$(\tilde{c}_{ih}^{n+1}, \mu_{ih}^{n+1}) = (0, 0).$$

V.5.2 Preuve du théorème V.10

Bornes sur les solutions discrètes

L'égalité (V.26) permet d'obtenir des bornes sur les solutions discrètes : nous pouvons obtenir une borne en norme $L^\infty(0, t_f, H^1(\Omega))$ discrète pour les paramètres d'ordre, en norme $L^2(0, t_f, H^1(\Omega))$ discrète pour les potentiels chimiques et en norme $L^2(0, t_f, (H^1(\Omega))')$ discrète pour les dérivées en temps des paramètres d'ordre.

De plus, la présence des termes de diffusion numérique dans l'égalité (V.26) permet de montrer que les dérivées en temps des paramètres d'ordre croissent au plus comme $\frac{1}{\sqrt{\Delta t}}$ en norme $L^2(0, t_f, H^1(\Omega))$.

Proposition V.25

Supposons que les hypothèses du théorème d'existence V.9 sont satisfaites. Alors, il existe $h_0 > 0$ et des constantes positives K_1, K_2 , indépendantes de Δt et h telle que, pour tout $h \leq h_0$, nous avons

$$\begin{aligned} & \left(\sup_{n \leq N} |\mathbf{c}_h^n|_{(H^1(\Omega))^3} \right) + \left(\sum_{n=0}^{N-1} \Delta t \sum_{i=1}^3 |\mu_{ih}^{n+1}|_{H^1(\Omega)}^2 \right) \leq K_1, \\ & \left(\sum_{n=0}^{N-1} \Delta t \sum_{i=1}^3 \left| \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \right|_{(H^1(\Omega))'}^2 \right) + \Delta t \left(\sum_{n=0}^{N-1} \Delta t \sum_{i=1}^3 \left| \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \right|_{H^1(\Omega)}^2 \right) \leq K_2. \end{aligned}$$

Démonstration : Soit $\Sigma_m = \min_{i=1,2,3} |\Sigma_i|$ et $\Sigma_M = \max_{i=1,2,3} |\Sigma_i|$.

(i) L'estimation d'énergie discrète (V.26), donne en particulier une borne uniforme sur l'énergie discrète :

$$\forall n \in \llbracket 0, N \rrbracket, \quad \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^n) \leq \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^0). \quad (\text{V.63})$$

De plus, grâce à l'hypothèse de croissance polynomiale (IV.20) de F , l'énergie initiale $\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^0)$ peut être majorée indépendamment de h :

$$\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^0) \leq B_1 \left(|\Omega| + |\mathbf{c}_h^0|_{L^p}^p \right) + \Sigma_M |\mathbf{c}_h^0|_{H^1}^2 \leq B_1 \left(|\Omega| + |\mathbf{c}^0|_{H^1}^p \right) + \Sigma_M |\mathbf{c}^0|_{H^1}^2 := K_0. \quad (\text{V.64})$$

Puisque F est positive et en utilisant la proposition IV.5, la borne (V.63) donne en particulier,

$$\forall n \in \llbracket 0, N \rrbracket, \quad \int_{\Omega} \sum_{i=1}^3 |\nabla c_{ih}^n|^2 dx \leq \frac{8}{3\varepsilon \Sigma} K_0. \quad (\text{V.65})$$

De plus, la forme discrète de la conservation du volume (V.9) mène à

$$\forall n \in \mathbb{N}, \quad |m(c_{ih}^n)| \leq |\Omega|^{-\frac{1}{2}} |c_{ih}^0|_{L^2} \leq |\Omega|^{-\frac{1}{2}} |c_i^0|_{H^1} \quad (\text{V.66})$$

Donc, en utilisant (V.65), (V.66) et l'inégalité de Poincaré (V.55), nous trouvons que

$$\forall n \in \llbracket 0, N \rrbracket, \quad |\mathbf{c}_h^n|_{H^1(\Omega)} \leq C_p \left(\frac{16}{3\varepsilon \Sigma_m} K_0 + \frac{2}{|\Omega|} \sum_{i=1}^3 |c_i^0|_{H^1}^2 \right)^{\frac{1}{2}} := K'_1. \quad (\text{V.67})$$

(ii) Maintenant nous ajoutons les équations (V.26) pour n de 0 à $N-1$:

$$\begin{aligned} & \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^N) - \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^0) \\ & + C \left[\sum_{n=0}^{N-1} \Delta t \sum_{i=1}^3 \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx + \frac{3}{8} (2\beta - 1) \varepsilon \int_{\Omega} \sum_{n=0}^{N-1} \sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 dx \right] \leq 0. \end{aligned} \quad (\text{V.68})$$

Puisque F est positive et que la mobilité est minorée, (V.68) donne en particulier

$$\sum_{n=0}^{N-1} \Delta t \sum_{i=1}^3 \int_{\Omega} |\nabla \mu_{ih}^{n+1}|^2 dx \leq \frac{2\Sigma_M}{M_1} K_0. \quad (\text{V.69})$$

Soit θ une fonction positive donnée de $H^1(\Omega)$ à support compact dans Ω . Nous notons θ_h sa projection H^1 sur $\mathcal{V}_{hD,0}^c$ et nous prenons $\nu_h^c = \theta_h$ comme fonction test dans la seconde équation de (V.8). Nous obtenons

$$|\Omega| m(\mu_{ih}^{n+1} \theta_h) = \int_{\Omega} D_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \theta_h dx + \int_{\Omega} \frac{3}{4} \Sigma_i \varepsilon [(1 - \beta) \nabla c_{ih}^n + \beta \nabla c_{ih}^{n+1}] \cdot \nabla \theta_h dx.$$

Donc, nous déduisons que

$$\begin{aligned} |\Omega| |m(\mu_{ih}^{n+1} \theta_h)| &\leq \frac{4\Sigma_T}{\varepsilon} \sum_{j \neq i} \left(\frac{1}{|\Sigma_j|} \left(\int_{\Omega} |d_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1})| |\theta_h| dx + \int_{\Omega} |d_j^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1})| |\theta_h| dx \right) \right) \\ &\quad + \frac{3}{4} |\Sigma_i| \varepsilon \left[(1 - \beta) \int_{\Omega} |\nabla c_{ih}^n| |\nabla \theta_h| dx + \beta \int_{\Omega} |\nabla c_{ih}^{n+1}| |\nabla \theta_h| dx \right], \end{aligned}$$

Le premier terme peut être borné en utilisant (V.20), (V.67) et l'inégalité de Hölder :

$$\begin{aligned} &\int_{\Omega} |d_k^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1})| |\theta_h| dx \\ &\leq B_1 \left(\int_{\Omega} |\theta_h| dx + \int_{\Omega} |\mathbf{c}_h^{n+1}|^{p-1} |\theta_h| dx + \int_{\Omega} |\mathbf{c}_h^n|^{p-1} |\theta_h| dx \right) \\ &\leq B_1 \left(|\mathbf{c}_h^{n+1}|_{L^6(\Omega)}^{p-1} + |\mathbf{c}_h^n|_{L^6(\Omega)}^{p-1} \right) |\theta_h|_{L^{\frac{6}{7-p}}(\Omega)} + B_1 |\theta_h|_{L^1(\Omega)} \\ &\leq B_1 C_{S,6} \left(|\mathbf{c}_h^{n+1}|_{H^1(\Omega)}^{p-1} + |\mathbf{c}_h^n|_{H^1(\Omega)}^{p-1} \right) |\theta_h|_{H^1(\Omega)} + B_1 C_{S,1} |\theta_h|_{H^1(\Omega)} \\ &\leq B_1 \left[2C_{S,6} (K'_1)^{p-1} + B_1 C_{S,1} \right] |\theta|_{H^1(\Omega)}, \end{aligned}$$

et nous obtenons

$$\begin{aligned} |m(\mu_{ih}^{n+1} \theta_h)| &\leq \frac{1}{|\Omega|} \frac{16\Sigma_T}{\varepsilon |\Sigma_j|} \left(2B_1 C_{S,6}^2 (K'_1)^{p-1} |\theta|_{H^1(\Omega)} \right) \\ &\quad + \frac{3}{4} |\Sigma_i| \varepsilon \left[(1 - \beta) |c_{ih}^n|_{H^1(\Omega)} |\theta_h|_{H^1(\Omega)} + \beta |c_{ih}^{n+1}|_{H^1(\Omega)} |\theta_h|_{H^1(\Omega)} \right] := M_1^\theta. \end{aligned}$$

Finalement, nous trouvons

$$|m(\mu_{ih}^{n+1} \theta)| \leq \frac{1}{|\Omega|} \int_{\Omega} |\mu_{ih}^{n+1}| |\theta - \theta_h| dx + |m(\mu_{ih}^{n+1} \theta_h)| \leq \frac{1}{|\Omega|} |\mu_{ih}^{n+1}|_{H^1(\Omega)} |\theta - \theta_h|_{L^2(\Omega)} + M_1^\theta,$$

et l'inégalité de Poincaré (V.55) donne

$$\left[1 - \frac{C_{p,\theta}}{|\Omega|} |\theta - \theta_h|_{L^2(\Omega)} \right] |\mu_{ih}^{n+1}|_{H^1(\Omega)} \leq C_{p,\theta} \left[|\nabla \mu_{ih}^{n+1}|_{L^2(\Omega)} + M_1^\theta \right].$$

D'après (V.5), nous pouvons prendre h_0 tel que pour tout $h \leq h_0$, nous ayons $C_{p,\theta} |\theta - \theta_h|_{L^2(\Omega)} \leq \frac{1}{2} |\Omega|$.

En utilisant (V.69), nous pouvons conclure que pour tout $h \leq h_0$,

$$\begin{aligned} \sum_{n=0}^{N-1} \Delta t \sum_{i=1}^3 |\mu_{ih}^{n+1}|_{H^1(\Omega)}^2 &\leq 8C_{p,\theta}^2 \left[\frac{2\Sigma_M}{M_1} K_0 + (M_1^\theta)^2 \right] := K_1'', \\ \sum_{n=0}^{N-1} \Delta t \sum_{i=1}^3 |\mu_{ih}^{n+1}|_{H^1(\Omega)}^2 &\leq K_2^\varepsilon, \end{aligned} \tag{V.70}$$

où $K_2^\varepsilon = 2C_p (K_1^\varepsilon + K_2^\varepsilon)$

(iv) De (V.64) et (V.68), nous déduisons

$$\frac{3}{8} (2\beta - 1) C_\varepsilon \int_{\Omega} \sum_{n=0}^{N-1} \sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 dx \leq K_0.$$

En définissant $K'_2 = \frac{8C_p^2}{3(2\beta - 1)\Sigma_\varepsilon} K_0$, en utilisant la proposition IV.5, l'inégalité de Poincaré et la propriété de conservation du volume (V.66), nous obtenons finalement

$$\sum_{n=0}^{N-1} \Delta t \sum_{i=1}^3 \left| \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \right|_{H^1(\Omega)}^2 \leq \frac{K'_2}{\Delta t}.$$

(v) Soit $\nu \in H^1(\Omega)$. Notons ν_h^μ la projection L^2 de ν dans $\mathcal{V}_h^\mu$. D'après (V.6), nous avons $|\nu_h^\mu|_{H^1(\Omega)} \leq C|\nu|_{H^1(\Omega)}$. En utilisant la première équation de (V.8), nous obtenons

$$\int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx = - \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} \nabla \mu_{ih}^{n+1} \cdot \nabla \nu_h^\mu dx.$$

Donc, nous trouvons

$$\left| \left(\frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t}, \nu \right)_{L^2(\Omega)} \right| = \left| \left(\frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t}, \nu_h^\mu \right)_{L^2(\Omega)} \right| \leq \frac{M_2 C}{\Sigma_m} |\nabla \mu_{ih}^{n+1}|_{L^2(\Omega)} |\nu|_{H^1(\Omega)}.$$

Puisque cette inégalité est vraie pour tout $\nu \in H^1(\Omega)$, nous avons

$$\left| \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \right|_{(H^1(\Omega))'} = \sup_{\nu \in H^1(\Omega)} \frac{\left| \left(\frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t}, \nu \right)_{L^2(\Omega)} \right|}{|\nu|_{H^1(\Omega)}} \leq \frac{M_2 C}{\Sigma_m} |\nabla \mu_{ih}^{n+1}|_{L^2(\Omega)},$$

et ainsi,

$$\sum_{n=0}^{N-1} \Delta t \sum_{i=1}^3 \left| \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \right|_{(H^1(\Omega))'}^2 \leq \left(\frac{M_2 C}{\Sigma_m} \right)^2 K_1'' := K_2''.$$

■

Estimation des résidus

Les bornes établies dans la proposition V.25 et des arguments de compacité permettent d'extraire des sous-suites convergentes à partir d'une suite de solutions approchées. Ensuite, il reste à montrer que la limite que nous obtenons est solution faible du modèle de Cahn-Hilliard triphasique (IV.9). Ainsi, la première étape est de spécifier le lien entre les équations satisfaites par les solutions approchées et celles satisfaites par la solution faible de (IV.9).

La proposition suivante donne les estimations des termes résiduels dus à la discrétisation en temps.

Proposition V.26

Soient $\tau \in C_c^\infty([0, t_f])$, $\nu_h^c \in \mathcal{V}_{Dh,0}^c$ et $\nu_h^\mu \in \mathcal{V}_h^\mu$. Les suites $(\mathbf{c}_h^N)_{N \in \mathbb{N}}$ et $(\boldsymbol{\mu}_h^N)_{N \in \mathbb{N}}$ satisfont les équations suivantes,

$$\begin{cases} \int_0^{t_f} \left(\int_{\Omega} \frac{dc_{ih}^N}{dt}(t, x) \nu_h^\mu(x) dx \right) \tau(t) dt = - \int_0^{t_f} \left(\int_{\Omega} \frac{M_{0h}^{N+\alpha}}{\Sigma_i} \nabla \mu_{ih}^N(t, x) \cdot \nabla \nu_h^\mu(x) dx \right) \tau(t) dt \\ \int_0^{t_f} \left(\int_{\Omega} \mu_{ih}^N(t, x) \nu_h^c(x) dx \right) \tau(t) dt = \int_0^{t_f} \left(\int_{\Omega} f_i^F(\mathbf{c}_h^N(t, x)) \nu_h^c(x) dx \right) \tau(t) dt \\ \quad + \int_0^{t_f} \left(\int_{\Omega} \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^N(t, x) \cdot \nabla \nu_h^c(x) dx \right) \tau(t) dt + R_{i1}(\nabla \nu_h^c, \Delta t) + R_{i2}(\nu_h^c, \Delta t), \end{cases} \quad (\text{V.71})$$

où $M_{0h}^{N+\alpha} = M_0((1-\alpha)\underline{\mathbf{c}}_h^N + \alpha\bar{\mathbf{c}}_h^N)$ et les termes résiduels R_{i1} et R_{i2} satisfont les estimations suivantes : il existe deux constantes K_3 et K_4 indépendantes de h et Δt telle que, pour tout $i \in \{1, 2, 3\}$,

$$|R_{i1}(\nu_h^c, \Delta t)| \leq K_3 |\nu_h^c|_{H^1(\Omega)} \sqrt{\Delta t}, \quad (\text{V.72})$$

$$|R_{i2}(\nabla \nu_h^c, \Delta t)| \leq K_4 |\nabla \nu_h^c|_{L^2(\Omega)} \Delta t. \quad (\text{V.73})$$

Démonstration : Nous prolongeons la fonction τ sur $\mathbb{R}$ par 0. La première équation de (V.71) est obtenue à partir de la première équation de (V.8) en utilisant les définitions (V.23), (V.22), (V.24), (V.25) de $\bar{\mathbf{c}}_h^N$, $\underline{\mathbf{c}}_h^N$, $\mathbf{c}_h^N$ et $\boldsymbol{\mu}_h^N$. De plus, en multipliant la seconde équation de (V.8) par la fonction τ et en intégrant sur l'intervalle

$[0, t_f]$, nous obtenons la seconde équation de (V.71) avec :

$$\begin{aligned} R_{i1} &= \sum_{n=0}^{N-1} \int_{t_n}^{t_{n+1}} \left(\int_{\Omega} [D_i^F(\mathbf{c}_h^n(x), \mathbf{c}_h^{n+1}(x)) - D_i^F(\mathbf{c}_h^N(t, x), \mathbf{c}_h^N(t, x))] \nu_h^c(x) dx \right) \tau(t) dt, \\ R_{i2} &= \sum_{n=0}^{N-1} \int_{t_n}^{t_{n+1}} \left(\int_{\Omega} \frac{3}{4} \Sigma_i \varepsilon [(1 - \beta) \nabla c_{ih}^n(x) + \beta \nabla c_{ih}^{n+1}(x) - \nabla c_{ih}^N(t, x)] \cdot \nabla \nu_h^c(x) dx \right) \tau(t) dt. \end{aligned}$$

Notons que nous utilisons ici l'hypothèse de consistance (V.3) qui implique que

$$D_i^F(\mathbf{c}_h^N(t, x), \mathbf{c}_h^N(t, x)) = f_i^F(\mathbf{c}_h^N(t, x)).$$

Il reste à prouver que R_{i1} et R_{i2} satisfont les bornes (V.72) et (V.73).

- (i) La borne pour R_{i1} est basée sur l'hypothèse (V.20). En effet, en appliquant le théorème des accroissements finis et en utilisant les hypothèses (IV.20) et (V.20) de croissance polynomiale de F et d_i^F nous montrons que : pour $(\mathbf{a}, \mathbf{b}) \in \mathcal{S}^2$, pour $\lambda \in [0, 1]$,

$$\begin{aligned} |d_k(\mathbf{a}, \mathbf{b}) - \partial_k F(\lambda \mathbf{a} + (1 - \lambda) \mathbf{b})| &\leq |d_k(\mathbf{a}, \mathbf{b}) - d_k(\mathbf{a}, \mathbf{a})| + |\partial_k F(\mathbf{a}) - \partial_k F(\lambda \mathbf{a} + (1 - \lambda) \mathbf{b})| \\ &\leq \left(\sup_{s \in [0, 1]} |D(d_k^F(\mathbf{a}, \cdot))(s \mathbf{a} + (1 - s) \mathbf{b})| \right) |\mathbf{b} - \mathbf{a}| \\ &\quad + \left(\sup_{s \in [0, \lambda]} |D^2 F(s \mathbf{a} + (1 - s) \mathbf{b})| \right) (1 - \lambda) |\mathbf{b} - \mathbf{a}| \\ &\leq B_1 \left(|\mathbf{a}|^{p-2} + 2 \sup_{s \in [0, 1]} |s \mathbf{a} + (1 - s) \mathbf{b}|^{p-2} + 2 \right) |\mathbf{b} - \mathbf{a}| \\ &\leq B_1 \left(|\mathbf{a}|^{p-2} + 2 (|\mathbf{b}| + |\mathbf{b} - \mathbf{a}|)^{p-2} + 2 \right) |\mathbf{b} - \mathbf{a}|. \end{aligned}$$

Nous en déduisons, grâce à l'inégalité de Young, qu'il existe une constante positive T_1 telle que, pour tout $(\mathbf{a}, \mathbf{b}) \in \mathcal{S}^2$, pour tout $\lambda \in [0, 1]$,

$$|d_k(\mathbf{a}, \mathbf{b}) - \partial_k F(\lambda \mathbf{a} + (1 - \lambda) \mathbf{b})| \leq T_1 |\mathbf{b} - \mathbf{a}| \left(1 + |\mathbf{b}|^{p-2} + |\mathbf{a}|^{p-2} \right). \quad (\text{V.74})$$

Puisque R_{i1} peut être écrit de la manière suivante :

$$\begin{aligned} R_{i1} &= \frac{4\Sigma_T}{\varepsilon} \sum_{j \neq i} \frac{1}{\Sigma_j} \sum_{n=0}^{N-1} \int_{t_n}^{t_{n+1}} \int_{\Omega} \left[(d_i^F(\mathbf{c}_h^n(x), \mathbf{c}_h^{n+1}(x)) - \partial_i F(\mathbf{c}_h^N(t, x))) \right. \\ &\quad \left. - (d_j^F(\mathbf{c}_h^n(x), \mathbf{c}_h^{n+1}(x)) - \partial_j F(\mathbf{c}_h^N(t, x))) \right] \nu_h^c(x) dx \tau(t) dt; \end{aligned}$$

et, d'après (V.74), nous avons

$$|d_k^F(\mathbf{c}_h^n(x), \mathbf{c}_h^{n+1}(x)) - \partial_k F(\mathbf{c}_h^N(t, x))| \leq T_1 |\mathbf{c}_h^{n+1}(x) - \mathbf{c}_h^n(x)| \left(1 + |\mathbf{c}_h^n(x)|^{p-2} + |\mathbf{c}_h^{n+1}(x)|^{p-2} \right).$$

Ainsi, puisque $2 \leq p \leq 6$, nous avons $1 \leq \frac{6}{7-p} \leq 6$, $\frac{6}{p-2} \geq 0$ et $\frac{7-p}{6} + \frac{p-2}{6} + \frac{1}{6} = 1$ et nous pouvons appliquer l'inégalité d'Hölder pour obtenir

$$\begin{aligned} &\left| \int_{\Omega} \left(d_k^F(\mathbf{c}_h^n(x), \mathbf{c}_h^{n+1}(x)) - \partial_k F(\mathbf{c}_h^N(t, x)) \right) \nu_h^c(x) dx \right| \\ &\leq T_1 \int_{\Omega} |\mathbf{c}_h^{n+1}(x) - \mathbf{c}_h^n(x)| \left(1 + |\mathbf{c}_h^n(x)|^{p-2} + |\mathbf{c}_h^{n+1}(x)|^{p-2} \right) |\nu_h^c(x)| dx \\ &\leq T_1 \left| \mathbf{c}_h^{n+1} - \mathbf{c}_h^n \right|_{L^{\frac{6}{7-p}}(\Omega)} \left| 1 + |\mathbf{c}_h^n|^{p-2} + |\mathbf{c}_h^{n+1}|^{p-2} \right|_{L^{\frac{6}{p-2}}(\Omega)} |\nu_h^c|_{L^6(\Omega)} \\ &\leq T_2 \left(1 + 2K_1^{p-2} \right) |\nu_h^c|_{H^1(\Omega)} |\mathbf{c}_h^{n+1} - \mathbf{c}_h^n|_{H^1(\Omega)}. \end{aligned}$$

En appliquant l'inégalité de Cauchy-Schwarz, nous trouvons :

$$\begin{aligned} & \left| \sum_{n=0}^{N-1} \int_{t_n}^{t_{n+1}} \int_{\Omega} (d_k^F(\mathbf{c}_h^n(x), \mathbf{c}_h^{n+1}(x)) - \partial_k F(\mathbf{c}_h^N(t, x))) \nu_h^c(x) dx \tau(t) dt \right| \\ & \leq T_2 \left(1 + 2K_1^{p-2}\right) |\nu_h^c|_{H^1(\Omega)} \left(\sup_{t \in [0, t_f]} |\tau(t)| \right) \sum_{n=0}^{N-1} \Delta t |\mathbf{c}_h^{n+1} - \mathbf{c}_h^n|_{H^1(\Omega)} \\ & \leq T_2 \left(1 + 2K_1^{p-2}\right) |\nu_h^c|_{H^1(\Omega)} \left(\sup_{t \in [0, t_f]} |\tau(t)| \right) \Delta t \left(\sum_{n=0}^{N-1} \Delta t \left| \frac{\mathbf{c}_h^{n+1} - \mathbf{c}_h^n}{\Delta t} \right|_{H^1(\Omega)}^2 \right)^{\frac{1}{2}}. \end{aligned}$$

En conclusion, en utilisant la troisième borne du théorème V.25, nous obtenons

$$|R_{i1}| \leq T_2 K_2 \left(1 + 2K_1^{p-2}\right) |\tau|_{L^\infty([0, t_f])} |\nu_h^c|_{H^1(\Omega)} \Delta t^{\frac{1}{2}}.$$

Donc, l'estimation (V.72) est satisfaite avec $K_3 := T_2 K_2 \left(1 + 2K_1^{p-2}\right) |\tau|_{L^\infty([0, t_f])}$.

(ii) Un changement de variable et une renumérotation des termes permettent d'obtenir :

$$\begin{aligned} R_{i2} &= \sum_{n=0}^{N-1} \int_{t_n}^{t_{n+1}} \left(\int_{\Omega} \frac{3}{4} \Sigma_i \varepsilon \left[\left(\beta - \frac{t - t_n}{\Delta t} \right) (\nabla c_{ih}^{n+1}(x) - \nabla c_{ih}^n(x)) \right] \cdot \nabla \nu_h^c(x) dx \right) \tau(t) dt \\ &= \frac{3}{4} \Sigma_i \varepsilon \sum_{n=0}^{N-1} \int_0^1 \left(\int_{\Omega} \Delta t \left[(\beta - u) (\nabla c_{ih}^{n+1}(x) - \nabla c_{ih}^n(x)) \right] \cdot \nabla \nu_h^c(x) dx \right) \tau((n+u)\Delta t) du \\ &= \frac{3}{4} \Sigma_i \varepsilon \sum_{n=0}^N \Delta t \left(\int_{\Omega} \nabla c_{ih}^n(x) \cdot \nabla \nu_h^c(x) dx \right) \left(\int_0^1 (\beta - u) \underbrace{(\tau((n-1+u)\Delta t) - \tau((n+u)\Delta t))}_{\leq \Delta t |\tau'|_{L^\infty(\mathbb{R})}} du \right), \end{aligned}$$

et grâce au théorème V.25, nous obtenons

$$|R_{i2}| \leq \frac{3}{4} \Sigma_M \varepsilon (N+1) \Delta t K_1 |\nabla \nu_h^c(x)|_{L^2(\Omega)} \Delta t |\tau'|_{L^\infty(\mathbb{R})} \leq \frac{3}{4} \Sigma_M \varepsilon 2t_f |\nabla \nu_h^c(x)|_{L^2(\Omega)} \Delta t |\tau'|_{L^\infty(\mathbb{R})}.$$

Donc, l'estimation (V.73) est satisfaite avec $K_4 = \frac{3}{2} K_1 t_f \Sigma_M \varepsilon |\tau'|_{L^\infty(\mathbb{R})}$. ■

Pour pouvoir montrer la convergence lorsque les pas de temps et pas d'espace tendent vers zéro, nous devons également estimer les termes résiduels dus à la discrétisation en espace.

Proposition V.27

Soient $\tau \in \mathcal{C}_c^\infty([0, t_f])$, $\nu^c \in \mathcal{V}_{D,0}^c$ et $\nu^\mu \in \mathcal{V}^\mu$. Les suites $(\mathbf{c}_h^N)_{N \in \mathbb{N}}$ et $(\mu_h^N)_{N \in \mathbb{N}}$ satisfont les équations suivantes,

$$\begin{aligned} \int_0^{t_f} \left(\int_{\Omega} \frac{dc_{ih}^N}{dt}(t, x) \nu^\mu(x) dx \right) \tau(t) dt &= - \int_0^{t_f} \left(\int_{\Omega} \frac{M_{0h}^{N+\alpha}}{\Sigma_i} \nabla \mu_{ih}^N(t, x) \cdot \nabla \nu^\mu(x) dx \right) \tau(t) dt + R_{i3}(h, \Delta t), \\ \int_0^{t_f} \left(\int_{\Omega} \mu_{ih}^N(t, x) \nu^c(x) dx \right) \tau(t) dt &= \int_0^{t_f} \left(\int_{\Omega} f_i^F(\mathbf{c}_h^N(t, x)) \nu^c(x) dx \right) \tau(t) dt \\ &\quad + \int_0^{t_f} \left(\int_{\Omega} \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^N(t, x) \cdot \nabla \nu^c(x) dx \right) \tau(t) dt + R_{i1}(h, \Delta t) + R_{i2}(h, \Delta t) + R_{i4}(h, \Delta t), \end{aligned}$$

où R_{i1} , R_{i2} , R_{i3} et R_{i4} satisfont les estimations suivantes : il existe 4 constantes K_5 , K_6 , K_7 et K_8

indépendantes de h et Δt telle que,

$$\begin{aligned} |R_{i1}(h, \Delta t)| &\leq K_5 \Delta t, \\ |R_{i2}(h, \Delta t)| &\leq K_6 \sqrt{\Delta t}, \\ |R_{i3}(h, \Delta t)| &\leq K_7 \inf_{\nu_h^\mu \in \mathcal{V}_h^\mu} |\nu^\mu - \nu_h^\mu|_{H^1(\Omega)}, \\ |R_{i4}(h, \Delta t)| &\leq K_8 \inf_{\nu_h^c \in \mathcal{V}_{Dh,0}^c} |\nu^c - \nu_h^c|_{H^1(\Omega)}. \end{aligned}$$

Démonstration : Soit ν_h^c , (resp. ν_h^μ), la projection H^1 de ν^c , (resp. ν^μ), sur $\mathcal{V}_{hD,0}^c$, (resp. $\mathcal{V}_h^\mu$). En utilisant le théorème V.26 et en notant $R_{i1}(h, \Delta t)$ et $R_{i2}(h, \Delta t)$ les termes $R_{i1}(\nu_h^c, \Delta t)$ et $R_{i2}(\nabla \nu_h^c, \Delta t)$, nous trouvons que les termes résiduels R_{i3} et R_{i4} sont donnés par

$$\begin{aligned} R_{i3}(h, \Delta t) &= \int_0^{t_f} \left(\int_\Omega \frac{dc_{ih}^N}{dt}(t, x) (\nu^\mu(x) - \nu_h^\mu(x)) dx \right) \tau(t) dt \\ &\quad + \int_0^{t_f} \left(\int_\Omega \frac{M_{0h}^{N+\alpha}}{\Sigma_i} \nabla \mu_{ih}^N(t, x) \cdot \nabla (\nu^\mu(x) - \nu_h^\mu(x)) dx \right) \tau(t) dt, \end{aligned}$$

et

$$\begin{aligned} R_{i4}(h, \Delta t) &= \int_0^{t_f} \left(\int_\Omega \mu_{ih}^N(t, x) (\nu^c(x) - \nu_h^c(x)) dx \right) \tau(t) dt - \int_0^{t_f} \left(\int_\Omega f_i^F(\mathbf{c}_h^N(t, x)) (\nu^c(x) - \nu_h^c(x)) dx \right) \tau(t) dt \\ &\quad - \int_0^{t_f} \left(\int_\Omega \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^N(t, x) \cdot \nabla (\nu^c(x) - \nu_h^c(x)) dx \right) \tau(t) dt. \end{aligned}$$

Les bornes pour R_{i1} et R_{i2} proviennent de l'inégalité $|\nu_h^c|_{H^1(\Omega)} \leq |\nu^c|_{H^1(\Omega)}$. La borne pour R_{i3} est obtenue de la manière suivante :

$$\begin{aligned} |R_{i3}| &\leq \left| \frac{dc_{ih}^N}{dt} \right|_{L^2(0, t_f, (H^1(\Omega))')} |\tau|_{L^2(0, t_f)} |\nu^\mu - \nu_h^\mu|_{H^1(\Omega)} + \frac{M_2}{\Sigma_m} |\mu_{ih}^N|_{L^2(0, t_f, H^1(\Omega))} |\tau|_{L^2(0, t_f)} |\nu^\mu - \nu_h^\mu|_{H^1(\Omega)} \\ &\leq K_7 |\nu^\mu - \nu_h^\mu|_{H^1(\Omega)}, \end{aligned}$$

avec $K_7 := \left(\frac{M_2}{\Sigma_m} \sqrt{K_1} + K_2 \right) |\tau|_{L^2(0, t_f)}$, et la borne R_{i4} est déduite de l'inégalité suivante :

$$\begin{aligned} |R_{i4}| &\leq |\mu_{ih}^N|_{L^2(0, t_f, L^2(\Omega))} |\tau|_{L^2(0, t_f)} |\nu^c - \nu_h^c|_{L^2(\Omega)} \\ &\quad + \frac{24 \Sigma_T}{\varepsilon \Sigma_m} \int_0^{t_f} B_1 \left(|c_{ih}^N(t, \cdot)|_{L^6(\Omega)}^{p-1} |\nu^c - \nu_h^c|_{L^{\frac{6}{7-p}}(\Omega)} + |\Omega|^{\frac{1}{2}} |\nu^c - \nu_h^c|_{L^2(\Omega)} \right) \tau(t) dt \\ &\quad + \frac{3}{4} \Sigma_M \varepsilon t_f |c_{ih}^N|_{L^\infty(0, t_f, H^1(\Omega))} |\tau|_{L^\infty(0, t_f)} |\nu^c - \nu_h^c|_{H^1(\Omega)} \\ &\leq \underbrace{\left[\sqrt{K_1} |\tau|_{L^2(0, t_f)} + \frac{24 \Sigma_T}{\varepsilon \Sigma_m} t_f |\tau|_{L^\infty(0, t_f)} B_1 \left(K_1^{p-1} + |\Omega|^{\frac{1}{2}} \right) + \frac{3}{4} \Sigma_M \varepsilon K_1 |\tau|_{L^\infty(0, t_f)} t_f \right]}_{:=K_8} |\nu^c - \nu_h^c|_{H^1(\Omega)}. \end{aligned}$$

■

Démonstration du théorème V.10

Le théorème V.25 donne les bornes suivantes :

$$|\mathbf{c}_{hK}^N|_{L^\infty(0, t_f, (H^1(\Omega))^3)} + |\boldsymbol{\mu}_{hK}^N|_{L^2(0, t_f, (H^1(\Omega))^3)}^2 + \left| \frac{\partial \mathbf{c}_{hK}^N}{\partial t} \right|_{L^2(0, t_f, (H^1(\Omega))')}^2 \leq K_1 + K_2, \quad (\text{V.76a})$$

$$|\mathbf{c}_{hK}^N - \bar{\mathbf{c}}_{hK}^N|_{L^2(0, t_f, (H^1(\Omega))^3)} + |\mathbf{c}_{hK}^N - \bar{\mathbf{c}}_{hK}^N|_{L^2(0, t_f, (H^1(\Omega))^3)} \leq 2\sqrt{K_2 \Delta t}. \quad (\text{V.76b})$$

En utilisant les estimations (V.76a), nous pouvons extraire des sous-suites de $(\mathbf{c}_{h_K}^N)_{(N,K)}$ et $(\boldsymbol{\mu}_{h_K}^N)_{(N,K)}$ (nous noterons encore ces sous-suites $(\mathbf{c}_{h_K}^N)_{(N,K)}$ et $(\boldsymbol{\mu}_{h_K}^N)_{(N,K)}$) telles que

$$\mathbf{c}_{h_K}^N \rightharpoonup \mathbf{c} \quad \text{dans } L^\infty(0, t_f, (H^1(\Omega))^3) \text{ faible-}, \quad (\text{V.77})$$

$$\boldsymbol{\mu}_{h_K}^N \rightharpoonup \boldsymbol{\mu} \quad \text{dans } L^2(0, t_f, (H^1(\Omega))^3) \text{ faible}, \quad (\text{V.78})$$

$$\frac{\partial \mathbf{c}_{h_K}^N}{\partial t} \rightharpoonup \frac{\partial \mathbf{c}}{\partial t} \quad \text{dans } L^2(0, t_f, (H^1(\Omega))') \text{ faible}. \quad (\text{V.79})$$

A partir de l'estimation (V.76a), nous pouvons utiliser le théorème de compacité d'Aubin–Lions–Simon [Sim87] pour obtenir, à sous-suites près,

$$\mathbf{c}_{h_K}^N \rightarrow \mathbf{c} \quad \text{dans } \mathcal{C}^0(0, t_f, (L^q(\Omega))^3) \text{ fort, pour tout } 1 \leq q < +\infty \text{ si } d = 2, \text{ ou } 1 \leq q < 6 \text{ si } d = 3. \quad (\text{V.80})$$

En particulier, (V.80) implique que

$$\mathbf{c}_{h_K}^N \rightarrow \mathbf{c} \quad \text{dans } L^2(0, t_f, (L^2(\Omega))^3) \text{ fort}, \quad (\text{V.81})$$

et l'estimation (V.76b) mène à

$$\underline{\mathbf{c}}_{h_K}^N \rightarrow \mathbf{c} \quad \text{dans } L^2(0, t_f, (L^2(\Omega))^3) \text{ fort}, \quad (\text{V.82})$$

$$\overline{\mathbf{c}}_{h_K}^N \rightarrow \mathbf{c} \quad \text{dans } L^2(0, t_f, (L^2(\Omega))^3) \text{ fort}. \quad (\text{V.83})$$

Soit $\tau \in \mathcal{C}_c^\infty([0, t_f])$, $\nu^c \in \mathcal{V}_{D,0}^c$ et $\nu^\mu \in \mathcal{V}^\mu$. Nous pouvons appliquer le théorème V.27 et passer à la limite dans (V.75) :

- (i) Les convergences (V.79), (V.78) et (V.77) autorisent à passer à la limite dans les termes linéaires.
- (ii) Les termes R_{i1} , R_{i2} , R_{i3} et R_{i4} tendent vers 0 grâce aux hypothèses (V.5).
- (iii) Soit $\eta > 0$. Puisque l'espace $\mathcal{C}^\infty(\overline{\Omega})$ est dense dans $\mathcal{V}^\mu = H^1(\Omega)$, nous pouvons prendre $\nu_\eta^\mu \in \mathcal{C}^\infty(\overline{\Omega})$ (dépendant de η) telle que

$$|\nu^\mu - \nu_\eta^\mu|_{H^1(\Omega)} < \frac{\Sigma_m}{M_2 K_1} (|\tau|_{L^2(0, t_f)})^{-1} \frac{\eta}{4},$$

(où M_2 et K_1 sont les constantes intervenant dans (IV.19) et dans le théorème V.25). Nous obtenons

$$\left| \int_0^{t_f} \left(\int_\Omega \frac{M_{0h_K}^{N+\alpha}}{\Sigma_i} \nabla \mu_{ih_K}^N(t, x) \cdot \nabla (\nu^\mu - \nu_\eta^\mu)(x) dx \right) \tau(t) dt \right| \leq \frac{\eta}{4}, \quad (\text{V.84})$$

et de manière similaire

$$\left| \int_0^{t_f} \left(\int_\Omega \frac{M_0(\mathbf{c})}{\Sigma_i} \nabla \mu_{ih_K}^N(t, x) \cdot \nabla (\nu^\mu - \nu_\eta^\mu)(x) dx \right) \tau(t) dt \right| \leq \frac{\eta}{4}. \quad (\text{V.85})$$

De plus, en utilisant les hypothèses (IV.19), nous avons

$$\begin{aligned} & \left| \int_0^{t_f} \left(\int_\Omega \frac{M_{0h_K}^{N+\alpha} - M_0(\mathbf{c})}{\Sigma_i} \nabla \mu_{ih_K}^N(t, x) \cdot \nabla \nu_\eta^\mu(x) dx \right) \tau(t) dt \right| \\ & \leq \frac{|\nabla \nu_\eta^\mu|_{L^\infty(\Omega)^3}}{\Sigma_m} \int_0^{t_f} \left(\int_\Omega M_3 |1 - \alpha| (\underline{\mathbf{c}}_{h_K}^N - \mathbf{c}) + \alpha (\overline{\mathbf{c}}_{h_K}^N - \mathbf{c}) \|\nabla \mu_{ih_K}^N(t, x)\| dx \right) |\tau(t)| dt \\ & \leq \frac{|\nabla \nu_\eta^\mu|_{L^\infty(\Omega)^3}}{\Sigma_m} M_3 K_1 |\tau|_{L^\infty(0, t_f)} \left[\|\underline{\mathbf{c}}_{h_K}^N - \mathbf{c}\|_{L^2(0, t_f, L^2(\Omega)^3)} + \|\overline{\mathbf{c}}_{h_K}^N - \mathbf{c}\|_{L^2(0, t_f, L^2(\Omega)^3)} \right]. \end{aligned}$$

D'après les convergences (V.82) et (V.83), il existe $P_1 \in \mathbb{N}$ (dépendant de η) tel que : $\forall (N, K) \in \mathbb{N}^2$ tel que $\min(N, K) \geq P_1$ nous avons

$$\left| \int_0^{t_f} \left(\int_\Omega \frac{M_{0h_K}^{N+\alpha} - M_0(\mathbf{c})}{\Sigma_i} \nabla \mu_{ih_K}^N(t, x) \cdot \nabla \nu_\eta^\mu(x) dx \right) \tau(t) dt \right| \leq \frac{\eta}{4}. \quad (\text{V.86})$$

De plus, l'hypothèse (IV.19) implique que $M_0(\mathbf{c}) \in L^\infty([0, t_f], L^\infty(\Omega))$, et ainsi $M_0(\mathbf{c}) \nabla \nu_\eta^\mu \tau$ appartient à $L^2(0, t_f, (L^2(\Omega))^d)$. La convergence (V.28) implique que

$$\int_0^{t_f} \left(\int_\Omega \frac{M_0(\mathbf{c})}{\Sigma_i} \nabla \mu_{ih_K}^N(t, x) \cdot \nabla \nu_\eta^\mu(x) dx \right) \tau(t) dt \xrightarrow{\min(N, K) \rightarrow \infty} \int_0^{t_f} \left(\int_\Omega \frac{M_0(\mathbf{c})}{\Sigma_i} \nabla \mu_i(t, x) \cdot \nabla \nu_\eta^\mu(x) dx \right) \tau(t) dt.$$

Donc, il existe $P_2 \in \mathbb{N}$ tel que : $\forall (N, K) \in \mathbb{N}^2$ tels que $\min(N, K) \geq P_2$ nous avons

$$\left| \int_0^{t_f} \left(\int_{\Omega} \frac{M_0(\mathbf{c})}{\Sigma_i} \nabla(\mu_{ih_K}^N - \mu_i)(t, x) \cdot \nabla \nu_{\eta}^{\mu}(x) dx \right) \tau(t) dt \right| \leq \frac{\eta}{4}. \quad (\text{V.87})$$

Finalement, en utilisant (V.84), (V.85), (V.86), (V.87) et l'inégalité triangulaire, nous obtenons : $\forall (N, K) \in \mathbb{N}^2$ tels que $\min(N, K) \geq \max(P_1, P_2)$,

$$\left| \int_0^{t_f} \left(\int_{\Omega} \frac{M_0(\mathbf{c})}{\Sigma_i} \nabla \mu_i(t, x) \cdot \nabla \nu^{\mu}(x) dx \right) \tau(t) dt - \int_0^{t_f} \left(\int_{\Omega} \frac{M_{0h_K}^{N+\alpha}}{\Sigma_i} \nabla \mu_{ih_K}^N(t, x) \cdot \nabla \nu^{\mu}(x) dx \right) \tau(t) dt \right| \leq \eta.$$

Nous concluons que

$$\int_0^{t_f} \left(\int_{\Omega} \frac{M_{0h_K}^{N+\alpha}}{\Sigma_i} \nabla \mu_{ih_K}^N(t, x) \cdot \nabla \nu_{\eta}^{\mu}(x) dx \right) \tau(t) dt \xrightarrow{\min(N, K) \rightarrow \infty} \int_0^{t_f} \left(\int_{\Omega} \frac{M_0(\mathbf{c})}{\Sigma_i} \nabla \mu_i(t, x) \cdot \nabla \nu_{\eta}^{\mu}(x) dx \right) \tau(t) dt.$$

(iv) En utilisant la réciproque du théorème de Lebesgue et la convergence (V.81), il existe une sous-suite de $(\mathbf{c}_{h_K}^N)_{(N, K)}$ (encore notée $(\mathbf{c}_{h_K}^N)_{(N, K)}$) et une fonction $\mathbf{S} \in L^q(0, t_f, L^q(\Omega))$ telles que :

$$\mathbf{c}_{h_K}^N \rightarrow \mathbf{c} \text{ presque partout,} \quad (\text{V.88})$$

et

$$|\mathbf{c}_{h_K}^N(t, x)| \leq \mathbf{S}(t, x) \text{ pour presque tout } (t, x) \in]0, t_f[\times \Omega, \text{ pour tout } (N, K) \in \mathbb{N}^2.$$

Grâce à l'hypothèse (IV.20) de croissance polynomiale des fonctions $\partial_i F$, nous avons,

$$|f_i^F(\mathbf{c}_{h_K}^N) \nu^c(x) \tau(t)| \leq \frac{16 \Sigma_T}{\Sigma_m \varepsilon} B_1 \left[|\mathbf{S}(x)|^{p-1} + 1 \right] \nu^c(x) \tau(t),$$

pour presque tout $(t, x) \in]0, t_f[\times \Omega$, pour tout $(N, K) \in \mathbb{N}^2$. Le membre de droite appartient à $L^1(0, t_f, L^1(\Omega))$. Donc, grâce à la convergence (V.88) et au théorème de Lebesgue, nous avons

$$\int_0^{t_f} \left(\int_{\Omega} f_i^F(\mathbf{c}_{h_K}^N(t, x)) \nu^c(x) dx \right) \tau(t) dt \rightarrow \int_0^{t_f} \left(\int_{\Omega} f_i^F(\mathbf{c}(t, x)) \nu^c(x) dx \right) \tau(t) dt$$

Cela montre l'existence d'une solution faible $(\mathbf{c}, \boldsymbol{\mu})$ au problème (IV.9) ainsi que les convergences (V.27) et (V.28). ■

V.6 Conclusion

Nous avons proposé dans ce chapitre une discrétisation en temps spécifique des termes non linéaires du système de Cahn-Hilliard. Cette discrétisation permet d'obtenir un schéma stable pour tout pas de temps. Il faut également noter que nous avons démontré que, dans les situations d'étalement partiel, le schéma implicite est stable lorsque le pas de temps est suffisamment petit. Cette condition sur le pas de temps n'est pas une condition de type CFL puisqu'elle n'est pas liée à la valeur du pas d'espace mais seulement aux paramètres du modèle : l'épaisseur d'interface et la mobilité. Dans les cas d'étalement total, lors de l'utilisation du schéma implicite nous avons rencontré dans la plupart des cas des difficultés de convergence dans la méthode de Newton. En fait, nous ne savons pas garantir que le schéma correspondant à une solution. L'utilisation de schémas semi-implicites en temps (CC. ou SImpl.) permet de résoudre ces problèmes, néanmoins le schéma CC. semble inutilisable dans la pratique au vu de l'erreur de troncature introduite. Le schéma SImpl apparaît comme un bon compromis en stabilité et précision. Ceci sera encore illustré par les simulations d'écoulements (par le couplage aux équations de Navier-Stokes) présentées dans la partie 3.

Chapitre VI

Discrétisation inconditionnellement stable du système Cahn-Hilliard/Navier-Stokes (CH/NS)

Dans ce chapitre nous proposons un schéma original pour la discrétisation du système complet Cahn-Hilliard/Navier-Stokes (IV.28) avec les conditions aux bords (IV.30)-(IV.31) et la condition initiale (IV.32). Nous montrons que ce schéma est inconditionnellement stable et qu'il préserve les propriétés essentielles du système de Cahn-Hilliard triphasique, à savoir la conservation du volume de chaque phase et le fait que la somme des trois paramètres d'ordre reste égale à 1 au cours du temps. En outre, le schéma présenté dans ce chapitre permet une résolution découplée des systèmes (discrets) de Cahn-Hilliard et de Navier-Stokes. Nous proposons une étude de convergence dans le cas homogène, *i.e.* lorsque les trois fluides en présence ont la même densité. Ceci permet en particulier d'un point de vue théorique de montrer l'existence d'une solution au problème (IV.28) (dans le cas homogène).

Le schéma numérique est présenté dans la section VI.1. Nous discutons également dans cette section des premières propriétés fondamentales satisfaites par le schéma : la conservation du volume de chaque phase et le fait que la somme des trois paramètres d'ordre reste égale à 1 au cours du temps, ce qui permet de ne résoudre que les équations portant sur deux des paramètres d'ordre, le dernier étant déduit *a posteriori*.

Nous consacrons ensuite la section VI.2 à l'étude du schéma lorsque l'une des phases n'est pas présente. Ceci permet d'obtenir une discrétisation inconditionnellement stable du système de Cahn-Hilliard diphasique pour lequel des schémas ont été proposés dans la littérature [Fen06, KSW08] (dans le cas de deux fluides homogènes). Dans [Fen06], les auteurs proposent l'analyse d'un schéma totalement implicite en temps, la résolution des équations de Cahn-Hilliard et Navier-Stokes étant alors couplée. Dans [KSW08], les auteurs ont proposé l'analyse du schéma plus classique où un traitement explicite des termes de transport dans Cahn-Hilliard permet le découplage des deux systèmes dans un pas de temps. La stabilité n'est alors obtenue que conditionnellement en supposant (essentiellement) que $\Delta t \leq Ch$ où C est une constante.

Dans la section VI.3, nous établissons l'estimation d'énergie. Celle-ci est la base du théorème d'existence des solutions approchées qui est énoncé et prouvé dans la section VI.4 ainsi que du théorème de convergence établi dans la section VI.5 pour le cas de trois fluides homogènes.

VI.1 Discrétisation du modèle CH/NS triphasique

VI.1.1 Discrétisation en temps

Soit $N \in \mathbb{N}^*$ et $t_f \in]0, +\infty[$. L'intervalle de temps $[0, t_f]$ est uniformément discrétisé avec un pas de temps fixe $\Delta t = \frac{t_f}{N}$ et nous définissons $t_n = n\Delta t$, pour tout $n \in \llbracket 0, N \rrbracket$.

Nous supposons que les fonctions $\mathbf{c}^n \in \mathcal{V}_S^c$ et $\mathbf{u}^n \in \mathcal{V}_0^u$ ($n \in \llbracket 0, N \rrbracket$) sont données et décrivons le système à résoudre pour pouvoir calculer les inconnues $\mathbf{c}^{n+1} \in \mathcal{V}_S^c$ et $\mathbf{u}^{n+1} \in \mathcal{V}_0^u$ à l'instant t^{n+1} .

Notre présentation se découpe de la manière suivante :

- Nous décrivons tout d'abord, dans deux paragraphes séparés, les schémas utilisés pour les systèmes de Cahn-Hilliard et de Navier-Stokes sans tenir compte des termes de couplage. Rappelons que l'étude des discrétisations du système de Cahn-Hilliard a été présentée dans le chapitre V, nous y renvoyons pour plus de détails et celle du système de Navier-Stokes a fait l'objet de nombreux développements dans la littérature (nous nous référons en particulier aux articles [GQ00] et [LW07] qui traitent des cas où la densité est variable).
- Nous expliquons ensuite, dans les deux paragraphes suivants, la démarche qui a permis d'aboutir à la discrétisation des termes de couplage avant de décrire le schéma complet dans le dernier paragraphe de cette section.

Système de Cahn-Hilliard

Nous considérons une discrétisation du système de Cahn-Hilliard de la forme : pour $i = 1, 2, 3$,

$$\begin{cases} \frac{c_i^{n+1} - c_i^n}{\Delta t} + \text{terme de transport} = \operatorname{div} \left(\frac{M_0^{n+\alpha}}{\Sigma_i} \nabla \mu_i^{n+1} \right), \\ \mu_i^{n+1} = D_i^F(\mathbf{c}^n, \mathbf{c}^{n+1}) - \frac{3}{4} \varepsilon \Sigma_i \Delta c_i^{n+\beta}. \end{cases}$$

Comme nous l'avons déjà expliqué, la discrétisation des termes de transport sera décrite ultérieurement.

Le type de discrétisation ci-dessus et différents choix du terme $D_i^F(\mathbf{c}^n, \mathbf{c}^{n+1})$ ont été présentés dans le chapitre V. Nous ne discuterons pas ce point dans cette partie où nous supposons que $D_i^F(\mathbf{c}^n, \mathbf{c}^{n+1})$ est donnée. Nous rappelons simplement que lorsque le terme de transport n'est pas présent, l'estimation d'énergie pour ce système discret est obtenue en multipliant la première équation par μ_i^{n+1} , la seconde par $\frac{c_i^{n+1} - c_i^n}{\Delta t}$ puis en écrivant l'égalité des membres de gauche et en sommant pour $i = 1, 2, 3$.

Système de Navier-Stokes

Nous présentons maintenant la discrétisation en temps de l'équation de bilan de quantité de mouvement (rappelée ci-dessous) du système de Navier-Stokes :

$$\underbrace{\sqrt{\varrho(\mathbf{c})} \frac{\partial}{\partial t} (\sqrt{\varrho(\mathbf{c})} \mathbf{u})}_{(1)} + \underbrace{(\varrho(\mathbf{c}) \mathbf{u} \cdot \nabla) \mathbf{u} + \frac{\mathbf{u}}{2} \operatorname{div} (\varrho(\mathbf{c}) \mathbf{u})}_{(2)} - \underbrace{\operatorname{div} (2\eta(\mathbf{c}) D(\mathbf{u}))}_{(3)} + \nabla p = \sum_{i=1}^3 \mu_i \nabla c_i + \varrho(\mathbf{c}) \mathbf{g}.$$

Nous distinguons dans notre présentation la discrétisation des différents termes (1), (2) et (3) intervenant dans l'équation ci-dessus ; pour chacun d'entre eux nous précisons leur contribution au bilan d'énergie obtenue au niveau discret en multipliant l'équation par $\mathbf{u}^{n+1}$.

Terme (1) : Tirant partie de l'égalité formelle

$$\sqrt{\varrho} \frac{\partial}{\partial t} (\sqrt{\varrho} \mathbf{u}) = \varrho \frac{\partial \mathbf{u}}{\partial t} + \frac{1}{2} \frac{\partial \varrho}{\partial t} \mathbf{u}, \quad (\text{VI.1})$$

nous donnons deux discrétisations possibles du terme (1) :

- la première (extraite de [GQ00]) exploite l'écriture donnée par le membre de gauche de VI.1 :

$$\frac{\sqrt{\varrho^{n+1}} \sqrt{\varrho^{n+1}} \mathbf{u}^{n+1} - \sqrt{\varrho^n} \mathbf{u}^n}{\Delta t} = \frac{\varrho^{n+1} \mathbf{u}^{n+1} - \sqrt{\varrho^n \varrho^{n+1}} \mathbf{u}^n}{\Delta t}.$$

Sa contribution à l'estimation d'énergie est donnée par :

$$\begin{aligned} \int_{\Omega} \sqrt{\varrho^{n+1}} \frac{\sqrt{\varrho^{n+1}} \mathbf{u}^{n+1} - \sqrt{\varrho^n} \mathbf{u}^n}{\Delta t} \cdot \mathbf{u}^{n+1} dx \\ = \frac{1}{2\Delta t} \left[\left| \sqrt{\varrho^{n+1}} \mathbf{u}^{n+1} \right|_{L^2(\Omega)}^2 - \left| \sqrt{\varrho^n} \mathbf{u}^n \right|_{L^2(\Omega)}^2 + \left| \sqrt{\varrho^{n+1}} \mathbf{u}^{n+1} - \sqrt{\varrho^n} \mathbf{u}^n \right|_{L^2(\Omega)}^2 \right]. \end{aligned}$$

– la seconde (extraite de [LW07]) exploite l'écriture donnée par le membre de droite de VI.1 :

$$\varrho^n \frac{\mathbf{u}^{n+1} - \mathbf{u}^n}{\Delta t} + \frac{1}{2} \frac{\varrho^{n+1} - \varrho^n}{\Delta t} \mathbf{u}^{n+1} = \frac{\frac{\varrho^n + \varrho^{n+1}}{2} \mathbf{u}^{n+1} - \varrho^n \mathbf{u}^n}{\Delta t}.$$

Sa contribution à l'estimation d'énergie est donnée par :

$$\begin{aligned} & \int_{\Omega} \varrho^n \frac{\mathbf{u}^{n+1} - \mathbf{u}^n}{\Delta t} \cdot \mathbf{u}^{n+1} dx + \int_{\Omega} \frac{1}{2} \frac{\varrho^{n+1} - \varrho^n}{\Delta t} \mathbf{u}^{n+1} \cdot \mathbf{u}^{n+1} dx \\ &= \frac{1}{2\Delta t} \left[\left| \sqrt{\varrho^{n+1}} \mathbf{u}^{n+1} \right|_{L^2(\Omega)}^2 - \left| \sqrt{\varrho^n} \mathbf{u}^n \right|_{L^2(\Omega)}^2 + \left| \sqrt{\varrho^n} (\mathbf{u}^{n+1} - \mathbf{u}^n) \right|_{L^2(\Omega)}^2 \right]. \end{aligned}$$

Remarque VI.1

Ces deux discrétisations sont de la forme $\frac{\tilde{\varrho}^{n+1} \mathbf{u}^{n+1} - \tilde{\varrho}^n \mathbf{u}^n}{\Delta t}$ où $\tilde{\varrho}^\ell$ ($\ell = n$ ou $n+1$) désigne soit ϱ^ℓ soit une certaine moyenne de ϱ^n et ϱ^{n+1} . En effet, la première forme correspond à $\tilde{\varrho}^{n+1} = \varrho^{n+1}$ et $\tilde{\varrho}^n = \sqrt{\varrho^n \varrho^{n+1}}$ (moyenne géométrique de ϱ^n et ϱ^{n+1}) alors que la seconde forme correspond $\tilde{\varrho}^{n+1} = \frac{\varrho^n + \varrho^{n+1}}{2}$ (moyenne arithmétique de ϱ^n et ϱ^{n+1}) et $\tilde{\varrho}^n = \varrho^n$. La discrétisation des termes de couplage que nous présentons dans la suite de ce chapitre fait intervenir le coefficient $\tilde{\varrho}^n$ dans le système de Cahn-Hilliard. Nous utiliserons donc la seconde forme pour éviter l'ajout d'une non linéarité à travers le coefficient ϱ^{n+1} (qui, rappelons-le, vaut $\varrho(\mathbf{c}^{n+1})$).

Terme (2) : Le terme (2) est linéarisé en explicitant la vitesse d'advection :

$$(\varrho^{n+1} \mathbf{u}^n \cdot \nabla) \mathbf{u}^{n+1} + \frac{\mathbf{u}^{n+1}}{2} \operatorname{div} (\varrho^{n+1} \mathbf{u}^n).$$

Sa contribution à l'estimation d'énergie est nulle. En effet, pour tout $\boldsymbol{\nu}^u \in \mathcal{V}_0^u$, nous avons

$$\begin{aligned} & \int_{\Omega} (\varrho^{n+1} \mathbf{u}^n \cdot \nabla) \mathbf{u}^{n+1} \cdot \boldsymbol{\nu}^u dx + \int_{\Omega} \frac{1}{2} \operatorname{div} (\varrho^{n+1} \mathbf{u}^n) \mathbf{u}^{n+1} \cdot \boldsymbol{\nu}^u dx \\ &= \frac{1}{2} \left[\int_{\Omega} (\varrho^{n+1} \mathbf{u}^n \cdot \nabla) \mathbf{u}^{n+1} \cdot \boldsymbol{\nu}^u dx - \int_{\Omega} (\varrho^{n+1} \mathbf{u}^n \cdot \nabla) \boldsymbol{\nu}^u \cdot \mathbf{u}^{n+1} dx \right]. \end{aligned}$$

En particulier, lorsque l'on prend $\boldsymbol{\nu}^u = \mathbf{u}^{n+1}$, le terme ci-dessus est nul.

Terme (3) : Le terme (3) est discrétisé de manière implicite

$$-\operatorname{div} (2\eta^{n+1} D\mathbf{u}^{n+1}).$$

Sa contribution à l'estimation d'énergie est la suivante :

$$\int_{\Omega} \eta^{n+1} |D\mathbf{u}^{n+1}|^2 dx.$$

Ainsi, nous considérons la discrétisation suivante des équations de Navier-Stokes :

$$\begin{cases} \varrho^n \frac{\mathbf{u}^{n+1} - \mathbf{u}^n}{\Delta t} + \frac{1}{2} \frac{\varrho^{n+1} - \varrho^n}{\Delta t} \mathbf{u}^{n+1} + (\varrho^{n+1} \mathbf{u}^n \cdot \nabla) \mathbf{u}^{n+1} + \frac{\mathbf{u}^{n+1}}{2} \operatorname{div} (\varrho^{n+1} \mathbf{u}^n) \\ \quad - \operatorname{div} (\eta^{n+1} D\mathbf{u}^{n+1}) + \nabla p^{n+1} = \text{terme de force capillaire} + \varrho^{n+1} \mathbf{g}, \\ \operatorname{div} (\mathbf{u}^{n+1}) = 0, \end{cases}$$

la discrétisation du terme de force capillaire étant décrite dans les paragraphes ci-après.

Prise en compte des termes de couplage

Nous précisons maintenant la discrétisation des termes relatifs au couplage des deux systèmes. Il s'agit des termes de transport $\mathbf{u} \cdot \nabla c_i$ dans les équations de Cahn-Hilliard et du terme de force capillaire $\sum_{i=1}^3 \mu_i \nabla c_i$ dans le bilan de quantité de mouvement de l'équation de Navier-Stokes. Au niveau continu, lors de l'écriture du

bilan d'énergie, les contributions de ces deux termes se compensent exactement (*cf* sections IV.1.2 et IV.2.2). Nous avons vu dans les paragraphes précédents que, lors de l'écriture du bilan d'énergie discret, les termes de transport sont multipliés par μ_i^{n+1} avant d'être sommés pour $i = 1, 2, 3$ et que le terme de force capillaire est multiplié par $\mathbf{u}^{n+1}$.

En conséquence, il est facile de voir que, lorsque tous les termes précités sont discrétisés de manière implicite (*cf* [Fen06] dans le cas diphasique), *i.e.* $\mathbf{u}^{n+1} \cdot \nabla c_i^{n+1}$ et $\sum_{i=1}^3 \mu_i^{n+1} \nabla c_i^{n+1}$, la compensation de leurs contributions au bilan d'énergie est encore vraie au niveau discret. Mais cette discrétisation introduit un couplage fort entre le système de Cahn-Hilliard et celui de Navier-Stokes, il est alors difficile de résoudre en pratique le problème discret obtenu.

Classiquement, le découplage s'effectue (*cf* [KSW08] dans le cas diphasique, [BLM⁺] dans le cas triphasique) en explicitant le terme de vitesse dans l'équation de Cahn-Hilliard : $\mathbf{u}^n \cdot \nabla c_i^{n+1}$ mais alors, les contributions des termes de transport dans le système de Cahn-Hilliard et de force capillaire dans les équations de Navier-Stokes ne se compensent plus lors de l'établissement du bilan d'énergie qui contient le terme additionnel $(\mathbf{u}^{n+1} - \mathbf{u}^n) \cdot \sum_{i=1}^3 \mu_i^{n+1} \nabla c_i^{n+1}$ (auquel il est difficile d'attribuer un signe). La stabilité du schéma n'est alors obtenue que conditionnellement (*cf* [KSW08] dans le cas diphasique), en supposant par exemple que le rapport du pas de temps sur le pas de maillage est borné.

Nous constatons qu'il est possible de découpler la résolution du système de Navier-Stokes de la prise en compte du terme de force capillaire. Cette prise en compte est alors effectuée dans une première étape qui fournit une vitesse intermédiaire $\mathbf{u}^*$ utilisée comme vitesse d'advection dans le système de Cahn-Hilliard. Le système de Navier-Stokes est résolu de manière totalement découplée dans une seconde étape.

(i) prise en compte des forces capillaires :

$$\begin{cases} \varrho^n \frac{\mathbf{u}^* - \mathbf{u}^n}{\Delta t} + \nabla p^* = \sum_{i=1}^3 \mu_i^{n+1} \nabla c_i^{n+1}, \\ \operatorname{div}(\mathbf{u}^*) = 0, \end{cases}$$

(ii) système de Cahn-Hilliard :

$$\begin{cases} \frac{c_i^{n+1} - c_i^n}{\Delta t} + \mathbf{u}^* \cdot \nabla c_i^{n+1} = \operatorname{div} \left(\frac{M_0^{n+\alpha}}{\Sigma_i} \nabla \mu_i^{n+1} \right), \\ \mu_i^{n+1} = D_i^F(\mathbf{c}^n, \mathbf{c}^{n+1}) - \frac{3}{4} \varepsilon \Sigma_i \Delta c_i^{n+\beta}, \end{cases}$$

(iii) système de Navier-Stokes :

$$\begin{cases} \varrho^n \frac{\mathbf{u}^{n+1} - \mathbf{u}^*}{\Delta t} + \frac{1}{2} \frac{\varrho^{n+1} - \varrho^n}{\Delta t} \mathbf{u}^{n+1} + (\varrho^{n+1} \mathbf{u}^n \cdot \nabla) \mathbf{u}^{n+1} + \frac{\mathbf{u}^{n+1}}{2} \operatorname{div}(\varrho^{n+1} \mathbf{u}^n) \\ \quad - \operatorname{div}(\eta^{n+1} D \mathbf{u}^{n+1}) + \nabla(p^{n+1} - p^*) = \varrho^{n+1} \mathbf{g}, \\ \operatorname{div}(\mathbf{u}^{n+1}) = 0. \end{cases}$$

Cette discrétisation est inconditionnellement stable mais le système de l'étape (i) (de type Darcy) reste couplé aux équations de Cahn-Hilliard (système (ii)).

Il est alors envisageable de ne pas tenir compte de la contrainte de divergence nulle imposée à $\mathbf{u}^*$ (et en conséquence du terme de pression ∇p^*) dans le système de l'étape (i) et ainsi poser directement

$$\varrho^n \frac{\mathbf{u}^* - \mathbf{u}^n}{\Delta t} = \sum_{i=1}^3 c_i^{n+1} \nabla \mu_i^{n+1}.$$

La définition de $\mathbf{u}^*$ est alors explicite et $\mathbf{u}^*$ peut être remplacé par son expression dans le système de Cahn-Hilliard supprimant tout couplage (au prix d'une non linéarité supplémentaire).

Le problème est alors que $\mathbf{u}^*$ n'est plus à divergence nulle, nous conservons néanmoins la propriété $\mathbf{u}^* \cdot \mathbf{n} = 0$ sur Γ . Est-il alors possible de discrétiser le terme de transport dans l'équation de Cahn-Hilliard de manière à conserver les propriétés du système de Cahn-Hilliard (conservation du volume et somme des trois paramètres d'ordre égale à 1)? Nous discutons cet aspect dans le paragraphe suivant.

Terme de transport du système de Cahn-Hilliard lorsque la vitesse n'est pas à divergence nulle

Dans ce paragraphe, nous nous intéressons à la forme que doit prendre le terme de transport dans l'équation de Cahn-Hilliard lorsque la vitesse d'advection, notée $\mathbf{u}^*$ n'est pas à divergence nulle mais satisfait $\mathbf{u}^* \cdot \mathbf{n} = 0$ sur le bord du domaine.

Remarque VI.2

Conserver les propriétés du système de Cahn-Hilliard lorsque le champ advectif n'est pas à divergence nulle peut être intéressant dans d'autres contextes. Par exemple, lors de l'utilisation d'une méthode de projection incrémentale (cf chapitre VII), la vitesse obtenue in fine n'est pas à divergence nulle.

Le terme de transport peut s'écrire sous forme conservative ou non conservative (ces formes n'étant plus équivalentes puisque *a priori* $\text{div}(\mathbf{u}^*) \neq 0$) :

- forme non conservative : $\mathbf{u}^* \cdot \nabla c_i$,
- forme conservative : $\text{div}(c_i \mathbf{u}^*)$.

La forme conservative permet de garantir la conservation du volume alors que la formulation non conservative ne le permet pas puisque *a priori* $\int_{\Omega} \mathbf{u}^* \cdot \nabla c_i dx \neq 0$. Par ailleurs, une condition nécessaire pour que la somme des c_i reste constante égale à 1 lorsque l'on utilise la forme conservative est que $\text{div}(\mathbf{u}^*) = 0$. Aucune de ces deux formes ne permet donc d'obtenir à la fois la conservation du volume et de garantir que la somme des paramètres d'ordre est égale à 1.

