

THÈSES D'ORSAY

CATERINA CALGARO

**Méthodes multi-résolution auto-adaptatives en éléments finis :
application aux équations de la mécanique des fluides**

Thèses d'Orsay, 1996

[<http://www.numdam.org/item?id=BJHTUP11_1996__0445__P0_0>](http://www.numdam.org/item?id=BJHTUP11_1996__0445__P0_0)

L'accès aux archives de la série « Thèses d'Orsay » implique l'accord avec les conditions générales d'utilisation (<http://www.numdam.org/conditions>). Toute utilisation commerciale ou impression systématique est constitutive d'une infraction pénale. Toute copie ou impression de ce fichier doit contenir la présente mention de copyright.



NUMDAM

*Thèse numérisée par la bibliothèque mathématique Jacques Hadamard - 2016
et diffusée dans le cadre du programme
Numérisation de documents anciens mathématiques
<http://www.numdam.org/>*

ORSAY
n^o d'ordre :

UNIVERSITÉ DE PARIS-SUD
CENTRE D'ORSAY

THÈSE

présentée
pour obtenir

Le GRADE de DOCTEUR EN SCIENCES
DE L'UNIVERSITÉ PARIS XI ORSAY

Spécialité : Mathématiques

PAR

Caterina CALGARO

Sujet : Méthodes multi-résolution auto-adaptatives en éléments finis :
application aux équations de la mécanique des fluides

Soutenue le : 18 Décembre 1996 devant la Commission d'examen

M.	BREZZI F.	Rapporteur
M.	CHABARD J.P.	
M.	LABROSSE G.	
M.	LAMINIE J.	
Mme	MORANDI CECCHI M.	
M.	PIRONNEAU O.	Rapporteur
M.	TEMAM R.	Président

A Massimo...

Al bimbo che porto in grembo!

J'adresse ma profonde reconnaissance à Monsieur Roger Temam, mon directeur de thèse. Il m'a accueillie au sein du Laboratoire d'Analyse Numérique d'Orsay, il a su encadrer mes premiers pas dans la recherche et il a continué à avoir une attention vigilante sur mon travail tout au long de ces années.

Je remercie Messieurs Franco Brezzi et Olivier Pironneau non seulement pour avoir accepté de rapporter sur ce travail, mais également pour s'être adaptés au temps restreint que je leur ai imposé et pour les suggestions précieuses qu'ils m'ont proposées.

Je tiens à exprimer ma gratitude à Jacques Laminie qui a pris le relais de la direction de ma recherche et qui fut constamment disponible pour me conseiller, me guider et me relancer dans mon travail. C'est également vrai pour Arnaud Debussche auquel j'adresse ma pensée et mes remerciements. Je compte, bien sûr, sur les fruits de cette enrichissante collaboration à trois.

Je tiens particulièrement à remercier Madame Maria Morandi Cecchi qui a cru en moi avant même le moindre épanouissement de mes capacités et qui m'honore de sa présence dans le jury.

Merci également à Messieurs Jean Paul Chabard et Gérard Labrosse pour avoir accepté de participer à mon jury.

Je n'oublie pas Frédéric Pascal pour son savoir-faire, sa constante disponibilité et son aide à maintes reprises. J'espère qu'une fructueuse collaboration pourra continuer dans l'avenir.

Qu'une pluie de remerciements tombe sur Claude Jouron pour ses conseils avisés, sur Danièle Lemeur pour sa grande gentillesse et sa disponibilité, sur Hervé Le Meur pour sa constante bonne humeur et sa relecture précieuse de mon français et sur tous les autres membres du Laboratoire d'Analyse Numérique et E.D.P. d'Orsay.

Je n'oublie pas de saluer et de remercier mes compagnons thésards pour les discussions enrichissantes, pour leur soutien essentiel et surtout pour l'ambiance sympathique qu'ils ont créée. Je ne cite que Catherine, Olivier, Pascal, Ezzedine, Farid et Guy, mais ma pensée s'oriente, d'une manière générale, vers tous les locataires actuels des salles 256 et 258 et aussi vers tous ceux qui ont déjà quitté le laboratoire.

Pour terminer, je m'adresse à Massimo, mon mari, qui a su m'encourager, me soutenir et me motiver tout au long de ces dernières années, à ma famille qui m'a laissé la liberté de choisir mon chemin et a accepté mon éloignement et à tous mes amis rencontrés ici en France et à ceux qui ne m'oublient pas en Italie.