Nous proposons d'utiliser la formulation suivante :

$$\text{div}((c_i - \alpha_i) \mathbf{u}^*),$$

où α_i est une constante à déterminer. Cette formulation permet de satisfaire les deux propriétés souhaitées si $\sum_{i=1}^3 \alpha_i = 1$. Pour des raisons de consistance algébrique, il est souhaitable que lorsque l'une des phases n'est pas présente la constante α_i soit nulle. Dans la suite, nous proposons de choisir

$$\alpha_i = \int_{\Omega} c_i^0 dx.$$

Remarque VI.3

Il serait peut-être envisageable d'adopter une définition locale des α_i par exemple sur le support de la fonction de base correspondant à l'équation mais il est alors plus délicat d'obtenir le bilan d'énergie.

Ainsi, cette formulation nous permet d'utiliser une vitesse qui n'est pas strictement à divergence nulle. Le terme $-\alpha_i \text{div}(\mathbf{u}^*)$ est ajouté dans le système de Cahn-Hilliard, moralement son rôle est de ré-équilibrer la valeur de chaque paramètre d'ordre pour garantir le fait que leur somme reste égale à 1. Nous montrerons dans la section VI.5 que ce terme est moralement en $O(h + \Delta t)$ et donc qu'il ne gêne pas la consistance du schéma.

Au vu de cette formulation du terme de transport, il semble naturel d'adopter, dans le système de Navier-Stokes, la définition suivante des forces capillaires $-\sum_{i=1}^3 (c_i - \alpha_i) \nabla \mu_i$. Ceci revient à changer la définition de la pression en y incluant le terme $\sum_{i=1}^3 (c_i - \alpha_i) \mu_i$.

Discrétisation en temps du système de Cahn-Hilliard/Navier-Stokes

Finalement, les différentes considérations exposées dans les paragraphes précédents nous amènent à proposer le schéma suivant :

Problème VI.4

– Etape 1 : résolution du système de Cahn-Hilliard

Trouver $(\mathbf{c}^{n+1}, \boldsymbol{\mu}^{n+1}) \in (\mathcal{V}^c)^3 \times (\mathcal{V}^\mu)^3$ tel que, pour $i = 1, 2$ et 3 ,

$$\begin{cases} \frac{c_i^{n+1} - c_i^n}{\Delta t} + \text{div} \left([c_i^n - \alpha_i] \left[\mathbf{u}^n - \frac{\Delta t}{\varrho^n} \sum_{j=1}^3 (c_j^n - \alpha_j) \nabla \mu_j^{n+1} \right] \right) = \text{div} \left(\frac{M_0^{n+\alpha}}{\Sigma_i} \nabla \mu_i^{n+1} \right), \\ \mu_i^{n+1} = D_i^F(\mathbf{c}^n, \mathbf{c}^{n+1}) - \frac{3}{4} \Sigma_i \varepsilon \Delta c_i^{n+\beta}, \end{cases} \quad (\text{VI.2})$$

avec α_j une constante : $\alpha_j = \int_{\Omega} c_j^0 dx$.

– Etape 2 : résolution des équations de Navier-Stokes

Trouver $(\mathbf{u}^{n+1}, p^{n+1}) \in \mathcal{V}_0^{\mathbf{u}} \times \mathcal{V}^p$ tel que,

$$\left\{ \begin{array}{l} \varrho^n \frac{\mathbf{u}^{n+1} - \mathbf{u}^n}{\Delta t} + \frac{1}{2} \frac{\varrho^{n+1} - \varrho^n}{\Delta t} \mathbf{u}^{n+1} \\ \quad + (\varrho^{n+1} \mathbf{u}^n \cdot \nabla) \mathbf{u}^{n+1} + \frac{\mathbf{u}^{n+1}}{2} \operatorname{div} (\varrho^{n+1} \mathbf{u}^n) + \operatorname{div} (2\eta^{n+1} D\mathbf{u}^{n+1}) \\ \quad + \nabla p^{n+1} = \varrho^{n+1} \mathbf{g} + \sum_{j=1}^3 (c_j^n - \alpha_j) \nabla \mu_j^{n+1}, \\ \operatorname{div} (\mathbf{u}^{n+1}) = 0, \end{array} \right. \quad (\text{VI.3})$$

où $\eta^{n+1} = \eta(\mathbf{c}_h^{n+1})$ et $\varrho^\ell = \varrho(\mathbf{c}^\ell)$, pour $\ell = n$ et $\ell = n+1$.

Remarque VI.5

Nous avons explicité le paramètre d'ordre dans les termes de transport du système de Cahn-Hilliard $\operatorname{div} \left([c_i^n - \alpha_i] [\mathbf{u}^n - \frac{\Delta t}{\varrho^n} \sum_{j=1}^3 (c_j^n - \alpha_j) \nabla \mu_j^{n+1}] \right)$ et dans le terme de force capillaire du système de Navier-Stokes $\sum_{j=1}^3 (c_j^n - \alpha_j) \nabla \mu_j^{n+1}$. Il est également possible d'utiliser une version implicite, i.e. $\operatorname{div} \left([c_i^{n+1} - \alpha_i] [\mathbf{u}^n - \frac{\Delta t}{\varrho^n} \sum_{j=1}^3 (c_j^{n+1} - \alpha_j) \nabla \mu_j^{n+1}] \right)$ et $\sum_{j=1}^3 (c_j^{n+1} - \alpha_j) \nabla \mu_j^{n+1}$, mais ceci introduit une non-linéarité supplémentaire dans le système de Cahn-Hilliard.

Pour finir, il est intéressant d'examiner le schéma que l'on obtient pour la résolution du système de Cahn-Hilliard lorsque la vitesse $\mathbf{u}^n$ est nulle :

$$\left\{ \begin{array}{l} \frac{c_i^{n+1} - c_i^n}{\Delta t} = \operatorname{div} \left([c_i^n - \alpha_i] \left[\frac{\Delta t}{\varrho^n} \sum_{j=1}^3 (c_j^n - \alpha_j) \nabla \mu_j^{n+1} \right] \right) + \operatorname{div} \left(\frac{M_0^{n+\alpha}}{\Sigma_i} \nabla \mu_i^{n+1} \right), \\ \mu_i^{n+1} = D_i^F(\mathbf{c}^n, \mathbf{c}^{n+1}) - \frac{3}{4} \Sigma_i \varepsilon \Delta c_i^{n+\beta}. \end{array} \right. \quad (\text{VI.4})$$

Ce schéma diffère de celui présenté dans la section V.1 par l'ajout d'une diffusion supplémentaire (de coefficient d'ordre Δt). L'estimation d'énergie sur ce système donne :

$$\begin{aligned} \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}^{n+1}) - \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}^n) + \Delta t \sum_{i=1}^3 \int_{\Omega} \frac{M_0^{n+\alpha}}{\Sigma_i} |\nabla \mu_i^{n+1}|^2 dx + \Delta t^2 \int_{\Omega} \frac{1}{\varrho^n} \left| \sum_{i=1}^3 (c_i^n - \alpha_i) \nabla \mu_i^{n+1} \right|^2 dx \\ + \frac{3}{8} (2\beta - 1) \varepsilon \int_{\Omega} \sum_{i=1}^3 \Sigma_i |\nabla c_i^{n+1} - \nabla c_i^n|^2 dx = \frac{12}{\varepsilon} \int_{\Omega} [F(\mathbf{c}^{n+1}) - F(\mathbf{c}^n) - \mathbf{d}^F(\mathbf{c}^n, \mathbf{c}^{n+1}) \cdot (\mathbf{c}^{n+1} - \mathbf{c}^n)] dx. \end{aligned}$$

Ainsi, le terme ajouté contribue à la décroissance de l'énergie. En particulier le schéma que nous proposons permet de calculer correctement les états d'équilibre.

VI.1.2 Discrétisation en espace

La discrétisation en espace est réalisée grâce à la méthode de Galerkin et à la méthode des éléments finis. Soient $\mathcal{V}_h^c$, $\mathcal{V}_h^\mu$, $\mathcal{V}_h^{\mathbf{u}}$ et $\mathcal{V}_h^p$ des espaces d'approximation éléments finis de $\mathcal{V}^c$, $\mathcal{V}^\mu$, $\mathcal{V}^{\mathbf{u}}$ et $\mathcal{V}^p$ respectivement. Puisque la vitesse vérifie des conditions de Dirichlet homogènes sur la frontière Γ , nous définissons l'espace d'approximation suivant :

$$\mathcal{V}_{h,0}^{\mathbf{u}} = \left\{ \nu_h^{\mathbf{u}} \in \mathcal{V}_h^{\mathbf{u}}; \nu_h^{\mathbf{u}} = 0 \text{ sur } \Gamma \right\}.$$

Enfin, pour simplifier les écritures, nous introduisons l'espace :

$$\mathcal{V}_{h,\mathcal{S}}^c = \left\{ \mathbf{c}_h = (c_{1h}, c_{2h}, c_{3h}) \in (\mathcal{V}_h^c)^3; \mathbf{c}_h(x) \in \mathcal{S} \text{ pour presque tout } x \in \Omega \right\}.$$

Les hypothèses générales requises sur les espaces d'approximation sont les suivantes :

$$\bullet 1 \in \mathcal{V}_h^c \quad \text{et} \quad 1 \in \mathcal{V}_h^\mu, \quad (\text{VI.5})$$

$$\bullet \forall \nu^c \in \mathcal{V}^c, \quad \inf_{\nu_h^c \in \mathcal{V}_h^c} |\nu^c - \nu_h^c|_{H^1(\Omega)} \xrightarrow{h \rightarrow 0} 0, \quad (\text{VI.6})$$

$$\bullet \forall \nu^\mu \in \mathcal{V}^\mu, \quad \inf_{\nu_h^\mu \in \mathcal{V}_h^\mu} |\nu^\mu - \nu_h^\mu|_{H^1(\Omega)} \xrightarrow{h \rightarrow 0} 0, \quad (\text{VI.7})$$

$$\bullet \forall \nu^u \in \mathcal{V}_0^u, \quad \inf_{\nu_h^u \in \mathcal{V}_{h,0}^u} |\nu^u - \nu_h^u|_{(H^1(\Omega))^d} \xrightarrow{h \rightarrow 0} 0, \quad (\text{VI.8})$$

$$\bullet \forall \nu^p \in \mathcal{V}^p, \quad \inf_{\nu_h^p \in \mathcal{V}_h^p} |\nu^p - \nu_h^p|_{(L^2(\Omega))^d} \xrightarrow{h \rightarrow 0} 0, \quad (\text{VI.9})$$

• il existe une constante strictement positive β (indépendante de h) telle que

$$\inf_{\nu_h^p \in \mathcal{V}_h^p} \sup_{\nu_h^u \in \mathcal{V}_{h,0}^u} \frac{\int_{\Omega} \nu_h^p \operatorname{div} \nu_h^u dx}{|\nu_h^p|_{L^2(\Omega)} |\nu_h^u|_{(H^1(\Omega))^d}} \geq \beta, \quad (\text{VI.10})$$

• il existe une constante strictement positive C indépendante de h telle que

$$\forall \nu^c \in \mathcal{V}^c, \quad \left| \Pi_0^{\mathcal{V}^c}(\nu^c) \right|_{H^1(\Omega)} \leq |\nu^c|_{H^1(\Omega)} \quad \text{et} \quad \forall \nu^\mu \in \mathcal{V}^\mu, \quad \left| \Pi_0^{\mathcal{V}^\mu}(\nu^\mu) \right|_{H^1(\Omega)} \leq C |\nu^\mu|_{H^1(\Omega)}, \quad (\text{VI.11})$$

où $\Pi_0^{\mathcal{V}^h}$ est la projection $L^2(\Omega)$ sur $\mathcal{V}_h$,

• il existe une fonction C_{inv} de h telle que

$$\forall \nu_h^c \in \mathcal{V}_h^c, \quad |\nu_h^c|_{L^\infty(\Omega)}^2 \leq C_{\text{inv}}(h) |\nu_h^c|_{H^1(\Omega)}^2, \quad (\text{VI.12})$$

$$\bullet \mathcal{V}_h^c \subset \mathcal{V}_h^\mu. \quad (\text{VI.13})$$

Remarque VI.6

En plus des hypothèses requises sur les espaces d'approximation des paramètres d'ordre et des potentiels chimiques énoncées au chapitre V, nous supposons que l'espace d'approximation des paramètres d'ordre satisfait l'inégalité inverse (VI.12). Celle-ci est par exemple satisfaite pour une famille de maillages quasi-uniformes et des espaces d'approximation associés à des éléments finis de Lagrange correspondants, dans ce cas on peut choisir $C_{\text{inv}}(h) = C(1 + \ln(h))$ si $d = 2$ et $C_{\text{inv}}(h) = Ch^{-1}$ si $d = 3$ où C est une constante ne dépendant que de la régularité du maillage et non de h (cf [BS08, 4.5.11 (p. 112) et 4.9.2 (p. 123)]). Il est en outre nécessaire que les espaces d'approximation de la vitesse et de la pression satisfassent la condition inf-sup.

Nous commençons par définir les fonctions discrètes $\mathbf{c}_h^0 \in \mathcal{V}_S^c$ et $\mathbf{u}_h^0 \in \mathcal{V}_{h,0}^u$ à l'instant initial de manière que :

$$\mathbf{c}_h^0(x) \in \mathcal{S}, \quad \forall h > 0, \quad \text{pour presque tout } x \in \Omega \quad \text{et} \quad |\mathbf{c}_h^0 - \mathbf{c}^0|_{(H^1(\Omega))^3} \xrightarrow{h \rightarrow 0} 0, \quad (\text{VI.14})$$

$$|\mathbf{u}_h^0 - \mathbf{u}^0|_{(H^1(\Omega))^d} \xrightarrow{h \rightarrow 0} 0. \quad (\text{VI.15})$$

Ces fonctions discrètes $\mathbf{c}_h^0$ et $\mathbf{u}_h^0$ peuvent être obtenues à partir des conditions initiales $\mathbf{c}^0$ et $\mathbf{u}^0$ par projection $H^1(\Omega)$, ou comme c'est le cas en pratique, par interpolation élément fini pourvu que $\mathbf{c}_i^0$ et $\mathbf{u}^0$ soient assez régulières.

Supposons que $\mathbf{c}_h^n \in \mathcal{V}_S^c$ et $\mathbf{u}_h^n \in \mathcal{V}_{h,0}^u$ sont donnés, l'approximation de Galerkin du problème (V.1) au temps t_{n+1} s'écrit de la manière suivante :

Problème VI.7 (Formulation avec trois paramètres d'ordre)

– Etape 1 : résolution du système de Cahn-Hilliard

Trouver $(\mathbf{c}_h^{n+1}, \boldsymbol{\mu}_h^{n+1}) \in (\mathcal{V}_{h,S}^c)^3 \times (\mathcal{V}_h^\mu)^3$ tels que $\forall \nu_h^c \in \mathcal{V}_h^c, \forall \nu_h^\mu \in \mathcal{V}_h^\mu$, nous avons, pour $i = 1, 2$ et 3 ,

$$\left\{ \begin{array}{l} \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx - \int_{\Omega} [c_{ih}^n - \alpha_{ih}] \left[\mathbf{u}_h^n - \frac{\Delta t}{\varrho_h^n} \sum_{j=1}^3 (c_{jh}^n - \alpha_{jh}) \nabla \mu_{jh}^{n+1} \right] \cdot \nabla \nu_h^\mu dx \\ \qquad \qquad \qquad = - \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} \nabla \mu_{ih}^{n+1} \cdot \nabla \nu_h^\mu dx, \quad (\text{VI.16}) \\ \int_{\Omega} \mu_{ih}^{n+1} \nu_h^c dx = \int_{\Omega} D_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \nu_h^c dx + \int_{\Omega} \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^{n+\beta} \nabla \nu_h^c dx, \end{array} \right.$$

où α_{jh} est la constante définie par $\alpha_{jh} = \int_{\Omega} c_{jh}^0 dx$.

– Etape 2 : résolution des équations de Navier-Stokes

Trouver $(\mathbf{u}_h^{n+1}, p_h^{n+1}) \in \mathcal{V}_{h,0}^u \times \mathcal{V}_h^p$ tels que $\forall \nu_h^u \in \mathcal{V}_{h,0}^u, \forall \nu_h^p \in \mathcal{V}_h^p$,

$$\left\{ \begin{array}{l} \int_{\Omega} \varrho_h^n \frac{\mathbf{u}_h^{n+1} - \mathbf{u}_h^n}{\Delta t} \cdot \nu_h^u dx + \frac{1}{2} \int_{\Omega} \frac{\varrho_h^{n+1} - \varrho_h^n}{\Delta t} \mathbf{u}_h^{n+1} \cdot \nu_h^u dx \\ \quad + \frac{1}{2} \int_{\Omega} \varrho_h^{n+1} (\mathbf{u}_h^n \cdot \nabla) \mathbf{u}_h^{n+1} \cdot \nu_h^u dx - \frac{1}{2} \int_{\Omega} \varrho_h^{n+1} (\mathbf{u}_h^n \cdot \nabla) \nu_h^u \cdot \mathbf{u}_h^{n+1} dx \\ \quad + \int_{\Omega} 2\eta_h^{n+1} D\mathbf{u}_h^{n+1} : D\nu_h^u dx - \int_{\Omega} p_h^{n+1} \operatorname{div}(\nu_h^u) dx \\ \qquad \qquad \qquad = \int_{\Omega} \varrho_h^{n+1} \mathbf{g} \cdot \nu_h^u dx - \int_{\Omega} \sum_{j=1}^3 (c_{jh}^n - \alpha_{jh}) \nabla \mu_{jh}^{n+1} \cdot \nu_h^u dx, \quad (\text{VI.17}) \\ \int_{\Omega} \nu_h^p \operatorname{div}(\mathbf{u}_h^{n+1}) dx = 0, \end{array} \right.$$

où $\eta_h^{n+1} = \eta(\mathbf{c}_h^{n+1})$ et $\varrho_h^\ell = \varrho(\mathbf{c}_h^\ell)$, pour $\ell = n$ et $\ell = n+1$.

VI.1.3 Equivalence du système de Cahn-Hilliard avec un système de deux équations

En pratique, pour la résolution du système de Cahn-Hilliard, nous résolvons seulement les équations satisfaites par (c_1, c_2, μ_1, μ_2) . En effet, le problème VI.7 est équivalent au suivant :

Problème VI.8 (Formulation avec deux paramètres d'ordre)

– Etape 1 : résolution du système de Cahn-Hilliard

Trouver $(c_{1h}^{n+1}, c_{2h}^{n+1}, \mu_{1h}^{n+1}, \mu_{2h}^{n+1}) \in (\mathcal{V}_h^c)^2 \times (\mathcal{V}_h^\mu)^2$ tels que $\forall \nu_h^c \in \mathcal{V}_h^c, \forall \nu_h^\mu \in \mathcal{V}_h^\mu$, pour $i = 1$ et 2 ,

$$\left\{ \begin{array}{l} \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx - \int_{\Omega} [c_{ih}^n - \alpha_i] \left[\mathbf{u}_h^n - \frac{\Delta t}{\varrho_h^n} \sum_{j=1}^3 (c_{jh}^n - \alpha_j) \nabla \mu_{jh}^{n+1} \right] \cdot \nabla \nu_h^\mu dx \\ \qquad \qquad \qquad = - \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} \nabla \mu_{ih}^{n+1} \cdot \nabla \nu_h^\mu dx, \quad (\text{VI.18}) \\ \int_{\Omega} \mu_{ih}^{n+1} \nu_h^c dx = \int_{\Omega} D_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \nu_h^c dx + \int_{\Omega} \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^{n+\beta} \nabla \nu_h^c dx. \end{array} \right.$$

avec $\mathbf{c}_h^{n+1} = (c_{1h}^{n+1}, c_{2h}^{n+1}, 1 - c_{1h}^{n+1} - c_{2h}^{n+1})$.

– Nous définissons ensuite :

$$c_{3h}^{n+1} = 1 - c_{1h}^{n+1} - c_{2h}^{n+1} \quad \text{et} \quad \mu_{3h}^{n+1} = - \left(\frac{\Sigma_3}{\Sigma_1} \mu_{1h}^{n+1} + \frac{\Sigma_3}{\Sigma_2} \mu_{2h}^{n+1} \right). \quad (\text{VI.19})$$

– Etape 2 : la résolution du problème de Navier-Stokes reste inchangée (cf problème VI.7).

Remarque VI.9

Comme nous l'avons déjà mentionné dans le chapitre V, dans les systèmes où seulement les inconnues $(c_{1h}^{n+1}, \mu_{1h}^{n+1}, c_{2h}^{n+1}, \mu_{2h}^{n+1})$ sont présentes, la notation $\mathbf{c}_h^{n+1}$ désigne le vecteur $(c_{1h}^{n+1}, c_{2h}^{n+1}, 1 - c_{1h}^{n+1} - c_{2h}^{n+1})$.

Théorème VI.10

Le problème VI.7 est équivalent au problème VI.8. En particulier, toute solution $(\mathbf{c}_h^{n+1}, \boldsymbol{\mu}_h^{n+1}, \mathbf{u}^{n+1}, p^{n+1})$ du problème VI.7 satisfait

$$\sum_{i=1}^3 c_{ih}^{n+1} = 1 \quad \text{et} \quad \sum_{i=1}^3 \frac{\mu_{ih}^{n+1}}{\Sigma_i} = 0. \quad (\text{VI.20})$$

Démonstration : Nous avons démontré, au cours de la preuve du théorème V.6, la relation suivante (cf (V.13)) (j et k désignant les deux indices différents de i) :

$$\sum_{i=1}^3 \frac{1}{\Sigma_i} D_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) = 0. \quad (\text{VI.21})$$

Supposons maintenant que le problème VI.8 soit satisfait. Alors, en ajoutant les équations du système (VI.18) pour $i = 1, 2$ et en utilisant (VI.19) et (VI.21), nous obtenons

$$\left\{ \begin{array}{l} \int_{\Omega} \frac{(1 - c_{3h}^{n+1}) - (1 - c_{3h}^n)}{\Delta t} \nu_h^\mu dx - \int_{\Omega} [(1 - c_{3h}^n) - (1 - \alpha_3)] [\mathbf{u}_h^n - \frac{\Delta t}{\varrho_h^n} \sum_{j=1}^3 (c_{jh}^n - \alpha_j) \nabla \mu_{jh}^{n+1}] \cdot \nabla \nu_h^\mu dx \\ \quad = - \int_{\Omega} M_{0h}^{n+\alpha} \nabla \left(-\frac{\mu_{3h}^{n+1}}{\Sigma_3} \right) \cdot \nabla \nu_h^\mu dx, \\ \int_{\Omega} \left(-\frac{\mu_{3h}^{n+1}}{\Sigma_3} \right) \nu_h^c dx = \int_{\Omega} \left(-\frac{1}{\Sigma_3} D_3(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \right) \nu_h^c dx + \frac{3}{4} \varepsilon \int_{\Omega} \nabla (1 - c_{3h}^{n+\beta}) \cdot \nabla \nu_h^c dx. \end{array} \right.$$

Cela prouve que c_{3h}^{n+1} satisfait (VI.2) pour $i = 3$.

Réciproquement, si nous supposons que le problème VI.7 est satisfait, alors en ajoutant les équations pour $i = 1, 2, 3$, grâce à l'égalité (VI.21), et puisque $\sum_{i=1}^3 c_{ih}^n \equiv 1$ ($\mathbf{c}_h^n \in \mathcal{V}_{Dh,S}$), nous obtenons

$$\left\{ \begin{array}{l} \int_{\Omega} \frac{S_h^{n+1} - S_h^n}{\Delta t} \nu_h^\mu dx = - \int_{\Omega} M_{0h}^{n+\alpha} \nabla \Theta_h^{n+1} \cdot \nabla \nu_h^\mu dx \\ \int_{\Omega} \Theta_h^{n+1} \nu_h^c dx = \frac{3}{4} \varepsilon \int_{\Omega} [(1 - \beta) \nabla S_h^n + \beta \nabla S_h^{n+1}] \cdot \nabla \nu_h^c dx, \end{array} \right. \quad (\text{VI.22})$$

où $S_h^\ell = \sum_{i=1}^3 c_{ih}^\ell$ et $\Theta_h^\ell = \sum_{i=1}^3 \frac{\mu_{ih}^\ell}{\Sigma_i}$ pour $\ell = n$ et $\ell = n+1$. Ce système d'équations est exactement celui obtenu

dans la preuve du théorème V.6 et les mêmes arguments permettent de conclure que $S_h^{n+1} \equiv 1$ et $\Theta_h^{n+1} \equiv 0$. Ainsi, le couple $(\mathbf{c}_h^{n+1}, \boldsymbol{\mu}_h^{n+1})$ satisfait (VI.20) et en conséquence le système (VI.18)-(VI.19). ■

VI.2 Schéma correspondant dans le cas diphasique

Considérons un système avec deux composants (noté ci-dessous avec les indices 1 et 2 respectivement) et supposons que l'évolution des paramètres d'ordre c_i , ($i = 1, 2$) et des potentiels chimiques $\tilde{\mu}_i$, ($i = 1, 2$) associés à ces deux phases est gouvernée par le modèle de Cahn-Hilliard diphasique :

$$\left\{ \begin{array}{l} \frac{\partial c_i}{\partial t} = \operatorname{div} (M(c_1, c_2) \nabla \tilde{\mu}_i), \quad \text{pour } i = 1, 2, \\ \tilde{\mu}_i = \frac{12}{\varepsilon} \sigma_{12} f'(c_i) - \frac{3}{2} \varepsilon \sigma_{12} \Delta c_i \quad \text{pour } i = 1, 2, \end{array} \right. \quad (\text{VI.23})$$

où ε est l'épaisseur d'interface, $M(c_1, c_2)$ la mobilité et σ_{12} la tension de surface entre les deux phases. Les inconnues sont liées par les relations suivantes : $c_1 + c_2 = 1$ et $\tilde{\mu}_1 + \tilde{\mu}_2 = 0$.

La consistance algébrique du modèle triphasique (cf section IV.2.1) garantit que le triplet $(c_1, c_2 = 1 - c_1, c_3 = 0)$ est une solution particulière du modèle de Cahn-Hilliard triphasique (IV.9) (avec $M_0(\mathbf{c}) = 2\sigma_{12}M(c_1, c_2)$) et pour tout choix des valeurs des tensions de surface σ_{13} et σ_{23} impliquant le troisième composant. Dans ce cas, les potentiels chimiques triphasiques sont donnés par $\mu_i = \frac{\Sigma_i}{2\sigma_{12}}\tilde{\mu}_i$ pour $i = 1, 2$ et $\mu_3 = 0$.

Nous donnons l'équivalent de ce résultat au niveau discret, en considérant la discrétisation suivante du modèle diphasique :

Problème VI.11

- *Etape 1 : résolution du système de Cahn-Hilliard*
Trouver $(c_{1h}^{n+1}, c_{2h}^{n+1}, \tilde{\mu}_{1h}^{n+1}, \tilde{\mu}_{2h}^{n+1}) \in (\mathcal{V}_h^c)^2 \times (\mathcal{V}_h^\mu)^2$ tels que $\forall \nu_h^c \in \mathcal{V}_h^c, \forall \nu_h^\mu \in \mathcal{V}_h^\mu$, pour $i = 1$ et 2 , nous avons,

$$\left\{ \begin{array}{l} \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx - \int_{\Omega} (c_{ih}^n - \alpha_i) \mathbf{u}_h^n \cdot \nabla \nu_h^\mu dx \\ \qquad \qquad \qquad = - \int_{\Omega} \left[M + \frac{\Delta t}{\varrho_h^n} (c_{ih}^n - \alpha_i)^2 \right] \nabla \tilde{\mu}_{ih}^{n+1} \cdot \nabla \nu_h^\mu dx, \\ \int_{\Omega} \mu_{ih}^{n+1} \nu_h^c dx = \int_{\Omega} D_i^F((c_{1h}^n, c_{2h}^n, 0), (c_{1h}^{n+1}, c_{2h}^{n+1}, 0)) \nu_h^c dx + \int_{\Omega} \frac{3}{4} \Sigma_i \nabla c_{ih}^{n+\beta} \nabla \nu_h^c dx. \end{array} \right. \quad (\text{VI.24})$$

- *Etape 2 : résolution du système de Navier-Stokes*
Trouver $(\mathbf{u}_h^{n+1}, p_h^{n+1}) \in \mathcal{V}_{h,0}^{\mathbf{u}} \times \mathcal{V}_h^p$ tels que $\forall \nu_h^{\mathbf{u}} \in \mathcal{V}_{h,0}^{\mathbf{u}}, \forall \nu_h^p \in \mathcal{V}_h^p$,

$$\left\{ \begin{array}{l} \int_{\Omega} \varrho_h^n \frac{\mathbf{u}_h^{n+1} - \mathbf{u}_h^n}{\Delta t} \cdot \nu_h^{\mathbf{u}} dx + \frac{1}{2} \int_{\Omega} \frac{\varrho_h^{n+1} - \varrho_h^n}{\Delta t} \mathbf{u}_h^{n+1} \cdot \nu_h^{\mathbf{u}} dx \\ \quad + \frac{1}{2} \int_{\Omega} \varrho_h^{n+1} (\mathbf{u}_h^n \cdot \nabla) \mathbf{u}_h^{n+1} \cdot \nu_h^{\mathbf{u}} dx - \frac{1}{2} \int_{\Omega} \varrho_h^{n+1} (\mathbf{u}_h^n \cdot \nabla) \nu_h^{\mathbf{u}} \cdot \mathbf{u}_h^{n+1} dx \\ \quad + \int_{\Omega} 2\eta_h^{n+1} D\mathbf{u}_h^{n+1} : D\nu_h^{\mathbf{u}} dx - \int_{\Omega} p_h^{n+1} \text{div}(\nu_h^{\mathbf{u}}) dx \\ \qquad \qquad \qquad = \int_{\Omega} \varrho_h^{n+1} \mathbf{g} \cdot \nu_h^{\mathbf{u}} dx - \frac{1}{2} \int_{\Omega} \sum_{j=1}^2 (c_{jh}^n - \alpha_{jh}) \nabla \tilde{\mu}_{jh}^{n+1} \cdot \nu_h^{\mathbf{u}} dx, \\ \int_{\Omega} \nu_h^p \text{div}(\mathbf{u}_h^{n+1}) dx = 0, \end{array} \right. \quad (\text{VI.25})$$

où $\eta_h^{n+1} = \eta(\mathbf{c}_h^{n+1})$ et $\varrho_h^\ell = \varrho(\mathbf{c}_h^\ell)$, pour $\ell = n$ et $\ell = n+1$.

Dans le cas diphasique la différence essentielle entre le schéma que nous proposons ici et le schéma plus classique de la littérature [KSW08] est l'ajout du terme $\frac{\Delta t}{\varrho_h^n} (c_{ih}^n - \alpha_i)^2$ dans le coefficient de mobilité. Le terme additionnel peut être interprété comme une diffusion supplémentaire (faible puisque proportionnelle à Δt) stabilisant le schéma plus standard de [KSW08].

Proposition VI.12

En définissant $M_0 = 2\sigma_{12}M$, $\mu_{ih}^{n+1} = \frac{\Sigma_i}{2\sigma_{12}}\tilde{\mu}_{ih}^{n+1}$ pour $i = 1, 2$ et $\mu_{3h}^{n+1} = 0$, nous avons le résultat suivant : si $((c_{1h}^{n+1}, \tilde{\mu}_{1h}^{n+1}), (c_{2h}^{n+1}, \tilde{\mu}_{2h}^{n+1}))$ est une solution du problème diphasique VI.11 alors $((c_{1h}^{n+1}, \mu_{1h}^{n+1}), (c_{2h}^{n+1}, \mu_{2h}^{n+1}), (0, 0))$ est une solution du problème triphasique VI.7. Inversement, si $((c_{1h}^{n+1}, \mu_{1h}^{n+1}), (c_{2h}^{n+1}, \mu_{2h}^{n+1}), (0, 0))$ est une solution du problème triphasique VI.7 alors $((c_{1h}^{n+1}, \tilde{\mu}_{1h}^{n+1}), (c_{2h}^{n+1}, \tilde{\mu}_{2h}^{n+1}))$ est une solution du problème diphasique VI.11.

VI.3 Stabilité inconditionnelle du schéma

Nous montrons dans cette section une égalité d'énergie qui assure la stabilité inconditionnelle du schéma.

Proposition VI.13 (Egalité d'énergie discrète)

Soient $\mathbf{c}_h^n \in \mathcal{V}_{h,S}^c$ et $\mathbf{u}_h^n \in \mathcal{V}_{h,0}^u$. Supposons qu'il existe une solution $(\mathbf{c}_h^{n+1}, \mu_h^{n+1}, \mathbf{u}_h^{n+1}, p_h^{n+1})$ au problème VI.7. Alors, nous avons l'égalité suivante :

$$\begin{aligned} & \left[\mathcal{F}_{\Sigma,\varepsilon}^{\text{triph}}(\mathbf{c}_h^{n+1}) + \frac{1}{2} \int_{\Omega} \varrho_h^{n+1} |\mathbf{u}_h^{n+1}|^2 dx \right] - \left[\mathcal{F}_{\Sigma,\varepsilon}^{\text{triph}}(\mathbf{c}_h^n) + \frac{1}{2} \int_{\Omega} \varrho_h^n |\mathbf{u}_h^n|^2 dx \right] \\ & + \Delta t \sum_{i=1}^3 \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx + \Delta t \int_{\Omega} 2\eta_h^{n+1} |D\mathbf{u}_h^{n+1}|^2 dx \\ & + \frac{3}{8} (2\beta - 1) \varepsilon \int_{\Omega} \sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 dx + \frac{1}{2} \int_{\Omega} \varrho_h^n [|\mathbf{u}_h^{n+1} - \mathbf{u}^*|^2 + |\mathbf{u}^* - \mathbf{u}_h^n|^2] dx \\ & = \Delta t \int_{\Omega} \varrho_h^{n+1} \mathbf{g} \cdot \mathbf{u}_h^{n+1} dx + \frac{12}{\varepsilon} \int_{\Omega} [F(\mathbf{c}_h^{n+1}) - F(\mathbf{c}_h^n) - \mathbf{d}^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \cdot (\mathbf{c}_h^{n+1} - \mathbf{c}_h^n)] dx, \end{aligned} \quad (\text{VI.26})$$

où $\mathbf{d}^F(\cdot, \cdot)$ est le vecteur $(d_i^F(\cdot, \cdot))_{i=1,2,3}$ et

$$\mathbf{u}^* = \mathbf{u}_h^n - \frac{\Delta t}{\varrho_h^n} \sum_{j=1}^3 (c_{jh}^n - \alpha_j) \nabla \mu_{jh}^{n+1}. \quad (\text{VI.27})$$

Démonstration : Le point clé de la démonstration est de constater que les systèmes de Cahn-Hilliard et Navier-Stokes peuvent se réécrire à l'aide de la fonction $\mathbf{u}^*$ définie par (VI.27). Ceci fait, les estimations standards sur les systèmes de Cahn-Hilliard et Navier-Stokes sont effectuées (étapes 1 et 3) et une estimation sur la norme L^2 de la fonction $\mathbf{u}^*$ permet de conclure (étape 2).

Etape 1 : En utilisant la fonction $\mathbf{u}^*$, on constate que le système VI.2 se réécrit :

$$\begin{cases} \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx - \int_{\Omega} [c_{ih}^n - \alpha_i] \mathbf{u}^* \cdot \nabla \nu_h^\mu dx = - \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} \nabla \mu_{ih}^{n+1} \cdot \nabla \nu_h^\mu dx, \\ \int_{\Omega} \mu_{ih}^{n+1} \nu_h^c dx = \int_{\Omega} D_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \nu_h^c dx + \int_{\Omega} \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^{n+\beta} \nabla \nu_h^c dx. \end{cases}$$

Prenons alors $\nu_h^\mu = \mu_{ih}^{n+1} - \mu_{ih}^n$ et $\nu_h^c = \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t}$ comme fonctions tests dans ce système. Quelques calculs similaires à ceux donnés dans la démonstration de la proposition V.7 nous permettent d'obtenir :

$$\begin{aligned} & \mathcal{F}_{\Sigma,\varepsilon}^{\text{triph}}(\mathbf{c}_h^{n+1}) - \mathcal{F}_{\Sigma,\varepsilon}^{\text{triph}}(\mathbf{c}_h^n) + \Delta t \sum_{i=1}^3 \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx \\ & + \frac{3}{8} (2\beta - 1) \varepsilon \int_{\Omega} \sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 dx = \Delta t \int_{\Omega} \mathbf{u}^* \cdot \sum_{i=1}^3 (c_{ih}^n - \alpha_i) \nabla \mu_{ih}^{n+1} dx \\ & + \frac{12}{\varepsilon} \int_{\Omega} [F(\mathbf{c}_h^{n+1}) - F(\mathbf{c}_h^n) - \mathbf{d}^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \cdot (\mathbf{c}_h^{n+1} - \mathbf{c}_h^n)] dx. \end{aligned} \quad (\text{VI.28})$$

Etape 2 : Il est alors possible d'obtenir une estimation du dernier terme du second membre de l'égalité ci-dessus. Par définition de $\mathbf{u}^*$, nous avons $\sqrt{\varrho_h^n} \mathbf{u}^* = \sqrt{\varrho_h^n} \mathbf{u}_h^n - \frac{\Delta t}{\sqrt{\varrho_h^n}} \sum_{j=1}^3 (c_{jh}^n - \alpha_j) \nabla \mu_{jh}^{n+1}$. En multipliant par la fonction $\sqrt{\varrho_h^n} \mathbf{u}^*$, et en intégrant sur Ω , il vient

$$\frac{1}{2} \int_{\Omega} \varrho_h^n |\mathbf{u}^*|^2 dx - \frac{1}{2} \int_{\Omega} \varrho_h^n |\mathbf{u}_h^n|^2 dx + \frac{1}{2} \int_{\Omega} \varrho_h^n |\mathbf{u}^* - \mathbf{u}_h^n|^2 dx = -\Delta t \int_{\Omega} \mathbf{u}^* \cdot \sum_{j=1}^3 (c_{jh}^n - \alpha_j) \nabla \mu_{jh}^{n+1} dx. \quad (\text{VI.29})$$

Etape 3 : Quant au système (VI.3), nous pouvons également le réécrire de la façon suivante grâce à la fonction $\mathbf{u}^*$:

$$\left\{ \begin{array}{l} \int_{\Omega} \varrho_h^n \frac{\mathbf{u}_h^{n+1} - \mathbf{u}^*}{\Delta t} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx + \frac{1}{2} \int_{\Omega} \mathbf{u}_h^{n+1} \frac{\varrho_h^{n+1} - \varrho_h^n}{\Delta t} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx \\ + \frac{1}{2} \int_{\Omega} \varrho_h^{n+1} (\mathbf{u}_h^n \cdot \nabla) \mathbf{u}_h^{n+1} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx - \frac{1}{2} \int_{\Omega} \varrho_h^{n+1} (\mathbf{u}_h^n \cdot \nabla) \boldsymbol{\nu}_h^{\mathbf{u}} \cdot \mathbf{u}_h^{n+1} dx \\ + \int_{\Omega} 2\eta_h^{n+1} D\mathbf{u}_h^{n+1} : D\boldsymbol{\nu}_h^{\mathbf{u}} dx - \int_{\Omega} p_h^{n+1} \operatorname{div}(\boldsymbol{\nu}_h^{\mathbf{u}}) dx = \int_{\Omega} \varrho_h^{n+1} \mathbf{g} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx, \\ \int_{\Omega} \nu_h^p \operatorname{div}(\mathbf{u}_h^{n+1}) dx = 0. \end{array} \right.$$

Prenons alors $\boldsymbol{\nu}_h^{\mathbf{u}} = \mathbf{u}_h^{n+1}$ et $\nu_h^p = p_h^{n+1}$ comme fonctions tests dans ce système. Il vient :

$$\begin{aligned} \frac{1}{2} \int_{\Omega} \varrho_h^{n+1} |\mathbf{u}_h^{n+1}|^2 dx - \frac{1}{2} \int_{\Omega} \varrho_h^n |\mathbf{u}^*|^2 dx + \frac{1}{2} \int_{\Omega} \varrho_h^n |\mathbf{u}_h^{n+1} - \mathbf{u}^*|^2 dx \\ + \Delta t \int_{\Omega} 2\eta_h^{n+1} |D\mathbf{u}_h^{n+1}|^2 dx = \Delta t \int_{\Omega} \varrho_h^{n+1} \mathbf{g} \cdot \mathbf{u}_h^{n+1} dx. \end{aligned} \quad (\text{VI.30})$$

Il ne reste plus qu'à sommer les équations (VI.28), (VI.29), (VI.30), pour obtenir la conclusion. $\blacksquare$

Remarque VI.14

Une différence importante avec les travaux de [KSW08] dans le cas du modèle Cahn-Hilliard/Navier-Stokes diphasique homogène est qu'aucune condition n'est requise sur le pas de temps pour obtenir la stabilité du schéma.

VI.4 Existence de solutions au problème discret

Nous montrons dans cette section l'existence de solutions au problème VI.7 discret non linéaire.

Théorème VI.15

- Etant donné $\mathbf{c}_h^n \in \mathcal{V}_S^c$, $\mathbf{u}_h^n \in \mathcal{V}_0^{\mathbf{u}}$, supposons que
- les coefficients $(\Sigma_1, \Sigma_2, \Sigma_3)$ satisfont (IV.14), la mobilité satisfait (IV.19), et que le potentiel de Cahn-Hilliard F satisfait (IV.20),
 - la discrétisation des termes non-linéaires $\mathbf{d}^F$ satisfait (V.20) et la propriété suivante : il existe $K_1^{\mathbf{c}_h^n} > 0$ (pouvant dépendre de $\mathbf{c}_h^n$) tel que

$$\int_{\Omega} [F(\mathbf{a}_h^{n+1}) - F(\mathbf{c}_h^n) - \mathbf{d}^F(\mathbf{c}_h^n, \mathbf{a}_h^{n+1}) \cdot (\mathbf{a}_h^{n+1} - \mathbf{c}_h^n)] dx \leq K_1^{\mathbf{c}_h^n}, \quad \forall \mathbf{a}_h^{n+1} \in \mathcal{V}_S^c. \quad (\text{VI.31})$$

Alors, il existe au moins une solution $(\mathbf{c}_h^{n+1}, \boldsymbol{\mu}_h^{n+1}, \mathbf{u}_h^{n+1}, p_h^{n+1})$ au problème VI.7.

A l'instar de la démonstration du théorème V.9 (existence de solutions au problème de Cahn-Hilliard discret) dans le chapitre V, la démonstration du théorème d'existence VI.15 repose sur le lemme suivant issu de la théorie de degré topologique [Dei85].

Lemme VI.16 (Degré topologique)

Soit W un espace vectoriel de dimension finie et G une fonction continue de W dans W . Supposons qu'il existe une fonction continue H de $W \times [0; 1]$ dans W satisfaisant

- (i) $H(\cdot, 1) = G$ et $H(\cdot, 0)$ est affine,
- (ii) $\exists R > 0$ tel que $\forall (w, \delta) \in W \times [0; 1]$, si $H(w, \delta) = 0$ alors $|w|_W < R$,
- (iii) l'équation $H(w, 0) = 0$ a une solution $w \in W$,

Alors il existe au moins une solution $w \in W$ telle que $G(w) = 0$ et $|w|_W < R$.

Le principe est de relier le problème discret non linéaire à un problème plus simple (linéaire) par homotopie (fonction H du lemme VI.16) pour lequel nous savons démontrer l'existence de solutions (hypothèse (ii) du lemme

VI.16). La théorie du degré topologique nous permet alors de déduire l'existence de solutions au problème non linéaire à partir d'estimations *a priori* essentiellement déduites de l'égalité d'énergie (VI.26) établie dans la proposition VI.13.

Démonstration du théorème VI.15 : Nous reformulons tout d'abord le problème VI.7 pour nous placer dans le cadre de l'énoncé du lemme VI.16, avant de valider une à une chacune des hypothèses (i), (ii) et (iii) de ce dernier.

Reformulation du problème

Soit W l'espace vectoriel de dimension finie $(\mathcal{V}^c)^2 \times (\mathcal{V}_h^\mu)^2 \times \mathcal{V}_{h,0}^u \times \mathcal{V}^p$. Nous définissons la norme suivante sur W :

$$\forall w = (c_{1h}, c_{2h}, \mu_{1h}, \mu_{2h}, \mathbf{u}_h, p_h) \in W, \\ |w|_W^2 = |c_{1h}|_{H^1(\Omega)}^2 + |c_{2h}|_{H^1(\Omega)}^2 + |\mu_{1h}|_{H^1(\Omega)}^2 + |\mu_{2h}|_{H^1(\Omega)}^2 + |\mathbf{u}_h|_{(H^1(\Omega))^d}^2 + |p_h|_{L^2(\Omega)}^2,$$

et nous introduisons la fonction H telle que

$$H : W \times [0; 1] \rightarrow W \\ (w^{n+1}, \delta) = (c_{1h}^{n+1}, c_{2h}^{n+1}, \mu_{1h}^{n+1}, \mu_{2h}^{n+1}, \mathbf{u}_h^{n+1}, p_h^{n+1}, \delta) \mapsto (\mathcal{R}_\delta^{\mu_1}, \mathcal{R}_\delta^{c_1}, \mathcal{R}_\delta^{\mu_2}, \mathcal{R}_\delta^{c_2}, \mathcal{R}_\delta^u, \mathcal{R}_\delta^p)$$

où $\mathcal{R}_\delta^{c_1}$ et $\mathcal{R}_\delta^{c_2}$, (resp. $\mathcal{R}_\delta^{\mu_1}$ et $\mathcal{R}_\delta^{\mu_2}$, resp. $\mathcal{R}_\delta^u$, resp. $\mathcal{R}_\delta^p$) sont définis par leurs coordonnées dans la base éléments finis $(\nu_I^c)_{I \in [1, \dim(\mathcal{V}_h^c)]}$ (resp. $(\nu_I^\mu)_{I \in [1, \dim(\mathcal{V}_h^\mu)]}$, resp. $(\nu_I^u)_{I \in [1, \dim(\mathcal{V}_{h,0}^u)]}$, resp. $(\nu_I^p)_{I \in [1, \dim(\mathcal{V}_h^p)]}$) de $\mathcal{V}_h^c$ (resp. $\mathcal{V}_h^\mu$, resp. $\mathcal{V}_{h,0}^u$, resp. $\mathcal{V}_h^p$) :

$$\forall I \in [1, \dim(\mathcal{V}_h^\mu)], \quad (\mathcal{R}_\delta^{\mu_i})_I = \int_\Omega \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_I^\mu dx - \delta \int_\Omega [c_{ih}^n - \alpha_{ih}] \left[\mathbf{u}_h^n - \frac{\Delta t}{\varrho_{h\delta}^n} \sum_{j=1}^3 (c_{jh}^n - \alpha_{jh}) \nabla \mu_{jh}^{n+1} \right] \cdot \nabla \nu_I^\mu dx \\ + \int_\Omega \frac{M_{0h\delta}^{n+\alpha}}{\Sigma_i} \nabla \mu_{ih}^{n+1} \cdot \nabla \nu_I^\mu dx, \\ \forall I \in [1, \dim(\mathcal{V}_h^c)], \quad (\mathcal{R}_\delta^{c_i})_I = \int_\Omega \mu_{ih}^{n+1} \nu_I^c dx - \int_\Omega \delta D_i(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \nu_I^c dx - \int_\Omega \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^{n+\beta} \cdot \nabla \nu_I^c dx, \\ \forall I \in [1, \dim(\mathcal{V}_{h,0}^u)], \quad (\mathcal{R}_\delta^u)_I = \int_\Omega \varrho_{h\delta}^n \frac{\mathbf{u}_h^{n+1} - \delta \mathbf{u}_h^n}{\Delta t} \cdot \nu_I^u dx + \frac{1}{2} \int_\Omega \frac{\varrho_{h\delta}^{n+1} - \varrho_{h\delta}^n}{\Delta t} \mathbf{u}_h^{n+1} \cdot \nu_I^u dx \\ + \frac{\delta}{2} \int_\Omega \varrho_h^{n+1} (\mathbf{u}_h^n \cdot \nabla) \mathbf{u}_h^{n+1} \cdot \nu_I^u dx - \frac{\delta}{2} \int_\Omega \varrho_h^{n+1} (\mathbf{u}_h^n \cdot \nabla) \nu_I^u \cdot \mathbf{u}_h^{n+1} dx \\ + \int_\Omega 2\eta_{h\delta}^{n+1} D\mathbf{u}_h^{n+1} : D\nu_I^u dx - \int_\Omega p_h^{n+1} \operatorname{div}(\nu_I^u) dx \\ - \int_\Omega \varrho_{h\delta}^{n+1} \mathbf{g} \cdot \nu_I^u dx + \delta \int_\Omega \sum_{j=1}^3 (c_{jh}^n - \alpha_{jh}) \nabla \mu_{jh}^{n+1} \cdot \nu_I^u dx, \\ \forall I \in [1, \dim(\mathcal{V}_h^p)], \quad (\mathcal{R}_\delta^p)_I = \int_\Omega \nu_I^p \operatorname{div}(\mathbf{u}_h^{n+1}) dx,$$

avec $M_{0h\delta}^{n+\alpha} = M_0((1 - \delta\alpha)\mathbf{c}_h^n + \delta\alpha\mathbf{c}_h^{n+1})$, $\varrho_{h\delta}^\ell = \varrho((1 - \delta)\mathbf{c}_h^{\ell-1} + \delta\mathbf{c}_h^\ell)$ pour $\ell = n$ ou $\ell = n + 1$ et $\eta_{h\delta}^{n+1} = \eta((1 - \delta)\mathbf{c}_h^n + \delta\mathbf{c}_h^{n+1})$. La fonction G est définie

$$G : W \rightarrow W \\ w \mapsto H(w, 1)$$

Le problème "Trouver w^{n+1} tel que $G(w^{n+1}) = 0$ " est équivalent (par définition de la fonction H) au problème VI.7. Pour démontrer le théorème, nous allons montrer que les fonctions H et G satisfont les hypothèses du lemme VI.16. La continuité de la fonction H est obtenue en utilisant la continuité des différentes fonctions non linéaires (D_i^F , ϱ et η) et le théorème de Lebesgue. La fonction $H(\cdot, 0)$ est clairement affine par construction.

Validation de l'hypothèse (ii) du lemme VI.16

Soit $(w^{n+1}, \delta) \in W \times [0; 1]$ tel que $H(w^{n+1}, \delta) = 0$. Remarquons que $H(w^{n+1}, \delta) = 0$ revient à dire que $w^{n+1} = (c_{1h}^{n+1}, c_{2h}^{n+1}, \mu_{1h}^{n+1}, \mu_{2h}^{n+1}, \mathbf{u}_h^{n+1}, p_h^{n+1})$ est solution d'un problème très similaire au problème VI.7. Ainsi

nous pouvons effectuer les mêmes calculs que dans la preuve de la proposition VI.13. En effet, en posant

$$\mathbf{u}_\delta^* = \delta \mathbf{u}_h^n - \delta \frac{\Delta t}{\varrho_{h\delta}^n} \sum_{j=1}^3 (c_{jh}^n - \alpha_{jh}) \nabla \mu_{jh}^{n+1},$$

l'égalité $H(w^{n+1}, \delta) = 0$ signifie exactement que nous avons

$$\left\{ \begin{array}{l} \forall \nu_h^\mu \in \mathcal{V}_h^\mu, \quad \int_\Omega \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx - \delta \int_\Omega (c_{ih}^n - \alpha_{ih}) \mathbf{u}_\delta^* \cdot \nabla \nu_h^\mu dx = - \int_\Omega \frac{M_{0h}^{n+\alpha}}{\Sigma_i} \nabla \mu_{ih}^{n+1} \cdot \nabla \nu_h^\mu dx, \\ \forall \nu_h^c \in \mathcal{V}_h^c, \quad \int_\Omega \mu_{ih}^{n+1} \nu_h^c dx = \int_\Omega \delta D_i(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \nu_h^c dx + \int_\Omega \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^{n+\beta} \cdot \nabla \nu_h^c dx, \\ \forall \nu_h^u \in \mathcal{V}_{h,0}^u, \quad \int_\Omega \varrho_{h\delta}^n \frac{\mathbf{u}_h^{n+1} - \mathbf{u}_\delta^*}{\Delta t} \cdot \nu_h^u dx + \frac{1}{2} \int_\Omega \frac{\varrho_{h\delta}^{n+1} - \varrho_{h\delta}^n}{\Delta t} \mathbf{u}_h^{n+1} \cdot \nu_h^u dx \\ \quad + \frac{\delta}{2} \int_\Omega \varrho_h^{n+1} (\mathbf{u}_h^n \cdot \nabla) \mathbf{u}_h^{n+1} \cdot \nu_h^u dx - \frac{\delta}{2} \int_\Omega \varrho_h^{n+1} (\mathbf{u}_h^n \cdot \nabla) \nu_h^u \cdot \mathbf{u}_h^{n+1} dx \\ \quad + \int_\Omega 2\eta_h^{n+1} D\mathbf{u}_h^{n+1} : D\nu_h^u dx - \int_\Omega p_h^{n+1} \operatorname{div}(\nu_h^u) dx = \int_\Omega \varrho_{h\delta}^{n+1} \mathbf{g} \cdot \nu_h^u dx, \\ \forall \nu_h^p \in \mathcal{V}_h^p, \quad \int_\Omega \nu_h^p \operatorname{div}(\mathbf{u}_h^{n+1}) dx = 0. \end{array} \right. \quad (\text{VI.32})$$

Nous prenons alors dans ce système les fonctions tests $\nu_h^\mu = \mu_{ih}^{n+1}$, $\nu_h^c = \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t}$, $\nu_h^u = \mathbf{u}_h^{n+1}$ et $\nu_h^p = p^{n+1}$ pour obtenir :

$$\begin{aligned} & \left[\mathcal{F}_{\Sigma, \varepsilon, \delta}^{\text{triph}}(\mathbf{c}_h^{n+1}) + \frac{1}{2} \int_\Omega \varrho_{h\delta}^{n+1} |\mathbf{u}_h^{n+1}|^2 dx \right] - \left[\mathcal{F}_{\Sigma, \varepsilon, \delta}^{\text{triph}}(\mathbf{c}_h^n) + \frac{1}{2} \int_\Omega \varrho_{h\delta}^n |\delta \mathbf{u}_h^n|^2 dx \right] \\ & + \Delta t \sum_{i=1}^3 \int_\Omega \frac{M_{0h}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx + \Delta t \int_\Omega 2\eta_h^{n+1} |D\mathbf{u}_h^{n+1}|^2 dx \\ & + \frac{3}{8} (2\beta - 1) \varepsilon \int_\Omega \sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 dx + \frac{1}{2} \int_\Omega \varrho_{h\delta}^n \left[|\mathbf{u}_h^{n+1} - \mathbf{u}_\delta^*|_{L^2(\Omega)}^2 + |\mathbf{u}_\delta^* - \delta \mathbf{u}_h^n|_{L^2(\Omega)}^2 \right] dx \\ & = \Delta t \int_\Omega \varrho_{h\delta}^{n+1} \mathbf{g} \cdot \mathbf{u}_h^{n+1} dx + \frac{12}{\varepsilon} \delta \int_\Omega \left[F(\mathbf{c}_h^{n+1}) - F(\mathbf{c}_h^n) - \mathbf{d}^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \cdot (\mathbf{c}_h^{n+1} - \mathbf{c}_h^n) \right] dx, \end{aligned}$$

où nous avons noté $\mathcal{F}_{\Sigma, \varepsilon, \delta}^{\text{triph}}(\mathbf{c}_h^\ell) = \int_\Omega \delta \frac{12}{\varepsilon} F(\mathbf{c}_h^\ell) + \sum_{i=1}^3 \frac{3}{8} \varepsilon \Sigma_i |\nabla c_{ih}^\ell|^2 dx$. En utilisant la remarque V.8, le fait que F est positive, les bornes inférieures $\varrho_{\min}$ et $\eta_{\min}$ (strictement positives) de la densité et la viscosité, le fait que la mobilité est minorée, le lemme de Korn (cf [BF06, lemme VII.3.5]) et l'hypothèse (VI.31), nous obtenons

$$\begin{aligned} & \frac{3}{8} \varepsilon \sum_{i=1}^3 |\nabla c_{ih}^{n+1}|_{(L^2(\Omega))^d}^2 + \frac{\varrho_{\min}}{2} |\mathbf{u}_h^{n+1}|_{(L^2(\Omega))^d}^2 \\ & + \frac{M_1 \underline{\Sigma} \Delta t}{\max_{i=1,2,3} (|\Sigma_i|)} \sum_{i=1}^3 |\nabla \mu_{ih}^{n+1}|_{(L^2(\Omega))^d}^2 + 2\Delta t \eta_{\min} C_k |\nabla \mathbf{u}_h^{n+1}|_{(L^2(\Omega))^d}^2 \\ & \leq \mathcal{F}_{\Sigma, \varepsilon, \delta}^{\text{triph}}(\mathbf{c}_h^n) + \frac{\varrho_{\max}}{2} |\delta \mathbf{u}_h^n|_{L^2(\Omega)}^2 + \Delta t \varrho_{\max} |g|_2 |\Omega|^{\frac{1}{2}} |\mathbf{u}_h^{n+1}|_{L^2(\Omega)^2} + \delta \frac{12}{\varepsilon} K_1^{\mathbf{c}_h^n}. \end{aligned} \quad (\text{VI.33})$$

Enfin en utilisant les inégalités de Poincaré et de Young, et puisque $\delta \leq 1$, il vient

$$\begin{aligned} & \frac{3}{8} \varepsilon \sum_{i=1}^3 |\nabla c_{ih}^{n+1}|_{(L^2(\Omega))^d}^2 + \frac{\varrho_{\min}}{2} |\mathbf{u}_h^{n+1}|_{(L^2(\Omega))^d}^2 \\ & + \frac{M_1 \underline{\Sigma} \Delta t}{\max_{i=1,2,3} (|\Sigma_i|)} \sum_{i=1}^3 |\nabla \mu_{ih}^{n+1}|_{(L^2(\Omega))^d}^2 + \Delta t \eta_{\min} C_k |\nabla \mathbf{u}_h^{n+1}|_{(L^2(\Omega))^d}^2 \\ & \leq \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^n) + \frac{\varrho_{\max}}{2} |\mathbf{u}_h^n|_{L^2(\Omega)}^2 + \frac{C_p^2 \Delta t \varrho_{\max}^2 |g|_2^2 |\Omega|}{4\eta_{\min} C_k} + \frac{12}{\varepsilon} K_1^{\mathbf{c}_h^n}. \end{aligned}$$

La constante $K_2^{\mathbf{c}_h^n} = \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^n) + \frac{\varrho_{\max}}{2} |\mathbf{u}_h^n|_{L^2(\Omega)}^2 + \frac{C_p^2 \Delta t \varrho_{\max}^2 |g|_2^2 |\Omega|}{4\eta_{\min} C_k} + \frac{12}{\varepsilon} K_1^{\mathbf{c}_h^n}$ est indépendante de δ et w^{n+1} et nous avons

$$\sum_{i=1}^3 |\nabla c_{ih}^{n+1}|_{(L^2(\Omega))^d}^2 \leq K_3^{\mathbf{c}_h^n}, \quad (\text{VI.34})$$

$$\sum_{i=1}^3 |\nabla \mu_{ih}^{n+1}|_{(L^2(\Omega))^d}^2 \leq K_4^{\mathbf{c}_h^n}, \quad (\text{VI.35})$$

$$|\mathbf{u}_h^{n+1}|_{(H^1(\Omega))^d} \leq K_5^{\mathbf{c}_h^n}, \quad (\text{VI.36})$$

avec $K_3^{\mathbf{c}_h^n} = \frac{8K_2^{\mathbf{c}_h^n}}{3\varepsilon \underline{\Sigma}}$, $K_4^{\mathbf{c}_h^n} = \frac{\max_{i=1,2,3} (|\Sigma_i|) K_2^{\mathbf{c}_h^n}}{M_1 \underline{\Sigma}}$, $K_5^{\mathbf{c}_h^n} = \max \left(\frac{2K_2^{\mathbf{c}_h^n}}{\varrho_{\min}}, \frac{K_2^{\mathbf{c}_h^n}}{\Delta t \eta_{\min} C_k} \right)^{\frac{1}{2}}$.

Nous utilisons maintenant la forme discrète de la conservation du volume : $m(c_{ih}^{n+1}) = m(c_{ih}^n)$ obtenue directement en choisissant $\nu_h^\mu \equiv 1$ dans le système (VI.32). Ainsi, grâce à l'inégalité de Poincaré (V.55) (avec $\theta \equiv 1$), il existe une constante positive $K_6^{\mathbf{c}_h^n} = C_p \left(K_3^{\mathbf{c}_h^n} + m(c_{ih}^n) \right)$ indépendante de δ et w^{n+1} telle que

$$|c_{ih}^{n+1}|_{H^1(\Omega)} \leq K_6^{\mathbf{c}_h^n}. \quad (\text{VI.37})$$

Pour estimer la moyenne $m(\mu_{ih}^{n+1})$, nous prenons $\nu_h^c \equiv 1$ dans le système d'équation (VI.32). Il vient

$$m(\mu_{ih}^{n+1}) = \int_{\Omega} \delta D_i^F(\mathbf{c}_h^{n+1}, \mathbf{c}_h^n) dx.$$

Ceci peut être contrôlé par $|\mathbf{c}_h^{n+1}|_{H^1(\Omega)}$ et $|\mathbf{c}_h^n|_{H^1(\Omega)}$ sous l'hypothèse (V.20). En effet, la croissance polynomiale (V.20) de d_i^F implique qu'il existe une constante positive $C_1 = \frac{16\Sigma_T}{3\Sigma_m} B_1$ telle que

$$|D_i^F(\mathbf{c}_h^{n+1}, \mathbf{c}_h^n)| \leq C_1 \left(1 + |\mathbf{c}_h^{n+1}|^{p-1} + |\mathbf{c}_h^n|^{p-1} \right).$$

Ainsi, puisque $\delta \leq 1$, et en utilisant (V.60), nous obtenons

$$m(\mu_{ih}^{n+1}) \leq C_1 \left(|\Omega| + |\mathbf{c}_h^{n+1}|_{L^{p-1}}^{p-1} + |\mathbf{c}_h^n|_{L^{p-1}}^{p-1} \right) \leq C_1 \left(|\Omega| + \left(K_6^{\mathbf{c}_h^n} \right)^{p-1} + |\mathbf{c}_h^n|_{H^1}^{p-1} \right) := K_7^{h, \mathbf{c}_h^n}.$$

Grâce à l'inégalité de Poincaré, il vient,

$$|\mu_{ih}^{n+1}|_{H^1(\Omega)} \leq C_p \left(K_4^{\mathbf{c}_h^n} + K_7^{h, \mathbf{c}_h^n} \right) := K_8^{\mathbf{c}_h^n}. \quad (\text{VI.38})$$

Enfin le contrôle sur la pression est obtenu en utilisant la borne sur la vitesse (VI.36) et la condition inf-sup (VI.10). Cette dernière assure (cf [BS08, 21.5.10, p. 344]) qu'il existe $\mathbf{v}_h \in \mathcal{V}_{h,0}^u$ tel que

$$\forall \nu_h^p \in \mathcal{V}^p, \quad \int_{\Omega} \nu_h^p \operatorname{div}(\mathbf{v}_h) dx = \int_{\Omega} \nu_h^p p_h^{n+1} dx \quad \text{et} \quad |\mathbf{v}_h|_{(H^1(\Omega))^d} \leq \frac{1}{\beta} |p_h^{n+1}|_{L^2(\Omega)}. \quad (\text{VI.39})$$

Ainsi en prenant, $\nu_h^u = \mathbf{v}_h$ dans le système (VI.32), il vient :

$$\begin{aligned} |p_h^{n+1}|_{L^2(\Omega)}^2 &= \int_{\Omega} \frac{\frac{\varrho_{h\delta}^n + \varrho_{h\delta}^{n+1}}{2} \mathbf{u}_h^{n+1} - \varrho_{h\delta}^n \delta \mathbf{u}_h^n}{\Delta t} \cdot \mathbf{v}_h dx + \int_{\Omega} 2\eta_{h\delta}^{n+1} D\mathbf{u}_h^{n+1} : D\mathbf{v}_h dx \\ &\quad + \frac{\delta}{2} \int_{\Omega} \varrho_h^{n+1} (\mathbf{u}_h^n \cdot \nabla) \mathbf{u}_h^{n+1} \cdot \mathbf{v}_h dx - \frac{\delta}{2} \int_{\Omega} \varrho_h^{n+1} (\mathbf{u}_h^n \cdot \nabla) \mathbf{v}_h \cdot \mathbf{u}_h^{n+1} dx \\ &\quad - \int_{\Omega} \varrho_{h\delta}^{n+1} \mathbf{g} \cdot \mathbf{v}_h dx + \delta \int_{\Omega} \sum_{j=1}^3 (c_{jh}^n - \alpha_{jh}) \nabla \mu_{jh}^{n+1} \cdot \mathbf{v}_h dx. \end{aligned}$$

Nous utilisons alors l'inégalité de Cauchy-Schwarz, les bornes supérieures $\varrho_{\max}$ et $\eta_{\max}$ de la densité et de la viscosité ainsi que les estimations (VI.36), (VI.37), (VI.38) et (VI.39) pour obtenir :

$$\begin{aligned}
|p_h^{n+1}|_{L^2(\Omega)}^2 &\leq \frac{\varrho_{\max}}{\Delta t} \left[|\mathbf{u}_h^n|_{(L^2(\Omega))^d} + |\mathbf{u}_h^{n+1}|_{(L^2(\Omega))^d} \right] |\mathbf{v}_h|_{(L^2(\Omega))^d} + 2\eta_{\max} |\nabla \mathbf{u}_h^{n+1}|_{(L^2(\Omega))^d} |\nabla \mathbf{v}_h|_{(L^2(\Omega))^d} \\
&\quad + \frac{\delta}{2} \varrho_{\max} |\mathbf{u}_h^n|_{(L^4(\Omega))^d} |\nabla \mathbf{u}_h^{n+1}|_{(L^2(\Omega))^d} |\mathbf{v}_h|_{(L^4(\Omega))^d} + \frac{\delta}{2} \varrho_{\max} |\mathbf{u}_h^n|_{(L^4(\Omega))^d} |\nabla \mathbf{v}_h|_{(L^2(\Omega))^d} |\mathbf{u}_h^{n+1}|_{(L^4(\Omega))^d} \\
&\quad + \varrho_{\max} |\Omega|^{\frac{1}{2}} |\mathbf{g}|_2 |\mathbf{v}|_{(L^2(\Omega))^d} + \delta \sum_{j=1}^3 |c_{jh}^n - \alpha_{jh}|_{(L^4(\Omega))^d} |\nabla \mu_{jh}^{n+1}|_{(L^2(\Omega))^d} |\mathbf{v}_h|_{(L^4(\Omega))^d} \\
&\leq \left[\frac{\varrho_{\max}}{\Delta t} \left[|\mathbf{u}_h^n|_{(L^2(\Omega))^d} + K_5^{\mathbf{c}_h^n} \right] + 2\eta_{\max} K_5^{\mathbf{c}_h^n} \right. \\
&\quad \left. + \varrho_{\max} |\mathbf{u}_h^n|_{(L^4(\Omega))^d} K_5^{\mathbf{c}_h^n} + \varrho_{\max} |\Omega|^{\frac{1}{2}} |\mathbf{g}|_2 + 3K_6^{\mathbf{c}_h^n} K_8^{\mathbf{c}_h^n} \right] |\mathbf{v}_h|_{(H^1(\Omega))^d} \\
&\leq \frac{1}{\beta} \left[\frac{\varrho_{\max}}{\Delta t} \left[|\mathbf{u}_h^n|_{(L^2(\Omega))^d} + K_5^{\mathbf{c}_h^n} \right] + 2\eta_{\max} K_5^{\mathbf{c}_h^n} \right. \\
&\quad \left. + \varrho_{\max} |\mathbf{u}_h^n|_{(L^4(\Omega))^d} K_5^{\mathbf{c}_h^n} + \varrho_{\max} |\Omega|^{\frac{1}{2}} |\mathbf{g}|_2 + 3K_6^{\mathbf{c}_h^n} K_8^{\mathbf{c}_h^n} \right] |p_h^{n+1}|_{L^2(\Omega)}.
\end{aligned}$$

En conclusion,

$$|p_h^{n+1}|_{L^2(\Omega)} \leq K_9^{\mathbf{c}_h^n}, \quad (\text{VI.40})$$

avec $K_9^{\mathbf{c}_h^n} = \frac{\varrho_{\max}}{\Delta t} \left[|\mathbf{u}_h^n|_{(L^2(\Omega))^d} + K_5^{\mathbf{c}_h^n} \right] + 2\eta_{\max} K_5^{\mathbf{c}_h^n} + \varrho_{\max} |\mathbf{u}_h^n|_{(L^4(\Omega))^d} K_5^{\mathbf{c}_h^n} + \varrho_{\max} |\Omega|^{\frac{1}{2}} |\mathbf{g}|_2 + 3K_6^{\mathbf{c}_h^n} K_8^{\mathbf{c}_h^n}$.