Enfin une tendre pensée est adressée à mon enfant qui m'a donné la détermination et l'énergie pour terminer cette thèse et qui va bientôt me combler de joie par sa naissance.

Caterina Calgaro :

Dynamical multilevel schemes in the Finite Element framework: applications to the Fluid Mechanics equations

Abstract :

This work deals with the splitting of a system of partial differential equations in small and large scale components (Incremental like method), in the framework of hierarchical finite element basis.

In the first part, we introduce a new approach by neglecting time variations of interaction terms between different structures. The numerical resolution is performed on the two-dimensional viscous Burgers' problem. In order to establish a space-time multilevel algorithm, the hierarchical components of the velocity field are treated differently. We derive several *a priori* estimates in order to study the perturbation introduced into the approximated equations and in order to deduce the V-cycles structure of the algorithm. With our adaptive multilevel algorithm, the gain on CPU time is significant.

In the last part, we generalize this method to the resolution of the incompressible Navier-Stokes equations, written in the primitive variables. In the first chapter we develop a Navier-Stokes solver with a penalty method approach. After various attempts, we choose the (P1-iso-P2) variant of the Hood-Taylor finite element and a direct method, at each time step, is applied to the linear system.

The last chapter is devoted to an extension of ideas presented in the first part. Noticing that the equations are not easily hierarchizable on more than two levels, we introduce an auxiliary pressure argument to extend the process to any level of grid. We perform a complete numerical study of the interaction terms depending upon the large scale components of the flow and their increments. Unfortunately we can not apply, on the discretized equations, the adaptive multilevel algorithm developed in the first part. Finally, we propose a new approach based on both the hierarchization and the domain decomposition method. This approach mainly allows to localize the treatment of interaction between large and small scales. From these results, it seems possible to develop a multilevel algorithm which is local in each subdomain.

Key words :

Finite elements ; hierarchical basis ; 2D Burgers equations ; auto-adaptive and multi-scale solvers ; long time integration ; incompressible Navier-Stokes equations ; penalty method ; domain decomposition method.

Introduction

L'intégration sur des longs intervalles de temps des équations évolutives issues de la mécanique des fluides est actuellement un problème encore difficile à résoudre. Pour simuler correctement des écoulements qui ne convergent pas simplement vers une solution stationnaire, mais qui présentent au contraire des dynamiques plus complexes (à partir de solutions périodiques jusqu'à la turbulence), il est nécessaire de considérer un nombre de degrés de liberté élevé. Ce nombre très élevé, qui est fonction par exemple du paramètre adimensionnel de Reynolds, amène rapidement à la limite des moyens matériels informatiques actuels.

Pour faire face à ce problème, plusieurs méthodes (voir par exemple la *Large Eddies Simulation*) envisagent un découpage des inconnues. Suivant cette philosophie, des nouveaux algorithmes de *type incrémental* ont été récemment introduits à partir de la théorie des systèmes dynamiques de dimension infinie (cf. [7]-[11], [16], [17] et [22] de la partie I). Tous ces algorithmes sont basés sur une décomposition de la solution approchée en petites et grandes structures. Il s'agit d'une décomposition *physique* de la solution dans le cadre d'une discrétisation spectrale (décomposition de la solution en grandes et petites échelles induite à partir du développement en séries de Fourier de la solution). En revanche, cette décomposition est *mathématique* dans le cadre d'une discrétisation par éléments finis ou en différences finies (décomposition dans les grandes structures de la solution et dans leurs corrections, obtenue respectivement grâce à l'utilisation des bases hiérarchiques et des inconnues incrémentales).

Afin de réaliser une approche de *type incrémental*, il est essentiel de traiter différemment les grandes structures de la solution et leurs corrections. Ceci a permis de mettre au point, d'abord en méthode spectrale, des stratégies effectives de résolution multi-échelles dynamique, afin d'approcher la solution de problèmes de turbulence homogène isotrope en dimension deux et trois avec des conditions aux limites périodiques (cf. [3]-[6] et [13] de la partie I). Dans les travaux en différences finies, les inconnues incrémentales jouent surtout le rôle de préconditionneurs pour les problèmes elliptiques linéaires en dimension deux et trois (cf. [8], [17] et [29] de la partie II). Il s'agit de préconditionneurs assez efficaces, qui sont bien adaptés aux calculateurs vectoriels. Les premiers travaux en éléments finis (cf. [14], [15], [17] et [18] de la partie I) sont réalisés pour une décomposition qui considère seulement deux niveaux hiérarchiques. Les approximations proposées sont justifiées lorsque le pas de discrétisation spatiale converge vers zéro. Cependant, du point de vue numérique, le pas en espace reste encore malheureusement trop grand, particulièrement pour des écoulements à nombre de Reynolds élevé où la solution présente

des zones à fort gradient. Dans ce cas, les expériences numériques ont montré que les termes dépendant des corrections ne peuvent plus être négligés.

Mon travail de thèse veut être la suite des résultats obtenus dans le cadre d'une discrétisation par éléments finis pour des problèmes de type Burgers généralisé (présenté dans la première partie de la thèse) ou de type Navier-Stokes évolutif (voir la deuxième partie).

Dans la première partie nous introduisons une approche nouvelle qui consiste à négliger les variations en temps des termes d'interactions entre les grandes structures et leurs corrections. L'application numérique est faite dans le cadre de la discrétisation du problème de Burgers visqueux bidimensionnel, pour une hiérarchisation de la vitesse effectuée sur plusieurs niveaux de grilles.

En détaillant un peu plus, si h dénote le paramètre de discrétisation spatiale, la méthode de Galerkin classique consiste à projeter les équations de Burgers sur un espace d'éléments finis V_h . La solution approchée $u_h \in V_h$ vérifie alors l'équation sous la forme variationnelle suivante :

$$\left(\frac{\partial u_h}{\partial t}, v_h\right) + \nu((u_h, v_h)) + (u_h \cdot \nabla u_h, v_h) = (f, v_h) ,$$

pour tout $v_h \in V_h$, $((.,.))$ étant une forme bilinéaire coercive. Les algorithmes de *type incrémental* se basent sur une décomposition hiérarchique de l'espace V_h :

$$V_h = V_{2h} + W_h ,$$

ce qui permet d'écrire la solution approchée $u_h \in V_h$ sous la forme

$$u_h = y_{2h} + z_h .$$

Cette décomposition peut être réitérée de façon récursive. Les y_{2h} représentent alors les grandes structures de la solution tandis que les z_h représentent leurs corrections ou incréments.

En projetant l'équation écrite ci-dessus sur la base hiérarchique, nous choisissons de ne pas résoudre l'équation en z_h et de considérer seulement celle projetée sur la grille grossière, soit, pour tout $v_{2h} \in V_{2h}$:

$$\left(\frac{\partial}{\partial t}(y_{2h} + z_h), v_{2h}\right) + \nu((y_{2h} + z_h, v_{2h})) + (y_{2h} \cdot \nabla y_{2h}, v_{2h}) + b_{int}(y_{2h}, z_h, v_{2h}) = (f, v_{2h}) .$$

Avec ces notations, $b_{int}(y_{2h}, z_h, v_{2h}) = (u_h \cdot \nabla u_h, v_{2h}) - (y_{2h} \cdot \nabla y_{2h}, v_{2h})$ représente le terme non linéaire d'interaction entre les grandes structures et leurs corrections ainsi que les interactions au sein des corrections.

Nous calculons d'abord l'ordre de grandeur de tous les termes appartenants aux équations projetées, ainsi que la variation sur un pas de temps de tous les termes d'interactions (il s'agit des termes linéaires et non linéaires en z_h présents dans l'équation projetée sur la grille grossière). Nous retrouvons les conclusions auxquelles parviennent les travaux

de Laminie, Pascal et Temam (cf. [14], [15] et [18] de la partie I); en effet les termes d'interactions, même si petits, ne sont pas négligeables. En revanche, les corrections z_h et les termes d'interactions linéaires et non linéaires peuvent être figés sur des courts intervalles de temps. Il s'agit, en effet, de négliger la variation de ces termes sur un pas de temps, variation qui est petite pour une précision de calcul donnée.