Ainsi, en combinant (VI.37), (VI.38), (VI.36) et (VI.40) nous obtenons une constante positive $K^{\mathbf{c}_h^n}$ indépendante de δ et $\mathbf{c}_h^{n+1}$ telle que

$$|w^{n+1}|_W \leq K^{\mathbf{c}_h^n}.$$

Donc, prendre $R > K^{\mathbf{c}_h^n} \geq 0$ garantit que pour tout $(w, \delta) \in W \times [0; 1]$, $H(w, \delta) = 0 \implies |w|_W < R$.

Validation de l'hypothèse (iii) du lemme VI.16

Nous devons montrer l'existence d'une solution au problème linéaire $H(w^{n+1}, 0) = 0$. Ce problème s'écrit sous la forme de trois problèmes totalement découplés :

(1-2) Trouver $(c_{ih}^{n+1}, c_{2h}^{n+1}, \mu_{1h}^{n+1}, \mu_{2h}^{n+1}) \in (\mathcal{V}_h^c)^2 \times (\mathcal{V}_h^\mu)^2$ tels que $\forall i = 1, 2, \forall \nu_h^\mu \in \mathcal{V}_h^\mu, \forall \nu_h^c \in \mathcal{V}_h^c$,

$$a_i((c_{ih}^{n+1}, \mu_{ih}^{n+1}), (\nu_h^c, \nu_h^\mu)) = \int_{\Omega} c_{ih}^n \nu_h^\mu dx,$$

où

$$a_i((c_{ih}^{n+1}, \mu_{ih}^{n+1}), (\nu_h^c, \nu_h^\mu)) = \int_{\Omega} \left[c_{ih}^{n+1} \nu_h^\mu + \frac{M_{0h}^n}{\Sigma_i} \Delta t \nabla \mu_{ih}^{n+1} \cdot \nabla \nu_h^\mu \right] dx + \int_{\Omega} \left[\frac{3}{4} \Sigma_i \varepsilon \beta \nabla c_{ih}^{n+1} \cdot \nabla \nu_h^c - \mu_{ih}^{n+1} \nu_h^c \right] dx,$$

avec $M_{0h}^n = M_0(\mathbf{c}_h^n)$.

(3) Trouver $(\mathbf{u}_h^{n+1}, p_h^{n+1}) \in \mathcal{V}_{h,0}^u \times \mathcal{V}_h^p$ tels que

$$\begin{aligned}
\forall \nu_h^u \in \mathcal{V}_{h,0}^u, \quad &\int_{\Omega} \frac{\varrho_h^{n-1} + \varrho_h^n}{2\Delta t} \mathbf{u}_h^{n+1} \cdot \nu_h^u dx + \int_{\Omega} 2\eta_h^n D\mathbf{u}_h^{n+1} : D\nu_h^u dx - \int_{\Omega} p_h^{n+1} \operatorname{div}(\nu_h^u) dx = \int_{\Omega} \varrho_h^n \mathbf{g} \cdot \nu_h^u dx, \\
\forall \nu_h^p \in \mathcal{V}_h^p, \quad &\int_{\Omega} \nu_h^p \operatorname{div}(\mathbf{u}_h^{n+1}) dx = 0.
\end{aligned}$$

Puisque les problèmes linéaires (1-2) sont posés en dimension finie, il est suffisant de montrer que, pour tout $(c_{ih}^{n+1}, \mu_{ih}^{n+1}) \in \mathcal{V}_h^c \times \mathcal{V}_h^\mu$:

$$(a_i((c_{ih}^{n+1}, \mu_{ih}^{n+1}), (\nu_h^c, \nu_h^\mu)) = 0, \quad \forall (\nu_h^c, \nu_h^\mu) \in \mathcal{V}_h^c \times \mathcal{V}_h^\mu) \implies (c_{ih}^{n+1}, \mu_{ih}^{n+1}) = (0, 0).$$

Soit $(c_{ih}^{n+1}, \mu_{ih}^{n+1}) \in \mathcal{V}^c \times \mathcal{V}_h^\mu$ tel que

$$(a_i((c_{ih}^{n+1}, \mu_{ih}^{n+1}), (\nu_h^c, \nu_h^\mu)) = 0, \quad \forall (\nu_h^c, \nu_h^\mu) \in \mathcal{V}_h^c \times \mathcal{V}_h^\mu), \quad (\text{VI.41})$$

Prenons $(\nu_h^c, \nu_h^\mu) = (c_{ih}^{n+1}, \mu_{ih}^{n+1})$ dans (VI.41), nous obtenons :

$$\int_{\Omega} c_{ih}^{n+1} \mu_{ih}^{n+1} dx + \int_{\Omega} \frac{M_{0h}^n}{\Sigma_i} \Delta t |\nabla \mu_{ih}^{n+1}|^2 dx + \frac{3}{4} \Sigma_i \varepsilon \beta \int_{\Omega} |\nabla c_{ih}^{n+1}|^2 dx - \int_{\Omega} \mu_{ih}^{n+1} c_{ih}^{n+1} dx = 0.$$

Ceci est équivalent à

$$\int_{\Omega} M_{0h}^n \Delta t |\nabla \mu_{ih}^{n+1}|^2 dx + \frac{3}{4} \Sigma_i \varepsilon \beta \int_{\Omega} |\nabla c_{ih}^{n+1}|^2 dx = 0.$$

Puisque la mobilité satisfait (IV.19), nous obtenons : $\nabla \mu_{ih}^{n+1} = \nabla c_{ih}^{n+1} = 0$. Donc, c_{ih}^{n+1} et μ_{ih}^{n+1} sont constant. En réinjectant ces constantes dans (VI.41), nous obtenons

$$(c_{ih}^{n+1}, \mu_{ih}^{n+1}) = (0, 0).$$

Le problème (3) admet une unique solution. En effet, les bornes inférieures sur la densité et la viscosité, le lemme de Korn (cf [BF06, lemme VII.3.5]) permettent de montrer que la forme bilinéaire continue

$$(\mathbf{u}, \mathbf{v}) \in \mathcal{V}_{h,0}^{\mathbf{u}} \rightarrow \int_{\Omega} \frac{\varrho_h^{n-1} + \varrho_h^n}{2\Delta t} \mathbf{u} \cdot \nu_h^{\mathbf{u}} dx + \int_{\Omega} 2\eta_h^n D\mathbf{u}_h^{n+1} : D\nu_h^{\mathbf{u}} dx$$

est coercive. La condition inf-sup (VI.10) permet alors de conclure. ■

VI.5 Convergence des solutions discrètes dans le cas homogène

Dans cette section, nous supposons que $\varrho_1 = \varrho_2 = \varrho_3 = \varrho_0 > 0$. Ceci a pour conséquence que la fonction $\varrho(\mathbf{c})$ est une fonction constante :

$$\varrho(\mathbf{c}) = \varrho_0, \quad \forall \mathbf{c} \in \mathcal{S}.$$

Dans ce cas particulier, le problème VI.7 s'écrit de la manière suivante :

Problème VI.17

– Etape 1 : résolution du système de Cahn-Hilliard

Trouver $(\mathbf{c}_h^{n+1}, \mu_h^{n+1}) \in \mathcal{V}_{h,\mathcal{S}}^c \times (\mathcal{V}_h^\mu)^3$ tel que $\forall \nu_h^c \in \mathcal{V}_{h,\mathcal{S}}^c, \forall \nu_h^\mu \in \mathcal{V}_h^\mu$, nous avons, pour $i = 1, 2$ et 3 ,

$$\left\{ \begin{array}{l} \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx - \int_{\Omega} [c_{ih}^n - \alpha_i] \left[\mathbf{u}_h^n - \frac{\Delta t}{\varrho_0} \sum_{j=1}^3 (c_{jh}^n - \alpha_j) \nabla \mu_{jh}^{n+1} \right] \cdot \nabla \nu_h^\mu dx \\ \hspace{15em} = - \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} \nabla \mu_{ih}^{n+1} \cdot \nabla \nu_h^\mu dx, \\ \int_{\Omega} \mu_{ih}^{n+1} \nu_h^c dx = \int_{\Omega} D_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \nu_h^c dx + \int_{\Omega} \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^{n+\beta} \nabla \nu_h^c dx, \end{array} \right. \quad (\text{VI.42})$$

avec α_j une constante : $\alpha_j = \int_{\Omega} c_{jh}^0 dx$.

– Etape 2 : résolution des équations de Navier-Stokes

Trouver $(\mathbf{u}_h^{n+1}, p_h^{n+1}) \in \mathcal{V}_{h,0}^{\mathbf{u}} \times \mathcal{V}_h^p$ tels que $\forall \nu_h^{\mathbf{u}} \in \mathcal{V}_{h,0}^{\mathbf{u}}, \forall \nu_h^p \in \mathcal{V}_h^p$,

$$\left\{ \begin{array}{l} \int_{\Omega} \varrho_0 \frac{\mathbf{u}_h^{n+1} - \mathbf{u}_h^n}{\Delta t} \cdot \nu_h^{\mathbf{u}} dx + \frac{1}{2} \int_{\Omega} \varrho_0 (\mathbf{u}_h^n \cdot \nabla) \mathbf{u}_h^{n+1} \cdot \nu_h^{\mathbf{u}} dx - \frac{1}{2} \int_{\Omega} \varrho_0 (\mathbf{u}_h^n \cdot \nabla) \nu_h^{\mathbf{u}} \cdot \mathbf{u}_h^{n+1} dx \\ \hspace{10em} + \int_{\Omega} 2\eta_h^{n+1} D\mathbf{u}_h^{n+1} : D\nu_h^{\mathbf{u}} dx - \int_{\Omega} p_h^{n+1} \operatorname{div}(\nu_h^{\mathbf{u}}) dx \\ \hspace{15em} = \int_{\Omega} \varrho_0 \mathbf{g} \cdot \nu_h^{\mathbf{u}} dx - \int_{\Omega} \sum_{j=1}^3 (c_{jh}^n - \alpha_j) \nabla \mu_{jh}^{n+1} \cdot \nu_h^{\mathbf{u}} dx, \\ \int_{\Omega} \nu_h^p \operatorname{div}(\mathbf{u}_h^{n+1}) dx = 0, \end{array} \right. \quad (\text{VI.43})$$

$$\text{où } \eta_h^{n+1} = \eta(\mathbf{c}_h^{n+1}).$$

L'existence de solutions à ce problème est donnée par le théorème VI.15.

Pour tout $N \in \mathbb{N}$, nous pouvons introduire les fonctions du temps $t \in [0, t_f]$ suivantes :

$$\underline{c}_{ih}^N(t, \cdot) = c_{ih}^n(\cdot), \quad \text{si } t \in]t_n, t_{n+1}[, \quad (\text{VI.44})$$

$$\bar{c}_{ih}^N(t, \cdot) = c_{ih}^{n+1}(\cdot), \quad \text{si } t \in]t_n, t_{n+1}[, \quad (\text{VI.45})$$

$$c_{ih}^N(t, \cdot) = \frac{t_{n+1} - t}{\Delta t} c_{ih}^n(\cdot) + \frac{t - t_n}{\Delta t} c_{ih}^{n+1}(\cdot), \quad \text{si } t \in]t_n, t_{n+1}[. \quad (\text{VI.46})$$

Pour les potentiels chimiques, pour tout $N \in \mathbb{N}$, nous introduisons les fonctions constantes par morceaux en temps suivantes :

$$\mu_{ih}^N(t, \cdot) = \mu_{ih}^{n+1}(\cdot), \quad \text{si } t \in]t_n, t_{n+1}[. \quad (\text{VI.47})$$

Et enfin pour la vitesse, nous introduisons les fonctions du temps $t \in [0, t_f]$ suivantes pour tout $N \in \mathbb{N}$:

$$\underline{\mathbf{u}}_h^N(t, \cdot) = \mathbf{u}_h^n(\cdot), \quad \text{si } t \in]t_n, t_{n+1}[, \quad (\text{VI.48})$$

$$\bar{\mathbf{u}}_h^N(t, \cdot) = \mathbf{u}_h^{n+1}(\cdot), \quad \text{si } t \in]t_n, t_{n+1}[, \quad (\text{VI.49})$$

$$\mathbf{u}_h^N(t, \cdot) = \frac{t_{n+1} - t}{\Delta t} \mathbf{u}_h^n(\cdot) + \frac{t - t_n}{\Delta t} \mathbf{u}_h^{n+1}(\cdot), \quad \text{si } t \in]t_n, t_{n+1}[. \quad (\text{VI.50})$$

Le résultat de convergence s'énonce de la manière suivante :

Théorème VI.18 (Théorème de convergence)

Nous supposons que les hypothèses du théorème VI.15 sont satisfaites, de manière qu'une solution $(\mathbf{c}_h^N, \boldsymbol{\mu}_h^N, \mathbf{u}_h^N, p_h^N)$ au problème VI.7 existe pour tout $N \in \mathbb{N}^*$ et pour tout $h > 0$. Nous supposons que $\beta \in]\frac{1}{2}, 1]$, que la propriété de consistance (V.3) est vérifiée et qu'il existe deux constantes $C > 0$ et $\Delta t_0 > 0$ telles que pour tout $\Delta t \leq \Delta t_0$ et pour tout $n \in \llbracket 0, N-1 \rrbracket$,

$$\begin{aligned} & \left[\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^{n+1}) + \frac{1}{2} \varrho_0 \int_{\Omega} |\mathbf{u}_h^{n+1}|^2 dx \right] - \left[\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^n) + \frac{1}{2} \varrho_0 \int_{\Omega} |\mathbf{u}_h^n|^2 dx \right] \\ & + C \left[\Delta t \sum_{i=1}^3 \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx + \frac{3}{8} (2\beta - 1) \varepsilon \int_{\Omega} \sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 dx \right] \\ & + \Delta t \int_{\Omega} 2\eta_h^{n+1} |D\mathbf{u}_h^{n+1}|^2 dx + \frac{1}{4} \varrho_0 \int_{\Omega} |\mathbf{u}_h^{n+1} - \mathbf{u}_h^n|^2 dx \leq \Delta t \varrho_0 \int_{\Omega} \mathbf{g} \cdot \mathbf{u}_h^{n+1} dx. \end{aligned} \quad (\text{VI.51})$$

Considérons le problème (IV.28), les conditions initiales (IV.32) et les conditions aux bords (IV.30)-(IV.31). Alors, il existe une solution faible $(\mathbf{c}, \boldsymbol{\mu}, \mathbf{u}, p)$ définie sur $[0, t_f]$ telle que

$$\mathbf{c} \in L^\infty(0, t_f; (H^1(\Omega))^3) \cap C^0([0, t_f]; (L^q(\Omega))^3), \quad \text{pour tout } q < 6$$

$$\boldsymbol{\mu} \in L^2(0, t_f; (H^1(\Omega))^3),$$

$$\mathbf{u} \in L^\infty(0, t_f; (L^2(\Omega))^d) \cap L^2(0, t_f; (H^1(\Omega))^d),$$

$$\mathbf{c}(t, x) \in \mathcal{S}, \quad \text{pour presque tout } (t, x) \in [0, t_f] \times \Omega.$$

De plus, pour toutes suites $(h_K)_{K \in \mathbb{N}^*}$ et $(N_K)_{K \in \mathbb{N}^*}$ vérifiant les propriétés suivantes :

$$\circ h_K \xrightarrow{K \rightarrow +\infty} 0 \quad \text{et} \quad N_K \xrightarrow{K \rightarrow +\infty} +\infty,$$

$$\circ \text{il existe une constante } A \text{ (indépendante de } K) \text{ telle que } \forall K \in \mathbb{N}^*, C_{\text{inv}}(h_K) \leq AN_K, \quad (\text{VI.52})$$

(rappelons que la fonction C_{inv} est définie par (VI.12)),

les suites $(\mathbf{c}_{h_K}^{N_K})_{K \in \mathbb{N}^*}$, $(\boldsymbol{\mu}_{h_K}^{N_K})_{K \in \mathbb{N}^*}$ et $(\mathbf{u}_{h_K}^{N_K})_{K \in \mathbb{N}^*}$ satisfont, à sous-suite près, les convergences suivantes lorsque $K \rightarrow +\infty$:

$$\mathbf{c}_{h_K}^{N_K} \rightarrow \mathbf{c} \quad \text{dans } C^0(0, t_f, (L^q)^3) \text{ fort}, \quad \text{pour tout } q < 6, \quad (\text{VI.53})$$

$$\mathbf{u}_{h_K}^{N_K} \rightarrow \mathbf{u} \quad \text{dans } L^2(0, t_f, (L^2)^d) \text{ fort}, \quad (\text{VI.54})$$

$$\boldsymbol{\mu}_{h_K}^{N_K} \rightharpoonup \boldsymbol{\mu} \quad \text{dans } L^2(0, t_f, (H^1)^3) \text{ faible}. \quad (\text{VI.55})$$

Remarque VI.19

L'hypothèse (VI.51) est obtenue dans la pratique par l'application de la proposition VI.13 et le contrôle du terme

$$\int_{\Omega} [F(\mathbf{c}_h^{n+1}) - F(\mathbf{c}_h^n) - \mathbf{d}^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \cdot (\mathbf{c}_h^{n+1} - \mathbf{c}_h^n)] dx,$$

du second membre de (VI.26). La manière d'obtenir ce contrôle diffère selon le choix du schéma $D_i^F(\mathbf{c}^n, \mathbf{c}^{n+1})$ de discrétisation des termes non linéaires du système de Cahn-Hilliard. Ceci a été discuté en détail dans le chapitre V. Ainsi, pour permettre son application aux différents schémas présentés dans le chapitre V, nous énonçons le théorème de convergence en supposant directement (VI.51).

Remarque VI.20

Dans l'énoncé du théorème VI.18, l'inégalité (VI.52) n'est pas une condition de stabilité (de type CFL par exemple), elle indique simplement que, pour obtenir la convergence vers la solution faible du problème continu, il faut que le pas de temps tende plus vite vers 0 que le pas d'espace. Ce résultat est donc différent du résultat de convergence établi dans le théorème V.10 concernant le système de Cahn-Hilliard pour lequel pas de temps et pas d'espace peuvent tendre vers 0 indépendamment l'un de l'autre.

La démonstration du théorème VI.18 s'inspire des travaux de [KSW08] (ainsi que de [Fen06]) effectués dans le cas du système de Cahn-Hilliard/Navier-Stokes diphasique. Mis à part le fait que nous traitons d'un modèle trois phases, la différence essentielle avec ces travaux réside dans la prise en compte du terme de transport de l'équation de Cahn-Hilliard faisant l'originalité de notre schéma. Il faut montrer que le terme supplémentaire ne nuit pas à la consistance. Ceci est vrai à condition que le pas de temps tende plus vite vers 0 que le pas d'espace, condition moins restrictive que la condition de stabilité introduite dans [KSW08] (cf remarque VI.20).

Classiquement, la démonstration du théorème VI.18 se déroule en trois grandes étapes :

- tout d'abord, l'égalité d'énergie (VI.51) permet de montrer que les suites $(\mathbf{c}_{h_K}^{N_K})_{K \in \mathbb{N}^*}$, $(\boldsymbol{\mu}_{h_K}^{N_K})_{K \in \mathbb{N}^*}$ et $(\mathbf{u}_{h_K}^{N_K})_{K \in \mathbb{N}^*}$ sont bornées pour certaines normes (précisées dans la suite).
- il est alors possible d'appliquer des théorèmes de compacité pour extraire de ces différentes suites des sous-suites convergentes.
- la troisième étape consiste à montrer que la limite obtenue est bien solution du problème (IV.28).

Nous détaillons séparément chacune de ces trois sous-étapes dans les sections ci-après.

Nous supposons donc dans ce qui suit (section VI.5.1, VI.5.2 et VI.5.3) que les hypothèses du théorème VI.18 sont satisfaites et qu'en particulier les notations $\mathbf{c}_h^n$, $\boldsymbol{\mu}_h^n$, $\mathbf{u}_h^n$, $p_h^n \dots$ désignent des solutions du problème discret VI.7, et $(\mathbf{c}_{h_K}^{N_K})_{K \in \mathbb{N}^*}$, $(\boldsymbol{\mu}_{h_K}^{N_K})_{K \in \mathbb{N}^*}$, $(\mathbf{u}_{h_K}^{N_K})_{K \in \mathbb{N}^*} \dots$ les suites de fonctions associées.

VI.5.1 Bornes sur les solutions discrètes

Dans cette section, nous supposons que K est fixé et pour simplifier les écritures nous omettons l'indice K dans la notation h_K et N_K .

Les premières estimations énoncées dans la proposition ci-après sont dérivées directement de l'estimation d'énergie (VI.51).

Proposition VI.21

Nous avons les inégalités suivantes :

$$\sup_{n \leq N} (|\mathbf{c}_h^n|_{(H^1(\Omega))^3}) + \sup_{n \leq N} (|\mathbf{u}_h^n|_{(L^2(\Omega))^d}) \leq K_1, \quad (\text{VI.56})$$

$$\left(\sum_{n=0}^{N-1} \Delta t \sum_{i=1}^3 |\mu_{ih}^{n+1}|_{H^1(\Omega)}^2 \right) + \left(\sum_{n=0}^{N-1} \Delta t |\mathbf{u}_h^{n+1}|_{(H^1(\Omega))^d}^2 \right) \leq K_2, \quad (\text{VI.57})$$

$$\Delta t \left(\sum_{n=0}^{N-1} \Delta t \sum_{i=1}^3 \left| \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \right|_{H^1(\Omega)}^2 \right) + \left(\sum_{n=0}^{N-1} |\mathbf{u}_h^{n+1} - \mathbf{u}_h^n|_{(L^2(\Omega))^d}^2 \right) \leq K_3, \quad (\text{VI.58})$$

où K_1 , K_2 et K_3 sont des constantes indépendantes de Δt et h .

Démonstration : Cette preuve est assez similaire à celle de la proposition V.25, elle utilise néanmoins quelques ingrédients supplémentaires (le lemme de Korn (cf [BF06, lemme VII.3.5]), la borne inférieure sur la

viscosité $\eta(c)$ et bien sûr le fait que la densité est constante) pour traiter les termes de vitesse qui n'étaient pas présents au chapitre V.

Nous notons $\Sigma_m = \min_{i=1,2,3} |\Sigma_i|$ et $\Sigma_M = \max_{i=1,2,3} |\Sigma_i|$.

(i) Tout d'abord, l'inégalité (VI.51) implique que

$$\forall n \in \llbracket 0, N-1 \rrbracket, \quad \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^{n+1}) + \frac{1}{2} \varrho_0 \int_{\Omega} |\mathbf{u}_h^{n+1}|^2 dx \leq \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^n) + \frac{1}{2} \varrho_0 \int_{\Omega} |\mathbf{u}_h^n|^2 dx.$$

Et ainsi, nous avons

$$\forall n \in \llbracket 0, N \rrbracket, \quad \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^n) + \frac{1}{2} \varrho_0 \int_{\Omega} |\mathbf{u}_h^n|^2 dx \leq \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^0) + \frac{1}{2} \varrho_0 \int_{\Omega} |\mathbf{u}_h^0|^2 dx. \quad (\text{VI.59})$$

De plus, grâce à l'hypothèse de croissance polynomiale (IV.20) de F et aux définitions de $\mathbf{c}_h^0$ et $\mathbf{u}_h^0$, l'énergie initiale peut être majorée indépendamment de h :

$$\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^0) + \frac{1}{2} \int_{\Omega} \varrho_0 |\mathbf{u}_h^0|^2 dx \leq B_1 \left(|\Omega| + |\mathbf{c}^0|_{H^1}^p \right) + \Sigma_M |\mathbf{c}^0|_{H^1}^2 + \frac{1}{2} \varrho_0 |\mathbf{u}^0|_{H^1(\Omega)}^2 := K_0.$$

Du fait que F est une fonction positive et en utilisant la proposition IV.5, l'inégalité (VI.59) nous permet alors de déduire que

$$\forall n \in \llbracket 0, N \rrbracket, \quad \frac{3}{8} \varepsilon \Sigma |\nabla \mathbf{c}_h^n|_{(L^2(\Omega))^3} + |\mathbf{u}_h^n|_{(L^2(\Omega))^d}^2 \leq K_0.$$

L'inégalité (VI.56) s'obtient alors en utilisant la conservation discrète du volume et l'inégalité de Poincaré moyenne (V.55).

(ii) Nous obtenons (VI.57) et (VI.58) en sommant les équations (VI.51) pour n allant de 0 à $N-1$:

$$\begin{aligned} & \left[\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^{n=N}) + \frac{1}{2} \int_{\Omega} \varrho_0 |\mathbf{u}_h^{n=N}|^2 dx \right] - \left[\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^0) + \frac{1}{2} \int_{\Omega} \varrho_0 |\mathbf{u}_h^0|^2 dx \right] \\ & + C \left[\sum_{n=0}^{N-1} \Delta t \sum_{i=1}^3 \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx + \frac{3}{8} (2\beta - 1) \varepsilon \sum_{n=0}^{N-1} \int_{\Omega} \sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 dx \right] \\ & + \sum_{n=0}^{N-1} \Delta t \int_{\Omega} 2\eta_h^{n+1} |D\mathbf{u}_h^{n+1}|^2 dx + \frac{1}{4} \sum_{n=0}^{N-1} \int_{\Omega} \varrho_0 |\mathbf{u}_h^{n+1} - \mathbf{u}_h^n|^2 dx \leq \sum_{n=0}^{N-1} \Delta t \int_{\Omega} \varrho_0 \mathbf{g} \cdot \mathbf{u}_h^{n+1} dx. \end{aligned}$$

Puisque l'énergie discrète est positive, en utilisant la proposition IV.5, les bornes inférieures de la mobilité et de la viscosité, nous obtenons :

$$\begin{aligned} & C \left[\frac{M_1 \underline{\Sigma}}{(\Sigma_M)^2} \sum_{n=0}^{N-1} \Delta t \sum_{i=1}^3 |\nabla \mu_{ih}^{n+1}|^2 + \frac{3}{8} (2\beta - 1) \varepsilon \underline{\Sigma} \sum_{n=0}^{N-1} \int_{\Omega} \sum_{i=1}^3 |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 dx \right] \\ & + 2\eta_{\min} \sum_{n=0}^{N-1} \Delta t \int_{\Omega} |D\mathbf{u}_h^{n+1}|^2 dx + \frac{1}{4} \varrho_0 \sum_{n=0}^{N-1} \int_{\Omega} |\mathbf{u}_h^{n+1} - \mathbf{u}_h^n|^2 dx \leq K_0 + \varrho_0 |\mathbf{g}|_2 |\Omega|^{1/2} \sum_{n=0}^{N-1} \Delta t |\mathbf{u}_h^{n+1}|_{(L^2(\Omega))^d}. \end{aligned}$$

En utilisant les lemmes de Poincaré, de Korn et l'inégalité de Young, nous trouvons :

$$\begin{aligned} & C \left[\frac{M_1 \underline{\Sigma}}{(\Sigma_M)^2} \sum_{n=0}^{N-1} \Delta t \sum_{i=1}^3 |\nabla \mu_{ih}^{n+1}|^2 + \frac{3}{8} (2\beta - 1) \varepsilon \underline{\Sigma} \sum_{n=0}^{N-1} \int_{\Omega} \sum_{i=1}^3 |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|^2 dx \right] \\ & + C_K \eta_{\min} \sum_{n=0}^{N-1} \Delta t \int_{\Omega} |\nabla \mathbf{u}_h^{n+1}|^2 dx + \frac{1}{4} \varrho_0 \sum_{n=0}^{N-1} \int_{\Omega} |\mathbf{u}_h^{n+1} - \mathbf{u}_h^n|^2 dx \leq K_0 + t_f \frac{\varrho_0^2 |\mathbf{g}|_2^2 |\Omega| (C_p)^2 C_K}{\eta_{\min}}. \end{aligned}$$

Cette inégalité donne à la fois (VI.57) et (VI.58). ■

Le passage à la limite dans les équations non linéaires (*cf* section VI.5.3) nécessite l'établissement de la convergence forte des sous-suites, pour cela il est utile d'obtenir des estimations plus fines.

Proposition VI.22

Il existe deux constantes K_4 et K_5 indépendantes de h et Δt telles que :

$$\left(\sum_{n=0}^{N-1} \Delta t \sum_{i=1}^3 \left| \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \right|_{(H^1(\Omega))'}^2 \right) + \left(\frac{1}{\sqrt{\Delta t}} \sum_{n=0}^{N-1} \sum_{i=1}^3 |c_{ih}^{n+1} - c_{ih}^n|_{(L^2(\Omega))}^2 \right) \leq K_4, \quad (\text{VI.60})$$

$$\sum_{n=0}^{N-i-1} \Delta t |\mathbf{u}^{n+i} - \mathbf{u}^n|_{(L^2(\Omega))^d}^2 \leq K_5 (t^i)^{\frac{1}{4}}, \quad \forall i \in \llbracket 0, N-1 \rrbracket. \quad (\text{VI.61})$$

Démonstration :

(i) L'estimation (VI.60) s'obtient à partir de la première équation du système de Cahn-Hilliard.

(α) Considérons $\nu_h^\mu \in \mathcal{V}_h^\mu$ pour l'instant quelconque. La première équation de (VI.2) s'écrit :

$$\int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx = \int_{\Omega} [c_{ih}^n - \alpha_i] [\mathbf{u}_h^n - \frac{\Delta t}{\varrho_0} \sum_{j=1}^3 (c_{jh}^n - \alpha_j) \nabla \mu_{jh}^{n+1}] \cdot \nabla \nu_h^\mu dx - \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} \nabla \mu_{ih}^{n+1} \cdot \nabla \nu_h^\mu dx.$$

Ainsi l'inégalité inverse (VI.12), il vient :

$$\begin{aligned} \left| \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx \right| &\leq (|\alpha_i| + |c_{ih}^n|_{L^4(\Omega)}) |\mathbf{u}_h^n|_{L^4(\Omega)} |\nabla \nu_h^\mu|_{L^2(\Omega)} \\ &\quad + \frac{\Delta t}{\varrho_0} (|\alpha_i| + |c_{ih}^n|_{L^\infty(\Omega)}) \sum_{j=1}^3 (|\alpha_j| + |c_{jh}^n|_{L^\infty(\Omega)}) \left| \nabla \mu_{jh}^{n+1} \right|_{L^2(\Omega)} |\nabla \nu_h^\mu|_{L^2(\Omega)} \\ &\quad + \frac{M_2}{\Sigma_m} |\nabla \mu_{ih}^{n+1}|_{L^2(\Omega)} |\nabla \nu_h^\mu|_{L^2(\Omega)} \\ &\leq (|\alpha_i| + |c_{ih}^n|_{H^1(\Omega)}) |\mathbf{u}_h^n|_{H^1(\Omega)} |\nu_h^\mu|_{H^1(\Omega)} \\ &\quad + \frac{\Delta t}{\varrho_0} (|\alpha_i| + C_{\text{inv}}(h)^{\frac{1}{2}} |c_{ih}^n|_{H^1(\Omega)}) \sum_{j=1}^3 (1 + C_{\text{inv}}(h)^{\frac{1}{2}} |c_{jh}^n|_{H^1(\Omega)}) \left| \mu_{jh}^{n+1} \right|_{H^1(\Omega)} |\nu_h^\mu|_{H^1(\Omega)} \\ &\quad + \frac{M_2}{\Sigma_m} |\mu_{ih}^{n+1}|_{H^1(\Omega)} |\nu_h^\mu|_{H^1(\Omega)} \end{aligned}$$

Finalement, grâce à (VI.52) et (VI.56), nous obtenons qu'il existe une constante K (indépendante de h et Δt) telle que :

$$\left| \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx \right| \leq K \left[|\mathbf{u}_h^n|_{H^1(\Omega)} + \sum_{i=1}^3 |\mu_{ih}^{n+1}|_{H^1(\Omega)} \right] |\nu_h^\mu|_{H^1(\Omega)}. \quad (\text{VI.62})$$

Nous allons maintenant utiliser cette inégalité intermédiaire afin de montrer (VI.60).

(β) Soit $\nu \in H^1(\Omega)$. On note ν_h^μ le projeté L^2 de ν sur $\mathcal{V}_h^\mu$. D'après (VI.11), nous avons

$$|\nu_h^\mu|_{H^1(\Omega)} \leq C |\nu|_{H^1(\Omega)}.$$

Ainsi, en utilisant (VI.62), nous obtenons

$$\begin{aligned} \left| \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu dx \right| &= \left| \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx \right| \leq K \left[|\mathbf{u}_h^n|_{H^1(\Omega)} + \sum_{i=1}^3 |\mu_{ih}^{n+1}|_{H^1(\Omega)} \right] |\nu_h^\mu|_{H^1(\Omega)} \\ &\leq KC \left[|\mathbf{u}_h^n|_{H^1(\Omega)} + \sum_{i=1}^3 |\mu_{ih}^{n+1}|_{H^1(\Omega)} \right] |\nu|_{H^1(\Omega)}. \end{aligned}$$

Puisque cette inégalité est vraie pour tout $\nu \in H^1(\Omega)$, nous avons

$$\begin{aligned} \left| \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \right|_{(H^1(\Omega))'} &= \sup_{\nu \in H^1(\Omega)} \frac{\left| \left(\frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t}, \nu \right)_{L^2(\Omega)} \right|}{\|\nu\|_{H^1(\Omega)}} \\ &\leq KC \left[\|\mathbf{u}_h^n\|_{H^1(\Omega)} + \sum_{i=1}^3 \|\mu_{ih}^{n+1}\|_{H^1(\Omega)} \right]. \end{aligned}$$

Et par suite, en utilisant (VI.57), il vient :

$$\begin{aligned} \sum_{n=0}^{N-1} \Delta t \sum_{i=1}^3 \left| \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \right|_{(H^1(\Omega))'}^2 &\leq 6K^2 C^2 \left[\sum_{n=0}^{N-1} \Delta t \|\mathbf{u}_h^n\|_{H^1(\Omega)}^2 + 3 \sum_{n=0}^{N-1} \Delta t \sum_{i=1}^3 \|\mu_{ih}^{n+1}\|_{H^1(\Omega)}^2 \right] \\ &\leq 18K^2 C^2 K_2. \end{aligned} \quad (\text{VI.63})$$

(γ) Prenons maintenant $\nu_h^\mu = \Delta t(c_{ih}^{n+1} - c_{ih}^n)$ dans (VI.62). Il vient :

$$\sum_{i=1}^3 \left| \int_{\Omega} |c_{ih}^{n+1} - c_{ih}^n|^2 dx \right| \leq K \Delta t \left[\|\mathbf{u}_h^n\|_{H^1(\Omega)} + \sum_{i=1}^3 \|\mu_{ih}^{n+1}\|_{H^1(\Omega)} \right] \sum_{i=1}^3 \|c_{ih}^{n+1} - c_{ih}^n\|_{H^1(\Omega)}.$$

et donc, en utilisant (VI.57) et (VI.58), nous avons

$$\begin{aligned} \sum_{n=0}^{N-1} \sum_{i=1}^3 \|c_{ih}^{n+1} - c_{ih}^n\|_{L^2(\Omega)}^2 &\leq K \left[\left(\sum_{n=0}^{N-1} \Delta t \|\mathbf{u}_h^n\|_{H^1(\Omega)}^2 \right)^{\frac{1}{2}} + \left(\sum_{n=0}^{N-1} \Delta t \sum_{i=1}^3 \|\mu_{ih}^{n+1}\|_{H^1(\Omega)}^2 \right)^{\frac{1}{2}} \right] \left(\sum_{n=0}^{N-1} \Delta t \sum_{i=1}^3 \|c_{ih}^{n+1} - c_{ih}^n\|_{H^1(\Omega)}^2 \right)^{\frac{1}{2}} \\ &\leq 2\sqrt{K_2} \sqrt{K_3} \sqrt{\Delta t}. \end{aligned} \quad (\text{VI.64})$$

L'inégalité (VI.60) se déduit immédiatement des équations (VI.63) et (VI.64) en définissant la constante $K_4 = \max(18K^2 C^2 K_2, 2\sqrt{3K_2} \sqrt{K_3})$.

(ii) Pour obtenir l'estimation (VI.61), nous commençons par estimer le terme $\|\mathbf{u}_h^{n+i} - \mathbf{u}_h^n\|_{(L^2(\Omega))^d}^2$ pour $n \in \llbracket 0, N-i-1 \rrbracket$. Pour cela, nous choisissons $\boldsymbol{\nu}_h^{\mathbf{u}} \in \mathcal{V}_{h,0}^{\mathbf{u}}$ tel que

$$\int_{\Omega} \nu_h^p \operatorname{div} \boldsymbol{\nu}_h^{\mathbf{u}} dx = 0, \quad \forall \nu_h^p \in \mathcal{V}^p, \quad (\text{VI.65})$$

comme fonction test dans (VI.43) et sommons les équations de manière à obtenir :

$$\begin{aligned} \int_{\Omega} \varrho_0 (\mathbf{u}_h^{n+i} - \mathbf{u}_h^n) \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx &+ \underbrace{\frac{1}{2} \sum_{k=n}^{n+i-1} \Delta t \int_{\Omega} \varrho_0 (\mathbf{u}_h^k \cdot \nabla) \mathbf{u}_h^{k+1} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx}_{(1)} - \underbrace{\frac{1}{2} \sum_{k=n}^{n+i-1} \Delta t \int_{\Omega} \varrho_0 (\mathbf{u}_h^k \cdot \nabla) \boldsymbol{\nu}_h^{\mathbf{u}} \cdot \mathbf{u}_h^{k+1} dx}_{(2)} \\ &+ \underbrace{\sum_{k=n}^{n+i-1} \Delta t \int_{\Omega} 2\eta_h^{k+1} D\mathbf{u}_h^{k+1} : D\boldsymbol{\nu}_h^{\mathbf{u}} dx}_{(3)} = \underbrace{\sum_{k=n}^{n+i-1} \Delta t \int_{\Omega} \varrho_0 \mathbf{g} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx}_{(4)} - \underbrace{\sum_{k=n}^{n+i-1} \Delta t \int_{\Omega} \sum_{j=1}^3 (c_{jh}^k - \alpha_j) \nabla \mu_{jh}^{k+1} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx}_{(5)}. \end{aligned}$$

Nous estimons alors chacun des termes numérotés de cette égalité séparément. Pour le terme (1), nous

obtenons :

$$\begin{aligned}
& \left| \frac{1}{2} \sum_{k=n}^{n+i-1} \Delta t \int_{\Omega} \varrho_0 (\mathbf{u}_h^k \cdot \nabla) \mathbf{u}_h^{k+1} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx \right| \\
& \leq \frac{1}{2} \varrho_0 \Delta t \sum_{k=n}^{n+i-1} \|\mathbf{u}_h^k\|_{(L^3(\Omega))^d} \|\mathbf{u}_h^{k+1}\|_{(H^1(\Omega))^d} \|\boldsymbol{\nu}_h^{\mathbf{u}}\|_{(L^6(\Omega))^d} \\
& \leq \frac{1}{2} \varrho_0 \|\boldsymbol{\nu}_h^{\mathbf{u}}\|_{(H^1(\Omega))^d} \Delta t \sum_{k=n}^{n+i-1} \|\mathbf{u}_h^k\|_{(L^2(\Omega))^d}^{\frac{1}{2}} \|\mathbf{u}_h^k\|_{(L^6(\Omega))^d}^{\frac{1}{2}} \|\mathbf{u}_h^{k+1}\|_{(H^1(\Omega))^d} \\
& \leq \frac{1}{2} \varrho_0 K_1^{\frac{1}{2}} \|\boldsymbol{\nu}_h^{\mathbf{u}}\|_{(H^1(\Omega))^d} \Delta t \sum_{k=n}^{n+i-1} \|\mathbf{u}_h^k\|_{(H^1(\Omega))^d}^{\frac{1}{2}} \|\mathbf{u}_h^{k+1}\|_{(H^1(\Omega))^d} \\
& \leq \frac{1}{2} \varrho_0 K_1^{\frac{1}{2}} \|\boldsymbol{\nu}_h^{\mathbf{u}}\|_{(H^1(\Omega))^d} \Delta t \sum_{k=n}^{n+i-1} \frac{2}{3} \left[\|\mathbf{u}_h^k\|_{(H^1(\Omega))^d}^{\frac{3}{2}} + \|\mathbf{u}_h^{k+1}\|_{(H^1(\Omega))^d}^{\frac{3}{2}} \right] \\
& \leq \frac{1}{3} \varrho_0 K_1^{\frac{1}{2}} \|\boldsymbol{\nu}_h^{\mathbf{u}}\|_{(H^1(\Omega))^d} \Delta t i^{\frac{1}{4}} \left[\left(\sum_{k=n}^{n+i-1} \|\mathbf{u}_h^k\|_{(H^1(\Omega))^d}^2 \right)^{\frac{3}{4}} + \left(\sum_{k=n}^{n+i-1} \|\mathbf{u}_h^{k+1}\|_{(H^1(\Omega))^d}^2 \right)^{\frac{3}{4}} \right] \\
& \leq \frac{2}{3} \varrho_0 K_1^{\frac{1}{2}} K_2^{\frac{3}{4}} \|\boldsymbol{\nu}_h^{\mathbf{u}}\|_{(H^1(\Omega))^d} (t_f)^{\frac{3}{4}} (t^i)^{\frac{1}{4}}.
\end{aligned}$$

Le terme (2) s'estime de la même manière :

$$\begin{aligned}
& \left| \frac{1}{2} \sum_{k=n}^{n+i-1} \Delta t \int_{\Omega} \varrho_0 (\mathbf{u}_h^k \cdot \nabla) \boldsymbol{\nu}_h^{\mathbf{u}} \cdot \mathbf{u}_h^{k+1} dx \right| \leq \frac{1}{2} \varrho_0 \Delta t \sum_{k=n}^{n+i-1} \|\mathbf{u}_h^k\|_{(L^3(\Omega))^d} \|\boldsymbol{\nu}_h^{\mathbf{u}}\|_{(H^1(\Omega))^d} \|\mathbf{u}_h^{k+1}\|_{(L^6(\Omega))^d} \\
& \leq \frac{1}{2} \varrho_0 \|\boldsymbol{\nu}_h^{\mathbf{u}}\|_{(H^1(\Omega))^d} \Delta t \sum_{k=n}^{n+i-1} \|\mathbf{u}_h^k\|_{(L^2(\Omega))^d}^{\frac{1}{2}} \|\mathbf{u}_h^k\|_{(L^6(\Omega))^d}^{\frac{1}{2}} \|\mathbf{u}_h^{k+1}\|_{(H^1(\Omega))^d} \\
& \leq \frac{2}{3} \varrho_0 K_1^{\frac{1}{2}} K_2^{\frac{3}{4}} \|\boldsymbol{\nu}_h^{\mathbf{u}}\|_{(H^1(\Omega))^d} (t_f)^{\frac{3}{4}} (t^i)^{\frac{1}{4}}.
\end{aligned}$$

Pour le terme visqueux (3), nous dérivons l'estimation suivante :

$$\begin{aligned}
& \left| \sum_{k=n}^{n+i-1} \Delta t \int_{\Omega} 2\eta_h^{k+1} D\mathbf{u}_h^{k+1} : D\boldsymbol{\nu}_h^{\mathbf{u}} dx \right| \leq 2\eta_{\max} \Delta t \sum_{k=n}^{n+i-1} \|\mathbf{u}_h^{k+1}\|_{(H^1(\Omega))^d} \|\boldsymbol{\nu}_h^{\mathbf{u}}\|_{(H^1(\Omega))^d} \\
& \leq 2\eta_{\max} \|\boldsymbol{\nu}_h^{\mathbf{u}}\|_{(H^1(\Omega))^d} \Delta t i^{\frac{1}{2}} \left(\sum_{k=n}^{n+i-1} \|\mathbf{u}_h^{k+1}\|_{(H^1(\Omega))^d}^2 \right)^{\frac{1}{2}} \\
& \leq 2\eta_{\max} K_2^{\frac{1}{2}} \|\boldsymbol{\nu}_h^{\mathbf{u}}\|_{(H^1(\Omega))^d} (t_f)^{\frac{1}{2}} (t^i)^{\frac{1}{2}}.
\end{aligned}$$

Enfin, il reste les termes (3) et (4) du second membre :

$$\left| \sum_{k=n}^{n+i-1} \Delta t \int_{\Omega} \varrho_0 \mathbf{g} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx \right| \leq \varrho_0 \|\mathbf{g}\|_2 \|\Omega\|^{\frac{1}{2}} \|\boldsymbol{\nu}_h^{\mathbf{u}}\|_{L^2(\Omega)} t^i,$$

et

$$\begin{aligned}
& \left| \sum_{k=n}^{n+i-1} \Delta t \int_{\Omega} \sum_{j=1}^3 (c_{jh}^k - \alpha_j) \nabla \mu_{jh}^{k+1} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx \right| \leq \sum_{k=n}^{n+i-1} \Delta t \sum_{j=1}^3 \|c_{jh}^k - \alpha_j\|_{L^4(\Omega)} \|\mu_{jh}^{k+1}\|_{H^1(\Omega)} \|\boldsymbol{\nu}_h^{\mathbf{u}}\|_{(L^4(\Omega))^d} \\
& \leq \|\Omega\| [K_1 + \max_{i=1,2,3} |\alpha_i|] K_2^{\frac{1}{2}} \|\boldsymbol{\nu}_h^{\mathbf{u}}\|_{(H^1(\Omega))^d} (t^i)^{\frac{1}{2}}.
\end{aligned}$$

Finalement, nous obtenons le résultat suivant : il existe une constante K strictement positive telle que, pour tout $\boldsymbol{\nu}_h^{\mathbf{u}} \in \mathcal{V}_{h,0}^{\mathbf{u}}$ satisfaisant (VI.65), nous avons

$$\left| \int_{\Omega} (\mathbf{u}_h^{n+i} - \mathbf{u}_h^n) \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx \right| \leq K \|\boldsymbol{\nu}_h^{\mathbf{u}}\|_{(H^1(\Omega))^d} (t^i)^{\frac{1}{4}}.$$

En particulier, pour $\nu_h^u = u_h^{n+i} - u_h^n$ (qui satisfait bien (VI.65) d'après (VI.43)), on trouve

$$|u_h^{n+i} - u_h^n|_{L^2(\Omega)}^2 \leq K |u_h^{n+i} - u_h^n|_{(H^1(\Omega))^d} (t^i)^{\frac{1}{4}}.$$

Nous obtenons ainsi :

$$\begin{aligned} \sum_{n=0}^{N-i-1} \Delta t |u_h^{n+i} - u_h^n|_{(L^2(\Omega))^d}^2 &\leq 2K (t^i)^{\frac{1}{4}} \sum_{n=0}^{N-1} \Delta t |u_h^n|_{(H^1(\Omega))^d}^2 \\ &\leq 2K (t_f)^{\frac{1}{2}} (t^i)^{\frac{1}{4}} \sum_{n=0}^{N-1} \Delta t |u_h^n|_{(H^1(\Omega))^d}^2 \\ &\leq 2KK_2 (t_f)^{\frac{1}{2}} (t^i)^{\frac{1}{4}}. \end{aligned}$$

Ce qui donne la conclusion en posant $K_5 = 2KK_2(t_f)^{\frac{1}{2}}$. ■

VI.5.2 Argument de compacité, convergence des sous-suites

Les estimations démontrées dans la section VI.5.1 (proposition VI.21 et VI.22), nous permettent d'obtenir (à sous-suite près) les convergences des suites $\mathbf{c}_{h_K}^{N_K}$, $\bar{\mathbf{c}}_{h_K}^{N_K}$, $\underline{\mathbf{c}}_{h_K}^{N_K}$, $\boldsymbol{\mu}_{h_K}^{N_K}$, $\mathbf{u}_{h_K}^{N_K}$, $\bar{\mathbf{u}}_{h_K}^{N_K}$ et $\underline{\mathbf{u}}_{h_K}^{N_K}$. Les propositions suivantes précisent en quel sens ont lieu ces convergences.

Proposition VI.23

A sous-suites près, nous avons les convergences suivantes lorsque $K \rightarrow +\infty$:

$$\mathbf{c}_{h_K}^{N_K} \rightharpoonup \mathbf{c} \quad \text{dans } L^\infty(0, t_f, (H^1(\Omega))^3) \text{ faible-}, \quad (\text{VI.66})$$

$$\boldsymbol{\mu}_{h_K}^{N_K} \rightharpoonup \boldsymbol{\mu} \quad \text{dans } L^2(0, t_f, (H^1(\Omega))^3) \text{ faible}, \quad (\text{VI.67})$$

$$\frac{\partial \mathbf{c}_{h_K}^{N_K}}{\partial t} \rightharpoonup \frac{\partial \mathbf{c}}{\partial t} \quad \text{dans } L^2(0, t_f, (H^1(\Omega))') \text{ faible}. \quad (\text{VI.68})$$

$$\mathbf{u}_{h_K}^{N_K} \rightharpoonup \mathbf{u} \quad \text{dans } L^2(0, t_f, (H^1(\Omega))^d) \text{ faible}. \quad (\text{VI.69})$$

Démonstration : Les convergences (VI.66), (VI.67), (VI.68) et (VI.69) sont des conséquences directes de la proposition VI.21. En effet, il est facile de vérifier que les estimations fournies dans cette proposition montrent que les suites $\mathbf{c}_{h_K}^{N_K}$, $\boldsymbol{\mu}_{h_K}^{N_K}$, $\partial_t \mathbf{c}_{h_K}^{N_K}$ et $\mathbf{u}_{h_K}^{N_K}$ sont bornées dans les normes $L^\infty(0, t_f, (H^1(\Omega))^3)$, $L^2(0, t_f, (H^1(\Omega))^3)$, $L^2(0, t_f, (H^1(\Omega))')$, $L^2(0, t_f, (H^1(\Omega))^d)$ respectivement. ■

Les convergences faibles que nous venons d'obtenir ne sont pas suffisantes pour passer à la limite dans les termes non linéaires des systèmes de Cahn-Hilliard et Navier-Stokes. Nous montrons dans les deux propositions ci-dessous qu'il est possible d'obtenir la convergence forte des paramètres d'ordre et de la vitesse dans certains espaces précisés ci-après.

Proposition VI.24

A sous-suites près, nous avons les convergences suivantes lorsque $K \rightarrow +\infty$:

$$\mathbf{c}_{h_K}^{N_K} \rightarrow \mathbf{c} \quad \text{dans } C^0(0, t_f, (L^q(\Omega))^3) \text{ fort, pour tout } 1 \leq q < +\infty \text{ si } d = 2, \text{ ou } 1 \leq q < 6 \text{ si } d = 3, \quad (\text{VI.70})$$

$$\mathbf{c}_{h_K}^{N_K} \rightarrow \mathbf{c} \quad \text{dans } L^2(0, t_f, (L^2(\Omega))^3) \text{ fort}, \quad (\text{VI.71})$$

$$\underline{\mathbf{c}}_{h_K}^{N_K} \rightarrow \mathbf{c} \quad \text{dans } L^2(0, t_f, (L^2(\Omega))^3) \text{ fort}, \quad (\text{VI.72})$$

$$\bar{\mathbf{c}}_{h_K}^{N_K} \rightarrow \mathbf{c} \quad \text{dans } L^2(0, t_f, (L^2(\Omega))^3) \text{ fort}. \quad (\text{VI.73})$$

Démonstration : La suite $\mathbf{c}_{h_K}^{N_K}$ est bornée dans $L^\infty(0, t_f, (H^1(\Omega))^3)$ et sa dérivée en temps $\partial_t \mathbf{c}_{h_K}^{N_K}$ dans $L^2(0, t_f, (H^1(\Omega))')$. A l'identique de ce qui a été fait dans le chapitre V, nous obtenons la convergence forte (VI.70) du paramètre d'ordre par application du théorème de compacité de Aubin–Lions–Simon [Sim87]. De

cette convergence se déduit la convergence forte (VI.71), puis en utilisant l'inégalité (VI.58) les convergences fortes (VI.72) et (VI.73) des fonctions $\underline{\mathbf{c}}_{h_K}^{N_K}$ et $\bar{\mathbf{c}}_{h_K}^{N_K}$. ■

Le résultat de convergence forte sur la vitesse nécessite l'application d'un résultat de compacité un peu plus fin puisque nous ne disposons pas d'estimation sur sa dérivée en temps. Nous appliquons alors un théorème de compacité dû à Simon [Sim87, Théorème 5, p.84] où la condition sur la dérivée en temps est remplacée par une estimation sur les translatés en temps.

Nous commençons dans le lemme suivant par réécrire le terme à estimer. Ce terme est défini à partir de la fonction discrète $\underline{\mathbf{u}}_h^N$ qui est constante par morceaux (en temps) et de ses translatés en temps. Nous le relierons aux valeurs $\mathbf{u}_h^n$ de la fonction sur chacun des intervalles de temps pour pouvoir utiliser les estimations données dans la section VI.5.1. Pour une meilleure lisibilité, nous omettons encore le temps de ce lemme, l'indice K dans les notations h_K et N_K .

Lemme VI.25

Soit $\tau \in]0, t_f[$. Nous notons $i \in \llbracket 0, N-1 \rrbracket$ l'unique indice tel que $t^i \leq \tau < t^{i+1}$. Alors, nous avons :

(i) si $\tau < \Delta t$ alors

$$\int_0^{t_f-\tau} |\underline{\mathbf{u}}_h^N(s+\tau, \cdot) - \underline{\mathbf{u}}_h^N(s, \cdot)|_{(L^2(\Omega))^d}^2 ds = \tau \sum_{n=0}^{N-2} |\mathbf{u}_h^{n+1} - \mathbf{u}_h^n|_{(L^2(\Omega))^d}^2, \quad (\text{VI.74})$$

(ii) dans tous les cas, nous avons :

$$\begin{aligned} \int_0^{t_f-\tau} |\underline{\mathbf{u}}_h^N(s+\tau, \cdot) - \underline{\mathbf{u}}_h^N(s, \cdot)|_{(L^2(\Omega))^d}^2 ds \\ \leq \sum_{n=0}^{N-i-1} \Delta t |\mathbf{u}_h^{n+i} - \mathbf{u}_h^n|_{(L^2(\Omega))^d}^2 + \sum_{n=0}^{N-i-2} \Delta t |\mathbf{u}_h^{n+i+1} - \mathbf{u}_h^n|_{(L^2(\Omega))^d}^2. \end{aligned} \quad (\text{VI.75})$$

Démonstration : Commençons par écrire le membre de gauche sous la forme :

$$\begin{aligned} \int_0^{t_f-\tau} |\underline{\mathbf{u}}_h^N(s+\tau, \cdot) - \underline{\mathbf{u}}_h^N(s, \cdot)|_{(L^2(\Omega))^d}^2 ds &= \sum_{n=0}^{N-i-2} \int_{t^n}^{t^{n+1}} |\underline{\mathbf{u}}_h^N(s+\tau, \cdot) - \underline{\mathbf{u}}_h^N(s, \cdot)|_{(L^2(\Omega))^d}^2 ds \\ &\quad + \int_{t^{N-i-1}}^{t_f-\tau} |\underline{\mathbf{u}}_h^N(s+\tau, \cdot) - \underline{\mathbf{u}}_h^N(s, \cdot)|_{(L^2(\Omega))^d}^2 ds. \end{aligned}$$

Il ne reste plus qu'à identifier la valeur prise par le translaté de la fonction sur les intervalles $]t^n, t^{n+1}[$ pour $n \in \llbracket 0, N-i-2 \rrbracket$ et $]t^{N-i-1}, t_f - \tau[$.

Pour cela, nous introduisons le réel $\bar{\tau}$ défini par $\bar{\tau} = \tau - t^i$, nous fixons $n \in \llbracket 0, N-i-2 \rrbracket$ et distinguons les deux cas suivants :

– soit $s \in [t^n, t^{n+1} - \bar{\tau}]$, nous avons alors $t^{n+i} \leq t^n + \tau \leq s + \tau \leq t^{n+1} - \bar{\tau} + \tau \leq t^{n+i+1}$ et par suite

$$\underline{\mathbf{u}}_h^N(s+\tau, \cdot) = \mathbf{u}_h^{n+i}(\cdot).$$

– soit $s \in [t^{n+1} - \bar{\tau}, t^{n+1}]$, nous avons alors $t^{n+i+1} \leq t^{n+1} + \tau - \bar{\tau} \leq s + \tau \leq t^{n+1} + \tau \leq t^{n+i+2}$ et par suite

$$\underline{\mathbf{u}}_h^N(s+\tau, \cdot) = \mathbf{u}_h^{n+i+1}(\cdot).$$

Finalement, considérons maintenant le cas où $s \in [t^{N-i-1}, t_f - \tau]$. Nous avons $t^{N-i-1} \leq s \leq t^{N-i}$ et $t^{N-1} \leq t^{N-1} + \bar{\tau} \leq t^{N-i-1} + \tau \leq s + \tau \leq t^N$. Ainsi, pour tout $s \in [t^{N-i-1}, t_f - \tau]$, nous avons

$$\underline{\mathbf{u}}_h^N(s+\tau, \cdot) = \mathbf{u}_h^{N-1}(\cdot).$$

Nous déduisons de ce qui précède les égalités suivantes :

$$\begin{aligned}
\int_0^{t_f-\tau} |\underline{\mathbf{u}}_h^N(s+\tau, \cdot) - \underline{\mathbf{u}}_h^N(s, \cdot)|_{(L^2(\Omega))^d}^2 ds &= \sum_{n=0}^{N-i-2} \int_{t^n}^{t^{n+1}} |\underline{\mathbf{u}}_h^N(s+\tau, \cdot) - \underline{\mathbf{u}}_h^N(s, \cdot)|_{(L^2(\Omega))^d}^2 ds \\
&\quad + \int_{t^{N-i-1}}^{t_f-\tau} |\underline{\mathbf{u}}_h^N(s+\tau, \cdot) - \underline{\mathbf{u}}_h^N(s, \cdot)|_{(L^2(\Omega))^d}^2 ds \\
&= \sum_{n=0}^{N-i-2} \left[(\Delta t - \bar{\tau}) |\mathbf{u}_h^{n+i} - \mathbf{u}_h^n|_{(L^2(\Omega))^d}^2 + \bar{\tau} |\mathbf{u}_h^{n+i+1} - \mathbf{u}_h^n|_{(L^2(\Omega))^d}^2 \right] \\
&\quad + (t_f - \tau - t^{N-i-1}) |\mathbf{u}_h^{N-1} - \mathbf{u}_h^{N-1-i}|_{(L^2(\Omega))^d}^2 \\
&= \sum_{n=0}^{N-i-1} (\Delta t - \bar{\tau}) |\mathbf{u}_h^{n+i} - \mathbf{u}_h^n|_{(L^2(\Omega))^d}^2 + \sum_{n=0}^{N-i-2} \bar{\tau} |\mathbf{u}_h^{n+i+1} - \mathbf{u}_h^n|_{(L^2(\Omega))^d}^2.
\end{aligned}$$

Examinons maintenant ce qui se passe dans les cas (i) et (ii) :

- (i) si $\tau < \Delta t$ alors nous avons $i = 0$ et $\bar{\tau} = \tau$. L'égalité ci-dessus donne alors exactement la conclusion.
- (ii) la deuxième conclusion découle également de l'égalité ci-dessus puisque $0 \leq \bar{\tau} \leq \Delta t$.

■

Nous pouvons maintenant énoncer la proposition donnant la convergence forte de la vitesse.

Proposition VI.26

A sous-suites près, nous avons les convergences suivantes lorsque $K \rightarrow +\infty$:

$$\mathbf{u}_{h_K}^{N_K} \rightarrow \mathbf{u} \quad \text{dans } L^2(0, t_f, (L^2(\Omega))^d) \text{ fort,} \quad (\text{VI.76})$$

$$\underline{\mathbf{u}}_{h_K}^{N_K} \rightarrow \mathbf{u} \quad \text{dans } L^2(0, t_f, (L^2(\Omega))^d) \text{ fort,} \quad (\text{VI.77})$$

$$\overline{\mathbf{u}}_{h_K}^{N_K} \rightarrow \mathbf{u} \quad \text{dans } L^2(0, t_f, (L^2(\Omega))^d) \text{ fort.} \quad (\text{VI.78})$$

Démonstration : La démonstration repose sur un théorème de compacité dû à Simon [Sim87, Théorème 5, p.84] qui permet d'obtenir le fait que l'injection

$$L^2([0, t_f[, (H^1(\Omega))^d) \cap N_2^{\frac{1}{8}}([0, t_f[, (L^2(\Omega))^d) \hookrightarrow L^2([0, t_f[, (L^2(\Omega))^d)$$

est compacte. L'espace de Nikolskii $N_2^{\frac{1}{8}}([0, t_f[, (L^2(\Omega))^d)$ est défini par

$$N_2^{\frac{1}{8}}([0, t_f[, (L^2(\Omega))^d) = \left\{ \mathbf{v} \in L^2([0, t_f[, (L^2(\Omega))^d), \exists C > 0, \forall \tau \in]0, t_f[, |\mathbf{v}(\cdot + \tau, \cdot) - \mathbf{v}|_{L^2([0, t_f-\tau[, (L^2(\Omega))^d)} \leq C\tau^{\frac{1}{8}} \right\}.$$

Il est muni de la norme

$$|\mathbf{v}|_{N_2^{\frac{1}{8}}([0, t_f[, (L^2(\Omega))^d)} = \left(|\mathbf{v}|_{L^2([0, t_f[, (L^2(\Omega))^d)}^2 + \sup_{0 < \tau < t_f} \left(\frac{1}{\tau^{\frac{1}{8}}} |\mathbf{v}(\cdot + \tau, \cdot) - \mathbf{v}|_{L^2([0, t_f-\tau[, (L^2(\Omega))^d)} \right)^2 \right)^{\frac{1}{2}}.$$

Ainsi puisque la suite $\underline{\mathbf{u}}_{h_K}^{N_K}$ est bornée dans $L^2([0, t_f[, (H^1(\Omega))^d)$ et $L^2([0, t_f[, (L^2(\Omega))^d)$ (cf équations (VI.56) et (VI.57)), il suffit de montrer qu'elle l'est dans l'espace $N_2^{\frac{1}{8}}([0, t_f[, (L^2(\Omega))^d)$ pour obtenir la conclusion. Nous nous donnons donc $\tau \in]0, t_f[$ et omettons l'indice K dans la notation h_K et N_K le temps du calcul.

- (i) Si $\tau < \Delta t$ alors d'après le lemme VI.25, il vient :

$$\begin{aligned}
\int_0^{t_f-\tau} |\underline{\mathbf{u}}_h^N(s+\tau, \cdot) - \underline{\mathbf{u}}_h^N(s, \cdot)|_{(L^2(\Omega))^d}^2 ds &= \tau \sum_{n=0}^{N-2} |\mathbf{u}_h^{n+1} - \mathbf{u}_h^n|_{(L^2(\Omega))^d}^2 \\
&\leq K_3 \tau.
\end{aligned}$$

(ii) Si $\tau \geq \Delta t$ alors d'après le lemme VI.25 puis en utilisant l'inégalité (VI.61), il vient :

$$\begin{aligned} & \int_0^{t_f - \tau} |\underline{\mathbf{u}}_h^N(s + \tau, \cdot) - \underline{\mathbf{u}}_h^N(s, \cdot)|_{(L^2(\Omega))^d}^2 ds \\ & \leq \sum_{n=0}^{N-i-1} \Delta t |\mathbf{u}_h^{n+i} - \mathbf{u}_h^n|_{(L^2(\Omega))^d}^2 + \sum_{n=0}^{N-i-2} \Delta t |\mathbf{u}_h^{n+i+1} - \mathbf{u}_h^n|_{(L^2(\Omega))^d}^2 \\ & \leq K_5 \left[(t^i)^{\frac{1}{4}} + (t^{i+1})^{\frac{1}{4}} \right] \\ & \leq K_5 \left[1 + 2^{\frac{1}{4}} \right] \tau^{\frac{1}{4}}, \end{aligned}$$

puisque nous avons $t^i \leq \tau$ et $t^{i+1} = t^i + \Delta t \leq 2\tau$.

Dans tous les cas, nous avons obtenu l'existence d'une constante K_6 strictement positive (indépendante de h et Δt) telle que, $\forall \tau \in]0, t_f[$,

$$\int_0^{t_f - \tau} |\underline{\mathbf{u}}_h^N(s + \tau, \cdot) - \underline{\mathbf{u}}_h^N(s, \cdot)|_{(L^2(\Omega))^d}^2 ds \leq K_6 \tau^{\frac{1}{4}}.$$

Ceci termine donc la preuve de la convergence (VI.77). Les convergences (VI.76) et (VI.78) s'obtiennent alors directement à partir de cette dernière grâce à l'inégalité (VI.58). ■

VI.5.3 Passage à la limite dans le schéma

Les convergences obtenues dans la section VI.5.2 permettent de passer à la limite dans le système discret. Pour le système de Cahn-Hilliard (en l'absence du terme de transport), ce travail a été déjà réalisé en détail dans le chapitre V. Nous nous focalisons donc dans cette section sur le terme de transport dans l'équation de Cahn-Hilliard ainsi que sur le système de Navier-Stokes.

Pour simplifier les écritures nous omettons encore une fois l'indice K des notations N_K et h_K mais il faut garder à l'esprit que lorsque nous parlons de convergence cela signifie $K \rightarrow +\infty$ (et en conséquence $N_K \rightarrow +\infty$ et $h_K \rightarrow 0$).

Terme de transport de l'équation de Cahn-Hilliard

Soient $\nu^\mu \in \mathcal{C}^\infty(\overline{\Omega})$ une fonction donnée et $\tau \in \mathcal{C}_c^\infty(]0, t_f[)$. Nous définissons ν_h^μ comme le projeté H^1 de la fonction ν^μ sur $\mathcal{V}_h^\mu$. En complément aux démonstrations du chapitre V, nous devons démontrer la convergence suivante :

$$\int_0^{t_f} \int_\Omega [\underline{\mathcal{C}}_{ih}^N - \alpha_i] \left[\underline{\mathbf{u}}_h^N - \frac{\Delta t}{\varrho_0} \sum_{j=1}^3 (\underline{\mathcal{C}}_{jh}^N - \alpha_j) \nabla \mu_{jh}^N \right] \cdot \nabla \nu_h^\mu dx \tau(t) dt \longrightarrow \int_0^{t_f} \int_\Omega (c_i - \alpha_i) \mathbf{u} \cdot \nabla \nu^\mu dx \tau(t) dt. \quad (\text{VI.79})$$

Nous procédons en deux étapes en considérant séparément deux termes du membre de gauche : le terme standard de transport puis le terme additionnel garantissant la stabilité inconditionnelle.

Les inégalités suivantes nous permettent d'identifier la limite du premier terme :

$$\begin{aligned} & \left| \int_0^{t_f} \int_\Omega (\underline{\mathcal{C}}_{ih}^N - \alpha_i) \underline{\mathbf{u}}_h^N \cdot \nabla \nu_h^\mu dx \tau(t) dt - \int_0^{t_f} \int_\Omega (c_i - \alpha_i) \mathbf{u} \cdot \nabla \nu^\mu dx \tau(t) dt \right| \\ & \leq \left| \int_0^{t_f} \int_\Omega (\underline{\mathcal{C}}_{ih}^N - \alpha_i) \underline{\mathbf{u}}_h^N \cdot \nabla (\nu_h^\mu - \nu^\mu) dx \tau(t) dt \right| \\ & \quad + \left| \int_0^{t_f} \int_\Omega (\underline{\mathcal{C}}_{ih}^N - \alpha_i) (\underline{\mathbf{u}}_h^N - \mathbf{u}) \cdot \nabla \nu^\mu dx \tau(t) dt \right| + \left| \int_0^{t_f} \int_\Omega (\underline{\mathcal{C}}_{ih}^N - c_i) \mathbf{u} \cdot \nabla \nu^\mu dx \tau(t) dt \right| \\ & \leq |\tau|_{L^\infty(0, t_f)} |\nu_h^\mu - \nu^\mu|_{H^1(\Omega)} |\underline{\mathcal{C}}_{ih}^N - \alpha_i|_{L^2(0, t_f, H^1(\Omega))} |\underline{\mathbf{u}}_h^N|_{L^2(0, t_f, (H^1(\Omega))^d)} \\ & \quad + |\tau|_{L^\infty(0, t_f)} |\nabla \nu^\mu|_{L^3(\Omega)} |\underline{\mathcal{C}}_{ih}^N - \alpha_i|_{L^2(0, t_f, H^1(\Omega))} |\underline{\mathbf{u}}_h^N - \mathbf{u}|_{L^2(0, t_f, (L^2(\Omega))^d)} \\ & \quad + |\tau|_{L^\infty(0, t_f)} |\nabla \nu^\mu|_{L^3(\Omega)} |\underline{\mathcal{C}}_{ih}^N - c_i|_{L^2(0, t_f, L^2(\Omega))} |\mathbf{u}|_{L^2(0, t_f, (H^1(\Omega))^d)} \\ & \longrightarrow 0, \end{aligned}$$

puisque les fonctions $\underline{c}_{ih}^N$ et $\underline{\mu}_{jh}^N$ sont bornées dans $L^2(0, t_f, H^1(\Omega))$ et $L^2(0, t_f, (H^1(\Omega))^d)$ respectivement et que de plus, $\underline{c}_{ih}^N$ converge fort dans $L^2(0, t_f, L^2(\Omega))$ vers c_i (cf équation (VI.72)), $\underline{\mu}_{jh}^N$ converge fort dans $L^2(0, t_f, (L^2(\Omega))^d)$ vers $\mathbf{u}$ (cf équation (VI.77)) et $|\nu^\mu - \nu_h^\mu|_{H^1(\Omega)} = \inf_{\nu_h \in \mathcal{V}^\mu} |\nu^\mu - \nu_h|_{H^1(\Omega)} \xrightarrow{h \rightarrow 0} 0$ par l'hypothèse (VI.7).

Nous utilisons maintenant le fait que les suites $\underline{c}_{ih}^N$ et μ_{jh}^N sont bornées en norme $L^\infty(0, t_f, H^1(\Omega))$ et $L^2(0, t_f, H^1(\Omega))$ respectivement, l'inégalité inverse (VI.12) et la condition (VI.52) sur les suites h_K et N_K pour montrer que le second terme tend vers 0 :

$$\begin{aligned}
& \left| \int_0^{t_f} \int_\Omega [\underline{c}_{ih}^N - \alpha_i] \left[\frac{\Delta t}{\varrho_0} \sum_{j=1}^3 (\underline{c}_{jh}^N - \alpha_j) \nabla \mu_{jh}^N \right] \cdot \nabla \nu_h^\mu dx \tau(t) dt \right| \\
& \leq \left| \int_0^{t_f} \int_\Omega [\underline{c}_{ih}^N - \alpha_i] \left[\frac{\Delta t}{\varrho_0} \sum_{j=1}^3 (\underline{c}_{jh}^N - \alpha_j) \nabla \mu_{jh}^N \right] \cdot \nabla (\nu_h^\mu - \nu^\mu) dx \tau(t) dt \right| \\
& \quad + \left| \int_0^{t_f} \int_\Omega [\underline{c}_{ih}^N - \alpha_i] \left[\frac{\Delta t}{\varrho_0} \sum_{j=1}^3 (\underline{c}_{jh}^N - \alpha_j) \nabla \mu_{jh}^N \right] \cdot \nabla \nu^\mu dx \tau(t) dt \right| \\
& \leq \frac{\Delta t}{\varrho_0} |\nabla (\nu_h^\mu - \nu^\mu)|_{L^2(\Omega)} \int_0^{t_f} |\underline{c}_{ih}^N - \alpha_i|_{L^\infty(\Omega)} \sum_{j=1}^3 |\underline{c}_{jh}^N - \alpha_j|_{L^\infty(\Omega)} |\nabla \mu_{jh}^N|_{L^2(\Omega)} \tau(t) dt \\
& \quad + \frac{\Delta t}{\varrho_0} |\nabla \nu^\mu|_{L^\infty(\Omega)} \int_0^{t_f} |\underline{c}_{ih}^N - \alpha_i|_{L^4(\Omega)} \sum_{j=1}^3 |\underline{c}_{jh}^N - \alpha_j|_{L^4(\Omega)} |\nabla \mu_{jh}^N|_{L^2(\Omega)} \tau(t) dt \\
& \leq \frac{\Delta t C_{\text{inv}}(h)}{\varrho_0} |\tau|_{L^2(0, t_f)} |\nu_h^\mu - \nu^\mu|_{H^1(\Omega)} |\underline{c}_{ih}^N - \alpha_i|_{L^\infty(0, t_f, H^1(\Omega))} \sum_{j=1}^3 |\underline{c}_{jh}^N - \alpha_j|_{L^\infty(0, t_f, H^1(\Omega))} |\mu_{jh}^N|_{L^2(0, t_f, H^1(\Omega))} \\
& \quad + \frac{\Delta t}{\varrho_0} |\tau|_{L^2(0, t_f)} |\nabla \nu^\mu|_{L^\infty(\Omega)} |\underline{c}_{ih}^N - \alpha_i|_{L^\infty(0, t_f, H^1(\Omega))} \sum_{j=1}^3 |\underline{c}_{jh}^N - \alpha_j|_{L^\infty(0, t_f, H^1(\Omega))} |\mu_{jh}^N|_{L^2(0, t_f, H^1(\Omega))} \\
& \longrightarrow 0.
\end{aligned}$$

Ainsi nous avons montré la convergence (VI.79). En réutilisant tels quels les raisonnements du chapitre V, nous pouvons passer à la limite dans les autres termes du système de Cahn-Hilliard.

Système de Navier-Stokes

Soient $\nu^\mathbf{u} \in \mathcal{C}_c^\infty(\Omega)$ vérifiant $\text{div}(\nu^\mathbf{u}) = 0$ et $\tau \in \mathcal{C}^1([0, t_f])$ telle que $\tau(t_f) = 0$.

Nous introduisons l'espace

$$Z_h = \left\{ \mathbf{z}_h \in \mathcal{V}_{h,0}^\mathbf{u}; \quad \forall \nu_h^p \in \mathcal{V}_h^p, \int_\Omega \text{div}(\mathbf{z}_h) \nu_h^p dx = 0 \right\}.$$

La condition inf-sup (VI.10) implique que la fonction $\nu^\mathbf{u} \in H_0^1(\Omega)$ à divergence nulle est "bien approchée" par des fonctions de Z_h . Ceci est détaillé dans la proposition ci-dessous.

Proposition VI.27 (Approximation des fonctions à divergence nulle, [BS08, 12.5.17, p.345])

Nous avons l'inégalité suivante :

$$\inf_{\mathbf{z}_h \in Z_h} |\nu^\mathbf{u} - \mathbf{z}_h|_{H^1(\Omega)} \leq \frac{1}{\beta} \inf_{\nu_h^\mathbf{u} \in \mathcal{V}_{h,0}^\mathbf{u}} |\nu^\mathbf{u} - \nu_h^\mathbf{u}|_{H^1(\Omega)}. \quad (\text{VI.80})$$

Démonstration : Considérons $\mathbf{v}_h \in \mathcal{V}_{h,0}^\mathbf{u}$ et notons $\pi_h \in \mathcal{V}_h^p$ le projeté L^2 de $\text{div}(\nu^\mathbf{u} - \mathbf{v}_h)$ sur $\mathcal{V}_h^p$ défini par :

$$\forall \nu_h^p \in \mathcal{V}_h^p, \quad \int_\Omega \pi_h \nu_h^p dx = \int_\Omega \nu_h^p \text{div}(\nu^\mathbf{u} - \mathbf{v}_h) dx.$$

Puisque $\pi_h \in \mathcal{V}_h^p$, la condition inf-sup (cf [BS08, 21.5.10, p. 344]) nous donne $\mathbf{w}_h \in \mathcal{V}_{h,0}^{\mathbf{u}}$ tel que

$$\forall \nu_h^p \in \mathcal{V}^p, \quad \int_{\Omega} \nu_h^p \operatorname{div}(\mathbf{w}_h) dx = \int_{\Omega} \nu_h^p \pi_h dx \quad \text{et} \quad |\mathbf{w}_h|_{(\mathbf{H}^1(\Omega))^d} \leq \frac{1}{\beta} |\pi_h|_{L^2(\Omega)}.$$

Par construction nous avons :

$$\forall \nu_h^p \in \mathcal{V}^p, \quad \int_{\Omega} \nu_h^p \operatorname{div}(\boldsymbol{\nu}^{\mathbf{u}} - \mathbf{v}_h) dx = \int_{\Omega} \nu_h^p \operatorname{div}(\mathbf{w}_h) dx.$$

Puisque $\operatorname{div}(\boldsymbol{\nu}^{\mathbf{u}}) = 0$, ceci signifie exactement que $\mathbf{v}_h + \mathbf{w}_h \in Z_h$. Nous avons également l'inégalité suivante :

$$|\mathbf{w}_h|_{\mathbf{H}^1(\Omega)} \leq \frac{1}{\beta} |\pi_h|_{L^2(\Omega)} \leq \frac{1}{\beta} |\operatorname{div}(\boldsymbol{\nu}^{\mathbf{u}} - \mathbf{v}_h)|_{L^2(\Omega)} \leq \frac{1}{\beta} |\boldsymbol{\nu}^{\mathbf{u}} - \mathbf{v}_h|_{\mathbf{H}^1(\Omega)}.$$

Nous obtenons alors la conclusion de la manière suivante :

$$\inf_{\mathbf{z}_h \in Z_h} |\boldsymbol{\nu}^{\mathbf{u}} - \mathbf{z}_h|_{\mathbf{H}^1(\Omega)} \leq |\boldsymbol{\nu}^{\mathbf{u}} - (\mathbf{v}_h + \mathbf{w}_h)|_{\mathbf{H}^1(\Omega)} \leq |\boldsymbol{\nu}^{\mathbf{u}} - \mathbf{v}_h|_{\mathbf{H}^1(\Omega)} + |\mathbf{w}_h|_{\mathbf{H}^1(\Omega)} \leq (1 + \frac{1}{\beta}) |\boldsymbol{\nu}^{\mathbf{u}} - \mathbf{v}_h|_{\mathbf{H}^1(\Omega)}.$$

■

Nous notons alors $\boldsymbol{\nu}_h^{\mathbf{u}}$ le projeté $\mathbf{H}^1$ de $\boldsymbol{\nu}^{\mathbf{u}}$ sur l'espace Z_h . La proposition VI.27 et l'hypothèse (VI.8) montrent que

$$\boldsymbol{\nu}_h^{\mathbf{u}} \rightarrow \boldsymbol{\nu}^{\mathbf{u}}, \quad \text{dans } (\mathbf{H}^1(\Omega))^d \text{ fort.} \quad (\text{VI.81})$$

Nous utilisons $\boldsymbol{\nu}_h^{\mathbf{u}}$ comme fonction test dans la première équation de (VI.43). Il vient :

$$\left\{ \begin{aligned} & \int_{\Omega} \varrho_0 \frac{\mathbf{u}_h^{n+1} - \mathbf{u}_h^n}{\Delta t} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx + \frac{1}{2} \int_{\Omega} \varrho_0 (\mathbf{u}_h^n \cdot \nabla) \mathbf{u}_h^{n+1} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx - \frac{1}{2} \int_{\Omega} \varrho_0 (\mathbf{u}_h^n \cdot \nabla) \boldsymbol{\nu}_h^{\mathbf{u}} \cdot \mathbf{u}_h^{n+1} dx \\ & + \int_{\Omega} 2\eta_h^{n+1} D\mathbf{u}_h^{n+1} : D\boldsymbol{\nu}_h^{\mathbf{u}} dx = \int_{\Omega} \varrho_0 \mathbf{g} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx - \int_{\Omega} \sum_{j=1}^3 (c_{jh}^n - \alpha_j) \nabla \mu_{jh}^{n+1} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx. \end{aligned} \right. \quad (\text{VI.82})$$

Nous multiplions cette équation par $\tau(t)$, $t \in]t^n, t^{n+1}[$ et intégrons entre t^n et t^{n+1} et sommons pour n allant de 0 à $N-1$ de manière à retrouver une formulation faible sur $]0, t_f[\times \Omega$. Avant de donner cette formulation, un calcul préliminaire est nécessaire pour transformer le terme instationnaire de manière à faire porter les dérivées en temps sur la fonction τ et non sur la vitesse :

$$\begin{aligned} \sum_{n=0}^{N-1} \int_{t^n}^{t^{n+1}} \int_{\Omega} \varrho_0 \frac{\mathbf{u}_h^{n+1} - \mathbf{u}_h^n}{\Delta t} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx \tau(t) dt &= \frac{\varrho_0}{\Delta t} \left[\sum_{n=1}^N \int_{t^{n-1}}^{t^n} \int_{\Omega} \mathbf{u}_h^n \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx \tau(t) dt - \sum_{n=0}^{N-1} \int_{t^n}^{t^{n+1}} \int_{\Omega} \mathbf{u}_h^n \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx \tau(t) dt \right] \\ &= -\varrho_0 \left[\sum_{n=0}^{N-1} \int_{t^n}^{t^{n+1}} \int_{\Omega} \mathbf{u}_h^n \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx \frac{\tau(t) - \tau(t - \Delta t)}{\Delta t} dt \right] \\ &\quad + \frac{\varrho_0}{\Delta t} \int_{t^{N-1}}^{t^N} \int_{\Omega} \mathbf{u}_h^{n=N} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx \tau(t) dt - \varrho_0 \int_{t^0}^{t^1} \int_{\Omega} \mathbf{u}_h^0 \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx \frac{\tau(t - \Delta t)}{\Delta t} dt \\ &= -\varrho_0 \int_0^{t_f} \int_{\Omega} \underline{\mathbf{u}}_h^N(t, x) \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx \frac{\tau(t) - \tau(t - \Delta t)}{\Delta t} dt \\ &\quad + \varrho_0 \int_{\Omega} \mathbf{u}_h^{n=N}(x) \cdot \boldsymbol{\nu}_h^{\mathbf{u}}(x) dx \int_0^1 \tau(t_f - t\Delta t) dt \\ &\quad - \varrho_0 \int_{\Omega} \mathbf{u}_h^0(x) \cdot \boldsymbol{\nu}_h^{\mathbf{u}}(x) dx \int_0^1 \tau(\Delta t(t-1)) dt. \end{aligned}$$

Ainsi à partir de (VI.82) et de l'égalité ci-dessus, nous obtenons la formulation suivante dans laquelle nous

pouvons passer à la limite :

$$\begin{aligned}
& - \underbrace{\varrho_0 \int_0^{t_f} \int_{\Omega} \underline{\mathbf{u}}_h^N(t, x) \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx \frac{\tau(t) - \tau(t - \Delta t)}{\Delta t} dt}_{(1)} - \underbrace{\varrho_0 \int_{\Omega} \mathbf{u}_h^0(x) \cdot \boldsymbol{\nu}_h^{\mathbf{u}}(x) dx \int_0^1 \tau(\Delta t(t - 1)) dt}_{(2)} \\
& + \underbrace{\frac{1}{2} \int_0^{t_f} \int_{\Omega} \varrho_0 (\underline{\mathbf{u}}_h^N \cdot \nabla) \bar{\mathbf{u}}_h^N \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx \tau(t) dt}_{(3)} - \underbrace{\frac{1}{2} \int_0^{t_f} \int_{\Omega} \varrho_0 (\underline{\mathbf{u}}_h^N \cdot \nabla) \boldsymbol{\nu}_h^{\mathbf{u}} \cdot \bar{\mathbf{u}}_h^N dx \tau(t) dt}_{(4)} \\
& + \underbrace{\int_0^{t_f} \int_{\Omega} 2\eta(\bar{\mathbf{c}}_h^N) D\bar{\mathbf{u}}_h^N : D\boldsymbol{\nu}_h^{\mathbf{u}} dx \tau(t) dt}_{(5)} = \underbrace{\int_0^{t_f} \int_{\Omega} \varrho_0 \mathbf{g} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx \tau(t) dt}_{(6)} \\
& - \underbrace{\int_0^{t_f} \int_{\Omega} \sum_{j=1}^3 (\underline{c}_{jh}^N - \alpha_j) \nabla \mu_{jh}^N \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx \tau(t) dt}_{(7)} - \underbrace{\varrho_0 \int_{\Omega} \mathbf{u}_h^{n=N}(x) \cdot \boldsymbol{\nu}_h^{\mathbf{u}}(x) dx \int_0^1 \tau(t_f - t\Delta t) dt}_{(8)}.
\end{aligned} \tag{VI.83}$$

La limite du terme (1) s'obtient facilement à partir des convergences fortes (VI.77), (VI.81) et de celle de la fonction $t \mapsto \frac{\tau(t) - \tau(t - \Delta t)}{\Delta t}$ vers τ' dans $L^2(0, t_f)$ (obtenue par exemple par convergence dominée puisque la fonction τ est $\mathcal{C}^1([0, t_f])$) :

$$(1) \longrightarrow \varrho_0 \int_0^{t_f} \int_{\Omega} \mathbf{u} \cdot \boldsymbol{\nu}^{\mathbf{u}} dx \tau'(t) dt.$$

Le terme (2) permet de montrer que $\mathbf{u}$ satisfait la condition initiale (IV.32) au sens faible. Les convergences (VI.15), (VI.81) et la convergence uniforme sur $[0, t_f]$ de la fonction $t \mapsto \tau(\Delta t(t - 1))$ vers la fonction constante égale à $\tau(0)$ permettent d'obtenir :

$$(2) \rightarrow \varrho_0 \int_{\Omega} \mathbf{u}^0(x) \cdot \boldsymbol{\nu}^{\mathbf{u}}(x) dx \tau(0) dt.$$

Concernant le terme (3), les inégalités suivantes nous permettent de conclure :

$$\begin{aligned}
& \left| \int_0^{t_f} \int_{\Omega} (\underline{\mathbf{u}}_h^N \cdot \nabla) \bar{\mathbf{u}}_h^N \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx \tau(t) dt - \int_0^{t_f} \int_{\Omega} (\mathbf{u} \cdot \nabla) \mathbf{u} \cdot \boldsymbol{\nu}^{\mathbf{u}} dx \tau(t) dt \right| \\
& \leq \left| \int_0^{t_f} \int_{\Omega} (\underline{\mathbf{u}}_h^N \cdot \nabla) \bar{\mathbf{u}}_h^N \cdot (\boldsymbol{\nu}_h^{\mathbf{u}} - \boldsymbol{\nu}^{\mathbf{u}}) dx \tau(t) dt \right| + \left| \int_0^{t_f} \int_{\Omega} ((\underline{\mathbf{u}}_h^N - \mathbf{u}) \cdot \nabla) \bar{\mathbf{u}}_h^N \cdot \boldsymbol{\nu}^{\mathbf{u}} dx \tau(t) dt \right| \\
& \quad + \left| \int_0^{t_f} \int_{\Omega} (\mathbf{u} \cdot \nabla) (\bar{\mathbf{u}}_h^N - \mathbf{u}) \cdot \boldsymbol{\nu}^{\mathbf{u}} dx \tau(t) dt \right| \\
& \leq |\bar{\mathbf{u}}_h^N|_{L^2(0, t_f, (H^1(\Omega))^d)} |(\boldsymbol{\nu}_h^{\mathbf{u}} - \boldsymbol{\nu}^{\mathbf{u}})|_{L^4(\Omega)} |\underline{\mathbf{u}}_h^N|_{L^2(0, t_f, (L^4(\Omega))^d)} |\tau|_{L^\infty(0, t_f)} \\
& \quad + |\tau(t)|_{L^\infty(0, t_f)} |\underline{\mathbf{u}}_h^N - \mathbf{u}|_{L^2(0, t_f, (L^2(\Omega))^d)} |\bar{\mathbf{u}}_h^N|_{L^2(0, t_f, (H^1(\Omega))^d)} |\boldsymbol{\nu}^{\mathbf{u}}|_{(L^\infty(\Omega))^d} \\
& \quad + \left| \int_0^{t_f} \int_{\Omega} (\mathbf{u} \cdot \nabla) (\bar{\mathbf{u}}_h^N - \mathbf{u}) \cdot \boldsymbol{\nu}^{\mathbf{u}} dx \tau(t) dt \right| \\
& \longrightarrow 0.
\end{aligned}$$

En effet, puisque les suites $(\underline{\mathbf{u}}_h^N)$ et $(\bar{\mathbf{u}}_h^N)$ sont bornées dans $L^2(0, t_f, (H^1(\Omega))^d)$, les convergences (VI.77) et (VI.81) montrent que les deux premiers termes du membre de droite ci-dessus tendent vers 0. Quant au dernier (le terme comportant l'intégrale), il tend également vers 0 par convergence faible de $\nabla \bar{\mathbf{u}}_h^N$ vers $\nabla \mathbf{u}$ dans $L^2(0, t_f, (L^2(\Omega))^d)$ (un raisonnement composante par composante donne le résultat puisque pour tout $1 \leq i, j \leq d$, la fonction $(t, x) \mapsto \mathbf{u}_i(x) \boldsymbol{\nu}_j^{\mathbf{u}}(x) \tau(t)$ est $L^2(0, t_f, L^2(\Omega))$)).

Le terme (4) se traite de manière similaire :

$$\begin{aligned}
& \left| \int_0^{t_f} \int_{\Omega} (\underline{\mathbf{u}}_h^N \cdot \nabla) \boldsymbol{\nu}_h^{\mathbf{u}} \cdot \bar{\mathbf{u}}_h^N dx \tau(t) dt - \int_0^{t_f} \int_{\Omega} (\mathbf{u} \cdot \nabla) \boldsymbol{\nu}^{\mathbf{u}} \cdot \mathbf{u} dx \tau(t) dt \right| \\
& \leq \left| \int_0^{t_f} \int_{\Omega} (\underline{\mathbf{u}}_h^N \cdot \nabla) (\boldsymbol{\nu}_h^{\mathbf{u}} - \boldsymbol{\nu}^{\mathbf{u}}) \cdot \bar{\mathbf{u}}_h^N dx \tau(t) dt \right| + \left| \int_0^{t_f} \int_{\Omega} (\underline{\mathbf{u}}_h^N \cdot \nabla) \boldsymbol{\nu}^{\mathbf{u}} \cdot (\bar{\mathbf{u}}_h^N - \mathbf{u}) dx \tau(t) dt \right| \\
& \quad + \left| \int_0^{t_f} \int_{\Omega} ((\underline{\mathbf{u}}_h^N - \mathbf{u}) \cdot \nabla) \boldsymbol{\nu}^{\mathbf{u}} \cdot \mathbf{u} dx \tau(t) dt \right| \\
& \leq \|\underline{\mathbf{u}}_h^N\|_{L^2(0,t_f,(L^4(\Omega))^d)} \|(\boldsymbol{\nu}_h^{\mathbf{u}} - \boldsymbol{\nu}^{\mathbf{u}})\|_{H^1(\Omega)} \|\bar{\mathbf{u}}_h^N\|_{L^2(0,t_f,(L^4(\Omega))^d)} \|\tau\|_{L^\infty(0,t_f)} \\
& \quad + \|\underline{\mathbf{u}}_h^N\|_{L^2(0,t_f,(L^6(\Omega))^d)} \|\nabla \boldsymbol{\nu}^{\mathbf{u}}\|_{(L^3(\Omega))^d} \|(\bar{\mathbf{u}}_h^N - \mathbf{u})\|_{L^2(0,t_f,(L^2(\Omega))^d)} \|\tau\|_{L^\infty(0,t_f)} \\
& \quad + \|\underline{\mathbf{u}}_h^N - \mathbf{u}\|_{L^2(0,t_f,(L^2(\Omega))^d)} \|\nabla \boldsymbol{\nu}^{\mathbf{u}}\|_{(L^3(\Omega))^d} \|\mathbf{u}\|_{L^2(0,t_f,(L^6(\Omega))^d)} \|\tau\|_{L^\infty(0,t_f)},
\end{aligned}$$

la conclusion étant obtenue cette fois en utilisant les convergences (VI.77), (VI.78), (VI.81) et le fait que les deux suites $\underline{\mathbf{u}}_h^N$ et $\bar{\mathbf{u}}_h^N$ sont bornées dans $L^2(0, t_f, (H^1(\Omega))^d)$.

La limite du terme (5) est obtenue en utilisant la convergence (à sous-suite près) suivante :

$$\eta(\bar{\mathbf{c}}_{h_K}^N) \rightarrow \eta(\mathbf{c}) \quad \text{dans } L^2(0, t_f, (L^2(\Omega))^d) \text{ fort.} \quad (\text{VI.84})$$

Cette convergence se montre par le théorème de convergence dominée (la viscosité η est une fonction continue bornée et $\bar{\mathbf{c}}_{h_K}^N$ converge fort dans $L^2(0, t_f, (L^2(\Omega))^3)$ donc presque partout à une nouvelle sous-suite près).

Ainsi, en utilisant les convergences (VI.81), (VI.84), le fait que la suite $\bar{\mathbf{u}}_h^N$ est bornée dans $L^2(0, t_f, (H^1(\Omega))^d)$, et la convergence faible de $D\bar{\mathbf{u}}_h^N$ vers $D\mathbf{u}$ dans $L^2(0, t_f, (L^2(\Omega))^d)$, nous obtenons

$$\begin{aligned}
& \left| \int_0^{t_f} \int_{\Omega} 2\eta(\bar{\mathbf{c}}_{h_K}^N) D\bar{\mathbf{u}}_h^N : D\boldsymbol{\nu}_h^{\mathbf{u}} dx \tau(t) dt - \int_0^{t_f} \int_{\Omega} 2\eta(\mathbf{c}) D\mathbf{u} : D\boldsymbol{\nu}^{\mathbf{u}} dx \tau(t) dt \right| \\
& \leq \left| \int_0^{t_f} \int_{\Omega} 2\eta(\bar{\mathbf{c}}_{h_K}^N) D\bar{\mathbf{u}}_h^N : D(\boldsymbol{\nu}_h^{\mathbf{u}} - \boldsymbol{\nu}^{\mathbf{u}}) dx \tau(t) dt \right| + \left| \int_0^{t_f} \int_{\Omega} 2(\eta(\bar{\mathbf{c}}_{h_K}^N) - \eta(\mathbf{c})) D\bar{\mathbf{u}}_h^N : D\boldsymbol{\nu}^{\mathbf{u}} dx \tau(t) dt \right| \\
& \quad + \left| \int_0^{t_f} \int_{\Omega} 2\eta(\mathbf{c}) D(\bar{\mathbf{u}}_h^N - \mathbf{u}) : D\boldsymbol{\nu}^{\mathbf{u}} dx \tau(t) dt \right| \\
& \leq 2\eta_{\max} \|\bar{\mathbf{u}}_h^N\|_{L^2(0,t_f,(H^1(\Omega))^d)} \|\boldsymbol{\nu}_h^{\mathbf{u}} - \boldsymbol{\nu}^{\mathbf{u}}\|_{(H^1(\Omega))^d} \|\tau\|_{L^2(0,t_f)} \\
& \quad + 2\|\eta(\bar{\mathbf{c}}_{h_K}^N) - \eta(\mathbf{c})\|_{L^2(0,t_f,L^2(\Omega))} \|\bar{\mathbf{u}}_h^N\|_{L^2(0,t_f,(H^1(\Omega))^d)} \|\nabla \boldsymbol{\nu}^{\mathbf{u}}\|_{L^\infty(\Omega)} \|\tau\|_{L^\infty(0,t_f)} \\
& \quad + \left| \int_0^{t_f} \int_{\Omega} 2\eta(\mathbf{c}) D(\bar{\mathbf{u}}_h^N - \mathbf{u}) : D\boldsymbol{\nu}^{\mathbf{u}} dx \tau(t) dt \right| \\
& \longrightarrow 0.
\end{aligned}$$

Par (VI.81), la convergence du terme (6) est immédiate :

$$(6) \rightarrow \int_{\Omega} \varrho_0 \mathbf{g} \cdot \boldsymbol{\nu}^{\mathbf{u}} dx \tau(t) dt.$$

La convergence pour le terme de force capillaire (7), s'obtient de la manière suivante :

$$\begin{aligned}
& \left| \int_0^{t_f} \int_{\Omega} \sum_{j=1}^3 (\underline{c}_{jh}^N - \alpha_j) \nabla \mu_{jh}^N \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx \tau(t) dt - \int_0^{t_f} \int_{\Omega} \sum_{j=1}^3 (c_{jh} - \alpha_j) \nabla \mu_{jh} \cdot \boldsymbol{\nu}^{\mathbf{u}} dx \tau(t) dt \right| \\
& \leq \left| \int_0^{t_f} \int_{\Omega} \sum_{j=1}^3 (\underline{c}_{jh}^N - \alpha_j) \nabla \mu_{jh}^N \cdot (\boldsymbol{\nu}_h^{\mathbf{u}} - \boldsymbol{\nu}^{\mathbf{u}}) dx \tau(t) dt \right| \\
& \quad + \left| \int_0^{t_f} \int_{\Omega} \sum_{j=1}^3 (\underline{c}_{jh}^N - c_{jh}) \nabla \mu_{jh}^N \cdot \boldsymbol{\nu}^{\mathbf{u}} dx \tau(t) dt \right| \\
& \quad + \left| \int_0^{t_f} \int_{\Omega} \sum_{j=1}^3 (c_{jh} - \alpha_j) \nabla (\mu_{jh}^N - \mu_{jh}) \cdot \boldsymbol{\nu}^{\mathbf{u}} dx \tau(t) dt \right| \\
& \leq \sum_{j=1}^3 \left| \underline{c}_{jh}^N - \alpha_j \right|_{L^{\infty}(0, t_f, L^4(\Omega))} \left| \nabla \mu_{jh}^N \right|_{L^2(0, t_f, (L^2(\Omega))^d)} \left| \boldsymbol{\nu}_h^{\mathbf{u}} - \boldsymbol{\nu}^{\mathbf{u}} \right|_{L^4(\Omega)} |\tau|_{L^2(0, t_f)} \\
& \quad + \sum_{j=1}^3 \left| \underline{c}_{jh}^N - c_{jh} \right|_{L^2(0, t_f, (L^2(\Omega))^d)} \left| \nabla \mu_{jh}^N \right|_{L^2(0, t_f, (L^2(\Omega))^d)} \left| \boldsymbol{\nu}^{\mathbf{u}} \right|_{L^{\infty}(\Omega)} |\tau|_{L^{\infty}(0, t_f)} \\
& \quad + \left| \int_{\Omega} \sum_{j=1}^3 (c_{jh} - \alpha_j) \nabla (\mu_{jh}^N - \mu_j) \cdot \boldsymbol{\nu}^{\mathbf{u}} dx \tau(t) dt \right| \\
& \longrightarrow 0.
\end{aligned}$$

Les deux premiers termes du membre de droite ci-dessus tendent vers 0 grâce aux convergences (VI.72) et (VI.81) puisque les suites $(\underline{c}_{jh}^N)$ et (μ_{jh}^N) sont bornées dans $L^2(0, t_f, H^1(\Omega))$ et $L^{\infty}(0, t_f, H^1(\Omega))$ respectivement. Le dernier terme tend vers 0 par convergence faible de $\nabla \mu_{jh}^N$ vers $\nabla \mu_j$ dans $L^2(0, t_f, (L^2(\Omega))^d)$.

Enfin, il ne reste plus qu'à montrer que le terme résiduel (8) tend vers 0. Ceci provient tout simplement du fait que

$$\left| \int_{\Omega} \mathbf{u}_h^N(x) \cdot \boldsymbol{\nu}_h^{\mathbf{u}}(x) dx \right| \leq \left| \mathbf{u}_h^N(\cdot) \right|_{L^2(\Omega)} \left| \boldsymbol{\nu}_h^{\mathbf{u}} \right|_{L^2(\Omega)} \leq K_1 \left| \boldsymbol{\nu}^{\mathbf{u}} \right|_{L^2(\Omega)},$$

et

$$\int_0^1 \tau(t_f - t \Delta t) dt \longrightarrow \tau(t_f) = 0.$$

En conclusion, nous venons donc de montrer que :

$$\begin{aligned}
& -\varrho_0 \int_0^{t_f} \int_{\Omega} \mathbf{u} \cdot \boldsymbol{\nu}^{\mathbf{u}} dx \tau' dt - \varrho_0 \int_{\Omega} \mathbf{u}^0 \cdot \boldsymbol{\nu}^{\mathbf{u}} dx \tau(0) \\
& \quad + \frac{1}{2} \int_0^{t_f} \int_{\Omega} \varrho_0 (\mathbf{u} \cdot \nabla) \mathbf{u} \cdot \boldsymbol{\nu}^{\mathbf{u}} dx \tau(t) dt - \frac{1}{2} \int_0^{t_f} \int_{\Omega} \varrho_0 (\mathbf{u} \cdot \nabla) \boldsymbol{\nu}^{\mathbf{u}} \cdot \mathbf{u} dx \tau(t) dt \\
& \quad + \int_0^{t_f} \int_{\Omega} 2\eta(\mathbf{c}) D\mathbf{u} : D\boldsymbol{\nu}^{\mathbf{u}} dx \tau(t) dt = \int_0^{t_f} \int_{\Omega} \varrho_0 \mathbf{g} \cdot \boldsymbol{\nu}^{\mathbf{u}} dx \tau(t) dt \\
& \quad - \int_0^{t_f} \int_{\Omega} \sum_{j=1}^3 (\underline{c}_j - \alpha_j) \nabla \mu_j \cdot \boldsymbol{\nu}^{\mathbf{u}} dx \tau(t) dt
\end{aligned}$$

Pour finir, le passage à la limite dans l'équation de contrainte permet d'obtenir :

$$\operatorname{div}(\mathbf{u}) = 0.$$

VI.6 Conclusion

Nous avons proposé dans ce chapitre un schéma original pour la discrétisation du modèle Cahn-Hilliard/Navier-Stokes triphasique.

Ce schéma est inconditionnellement stable et préserve, au niveau discret, les propriétés essentielles du modèle, à savoir la conservation du volume et le fait que la somme des trois paramètres d'ordre reste égale à 1 au cours du temps.

Nous avons démontré l'existence d'au moins une solution approchée, et, dans le cas homogène (*i.e.* trois phases de même densité), nous avons fait l'étude de convergence des solutions discrètes vers une solution faible du modèle (dont nous prouvons par ce procédé l'existence).

La principale perspective concerne l'étude de convergence dans le cas où les trois fluides en présence ont des densités différentes. Même si l'estimation d'énergie (et l'existence de solutions discrètes) restent vraies dans ce cas là, il est alors plus délicat d'obtenir les estimations donnant par compacité la convergence forte sur la vitesse qui est nécessaire pour passer à la limite dans les termes non linéaires. En effet, les équations de Navier-Stokes comportent alors un terme de la forme :

$$\mathbf{u} \cdot \partial_t \varrho.$$

La dérivée en temps de la densité est peu régulière puisque celle-ci est une fonction des paramètres d'ordre dont la dérivée en temps est seulement $L^2(0, t_f, (H^1(\Omega))')$.

Chapitre VII

Méthode de projection incrémentale

La méthode de projection incrémentale [God79] est une discrétisation en temps du système de Navier-Stokes incompressible. Elle permet de découpler la résolution du bilan de quantité de mouvement (problème de convection-diffusion non linéaire) de la prise en compte de la contrainte d'incompressibilité, en utilisant une stratégie à pas fractionnaires.

Pour décrire de manière simple son fonctionnement nous considérons, dans un premier temps, le problème de Stokes sur un domaine Ω ouvert connexe borné régulier de $\mathbb{R}^d$ ($d = 2$ ou 3) et un intervalle de temps fini $]0, T[$ ($T > 0$) :

$$\begin{cases} \frac{\partial \mathbf{u}}{\partial t} - \Delta \mathbf{u} + \nabla p = \mathbf{f}, & \text{dans }]0, T[\times \Omega, \\ \operatorname{div}(\mathbf{u}) = 0, & \text{dans }]0, T[\times \Omega, \end{cases} \quad (\text{VII.1})$$

où la vitesse $\mathbf{u} :]0, T[\times \Omega \rightarrow \mathbb{R}^d$ et la pression $p :]0, T[\times \Omega \rightarrow \mathbb{R}$ sont les inconnues du système et $\mathbf{f} :]0, T[\times \Omega \rightarrow \mathbb{R}^d$ est un terme source supposé donné et régulier.

Nous supposons que la frontière Γ de Ω est l'union de deux parties disjointes Γ_D et Γ_N sur lesquelles nous imposons respectivement des conditions aux bords de type Dirichlet et Neumann :

$$\begin{cases} \mathbf{u} = \mathbf{u}_D, & \text{sur }]0, T[\times \Gamma_D, \\ \nabla \mathbf{u} \cdot \mathbf{n} - p \mathbf{n} = \mathbf{f}_N, & \text{sur }]0, T[\times \Gamma_N, \end{cases}$$

où $\mathbf{n}$ représente la normale à la frontière Γ extérieure au domaine Ω et les fonctions $\mathbf{u}_D$ et $\mathbf{f}_N$ sont données.

Enfin, nous supposons que la condition initiale

$$\mathbf{u}(0, \cdot) = \mathbf{u}^0, \quad \text{dans } \Omega,$$

est donnée.

Nous considérons une discrétisation uniforme $0 = t^0 < t^1 < \dots < t^N = T$ de l'intervalle de temps $]0, T[$ et nous notons $\Delta t = t^{n+1} - t^n$ ($0 \leq n \leq N-1$) le pas de temps. Par ailleurs, dans la suite de ce chapitre, lorsqu'une fonction $\mathbf{f}$ est donnée, la notation $\mathbf{f}^n$ ($0 \leq n \leq N$) désigne la valeur $\mathbf{f}(t^n, \cdot)$ de la fonction $\mathbf{f}$ au temps t^n .

Nous initialisons l'algorithme de projection par la donnée initiale $\mathbf{u}^0$ pour la vitesse et par une donnée initiale p^0 arbitraire pour la pression (en pratique nous utilisons $p^0 = 0$).

Etant donné une approximation $(\mathbf{u}^n, p^n)$ du couple vitesse-pression à l'instant t^n , la première étape de la méthode de projection consiste à produire une approximation intermédiaire $\tilde{\mathbf{u}}^{n+1}$ de la vitesse à l'instant t^{n+1} en ignorant la contrainte d'incompressibilité (le terme de pression pouvant alors être explicité) :

$$\begin{cases} \frac{\tilde{\mathbf{u}}^{n+1} - \mathbf{u}^n}{\Delta t} - \Delta \tilde{\mathbf{u}}^{n+1} + \nabla p^n = \mathbf{f}^{n+1}, & \text{dans } \Omega, \\ \tilde{\mathbf{u}}^{n+1} = \mathbf{u}_D^{n+1}, & \text{sur } \Gamma_D, \\ \nabla \tilde{\mathbf{u}}^{n+1} \cdot \mathbf{n} - p^n \mathbf{n} = \mathbf{f}_N^{n+1}, & \text{sur } \Gamma_N. \end{cases} \quad (\text{VII.2})$$

Cette vitesse prédite $\tilde{\mathbf{u}}^{n+1}$ est ensuite corrigée par la résolution du problème (de type Darcy) suivant, permettant d'obtenir les approximations $\mathbf{u}^{n+1}$ de la vitesse et p^{n+1} de la pression à l'instant t^{n+1} :

$$\begin{cases} \frac{\mathbf{u}^{n+1} - \tilde{\mathbf{u}}^{n+1}}{\Delta t} + \nabla(p^{n+1} - p^n) = 0, & \text{dans } \Omega, \\ \operatorname{div}(\mathbf{u}^{n+1}) = 0, & \text{dans } \Omega, \\ \mathbf{u}^{n+1} \cdot \mathbf{n} = \mathbf{u}_D^{n+1} \cdot \mathbf{n}, & \text{sur } \Gamma_D, \\ p^{n+1} = p^n, & \text{sur } \Gamma_N. \end{cases} \quad (\text{VII.3})$$

Cet algorithme est consistant avec le problème continu au sens où, lorsque nous sommes les deux systèmes (VII.2) et (VII.3) précédents, nous obtenons :

$$\frac{\mathbf{u}^{n+1} - \mathbf{u}^n}{\Delta t} - \Delta \tilde{\mathbf{u}}^{n+1} + \nabla p^{n+1} = \mathbf{f}^{n+1}, \text{ dans } \Omega. \quad (\text{VII.4})$$

Par ailleurs, les conditions aux bords imposées dans la deuxième étape peuvent se justifier de la manière suivante :

- la condition $\mathbf{u}^{n+1} \cdot \mathbf{n} = \mathbf{u}_D^{n+1} \cdot \mathbf{n}$ sur Γ_D semble raisonnable au vu de la condition de Dirichlet imposée à la vitesse sur le bord Γ_D ,
- la condition $p^{n+1} = p^n$ sur Γ_N permet d'obtenir, sur Γ_N , l'égalité $\nabla \tilde{\mathbf{u}}^{n+1} \cdot \mathbf{n} - p^{n+1} \mathbf{n} = \mathbf{f}_N^{n+1}$ qui paraît également raisonnable au vu de la condition de type Neumann imposée sur Γ_N . Par ailleurs, cette condition aux bords est une condition naturelle associée à l'équation (VII.4), au sens où les termes $\nabla \tilde{\mathbf{u}}^{n+1} \cdot \mathbf{n} - p^{n+1} \mathbf{n}$ apparaissent lors de l'intégration par partie des termes $-\Delta \tilde{\mathbf{u}}^{n+1} + \nabla p^{n+1}$ contre une fonction test.

En outre, ces conditions aux bords permettent d'identifier $\mathbf{u}^{n+1}$ au projeté L^2 de $\tilde{\mathbf{u}}^{n+1}$ sur l'espace (affine) des fonctions à divergence nulle de trace normale $\mathbf{u}_D^{n+1} \cdot \mathbf{n}$ sur Γ_D suivant la décomposition de Leray :

$$(L^2(\Omega))^d = \{\mathbf{v} \in (L^2(\Omega))^d, \operatorname{div} \mathbf{v} = 0 \text{ dans } \Omega \text{ et } \mathbf{v} \cdot \mathbf{n} = 0 \text{ sur } \Gamma_D\} \oplus \nabla\{q \in H^1(\Omega); q = 0 \text{ sur } \Gamma_N\}. \quad (\text{VII.5})$$

Il est néanmoins important de remarquer (c'est le principal inconvénient de la méthode de projection incrémentale) que ces conditions aux bords imposent, pour tout $n \in \llbracket 0, N-1 \rrbracket$, les égalités suivantes :

$$\begin{cases} \nabla p^n \cdot \mathbf{n} = \nabla p^0 \cdot \mathbf{n}, & \text{sur } \Gamma_D, \\ p^n = p^0, & \text{sur } \Gamma_N. \end{cases}$$

Ces conditions sont artificielles (au sens où elles ne sont, en général, pas vérifiées par les solutions du problème continu) et conduisent à une perte de précision [GMS05].

Le système (VII.3) (*i.e.* étape de projection) peut être résolu en deux sous-étapes successives. En effet, en prenant la divergence de la première équation, la vitesse $\mathbf{u}^{n+1}$ est éliminée et nous obtenons une équation elliptique sur l'incrément de pression $\Phi^{n+1} = p^{n+1} - p^n$. Ainsi, formellement, la résolution du système (VII.3) est équivalente à :

$$\begin{cases} -\Delta \Phi^{n+1} = -\frac{1}{\Delta t} \operatorname{div}(\tilde{\mathbf{u}}^{n+1}), & \text{dans } \Omega, \\ \nabla \Phi^{n+1} \cdot \mathbf{n} = 0, & \text{sur } \Gamma_D, \\ \Phi^{n+1} = 0, & \text{sur } \Gamma_N, \end{cases} \quad \text{et} \quad \mathbf{u}^{n+1} = \tilde{\mathbf{u}}^{n+1} - \Delta t \nabla \Phi^{n+1} \quad \text{dans } \Omega. \quad (\text{VII.6})$$

La généralisation de cette méthode aux équations de Navier-Stokes incompressibles avec une densité variable a été proposée dans [GQ00]. Nous nous inspirons largement de cet article dans la section VII.1.3.

En pratique, la méthode de projection est utilisée en combinaison avec une discrétisation spatiale. Dans cette partie, nous étudions deux cadres assez différents : le premier est celui qui a été présenté dans les parties précédentes (éléments finis H^1 -conformes, raffinement local) et le deuxième est une discrétisation en espace avec des éléments finis non conformes de bas degré (de type Rannacher-Turek).

Dans le cadre des éléments finis conformes nous étudions les deux problématiques suivantes :

- d'une part, nous nous intéressons au cas particulier où le second membre $\mathbf{f}$ du bilan de quantité de mouvement s'écrit comme un gradient $\mathbf{f} = \nabla Q$.
- d'autre part, nous montrons qu'il est possible d'utiliser la méthode de projection pour le couplage avec les équations de Cahn-Hilliard triphasique tout en conservant l'estimation d'énergie obtenue dans le chapitre VI.

La section VII.2 est ensuite consacrée à la présentation d'un travail effectué en collaboration avec F. Dardalhon et J.C. Latché. Ce travail, en marge de la problématique Cahn-Hilliard/Navier-Stokes, a été effectué pendant le stage de master 2 de F. Dardalhon que j'ai eu l'occasion d'encadrer au cours de ma thèse. Il concerne l'étude de la méthode de projection incrémentale combinée à une discrétisation spatiale effectuée par éléments finis non conformes de bas degré de type Rannacher-Turek. Il faut alors donner un sens à l'opérateur elliptique portant sur l'incrément de pression. Classiquement, l'étape d'élimination de la pression est réalisée de manière algébrique après un lumping de la matrice masse de vitesse. Il est alors intéressant de remarquer que l'opérateur obtenu sur la pression est semblable à un Laplacien volumes finis contenant les conditions aux bords prescrites lors de l'étape (VII.3). Par ailleurs, il est possible d'écrire le problème totalement discret sous forme variationnelle en définissant des produits scalaires et normes dépendant du maillage. Ceci permet d'adapter à ce cas, les démonstrations des estimations d'erreurs obtenues dans le cas semi-discret [She92, Gue99] ou pour des éléments finis conformes [Gue96].

VII.1 Éléments finis conformes

Dans cette section, le cadre est donc celui des chapitres précédents :

- Nous supposons que les conditions aux bords sont de type Dirichlet sur l'ensemble du bord du domaine, *i.e.* $\Gamma_N = \emptyset$.
- La discrétisation en espace est réalisée grâce à la méthode de Galerkin et à la méthode des éléments finis. Nous utilisons les notations des chapitres précédents, celles-ci étant rappelées brièvement ci-dessous.

Nous considérons $\mathcal{V}_h^u$ et $\mathcal{V}_h^p$ des espaces d'approximation éléments finis de $\mathcal{V}^u = H^1(\Omega)$ et $\mathcal{V}^p = \{\nu^p \in L^2(\Omega); \int_{\Omega} \nu^p = 0\}$ respectivement. Puisque la vitesse vérifie des conditions de Dirichlet homogènes sur la frontière Γ , nous définissons l'espace d'approximation suivant :

$$\mathcal{V}_{h,0}^u = \{\nu_h^u \in \mathcal{V}_h^u; \nu_h^u = 0 \text{ sur } \Gamma\}.$$

Nous supposons que ces espaces d'approximation vérifient la condition inf-sup uniforme : il existe une constante strictement positive β (indépendante de h) telle que

$$\inf_{\nu_h^p \in \mathcal{V}_h^p} \sup_{\nu_h^u \in \mathcal{V}_{h,0}^u} \frac{\int_{\Omega} \nu_h^p \operatorname{div} \nu_h^u dx}{|\nu_h^p|_{L^2(\Omega)} |\nu_h^u|_{(H^1(\Omega))^d}} \geq \beta.$$

Enfin, nous supposons que l'espace d'approximation $\mathcal{V}_h^p$ est H^1 -conforme :

$$\mathcal{V}_h^p \subset H^1(\Omega).$$

En particulier, le problème elliptique (VII.6) peut être naturellement discrétisé dans cet espace.

Remarque VII.1

Ces hypothèses sont par exemple satisfaites par des éléments finis de type Lagrange $\mathbb{P}_{k+1}/\mathbb{P}_k$ pour $k \geq 1$. Nous renvoyons à [EG04] pour d'autres exemples.

VII.1.1 Problème de Stokes

Nous commençons par nous intéresser au problème de Stokes (VII.1). Dans ce cadre, la méthode de projection (sous forme variationnelle) s'écrit :

Problème VII.2 (Méthode de projection incrémentale standard)

Supposons que $(\mathbf{u}_h^n, p_h^n) \in \mathcal{V}_{h,0}^u \times \mathcal{V}^p$ sont donnés.

- **Étape 1** : Prédiction de vitesse

Trouver $\tilde{\mathbf{u}}_h^{n+1} \in \mathcal{V}_{h,0}^u$ tel que

$$\begin{aligned} \forall \nu_h^u \in \mathcal{V}_{h,0}^u, \quad & \int_{\Omega} \frac{\tilde{\mathbf{u}}_h^{n+1} - \mathbf{u}_h^n}{\Delta t} \cdot \nu_h^u dx + \int_{\Omega} \nabla \tilde{\mathbf{u}}_h^{n+1} : \nabla \nu_h^u dx \\ & + \int_{\Omega} \nu_h^u \cdot \nabla p_h^n dx = \int_{\Omega} \mathbf{f}^{n+1} \cdot \nu_h^u dx. \end{aligned} \tag{VII.7}$$

– **Etape 2.1** : Calcul de l'incrément de pression

Trouver $\Phi_h^{n+1} \in \mathcal{V}_h^p$ tel que

$$\forall \nu_h^p \in \mathcal{V}_h^p, \quad \int_{\Omega} \nabla \Phi_h^{n+1} \cdot \nabla \nu_h^p dx = \frac{1}{\Delta t} \int_{\Omega} \tilde{\mathbf{u}}_h^{n+1} \cdot \nabla \nu_h^p dx. \quad (\text{VII.8})$$

– **Etape 2.2** : Correction de la pression

$$p_h^{n+1} = p_h^n + \Phi_h^{n+1}. \quad (\text{VII.9})$$

– **Etape 2.3** : Correction de la vitesse

Trouver $\mathbf{u}_h^{n+1} \in \mathcal{V}_{h,0}^u$

$$\forall \nu_h^u \in \mathcal{V}_{h,0}^u, \quad \int_{\Omega} \mathbf{u}_h^{n+1} \cdot \nu_h^u dx = \int_{\Omega} \tilde{\mathbf{u}}_h^{n+1} \cdot \nu_h^u dx - \Delta t \int_{\Omega} \nu_h^u \cdot \nabla \Phi_h^{n+1} dx. \quad (\text{VII.10})$$

Afin d'introduire au mieux les sections à venir, nous rappelons l'analyse de stabilité de ce schéma. Celle-ci, ainsi que des estimations d'erreur peuvent être trouvées dans les articles [GQ98], [Gue99], [AJL09].

Théorème VII.3

Etant donnés $\mathbf{u}_h^n$ et p_h^n , supposons que le triplet $(\mathbf{u}_h^{n+1}, \tilde{\mathbf{u}}_h^{n+1}, p_h^{n+1})$ est solution du problème VII.2. Alors, nous avons l'inégalité suivante :

$$\begin{aligned} \frac{1}{2} |\mathbf{u}_h^{n+1}|_{(L^2(\Omega))^d}^2 - \frac{1}{2} |\mathbf{u}_h^n|_{(L^2(\Omega))^d}^2 + \frac{1}{2} |\tilde{\mathbf{u}}_h^{n+1} - \mathbf{u}_h^n|_{(L^2(\Omega))^d}^2 + \Delta t |\nabla \tilde{\mathbf{u}}_h^{n+1}|_{(L^2(\Omega))^d}^2 \\ + \frac{1}{2} \left[\Delta t^2 |\nabla p_h^{n+1}|_{(L^2(\Omega))^d}^2 - \Delta t^2 |\nabla p_h^n|_{(L^2(\Omega))^d}^2 \right] \leq \Delta t \int_{\Omega} \mathbf{f}^{n+1} \cdot \tilde{\mathbf{u}}_h^{n+1} dx. \end{aligned}$$

Démonstration :

(i) Nous prenons $\nu_h^u = \Delta t \tilde{\mathbf{u}}_h^{n+1}$ dans le système (VII.7) de l'étape de prédiction de vitesse pour obtenir :

$$\begin{aligned} \frac{1}{2} |\tilde{\mathbf{u}}_h^{n+1}|_{(L^2(\Omega))^d}^2 - \frac{1}{2} |\mathbf{u}_h^n|_{(L^2(\Omega))^d}^2 + \frac{1}{2} |\tilde{\mathbf{u}}_h^{n+1} - \mathbf{u}_h^n|_{(L^2(\Omega))^d}^2 \\ + \Delta t |\nabla \tilde{\mathbf{u}}_h^{n+1}|_{(L^2(\Omega))^d}^2 + \Delta t \int_{\Omega} \tilde{\mathbf{u}}_h^{n+1} \cdot \nabla p_h^n dx = \Delta t \int_{\Omega} \mathbf{f}^{n+1} \cdot \tilde{\mathbf{u}}_h^{n+1} dx. \end{aligned} \quad (\text{VII.11})$$

Le terme $\int_{\Omega} \tilde{\mathbf{u}}_h^{n+1} \cdot \nabla p_h^n dx$ n'est pas nul puisque $\tilde{\mathbf{u}}_h^{n+1}$ ne satisfait pas la contrainte de divergence nulle au sens faible (*i.e.* contre les fonctions tests de l'espace d'approximation $\mathcal{V}_h^p$ de la pression). Cette contrainte est ici imposée à la fonction $\hat{\mathbf{u}}_h = \tilde{\mathbf{u}}_h^{n+1} - \Delta t \nabla \Phi_h^{n+1}$ de $L^2(\Omega)$. En effet, l'étape de calcul (VII.8) de l'incrément de pression peut s'écrire :

$$\forall \nu_h^p \in \mathcal{V}_h^p, \quad \int_{\Omega} \hat{\mathbf{u}}_h \cdot \nabla \nu_h^p dx = 0. \quad (\text{VII.12})$$

C'est cette dernière relation que nous exploitons pour trouver l'expression de $\int_{\Omega} \tilde{\mathbf{u}}_h^{n+1} \cdot \nabla p_h^n dx$.

(ii) Puisque $\Phi_h^{n+1} = p_h^{n+1} - p_h^n$, nous avons par définition de $\hat{\mathbf{u}}_h$:

$$\hat{\mathbf{u}}_h + \Delta t \nabla p_h^{n+1} = \tilde{\mathbf{u}}_h^{n+1} + \Delta t \nabla p_h^n.$$

Nous évaluons la norme L^2 des deux membres de cette égalité (autrement dit nous élevons au carré et intégrons sur Ω) pour faire apparaître le terme $\int_{\Omega} \tilde{\mathbf{u}}_h^{n+1} \cdot \nabla p_h^n dx$:

$$|\hat{\mathbf{u}}_h|_{(L^2(\Omega))^d}^2 + \Delta t^2 |\nabla p_h^{n+1}|_{(L^2(\Omega))^d}^2 = |\tilde{\mathbf{u}}_h^{n+1}|_{(L^2(\Omega))^d}^2 + \Delta t^2 |\nabla p_h^n|_{(L^2(\Omega))^d}^2 + 2\Delta t \int_{\Omega} \tilde{\mathbf{u}}_h^{n+1} \cdot \nabla p_h^n dx,$$

puisque le double produit $2\Delta t \int_{\Omega} \widehat{\mathbf{u}}_h \cdot \nabla p_h^{n+1} dx$ est nul d'après l'égalité (VII.12). Ainsi, nous trouvons :

$$\Delta t \int_{\Omega} \widetilde{\mathbf{u}}_h^{n+1} \nabla p_h^n dx = \frac{1}{2} \left[|\widehat{\mathbf{u}}_h|_{(L^2(\Omega))^d}^2 - |\widetilde{\mathbf{u}}_h^{n+1}|_{(L^2(\Omega))^d}^2 + \Delta t^2 |\nabla p_h^{n+1}|_{(L^2(\Omega))^d}^2 - \Delta t^2 |\nabla p_h^n|_{(L^2(\Omega))^d}^2 \right]. \quad (\text{VII.13})$$

(iii) En sommant (VII.11) et (VII.13), nous obtenons :

$$\begin{aligned} \frac{1}{2} |\widehat{\mathbf{u}}_h|_{(L^2(\Omega))^d}^2 - \frac{1}{2} |\mathbf{u}_h^n|_{(L^2(\Omega))^d}^2 + \frac{1}{2} |\widetilde{\mathbf{u}}_h^{n+1} - \mathbf{u}_h^n|_{(L^2(\Omega))^d}^2 + \Delta t |\nabla \widetilde{\mathbf{u}}_h^{n+1}|_{(L^2(\Omega))^d}^2 \\ + \frac{1}{2} \left[\Delta t^2 |\nabla p_h^{n+1}|_{(L^2(\Omega))^d}^2 - \Delta t^2 |\nabla p_h^n|_{(L^2(\Omega))^d}^2 \right] = \Delta t \int_{\Omega} \mathbf{f}^{n+1} \cdot \widetilde{\mathbf{u}}_h^{n+1} dx. \end{aligned}$$

(iv) La dernière étape consiste à remarquer que l'étape de correction de vitesse définit $\mathbf{u}_h^{n+1}$ comme la projection $L^2(\Omega)$ de $\widehat{\mathbf{u}}_h$ dans $\mathcal{V}_{h,0}^{\mathbf{u}}$. Ainsi, nous avons

$$|\mathbf{u}_h^{n+1}|_{L^2(\Omega)} \leq |\widehat{\mathbf{u}}_h|_{L^2(\Omega)}.$$

(v) Nous pouvons alors conclure à l'inégalité d'énergie suivante :

$$\begin{aligned} \frac{1}{2} |\mathbf{u}_h^{n+1}|_{(L^2(\Omega))^d}^2 - \frac{1}{2} |\mathbf{u}_h^n|_{(L^2(\Omega))^d}^2 + \frac{1}{2} |\widetilde{\mathbf{u}}_h^{n+1} - \mathbf{u}_h^n|_{(L^2(\Omega))^d}^2 + \Delta t |\nabla \widetilde{\mathbf{u}}_h^{n+1}|_{(L^2(\Omega))^d}^2 \\ + \frac{1}{2} \left[\Delta t^2 |\nabla p_h^{n+1}|_{(L^2(\Omega))^d}^2 - \Delta t^2 |\nabla p_h^n|_{(L^2(\Omega))^d}^2 \right] \leq \Delta t \int_{\Omega} \mathbf{f}^{n+1} \cdot \widetilde{\mathbf{u}}_h^{n+1} dx. \end{aligned}$$

■

VII.1.2 Calcul d'un état d'équilibre : $\mathbf{f} = \nabla Q$

Supposons que le second membre de $\mathbf{f}$ du bilan de quantité de mouvement s'écrive comme le gradient d'une fonction $Q \in L_0^2(\Omega)$. La solution exacte du problème de Stokes (VII.1) est triviale : $\mathbf{u} = 0$, $p = Q$.

Nous nous intéressons alors à la problématique suivante : la méthode de projection incrémentale (cf problème VII.2) permet-elle de calculer la solution triviale exacte, lorsque la fonction Q appartient à $\mathcal{V}_h^p$, l'espace d'approximation des pressions ?

Nous commençons par donner un exemple très simple. Il s'agit de la simulation (en 2D) d'un fluide de densité et viscosité constante égale à 1, prisonnier dans une boîte $\Omega =]0, 1[^2$, soumis à la seule gravité $\mathbf{f} = \mathbf{g}$ où $\mathbf{g}$ est un vecteur constant. Nous résolvons numériquement par la méthode de projection incrémentale (cf problème VII.2) le problème de Stokes :

$$\begin{cases} \frac{\partial \mathbf{u}}{\partial t} - \Delta \mathbf{u} + \nabla p = \mathbf{g}, & \text{dans }]0, T[\times \Omega, \\ \operatorname{div}(\mathbf{u}) = 0, & \text{dans }]0, T[\times \Omega, \\ \mathbf{u} = 0, & \text{sur }]0, T[\times \Gamma, \\ \mathbf{u}(0, \cdot) = 0 & \text{dans } \Omega. \end{cases} \quad (\text{VII.14})$$

La simulation est effectuée avec un pas de temps Δt égal à 1. Nous utilisons des éléments finis de Taylor-Hood ($\mathbb{P}_2 - \mathbb{P}_1$ sur maillages triangles et $\mathbb{Q}_2 - \mathbb{Q}_1$ sur maillages quadrangles). La méthode de projection est initialisée en choisissant $\mathbf{u}^0 = 0$ et $p^0 = 0$. Les maillages utilisés ainsi que la solution $\mathbf{u}_h^1$ discrète obtenue à la fin du premier pas de temps sont présentés sur la figure VII.1 pour un maillage carré 20x20 structuré (à droite) et un maillage triangle non structuré (à gauche).

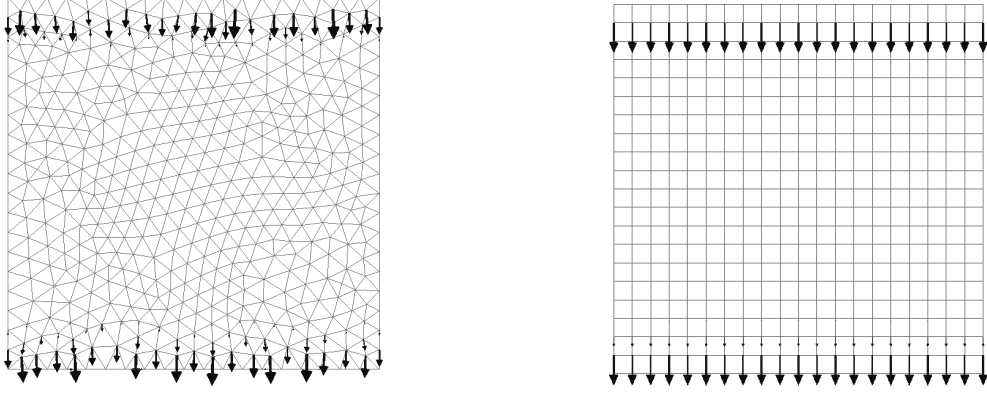


FIG. VII.1 – Exemple de vitesses parasites, $|\mathbf{u}_h^1|_{L^\infty(\Omega)} = 1.16 \times 10^{-3}$

La solution exacte ($\mathbf{u} \equiv 0$, $p = \varrho(\mathbf{g}_1 x + \mathbf{g}_2 y)$) du problème continu (VII.14) appartient à l'espace discret $\mathcal{V}_{h,0}^{\mathbf{u}} \times \mathcal{V}_h^p$. Pourtant, nous constatons que dès la première itération la vitesse discrète n'est plus nulle : $|\mathbf{u}_h^1|_{L^\infty(\Omega)} \sim 10^{-3}$. Le couple $(\mathbf{u}_h^n \equiv 0, p_h^n = \varrho(\mathbf{g}_1 x + \mathbf{g}_2 y))$ pour tout $n \in \llbracket 1, N \rrbracket$, n'est pas solution du problème VII.2, à cause de l'initialisation ($\mathbf{u}_h^0 \equiv 0, p_h^0 = 0$). Dans cet exemple simple, les vitesses non nulles sont de faible amplitude et sont localisées au voisinage des bords horizontaux (la condition au bord $\nabla p_h^n \cdot \mathbf{n} = 0, \forall n \in \llbracket 0, N \rrbracket$ imposée par le schéma n'est pas satisfaite par la solution exacte sur les bords horizontaux); mais ce phénomène peut prendre de l'ampleur lorsque le système de Navier-Stokes est couplé à d'autres équations comme le système de Cahn-Hilliard par exemple (cf section VII.1.3).

Il serait souhaitable que la vitesse prédite soit nulle lorsque le second membre s'écrit exactement comme le gradient d'une fonction (dépendant éventuellement du temps) de l'espace d'approximation $\mathcal{V}_h^p$ de la pression.

L'idée est alors d'appliquer la décomposition de Leray (VII.5) au second membre $\mathbf{f}$ de l'équation :

$$\begin{cases} \mathbf{f} = \mathbf{u}_f + \nabla p_f, \\ \operatorname{div}(\mathbf{u}_f) = 0, \\ \mathbf{u}_f \cdot \mathbf{n} = 0 \text{ sur } \Gamma. \end{cases} \quad (\text{VII.15})$$

et de discrétiser le système suivant (en lieu et place du système (VII.1)) à l'aide de la méthode de projection :

$$\begin{cases} \frac{\partial \mathbf{u}}{\partial t} - \Delta \mathbf{u} + \nabla q = \mathbf{f} - \nabla p_f, & \text{dans }]0, T[\times \Omega, \\ \operatorname{div}(\mathbf{u}) = 0, & \text{dans }]0, T[\times \Omega, \end{cases} \quad (\text{VII.16})$$

quitte à poser ensuite $p = q + p_f$.