A partir de l'étude faite *a posteriori* sur la décomposition de la solution nodale u_h sur plusieurs niveaux hiérarchiques, nous élaborons et implémentons une nouvelle méthode multi-résolution auto-adaptative qui réalise des V-cycles en espace et en temps. Les diverses composantes hiérarchiques du champ de vitesse sont alors traitées de manière différente. La structure en V-cycles permet de résoudre les équations de Burgers en les restreignant, durant une ou plusieurs itérations en temps, sur une grille plus grossière que la grille nodale et en approchant numériquement une seule équation, celle relative aux grandes structures où tous les termes d'interactions sont figés. La mise au point de la stratégie multi-échelles passe par l'étude de plusieurs critères qui permettent de déterminer la profondeur des V-cycles et l'intervalle de temps maximal pendant lequel le terme d'interaction non linéaire est figé (c'est-à-dire le nombre maximal de V-cycles effectués dans une même série). Des estimations *a priori* permettent de majorer l'erreur introduite dans les équations approchées, afin d'évaluer les paramètres déterminant la structure de la série de V-cycles. L'auto-adaptativité de notre méthode réside dans l'évaluation numérique des conditions issues des estimations *a priori*, évaluation qui sera effectuée avant le début de la série de V-cycles, et dans un contrôle *a posteriori* des quantités calculées à la fin de chaque V-cycle. Ces critères dépendent de la variation des incréments z_h ainsi que des termes linéaires et du terme non linéaire d'interaction, ce dernier étant le plus coûteux en temps calcul. La notion de corrélation est alors introduite afin d'obtenir des tests précis mais qui restent en même temps peu coûteux.

Les incréments z_h sont réévalués à chaque V-cycle lors de la remontée sur la grille fine ainsi qu'à la fin de la série. En effet, à la fin de la série de V-cycles peu d'itérations en temps sur la grille nodale permettent de lisser la solution. Cette dernière procédure est comparable à la projection de la solution calculée sur une Variété Inertielle Approchée adéquate. Nous montrons que pour le problème de Burgers le gain en temps est significatif et que la solution déterminée reste assez proche de celle obtenue par la méthode de Galerkin classique.

Je précise ici que ces algorithmes numériques ont été adaptés aux super-calculateurs vectoriels (voir CRAY C98). En effet, nous sommes parvenus à un fort taux de vectorisation grâce à l'application de techniques de coloriage qui consistent à renuméroter les nœuds du maillage afin d'éliminer les conflits de dépendance dus aux adressages indirects.

Dans la deuxième partie de la thèse je me propose d'étudier la généralisation de cette méthode dans le cadre de l'approximation numérique des équations de Navier-Stokes évolutives, incompressibles, en formulation vitesse-pression. Cette partie est subdivisée en deux chapitres.

Le premier chapitre est consacré à la construction d'un solveur permettant de résoudre

les équations de Navier-Stokes bidimensionnelles. Afin d'étendre le schéma multi-niveaux auto-adaptatif à ces équations et en considérant les difficultés rencontrées dans les premières applications numériques (cf. [23] et [27] pour les éléments finis et [29] pour les différences finies), nous nous sommes orientés vers une simplification du problème de départ. Puisque ces difficultés sont dues principalement au traitement hiérarchique lié à la contrainte d'incompressibilité, nous étudions les équations de Stokes et de Navier-Stokes sous leur forme pénalisée, afin d'éliminer la pression, une inconnue du problème. Nous obtenons ainsi une seule équation en vitesse qui ressemble à celle étudiée dans la première partie de la thèse.

La discrétisation en espace sera faite à l'aide des éléments finis mixtes. Notre choix de l'élément se base sur un certain nombre d'exigences. Afin d'approcher correctement nos équations, il faut dans un premier temps discrétiser le problème en vitesse et en pression et ensuite supprimer la pression dans la formulation matricielle associée au problème étudié. Pour cela il faut considérer des éléments finis mixtes qui vérifient la condition Inf-Sup. De plus, ces éléments doivent être facilement hiérarchisables, car notre stratégie multi-niveaux porte sur la résolution des mêmes équations sur différents niveaux hiérarchiques, équations qui possèdent un second membre modifié par les termes que l'on fixe. Enfin, nous voulons des éléments qui restent relativement précis pour les calculs et qui donnent en même temps des matrices relativement peu encombrantes en espace mémoire. Pour toutes ces raisons nous nous sommes orientés dans un premier temps vers un nouvel élément, le 4P1-P0, ce choix étant le plus naturel et simple, car la pression peut ainsi être éliminée de façon locale. Cependant, le choix d'une discrétisation P0 pour la pression présente des gros problèmes dans la résolution des équations de Stokes généralisées, dus à l'apparition des nombreuses oscillations dans le champ de vitesse approchée. Cet élément mixte se révèle donc peu adapté au cas évolutif, vu que la résolution du problème de Stokes généralisé représente une itération en temps dans l'approximation des équations de Navier-Stokes évolutives. L'élément P2-P0 sera également exclu car il présente des problèmes analogues. Le choix initial sera donc abandonné au profit d'une formulation 4P1-P1 qui n'engendre plus ces oscillations, tout en restant facilement hiérarchisable, bien que plus difficile à mettre en œuvre. En effet, l'élimination de la pression ne se fait plus de façon locale, ce qui nous amène à des matrices moins creuses qu'auparavant.