Ainsi, nous obtenons l'algorithme (1)-(4)

- (1) Trouver $\tilde{\mathbf{u}}_h^{n+1} \in \mathcal{V}_{h,0}^{\mathbf{u}}$ tel que $\forall \boldsymbol{\nu}_h^{\mathbf{u}} \in \mathcal{V}_{h,0}^{\mathbf{u}}, \quad \int_{\Omega} \frac{\tilde{\mathbf{u}}_h^{n+1} - \mathbf{u}_h^n}{\Delta t} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx + \int_{\Omega} \nabla \tilde{\mathbf{u}}_h^{n+1} : \nabla \boldsymbol{\nu}_h^{\mathbf{u}} dx + \int_{\Omega} \boldsymbol{\nu}_h^{\mathbf{u}} \cdot \nabla q_h^n dx = \int_{\Omega} (\mathbf{f}^{n+1} - \nabla p_f^{n+1}) \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx.$
- (2) Trouver $\Phi_h^{n+1} \in \mathcal{V}_h^p$ tel que $\forall \nu_h^p \in \mathcal{V}_h^p, \quad \int_{\Omega} \nabla \Phi_h^{n+1} \cdot \nabla \nu_h^p dx = \frac{1}{\Delta t} \int_{\Omega} \tilde{\mathbf{u}}_h^{n+1} \cdot \nabla \nu_h^p dx.$
- (3) Poser $q_h^{n+1} = q_h^n + \Phi_h^{n+1}$.
- (4) Trouver $\mathbf{u}_h^{n+1} \in \mathcal{V}_{h,0}^{\mathbf{u}}, \forall \boldsymbol{\nu}_h^{\mathbf{u}} \in \mathcal{V}_{h,0}^{\mathbf{u}}, \quad \int_{\Omega} \mathbf{u}_h^{n+1} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx = \int_{\Omega} \tilde{\mathbf{u}}_h^{n+1} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx - \Delta t \int_{\Omega} \boldsymbol{\nu}_h^{\mathbf{u}} \cdot \nabla \Phi_h^{n+1} dx.$

Nous pouvons maintenant revenir à la variable $p_h^n = q_h^n + p_f^n$. En posant $\tilde{p}_h^{n+1} = p_h^n + p_f^{n+1} - p_f^n$, les étapes (1) et (3) s'écrivent de la manière suivante (les étapes (2) et (4) restant inchangée) :

- (1') Trouver $\tilde{\mathbf{u}}_h^{n+1} \in \mathcal{V}_{h,0}^{\mathbf{u}}$ tel que $\forall \boldsymbol{\nu}_h^{\mathbf{u}} \in \mathcal{V}_{h,0}^{\mathbf{u}}, \quad \int_{\Omega} \frac{\tilde{\mathbf{u}}_h^{n+1} - \mathbf{u}_h^n}{\Delta t} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx + \int_{\Omega} \nabla \tilde{\mathbf{u}}_h^{n+1} : \nabla \boldsymbol{\nu}_h^{\mathbf{u}} dx + \int_{\Omega} \boldsymbol{\nu}_h^{\mathbf{u}} \cdot \nabla \tilde{p}_h^{n+1} dx = \int_{\Omega} \mathbf{f}^{n+1} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx.$
- (3') Poser $p_h^{n+1} = \tilde{p}_h^{n+1} + \Phi_h^{n+1}$.

Il ne reste plus qu'à remarquer que, d'après (VII.15), nous pouvons obtenir p_f en résolvant :

$$\begin{cases} -\Delta p_f = -\operatorname{div}(\mathbf{f}), & \text{dans }]0, T[\times \Omega, \\ \nabla p_f \cdot \mathbf{n} = \mathbf{f} \cdot \mathbf{n}, & \text{sur }]0, T[\times \Gamma, \end{cases}$$

et que donc $\tilde{p}^{n+1}$ est la solution du système suivant :

$$\begin{cases} \Delta \tilde{p}^{n+1} = \Delta p^n + \operatorname{div}(\mathbf{f}^{n+1}) - \operatorname{div}(\mathbf{f}^n), & \text{dans } \Omega, \\ \nabla \tilde{p}^{n+1} \cdot \mathbf{n} = \nabla p^n \cdot \mathbf{n} + \mathbf{f}^{n+1} \cdot \mathbf{n} - \mathbf{f}^n \cdot \mathbf{n}, & \text{sur } \Gamma. \end{cases}$$

Nous proposons donc la variante suivante de la méthode de projection incrémentale standard (problème VII.2) :

Problème VII.4 (Variante de la méthode de projection incrémentale)

- **Initialisation** : Soit $\mathbf{u}_h^0 = 0$ et p_h^0 solution du problème suivant :
Trouver $p_h^0 \in \mathcal{V}_h^p$ tel que

$$\forall \nu_h^p \in \mathcal{V}^p, \quad \int_{\Omega} \nabla p_h^0 \cdot \nabla \nu_h^p dx = \int_{\Omega} \mathbf{f}^0 \cdot \nabla \nu_h^p dx.$$

- **Étape 0** : Prédiction de pression
Trouver $\tilde{p}_h^{n+1} \in \mathcal{V}_h^p$ tel que

$$\forall \nu_h^p \in \mathcal{V}^p, \quad \int_{\Omega} \nabla \tilde{p}_h^{n+1} \cdot \nabla \nu_h^p dx = \int_{\Omega} \nabla p_h^n \cdot \nabla \nu_h^p dx + \int_{\Omega} (\mathbf{f}^{n+1} - \mathbf{f}^n) \cdot \nabla \nu_h^p dx.$$

- **Étape 1** : Prédiction de vitesse
Trouver $\tilde{\mathbf{u}}_h^{n+1} \in \mathcal{V}_{h,0}^{\mathbf{u}}$ tel que

$$\forall \boldsymbol{\nu}_h^{\mathbf{u}} \in \mathcal{V}_{h,0}^{\mathbf{u}}, \quad \int_{\Omega} \frac{\tilde{\mathbf{u}}_h^{n+1} - \mathbf{u}_h^n}{\Delta t} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx + \int_{\Omega} \nabla \tilde{\mathbf{u}}_h^{n+1} : \nabla \boldsymbol{\nu}_h^{\mathbf{u}} dx - \int_{\Omega} \tilde{p}_h^{n+1} \operatorname{div}(\boldsymbol{\nu}_h^{\mathbf{u}}) dx = \int_{\Omega} \mathbf{f}^{n+1} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx.$$

- **Étape 2.1** : Calcul de l'incrément de pression
Trouver $\Phi_h^{n+1} \in \mathcal{V}_h^p$ tel que

$$\int_{\Omega} \nabla \Phi_h^{n+1} \cdot \nabla \nu_h^p dx = \frac{1}{\Delta t} \int_{\Omega} \tilde{\mathbf{u}}_h^{n+1} \cdot \nabla \nu_h^p dx.$$

- **Étape 2.2** : Correction de la pression

$$p_h^{n+1} = \tilde{p}_h^{n+1} + \Phi_h^{n+1}.$$

- **Étape 2.3** : Correction de la vitesse
Trouver $\mathbf{u}_h^{n+1} \in \mathcal{V}_{h,0}^{\mathbf{u}}$ tel que

$$\forall \boldsymbol{\nu}_h^{\mathbf{u}} \in \mathcal{V}_{h,0}^{\mathbf{u}}, \quad \int_{\Omega} \mathbf{u}_h^{n+1} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx = \int_{\Omega} \tilde{\mathbf{u}}_h^{n+1} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx + \Delta t \int_{\Omega} \Phi_h^{n+1} \operatorname{div}(\boldsymbol{\nu}_h^{\mathbf{u}}) dx.$$

L'avantage de cet algorithme est de permettre le calcul de la solution exacte dans les cas particulier où le second membre s'écrit comme le gradient d'une fonction de l'espace d'approximation de la pression. Ceci est énoncé dans la proposition suivante :

Proposition VII.5

Supposons que, pour tout $n \in \llbracket 0, N \rrbracket$, $\mathbf{f}^n = \nabla q_h^n$ avec $q_h^n \in \mathcal{V}_h^p$ et notons $(\mathbf{u}_h^n, p_h^n)_{n \in \llbracket 0, N \rrbracket}$ la solution approchée donnée par l'algorithme VII.4. Alors

$$\forall n \in \llbracket 0, N \rrbracket, \quad \mathbf{u}_h^n = 0 \quad \text{et} \quad p_h^n = q_h^n.$$

Démonstration : La preuve s'effectue par récurrence.

- L'étape d'initialisation donne $\mathbf{u}_h^0 = 0$ et $p_h^0 = q_h^0$.
- Supposons que pour un n donné nous avons $\mathbf{u}_h^n = 0$ et $p_h^n = q_h^n$. Alors nous obtenons après l'étape 0 :

$$\tilde{p}_h^{n+1} = q_h^{n+1}.$$

Et ainsi l'étape 1 devient :

$$\int_{\Omega} \frac{\tilde{\mathbf{u}}_h^{n+1}}{\Delta t} \cdot \nu_h^{\mathbf{u}} dx + \int_{\Omega} \nabla \tilde{\mathbf{u}}_h^{n+1} : \nabla \nu_h^{\mathbf{u}} dx = 0, \quad \forall \nu_h^{\mathbf{u}} \in \mathcal{V}_{h,0}^{\mathbf{u}}.$$

Ceci montre que $\tilde{\mathbf{u}}_h^{n+1} = 0$, par suite l'étape 2.1 donne $\Phi^{n+1} = 0$. Enfin les étapes 2.2 et 2.3 donnent $p_h^{n+1} = q_h^{n+1}$ et $\mathbf{u}_h^{n+1} = 0$. ■

Reprenons le cas test de la simulation d'un fluide au repos soumis à la gravité, présenté en début de cette section. La proposition VII.5 affirme qu'aucune vitesse parasite n'est créée par l'utilisation de l'algorithme VII.4. Ceci est confirmé par les tests numériques : $\|\mathbf{u}\|_{L^\infty} \sim 10^{-9}$.

Remarque VII.6

- Lorsque le second membre $\mathbf{f}$ ne dépend pas du temps, le problème VII.4 ne diffère de l'algorithme standard (problème VII.2) que par l'initialisation de la pression. Ce n'est bien sûr plus le cas lorsque le second membre dépend du temps.
- Les solutions du problème VII.4 satisfont la condition aux bords artificielle suivante :

$$\forall n \in \llbracket 0, N \rrbracket, \quad \nabla p_h^n \cdot \mathbf{n} = \mathbf{f}^n \cdot \mathbf{n}.$$

- La même idée appliquée à la version non-incrémentale (moins précise) de la méthode de projection conduit à un algorithme similaire dans lequel l'étape de prédiction de pression est remplacée par la suivante : Trouver $\tilde{p}_h^{n+1} \in \mathcal{V}_h^p$ tel que

$$\forall \nu_h^p \in \mathcal{V}^p, \quad \int_{\Omega} \nabla \tilde{p}_h^{n+1} \cdot \nabla \nu_h^p dx = \int_{\Omega} \mathbf{f}^{n+1} \cdot \nabla \nu_h^p dx.$$

Remarque VII.7

- Le raisonnement effectué dans cette section présente des liens étroits avec les algorithmes dits de séparation de pression (cf [GJ05] et [TOH09]) pour la résolution des équations de Navier-Stokes puisque ceux-ci consistent à soustraire une approximation de la pression au deux membres du bilan de quantité de mouvement avant d'effectuer sa résolution. C'est ce principe qui a conduit à l'écriture de l'équation (VII.16).
- L'idée sous-jacente est également proche des travaux effectués dans [GLBB97] permettant de limiter l'apparition de vitesses parasites. Dans cet article, les auteurs utilisent la décomposition (VII.15) et proposent de calculer dans un premier temps une approximation q_h de la "partie gradient" p_f du second membre $\mathbf{f}$, puis de remplacer, dans la résolution du problème de Stokes, le second membre $\mathbf{f}$ par $(\mathbf{f} - \nabla q_h) + \nabla(\Pi_h q_h)$ où Π_h est la projection L^2 sur l'espace d'approximation des pressions. Le calcul de q_h doit être effectué dans un espace d'approximation plus grand que celui des pressions, l'objectif étant que les deux termes $\mathbf{f} - \nabla q_h$ et $\nabla(\Pi_h q_h)$ génèrent le moins possible de vitesses parasites : le premier parce qu'il est proche de la "partie solénoïdale" $\mathbf{u}_f$ et le second parce qu'il s'écrit comme le gradient d'une fonction $(\Pi_h q_h)$ de l'espace d'approximation des pressions. Ceci suppose donc que la méthode de résolution des équations ne génère pas de vitesses parasites dans le cas particulier où le second membre s'écrit comme le gradient d'une fonction de l'espace d'approximation des pressions. Ainsi, la variante de la méthode de projection proposée ci-dessus pourrait être utilisée pour résoudre le système de Stokes en combinaison à de telles méthodes.

VII.1.3 Problème de Navier-Stokes à densité variable. Système couplé CH/NS

Nous revenons maintenant au modèle de Cahn-Hilliard/Navier-Stokes. Nous avons présenté dans le chapitre VI deux discrétisations semi-implicites en temps permettant de découpler la résolution du système de Cahn-Hilliard de celle du système de Navier-Stokes. Le premier schéma utilise une vitesse explicite $\mathbf{u}^n$ dans le terme

de transport de l'équation de Cahn-Hilliard permettant de calculer les paramètres d'ordre $\mathbf{c}_h^{n+1}$ et potentiels chimiques μ_h^{n+1} au temps $n + 1$ par la résolution du système de Cahn-Hilliard ; les équations de Navier-Stokes étant résolues dans un second temps pour obtenir la vitesse $\mathbf{u}_h^{n+1}$ et la pression p_h^{n+1} . Nous utilisons cette discrétisation dans les expérimentations numériques présentées dans la partie 3. Cependant, cette discrétisation ne permet pas de garantir la décroissance de l'énergie totale (somme de l'énergie libre et de l'énergie cinétique). Nous avons alors proposé un autre type de couplage permettant d'obtenir les bonnes propriétés théoriques néanmoins ce schéma n'a pas été testé numériquement.

Dans cette section, nous consacrons un paragraphe pour l'étude de chacun de ces deux schémas.

Avec l'utilisation du schéma standard nous avons été confrontés à la problématique des courants parasites (cf [SZ99, JTB02]). Il s'agit de vitesses de faible amplitude localisées au voisinage de l'interface. Nous allons montrer que ce phénomène est en fait lié au problème étudié dans la section VII.1.2.

Nous montrons dans un second paragraphe que nous pouvons utiliser la méthode de projection avec le schéma inconditionnellement stable présenté dans le chapitre VI en préservant la stabilité.

Méthode de projection, schéma de couplage standard

Problème VII.8

On suppose que $(c_h^n, \mathbf{u}_h^n, p_h^n) \in \mathcal{V}_h^c \times \mathcal{V}_{h,0}^u \times \mathcal{V}_h^p$ sont donnés.

– **Étape I : Système de Cahn-Hilliard**

Trouver $(\mathbf{c}_h^{n+1}, \mu_h^{n+1}) \in (\mathcal{V}_{h,S}^c)^3 \times (\mathcal{V}_h^\mu)^3$ tels que $\forall \nu_h^c \in \mathcal{V}_h^c, \forall \nu_h^\mu \in \mathcal{V}_h^\mu$, pour $i = 1, 2$ et 3 ,

$$\begin{cases} \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx - \int_{\Omega} \nu_h^\mu \mathbf{u}_h^n \cdot \nabla c_{ih}^{n+1} dx = - \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} \nabla \mu_{ih}^{n+1} \cdot \nabla \nu_h^\mu dx, \\ \int_{\Omega} \mu_{ih}^{n+1} \nu_h^c dx = \int_{\Omega} D_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \nu_h^c dx + \int_{\Omega} \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^{n+\beta} \nabla \nu_h^c dx, \end{cases} \quad (\text{VII.17})$$

– **Étape II.0 : Renormalisation de la pression**

Trouver $\tilde{p}_h^{n+1} \in \mathcal{V}_h^p$ tel que

$$\forall \nu_h^p \in \mathcal{V}_h^p, \quad \int_{\Omega} \frac{\nabla \tilde{p}_h^{n+1}}{\sqrt{\varrho_h^{n+1}}} \frac{\nabla \nu_h^p}{\sqrt{\varrho_h^{n+1}}} dx = \int_{\Omega} \frac{\nabla p_h^n}{\sqrt{\varrho_h^n}} \frac{\nabla \nu_h^p}{\sqrt{\varrho_h^{n+1}}} dx.$$

– **Étape II.1 : Prédiction de vitesse**

Trouver $\tilde{\mathbf{u}}_h^{n+1} \in \mathcal{V}_{h,0}^u$ tel que, $\forall \nu_h^u \in \mathcal{V}_{h,0}^u$,

$$\begin{aligned} \int_{\Omega} \sqrt{\varrho^{n+1}} \frac{\sqrt{\varrho^{n+1}} \tilde{\mathbf{u}}_h^{n+1} - \sqrt{\varrho^n} \mathbf{u}_h^n}{\Delta t} \cdot \nu_h^u dx + \frac{1}{2} \int_{\Omega} (\varrho^{n+1} \mathbf{u}_h^n \cdot \nabla) \tilde{\mathbf{u}}_h^{n+1} \cdot \nu_h^u - (\varrho^{n+1} \mathbf{u}_h^n \cdot \nabla) \nu_h^u \cdot \tilde{\mathbf{u}}_h^{n+1} dx \\ + \int_{\Omega} 2\eta^{n+1} D \tilde{\mathbf{u}}_h^{n+1} : D \nu_h^u dx - \int_{\Omega} \tilde{p}_h^{n+1} \operatorname{div}(\nu_h^u) dx = \int_{\Omega} \sum_{i=1}^3 \mu_{ih}^{n+1} \nabla c_{ih}^{n+1} \cdot \nu_h^u dx, \end{aligned}$$

où ϱ^{n+1} et η^{n+1} désignent respectivement $\varrho(\mathbf{c}_h^{n+1})$ et $\eta(\mathbf{c}_h^{n+1})$.

– **Étape II.2.1 : Calcul de l'incrément de pression**

Trouver $\Phi_h^{n+1} \in \mathcal{V}_h^p$ tel que, $\forall \nu_h^p \in \mathcal{V}_h^p$,

$$\int_{\Omega} \frac{1}{\varrho_h^{n+1}} \nabla \Phi_h^{n+1} \nabla \nu_h^p dx = \frac{1}{\Delta t} \int_{\Omega} \tilde{\mathbf{u}}_h^{n+1} \cdot \nabla \nu_h^p dx$$

– **Étape II.2.2 : Correction de la pression**

$$p_h^{n+1} = \tilde{p}_h^{n+1} + \Phi_h^{n+1}.$$

– **Étape II.2.3 : Correction de la vitesse**

Trouver $\mathbf{u}_h^{n+1} \in \mathcal{V}_{h,0}^{\mathbf{u}}$ tel que, $\forall \boldsymbol{\nu}_h^{\mathbf{u}} \in \mathcal{V}_{h,0}^{\mathbf{u}}$,

$$\int_{\Omega} \varrho_h^{n+1} \mathbf{u}_h^{n+1} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx = \int_{\Omega} \varrho_h^{n+1} \tilde{\mathbf{u}}_h^{n+1} \cdot \boldsymbol{\nu}_h^{\mathbf{u}} dx + \Delta t \int_{\Omega} \Phi_h^{n+1} \operatorname{div} \boldsymbol{\nu}_h^{\mathbf{u}} dx.$$

Remarque VII.9

L'étape II.0 de renormalisation de la pression, introduite initialement dans [GQ00], a deux intérêts :

- comme dans [GQ00] (cf paragraphe suivant), elle permet d'obtenir la stabilité de la méthode de projection incrémentale dans les cas où la densité est variable,
- elle permet également de donner un sens à la méthode de projection lorsque les espaces d'approximation changent d'une itération en temps à l'autre à cause de l'adaptation de maillage. En effet, l'étape de correction de pression (VII.9) : $p_h^{n+1} = p_h^n + \Phi_h^{n+1}$ n'a plus de sens clair puisqu'elle consiste à ajouter algébriquement deux fonctions discrètes qui n'appartiennent pas aux mêmes espaces d'approximation. L'introduction de l'étape II.0, formulée de manière variationnelle, permet de corriger ce problème puisqu'alors la pression prédite $\tilde{p}_h^{n+1}$ appartient à l'espace d'approximation au temps t^{n+1} .

Nous effectuons le cas test de Laplace : il s'agit de la simulation d'une bulle à l'équilibre en deux dimensions. Les paramètres du cas test sont donnés dans la table VII.1.

R	Ω	σ	ϱ_b	ϱ_l	η_b	η_l
10^{-2}	$[0, 4R] \times [0, 4R]$	4	1	1000	1.5×10^{-3}	150×10^{-3}

ε	h_{fin}	$\varepsilon/h_{\text{fin}}$	Δt	M_{deg}
$R/10$	$4R/320$	8	10^{-3}	10^{-6}

TAB. VII.1 – Les paramètres du cas test.

A l'équilibre, nous devons obtenir : $\frac{\partial c}{\partial t} = 0$ et $\mathbf{u} = 0$. Les équations donnent alors : $\mu = \text{constante}$ et $\nabla p = \mu \nabla c$. Ainsi, nous nous attendons à trouver :

$$\begin{cases} \mathbf{u} = 0, \\ p = \mu c \text{ (à une constante près)}, \\ \mu = \text{constante}. \end{cases}$$

Par ailleurs, la relation de Laplace nous donne le saut de pression attendu à l'équilibre :

$$[p] = \frac{\sigma}{R},$$

où R est le rayon de la bulle.

L'initialisation est effectuée suivant le profil d'une interface plane à l'équilibre :

$$c^0(x, y) = \frac{1}{2} - \frac{1}{2} \tanh \left(\frac{2(\sqrt{x^2 + y^2} - R)}{\varepsilon} \right).$$

Le résultat obtenu est présenté à différents instants sur les figures VII.2 et VII.3. Nous observons l'apparition de vitesses parasites au voisinage de l'interface. L'interface se déstabilise après quelques pas de temps.

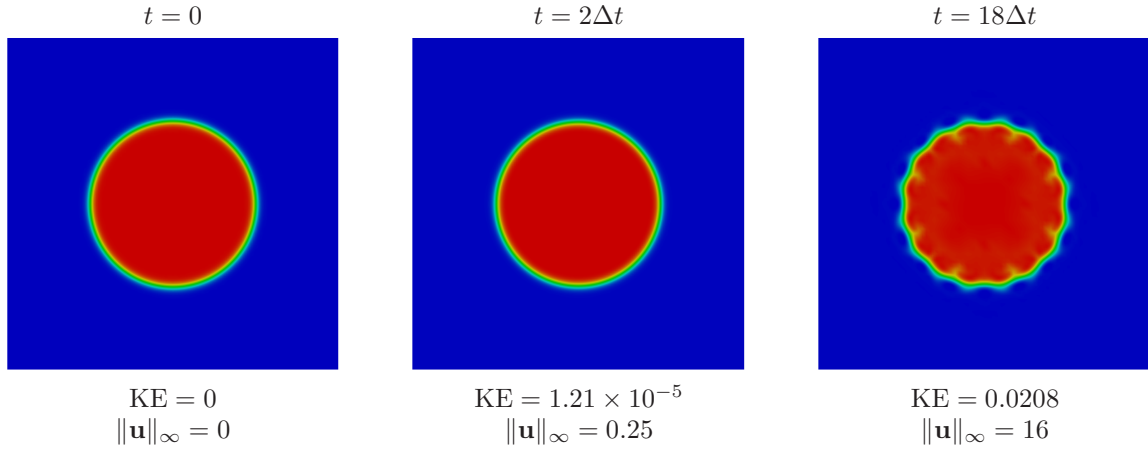
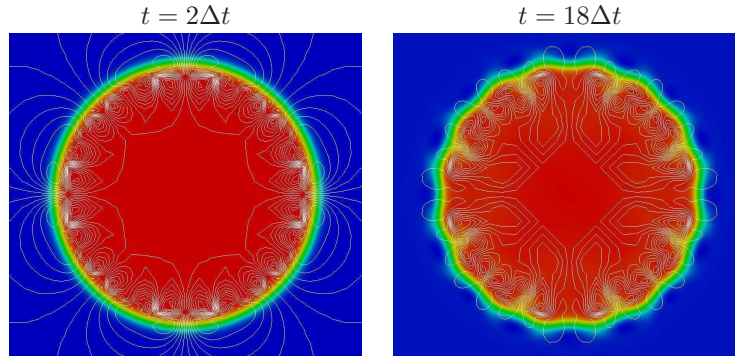
FIG. VII.2 – Paramètre d'ordre c , Energie cinétique (KE) et norme infinie de la vitesse sur le domaine.

FIG. VII.3 – Zoom et lignes de courant de la vitesse

Ce phénomène s'explique par les deux raisons suivantes :

- A l'initialisation, le profil du paramètre d'ordre choisi ne permet pas d'obtenir μ constant. Par conséquent, le second membre n'est pas un gradient. Nous avons alors aucune chance d'obtenir une solution avec une vitesse nulle.
- La méthode de projection ne permet pas de calculer les solutions $\mathbf{u}_h = 0$ et $p_h = q_h$ de l'équation de Navier-Stokes dans le cas particulier où le second membre s'écrit comme un gradient ∇q_h , $q_h \in \mathcal{V}_h^p$.

Pour le problème de l'initialisation, il est difficile d'y répondre puisque l'expression analytique du profil d'équilibre n'est connue que dans le cas d'une interface plane sur un domaine infini. Nous proposons de la chercher numériquement.

Dans le cas diphasique, supposons qu'un équilibre soit atteint et que le paramètre d'ordre soit constant loin des interfaces, notons c_0 et c_∞ les valeurs qu'il prend dans chacune des phases. Le système de Cahn-Hilliard diphasique donne alors :

$$f'(c_0) = f'(c_\infty) = \frac{\varepsilon}{12R}.$$

Cette équation polynomiale est résolue numériquement et les valeurs obtenues sont utilisées pour l'initialisation des paramètres d'ordre. Il est ensuite nécessaire en début de calcul d'effectuer quelques itérations en temps de résolution du système de Cahn-Hilliard à mobilité constante pour obtenir une solution numérique proche de la solution stationnaire.

Nous arrivons à obtenir une solution discrète telle que $\mu_{\max} - \mu_{\min} \sim \text{erreur machine}$ quel que soit le maillage.

Concernant la méthode de projection, en s'inspirant de ce qui a été présenté dans la section VII.1.2, nous proposons la variante suivante :

Problème VII.10

- **Étape I.** : Résolution du système de Cahn-Hilliard

Cette étape est inchangée (cf problème VII.8).

- **Étape II.0** : Prédiction de pression

Trouver $\tilde{p}_h^{n+1} \in \mathcal{V}_h^p$ tel que,

$$\forall \nu_h^p \in \mathcal{V}_h^p, \int_{\Omega} \frac{1}{\varrho^{n+1}} \nabla \tilde{p}_h^{n+1} \nabla \nu_h^p dx = \int_{\Omega} \frac{1}{\sqrt{\varrho^n} \sqrt{\varrho^{n+1}}} \nabla p_h^n \nabla \nu_h^p dx + \int_{\Omega} \left(\frac{\mathbf{f}^{n+1}}{\sqrt{\varrho^{n+1}}} - \frac{\mathbf{f}^n}{\sqrt{\varrho^n}} \right) \cdot \frac{\nabla \nu_h^p}{\sqrt{\varrho^{n+1}}} dx,$$

$$\text{où l'on note } \mathbf{f}^n = \sum_{i=1}^3 \mu_{ih}^n \nabla c_{ih}^n + \varrho(\mathbf{c}^{n+1}) \mathbf{g}.$$

- **Étape II.1** : Prédiction de vitesse

Trouver $\tilde{\mathbf{u}}_h^{n+1} \in \mathcal{V}_{h,0}^{\mathbf{u}}$ tel que

$$\forall \nu_h^{\mathbf{u}} \in \mathcal{V}_{h,0}^{\mathbf{u}},$$

$$\begin{aligned} \int_{\Omega} \sqrt{\varrho^{n+1}} \frac{\sqrt{\varrho^{n+1}} \tilde{\mathbf{u}}_h^{n+1} - \sqrt{\varrho^n} \mathbf{u}_h^n}{\Delta t} \cdot \nu_h^{\mathbf{u}} dx &+ \frac{1}{2} \int_{\Omega} (\varrho^{n+1} \mathbf{u}_h^n \cdot \nabla) \tilde{\mathbf{u}}_h^{n+1} \cdot \nu_h^{\mathbf{u}} - (\varrho^{n+1} \mathbf{u}_h^n \cdot \nabla) \nu_h^{\mathbf{u}} \cdot \tilde{\mathbf{u}}_h^{n+1} dx \\ &+ \int_{\Omega} 2\eta^{n+1} D\tilde{\mathbf{u}}_h^{n+1} : D\nu_h^{\mathbf{u}} dx - \int_{\Omega} \tilde{p}_h^{n+1} \operatorname{div}(\nu_h^{\mathbf{u}}) dx = \int_{\Omega} \mathbf{f}^{n+1} \cdot \nu_h^{\mathbf{u}} dx, \end{aligned}$$

où ϱ^{n+1} et η^{n+1} désignent respectivement $\varrho(c_h^{n+1})$ et $\eta(c_h^{n+1})$.

- **Étape II.2.1** : Calcul de l'incrément de pression

Trouver $\Phi_h^{n+1} \in \mathcal{V}_h^p$ tel que

$$\forall \nu_h^p \in \mathcal{V}_h^p, \int_{\Omega} \frac{1}{\varrho_h^{n+1}} \nabla \Phi_h^{n+1} \nabla \nu_h^p dx = \frac{1}{\Delta t} \int_{\Omega} \tilde{\mathbf{u}}_h^{n+1} \cdot \nabla \nu_h^p dx.$$

- **Étape II.2.2** : Correction de la pression

$$p_h^{n+1} = \tilde{p}_h^{n+1} + \Phi_h^{n+1}.$$

- **Étape II.2.3** : Correction de la vitesse

Trouver $\mathbf{u}_h^{n+1} \in \mathcal{V}_{h,0}^{\mathbf{u}}$ tel que

$$\forall \nu_h^{\mathbf{u}} \in \mathcal{V}_{h,0}^{\mathbf{u}}, \int_{\Omega} \varrho_h^{n+1} \mathbf{u}_h^{n+1} \cdot \nu_h^{\mathbf{u}} dx = \int_{\Omega} \varrho_h^{n+1} \tilde{\mathbf{u}}_h^{n+1} \cdot \nu_h^{\mathbf{u}} dx + \Delta t \int_{\Omega} \Phi_h^{n+1} \operatorname{div} \nu_h^{\mathbf{u}} dx.$$

L'avantage de cette méthode est précisément le même que dans la section VII.1.2. Supposons que $\mathbf{f}^n = \nabla q_h^n$, $\forall n \in \mathbb{N}$, et que de plus $\mathbf{u}_h^n = 0$ et $p_h^n = q_h^n$ alors nous obtenons après l'étape II.0 :

$$\tilde{p}_h^{n+1} = q_h^{n+1}.$$

Et ainsi l'étape II.1 devient :

$$\int_{\Omega} \varrho^{n+1} \frac{\tilde{\mathbf{u}}_h^{n+1}}{\Delta t} \cdot \mathbf{v}_h dx + \int_{\Omega} 2\eta^{n+1} D\tilde{\mathbf{u}}_h^{n+1} : D\mathbf{v}_h dx = 0, \quad \forall \mathbf{v}_h \in \mathcal{V}_{h,0}^{\mathbf{u}},$$

La solution est $\tilde{\mathbf{u}}_h^{n+1} = 0$ puis l'étape II.2.1 implique que $\Phi_h^{n+1} = 0$ et les étapes II.2.2 et II.2.3 donnent $p_h^{n+1} = q_h^{n+1}$ et $\mathbf{u}_h^{n+1} = 0$.

Revenons au cas test de Laplace. Les résultats obtenus avec cette variante de la méthode de projection incrémentale sont présentés sur la figure VII.2. Sur ce cas test académique, le phénomène de courants parasites est complètement éliminé.

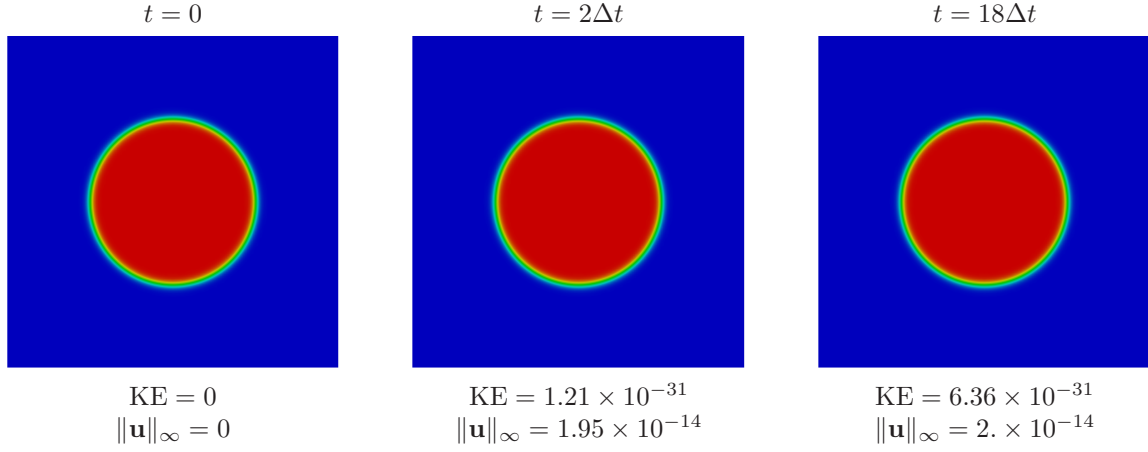


FIG. VII.4 – Paramètre d'ordre c , Energie cinétique (KE) et norme infinie de la vitesse sur le domaine, variante de la méthode de projection

Méthode de projection, schéma inconditionnellement stable

Nous utilisons maintenant la méthode de projection en lieu et place du schéma implicite (VI.3) dans le problème VI.7.

Problème VII.11

On suppose que $(c_h^n, \mathbf{u}_h^n, p_h^n) \in \mathcal{V}_h^c \times \mathcal{V}_{h,0}^u \times \mathcal{V}_h^p$ sont donnés.

– **Etape I** : Système de Cahn-Hilliard

Trouver $(\mathbf{c}_h^{n+1}, \boldsymbol{\mu}_h^{n+1}) \in (\mathcal{V}_{h,S}^c)^3 \times (\mathcal{V}_h^\mu)^3$ tels que $\forall \nu_h^c \in \mathcal{V}_h^c, \forall \nu_h^\mu \in \mathcal{V}_h^\mu$, pour $i = 1, 2$ et 3 ,

$$\left\{ \begin{array}{l} \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx - \int_{\Omega} [c_{ih}^n - \alpha_{ih}] \left[\mathbf{u}_h^n - \frac{\Delta t}{\varrho_h^n} \sum_{j=1}^3 (c_{jh}^n - \alpha_{jh}) \nabla \mu_{jh}^{n+1} \right] \cdot \nabla \nu_h^\mu dx \\ \hspace{15em} = - \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} \nabla \mu_{ih}^{n+1} \cdot \nabla \nu_h^\mu dx, \quad (\text{VII.18}) \\ \int_{\Omega} \mu_{ih}^{n+1} \nu_h^c dx = \int_{\Omega} D_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \nu_h^c dx + \int_{\Omega} \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^{n+\beta} \nabla \nu_h^c dx, \end{array} \right.$$

où α_{jh} est la constante définie par $\alpha_{jh} = \int_{\Omega} c_{jh}^0 dx$.

– **Etape II.0** : Renormalisation de la pression

Trouver $\tilde{p}_h^{n+1} \in \mathcal{V}_h^p$ tel que

$$\forall \nu_h^p \in \mathcal{V}_h^p, \quad \int_{\Omega} \frac{\nabla \tilde{p}_h^{n+1}}{\sqrt{\varrho_h^{n+1}}} \frac{\nabla \nu_h^p}{\sqrt{\varrho_h^{n+1}}} dx = \int_{\Omega} \frac{\nabla p_h^n}{\sqrt{\varrho_h^n}} \frac{\nabla \nu_h^p}{\sqrt{\varrho_h^{n+1}}} dx.$$

– **Etape II.1** : Prédiction de vitesse

Trouver $\tilde{\mathbf{u}}_h^{n+1} \in \mathcal{V}_{h,0}^u$ tel que, $\forall \boldsymbol{\nu}_h^u \in \mathcal{V}_{h,0}^u$,

$$\begin{aligned} & \int_{\Omega} \varrho_h^n \frac{\tilde{\mathbf{u}}_h^{n+1} - \mathbf{u}_h^n}{\Delta t} \cdot \boldsymbol{\nu}_h^u dx + \frac{1}{2} \int_{\Omega} \frac{\varrho_h^{n+1} - \varrho_h^n}{\Delta t} \tilde{\mathbf{u}}_h^{n+1} \cdot \boldsymbol{\nu}_h^u dx \\ & + \frac{1}{2} \int_{\Omega} (\varrho_h^{n+1} \mathbf{u}_h^n \cdot \nabla) \tilde{\mathbf{u}}_h^{n+1} \cdot \boldsymbol{\nu}_h^u - (\varrho_h^{n+1} \mathbf{u}_h^n \cdot \nabla) \boldsymbol{\nu}_h^u \cdot \tilde{\mathbf{u}}_h^{n+1} dx + \int_{\Omega} 2\eta_h^{n+1} D\tilde{\mathbf{u}}_h^{n+1} : D\boldsymbol{\nu}_h^u dx \\ & - \int_{\Omega} \tilde{p}_h^{n+1} \operatorname{div}(\boldsymbol{\nu}_h^u) dx = - \int_{\Omega} \sum_{i=1}^3 (c_{ih}^n - \alpha_{ih}) \nabla \mu_{ih}^{n+1} \cdot \boldsymbol{\nu}_h^u dx + \int_{\Omega} \varrho_h^{n+1} \mathbf{g} \cdot \boldsymbol{\nu}_h^u dx, \end{aligned}$$

où ϱ_h^{n+1} et η_h^{n+1} désignent respectivement $\varrho(c_h^{n+1})$ et $\eta(c_h^{n+1})$.

– **Etape II.2.1** : Calcul de l'incrément de pression

Trouver $\Phi_h^{n+1} \in \mathcal{V}_h^p$ tel que, $\forall \nu_h^p \in \mathcal{V}_h^p$,

$$\int_{\Omega} \frac{1}{\varrho_h^{n+1}} \nabla \Phi_h^{n+1} \nabla \nu_h^p dx = \frac{1}{\Delta t} \int_{\Omega} \tilde{\mathbf{u}}_h^{n+1} \cdot \nabla \nu_h^p dx.$$

– **Etape II.2.2** : Correction de la pression

$$p_h^{n+1} = \tilde{p}_h^{n+1} + \Phi_h^{n+1}.$$

– **Etape II.2.3** : Correction de la vitesse

Trouver $\mathbf{u}_h^{n+1} \in \mathcal{V}_{h,0}^u$ tel que, $\forall \nu_h^u \in \mathcal{V}_{h,0}^u$,

$$\int_{\Omega} \varrho_h^{n+1} \mathbf{u}_h^{n+1} \cdot \nu_h^u dx = \int_{\Omega} \varrho_h^{n+1} \tilde{\mathbf{u}}_h^{n+1} \cdot \nu_h^u dx + \Delta t \int_{\Omega} \Phi_h^{n+1} \operatorname{div} \nu_h^u dx.$$

Nous pouvons maintenant démontrer la stabilité de ce schéma en adaptant les arguments utilisés dans les démonstrations des théorèmes VI.13 et VII.3.

Théorème VII.12

Supposons que $(c_{ih}^{n+1}, \mu_{ih}^{n+1})$, $i = 1, 2, 3$ ($\mathbf{u}_h^{n+1}, \tilde{\mathbf{u}}_h^{n+1}, p_h^{n+1}$) sont solutions du problème VII.11. Alors, nous avons l'inégalité suivante :

$$\begin{aligned} & \left[\mathcal{F}_{\Sigma, \varepsilon}^{triph}(\mathbf{c}_h^{n+1}) + \frac{1}{2} \left| \sqrt{\varrho_h^{n+1}} \mathbf{u}_h^{n+1} \right|_{L^2(\Omega)}^2 + \frac{1}{2} \Delta t^2 \left| \frac{\nabla p_h^{n+1}}{\sqrt{\varrho_h^{n+1}}} \right|_{L^2(\Omega)}^2 \right] \\ & - \left[\mathcal{F}_{\Sigma, \varepsilon}^{triph}(\mathbf{c}_h^n) + \frac{1}{2} \left| \sqrt{\varrho_h^n} \mathbf{u}_h^n \right|_{L^2(\Omega)}^2 + \frac{1}{2} \Delta t^2 \left| \frac{\nabla p_h^n}{\sqrt{\varrho_h^n}} \right|_{L^2(\Omega)}^2 \right] + \Delta t \int_{\Omega} \sum_{i=1}^3 \frac{M_{0h}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx \\ & + \Delta t \int_{\Omega} 2\eta_h^{n+1} |D\tilde{\mathbf{u}}_h^{n+1}|^2 dx + \frac{3}{8} \varepsilon (2\beta - 1) \sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|_{L^2(\Omega)}^2 + \frac{1}{2} \left| \sqrt{\varrho_h^n} (\tilde{\mathbf{u}}_h^{n+1} - \mathbf{u}^*) \right|^2 \\ & + \frac{1}{2} \left| \sqrt{\varrho_h^n} (\mathbf{u}_h^* - \mathbf{u}_h^n) \right|_{L^2(\Omega)}^2 \leq \Delta t \int_{\Omega} \varrho_h^{n+1} \mathbf{g} \cdot \tilde{\mathbf{u}}_h^{n+1} dx + \frac{12}{\varepsilon} \int_{\Omega} F(\mathbf{c}_h^{n+1}) - F(\mathbf{c}_h^n) - \mathbf{d}^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \cdot (\mathbf{c}_h^{n+1} - \mathbf{c}_h^n) dx. \end{aligned}$$

Démonstration :

(i) Nous constatons que si nous posons

$$\mathbf{u}_h^* = \mathbf{u}_h^n - \frac{\Delta t}{\varrho_h^n} \sum_{j=1}^3 (c_{jh}^n - \alpha_j) \nabla \mu_{jh}^{n+1},$$

les étapes I et II.1 du problème VII.11 peuvent se réécrire de la manière suivante : $\forall \nu_h^\mu \in \mathcal{V}_h^\mu$, $\forall \nu_h^c \in \mathcal{V}_h^c$,

$$\begin{cases} \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx - \int_{\Omega} (c_{ih}^n - \alpha_{ih}) \mathbf{u}_h^* \cdot \nabla \nu_h^\mu dx = - \int_{\Omega} \frac{M_0}{\Sigma_i} \nabla \mu_{ih}^{n+1} \cdot \nabla \nu_h^\mu dx, \\ \int_{\Omega} \mu_{ih}^{n+1} \nu_h^c dx = \int_{\Omega} D_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \nu_h^c dx + \int_{\Omega} \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^{n+\beta} \nabla \nu_h^c dx, \end{cases} \quad (\text{VII.19})$$

et $\forall \nu_h^u \in \mathcal{V}_{h,0}^u$, $\forall \nu_h^p \in \mathcal{V}_h^p$,

$$\begin{aligned} & \int_{\Omega} \varrho_h^n \frac{\tilde{\mathbf{u}}_h^{n+1} - \mathbf{u}_h^*}{\Delta t} \nu_h^u dx + \frac{1}{2} \int_{\Omega} \frac{\varrho_h^{n+1} - \varrho_h^n}{\Delta t} \tilde{\mathbf{u}}_h^{n+1} \cdot \nu_h^u dx \\ & + \frac{1}{2} \int_{\Omega} (\varrho_h^{n+1} \mathbf{u}_h^{n+1} \cdot \nabla) \tilde{\mathbf{u}}_h^{n+1} \cdot \nu_h^u - (\varrho_h^{n+1} \mathbf{u}_h^n \cdot \nabla) \nu_h^u \cdot \tilde{\mathbf{u}}_h^{n+1} dx \\ & + \int_{\Omega} 2\eta_h^{n+1} D\tilde{\mathbf{u}}_h^{n+1} : D\nu_h^u dx - \int_{\Omega} \tilde{p}_h^{n+1} \operatorname{div}(\nu_h^u) dx = \int_{\Omega} \varrho_h^{n+1} \mathbf{g} \cdot \nu_h^u dx. \end{aligned}$$

- (ii) L'estimation standard sur le système de Cahn-Hilliard ($\nu_h^\mu = \mu_{ih}^{n+1}$, $\nu_h^c = \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t}$ et somme sur i des équations (VII.19)) donne alors

$$\begin{aligned} \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^{n+1}) - \mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^n) + \Delta t \int_{\Omega} \sum_{i=1}^3 \frac{M_{0h}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx \\ + \frac{3}{8} \varepsilon (2\beta - 1) \sum_{i=1}^3 \Sigma_i |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|_{L^2(\Omega)}^2 = \Delta t \int_{\Omega} \mathbf{u}_h^* \cdot \sum_{i=1}^3 (c_{ih}^{n+1} - \alpha) \nabla \mu_{ih}^{n+1} dx \\ + \frac{12}{\varepsilon} \int_{\Omega} [F(\mathbf{c}_h^{n+1}) - F(\mathbf{c}_h^n) - \mathbf{d}^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \cdot (\mathbf{c}_h^{n+1} - \mathbf{c}_h^n)] dx. \end{aligned} \quad (\text{VII.20})$$

- (iii) Par ailleurs, par définition de $\mathbf{u}_h^*$, nous avons

$$\sqrt{\varrho_h^n} \mathbf{u}_h^* = \sqrt{\varrho_h^n} \mathbf{u}_h^n - \frac{\Delta t}{\sqrt{\varrho_h^n}} \sum_{j=1}^3 (c_{jh}^n - \alpha_j) \nabla \mu_{jh}^{n+1}.$$

En multipliant cette égalité par $\sqrt{\varrho_h^n} \mathbf{u}_h^*$, puis en intégrant sur Ω nous obtenons :

$$\begin{aligned} -\Delta t \int_{\Omega} \sum_{j=1}^3 (c_{jh}^n - \alpha_j) \mathbf{u}_h^* \cdot \nabla \mu_{jh}^{n+1} dx \\ = \frac{1}{2} |\sqrt{\varrho_h^n} \mathbf{u}_h^*|_{L^2(\Omega)}^2 - \frac{1}{2} |\sqrt{\varrho_h^n} \mathbf{u}_h^n|_{L^2(\Omega)}^2 + \frac{1}{2} |\sqrt{\varrho_h^n} (\mathbf{u}_h^* - \mathbf{u}_h^n)|_{L^2(\Omega)}^2. \end{aligned} \quad (\text{VII.21})$$

- (iv) Le bilan de quantité de mouvement avec $\nu_h^u = \Delta t \tilde{\mathbf{u}}_h^{n+1}$ permet d'obtenir l'estimation :

$$\begin{aligned} \frac{1}{2} \left| \sqrt{\varrho_h^{n+1}} \tilde{\mathbf{u}}_h^{n+1} \right|_{L^2(\Omega)}^2 - \frac{1}{2} |\sqrt{\varrho_h^n} \mathbf{u}_h^*|_{L^2(\Omega)}^2 + \frac{1}{2} |\sqrt{\varrho_h^n} (\tilde{\mathbf{u}}_h^{n+1} - \mathbf{u}^*)|^2 \\ + \Delta t \int_{\Omega} 2\eta_h^{n+1} |D\tilde{\mathbf{u}}_h^{n+1}|^2 dx + \Delta t \int_{\Omega} \tilde{\mathbf{u}}_h^{n+1} \cdot \nabla \tilde{p}_h^{n+1} dx = \Delta t \int_{\Omega} \varrho_h^{n+1} \mathbf{g} \cdot \tilde{\mathbf{u}}_h^{n+1} dx. \end{aligned} \quad (\text{VII.22})$$

- (v) Comme dans la démonstration du théorème VII.3, le terme $\int_{\Omega} \tilde{\mathbf{u}}_h^{n+1} \cdot \nabla \tilde{p}_h^{n+1} dx$ est éliminé en introduisant la fonction

$$\hat{\mathbf{u}}_h = \tilde{\mathbf{u}}_h^{n+1} - \frac{\Delta t}{\varrho_h^{n+1}} \nabla \Phi_h^{n+1}.$$

L'étape du calcul d'incrément II.2.1 se réécrit de la façon suivante :

$$\forall \nu_h^p \in \mathcal{V}_h^p, \quad \int_{\Omega} \hat{\mathbf{u}}_h \cdot \nabla \nu_h^p dx = 0.$$

Il suffit alors d'élever au carré les deux membres de l'égalité

$$\sqrt{\varrho_h^{n+1}} \hat{\mathbf{u}}_h + \frac{\Delta t}{\sqrt{\varrho_h^{n+1}}} \nabla \tilde{p}_h^{n+1} = \sqrt{\varrho_h^{n+1}} \tilde{\mathbf{u}}_h^{n+1} + \frac{\Delta t}{\sqrt{\varrho_h^{n+1}}} \nabla \tilde{p}_h^{n+1},$$

pour obtenir

$$\begin{aligned} \frac{1}{2} \left| \sqrt{\varrho_h^{n+1}} \hat{\mathbf{u}}_h \right|_{L^2(\Omega)}^2 + \frac{1}{2} \Delta t^2 \left| \frac{\nabla \tilde{p}_h^{n+1}}{\sqrt{\varrho_h^{n+1}}} \right|_{L^2(\Omega)}^2 \\ = \frac{1}{2} \left| \sqrt{\varrho_h^{n+1}} \tilde{\mathbf{u}}_h^{n+1} \right|_{L^2(\Omega)}^2 + \frac{1}{2} \Delta t^2 \left| \frac{\nabla \tilde{p}_h^{n+1}}{\sqrt{\varrho_h^{n+1}}} \right|_{L^2(\Omega)}^2 + \Delta t \int_{\Omega} \tilde{\mathbf{u}}_h^{n+1} \cdot \nabla \tilde{p}_h^{n+1} dx. \end{aligned} \quad (\text{VII.23})$$

(vi) La combinaison des égalités ou inégalités (VII.20), (VII.21), (VII.22) et (VII.23) permet d'obtenir

$$\begin{aligned}
& \left[\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^{n+1}) + \frac{1}{2} \left| \sqrt{\varrho_h^{n+1}} \hat{\mathbf{u}}_h \right|_{L^2(\Omega)}^2 + \frac{1}{2} \Delta t^2 \left| \frac{\nabla p_h^{n+1}}{\sqrt{\varrho_h^{n+1}}} \right|_{L^2(\Omega)}^2 \right] \\
& - \left[\mathcal{F}_{\Sigma, \varepsilon}^{\text{triph}}(\mathbf{c}_h^n) + \frac{1}{2} \left| \sqrt{\varrho_h^n} \mathbf{u}_h^n \right|_{L^2(\Omega)}^2 + \frac{1}{2} \Delta t^2 \left| \frac{\nabla \tilde{p}_h^{n+1}}{\sqrt{\varrho_h^{n+1}}} \right|_{L^2(\Omega)}^2 \right] + \Delta t \int_{\Omega} \sum_{i=1}^3 \frac{M_{\theta h}^{n+\alpha}}{\Sigma_i} |\nabla \mu_{ih}^{n+1}|^2 dx \\
& + \Delta t \int_{\Omega} 2\eta_h^{n+1} |D\tilde{\mathbf{u}}_h^{n+1}|^2 dx + \frac{3}{8} \varepsilon (2\beta - 1) \sum_{i=1}^3 |\nabla c_{ih}^{n+1} - \nabla c_{ih}^n|_{L^2(\Omega)}^2 + \frac{1}{2} |\sqrt{\varrho_h^n} (\mathbf{u}_h^* - \mathbf{u}_h^n)|_{L^2(\Omega)}^2 \\
& + \frac{1}{2} |\sqrt{\varrho_h^n} (\tilde{\mathbf{u}}_h^{n+1} - \mathbf{u}^*)|^2 = \Delta t \int_{\Omega} \varrho_h^{n+1} \mathbf{g} \cdot \tilde{\mathbf{u}}_h^{n+1} dx + \frac{12}{\varepsilon} \int_{\Omega} F(\mathbf{c}_h^{n+1}) - F(\mathbf{c}_h^n) - \mathbf{d}^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \cdot (\mathbf{c}_h^{n+1} - \mathbf{c}_h^n) dx.
\end{aligned}$$

(vii) Enfin, nous utilisons l'étape de correction de vitesse avec $\boldsymbol{\nu}_h^{\mathbf{u}} = \mathbf{u}_h^{n+1}$ pour obtenir

$$\left| \sqrt{\varrho_h^{n+1}} \mathbf{u}_h^{n+1} \right|_{L^2(\Omega)} \leq \left| \sqrt{\varrho_h^{n+1}} \hat{\mathbf{u}}_h \right|_{L^2(\Omega)}$$

et l'étape de renormalisation de la pression avec $\nu_h^p = \tilde{p}_h^{n+1}$:

$$\left| \frac{\nabla \tilde{p}_h^{n+1}}{\sqrt{\varrho_h^{n+1}}} \right|_{L^2(\Omega)}^2 \leq \left| \frac{\nabla p_h^n}{\sqrt{\varrho_h^n}} \right|_{L^2(\Omega)}^2.$$

Ceci donne la conclusion. ■

VII.2 Éléments finis non conformes

Nous présentons, dans cette section, une étude de la méthode de projection incrémentale pour résoudre les équations de Stokes incompressibles discrétisées en espace par des éléments finis non conformes de bas degré de Rannacher-Turek [RT92] (avec en particulier, une pression constante par maille). Ce travail a été effectué en collaboration avec F. Dardalhon et J.C. Latché.

Nous considérons les équations de Stokes incompressibles instationnaires, posées sur un intervalle de temps fini $]0, T[$ et sur un ouvert Ω polygonal ou polyédral borné. Le système s'écrit :

$$\begin{cases} \frac{\partial \mathbf{u}}{\partial t} - \Delta \mathbf{u} + \nabla p = \mathbf{f}, & \text{dans }]0, T[\times \Omega, \\ \operatorname{div} \mathbf{u} = 0, & \text{dans }]0, T[\times \Omega. \end{cases} \quad (\text{VII.24})$$

La frontière Γ de Ω est partagée en deux parties $\Gamma = \Gamma_D \cup \Gamma_N$, avec $\Gamma_D \neq \emptyset$. La vitesse est prescrite sur Γ_D et des conditions de Neumann sont imposées sur Γ_N :

$$\mathbf{u} = \mathbf{u}_{\Gamma_D} \text{ sur }]0, T[\times \Gamma_D, \quad -p\mathbf{n} + \nabla \mathbf{u} \cdot \mathbf{n} = \mathbf{f}_N \text{ dans }]0, T[\times \Gamma_N. \quad (\text{VII.25})$$

Nous ajoutons finalement au système la condition initiale $\mathbf{u} = \mathbf{u}_0$ sur Ω , pour $t = 0$. Les champs de vecteurs $\mathbf{f}$, $\mathbf{u}_{\Gamma_D}$, $\mathbf{f}_N$ et $\mathbf{u}_0$ sont supposés donnés et réguliers.

Puisque la pression est approchée par des fonctions constantes par cellules, l'étape de projection doit s'écrire comme un système de Darcy (*cf* (VII.3)). Ainsi, nous choisissons d'utiliser une discrétisation "lumpée" des termes de dérivées en temps, qui nous permet d'obtenir le problème elliptique pour la pression de manière explicite.

Tout d'abord, nous montrons qu'il est possible d'écrire le schéma obtenu sous forme variationnelle grâce à des produits scalaires, des opérateurs et des normes dépendant du maillage. Ceci autorise, pour le problème discret

que nous considérons, à adapter les analyses d'erreur réalisées dans le cas semi-discret en temps [She92, Gue99] ou pour des éléments finis conformes [Gue96]. Nous obtenons ainsi, pour des conditions de Dirichlet homogènes, une estimation d'ordre 2 (par rapport au pas de temps) pour l'erreur de splitting. En outre, nous donnons l'expression explicite de l'opérateur discret appliqué à l'incrément de pression dans l'étape de projection. Cette construction apporte quelques éléments nouveaux au problème plutôt controversé (dans le cas des méthodes algébriques) des conditions aux bords artificielles en pression [GMS05] : en effet, nous montrons que nous obtenons une discrétisation de type volumes finis de l'opérateur de Laplace, avec les conditions aux bords attendues : conditions de Neumann homogène sur Γ_D et de Dirichlet homogène sur Γ_N ; cependant, puisque ces conditions aux bords sont imposées en un sens faible, nous observons que leur influence diminue lorsque le pas de temps tend vers 0 et nous retrouvons les ordres optimaux de convergence par rapport à la taille du maillage, même en norme L^∞ pour la pression dans le cas de conditions aux bords ouvertes.

Nous décrivons tout d'abord le schéma (section VII.2.1), nous donnons ensuite l'expression de l'opérateur elliptique de pression (section VII.2.2), nous donnons les estimations d'erreurs (section VII.2.3), et enfin nous décrivons quelques cas tests numériques pour illustrer notre analyse (section VII.2.4).

VII.2.1 Discrétisation

Soit $\mathcal{T}$ une décomposition du domaine Ω en quadrangles ($d = 2$) ou en hexaèdres ($d = 3$), supposée régulière au sens usuel de la littérature éléments finis. Nous notons $\mathcal{E}$ l'ensemble des faces σ du maillage ; $\mathcal{E}_{\text{ext}}$ les faces de la frontière de Ω , $\mathcal{E}_{\text{int}}$ les faces intérieures (*i.e.* $\mathcal{E} \setminus \mathcal{E}_{\text{ext}}$) et $\mathcal{E}(K)$ les faces d'une cellule donnée $K \in \mathcal{T}$. La face intérieure séparant deux cellules voisines K et L est notée $\sigma = K|L$. Pour chaque cellule $K \in \mathcal{T}$ et chaque face $\sigma \in \mathcal{E}(K)$, $\mathbf{n}_{K,\sigma}$ désigne le vecteur normal à σ sortant de K . Nous notons $|K|$ et $|\sigma|$ les mesures de la cellule K et de la face σ respectivement.

La vitesse et la pression sont discrétisées en utilisant l'élément fini de Rannacher-Turek [RT92]. L'approximation de la vitesse est ainsi non-conforme : l'espace $\mathcal{V}^{\mathbf{u}}$ est composé de fonctions discrètes discontinues à travers les arêtes mais le saut de leur intégrale le long des arêtes est nul ; les degrés de liberté sont localisés au centre des arêtes du maillage et nous choisissons la version de l'élément où ils représentent la moyenne de la vitesse à travers une arête. L'ensemble des degrés de liberté s'écrit :

$$\{\mathbf{u}_{\sigma,i}, \sigma \in \mathcal{E}, 1 \leq i \leq d\}.$$

Nous notons $\varphi_\sigma^{(i)}$ la fonction de forme vectorielle associée à $\mathbf{u}_{\sigma,i}$. Par définition, nous avons $\varphi_\sigma^{(i)} = \varphi_\sigma \mathbf{e}^{(i)}$, où φ_σ est la fonction de forme scalaire de Rannacher-Turek et $\mathbf{e}^{(i)}$ est le i^{e} vecteur de la base canonique de $\mathbb{R}^d$, et nous définissons $\mathbf{u}_\sigma$ par $\mathbf{u}_\sigma = \sum_{i=1}^d \mathbf{u}_{\sigma,i} \mathbf{e}^{(i)}$. Avec ces identités, nous avons

$$\mathbf{u} = \sum_{\sigma \in \mathcal{E}} \sum_{i=1}^d \mathbf{u}_{\sigma,i} \varphi_\sigma^{(i)}(\mathbf{x}) = \sum_{\sigma \in \mathcal{E}} \mathbf{u}_\sigma \varphi_\sigma(\mathbf{x}), \quad \text{pour p.p. } \mathbf{x} \in \Omega.$$

Soit $\mathcal{E}_D \subset \mathcal{E}_{\text{ext}}$ l'ensemble des arêtes où la vitesse est prescrite, disons $\mathbf{u} = \mathbf{u}_D$. Alors, classiquement, ces conditions de Dirichlet sont imposées dans la définition de l'espace discret :

$$\forall \sigma \in \mathcal{E}_D, \text{ pour } 1 \leq i \leq d, \quad \mathbf{u}_{\sigma,i} = \frac{1}{|\sigma|} \int_\sigma \mathbf{u}_{D,i}, \quad (\text{VII.26})$$

où $\mathbf{u}_{D,i}$ est la i^{e} composante de $\mathbf{u}_D$. Pour $\mathbf{v} \in \mathcal{V}^{\mathbf{u}}$, nous notons $\nabla_h \mathbf{v}$ et $\text{div}_h \mathbf{v}$ les fonctions de $L^2(\Omega)^{d \times d}$ et $L^2(\Omega)$ respectivement égales à $\nabla \mathbf{v}$ et $\text{div } \mathbf{v}$ presque partout dans Ω .

La pression est constante par cellule, et ses degrés de liberté sont notés p_K pour toute cellule $K \in \mathcal{T}$. Nous notons $\mathcal{V}^p$ l'espace des pressions discrètes.

Pour obtenir notre algorithme à pas fractionnaires, nous effectuons la résolution en deux étapes : nous supposons que la vitesse $\mathbf{u}^n \in \mathcal{V}^{\mathbf{u}}$ et la pression $p^n \in \mathcal{V}^p$ sont connues, nous réalisons tout d'abord une étape de prédiction pour obtenir une vitesse (à divergence non nulle) $\tilde{\mathbf{u}}^{n+1} \in \mathcal{V}^{\mathbf{u}}$, ensuite nous calculons la pression $p^{n+1} \in \mathcal{V}^p$ et la vitesse (à divergence nulle) $\mathbf{u}^{n+1} \in \mathcal{V}^{\mathbf{u}}$ dans une seconde étape. Nous obtenons, pour $0 \leq n < N$:

1 - Etape de prédiction de vitesse :

Trouver $\tilde{\mathbf{u}}^{n+1} \in \mathcal{V}^{\mathbf{u}}$ tel que (VII.26) est satisfait avec $\mathbf{u}_D = \mathbf{u}_{\Gamma_D}$ et, pour toute face $\sigma \in \mathcal{E} \setminus \mathcal{E}_D$, pour tout entier i dans $\{1, \dots, d\}$:

$$\frac{|D_\sigma|}{\Delta t} [\tilde{\mathbf{u}}_{\sigma,i}^{n+1} - \mathbf{u}_{\sigma,i}^n] + \int_\Omega \nabla_h \tilde{\mathbf{u}}^{n+1} : \nabla_h \varphi_\sigma^{(i)} - \int_\Omega p^n \text{div}_h \varphi_\sigma^{(i)} = \int_\Omega \mathbf{f}^{n+1} \cdot \varphi_\sigma^{(i)} + \int_{\Gamma_N} \mathbf{f}_N^{n+1} \cdot \varphi_\sigma^{(i)},$$

où $|D_\sigma| = \int_\Omega \varphi_\sigma$.

Le terme de gradient de pression dans cette relation peut de manière équivalente s'écrire, pour $1 \leq i \leq d$:

$$\begin{aligned} \forall \sigma \in \mathcal{E}_{\text{int}}, \sigma = K|L, \quad & \int_\Omega p^n \operatorname{div}_h \varphi_\sigma^{(i)} = |\sigma| (p_K^n - p_L^n) \mathbf{n}_{K,\sigma}^{(i)}, \\ \forall \sigma \in \mathcal{E}_{\text{ext}} \setminus \mathcal{E}_D, \sigma \in \mathcal{E}(K), \quad & \int_\Omega p^n \operatorname{div}_h \varphi_\sigma^{(i)} = |\sigma| p_K^n \mathbf{n}_{K,\sigma}^{(i)}. \end{aligned} \quad (\text{VII.27})$$

2 - Etape de projection de vitesse :

Trouver $\mathbf{u}^{n+1} \in \mathcal{V}^{\mathbf{u}}$ et p^{n+1} dans $\mathcal{V}^p$ tels que (VII.26) est satisfait avec $\mathbf{u}_D = \mathbf{u}_{\Gamma_D}$ et :

$$\begin{aligned} \forall \sigma \in \mathcal{E} \setminus \mathcal{E}_D, \text{ pour } 1 \leq i \leq d, \quad & \frac{|D_\sigma|}{\Delta t} [\mathbf{u}_{\sigma,i}^{n+1} - \tilde{\mathbf{u}}_{\sigma,i}^{n+1}] - \int_\Omega (p^{n+1} - p^n) \operatorname{div}_h \varphi_\sigma^{(i)} = 0, \\ \forall K \in \mathcal{T}, \quad & \sum_{\sigma \in \mathcal{E}(K)} |\sigma| \mathbf{u}_\sigma^{n+1} \cdot \mathbf{n}_{K,\sigma} = 0. \end{aligned} \quad (\text{VII.28})$$

En comparaison avec la version semi-discrète de l'algorithme de projection incrémentale, il peut sembler étrange que les conditions aux bords de Dirichlet (VII.26) soient imposées aux deux composantes de la vitesse $\mathbf{u}^{n+1}$. En fait, l'expression du gradient discret (VII.27) montre que, pour la discrétisation considérée ici, le gradient discret de pression sur une face σ est colinéaire à la normale et en conséquence, les vitesses tangentes aux faces (et ainsi à la frontière du domaine) sont laissées inchangées par l'étape de projection (*i.e.* elles peuvent être prescrites ou non).

VII.2.2 Opérateur elliptique discret pour la pression et conditions aux bords artificielles

Puisque le problème elliptique de pression n'est pas posé explicitement (contrairement au niveau continu), les conditions aux bords associées ne le sont pas non plus. Nous allons montrer que ces conditions aux bords sont retrouvées lorsque l'on calcule l'opérateur.

Nous multiplions la première équation de l'étape de projection de vitesse (VII.28) par $\frac{1}{|D_\sigma|} |\sigma| \mathbf{n}_{K,\sigma}^{(i)}$ et nous sommions les équations obtenues pour $1 \leq i \leq d$ et $\sigma \in \mathcal{E}(K)$. Nous avons, pour tout $K \in \mathcal{T}$:

$$\sum_{\sigma \in \mathcal{E}(K), \sigma = K|L} \frac{|\sigma|^2}{|D_\sigma|} [\phi_K^{n+1} - \phi_L^{n+1}] + \sum_{\sigma \in (\mathcal{E}_{\text{ext}} \setminus \mathcal{E}_D) \cap \mathcal{E}(K)} \frac{|\sigma|^2}{|D_\sigma|} \phi_K^{n+1} = \frac{1}{\Delta t} \sum_{\sigma \in \mathcal{E}(K)} |\sigma| \tilde{\mathbf{u}}_\sigma^{n+1} \cdot \mathbf{n}_{K,\sigma},$$

où nous avons posé $\phi_K^{n+1} = p_K^{n+1} - p_K^n$, $\forall K \in \mathcal{T}$. Nous reconnaissons dans le membre de gauche de cette relation une approximation de l'opérateur de Laplace de type volumes finis, cependant inconsistante puisque sur un maillage uniforme nous pouvons facilement montrer que le coefficient $|\sigma|^2/|D_\sigma|$ est d fois plus grand que celui des schémas volumes finis. Ceci est probablement relié au fait que les éléments de Rannacher-Turek sont connus pour donner des approximations inconsistantes au problème de Darcy. Les conditions aux bords artificielles (celles de l'algorithme semi-discrét en temps), *i.e.* conditions de type Neumann homogènes sur tout $\sigma \in \mathcal{E}_D$ et des conditions aux bords de type Dirichlet homogènes sur tout $\sigma \in \mathcal{E}_{\text{ext}} \setminus \mathcal{E}_D$, font partie intégrante de l'opérateur. Cependant, sur Γ_N , les conditions sont imposées en un sens plus faible que pour les approximations conformes où les degrés de liberté de pression sont sur la frontière (dans ce dernier cas, les incréments de pression sur la frontière sont exactement mis à zero). Ceci laisse espérer que la condition au bord sera relaxée lorsque le pas de temps tend vers 0.

VII.2.3 Formulation variationnelle et estimation d'erreur

Nous supposons dans cette section que les conditions aux bords sont de type Dirichlet homogènes sur l'intégralité de la frontière, *i.e.* $\Gamma_N = \emptyset$ et $\mathbf{u}_{\Gamma_D} = \mathbf{u}_{\Gamma_D} = 0$. Nous considérons le schéma implicite comme schéma de référence et notons $(\bar{\mathbf{u}}^n, \bar{p}^n) \in \mathcal{V}^{\mathbf{u}} \times \mathcal{V}^p$, $1 \leq n \leq N$, sa solution. Nous définissons l'erreur de splitting par

$$\mathbf{e}_h^n = \mathbf{u}_h^n - \bar{\mathbf{u}}_h^n \in \mathcal{V}^{\mathbf{u}}, \quad \tilde{\mathbf{e}}_h^n = \tilde{\mathbf{u}}_h^n - \bar{\mathbf{u}}_h^n \in \mathcal{V}^{\mathbf{u}}, \quad \text{et} \quad \epsilon_h^n = p_h^n - \bar{p}_h^n \in \mathcal{V}^p.$$

En prenant la différence des équations des deux schémas, en sommant sur $\sigma \in \mathcal{E} \setminus \mathcal{E}_D$ et $1 \leq i \leq d$, nous déduisons la formulation variationnelle discrète :

$$\begin{aligned}
& \forall (\mathbf{v}, q) \in \mathcal{V}^{\mathbf{u}} \times \mathcal{V}^p, \\
& \sum_{\sigma \in \mathcal{E}} \frac{|D_\sigma|}{\Delta t} [\tilde{\mathbf{e}}_\sigma^{n+1} - \mathbf{e}_\sigma^n] \cdot \mathbf{v}_\sigma + \int_\Omega \nabla_h \tilde{\mathbf{e}}^{n+1} : \nabla_h \mathbf{v} - \int_\Omega \epsilon^n \operatorname{div}_h \mathbf{v} = \int_\Omega (\bar{p}^{n+1} - \bar{p}^n) \operatorname{div}_h \mathbf{v}, \\
& \sum_{\sigma \in \mathcal{E}} \frac{|D_\sigma|}{\Delta t} [\mathbf{e}_\sigma^{n+1} - \tilde{\mathbf{e}}_\sigma^{n+1}] \cdot \mathbf{v}_\sigma - \int_\Omega (\epsilon^{n+1} - \epsilon^n) \operatorname{div}_h \mathbf{v} = \int_\Omega (\bar{p}^{n+1} - \bar{p}^n) \operatorname{div}_h \mathbf{v}, \\
& \int_\Omega q \operatorname{div}_h \mathbf{e}^{n+1} = 0,
\end{aligned} \tag{VII.29}$$

où nous avons supposé que toutes les fonctions de $\mathcal{V}^{\mathbf{u}}$ sont nulles sur la frontière.

Nous définissons maintenant les produits scalaires, normes et semi-normes discrets :

$$\begin{aligned}
& \forall (\mathbf{u}, \mathbf{v}) \in (\mathcal{V}^{\mathbf{u}})^2, \quad (\mathbf{u}, \mathbf{v})_h = \sum_{\sigma \in \mathcal{E}} |D_\sigma| \mathbf{u}_\sigma \cdot \mathbf{v}_\sigma, \quad \|\mathbf{u}\|_{0,h}^2 = (\mathbf{u}, \mathbf{u})_h, \\
& \forall (p, q) \in (\mathcal{V}^p)^2, \quad \langle p, q \rangle_h = \sum_{\substack{\sigma \in \mathcal{E}_{\text{int}} \\ (\sigma=K|L)}} \frac{|\sigma|^2}{|D_\sigma|} (p_K - p_L) (q_K - q_L), \quad \|p\|_{1,h}^2 = \langle p, p \rangle_h.
\end{aligned}$$

Avec ces notations, nous retrouvons la structure utilisée dans l'analyse des schémas de correction de pression [Gue99], et en conséquence, avec quelques adaptations techniques (en particulier, la preuve de quelques propriétés de l'inverse de l'opérateur de Stokes discret, en supposant que le problème de Stokes est régularisant, ce qui dans notre cas se résume à la convexité du domaine), nous sommes capables de dériver une estimation d'ordre 2 pour la vitesse et d'ordre 1 pour la pression (par rapport au pas de temps). Sous des hypothèses de régularité des solutions du schéma implicite et l'effet régularisant du problème de Stokes, nous prouvons les résultats suivants dans le cas de conditions aux bord de Dirichlet sur $\partial\Omega$.

Théorème VII.13

Supposons que le problème implicite est régulier, au sens où il existe $C > 0$ tel que, pour $1 \leq n \leq N - 1$:

$$|\bar{p}^{n+1} - 2\bar{p}^n + \bar{p}^{n-1}|_{1,h} \leq C\Delta t^2, \quad \sum_{k=1}^n |\bar{p}^{k+1} - \bar{p}^k|_{1,h}^2 \leq C\Delta t,$$

qui signifie que la dérivée seconde par rapport au temps du gradient de pression est uniformément bornée.

Alors, il existe $c > 0$ tel que, pour $1 \leq n \leq N$:

$$\left(\sum_{k=0}^n \Delta t \|\mathbf{e}_h^k\|_{0,h}^2 \right)^{1/2} + \left(\sum_{k=0}^n \Delta t \|\tilde{\mathbf{e}}_h^k\|_{0,h}^2 \right)^{1/2} + \Delta t \left(\sum_{k=0}^n \Delta t \|\epsilon_h^k\|_0^2 \right)^{1/2} \leq c\Delta t^2.$$

Nous ne donnons pas ici la preuve de ce théorème, nous renvoyons pour cela à [DLM10a].

VII.2.4 Tests numériques

Nous présentons un problème avec des conditions aux bords ouvertes.

Le domaine de calcul Ω est le carré unité $]0, 1[^2$, avec Γ_N égal au côté vertical gauche et en conséquence $\Gamma_D = \partial\Omega \setminus \Gamma_N$ égal à l'union des trois autres côtés. Nous calculons le terme source $\mathbf{f}$ de manière à ce que la solution exacte $(\mathbf{u}, p)$ soit :

$$\mathbf{u}(x, y, t) = \begin{bmatrix} \sin(x) \sin(y + t) \\ \cos(x) \cos(y + t) \end{bmatrix}, \quad p(x, y, t) = \cos(x) \sin(y + t),$$

i.e.

$$\mathbf{f} = \begin{bmatrix} \sin(x)(\cos(y + t) + \sin(y + t)) \\ \cos(x)(3 \cos(y + t) + \sin(y + t)) \end{bmatrix}.$$

La condition initiale et les conditions aux bords sont données par la solution exacte, en particulier nous avons sur Γ_N :

$$\nabla \mathbf{u} \cdot \mathbf{n} - p\mathbf{n} = 0.$$

Nous traçons sur la figure VII.5 l'erreur numérique en fonction du pas de temps. Cette erreur est mesurée en norme L^2 et calculée à un temps fixe, pour 20×20 , 40×40 et 80×80 maillages uniformes structurés. Les erreurs diminuent dans un premier temps avec le pas de temps jusqu'à atteindre un plateau qui correspond à l'erreur résiduelle en espace. L'ordre de convergence en temps est proche de 2 pour la vitesse (pente de la courbe de gauche) et 1 pour la pression (à droite). Cela peut être surprenant puisque nous utilisons seulement un schéma d'Euler implicite d'ordre 1, mais peut être expliqué par le fait que l'erreur est essentiellement due au "splitting" en temps. Sur le plateau, nous observons une convergence d'ordre 2 en espace pour la vitesse et d'ordre 1 pour la pression, ce qui correspond à l'ordre optimal de convergence pour l'approximation de bas degré que nous utilisons. En comparaison, dans [JLL⁺06], les auteurs observent seulement, pour une approximation de Taylor-Hood (*i.e.* $\mathbb{P}_2 - \mathbb{P}_1$), de l'ordre 1 pour la vitesse et de l'ordre 1/2 pour la pression. En conséquence, pour la vitesse les calculs présentés ici deviennent déjà plus précis pour le maillage 80×80 .

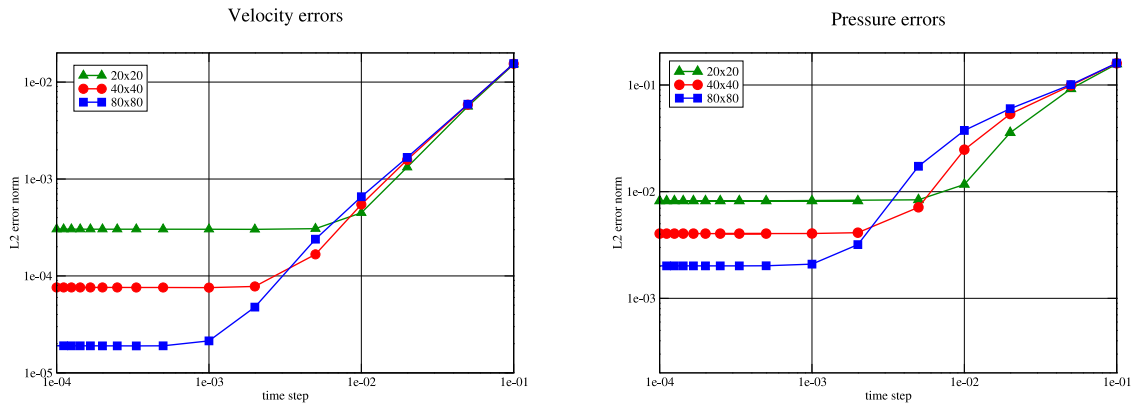


FIG. VII.5 – Erreur d'approximation (en norme L^2) en fonction de Δt .

Partie 3

Expérimentations numériques

Dans cette partie, nous décrivons les résultats de diverses expérimentations numériques. Nous présentons à la fois des simulations diphasiques et triphasiques, bidimensionnelles et tridimensionnelles. Les objectifs visés sont les suivants :

- comparer les différents schémas proposés dans la partie 2 sur des tests moins académiques que ceux donnés précédemment,
- étudier numériquement l'influence de certains paramètres du modèle (mobilité, épaisseur d'interface) ou des schémas numériques (pas de temps, pas d'espace),
- illustrer quelques possibilités de l'ensemble de la méthode de simulation (modèle et schémas numériques).

Nous présentons les tests suivants :

- Benchmark [HTK⁺09] : il s'agit de la simulation diphasique d'une bulle (se déformant peu) montant dans un liquide sous l'action de la gravité. Nous comparons les résultats que nous obtenons aux solutions de référence données dans le benchmark. Nous donnons également une étude de l'influence de l'épaisseur d'interface et de la valeur du coefficient de mobilité.
- Nous donnons une série de simulations (toujours diphasiques) d'une bulle de gaz montant dans une colonne de liquide pour une large gamme de paramètres physiques. Ces simulations reprennent les paramètres de celles effectuées dans [BM07] avec lesquelles une comparaison est donc possible.
- Nous considérons la simulation de la traversée d'une interface liquide-liquide par une bulle de gaz. Sur ce cas triphasique, nous reprenons la comparaison des diverses méthodes numériques présentées dans la partie 2, nous illustrons l'influence des différents paramètres (épaisseur d'interface, pas de temps et d'espace) et enfin nous donnons des statistiques sur les nombres d'itérations effectuées par les solveurs linéaires itératifs et par la méthode de Newton.
- Nous illustrons les possibilités de la méthode par la présentation de deux cas tests tridimensionnels : un cas test diphasique simulant la montée de trois bulles dans une colonne de liquide, un cas test triphasique simulant la traversée d'une interface liquide-liquide par une bulle de gaz (avec entrainement du liquide lourd dans le liquide léger).

Tous les développements théoriques présentés dans les parties précédentes n'ont pu être testés. Par soucis de clarté et de précision, nous profitons de cette introduction pour faire une synthèse des schémas numériques utilisés dans cette partie (en renvoyant le lecteur aux chapitres concernés pour une description plus précise). Ceci nous permet également de définir des notations qui nous permettront d'identifier ces différents schémas dans la suite de cette partie.

Modèle

Rappelons le modèle (*cf* section IV) dont nous souhaitons approcher numériquement les solutions :

$$\left\{ \begin{array}{l} \frac{\partial c_i}{\partial t} + \mathbf{u} \cdot \nabla c_i = \operatorname{div} \left(\frac{M_0}{\Sigma_i} \nabla \mu_i \right), \quad \forall i = 1, 2, 3, \\ \mu_i = \frac{4\Sigma_T}{\varepsilon} \sum_{j \neq i} \left(\frac{1}{\Sigma_j} (\partial_i F(\mathbf{c}) - \partial_j F(\mathbf{c})) \right) - \frac{3}{4} \varepsilon \Sigma_i \Delta c_i, \quad \forall i = 1, 2, 3, \\ \sqrt{\varrho(\mathbf{c})} \frac{\partial}{\partial t} (\sqrt{\varrho(\mathbf{c})} \mathbf{u}) + (\varrho(\mathbf{c}) \mathbf{u} \cdot \nabla) \mathbf{u} + \frac{\mathbf{u}}{2} \operatorname{div} (\varrho(\mathbf{c}) \mathbf{u}) - \operatorname{div} (2\eta(\mathbf{c}) D(\mathbf{u})) + \nabla p = \sum_{i=1}^3 \mu_i \nabla c_i + \varrho(\mathbf{c}) \mathbf{g}, \\ \operatorname{div} \mathbf{u} = 0, \end{array} \right.$$

où le vecteur $\mathbf{g}$ représente la gravité, la densité ϱ et la viscosité η sont des fonctions régulières de $\mathbf{c}$ définies par la formule (IV.29), les coefficients Σ_i sont définis par l'équation (IV.12) et Σ_T est défini dans la formule (IV.10). Nous ne donnons pas de simulations dans le cas d'étalement total. Nous choisissons donc $F = F_0$ (*cf* (IV.21)). Enfin, nous utilisons une mobilité dégénérée donnée par l'expression suivante :

$$M_0(\mathbf{c}) = M_{\text{deg}} \prod_{i=1}^3 (1 - c_i)^2.$$

Les conditions aux bords sont de type Neumann pour les paramètres d'ordre et les potentiels chimiques pour toutes les simulations. Pour le système de Navier-Stokes, nous utilisons des conditions de non-glissement ($\mathbf{u} = 0$) ou des conditions de type glissement ($\mathbf{u} \cdot \mathbf{n} = 0$ et $[2\eta D\mathbf{u} \cdot \mathbf{n} - p\mathbf{n}] \cdot \boldsymbol{\tau} = 0$) selon les différents cas test.

Discrétisation en espace

La discrétisation en espace repose sur la méthode de raffinement local décrite dans la partie 1. Nous utilisons des espaces d'approximation multiniveaux associés à :

- des éléments finis $\mathbb{Q}_1$ pour les paramètres d'ordre c_i , les potentiels chimiques μ_i et la pression p .
- des éléments finis $\mathbb{Q}_2$ pour la vitesse $\mathbf{u}$.

Le critère de raffinement utilisé est le même pour toutes les simulations. Il est basé sur la position de l'interface indiquée par les valeurs du paramètre d'ordre. A une itération en temps n donnée, pour chaque cellule active K , nous définissons l'indicateur (par cellule) suivant :

$$\eta_K = \max\left(\frac{1}{|K|} \int_K c_1^n, \frac{1}{|K|} \int_K c_2^n, \frac{1}{|K|} \int_K c_3^n \right).$$

La définition de cet indicateur par cellule peut-être interprétée comme suit :

- $\eta_K = 1$ signifie que la cellule K est entièrement contenue dans une phase pure.
- $\eta_K < 1$ signifie que la cellule K contient une partie d'interface.

Nous déduisons de cet indicateur par cellule un critère pour décider si une fonction de base donnée doit être (dé)raffinée par l'intermédiaire d'un indicateur par fonction de base. A une itération en temps n donnée, pour une fonction de base φ , nous définissons l'indicateur (par fonction de base) suivant :

$$\eta_\varphi = \frac{1}{|\text{supp}[\varphi]|} \sum_{\substack{K \in \mathcal{T}, \\ K \cap \text{supp}[\varphi] \neq \emptyset}} |K| \eta_K,$$

où $\mathcal{T}$ représente l'ensemble des cellules actives.

Critère VII.14 (Critère de (dé)raffinement)

Etant donnée une taille de cellule $h_{\text{interface}}$, les deux critères suivants nous permettent de décider si une fonction de base φ doit être raffinée ou non.

- Critère de raffinement :

$$\eta_\varphi < 0.90 \quad \text{et} \quad \text{diam}(K) > h_{\text{interface}} \quad \text{pour au moins une cellule } K \subset \text{supp}[\varphi].$$

- Critère de déraffinement :

$$\eta_\varphi > 0.95.$$

Discrétisation en temps

Les résolutions des systèmes de Cahn-Hilliard et de Navier-Stokes discrets sont complètement découplées par l'utilisation d'une vitesse explicite (au temps t^n) dans le terme de transport des équations de Cahn-Hilliard.

Schémas pour le système de Cahn-Hilliard Trouver $(\mathbf{c}_h^{n+1}, \boldsymbol{\mu}_h^{n+1}) \in (\mathcal{V}_{h,S}^c)^3 \times (\mathcal{V}_h^\mu)^3$ tels que $\forall \nu_h^c \in \mathcal{V}_h^c$, $\forall \nu_h^\mu \in \mathcal{V}_h^\mu$, pour $i = 1, 2$,

$$\left\{ \begin{array}{l} \int_{\Omega} \frac{c_{ih}^{n+1} - c_{ih}^n}{\Delta t} \nu_h^\mu dx - \int_{\Omega} \nu_h^\mu \mathbf{u}_h^n \cdot \nabla c_{ih}^{n+1} dx \\ \quad = - \int_{\Omega} \frac{M_{0h}^{n+\alpha}}{\Sigma_i} \nabla \mu_{ih}^{n+1} \cdot \nabla \nu_h^\mu dx - m \int_{\Omega} \frac{|M_0|_\infty - M_{0h}^{n+\alpha}}{\Sigma_i} (\nabla \mu_{ih}^{n+1} - \nabla \mu_{ih}^n) \cdot \nabla \nu_h^\mu dx, \quad (\text{VII.30}) \\ \int_{\Omega} \mu_{ih}^{n+1} \nu_h^c dx = \int_{\Omega} D_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1}) \nu_h^c dx + \int_{\Omega} \frac{3}{4} \Sigma_i \varepsilon \nabla c_{ih}^{n+\beta} \nabla \nu_h^c dx, \end{array} \right.$$

Ce type de schéma a été décrit dans le chapitre V. L'ajout du dernier terme dans la première équation ci-dessus permet de formuler dans une même écriture le schéma (V.8) ($m = 0$) et sa variante (V.53) ($m = 1$) que nous utilisons lorsque la mobilité est dégénérée. Dans la suite de cette partie, nous utilisons les notations suivantes :

- Discrétisation des termes non linéaires (choix de $D_i^F(\mathbf{c}_h^n, \mathbf{c}_h^{n+1})$) : nous utilisons les trois schémas présentés dans la section V.2. Nous adoptons les mêmes notations que dans la section V.3, à savoir :
- Impl. pour la discrétisation implicite (V.32).
- CC. pour la discrétisation convexe-concave (V.39).

- SImpl. pour la discrétisation semi-implicite (V.44).
- Discrétisation du terme capillaire $c_{ih}^{n+\beta}$: lorsqu'elle est différente de 1, la valeur du paramètre β est précisée entre parenthèses à la suite des noms précédemment introduits Impl., SImpl. ou CC.
- Discrétisation de la mobilité : Le paramètre α est toujours choisi égal à 0. Ceci signifie qu'en pratique nous discrétisons toujours la mobilité de manière explicite. Lorsque nous utilisons la valeur $m = 1$, ceci sera noté par un "m" entre parenthèses ; dans le cas contraire, nous utilisons $m = 0$.

Par exemple, SImpl(m,0.5) signifie que nous utilisons le schéma semi-implicite avec $\beta = 0.5$ et $m = 1$.

Notons que même dans les situations diphasiques, nous utilisons le schéma triphasique. Ceci est possible grâce aux propriétés de consistance.

Schémas pour le système de Navier-Stokes Pour le système de Navier-Stokes, nous utilisons deux types de méthodes de résolution : une méthode de type Lagrangien augmenté et la méthode de projection incrémentale.

Nous ne décrivons pas la méthode de Lagrangien augmenté, il s'agit d'une méthode de résolution itérative du système de Navier-Stokes. L'algorithme que nous utilisons a entièrement été décrit dans la thèse [Lap06, section III.3.2.a] à laquelle nous renvoyons (des références y sont aussi disponibles). Cette méthode comporte un paramètre d'augmentation r qui pour l'ensemble des simulations présentées dans la suite est fixé à $r = 50000$.

La méthode de projection incrémentale a été décrite dans la section VII.1.3. Nous considérons les deux variantes données par les problèmes VII.8 et VII.10.

Algorithme complet de résolution

La phase d'initialisation est réalisée en trois étapes : tout d'abord le maillage initial est créé, les champs discrets sont ensuite initialisés en fonction des valeurs données dans le jeu de données, et enfin un cycle "adaptation + initialisation des nouveaux degrés de liberté" est répété jusqu'à que le critère de raffinement soit respecté par la donnée initiale ou qu'un nombre maximum d'itérations (donné dans le jeu de données) soit atteint. L'initialisation des nouveaux degrés de liberté se fait encore une fois en utilisant la formule donnée dans le jeu de données. Les étapes de la procédure d'adaptation sont décrites en détail dans l'algorithme III.5.

Ensuite, les cycles de la marche en temps commencent. Les opérations suivantes sont effectuées (dans cet ordre) : adaptation, initialisation des nouveaux degrés de liberté, résolution du système de Cahn-Hilliard, résolution du système de Navier-Stokes, copie de champs discrets implicites dans les champs discrets explicites, suppression des mailles devenues inactives, sauvegardes (éventuelles) pour le post-traitement des résultats.

Chapitre VIII

Etude de configurations diphasiques

VIII.1 Benchmark : une bulle immergée dans un liquide

Ce chapitre présente l'étude d'un cas test proposé comme benchmark dans [HTK⁺09]. Il s'agit de la simulation bidimensionnelle diphasique d'une bulle montant dans une colonne de liquide sous l'effet de la gravité. Les propriétés des fluides en présence ne sont pas réalistes, l'objectif est uniquement de comparer entre eux les résultats obtenus par différents codes de calcul (*i.e.* différents types de modélisation ou différents schémas numériques).

Les points de comparaison portent sur trois quantités : la circularité (rapport du périmètre de la bulle à celui d'une bulle circulaire de même volume), la position du centre de masse et la vitesse moyenne d'évolution de la bulle.

Trois groupes (notés G1, G2 et G3 dans la suite) participent initialement à ce benchmark mais d'autres résultats sont également disponibles dans [vTM09]. Les groupes G1 et G2 utilisent une méthode de type level-set pour décrire l'interface et une représentation volumique du terme de force capillaire alors que le groupe G3 adopte une approche ALE (arbitrary Lagrangian-Eulerian moving grid) et introduit l'opérateur de Laplace-Beltrami pour prendre en compte les effets de tension de surface.

Le groupe G1 utilise une discrétisation éléments finis de type Rannacher-Turek [RT92] pour la vitesse et la pression, une discrétisation éléments finis conformes $\mathbb{Q}_1$ pour la fonction level-set et une stratégie à pas fractionnaires pour la discrétisation en temps. Le groupe G2 utilise des éléments finis de type $\mathbb{P}_1 - \text{iso}\mathbb{P}_2$ pour la vitesse et $\mathbb{P}_1$ pour la pression ; la résolution de l'équation de transport de la fonction level set étant stabilisée en utilisant la sous-grille associée aux éléments $\mathbb{P}_1 - \text{iso}\mathbb{P}_2$. Le groupe G3 discrétise les composantes de la vitesse (sur des grilles triangulaires) dans des espaces d'approximation engendrés par des fonctions de base quadratiques enrichis de fonctions de base cubiques et la pression par des fonctions de base discontinues linéaires par éléments.

Les résultats obtenus par les trois groupes sont très similaires et déterminent ainsi des valeurs de référence pour les trois quantités précitées.

Nous utilisons ici ces solutions de référence afin d'étudier l'influence des paramètres "non-objectifs" du modèle Cahn-Hilliard/Navier-Stokes (IV.5) à savoir l'épaisseur d'interface ε et l'ordre de grandeur de la mobilité M_{deg} . De cette étude se dégagent deux conclusions principales :

- la valeur de l'épaisseur d'interface ε influe peu sur les résultats obtenus,
- l'ordre de grandeur de la mobilité peut considérablement changer les résultats obtenus mais il existe une plage de mobilité pour laquelle ceux-ci restent très similaires. Par ailleurs, les résultats obtenus pour ces valeurs de M_{deg} sont très proches (*cf* section VIII.1.3) des résultats de référence obtenus dans [vTM09].

VIII.1.1 Définition du cas test

La configuration initiale du cas test est présentée sur la figure VIII.1. Il s'agit d'une bulle bidimensionnelle immergée dans une colonne rectangulaire de liquide.

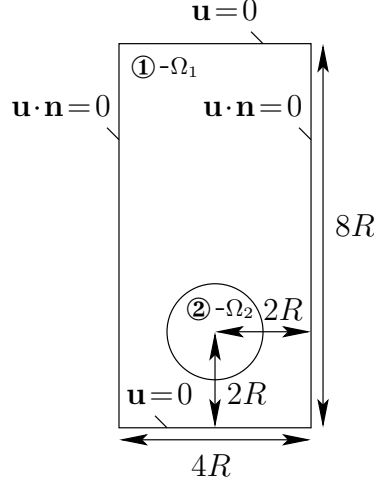


FIG. VIII.1 – Configuration initiale

Le domaine de calcul est $] -0.5, 0.5[\times]0, 2[$ et le temps final est $T = 3$. La bulle est de diamètre 0.5 et son centre est placé, à l'instant initial, à la position $\mathbf{x}_c = (x_c, y_c) = (0, 0.5)$. Pour pouvoir désigner plus facilement les deux phases (notamment dans la donnée de paramètres physiques), nous les numérotons : la phase constituant la bulle est désignée comme phase 2 et nous notons Ω_2 le volume occupé par la bulle, l'autre phase sera alors désignée comme phase 1 et nous notons $\Omega_1 = \Omega \setminus \Omega_2$. Les propriétés des fluides en présence sont indiquées dans la table VIII.1.

R	ϱ_1	ϱ_2	η_1	η_2	σ	$ \mathbf{g} $	T
0.25	1000	1	10	1	24.5	0.98	3

TAB. VIII.1 – Paramètres physiques

Les conditions aux bords sont également décrites sur la figure VIII.1 : une condition de non-glissement ($\mathbf{u} = 0$) est imposée sur les frontières horizontales du domaine alors qu'une condition de glissement est imposée sur les frontières verticales ($\mathbf{u} \cdot \mathbf{n} = 0$, $[2\eta D\mathbf{u} \cdot \mathbf{n} - p\mathbf{n}] \cdot \boldsymbol{\tau} = 0$). Rappelons que nous avons choisi d'imposer des conditions aux bords de type Neumann pour le paramètre d'ordre et le potentiel chimique sur tout le bord du domaine.

VIII.1.2 Quantités du benchmark

Cette section rappelle les définitions des trois quantités qui feront l'objet des comparaisons dans la section VIII.1.3 : la position du centre de masse, la vitesse moyenne de montée de la bulle et la circularité ; puis décrit les formules utilisées pour approcher ces quantités au niveau discret.

Au cours de chaque pas de temps, les calculs des valeurs moyennes (position du centre de gravité et vitesse moyenne) se font sur le domaine Ω_2 modélisant la bulle. Il faut alors adopter une représentation discrète de ce domaine. Nous le décrivons par un ensemble de cellules, noté B , défini, au pas de temps $n + 1$, par :

$$B = \left\{ K \in \mathcal{T}; \int_K c^{n+1} dx \geq \underline{c}|K| \right\},$$

où c^{n+1} est le paramètre d'ordre associé à la bulle à l'instant t^{n+1} , $\mathcal{T}$ l'ensemble des cellules actives (*cf* section II.21) à l'instant t^{n+1} et $\underline{c}$ un seuil fixé par l'utilisateur. Nous utilisons, dans les calculs présentés ci-dessous, la valeur $\underline{c} = 0.5$.

Au niveau discret, le volume de la bulle est alors approché par la quantité V_b définie de la manière suivante :

$$V_b = \sum_{K \in B} |K|.$$

Position du centre de masse

La position du centre de masse est définie par la formule :

$$\mathbf{x}_c = (x_c, y_c) = \frac{\int_{\Omega_2} \mathbf{x} dx}{\int_{\Omega_2} 1 dx}.$$

Nous approchons, au niveau discret, cette quantité par :

$$\mathbf{x}_c \sim \frac{1}{V_b} \sum_{K \in B} |K| \mathbf{x}_K,$$

où $\mathbf{x}_K$ est le centre de gravité de K .

Vitesse moyenne de montée de bulle

La vitesse moyenne de la bulle est définie par la formule suivante :

$$\mathbf{u}_c = \frac{\int_{\Omega_2} \mathbf{u} dx}{\int_{\Omega_2} 1 dx}.$$

Au niveau discret, nous utilisons la formule :

$$\mathbf{u}_c \sim \frac{1}{V_b} \sum_{K \in B} |K| \int_K \mathbf{u} dx.$$

Circularité

La circularité est définie comme :

$$\mathfrak{c} = \frac{P_a}{P_b} = \frac{\text{périmètre du disque d'aire équivalente}}{\text{périmètre de la bulle}},$$

où le réel P_a désigne le périmètre du disque ayant une aire égale à celle de la bulle et le réel P_b désigne le périmètre de la bulle.

Lorsque la bulle est parfaitement circulaire, la circularité est égale à 1. Dans le cas contraire, elle est strictement inférieure à 1. Au niveau discret, nous approchons cette quantité en utilisant la formule : (*cf* [vTM09]) :

$$P_a \sim \int_{\Omega} |\nabla c| dx. \quad (\text{VIII.1})$$

Le valeur de P_b peut être facilement déduite en fonction de la valeur que nous donnons à l'initialisation du paramètre d'ordre.

VIII.1.3 Paramètres numériques et résultats

Nous utilisons un maillage initial (*i.e.* avant raffinement) constitué de 4×8 cellules. Le raffinement est ensuite effectué selon le critère VII.14 avec $h_{\text{interface}} = \frac{\varepsilon}{2}$. Ceci permet d'avoir entre trois et quatre mailles dans l'interface. A titre d'exemple, la figure VIII.2 montre le maillage utilisé pour calculer la solution discrète à l'instant final $T = 3$ ainsi que l'allure du paramètre d'ordre obtenu. Le pas de temps (identique dans toutes les simulations effectuées) est fixé à $\Delta t = 10^{-3}$.

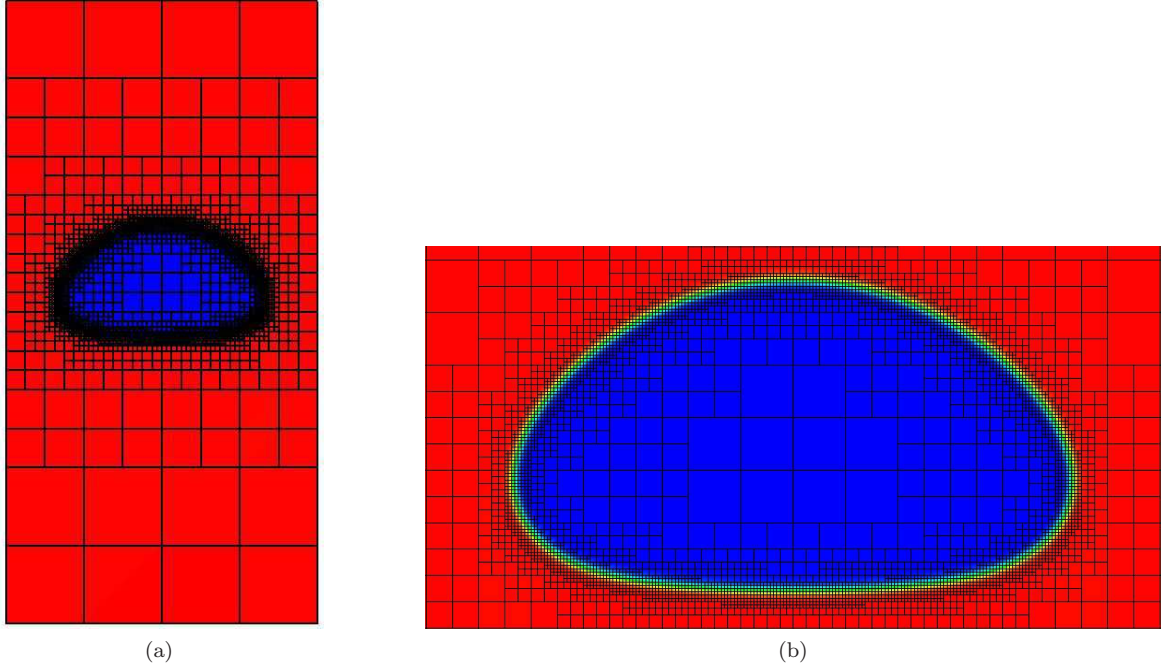


FIG. VIII.2 – Maillage à l'instant final $T = 3$, domaine complet (à gauche), zoom sur l'interface (à droite)

L'étude porte sur les valeurs de l'épaisseur d'interface ε et sur l'ordre de grandeur du coefficient de mobilité M_{deg} . Nous effectuons 24 simulations pour $\varepsilon \in \{\frac{R}{12}, \frac{R}{16}, \frac{R}{24}\}$ et $10^{-6} \leq M_{\text{deg}} \leq 10$. Nous utilisons le schéma Impl pour la discrétisation du système de Cahn-Hilliard et la méthode de Lagrangien augmenté pour la résolution du système de Navier-Stokes. Les systèmes linéaires sont résolus par un solveur direct de la librairie UMFPACK.

Nous comparons, aux valeurs de références, les valeurs des quantités du benchmark que nous avons calculées. Les résultats obtenus pour les différents couples $(\varepsilon, M_{\text{deg}})$ testés ainsi que les valeurs de référence sont présentés dans la table VIII.2 dans laquelle figure :

- la valeur de la vitesse moyenne maximale de la bulle,
- la position du centre de masse de la bulle à l'instant final $T = 3$,
- la circularité minimale de la bulle.

$M_{\text{deg}} \backslash \varepsilon$	$\frac{R}{20}$	$\frac{R}{16}$	$\frac{R}{12}$
10	0.2973	0.3020	0.3082
1	0.2481	0.2490	0.2514
10^{-1}	0.2419	0.2413	0.2412
10^{-2}	0.2417	0.2403	0.2404
10^{-3}	0.2414	0.2400	0.2390
10^{-4}	0.2389	0.2361	0.2326
10^{-5}	0.2289	0.2215	0.2135
10^{-6}	0.2127	0.2050	0.1982
valeur de référence : 0.2419 ± 0.0002			

(a) vitesse moyenne maximale de montée de bulle

$M_{\text{deg}} \backslash \varepsilon$	$\frac{R}{20}$	$\frac{R}{16}$	$\frac{R}{12}$
10	1.201	1.210	1.225
1	1.129	1.139	1.152
10^{-1}	1.089	1.088	1.090
10^{-2}	1.084	1.082	1.082
10^{-3}	1.084	1.081	1.080
10^{-4}	1.081	1.076	1.069
10^{-5}	1.059	1.043	1.022
10^{-6}	1.010	1.000	0.9806
valeur de référence : 1.081 ± 0.001			

(b) position du centre de masse à l'instant $T = 3$

$M_{\text{deg}} \backslash \varepsilon$	$\frac{R}{20}$	$\frac{R}{16}$	$\frac{R}{12}$
10	0.9927	0.9915	0.9900
1	0.9491	0.9527	0.9578
10^{-1}	0.9197	0.9183	0.9165
10^{-2}	0.9097	0.9056	0.8991
10^{-3}	0.8989	0.8911	0.8786
10^{-4}	0.8815	0.8716	0.8626
10^{-5}	0.8882	0.8855	0.8928
10^{-6}	0.9094	0.9180	0.9358
valeur de référence : 0.9012 ± 0.0001			

(c) circularité minimale

TAB. VIII.2 – Quantités du benchmark : valeurs de référence et valeurs obtenues en fonction de l'ordre de grandeur de la mobilité et de l'épaisseur d'interface

Pour faciliter la lecture de la table VIII.2, nous avons rapporté en gras les valeurs proches de la référence : à 2×10^{-3} près pour la vitesse moyenne maximale de montée de bulle et pour la position du centre de masse à l'instant final et à 1×10^{-2} près pour la circularité minimale.

Nous constatons que les résultats (table VIII.2) dépendent fortement de l'ordre de grandeur de la mobilité mais il apparaît néanmoins une plage de valeurs (entre 10^{-1} et 10^{-3}) pour lesquelles les résultats obtenus sont très proches des valeurs de référence.

En outre, nous observons que plus la valeur du coefficient de mobilité M_{deg} augmente plus la vitesse de montée de la bulle est importante et plus la bulle occupe une position haute dans le domaine à l'instant final.

La valeur de la mobilité influe également sur l'évolution de la forme de la bulle. En effet, le coefficient de circularité minimale permet d'évaluer la déformation maximale de la forme de la bulle. Si celui-ci reste très proche de 1, la bulle conserve une forme proche d'un disque. Ainsi, la table VIII.2 montre que plus la valeur de la mobilité est importante moins la bulle se déforme.

Ces observations sont illustrées par les figures VIII.3 et VIII.4.

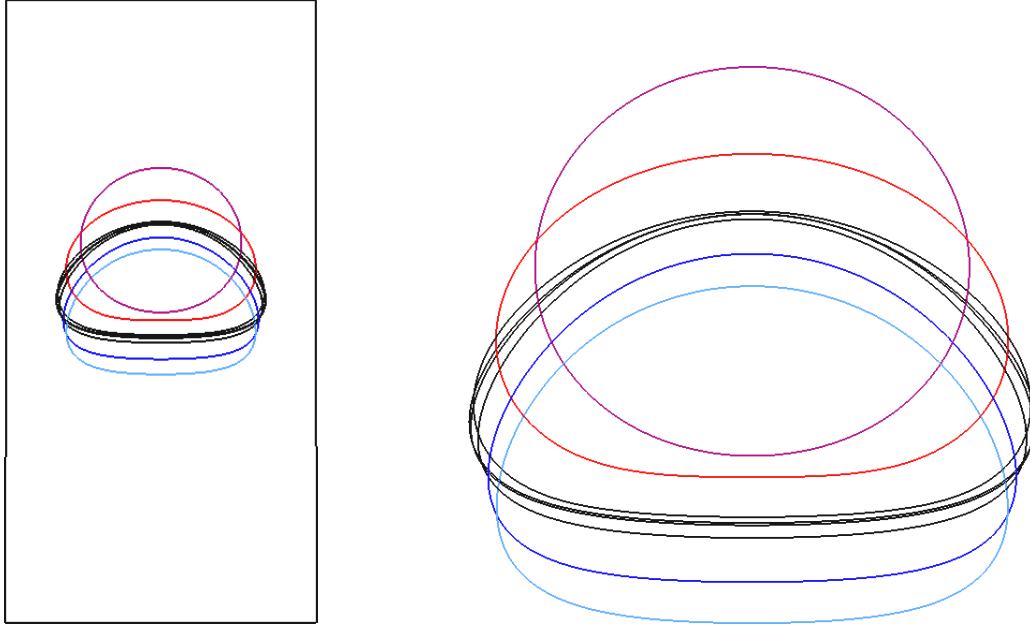


FIG. VIII.3 – Influence du coefficient de mobilité M_{deg} sur la forme de la bulle ($10^{-6} \leq M_{\text{deg}} \leq 10^{-1}$, $\varepsilon = \frac{R}{20}$), lignes de niveau des paramètres d'ordre au même instant physique $T = 3$

La figure VIII.3 permet de comparer les formes de bulle obtenues à l'instant final $T = 3$ pour différents ordres de grandeur de la mobilité, la valeur de ε étant fixée à $\frac{R}{20}$. Cette figure superpose les lignes de niveau $c = 0.5$ du paramètre d'ordre obtenu en donnant au coefficient M_{deg} chacune des huit valeurs suivantes : 10^{-6} , 10^{-5} , 10^{-4} , 10^{-3} , 10^{-2} , 10^{-1} , 1 et 10. Quatre de ces huit valeurs (10^{-4} , 10^{-3} , 10^{-2} , 10^{-1}) conduisent à des formes de bulles très semblables, correspondant aux tracés en noir sur la figure VIII.3. Cependant, pour les valeurs de mobilité les plus importantes 1 et 10 nous observons des formes de bulles tout à fait différentes, quasiment circulaires, représentées en rouge et mauve sur la figure VIII.3. Pour les valeurs de mobilité les plus faibles 10^{-5} et 10^{-6} , nous observons une forme de bulle correcte (représentée en bleu foncé et clair sur la figure) mais une sous-estimation importante de la vitesse de montée de bulle. De plus, pour ces valeurs de M_{deg} le profil du paramètre d'ordre est modifié. Ceci apparaît sur la figure VIII.4 où sont représentés les paramètres d'ordre obtenus pour des valeurs de mobilité égales à 10^{-6} , 10^{-1} et 10. La figure du haut montre les lignes de niveau (entre 0.01 et 0.99) du paramètre d'ordre et la figure du bas montre la valeur du paramètre d'ordre en dégradé de couleurs. Lorsque $M_{\text{deg}} = 10^{-6}$, nous constatons que la zone interfaciale a tendance à s'élargir au dessous de la bulle alors que ceci ne se produit pas pour des valeurs plus importantes du coefficient de mobilité M_{deg} puisque les lignes de niveau restent bien resserrées.

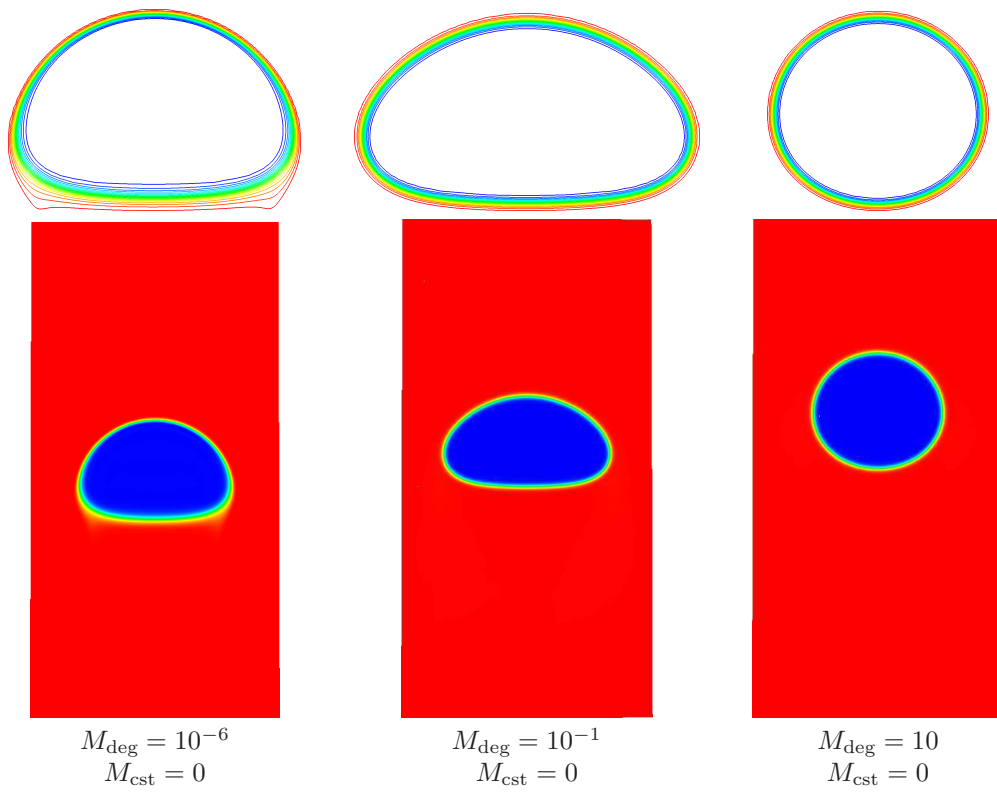


FIG. VIII.4 – Représentation du paramètre d'ordre associé à la bulle pour différents ordre de grandeur de la mobilité

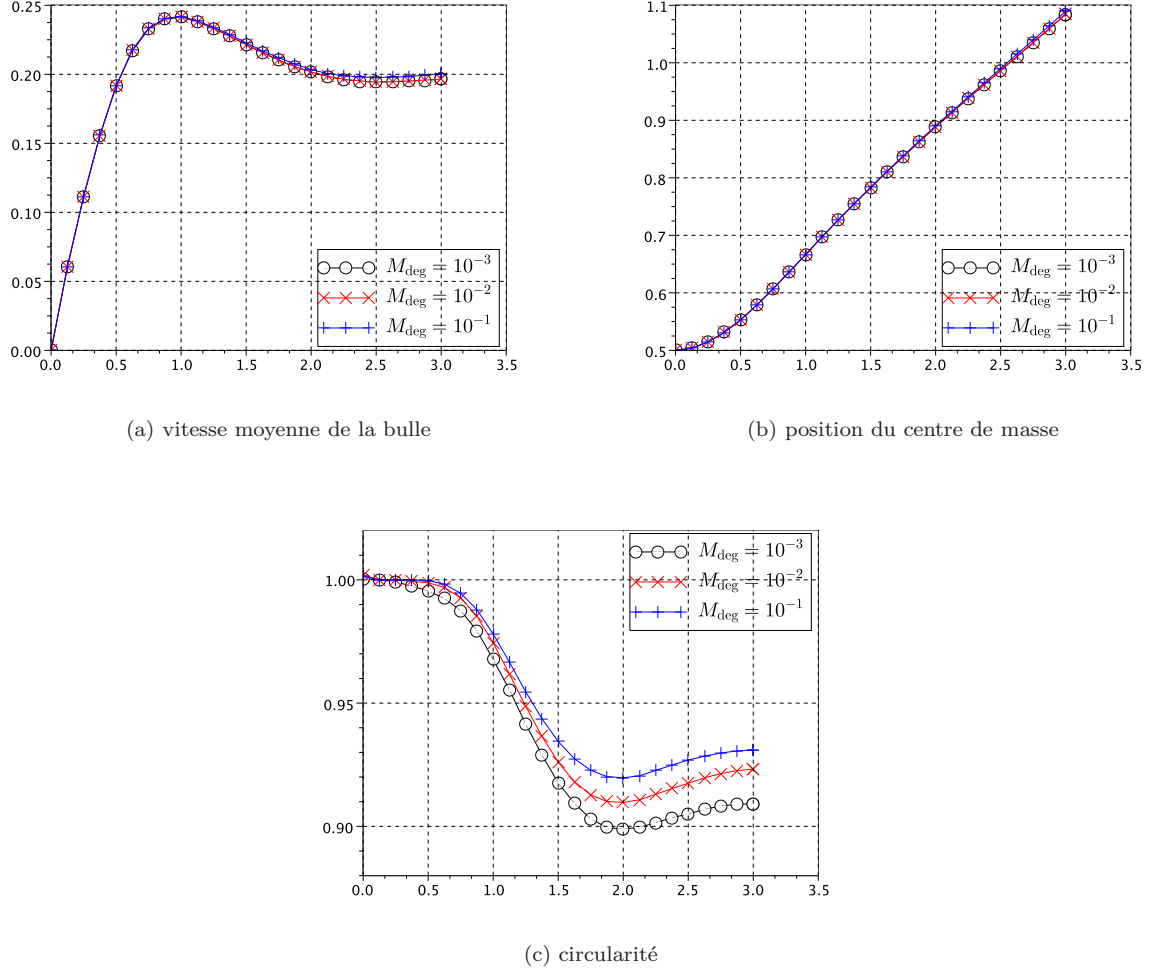


FIG. VIII.5 – Evolution temporelle des quantités du benchmark pour différentes valeurs de la mobilité ($\varepsilon = \frac{R}{20}$)

Nous examinons maintenant plus en détail les résultats obtenus pour un coefficient de mobilité M_{deg} appartenant à la plage de “bonnes valeurs” identifiée ci-avant. Nous comparons sur la figure VIII.5 l’évolution temporelle des quantités du benchmark obtenues pour ces différentes valeurs de la mobilité, l’épaisseur d’interface étant fixée à $\varepsilon = \frac{R}{20}$. Nous constatons que les courbes d’évolution de la vitesse moyenne et de la position du centre de masse de la bulle se superposent, l’évolution de la circularité est légèrement modifiée par les changements de valeur de la mobilité, ceci pouvant s’expliquer par la formule VIII.1 de calcul de la circularité qui dépend du profil de c (les différences observées restent néanmoins inférieures à 2%). Ainsi, ces résultats semblent confirmer que pour cet ensemble de “bonnes valeurs” du coefficient de mobilité, l’évolution de la bulle est qualitativement la même. Il n’est cependant pas évident d’identifier la dépendance entre cette plage de mobilité et les autres grandeurs du système (comme l’épaisseur d’interface, la vitesse de montée de bulle ...)

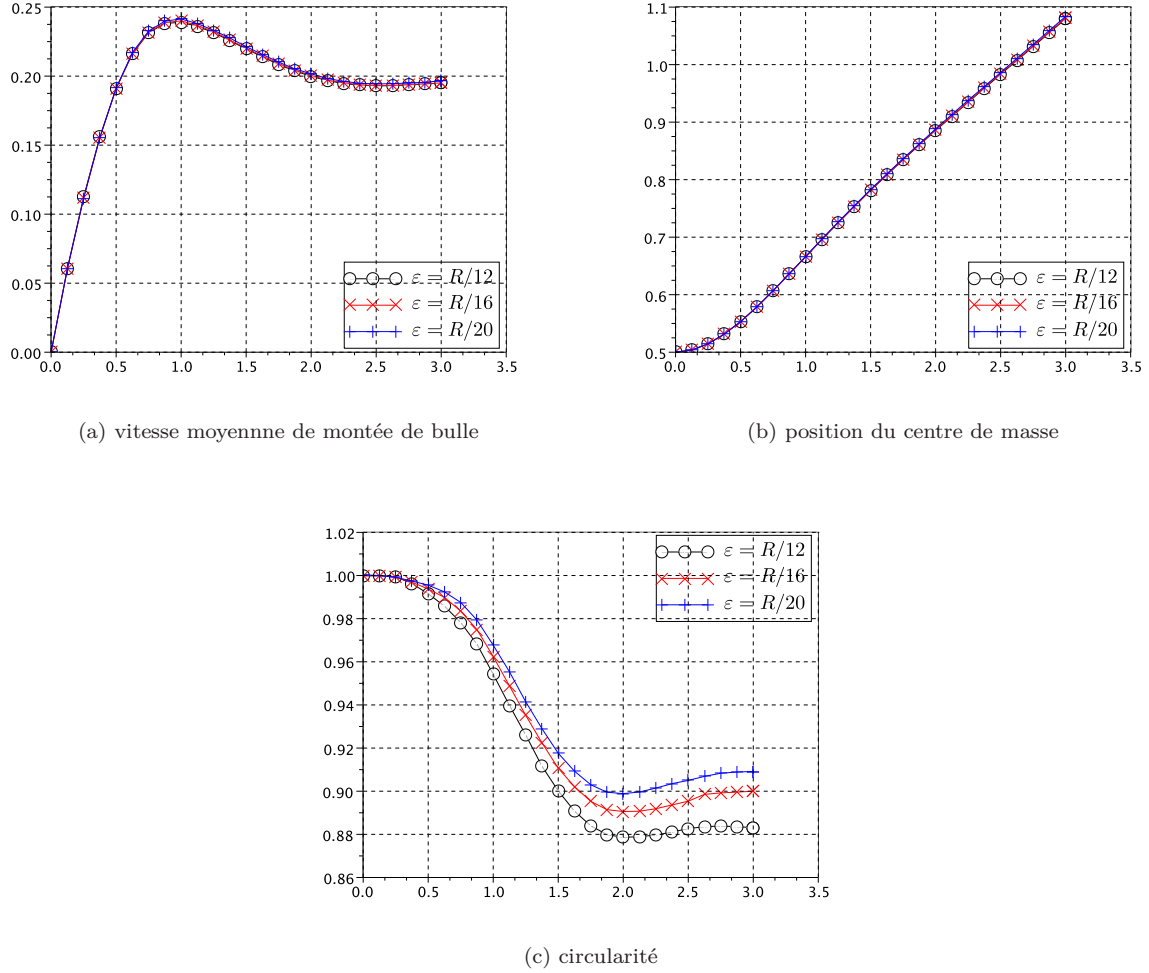


FIG. VIII.6 – Evolution temporelle des quantités du benchmark pour différentes valeurs de l'épaisseur d'interface, $M_{\text{deg}} = 10^{-3}$

La figure VIII.6 montre que les résultats sont très similaires lorsque nous faisons varier l'épaisseur d'interface pour une valeur de la mobilité fixée $M_{\text{deg}} = 10^{-3}$. Nous observons cette fois encore de légères différences sur les courbes de circularité. Nous pouvons néanmoins conclure que l'épaisseur d'interface ε impacte peu sur la dynamique de la bulle.

Enfin, nous terminons cette section en présentant sur la figure VIII.7 la comparaison entre les résultats obtenus pour une mobilité $M_{\text{deg}} = 10^{-3}$ et une épaisseur d'interface égale à $\varepsilon = \frac{R}{20}$ avec les courbes obtenues par les différents participants initiaux du benchmark.

Les résultats coïncident avec les “solutions” de référence établies dans le benchmark. La courbe de circularité n'est pas parfaitement superposée mais les différences restent inférieures à 2%.

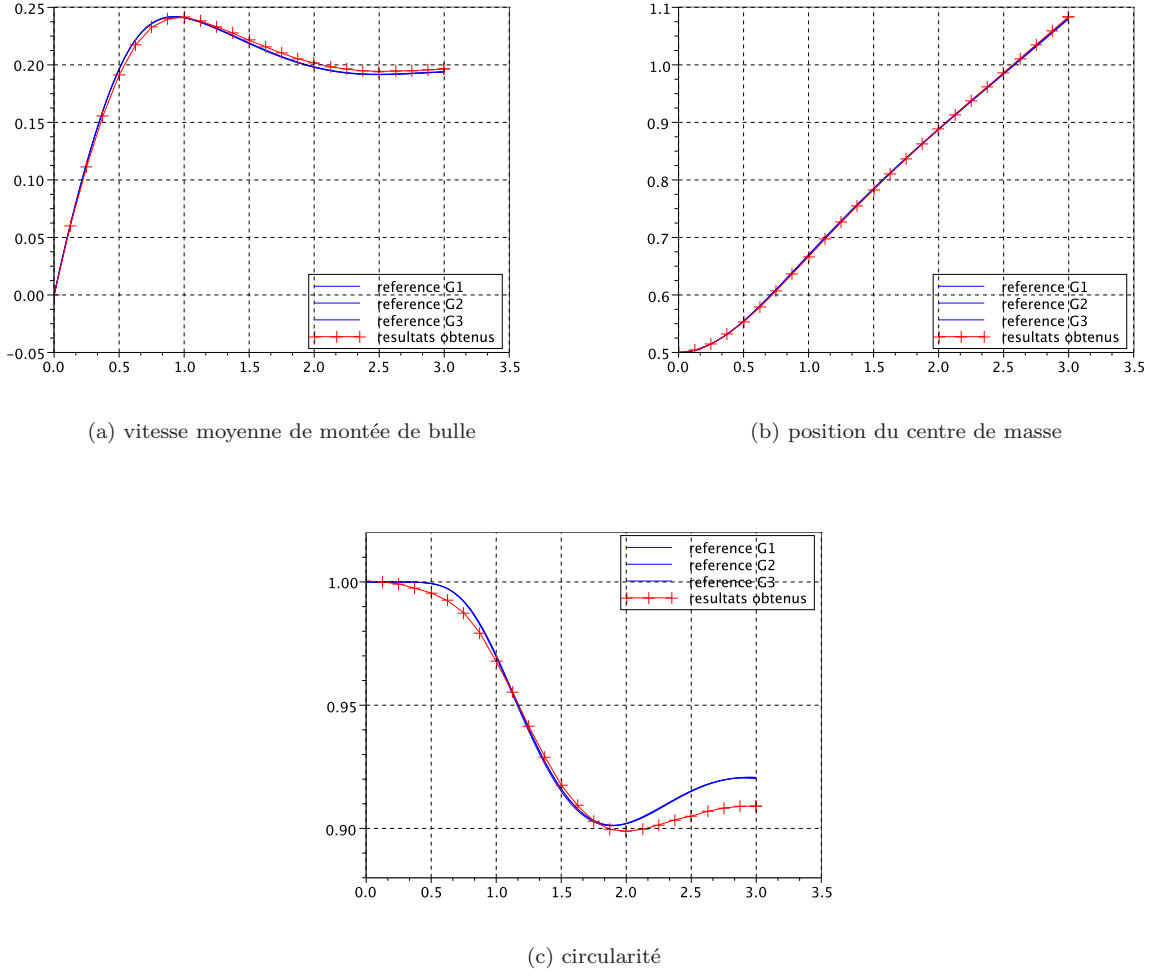


FIG. VIII.7 – Evolution temporelle des quantités du benchmark pour différentes valeurs de la mobilité, $\varepsilon = \frac{R}{20}$

VIII.2 Forme d'une bulle montant dans un liquide

Nous présentons dans cette section, une série de simulations de l'ascension d'une bulle (sous l'effet de la gravité) dans un liquide incompressible. Cet ensemble de simulations recouvre un grand nombre de situations d'intérêt physique.

Un tel écoulement est caractérisé par les trois nombres adimensionnels suivants :

- le nombre de Bond

$$Bo = \frac{\varrho_1 |\mathbf{g}| D^2}{\sigma},$$

- le nombre de Morton

$$Mo = \frac{|\mathbf{g}| \eta_1^4}{\varrho_1 \sigma^3},$$

- le nombre de Reynolds

$$Re = \frac{\varrho_1 v_{c,z} D}{\eta_1},$$

où D désigne le diamètre de la bulle, $v_{c,z}$ la composante verticale de la vitesse moyenne de montée de la bulle, ϱ_1 et η_1 les densité et viscosité du liquide (dans lequel est immergée la bulle), σ la tension de surface entre les deux phases et $|\mathbf{g}|$ la norme du vecteur gravité.

Les nombres de Bond et de Morton dépendent uniquement des propriétés des fluides en présence alors que le nombre de Reynolds est un résultat de la dynamique de la bulle (puisqu'il fait intervenir la vitesse moyenne de montée de la bulle).

Dans [CGW78], les auteurs ont établi expérimentalement une correspondance entre les valeurs des nombres adimensionnels ci-dessus et la forme qu'adopte la bulle lors de son ascension. Ainsi, ils ont pu proposer une cartographie des formes de bulles par l'intermédiaire d'un diagramme (*cf* figure VIII.8) dont les axes sont les nombres de Bond (également appelé nombre d'Eötvös (EO)) et de Reynolds, les isovaleurs du nombre de Morton figurant en arrière plan.

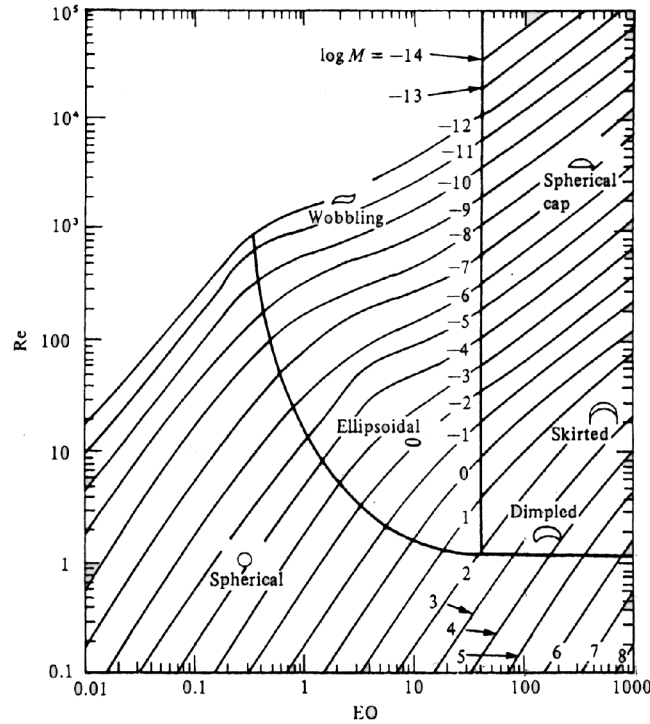


FIG. VIII.8 – Diagramme de Clift, Grace et Weber [CGW78]

Il s'agit ici de reproduire numériquement les différents types d'écoulements figurant sur ce diagramme. Pour faciliter les comparaisons, nous reprenons pour notre étude (un sous-ensemble) des simulations effectuées dans [BM07] (avec une méthode de type level-set) pour une large gamme de nombres de Bond ($1 \leq Bo \leq 100$) et de Morton ($5 \times 10^{-8} \leq Mo \leq 5 \times 10^6$) conduisant *a posteriori* à des nombres de Reynolds compris entre 1 et 100.

Nous utilisons le schéma $SI_{impl}(m, 0.5)$ pour le système de Cahn-Hilliard et la variante de la méthode de projection incrémentale donnée par le problème VII.10 pour le système de Navier-Stokes. Par ailleurs, nous initialisons numériquement les paramètres d'ordre (*cf* chapitre VII) par quelques itérations du système de Cahn-Hilliard avec une mobilité constante. Les simulations sont effectuées en géométrie axisymétrique mais il est important de noter que nous utilisons des solveurs itératifs pour chacune des étapes de résolution, et qu'il est donc envisageable de reproduire ces simulations en vraie géométrie tridimensionnelle (*cf* sections VIII.3 et IX.2). Nous utilisons le solveur GMRES préconditionné par ILU0 pour la résolution des systèmes linéaires de Cahn-Hilliard et de prédiction de vitesse. Nous utilisons la méthode du gradient conjugué préconditionnée par la méthode multigrilles II.13 (l'ensemble des niveaux disponibles est utilisé, le nombre de pré et post lissages est fixé à 2) pour la résolution des étapes de prédiction de pression et de calcul de l'incrément de pression. Et enfin, nous utilisons la méthode du gradient conjugué préconditionnée par la diagonale de la matrice (méthode de Jacobi) pour la résolution de l'étape de correction de vitesse.

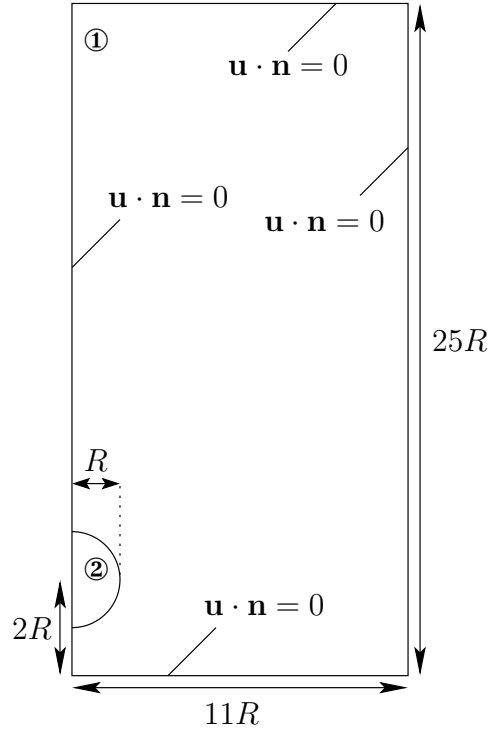


FIG. VIII.9 – Configuration initiale

La configuration initiale, décrite par la figure VIII.9, est identique pour toutes les simulations. Il s'agit d'une bulle sphérique de rayon R immergée dans un liquide occupant un domaine cylindrique (rappelons que nous supposons que la géométrie est axisymétrique) de hauteur $25R$ et de rayon $11R$. Le centre de la bulle (situé sur l'axe de symétrie) est initialement placé à une hauteur de $2R$ par rapport au bas du domaine. Pour les désigner plus facilement, nous numérotions les phases : le numéro 2 fait référence à la phase constituant la bulle et le numéro 1 à la phase entourant la bulle.

Pour toutes les simulations, le maillage initial (avant raffinement) est rectangle structuré et constitué de 5 cellules sur le rayon du cylindre et 10 cellules dans la hauteur. Le raffinement est ensuite dirigé par le critère VII.14 avec $h_{\min} = \varepsilon$.

A l'instar des références [Lap06, BM07], nous adoptons la démarche suivante :

(i) Les densités ϱ_1 , ϱ_2 , la tension de surface σ entre les deux fluides et le rapport des viscosités sont fixés :

$$\varrho_1 = 1000 \quad \text{et} \quad \varrho_2 = 1,$$

$$\sigma = 0.07,$$

$$\eta_1 = 100\eta_2.$$

(ii) L'étude est réalisée en faisant varier les nombres de Bond et de Morton. Pour rendre les comparaisons possibles, nous choisissons des valeurs proposées dans [BM07]. Celles-ci sont reportées dans la table VIII.3.

Test	(a)	(b)	(c)	(d)	(e)	(f)	(g)	(h)	(i)	(j)	(k)	(l)
Bo	1	10	100	1000	1	10	100	1000	1	10	100	1000
Mo	5.e-3	5	5.e+3	5.e+6	1.e-5	1.e-2	10	1.e+4	5.e-8	5.e-6	1.e-3	1

TAB. VIII.3 – Valeurs des nombres adimensionnels, nombre de Morton (Mo) et nombre de Bond (Bo) pour chacun des cas tests

(iii) Pour chaque couple de valeurs (Bo, Mo) des nombres de Bond et de Morton nous déduisons les valeurs

de la viscosité et du diamètre de bulle :

$$\eta_1 = \left(\frac{Mo\varrho_1\sigma^3}{|\mathbf{g}|} \right)^{\frac{1}{4}} \quad \text{et} \quad D = \left(\frac{Bo\sigma}{\varrho_1|\mathbf{g}|} \right)^{\frac{1}{2}}.$$

Ceci détermine l'ensemble des paramètres physiques nécessaires à la simulation. Le nombre de Reynolds est calculé *a posteriori* en évaluant la vitesse moyenne de montée de la bulle (*cf* section VIII.1.2).

Pour être complet dans la définition des paramètres utilisés pour nos simulations nous reportons dans la table VIII.4, pour chacune d'entre elles, la valeur du pas de temps Δt et du coefficient de mobilité M_{deg} utilisés.