Des nombreux problèmes sont aussi liés au mauvais conditionnement de la matrice issue de la pénalisation de la pression. Une résolution performante du système linéaire auquel nous parvenons n'apparaît pas très évidente. Nous montrerons que les préconditionnements classiques des méthodes itératives du type gradient conjugué ne donnent pas de bons résultats. Après avoir effectué une brève étude des valeurs propres de la matrice de notre système linéaire, nous optons pour une résolution directe de ce système. Seulement la méthode directe constitue un solveur performant en temps de calcul, même si la taille de la discrétisation spatiale du problème est limitée par l'encombrement mémoire dû au stockage de la matrice factorisée avec coefficients réordonnés. Ce solveur permet d'obtenir des simulations numériques des équations de Navier-Stokes qui sont satisfaisantes, comparées aux résultats obtenus par d'autres auteurs. Les tests portent sur l'approximation de l'écoulement dans la cavité entraînée régularisée ou non régularisée

pour des nombres de Reynolds jusqu'à 5000. Pour ces valeurs du paramètre la solution converge vers une solution stationnaire. Néanmoins, avant que le régime stationnaire soit définitivement établi, l'évolution apparaissant dans l'écoulement nous permet de réaliser une première série de tests, qui devront ensuite être confirmés par des valeurs de Reynolds supérieures. Le solveur mis au point dans ce chapitre est donc destiné à être l'outil nécessaire aux développements présentés dans le chapitre suivant.

Ce chapitre sur Navier-Stokes est une extension des idées et des techniques utilisées pour le problème de Burgers. La hiérarchisation des équations, réalisée pour deux niveaux, ne pose aucun problème. En revanche, il s'avère difficile de réitérer la hiérarchisation sur plusieurs niveaux. En effet, une décomposition de la pression analogue à celle de la vitesse nous amène à des équations trop compliquées. Une nouvelle décomposition de la pression sera alors proposée : il s'agit d'introduire une pression auxiliaire, que l'on nommera *mathématique*. Cette pression auxiliaire permet de hiérarchiser correctement le problème de Navier-Stokes sur plusieurs niveaux. Il est nécessaire d'introduire aussi une norme convenable, indépendante du pas de discrétisation, afin d'évaluer tous les termes apparaissant dans les équations projetées sur la grille grossière.

L'analyse faite pour les différents termes est analogue à celle effectuée pour Burgers. Des arguments numériques permettent de conclure que les variations des termes de couplage (à figer) sont inférieures à celles des termes dépendant seulement des grandes structures. Cependant, leurs variations n'apparaissent pas suffisamment petites pour appliquer notre stratégie de V-cycles ; ces variations ne sont négligeables que sur des très courts intervalles de temps. Nous constatons aussi que la pression *mathématique* ne nous convient pas, car ses valeurs sont trop élevées et surtout ses variations sur un pas de temps sont trop importantes par rapport aux autres termes en vitesse projetés sur la grille de même niveau. Enfin, nous remarquons aussi des fortes variations dans les termes de pénalisation, dues à l'inverse du coefficient de pénalisation présent dans ces termes. De plus, les termes de pénalisation étant strictement liés aux termes de divergence, la variation du terme dépendant des corrections z_h devient du même ordre que celle relative au terme en y_{2h} . Toutefois, il conviendra de relativiser ce problème, car ces termes sont figés seulement pendant une ou, au plus, quelques itérations en temps.