Test	(a)	(b)	(c)	(d)	(e)	(f)	(g)	(h)	(i)	(j)	(k)	(l)
Δt	5.e-6	5.e-4	1.e-5	5.e-5	5.e-5	1.e-4	2.e-5	1.e-4	2.e-5	5.e-5	5.e-6	1.e-5
M_{deg}	1.e-5	1.e-5	1.e-5	1.e-4	1.e-5	1.e-5	1.e-4	1.e-3	1.e-5	1.e-5	1.e-5	5.e-4

TAB. VIII.4 – Valeurs du pas de temps Δt et du coefficient de mobilité M_{deg} utilisés pour chacune des simulations

Les résultats que nous obtenons sont présentés dans les deux figures VIII.10 et VIII.12. Ces figures sont à comparer aux deux figures 15. et 16. de la référence [BM07].

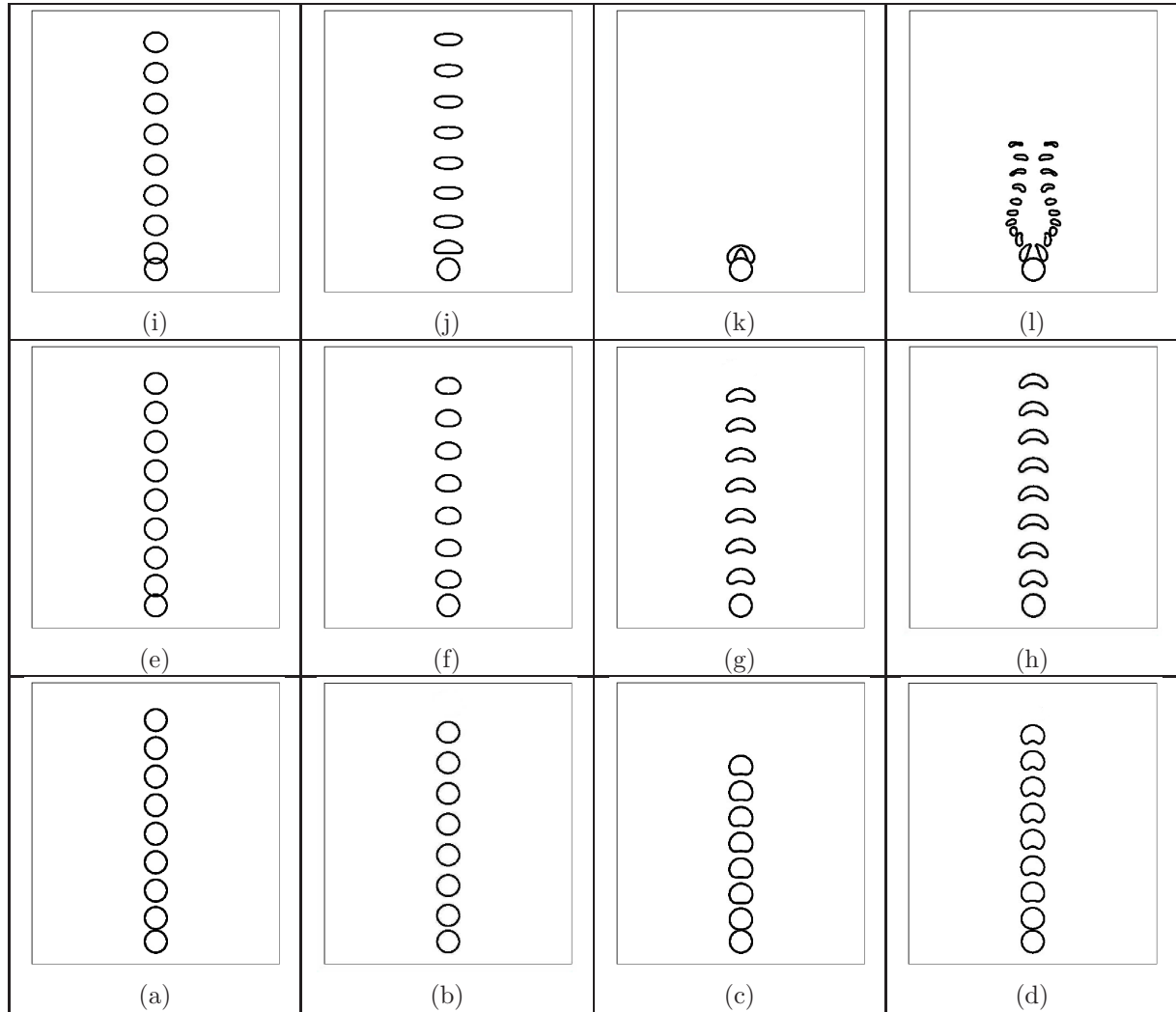


FIG. VIII.10 – Ascension d'une bulle dans un liquide incompressible

La figure VIII.10 présente l'évolution dans le temps de la forme des bulles. Chacune des sous figures correspond à une simulation, les lignes de niveau du paramètre d'ordre (entre 0.05 et 0.95) y sont représentées à différents instants. Ces instants sont régulièrement espacés pour chacune des sous-figures mais les échelles de temps sont différentes (et donc non comparable) entre deux sous-figures puisque les vitesses de montée de bulle sont différentes.

La simulation (k) n'a pas abouti à cause d'un problème de convergence du solveur linéaire interne à la méthode de Newton pour la résolution du système de Cahn-Hilliard. Ceci intervient au moment où le tore doit se former. Par ailleurs les simulations (a) et (h) ont posé des problèmes du même type. En ce qui concerne la simulation (a), nous observons une création de vitesses parasites importantes : les développements proposés dans la section VII.1.3 ont permis de résorber ces phénomènes néanmoins le pas de temps a dû être choisi très petit (*cf* table VIII.4) pour que le calcul se déroule bien. Pour le cas (h), nous avons augmenté le coefficient de mobilité M_{deg} jusqu'à 10^{-3} . Ceci a un impact sur les résultats obtenus : la figure VIII.11 montre que la forme obtenue avec une mobilité plus faible $M_{\text{deg}} = 5 \cdot 10^{-5}$ semble plus proche de celle observée expérimentalement (à noter que les paramètres physiques de l'expérimentation sont légèrement différents de ceux du calcul mais reste dans la même gamme). Cependant pour une telle mobilité, la déformation de l'interface est telle que le calcul n'aboutit pas. On remarquera que dans [BM07], les auteurs observent également une déformation de la zone de transition entre les phases. Pour ce type de simulations, il est peut être nécessaire de considérer une mobilité variable (par exemple dépendant de la vitesse de montée de bulle).

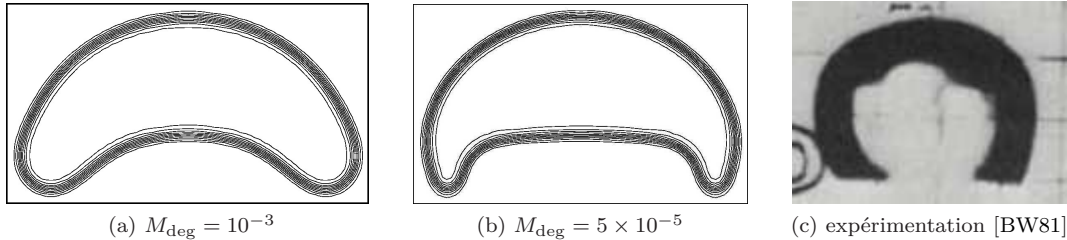


FIG. VIII.11 – Comparaison des simulations cas (h) aux expérimentations de [BW81] ($Bo = 339$, $Mo = 43.1$, $Re = 18.3$, $\frac{\rho_1}{\rho_2} = 1050$, $\frac{\eta_1}{\eta_2} \geq 4 \times 10^3$)

L'ensemble des autres simulations permet néanmoins d'obtenir des formes de bulle assez variées. Les différentes formes répertoriées dans la figure VIII.8 sont représentées. Nous observons des bulles sphériques (simulation (a)), des bulles ellipsoïdales (simulations (f), et (i)), des bulles à jupe dont nous avons déjà parlé (simulations (g) et (h)), des bulles toriques (simulation (l))...

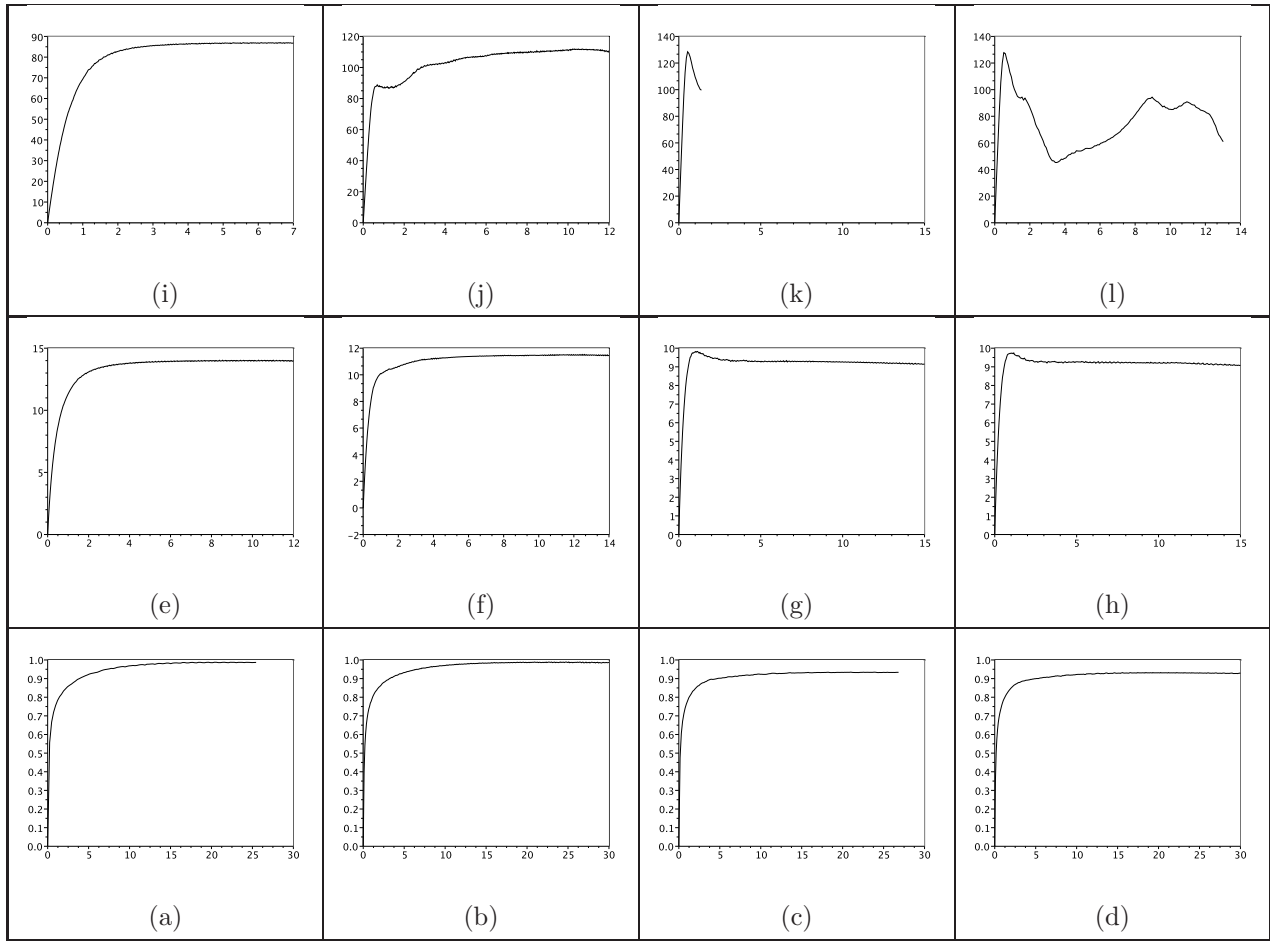


FIG. VIII.12 – Evolution en temps du nombre de Reynolds de bulle

Enfin, nous présentons sur la figure VIII.12, les courbes d'évolution du nombre de Reynolds.

VIII.3 Calcul tridimensionnel

Nous présentons maintenant un calcul tridimensionnel simulant l'ascension de trois bulles immergées dans un liquide.

Le domaine de calcul est $] - 0.032; 0.032[\times] - 0.032; 0.032[\times] 0; 0.096[$. Les trois bulles sont initialement sphériques, leur position est indiquée dans la table VIII.5a qui contient les coordonnées des centres de chacune des bulles $\mathbf{x}_1$, $\mathbf{x}_2$ et $\mathbf{x}_3$ ainsi que leur rayons R_1 , R_2 et R_3 . Les conditions aux bords pour le système de Navier-Stokes sont de type glissement sur tous les bords du domaine. La table VIII.5b donne les propriétés physiques des fluides en présence et enfin la table VIII.5c présente les valeurs des paramètres numériques utilisées pour le calcul présenté.

$\mathbf{x}_1$	$\mathbf{x}_2$	$\mathbf{x}_3$
$(-0.008, 0.008, 0.016)$	$(0.004, -0.004, 0.024)$	$(0.012, 0.012, 0.032)$
R_1	R_2	R_3
0.01	0.006	0.004

(a) positions des bulles

ϱ_1	ϱ_2	η_1	η_2	σ
1	1000	10^{-4}	0.1	0.07

(b) paramètres physiques

ε	$h_{\text{interface}}$	M_{deg}
1.77×10^{-3}	ε	1.1×10^{-5}

(c) paramètres numériques

TAB. VIII.5 – Configuration initiale et paramètres du cas test

Le maillage initial (avant raffinement) est constitué de $8 \times 8 \times 12$ cellules.

Nous utilisons le schéma $\text{SIMPL}(m, 0.5)$ et la méthode de projection décrite dans le problème VII.8. Nous utilisons le préconditionneur ILU0 pour la résolution de tous les systèmes linéaires.

La figure VIII.13 montre l'évolution en temps du système à travers la représentation de la ligne de niveau $c = 0.5$ du paramètre d'ordre.

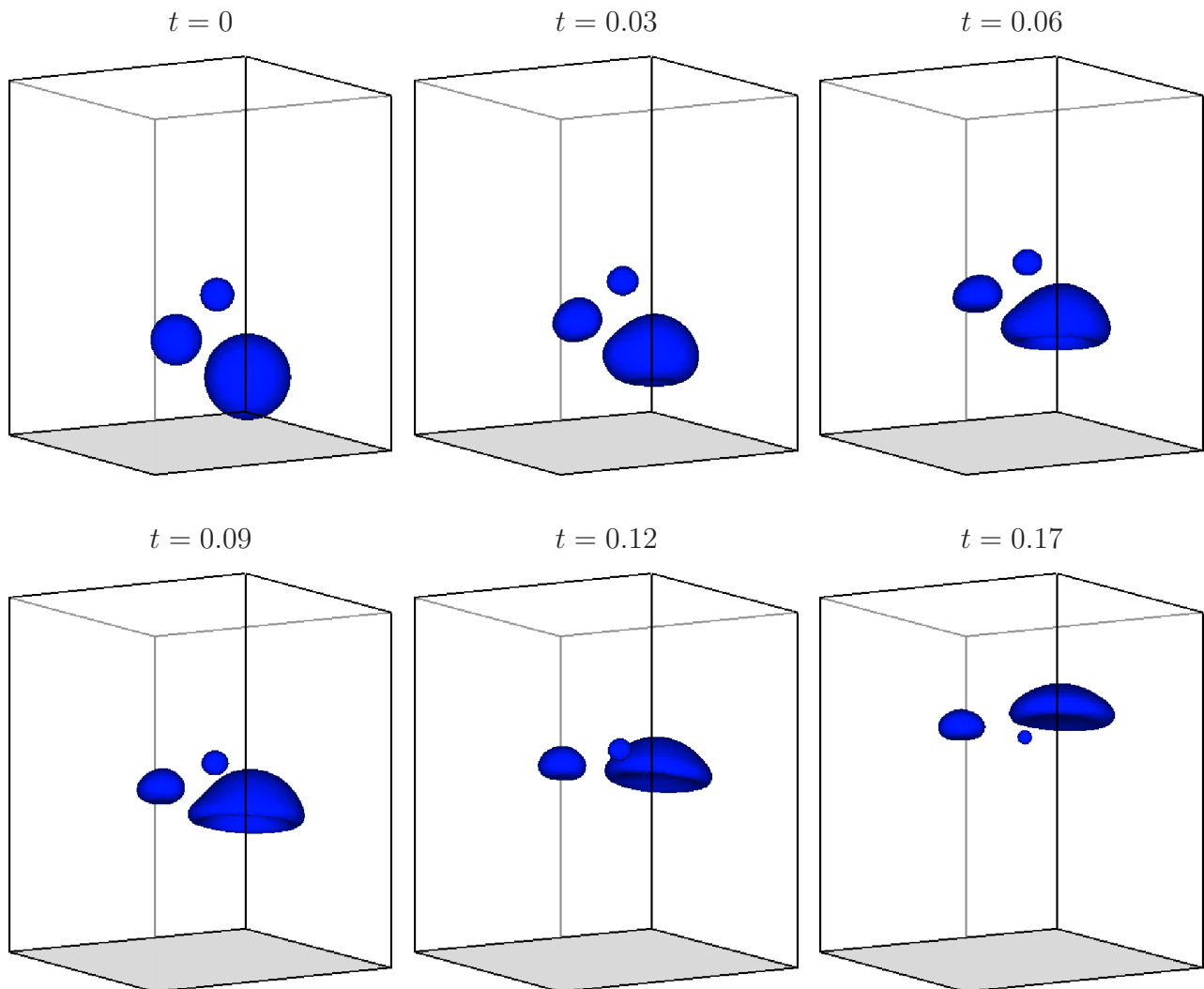


FIG. VIII.13 – Exemple de calcul diphasique tridimensionnel avec trois bulles

Nous constatons qu'au cours de cette évolution la bulle 3 (celle de plus petit rayon) a tendance à disparaître. Ceci provient de notre choix des paramètres : l'épaisseur d'interface ε est à peine plus que deux fois inférieures

au rayon de cette bulle. La petite bulle a alors tendance à “diffuser” dans les grandes. Il faut noter que la modélisation impose au paramètre ε d’être fixé de manière identique pour toutes les interfaces.

Le nombre d’inconnues maximal pour un champ scalaire $\mathbb{Q}_1$ a été au cours de ce calcul de 55 778. Sur un maillage globalement raffiné de résolution identique, il aurait été de 409 825.

La charge de calcul a été répartie entre 16 processus, le découpage initial attribuant à chacun des processus une colonne de 2x2x12 cellules grossières. Le temps de calcul observé (*i.e.* le temps écoulé entre la première et la dernière sauvegarde) a été de 9 jours et 6 heures (ce qui correspond en moyenne à environ 35 minutes par itérations en temps). Ce temps de calcul n’est donné ici qu’à titre indicatif puisqu’il dépend bien sûr de la machine utilisée et de l’encombrement du réseaux.

Chapitre IX

Etude de configurations triphasiques

IX.1 Bulle traversant une interface liquide-liquide

Nous présentons dans cette section la simulation d'une bulle traversant une interface liquide. Nous comparons les différents schémas en temps pour le système de Cahn-Hilliard et la différence observée entre une résolution des équations de Navier-Stokes par une méthode de type Lagrangien augmenté et la méthode de projection incrémentale décrite par le problème VII.8.

IX.1.1 Présentation du cas test

Nous reprenons une configuration proposée dans [BLP07]. Celle-ci est décrite dans la figure IX.1. Il s'agit d'une bulle (bidimensionnelle) circulaire totalement immergée dans la phase inférieure dans un bain stratifié. Les trois phases sont numérotées par convention de la manière suivante : phase 1 désigne la bulle, phase 2 désigne la phase inférieure et phase 3 désigne la phase supérieure. Les propriétés physiques des fluides en présence sont données dans le tableau IX.1. En particulier, nous avons un contraste de densité et de viscosité de l'ordre de 10^3 entre la phase 1 et les deux autres phases. Les tensions de surface sont différentes. Les conditions aux bords sont de type glissement sur tous les bords du domaine.

R	T	
0.006	0.8	
σ_{12}	σ_{13}	σ_{23}
0.07	0.07	0.05
ϱ_1	ϱ_2	ϱ_3
1	1200	1000
η_1	η_2	η_3
10^{-4}	0.15	0.1

TAB. IX.1 – Paramètres physiques

Le maillage initial (avant raffinement) est un maillage carré 16×80 et le temps final vaut $T = 0.8$. Le coefficient de mobilité est fixé à $M_{\text{deg}} = 10^{-5}$.

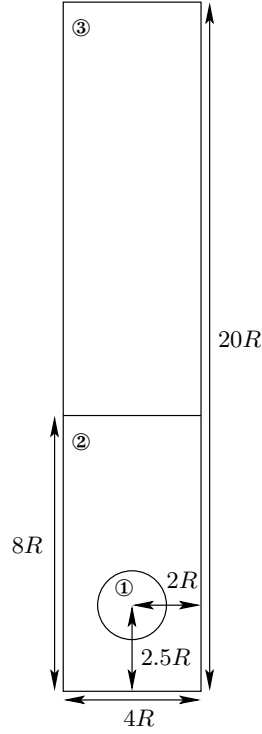


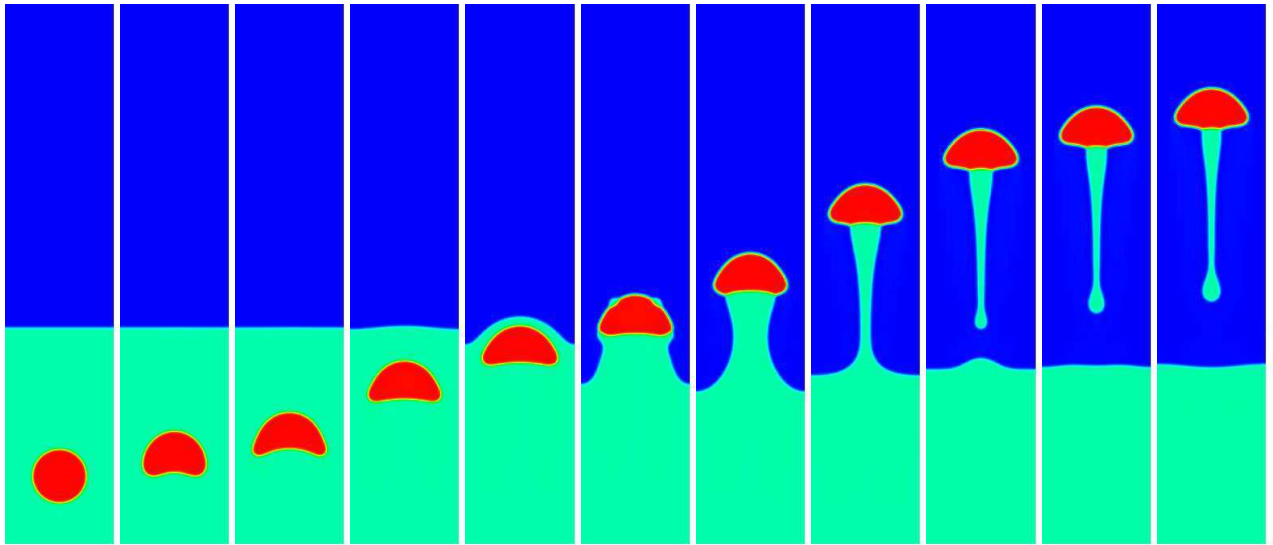
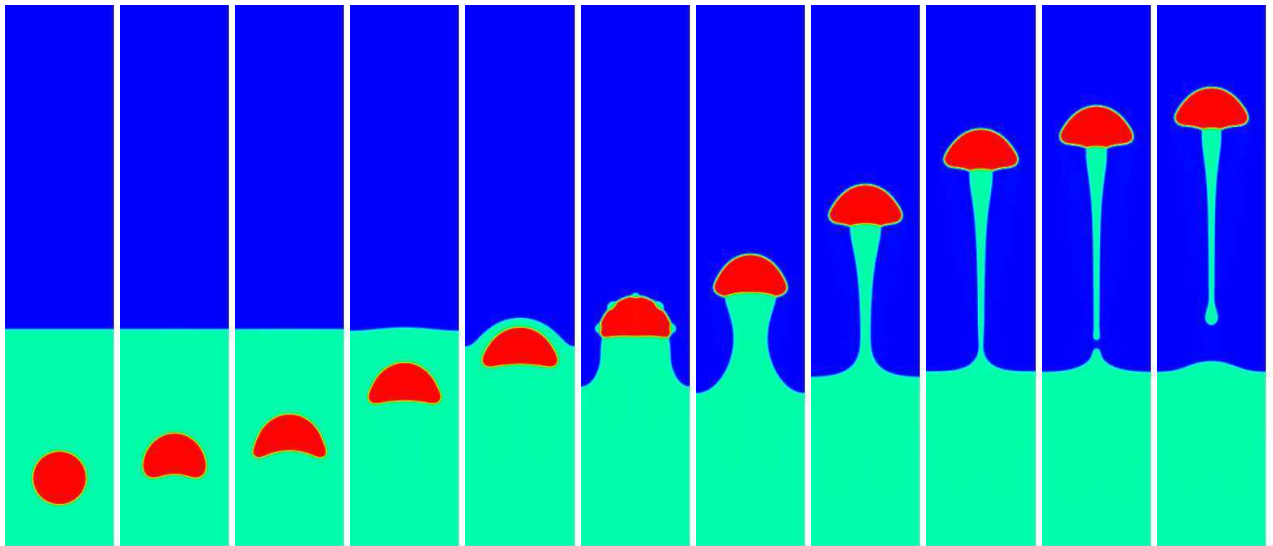
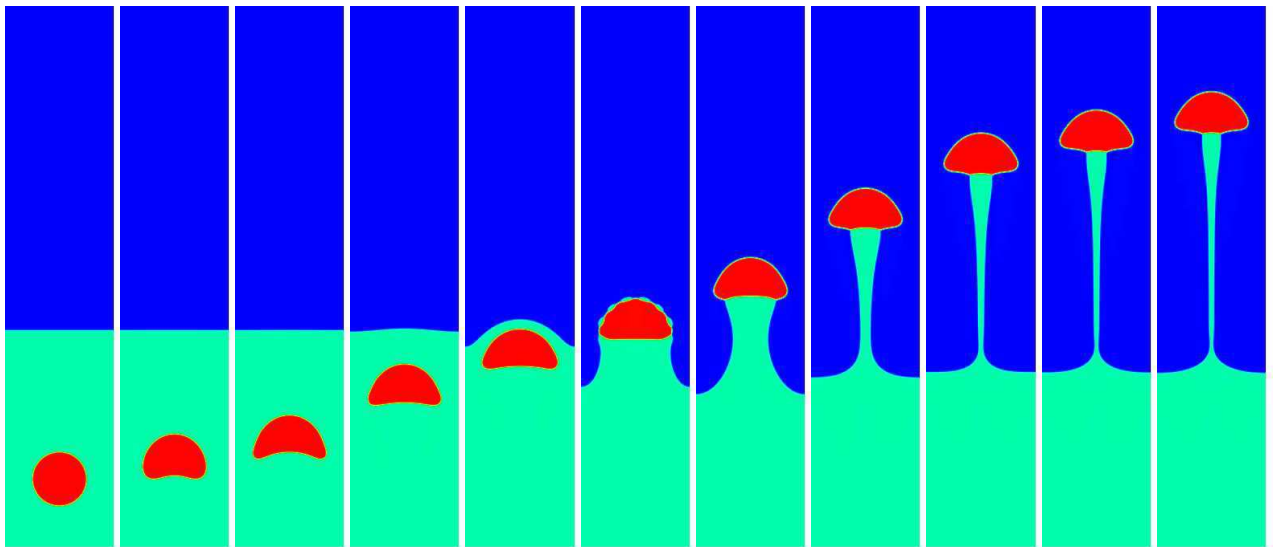
FIG. IX.1 – Configuration initiale du cas test

IX.1.2 Influence de la valeur de l'épaisseur d'interface ε

Nous comparons les résultats obtenus pour différentes valeurs $\frac{R}{10}$, $\frac{R}{15}$ et $\frac{R}{20}$ de l'épaisseur d'interface ε , le pas de temps étant fixé à $\Delta t = 5 \times 10^{-4}$ et le critère de raffinement VII.14 à $h_{\text{interface}} = \frac{\varepsilon}{4}$ (l'interface contient alors entre 6 et 8 mailles). Nous pouvons raisonnablement supposer que les différences observées sont dues à la valeur du paramètre ε et non à un manque de résolution dans l'interface.

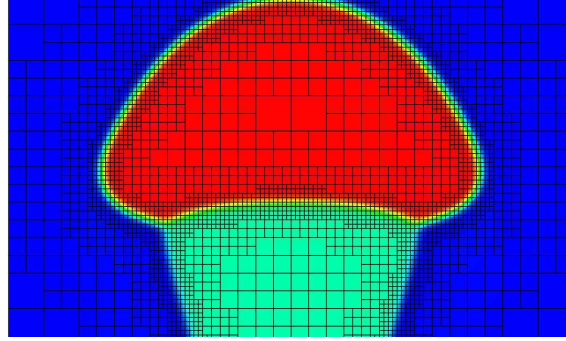
Le système de Cahn-Hilliard est ici discrétisé à l'aide du schéma Impl. et le système de Navier-Stokes résolu par la méthode de Lagrangien augmenté. Les systèmes linéaires sont tous résolus par un solveur direct de la librairie UMFPACK.

Nous observons des résultats (figure IX.2) très similaires pour les différentes valeurs de ε comme nous l'avons déjà constaté dans la section VIII.1.3. Cependant, des différences apparaissent lors de la rupture de la colonne de fluide entraîné. De manière assez intuitive, plus l'épaisseur d'interface ε est petite, plus l'instant de rupture est retardé. En effet, une valeur petite de ε autorise la représentation de films plus minces.

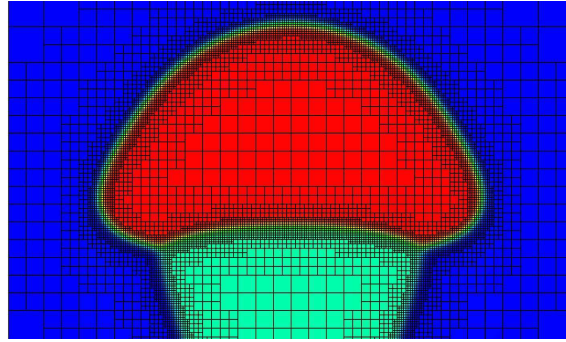
(a) $\varepsilon = \frac{R}{10}$ (b) $\varepsilon = \frac{R}{15}$ (c) $\varepsilon = \frac{R}{20}$ FIG. IX.2 – Influence de la valeur de l'épaisseur d'interface ε

IX.1.3 Influence de la valeur du nombre de mailles dans l'interface

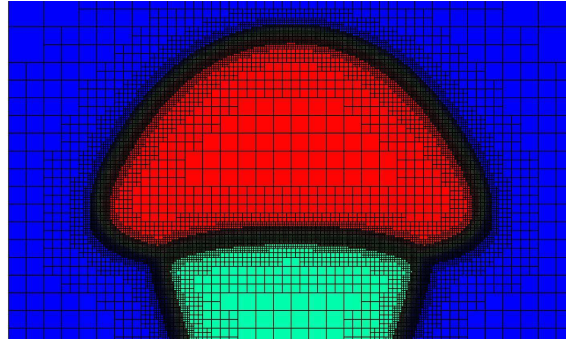
Nous fixons maintenant $\varepsilon = \frac{R}{15}$ et nous intéressons aux résultats obtenus lorsque la valeur de $h_{\text{interface}}$ varie. Nous avons effectué la simulation avec $h_{\text{interface}} = \varepsilon, \frac{\varepsilon}{2}$ et $\frac{\varepsilon}{4}$. Les schémas utilisés et le pas de temps sont toujours les mêmes. Les maillages obtenus à l'instant $t = 0.44$ sont présentés sur la figure IX.3.



(a) $h_{\text{interface}} = \varepsilon$



(b) $h_{\text{interface}} = \frac{\varepsilon}{2}$

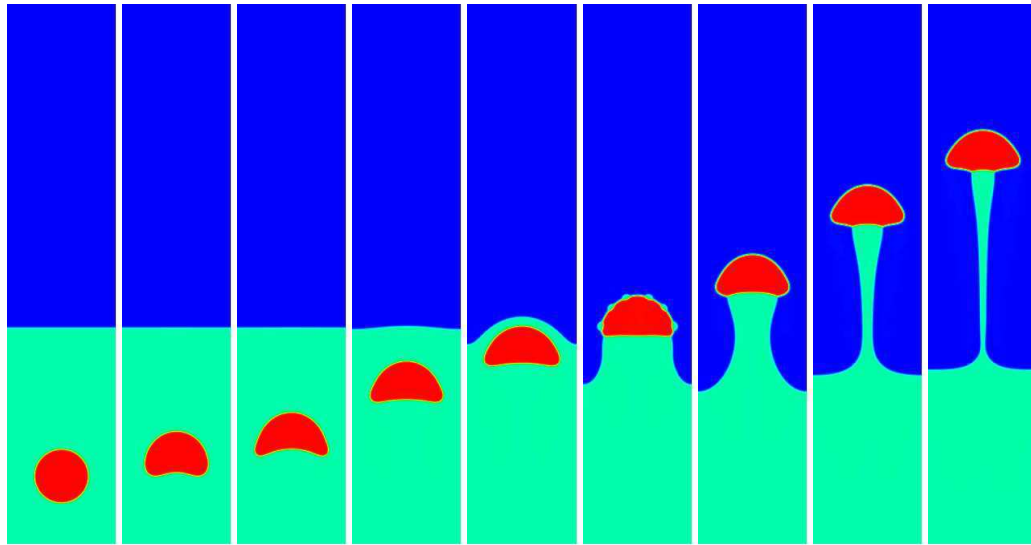


(c) $h_{\text{interface}} = \frac{\varepsilon}{4}$

FIG. IX.3 – Maillages obtenus à l'instant $t = 0.44$ pour différentes valeurs de $h_{\text{interface}}$

La figure IX.4 montre l'évolution des paramètres d'ordre pour les différentes valeurs $h_{\text{interface}} = \varepsilon, \frac{\varepsilon}{2}$ et $\frac{\varepsilon}{4}$. Le calcul avec $h_{\text{interface}} = \frac{\varepsilon}{4}$ ne s'est pas terminé à cause d'une surcharge de la mémoire au cours de la factorisation par le solveur direct de la matrice de rigidité $\mathbb{Q}_2$ dans la méthode de Lagrangien augmenté (le nombre d'inconnues augmente fortement avec l'étirement de l'interface puisque la zone de raffinement s'étale, cf section IX.1.7).

Nous observons des résultats très analogues à ceux de la section précédente (figure IX.2) : une résolution moins fine de l'interface a peu d'influence sur le calcul si ce n'est au moment de la rupture de la colonne de fluide entraîné. Ainsi, il n'est pas nécessaire d'avoir une résolution très importante de l'interface (2 ou 3 mailles dans l'interface suffisent, *i.e.* $h_{\text{interface}} = \varepsilon$) pour obtenir des solutions ayant un bon comportement qualitatif.



Non disponible

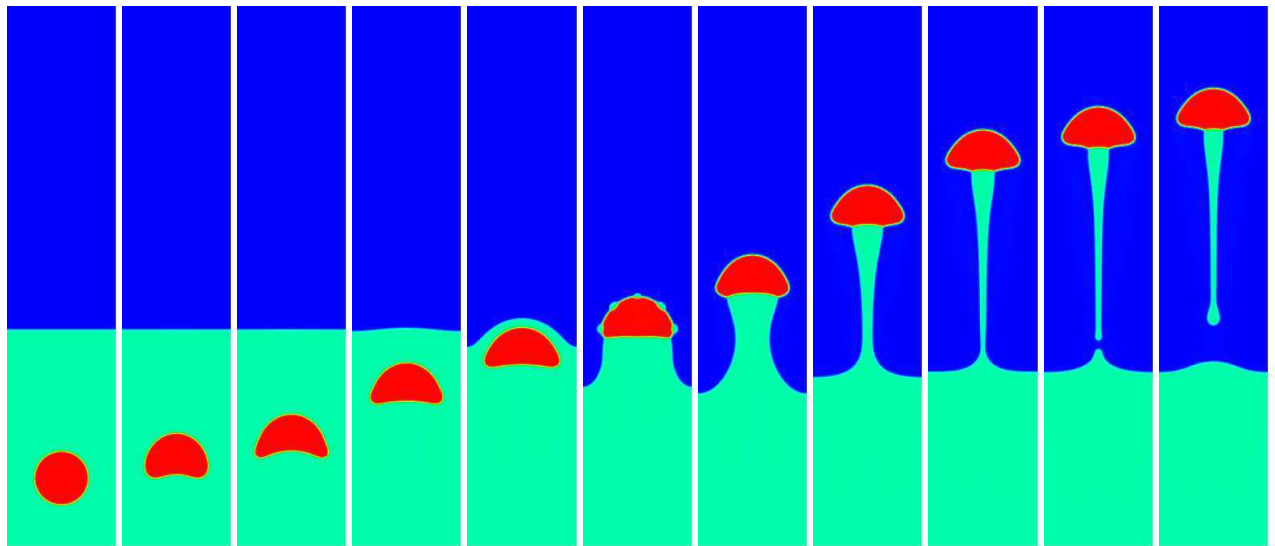
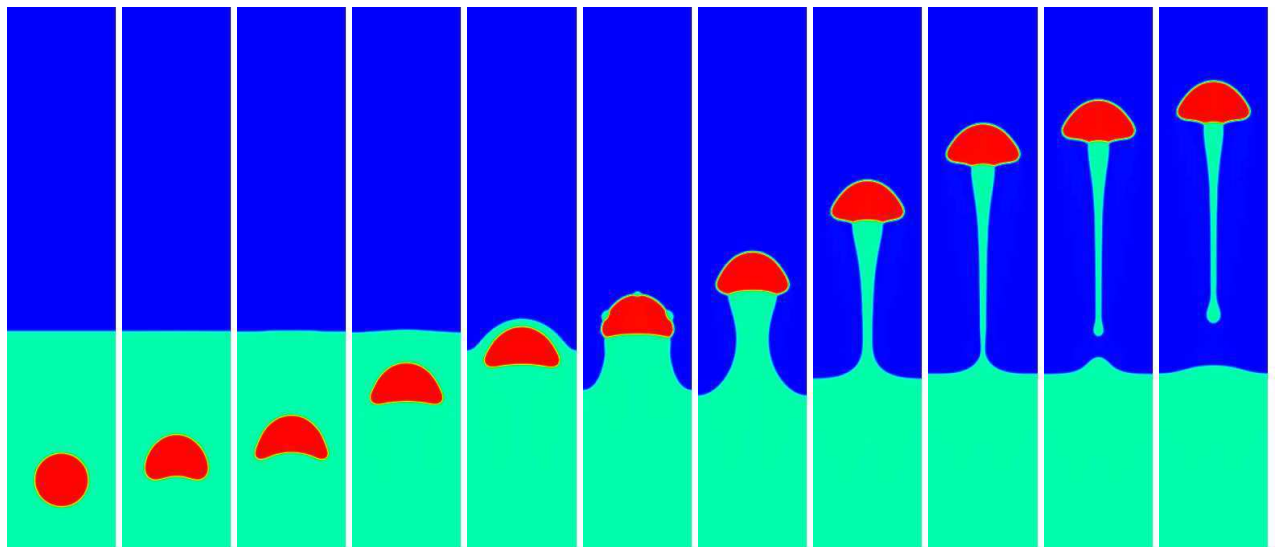
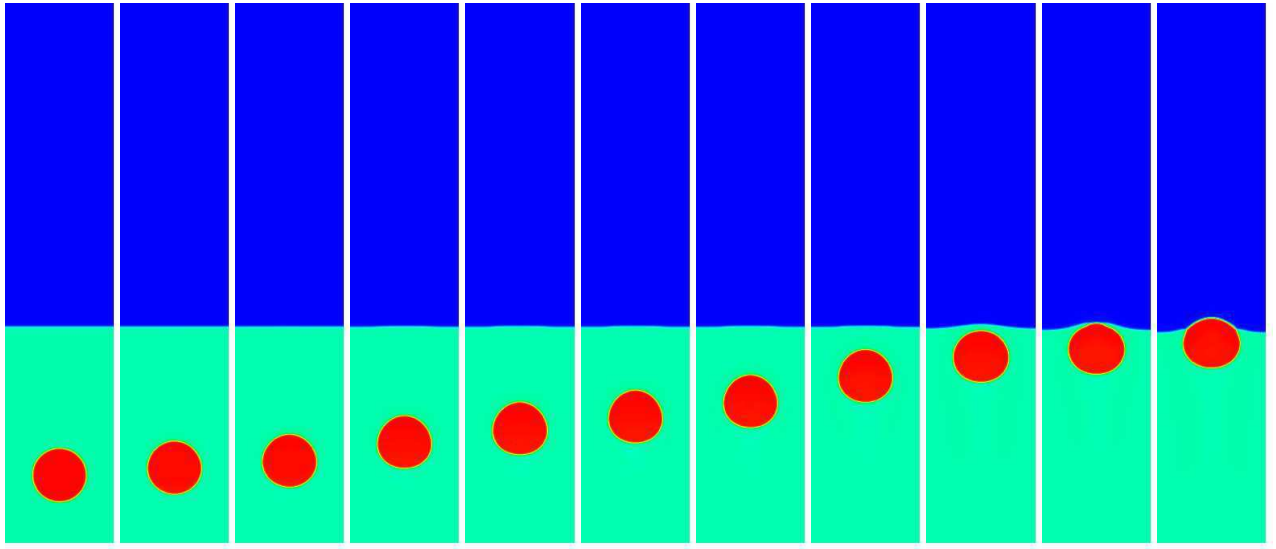
(a) $h_{\text{interface}} = \frac{\varepsilon}{4}$ (b) $h_{\text{interface}} = \frac{\varepsilon}{2}$ (c) $h_{\text{interface}} = \varepsilon$

FIG. IX.4 – Influence de la valeur du nombre de mailles dans l'interface

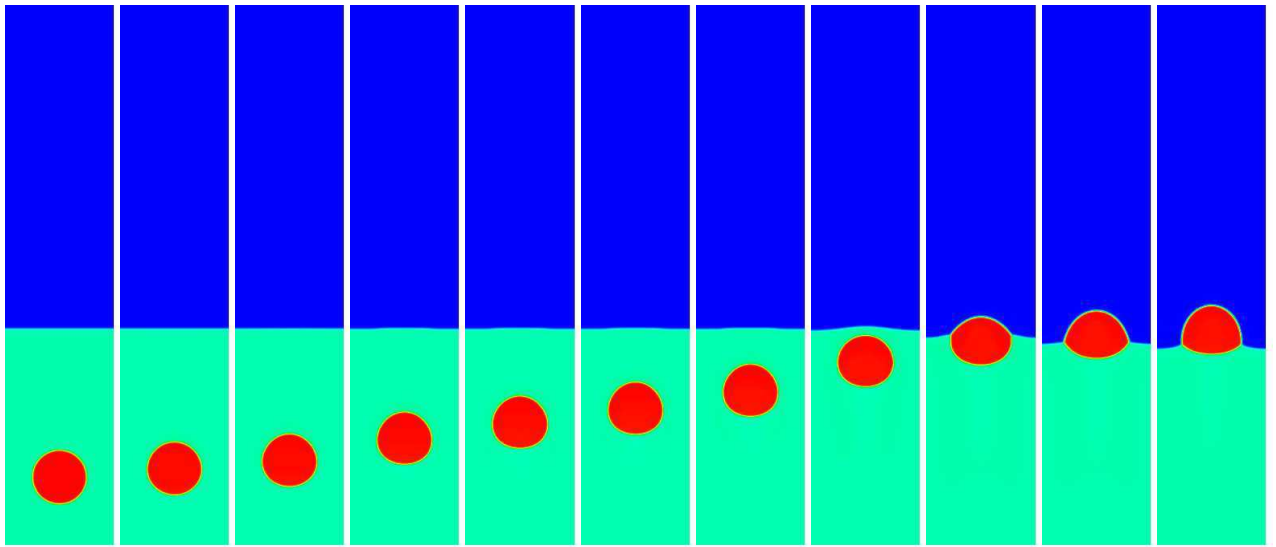
Au vu de ces premiers résultats, nous fixons $\varepsilon = \frac{R}{15}$ et $h_{\text{interface}} = \varepsilon$ pour la suite et comparons les différents schémas.

IX.1.4 Influence des schémas en temps sur le système Cahn-Hilliard

Nous discutons, dans cette section, l'influence de la discrétisation des termes non linéaires F_0 du système de Cahn-Hilliard. Nous comparons les différents schémas présentés au chapitre V. Nous fixons le pas de temps ($\Delta t = 5 \times 10^{-4}$) et présentons dans la figure IX.5, les résultats obtenus avec les schémas CC, CC(0.5), SImpl et SImpl(0.5), le système de Navier-Stokes étant toujours résolu par une méthode de type Lagrangien augmenté (et les systèmes linéaires par une méthode directe). Ces résultats sont à comparer avec celui montré sur la figure IX.4a. Il faut noter que le schéma Impl(0.5) ne permet pas d'aboutir à un résultat (problème de convergence dans l'algorithme de Newton). Ceci est cohérent avec les résultats théoriques présentés dans la section V.2.2 (le terme de diffusion numérique qui disparaît dans le cas $\beta = 0.5$ est utilisé pour compenser des termes dans l'estimation d'énergie, nous sommes donc incapables de garantir que le problème discret a bien une solution).

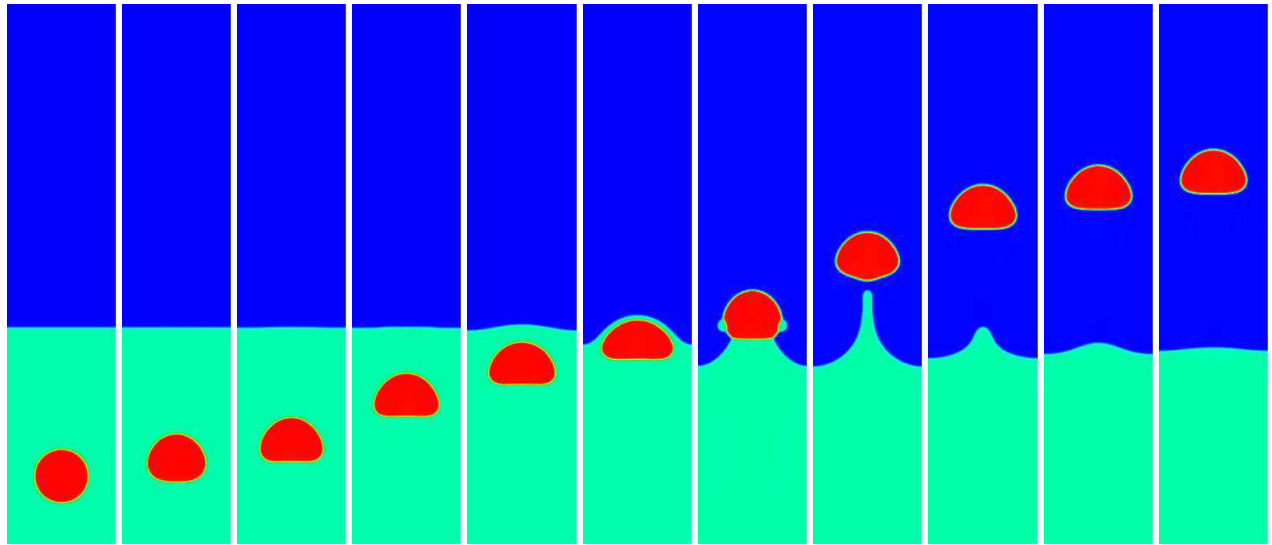


(a) Schéma CC.

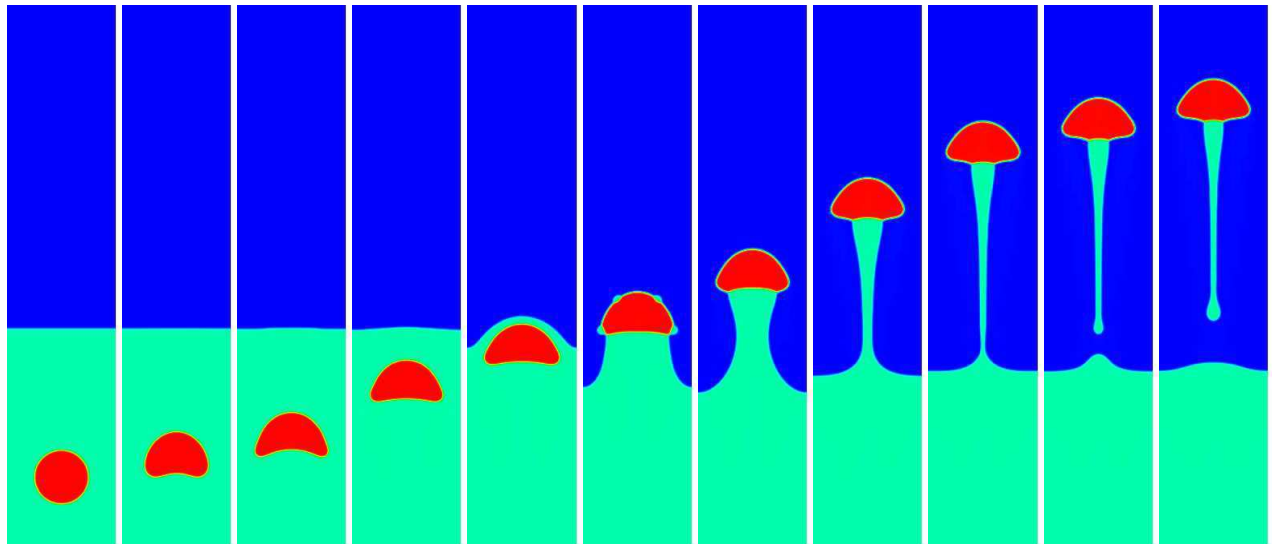


(b) Schéma CC(0.5)

FIG. IX.5 – Influence des schémas en temps pour les termes non linéaires F_0 de l'équation de Cahn-Hilliard



(c) Schéma SImpl



(d) Schéma SImpl(0.5)

FIG. IX.5 – Influence des schémas en temps pour les termes non linéaires F_0 de l'équation de Cahn-Hilliard

Seul le schéma SImpl(0.5) permet de retrouver des résultats qualitativement comparables à ceux observés lors de l'utilisation du schéma Impl. Le schéma CC quant à lui, sous-estime très fortement la vitesse de montée de la bulle et ce, même lorsque le paramètre β vaut 0.5.

Ces résultats sont analogues à ceux obtenus dans la section V.3 où des courbes de convergence montrant la précision des différents schémas ont été présentées.

IX.1.5 Influence de la valeur du pas de temps

Pour conforter les conclusions de la section précédente nous faisons maintenant une étude de l'influence de la valeur du pas de temps pour les schémas Impl et SImpl(0.5). Pour chacun de ces schémas nous effectuons deux simulations supplémentaires avec $\Delta t = 7 \times 10^{-4}$ et $\Delta t = 10^{-3}$ (rappelons que dans la figure IX.5, le pas de temps était fixé à $\Delta t = 5 \times 10^{-4}$).

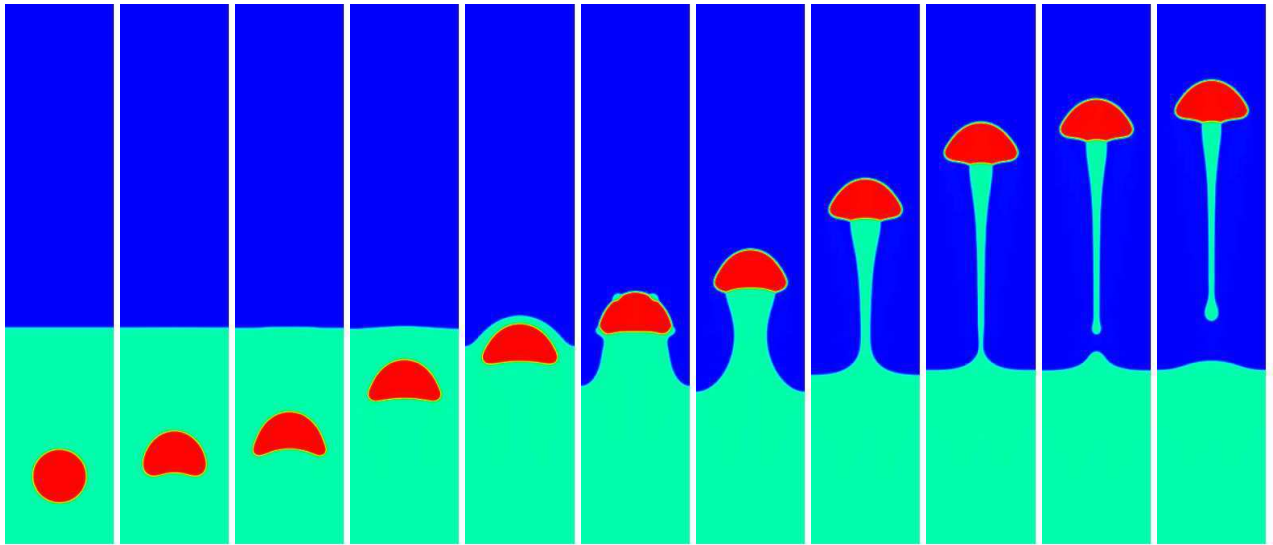
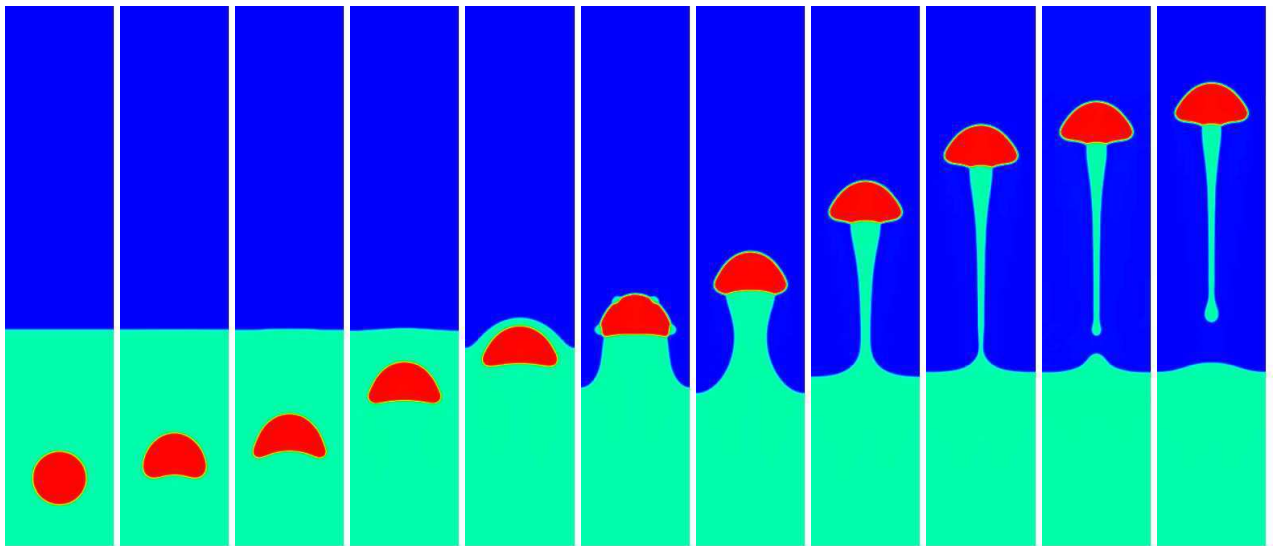
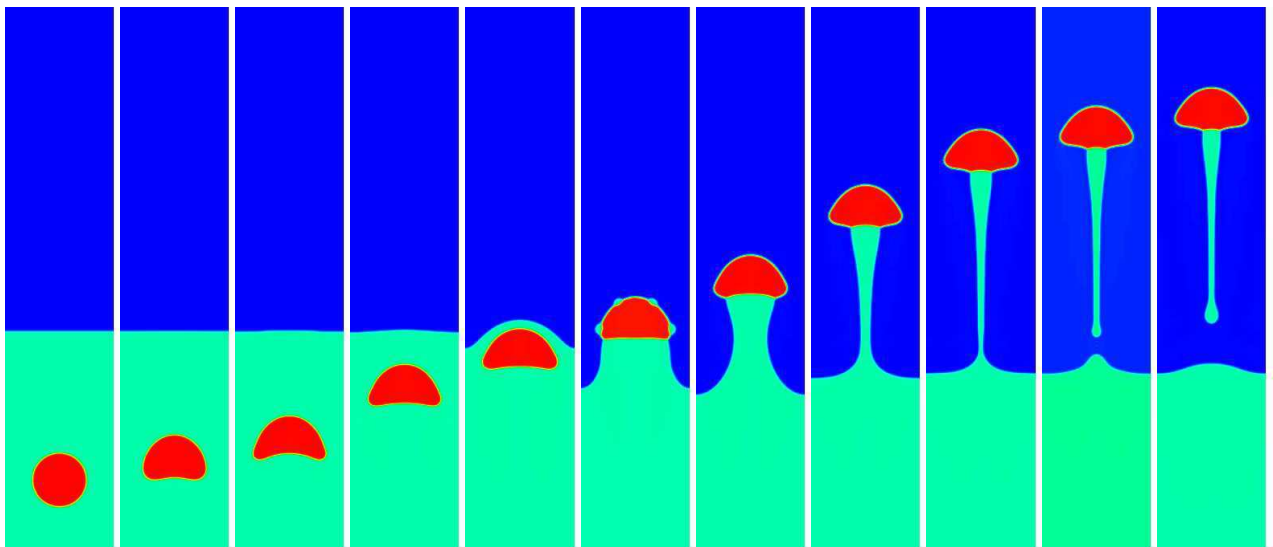
(e) Schéma Implicite, $\Delta t = 7 \times 10^{-4}$ (f) Schéma SIMPL(0.5), $\Delta t = 7 \times 10^{-4}$ (g) Schéma SIMPL(0.5), $\Delta t = 10^{-3}$

FIG. IX.6 – Influence du pas de temps

Encore une fois les résultats présentés dans la figure IX.6 sont très semblables qualitativement. Notons que nous avons rencontré un problème de convergence dans la méthode de Newton avec le schéma Impl pour le pas de temps $\Delta t = 10^{-3}$.

En conclusion, le schéma SImpl même s'il est moins précis que le schéma Impl permet d'obtenir des résultats qualitativement correct tout en garantissant une meilleure robustesse des simulations (notamment pour les grands pas de temps).

IX.1.6 Méthode de projection et Lagrangien augmenté

Nous comparons maintenant les résultats obtenus précédemment (avec une méthode de Lagrangien augmenté) et le résultat obtenu avec la méthode de projection pour un pas de temps de $\Delta t = 5 \times 10^{-4}$. Le système de Cahn-Hilliard est discrétisé par le schéma SImpl(0.5), les solveurs linéaires utilisés sont encore des solveurs directs de la librairie UMFPACK.

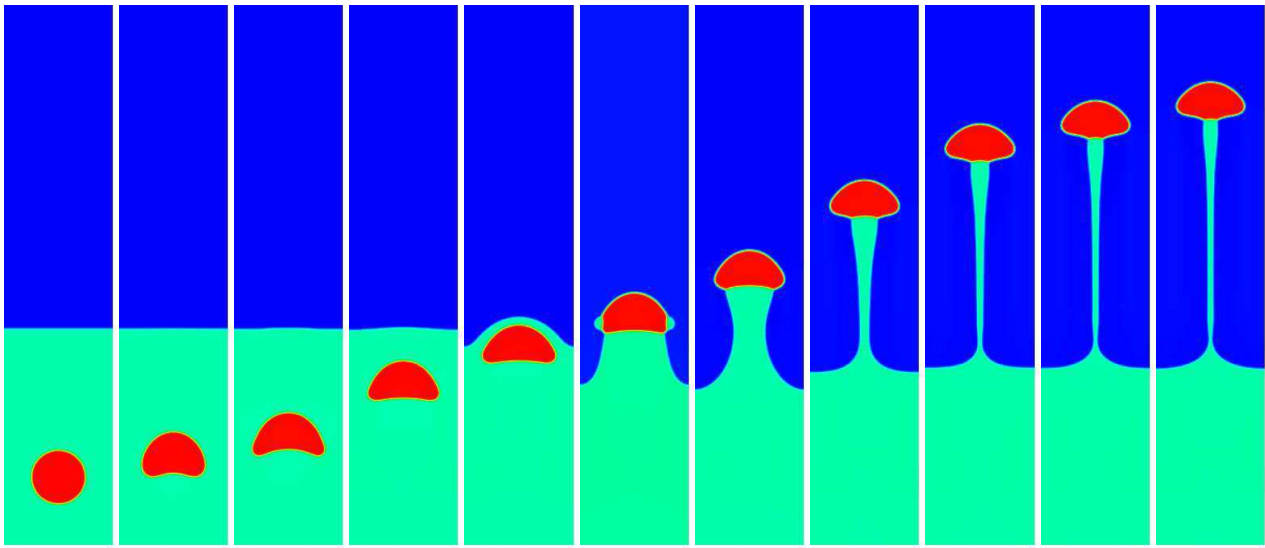


FIG. IX.7 – Méthode de projection

Nous observons des différences plus importantes (temps de rupture de la colonne de liquide entraîné par exemple) même si les résultats restent qualitativement corrects. Ceci s'explique par l'erreur de fractionnement en temps, due au découplage de l'équation de bilan de quantité de mouvement de celle de contrainte de divergence nulle dans la méthode de projection.

Néanmoins, il est à noter que la méthode de projection permet une économie très importante en temps de calcul :

- 14h 15min pour le calcul présenté en figure IX.7 (soit en moyenne 25 secondes par itération en temps),
- 27h 21min pour celui présenté en figure IX.5 (soit en moyenne 50 secondes par itération en temps).

Ceci s'explique par l'ajout du terme d'augmentation qui altère considérablement le stencil de la matrice de rigidité A (de l'équation de bilan de quantité de mouvement). En moyenne, la matrice A comporte 35 coefficients non nuls par ligne alors que la matrice augmentée en comporte 175.

De plus, l'ajout du terme d'augmentation change considérablement le conditionnement de la matrice A et l'utilisation de solveurs itératifs devient difficile (aujourd'hui impossible avec les préconditionneurs dont nous disposons).

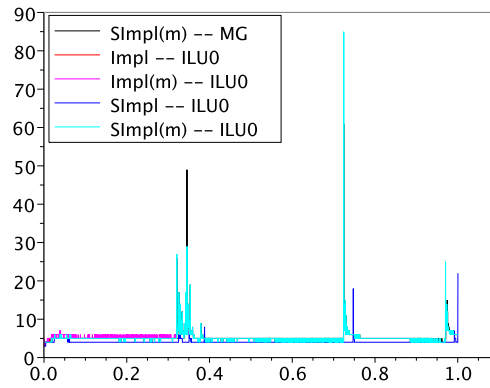
Pour ces raisons, nous choisissons d'utiliser la méthode de projection.

IX.1.7 Solveurs itératifs, raffinement local adaptatif et parallélisme

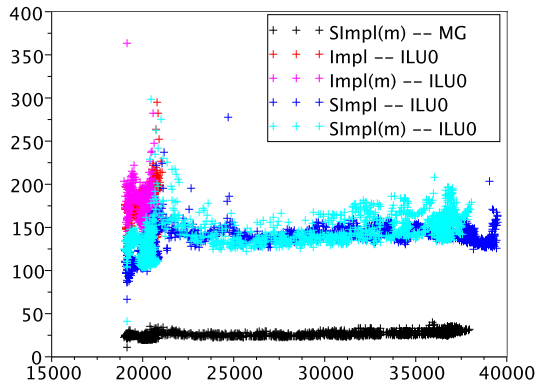
Nous présentons maintenant quelques statistiques pour illustrer le gain apporté par le raffinement local et les préconditionneurs multigrilles.

Résolution du système de Cahn-Hilliard

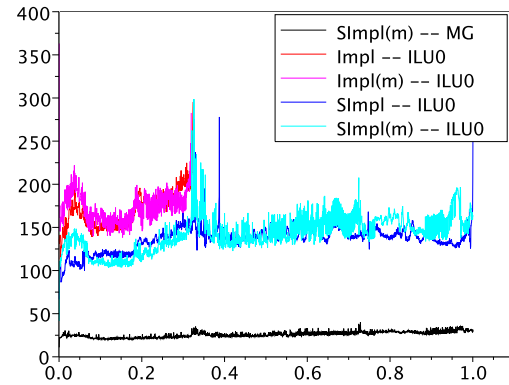
Nous nous intéressons tout d'abord à la résolution du système de Cahn-Hilliard. Nous utilisons le solveur linéaire GMRES, le critère d'arrêt est donné par la norme euclidienne du résidu préconditionné (norme relative inférieure à 10^{-14} ou norme absolue inférieure à 10^{-12}). Nous utilisons le préconditionneur ILU0 de la librairie PELICANS et le préconditionneur multigrille (II.13) (toutes les sous-grilles disponibles sont utilisées et le nombre de pré et post lissages est fixé à 2). La figure IX.8 permet d'identifier l'influence de la discrétisation du terme non linéaire F_0 et des préconditionnements des solveurs linéaires. Nous avons effectué différentes simulations avec les schémas Impl, SImpl(0.5), Impl(m), SImpl(m,0.5) le système de Navier-Stokes étant résolu par la méthode de type Lagrangien augmenté (et à l'aide d'un solveur direct).



(a) Nombre d'itérations du solveur non linéaire (méthode de Newton) en fonction du temps



(b) Nombre d'itérations moyen (pour chaque résolution non linéaire) du solveur linéaire (GMRES) en fonction du nombre d'inconnues



(c) Nombre d'itérations moyen (pour chaque résolution non linéaire) du solveur linéaire (GMRES) en fonction du temps

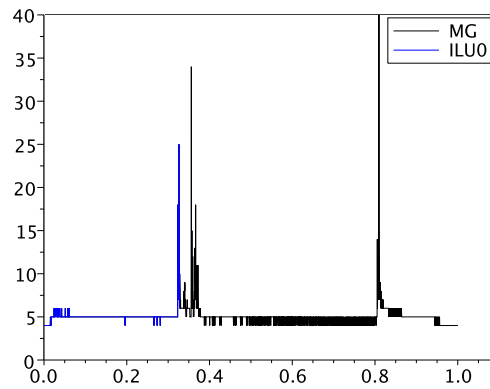
FIG. IX.8 – Influence de la discrétisation de F_0 et du préconditionnement du solveur linéaire (indiqués en légende) sur les statistiques de résolution du système de Cahn-Hilliard (le système de Navier-Stokes est résolu par une méthode de type Lagrangien augmenté)

La figure IX.8a rapporte l'évolution du nombre d'itérations du solveur non linéaire (la méthode de Newton) au cours du temps (*i.e.* le nombre d'itérations effectuées à l'itération n en temps en fonction de t^n). Il est intéressant

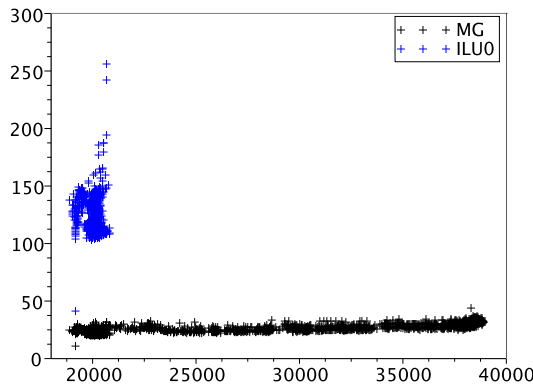
de constater que, pour la plupart des pas de temps n , très peu d'itérations (environ 5) sont nécessaires pour satisfaire le critère de convergence de la méthode de Newton, mais que néanmoins deux instants ($t = 3.2$ et $t = 0.8$) se distinguent par l'apparition d'un pic dans le nombre d'itérations. Ces instants correspondent à la traversée de l'interface liquide-liquide par la bulle pour le premier et à la rupture de la colonne de fluide entraînée pour le second.

Les figures IX.8b et IX.8c donnent le nombre d'itérations dans les solveurs linéaires cette fois-ci. A chaque pas de temps n , une moyenne de ce nombre d'itérations est faite sur l'ensemble des itérations de la méthode de Newton. C'est ce nombre qui est rapporté sur les figures IX.8b et IX.8c, en fonction du nombre d'inconnues à l'instant t^n pour la première et du temps t^n pour la seconde et ce pour les différents schémas numériques testés. Nous constatons qu'avec le schéma Impl, la résolution des systèmes linéaires est plus délicate qu'avec le schéma SImpl. En effet, lorsque nous utilisons les schémas Impl ou Impl(m), le solveur linéaire (GMRES) atteint le nombre d'itérations maximum (fixée à 2000) lorsque la bulle perce l'interface et le calcul ne se déroule pas jusqu'au bout. Le résidu au bout de 2000 itérations est, par exemple pour le schéma Impl., de l'ordre de 10^{-7} .

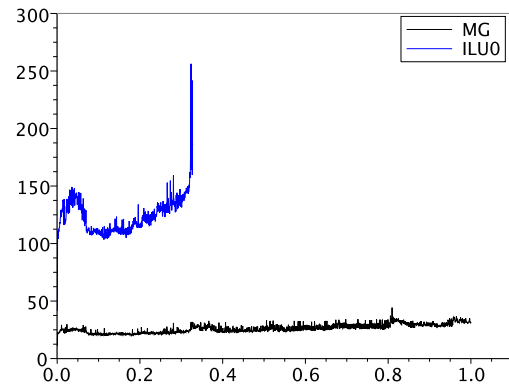
Le schéma SImpl permet de corriger ce problème. Nous observons également que le préconditionneur multigrille permet de réduire considérablement le nombre d'itérations (de 150 à environ 25).



(a) Nombre d'itérations de la méthode de Newton en fonction du temps



(b) Nombre d'itérations moyen (pour chaque résolution non linéaire) du solveur linéaire (GMRES) en fonction du nombre d'inconnues



(c) Nombre d'itérations moyen (pour chaque résolution non linéaire) du solveur linéaire (GMRES) en fonction du temps

FIG. IX.9 – Influence du préconditionnement du solveur linéaire (indiqué en légende) sur les statistiques de résolution du système de Cahn-Hilliard (utilisation du schéma SImpl(m) et résolution du système de Navier-Stokes par la méthode de projection incrémentale)

La fragilité du préconditionneur ILU0 se confirme lorsque nous passons en méthode de projection. La figure

IX.9 présente les mêmes résultats que ceux de la figure IX.8 (en se limitant au schéma $\text{SImpl}(m,0.5)$) mais cette fois le système de Navier-Stokes est résolu par la méthode de projection. Nous observons (simplement en changeant la méthode de résolution du système de Navier-Stokes) une non-convergence du solveur linéaire interne à la résolution du système de Cahn-Hilliard. Le préconditionneur multigrille permet d'éviter ce problème.

Résolution du système de Navier-Stokes

Nous nous intéressons maintenant aux étapes de prédiction de pression et de calcul de l'incrément de pression dans la méthode de projection. Nous donnons dans la figure IX.10, le nombre d'itérations du solveur linéaire (méthode de gradient conjugué) préconditionné par ILU0 ou par la méthode multigrille (II.13) (toutes les sous-grilles disponibles sont utilisées et le nombre de pré et post lissages est fixé à 2) en fonction du nombre d'inconnues. Le critère d'arrêt est donné par la norme euclidienne du résidu préconditionné (norme relative inférieure à 10^{-10} ou norme absolue inférieure à 10^{-15})

Comme attendu, la méthode multigrille permet de diminuer fortement le nombre d'itérations requises pour arriver à convergence (ce nombre d'itérations étant de plus indépendant du nombre d'inconnues).

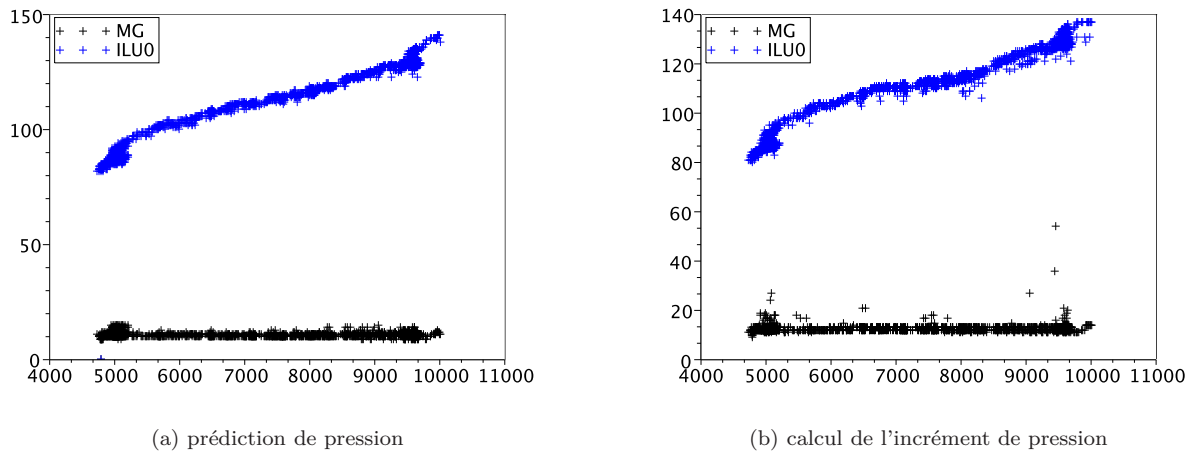


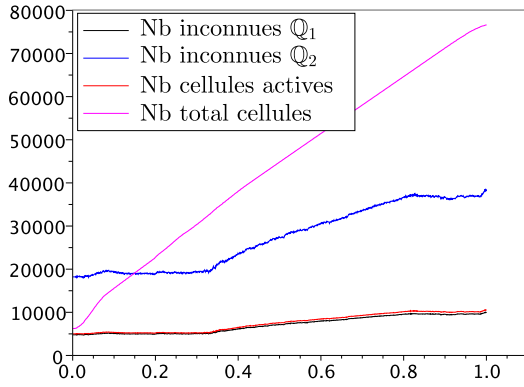
FIG. IX.10 – Nombre d'itérations du solveur linéaire (CG) dans les étapes de prédiction de pression et de calcul de l'incrément de pression

Statistiques complètes

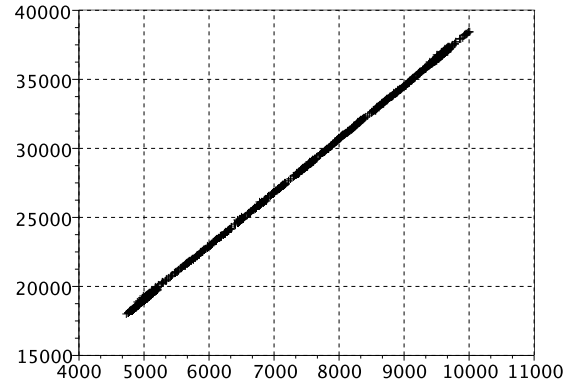
Enfin, nous présentons les statistiques complètes pour une simulation effectuée avec le schéma $\text{SImpl}(m,0.5)$ et la méthode de projection incrémentale donnée par le problème VII.8. Nous utilisons la méthode multigrille (II.13) (toutes les sous-grilles disponibles sont utilisées et le nombre de pré et post lissages est fixé à 2) pour préconditionner le solveur GMRES dans l'inversion des systèmes linéaires à chaque itération de la méthode de Newton dans la résolution du système discret de Cahn-Hilliard, ainsi que pour préconditionner la méthode de gradient conjugué lors de la résolution des étapes de prédiction de pression et de calcul de l'incrément de pression de la méthode de projection. Enfin, l'étape de prédiction de vitesse est résolue en utilisant la méthode GMRES préconditionnée par ILU0 et l'étape de correction de vitesse par la méthode du gradient conjugué préconditionnée par la méthode de Jacobi. Les critères de convergences sont basés sur la norme euclidienne du résidu préconditionné :

- Cahn-Hilliard : norme relative inférieure à 10^{-14} , norme absolue inférieure à 10^{-12} ,
- Prédiction de vitesse : norme relative inférieure à 10^{-8} , norme absolue inférieure à 10^{-15} ,
- Prédiction de pression et calcul de l'incrément de pression : norme relative inférieure à 10^{-10} , norme absolue inférieure à 10^{-15} ,
- Correction de vitesse : norme relative inférieure à 10^{-10} , norme absolue inférieure à 10^{-15} ,

La figure IX.11 permet, tout d'abord, de voir l'évolution du nombre d'inconnues.



(a) Nombre d'inconnues Q_1 et Q_2 (pour un unique champ scalaire) en fonction du temps



(b) Nombre d'inconnues Q_2 en fonction du nombre d'inconnues Q_1

FIG. IX.11 – Nombre d'inconnues

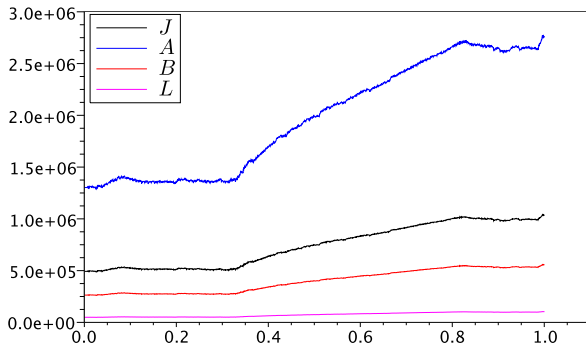
La figure IX.11a (à gauche) représente l'évolution du nombre d'inconnues pour un champ scalaire Q_1 et pour un champ scalaire Q_2 . Pour obtenir le nombre total d'inconnues, il faut multiplier la première quantité par 5 (puisque nous avons 5 champs scalaires Q_1 : 2 paramètres d'ordre, 2 potentiels chimiques et la pression) et ajouter la seconde multipliée par 2 (puisque la vitesse a deux composantes). Ce graphique montre également le nombre de cellules actives (*cf* section II.1.5) qui se comporte essentiellement comme le nombre d'inconnues d'un champ scalaire Q_1 ainsi que le nombre total de cellules créées au cours du calcul (rappelons que chacune des cellules construites n'est jamais détruite mais simplement désactivée).

La figure IX.11b (à droite) représente le nombre d'inconnues associées à un champ scalaire Q_2 en fonction du nombre de celles associées à un champ Q_1 . Nous constatons que le nombre d'inconnues Q_2 est 4 fois plus important que le nombre d'inconnues Q_1 (à l'identique de ce qui se passe sur un maillage conforme sans raffinement local).

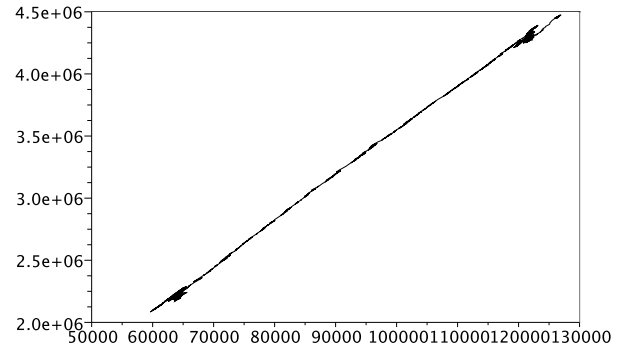
Il est également intéressant de noter la présence de trois comportements distincts dans les courbes présentées dans la figure IX.11a. Pendant une première période, le nombre d'inconnues et le nombre de cellules actives augmentent très peu alors que le nombre total de cellule créées ne cesse d'augmenter. Il s'agit de la période de simulation avant le contact entre la bulle et l'interface. En effet, la zone devant la bulle est raffinée mais ceci (en terme d'inconnues) est compensé par le déraffinement de la zone se trouvant derrière la bulle. Pendant une seconde période, le nombre d'inconnues ainsi que le nombre de cellules actives augmentent, ceci est dû à l'étirement de la zone de raffinement le long de la colonne de fluide entraîné. Une troisième période intervient après la rupture de la colonne, où la situation redevient identique à celle rencontrée lors de la première période.

Pour compléter ces informations, le nombre de cellules actives maximal, *i.e.* 10559, peut être comparé au nombre total de cellules qui auraient été nécessaires pour atteindre la même résolution avec un raffinement global : 81920.

Nous nous intéressons maintenant au nombre d'éléments stockés dans les différentes matrices. La figure IX.12 présente d'une part (à gauche) le nombre d'inconnues stockées dans chacune des matrices : J désigne la Jacobienne associée au système de Cahn-Hilliard, A la matrice de rigidité des équations de Navier-Stokes (utilisée pour l'étape de prédiction de vitesse), L la matrice utilisée dans l'étape de calcul de l'incrément de pression et B la matrice de l'opérateur de "divergence discrète"; d'autre part (à droite) le nombre total d'éléments stockés en fonction du nombre total d'inconnues. Nous constatons que le nombre total d'éléments stockés est bien proportionnel (le coefficient de proportionnalité est environ 35) au nombre d'inconnues (le stencil étant borné indépendamment du nombre de niveaux).



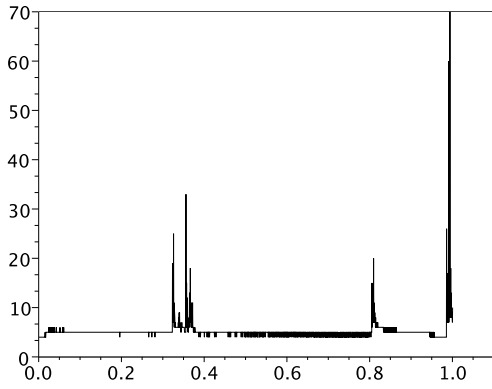
(a) en fonction du temps



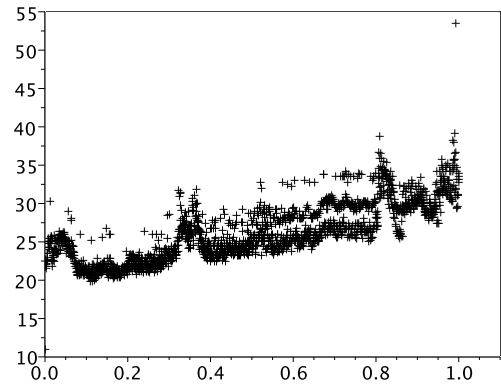
(b) en fonction du nombre total d'inconnues

FIG. IX.12 – Nombre d'éléments stockés dans les matrices

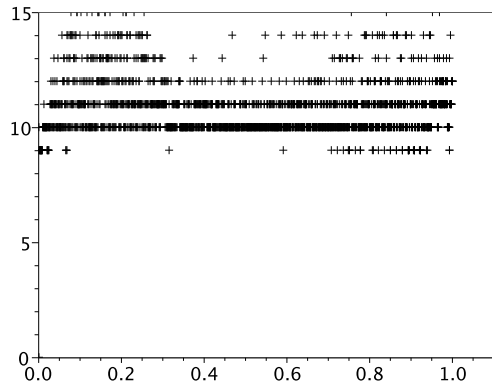
Nous présentons ensuite les statistiques de résolution des différentes étapes en fonction du temps. La figure IX.13 récapitule les itérations nécessaires dans chacune des étapes pour satisfaire les critères de convergence. La figure IX.13a rapporte le nombre d'itérations dans la méthode de Newton. Comme dans le paragraphe "résolution de Cahn-Hilliard" de cette section, nous constatons que ce nombre est très faible en général mais que deux pics apparaissent au moment de la traversée de l'interface par la bulle et de la rupture du film de fluide entraîné (le pic observé en fin de calcul est sans signification puisque la bulle a alors quitté le domaine de calcul). Les figures IX.13b-IX.13f montrent le nombre d'itérations effectuées par les différents solveurs linéaires. Nous constatons que, pour l'ensemble des résolutions, le nombre d'itérations reste faible.



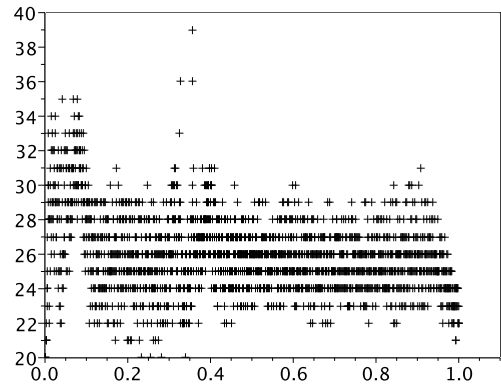
(a) Résolution du système Cahn-Hilliard – Nombre d'itérations du solveur non linéaire (méthode de Newton) en fonction du temps



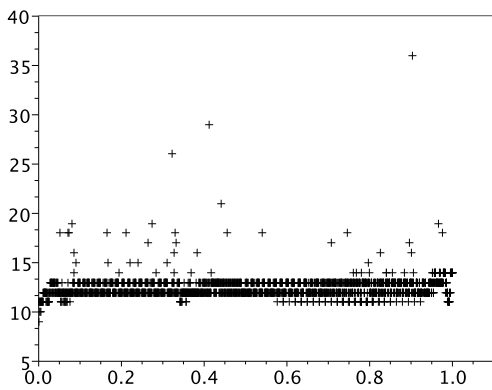
(b) Résolution du système Cahn-Hilliard – Nombre d'itérations moyen (pour chaque résolution non linéaire) du solveur linéaire (GMRES) en fonction du temps



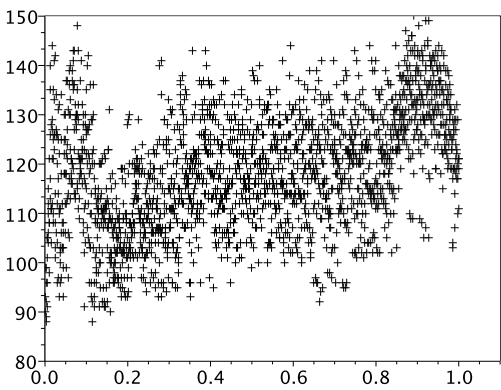
(c) Résolution du système Navier-Stokes – Nombre d'itérations du solveur linéaire (CG) de l'étape de prédiction de pression en fonction du temps



(d) Résolution du système Navier-Stokes – Nombre d'itérations du solveur linéaire (GMRES) de l'étape de prédiction de vitesse en fonction du temps



(e) Résolution du système Navier-Stokes – Nombre d'itérations du solveur linéaire (CG) de l'étape de calcul de l'incrément de pression



(f) Résolution du système Navier-Stokes – Nombre d'itérations du solveur linéaire (CG) de l'étape de correction de vitesse en fonction du temps

FIG. IX.13 – Statistiques de résolution des systèmes de Cahn-Hilliard et Navier-Stokes

Enfin, nous terminons cette section par quelques statistiques concernant l'utilisation de la mémoire et sur la répartition des inconnues entre les différents processus dans le cas d'utilisation des fonctionnalités de parallélisme.

La figure IX.14 présente la mémoire consommée en fonction du temps (figure IX.14a), en fonction du nombre total d'inconnues (figure IX.14b), et en fonction du nombre total de cellules créées (figure IX.14c). La taille mémoire affichée ici est obtenue en fin de pas de temps (après libération de la mémoire occupée par l'ensemble des matrices) par lecture dans le fichier `/proc/pid/status` le champ `VmSize` (correspondant à la mémoire virtuelle occupée par le programme). Il faut être prudent dans l'interprétation de ces courbes car la gestion de la mémoire par le système n'est pas toujours intuitive notamment en ce qui concerne les libérations de l'espace qui ne sont pas toujours visibles, le système gardant à disposition la place correspondante tant qu'un autre programme n'en fait pas la demande. Néanmoins, on constate qu'au début du calcul, la mémoire augmente (de manière non négligeable) alors que le nombre d'inconnues reste fixe. Cette augmentation de presque 500 Mo ne semble alors concerner que l'évolution des structures de données PELICANS nécessaires au fonctionnement du raffinement local.

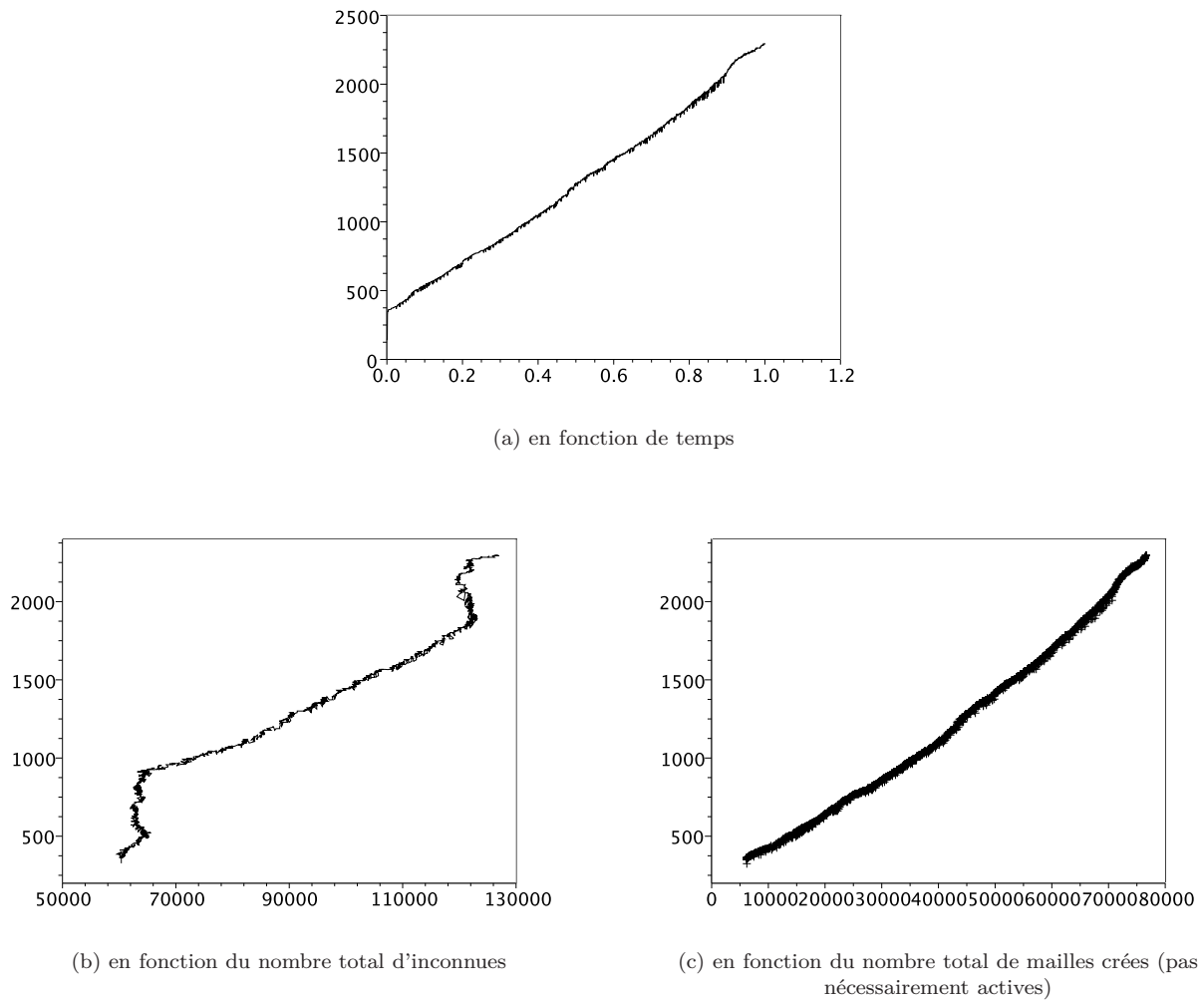


FIG. IX.14 – Utilisation de la mémoire (en Mo)

La figure IX.15 montre la répartition des cellules actives entre les différents processus lors du fonctionnement du calcul en parallèle dans le cas d'un partage entre deux processus et quatre processus.

Rappelons que la répartition est effectuée à l'instant initial par l'utilisateur et qu'ensuite chacun des processus gère le raffinement local de la partie qui lui est attachée.

Dans le cas d'un partage entre deux processus, la répartition est effectuée à l'instant initial suivant l'axe des ordonnées. Nous observons sur la figure IX.15a que cette répartition reste équilibrée tout au long de la simulation puisque l'écoulement (et donc par suite le raffinement local) est symétrique par rapport à cet axe.

Par contre, lorsque la répartition est effectuée sur quatre processus le long cette fois des axes des abscisses et des ordonnées, la répartition devient inégale au cours du temps : en début de calcul ce sont les processus 0 et 1 chargés des parties basses du domaine qui gèrent le plus grand nombre de cellules actives, alors qu'en fin de calcul lorsque la bulle entre dans les parties hautes du domaine, gérées par les processus 3 et 4, la tendance s'inverse.

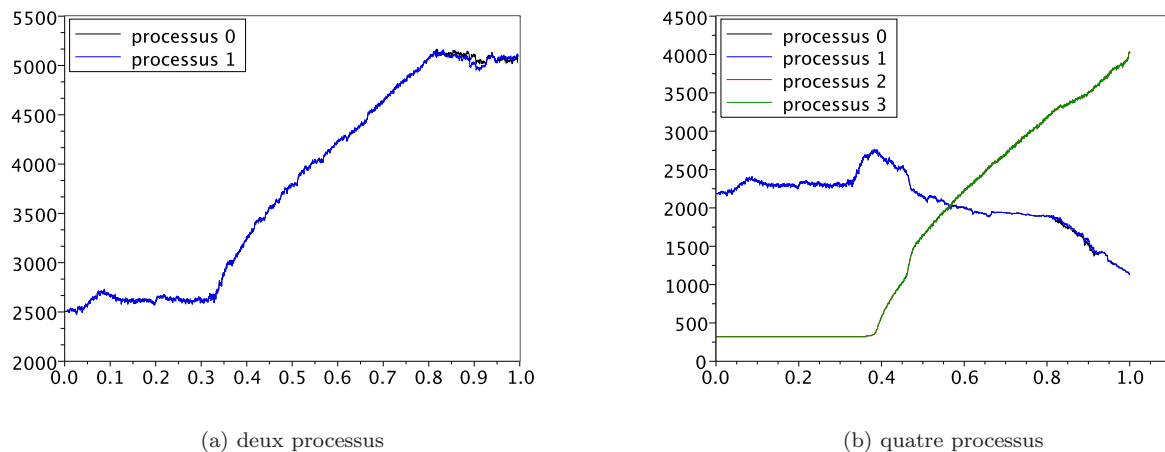


FIG. IX.15 – Répartition du nombre de cellules actives entre les processus en fonction du temps

La table IX.2 donne les temps de calculs observés pour cette simulation effectuée sur 1 processus, 2 processus et 4 processus. Les temps de calcul comprennent les temps d'attente des processus (lorsque des synchronisations sont nécessaires par exemple). Nous constatons que les temps de calculs sont bien réduits grâce à l'utilisation du parallélisme, ils ne sont cependant pas divisés par deux lorsque l'on multiplie le nombre de processus par deux. Ceci s'explique par le mauvais équilibrage de la charge de travail entre les différents processus en cas de raffinement local. Il faut également noter que, dans la préparation des algorithmes multigrilles, la reconstruction des opérateurs sur chacun des niveaux se fait par le calcul de produits de matrices. Ces opérations nécessitent de nombreuses communications entre les processus. Il serait plus efficace d'assembler directement ces opérateurs sur chacun des niveaux puisque ceci nécessite très peu de communication.

	Adaptation	Initialisation	Cahn-Hilliard	Navier-Stokes	Total
1 processus	2h24	1h27	14h47	4h	22h45
2 processus	1h29	54mn	13h47	2h31	18h45
4 processus	1h15	48mn	10h33	1h46	14h27

TAB. IX.2 – Temps de calcul

IX.2 Calcul tridimensionnel

Nous présentons dans cette section la simulation d'une bulle tridimensionnelle traversant une interface liquide-liquide. La configuration initiale est indiquée dans la figure IX.16 et les paramètres physiques rapportés dans la table IX.3. Les conditions aux bords sont de type glissement sur tous les bords du domaine.

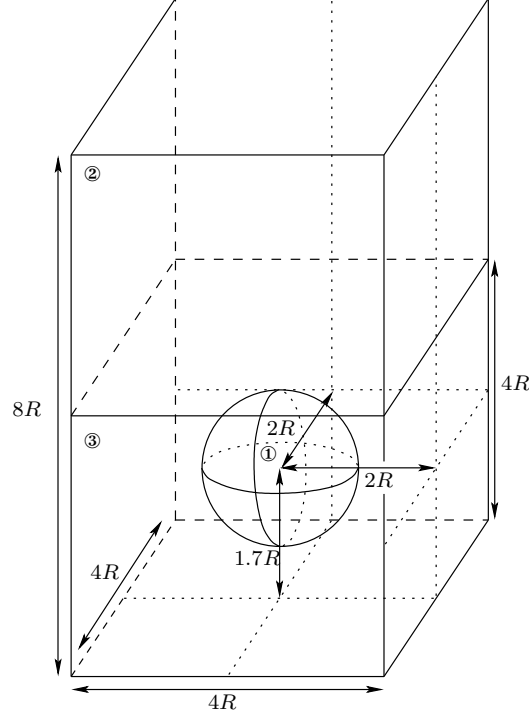


FIG. IX.16 – Configuration initiale

R	T	
0.008	1.	
σ_{12}	σ_{13}	σ_{23}
0.07	0.07	0.05
ϱ_1	ϱ_2	ϱ_3
1	1200	1000
η_1	η_2	η_3
10^{-4}	0.15	0.1

(a) paramètres physiques

ε	Δt	$h_{\text{interface}}$	M_{deg}
$\frac{R}{10}$	10^{-3}	ε	1.1×10^{-5}

(b) paramètres numériques

TAB. IX.3 – Paramètres physiques et numériques

Le domaine de calcul est $] -1.6e - 3; 1.6e - 3[\times] -1.6e - 3; 1.6e - 3[\times] 0; 6.4e - 3[$. Initialement, la bulle est sphérique, son rayon est $R = 8 \times 10^{-3}$ et son centre est placé au point $(0, 0, 0.0136)$. Le maillage initial (structuré) est formé de $4 \times 4 \times 8$ cellules cubiques. Le nombre de niveaux de raffinement est limité à 4. Nous utilisons le schéma SImpl(m,0.5) et la méthode de projection décrite dans le problème VII.8. Nous utilisons le préconditionneur ILU0 pour tous les systèmes linéaires sauf pour le calcul de l'incrément de pression pour lequel nous utilisons le préconditionneur multigrille (II.13) (toutes les sous-grilles disponibles sont utilisées et le nombre de pré et post lissages est fixé à 2).

La figure IX.17 présente le résultat obtenu en proposant trois vues différentes : une vue en coupe des paramètres d'ordre accompagnée du maillage ; les deux autres vues montrent les lignes de niveau ($c_i = 0.5$) des paramètres d'ordre, une vue de dessus au milieu et une vue de dessous en bas.

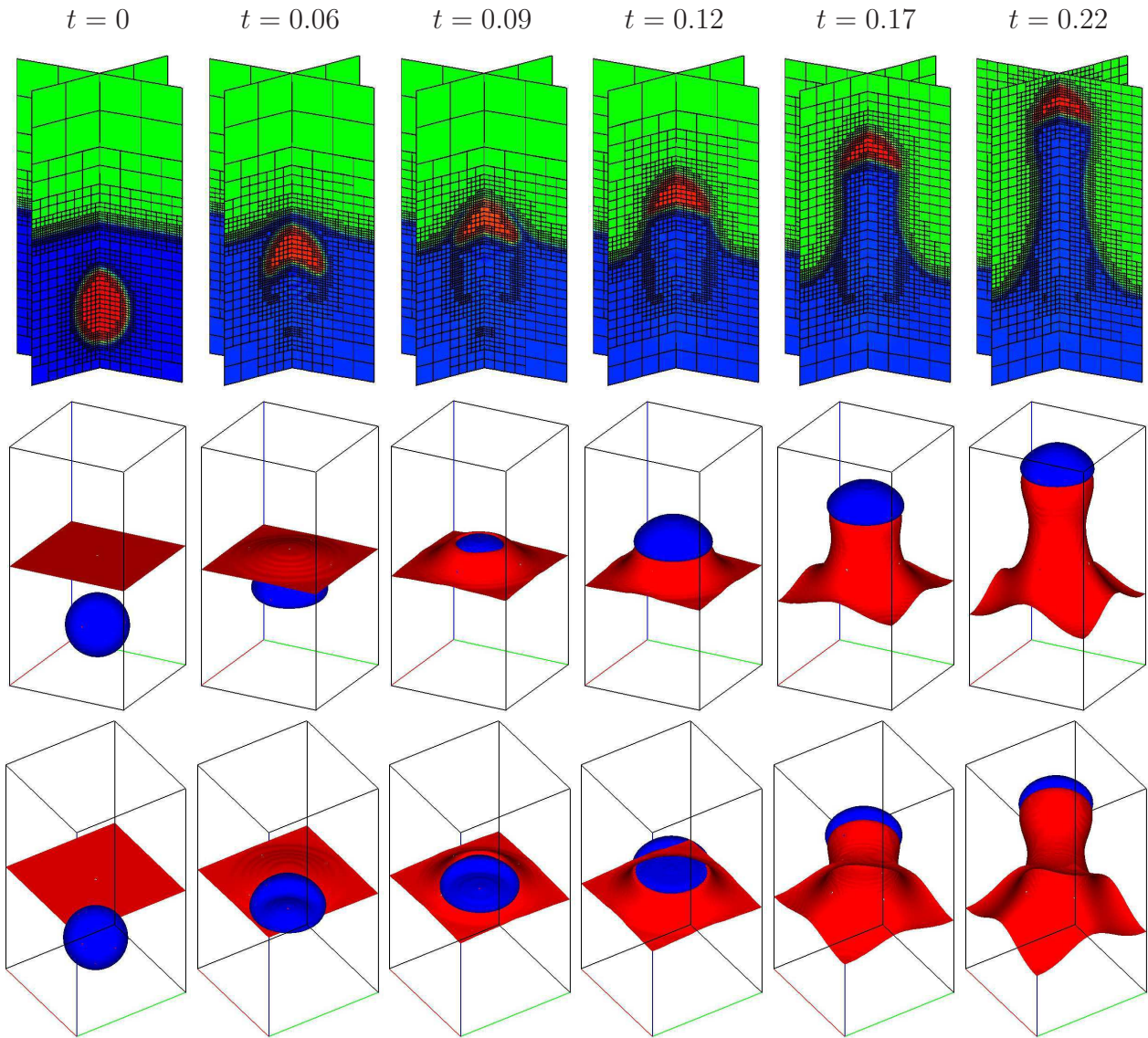


FIG. IX.17 – Exemple de montée de bulle 3D triphasiques, paramètres d'ordre : vue en coupes (en haut), vue de dessus (au milieu), vue de dessous (en bas)

Le nombre d'inconnues maximal pour un champ scalaire $\mathbb{Q}_1$ est ici de 116 810. Le raffinement local permet d'apporter un gain important puisque une simulation sur un maillage globalement raffiné de résolution équivalente dans l'interface aurait comporté 545 052 inconnues.

La charge de travail a été partagée entre 32 processus chacun gérant initialement 4 cellules grossières et le temps d'exécution observé (*i.e.* le temps écoulé pour obtenir les sauvegardes) a été de 4 jours et 8 heures (ce qui correspond en moyenne à environ 25 minutes par itération en temps). Ici encore, le temps de calcul n'est donné qu'à titre indicatif puisqu'il dépend bien sûr de la machine utilisée et de l'encombrement du réseau.

Conclusions et perspectives

Nous avons abordé dans ce manuscrit différents aspects de la simulation d'écoulements incompressibles triphasiques à l'aide d'un modèle à interfaces diffuses.

Modèle triphasique de type Cahn-Hilliard/Navier-Stokes

Le modèle de type Cahn-Hilliard/Navier-Stokes triphasique (introduit dans la thèse de C. Lapuerta [Lap06]) que nous avons étudié dans ce manuscrit comporte différents paramètres “non-objectifs” : l'épaisseur d'interface ε , la mobilité M_0 , et la forme des termes d'ordre élevé $\Lambda(\mathbf{c})c_1^2c_2^2c_3^2$ du potentiel F de Cahn-Hilliard.

▷ L'épaisseur d'interface ε a peu d'influence sur les résultats obtenus, de manière assez intuitive son influence se ressent surtout sur les instants de coalescence ou de rupture.

▷ La mobilité M_0 est sans aucun doute le paramètre du modèle le plus délicat à fixer. Nous avons proposé une étude de son influence sur un cas test diphasique issu du benchmark [HTK⁺09]. Lorsque le coefficient de mobilité a une valeur trop importante, la forme de bulle observée n'est pas correcte (la bulle se déforme peu, elle conserve une forme circulaire). Lorsque sa valeur est trop faible, le profil de l'interface n'est pas maintenu à une épaisseur fixe. En revanche, il apparaît une plage de valeurs de la mobilité où les résultats obtenus sont qualitativement et quantitativement très proches et en correspondance avec ceux obtenus par les participants au benchmark. Ces observations sont à confirmer sur des cas tests où les déformations des interfaces sont plus importantes (même si le travail est alors plus délicat parce que nous ne disposons plus de solution de référence pour valider les résultats). Dans ces cas, il peut être utile de faire varier la mobilité en fonction de la valeur de la norme de la vitesse en chaque point du domaine.

▷ Concernant la forme des termes d'ordre élevé $\Lambda(\mathbf{c})c_1^2c_2^2c_3^2$ du potentiel de Cahn-Hilliard, nous avons proposé une expression qui permet de réduire l'influence de ce terme tout en maintenant les propriétés mathématiques du modèle. Ce travail en collaboration avec R. Bonhomme est encore en cours. Les premiers résultats obtenus sur des tests académiques sont encourageants. Il reste à effectuer un travail sur la discrétisation en temps de ces termes (des difficultés de convergence dans la méthode de Newton apparaissent parfois) et surtout un travail de comparaison des résultats obtenus avec ceux d'expérimentations pour confirmer que cette approche convient.

Raffinement local adaptatif pour des éléments finis conformes de type Lagrange

Nous avons mis au point et développé dans la plateforme PELICANS une méthode de raffinement local adaptatif pour des éléments finis conformes de type Lagrange. Au vu des fortes variations des paramètres d'ordre c_i , celle-ci permet d'apporter un gain conséquent sur le nombre de degrés de liberté nécessaires pour effectuer une simulation.

▷ Un point que nous n'avons pas abordé est d'envisager des algorithmes d'assemblage plus efficaces tenant compte du fait que, à chaque pas de temps, la procédure d'adaptation modifie peu les espaces d'approximation. Il est sûrement envisageable de ne pas recalculer l'ensemble des coefficients assemblés et ainsi gagner en performance.

▷ Par ailleurs, un problème important dans l'exécution du raffinement local en parallèle est la répartition de la charge de travail entre les processus. Celle-ci est aujourd'hui effectuée par un partitionnement du maillage

initial. Ce partitionnement est figé pour toute la durée du calcul et en conséquence la charge de travail peut se déséquilibrer si plus d'étapes de raffinement sont effectuées dans une partie du domaine affectée à un processus que dans celles affectées aux autres. Il pourrait être envisageable de redistribuer la charge de travail en cours de simulation, en fonction de l'évolution de l'écoulement.

Préconditionnement multigrille

Nous avons écrit un algorithme qui permet à partir d'un espace d'approximation multiniveau de reconstruire une hiérarchie d'espaces pouvant être utilisée pour la construction de preconditionneurs multigrilles. Cette méthode permet de réduire considérablement le nombre d'itérations dans la résolution des systèmes linéaires (sous itérations dans le solveur de Newton pour la résolution du système de Cahn-Hilliard ou calcul de l'incrément de pression dans la méthode de projection incrémentale).

▷ La structure du module multigrille dans la librairie PELICANS permet aujourd'hui d'implémenter d'autres lisseurs à moindre frais. Par exemple, il pourrait être envisageable d'utiliser pour l'équation de Cahn-Hilliard un lisseur par blocs tenant compte du fait que le système couple plusieurs inconnues.

▷ Cependant, avec l'implémentation actuelle, le multigrille n'a pas fait ses preuves en ce qui concerne les temps de calcul. Cette méthode a en effet des coûts de calcul cachés : reconstruction de la hiérarchie mais surtout calcul des opérateurs approchés sur chacun des sous-niveaux. Dans l'implémentation actuelle, ces calculs sont effectués à l'aide de produits matrice-matrice, ceci est particulièrement coûteux (surtout en parallèle). De nombreux développements pourraient être envisagés :

- une première possibilité est d'assembler directement les opérateurs approchés sur les différents niveaux,
- il est également envisageable de conserver des informations d'un pas de temps à l'autre. Les sous-niveaux changent sûrement très peu et il doit être possible de faire un calcul incrémental des opérateurs,
- une autre stratégie possible est de ne jamais construire les opérateurs approchés en implémentant seulement leur action sur des vecteurs. Ceci peut alors être effectué par des transferts de vecteurs entre les différents niveaux et des produits matrice-vecteur sur le niveau le plus fin.

Discretisation en espace du système de Cahn-Hilliard/Navier-Stokes

Il pourrait être utile de s'intéresser à d'autres types de discrétisation en espace afin de mieux équilibrer l'ordre d'approximation de la vitesse et celui des paramètres d'ordre. Les choix effectués actuellement reposent sur les deux arguments suivants :

- la conservation du volume ou le fait que la somme des paramètres d'ordre reste égale à 1 repose sur l'égalité :

$$\int_{\Omega} c_i^{n+1} \operatorname{div}(\mathbf{u}^n) dx = 0.$$

Autrement dit, il faut que la vitesse soit à divergence nulle contre les paramètres d'ordre. Par conséquent, ceux-ci doivent appartenir à l'espace d'approximation de la pression. Cette condition n'est plus strictement nécessaire avec le schéma inconditionnellement stable proposé dans le chapitre VI.

- l'équilibrage entre le terme de pression et celui de force capillaire (lié au problème des courants parasites) est vrai seulement si, encore une fois, les paramètres d'ordre appartiennent à l'espace d'approximation de la pression. Des techniques du type de celles présentées dans l'article [GLBB97] pourraient réduire l'importance de ce problème lorsque les espaces d'approximation des paramètres d'ordre et de la pression sont différents.

▷ La possibilité d'utiliser des éléments finis $\mathbb{Q}_2$ (ou $\mathbb{P}_2$) pour approcher les paramètres d'ordre et les potentiels chimiques tout en conservant une approximation $\mathbb{P}_1$ pour la pression est donc maintenant envisageable. Ceci pourrait mener à considérer des maillages plus grossiers et ainsi obtenir un coût de résolution du système de Navier-Stokes moins important.

▷ *A contrario*, il pourrait être également possible de réduire l'ordre d'approximation du couple vitesse/pression en considérant des éléments finis de degré moins élevé (par exemple $\mathbb{Q}_2 \text{iso} \mathbb{Q}_1$ ou $\mathbb{P}_2 \text{iso} \mathbb{P}_1$), ou des techniques du type volumes finis.

Discretisation en temps du système de Cahn-Hilliard/Navier-Stokes

▷ La discrétisation semi-implicite des termes non linéaires du système de Cahn-Hilliard semble être un bon compromis entre stabilité et précision. Ce schéma est, dans les cas trois phases, moins précis que le schéma

implicite, mais son utilisation permet de réduire les difficultés de convergence observées dans la méthode de Newton et mène à des systèmes linéaires (sous-itérations de la méthode de Newton) plus faciles à résoudre à l'aide de solveurs itératifs (dont l'utilisation est nécessaire en trois dimensions).

▷ La discrétisation du système couplé Cahn-Hilliard/Navier-Stokes inconditionnellement stable proposée dans le chapitre VI permet d'obtenir les propriétés théoriques que nous souhaitons voir vérifier par les solutions discrètes tout en préservant la possibilité de résoudre les problèmes de Cahn-Hilliard et Navier-Stokes discrets de manière découplée. Néanmoins, ce schéma reste à tester numériquement. Par ailleurs, d'un point de vue théorique, une perspective intéressante est l'étude de convergence du schéma dans le cas non homogène.

Perspectives plus générales

Enfin, nous terminons par donner quelques perspectives plus larges.

▷ La validation de l'approche présentée ici par la comparaison aux expériences ou aux résultats obtenus avec d'autres codes de calcul reste à faire dans le cas triphasique. A ce sujet, nous mentionnons la thèse de R. Bonhomme qui est en cours et dans laquelle des expérimentations de bulles de gaz traversant une interface liquide-liquide sont réalisées.

▷ L'extension du principe de consistance (idée directrice de l'obtention du modèle triphasique) au cas d'écoulement à quatre phases est encore une question ouverte. De même que la pertinence physique de la condition sur les Σ_i : $\Sigma_1\Sigma_2 + \Sigma_1\Sigma_3 + \Sigma_2\Sigma_3 > 0$.

Bibliographie

- [Abe09a] Helmut Abels, *Existence of weak solutions for a diffuse interface model for viscous, incompressible fluids with general densities*, Comm. Math. Phys. **289** (2009), no. 1, 45–73.
- [Abe09b] ———, *On a diffuse interface model for two-phase flows of viscous, incompressible fluids with matched densities*, Arch. Ration. Mech. Anal. **194** (2009), no. 2, 463–506.
- [Ada01] Mark F. Adams, *A distributed memory unstructured gauss-seidel algorithm for multigrid smoothers*, Supercomputing '01 : Proceedings of the 2001 ACM/IEEE conference on Supercomputing (CDROM), 2001, pp. 4–4.
- [AJL09] Ph. Angot, M. Jobelin, and J.-C. Latché, *Error analysis of the penalty-projection method for the time dependent Stokes equations*, Int. J. Finite Vol. **6** (2009), no. 1, 26.
- [AMW98] D. M. Anderson, G. B. McFadden, and A. A. Wheeler, *Diffuse-interface methods in fluid mechanics*, Annual review of fluid mechanics, Vol. 30, Annu. Rev. Fluid Mech., vol. 30, Annual Reviews, Palo Alto, CA, 1998, pp. 139–165.
- [Bän91] E. Bänsch, *Local mesh refinement in 2 and 3 dimensions*, Impact Comput. Sci. Engrg. **3** (1991), no. 3, 181–191.
- [BB99] John W. Barrett and James F. Blowey, *Finite element approximation of the Cahn-Hilliard equation with concentration dependent mobility*, Math. Comp. **68** (1999), no. 226, 487–517.
- [BBG99] John W. Barrett, James F. Blowey, and Harald Garcke, *Finite element approximation of the Cahn-Hilliard equation with degenerate mobility*, SIAM J. Numer. Anal. **37** (1999), no. 1, 286–318 (electronic).
- [BBG01] J. W. Barrett, J. F. Blowey, and H. Garcke, *On fully practical finite element approximations of degenerate Cahn-Hilliard systems*, Math. Model. Numer. Anal. **35** (2001), no. 4, 713–748.
- [BDY88] R. E. Bank, T. F. Dupont, and H. Yserentant, *The hierarchical basis multigrid method*, Numerische Mathematik **52** (1988), 427–458.
- [BE92] J. F. Blowey and C. M. Elliott, *The Cahn-Hilliard gradient theory for phase separation with nonsmooth free energy. II. Numerical analysis*, European J. Appl. Math. **3** (1992), no. 2, 147–179.
- [Bey95] J. Bey, *Tetrahedral grid refinement*, Computing **55** (1995), no. 4, 355–378.
- [BF06] Franck Boyer and Pierre Fabrie, *Éléments d'analyse pour l'étude de quelques modèles d'écoulements de fluides visqueux incompressibles*, Mathématiques & Applications (Berlin) [Mathematics & Applications], vol. 52, Springer-Verlag, Berlin, 2006.
- [BL06] F. Boyer and C. Lapuerta, *Study of a three component Cahn-Hilliard flow model*, M2AN **40** (2006), no. 4, 653–687.
- [BLM⁺] F. Boyer, C. Lapuerta, S. Minjeaud, B. Piar, and M. Quintard, *Cahn-Hilliard / Navier-Stokes model for the simulation of three-phase flows*, Transp. Porous Media, to appear.
- [BLMP09] F. Boyer, C. Lapuerta, S. Minjeaud, and B. Piar, *A local adaptive refinement method with multigrid preconditioning illustrated by multiphase flows simulations*, ESAIM Proceedings **27** (2009), 15–53.
- [BLP07] F. Boyer, C. Lapuerta, and B. Piar, *Simulation cahn-hilliard/navier-stokes du passage d'une bulle à travers une interface liquide-liquide*, 18^e congrès de mécanique française (2007).

- [BM] F. Boyer and S. Minjeaud, *Fully discrete approximation of a three component Cahn-Hilliard model*, ALGORITHMY, Slovakia, March 2009.
- [BM00] A. Benkenida and J. Magnaudet, *A method for the simulation of two-phase flows without interface reconstruction*, Comptes Rendus de l'Académie des Sciences - Series IIB - Mechanics-Physics-Astronomy **328** (2000), no. 1, 25–32.
- [BM07] Thomas Bonometti and Jacques Magnaudet, *An interface-capturing method for incompressible two-phase flows. validation and application to bubble dynamics*, International Journal of Multiphase Flow **33** (2007), no. 2, 109–133.
- [BM10] F. Boyer and S. Minjeaud, *Numerical schemes for a three component Cahn-Hilliard model*, M2AN (2010), in revision.
- [Boy02] F. Boyer, *A theoretical and numerical model for the study of incompressible mixture flows*, Comput. Fluids **31** (2002), no. 1, 41–68.
- [BPX90] J. H. Bramble, J. E. Pasciak, and J. Xu, *Parallel multilevel preconditioners*, Mathematics of Computation **55** (1990), no. 191, 1–22.
- [BS08] Susanne C. Brenner and L. Ridgway Scott, *The mathematical theory of finite element methods*, third ed., Texts in Applied Mathematics, vol. 15, Springer, New York, 2008.
- [BSW83] R. E. Bank, A. H. Sherman, and A. Weiser, *Some refinement algorithms and data structures for regular local mesh refinement*, Scientific Computing, Applications of Mathematics and Computing to the Physical Sciences, 1983.
- [BW81] D. Bhaga and M. E. Weber, *Bubbles in viscous liquids : shapes, wakes and velocities*, Journal of Fluid Mechanics **105** (1981), no. -1, 61–85.
- [BZ00] J. H. Bramble and X. Zhang, *The analysis of multigrid methods*, Handbook of numerical analysis, VOL. VII (P.G. Ciarlet and J.L. Lions), 2000.
- [CGW78] R. Clift, J. R. Grace, and M. E. Weber, *Bubbles, drops and particles*, Academic Press, New York (1978).
- [Cra07] M. Cranga, *Status of knowledge and updating of the strategy of r&d on molten corium concrete interaction.*, Tech. report, SEMIC/LETR, 2007.
- [Dei85] K. Deimling, *Nonlinear functional analysis*, Springer-Verlag, 1985.
- [DLM10a] F. Dardalhon, J.C. Latché, and S. Minjeaud, *Analysis of a projection method for low order nonconforming finite elements*, in preparation.
- [DLM10b] ———, *On a projection method for piecewise-constant pressure nonconforming finite elements*, Proceedings of MFD2010 the International congress in mathematical fluid dynamics and its applications.
- [EG97] Charles M. Elliott and Harald Garcke, *Diffusional phase transitions in multicomponent systems with a concentration dependent mobility matrix*, Phys. D **109** (1997), no. 3-4, 242–256.
- [EG04] A. Ern and J.-L. Guermond, *Theory and practice of finite elements*, Applied Mathematical Sciences, vol. 159, Springer, 2004.
- [EL91] C.M. Elliott and S. Luckhaus, *A generalised diffusion equation for phase separation of a multi-component mixture with interfacial free energy*, IMA Preprint Series # **887** (1991).
- [Ell89] C. M. Elliott, *The Cahn-Hilliard model for the kinetics of phase separation*, Mathematical models for phase change problems (Óbidos, 1988), Internat. Ser. Numer. Math., vol. 88, Birkhäuser, Basel, 1989, pp. 35–73.
- [Eyr93] D. J. Eyre, *Systems of Cahn-Hilliard equations*, SIAM J. Appl. Math. **53** (1993), no. 6, 1686–1712.
- [Eyr98] ———, *Unconditionally gradient stable time marching the Cahn-Hilliard equation*, Computational and mathematical models of microstructural evolution (San Francisco, CA, 1998), Mater. Res. Soc. Sympos. Proc., vol. 529, MRS, Warrendale, PA, 1998, pp. 39–46.
- [Fen06] X. Feng, *Fully discrete finite element approximations of the Navier-Stokes-Cahn-Hilliard diffuse interface model for two-phase fluid flows*, SIAM J. Numer. Anal. **44** (2006), no. 3, 1049–1072.
- [GJ05] S. Ganesan and V. John, *Pressure separation—a technique for improving the velocity error in finite element discretisations of the Navier-Stokes equations*, Appl. Math. Comput. **165** (2005), no. 2, 275–290.

- [GKS02] Eitan Grinspun, Petr Krysl, and Peter Schröder, *CHARMS : A simple framework for adaptive simulation*, 29th International Conference on Computer Graphics and Interactive Techniques SIGGRAPH, 2002.
- [GLBB97] J.-F. Gerbeau, C. Le Bris, and M. Bercovier, *Spurious velocities in the steady flow of an incompressible fluid subjected to external forces*, Internat. J. Numer. Methods Fluids **25** (1997), no. 6, 679–695.
- [GMS05] J.-L. Guermond, P. Mineev, and J. Shen, *Error analysis of pressure-correction schemes for the time-dependent Stokes equations with open boundary conditions*, SIAM Journal on Numerical Analysis **43** (2005), no. 1, 239–258.
- [GNS00] H. Garcke, B. Nestler, and B. Stoth, *A multiphase field concept : numerical simulations of moving phase boundaries and multiple junctions*, SIAM J. Appl. Math. **60** (2000), no. 1, 295–315.
- [God79] K. Goda, *A multistep technique with implicit difference schemes for calculating two- or three-dimensional cavity flows*, Journal of Computational Physics **30** (1979), 76–95.
- [GQ98] J.-L. Guermond and L. Quartapelle, *On the approximation of the unsteady Navier-Stokes equations by finite element projection methods*, Numer. Math. **80** (1998), no. 2, 207–238.
- [GQ00] ———, *A projection FEM for variable density incompressible flows*, J. Comput. Phys. **165** (2000), no. 1, 167–188.
- [GS06] H. Garcke and B. Stinner, *Second order phase field asymptotics for multi-component systems*, Interface and Free boundaries **8** (2006), 131–157.
- [Gue96] J.-L. Guermond, *Some implementations of projection methods for Navier-Stokes equations*, Mathematical Modelling and Numerical Analysis **30** (1996), no. 5, 637–667.
- [Gue99] ———, *Un résultat de convergence d'ordre deux en temps pour l'approximation des équations de Navier-Stokes par une technique de projection incrémentale*, Mathematical Modelling and Numerical Analysis **33** (1999), no. 1, 169–189.
- [Hac85] W. Hackbusch, *Multigrid methods and applications*, Computational Mathematics, vol. 4, Springer-Verlag, Berlin, 1985.
- [HK03] P. Hammon and P. Krysl, *Implementation of a general mesh refinement technique*, Ctu reports, Czech Technical University, Prague, 2003.
- [HTK⁺09] S. Hysing, S. Turek, D. Kuzmin, N. Parolini, E. Burman, S. Ganesan, and L. Tobiska, *Quantitative benchmark computations of two-dimensional bubble dynamics*, Internat. J. Numer. Methods Fluids **60** (2009), no. 11, 1259–1288.
- [Jac99] D. Jacqmin, *Calculation of two-phase Navier-Stokes flows using phase-field modeling*, J. Comput. Phys. **155** (1999), no. 1, 96–127.
- [JLL⁺06] M. Jobelin, C. Lapuerta, J.-C. Latché, Ph. Angot, and B. Piar, *A finite element penalty-projection method for incompressible flows*, Journal of Computational Physics **217** (2006), 502–518.
- [JTB02] D. Jamet, D. Torres, and J. U. Brackbill, *On the theory and computation of surface tension : The elimination of parasitic currents through energy conservation in the second-gradient method*, Journal of Computational Physics **182** (2002), no. 1, 262 – 276.
- [KGS03] P. Krysl, E. Grinspun, and P. Schröder, *Natural hierarchical refinement for finite element methods*, Internat. J. Numer. Methods Engrg. **56** (2003), no. 8, 1109–1124.
- [KKL04a] J. Kim, K. Kang, and J. Lowengrub, *Conservative multigrid methods for Cahn-Hilliard fluids*, J. Comput. Phys. **193** (2004), no. 2, 511–543.
- [KKL04b] ———, *Conservative multigrid methods for ternary Cahn-Hilliard systems*, Commun. Math. Sci. **2** (2004), no. 1, 53–77.
- [KL05] J. Kim and J. Lowengrub, *Phase field modeling and simulation of three-phase flows*, Interfaces Free Bound. **7** (2005), no. 4, 435–466.
- [KSW08] D. Kay, V. Styles, and R. Welford, *Finite element approximation of a Cahn-Hilliard-Navier-Stokes system*, Interfaces Free Bound. **10** (2008), no. 1, 15–43.
- [KTZ04] P. Krysl, A. Trivedi, and B. Zhu, *Object-oriented hierarchical mesh refinement with CHARMS*, International Journal for Numerical Methods in Engineering **60** (2004), 1401–1424.

- [Lap06] C. Lapuerta, *Echanges de masse et de chaleur entre deux phases liquide stratifiées dans un écoulement à bulles*, Thèse de Mathématiques Appliquées, <http://tel.archives-ouvertes.fr/tel-00132564>.
- [LS03] Chun Liu and Jie Shen, *A phase field model for the mixture of two incompressible fluids and its approximation by a Fourier-spectral method*, Phys. D **179** (2003), no. 3-4, 211–228.
- [LW07] Chun Liu and Noel J. Walkington, *Convergence of numerical approximations of the incompressible Navier-Stokes equations with variable density and viscosity*, SIAM J. Numer. Anal. **45** (2007), no. 3, 1287–1304 (electronic).
- [Mau95] J. M. Maubach, *Local bisection refinement for n -simplicial grids generated by reflection*, SIAM Journal on Scientific Computing **16** (1995), no. 1, 210–227.
- [Min10] S. Minjeaud, *An unconditionally stable uncoupled scheme for the approximation of a triphasic cahn-hilliard/navier-stokes system*, in preparation.
- [Mit91] W. F. Mitchell, *Adaptive refinement for arbitrary finite-element spaces with hierarchical bases*, J. Comput. Appl. Math. **36** (1991), no. 1, 65–78.
- [MP10] S. Minjeaud and B. Piar, *An adaptive pressure correction method without spurious velocities for diffuse-interface models of incompressible flows*, in preparation.
- [NGS05] Britta Nestler, Harald Garcke, and Björn Stinner, *Multicomponent alloy solidification : Phase-field modeling and simulations*, Phys. Rev. E **71** (2005), no. 4, 041609.
- [RI96] M.-C. Rivara and G. Iribarren, *The 4-triangles longest-side partition of triangles and linear refinement algorithms*, Math. Comp. **65** (1996), no. 216, 1485–1502.
- [Riv84] M.-C. Rivara, *Mesh refinement processes based on the generalized bisection of simplices*, SIAM J. Numer. Anal. **21** (1984), no. 3, 604–613.
- [RT88] P.A. Raviart and J.-M. Thomas, *Introduction à l'analyse numérique des équations aux dérivées partielles*, 2e tirage ed., MASSON, 1988.
- [RT92] R. Rannacher and S. Turek, *Simple nonconforming quadrilateral Stokes element*, Numerical Methods for Partial Differential Equations **8** (1992), 97–111.
- [RW82] J. S. Rowlinson and B. Widom, *Molecular theory of capillarity*, Clarendon Press, 1982.
- [She92] J. Shen, *On error estimates of projection methods for Navier-Stokes equations : First-order schemes*, SIAM Journal on Numerical Analysis **29** (1992), no. 1, 57–77.
- [Sim87] J. Simon, *Compact sets in the space $L^p(0, T; B)$* , Ann. Mat. Pura Appl. (4) **146** (1987), 65–96.
- [SVdV00] Y. Saad and H. A. Van de Vorst, *Iterative solution of linear systems in the 20th century*, Journal of Computational and Applied Mathematics **123** (2000), 1–33.
- [SZ99] Ruben Scardovelli and Stéphane Zaleski, *Direct numerical simulation of free-surface and interfacial flow*, Annual review of fluid mechanics, Vol. 31, Annu. Rev. Fluid Mech., vol. 31, Annual Reviews, Palo Alto, CA, 1999, pp. 567–603.
- [TOH09] S. Turek, A. Ouazzi, and J. Hron, *On pressure separation algorithms (PSepA) for improving the accuracy of incompressible flow simulations*, Internat. J. Numer. Methods Fluids **59** (2009), no. 4, 387–403.
- [TOS01] U. Trottenberg, C. W. Oosterlee, and A. Schüller, *Multigrid*, Academic Press Inc., San Diego, CA, 2001, With contributions by A. Brandt, P. Oswald and K. Stüben.
- [vTM09] L. Štrubelj, I. Tiselj, and B. Mavkob, *Simulations of free surface flows with implementation of surface tension and interface sharpening in the two-fluid model*, International Journal of Heat and Fluid Flow **30** (2009), no. 4, 741–750.
- [Wes92] Pieter Wesseling, *An introduction to multigrid methods*, Pure and Applied Mathematics (New York), John Wiley & Sons Ltd., Chichester, 1992.
- [Xu97] J. Xu, *An introduction to multilevel methods*, Wavelets, multilevel methods and elliptic PDEs (Leicester, 1996), Numer. Math. Sci. Comput., Oxford Univ. Press, New York, 1997, pp. 213–302.
- [Yse92] H. Yserentant, *Hierarchical bases*, ICIAM 91 (Washington, DC, 1991), SIAM, Philadelphia, PA, 1992, pp. 256–276.
- [Yse93] ———, *Old and new convergence proofs for multigrid methods*, Acta Numerica (1993), 285–326.
- [Zha88] S. Zhang, *Multi-level iteratives techniques*, Research report no. 88020, Penn State University, 1988.

Raffinement local adaptatif et méthodes multiniveaux pour la simulation d'écoulements multiphasiques.

Résumé : Ce manuscrit de thèse décrit certains aspects numériques et mathématiques liés à la simulation d'écoulements incompressibles triphasiques à l'aide d'un modèle à interfaces diffuses de type Cahn-Hilliard (les interfaces sont représentées par des zones d'épaisseur faible mais non nulle) couplé aux équations de Navier-Stokes. La discrétisation en espace est effectuée par approximation variationnelle de Galerkin et la méthode des éléments finis. La présence d'échelles très différentes dans le système (les épaisseurs d'interfaces étant très petites devant les tailles caractéristiques du domaine) suggère l'utilisation d'une méthode de raffinement local adaptatif. La procédure que nous avons mise en place, permet de prendre en compte implicitement les non conformités des maillages générés, pour produire *in fine* des espaces d'approximation éléments finis conformes. Le principe est de raffiner en premier lieu les fonctions de base plutôt que le maillage. Le raffinement d'une fonction de base est rendu possible par l'existence conceptuelle d'une suite emboîtée de grilles uniformément raffinées, desquelles sont déduites des relations "parents-enfants" reliant les fonctions de bases de deux niveaux successifs de raffinement. Nous montrons, en outre, comment exploiter cette méthode pour construire des préconditionneurs multigrilles. A partir d'un espace d'approximation éléments finis composite (contenant plusieurs niveaux de raffinement), il est en effet possible par "coarsening" de reconstruire une suite d'espaces emboîtés auxiliaires, permettant ainsi d'entrer dans le cadre abstrait multigrille. Concernant la discrétisation en temps, notre étude a commencé par celle du système de Cahn-Hilliard. Pour remédier aux problèmes de convergence de la méthode de Newton utilisée pour résoudre ce système (non linéaire), un schéma semi-implicite a été proposé. Il permet de garantir la décroissance de l'énergie libre discrète assurant ainsi la stabilité du schéma. Nous montrons l'existence et la convergence des solutions discrètes vers une solution faible du système. Nous poursuivons ensuite cette étude en donnant une discrétisation en temps inconditionnellement stable du modèle complet Cahn-Hilliard/Navier-Stokes. Un point important est que cette discrétisation ne couple pas fortement les systèmes de Cahn-Hilliard et Navier-Stokes, autorisant une résolution découplée des deux systèmes dans chaque pas de temps. Nous montrons l'existence des solutions discrètes et, dans le cas où les trois fluides ont la même densité, nous montrons leur convergence vers des solutions faibles. Nous étudions, pour terminer cette partie, diverses problématiques liées à l'utilisation de la méthode de projection incrémentale. Enfin, la dernière partie présente plusieurs exemples de simulations numériques, diphasiques et triphasiques, en deux et trois dimensions.

Mots-clés : Éléments finis. Raffinement local adaptatif. Préconditionneurs multigrilles. Modèle de Cahn-Hilliard/Navier-Stokes. Schémas numériques. Estimations d'énergie.

Abstract : This manuscript describes some numerical and mathematical aspects of incompressible multiphase flows simulations with a diffuse interface Cahn-Hilliard/Navier-Stokes model (interfaces have a small but a positive thickness). The space discretisation is performed thanks to a Galerkin formulation and the finite elements method. The presence of different scales in the system (interfaces have a very small thickness compared to the characteristic lengths of the domain) suggests the use of a local adaptive refinement method. The algorithm, that we introduced, allows to implicitly handle the non conformities of the generated meshes to produce conformal finite elements approximation spaces. It consists in refining basis functions instead of cells. The refinement of a basis function is made possible by the conceptual existence of a nested sequence of uniformly refined grids from which "parent-child" relationships are deduced, linking the basis functions of two consecutive refinement levels. Moreover, we show how this method can be exploited to build multigrid preconditioners. From a composite finite elements approximation space, it is indeed possible to rebuild, by "coarsening", a sequence of auxiliary nested spaces which allows to enter in the abstract multigrid framework. Concerning the time discretization, we begin by the study of the Cahn-Hilliard system. A semi-implicit scheme is proposed to remedy to convergence failures of the Newton method used to solve this (non linear) system. It guarantees the decrease of the discrete free energy ensuring the stability of the scheme. We show existence and convergence of discrete solutions towards the weak solution of the system. We then continue this study by providing an unconditionally stable time discretization of the complete Cahn-Hilliard/Navier-Stokes model. An important point is that this discretization does not strongly couple the Cahn-Hilliard and Navier-Stokes systems allowing to independently solve the two systems in each time step. We show the existence of discrete solutions and, in the case where the three fluids have the same densities, we show their convergence towards weak solutions. We study, to finish this part, different issues linked to the use of the incremental projection method. Finally, the last part presents several examples of numerical simulations, diphasic and triphasic, in two and three dimensions.

Keywords : Finite elements. Adaptive local refinement. Multigrid preconditioners. Cahn-Hilliard/Navier-Stokes model. Numerical schemes. Energy estimates.

AMS Classification : 35K55, 65M60, 65M12, 76T30, 65M50, 65M55