Une approche totalement différente à celle utilisée pour le problème de Burgers est ensuite proposée. L'idée porte sur un couplage de la hiérarchisation avec la méthode de décomposition de domaine. Cette approche nous permet de localiser le traitement des interactions entre les grandes échelles et leurs corrections. Le plus simple est de considérer un découpage de Ω en quatre sous-domaines sans recouvrement. Les grandes échelles y_{2h} , ainsi que les corrections z_h , sont décomposées en quatre contributions, une dans chaque sous-domaine. Tous les termes d'interactions sont alors traités de façon locale et leurs variations sont analysées dans chacun des sous-domaines. L'analyse, faite zone par zone, est analogue à celle réalisée sur Ω tout entier et permet de déceler des comportements physiques locaux qui étaient noyés dans une séparation d'échelle globale. Nous observons que les variations en temps des termes de couplage sont inférieures à celle des termes dépendant des grandes structures, mais ceci ne se réalise pas en même temps dans tous

les sous-domaines. Il est naturel, à partir des résultats numériques trouvés, de concevoir une stratégie multi-échelle dans la zone (resp. les zones) à l'intérieur de laquelle (resp. desquelles) les variations des termes en z_h sont petites et, par contre, de continuer avec une résolution fine dans les sous-domaines où les variations en temps sont trop fortes. Enfin, nous remarquons que dans chaque zone il existe une très bonne corrélation entre les différents termes. Ce comportement nous permettra de deviner l'évolution du terme non linéaire d'interaction à partir de termes plus simples et moins coûteux en temps calcul.

La généralisation pour les équations de Navier-Stokes de la méthode mise en place pour résoudre le problème de Burgers n'est pas encore terminée. Toutefois, les premiers résultats laissent entrevoir le développement rapide d'une résolution des équations de Navier-Stokes par une méthode de type incrémentale couplée avec une méthode de décomposition de domaine. La première étape de ce couplage, qui porte sur la parallélisation du code séquentiel, est déjà en cours d'élaboration. L'attrait d'un code parallèle réside essentiellement dans l'accès à des maillages plus fins. La détermination de solutions hautement instationnaires (à grand nombre de Reynolds) devient ainsi envisageable.

Table des matières

Introduction	i
I Résolution du problème de BURGERS visqueux 2D	1
1 Dynamical multilevel schemes for the solution of evolution equations by hierarchical finite element discretization	3
Introduction	5
1.1 Test problem and weak formulation	7
1.2 The hierarchical finite element basis	8
1.3 Discretization on the nodal basis	10
1.4 Analysis of the decomposition based on the hierarchical basis	11
1.4.1 Small and large scale structures	12
1.4.2 The time derivative terms	13
1.4.3 The diffusive viscous terms	14
1.4.4 The nonlinear terms	15
1.5 The space-time multilevel algorithm	21
1.5.1 The multilevel procedure	23
1.5.2 Time scale estimate for the small structures and for the linear and nonlinear interaction terms	27
1.5.2.1 Estimate of p_V , the depth of a V-cycle	27
1.5.2.2 Estimate of $\tau(t_0)$, the time interval for the transfer terms	29
1.6 Numerical results	34
Concluding Remarks	45
Bibliography	47
II Résolution des équations de NAVIER-STOKES	51
1 Résolution des équations de Stokes et de Navier-Stokes	53
Introduction	53
1.1 Les équations considérées	54
1.1.1 Les équations en variables primitives	54

1.1.2	Cadre fonctionnel	56
1.1.3	Le problème de Stokes : formulations équivalentes	57
1.1.4	Les problèmes de Stokes et Navier-Stokes pénalisés	59
1.2	Discrétisation du problème de Stokes pénalisé	62
1.2.1	Espaces d'approximation	62
1.2.2	Formulation discrète du problème	63
1.2.3	Forme matricielle du problème	65
1.3	Résolution du système linéaire issu du problème de Stokes	67
1.3.1	La méthode itérative : une étude de différents préconditionnements	69
1.3.2	Les valeurs propres de la matrice à inverser	73
1.3.3	La méthode directe : mise en œuvre de la factorisation complète	78
1.4	Les résultats numériques pour Stokes	84
1.4.1	Influence de la valeur de ϵ	84
1.4.2	Perturbations dans le champ de vitesse	86
1.4.3	Etude numérique des perturbations	91
1.5	Résolution du problème de Navier-Stokes	95
1.5.1	Discrétisation en espace et en temps	95
1.5.2	Résultats numériques	96
1.6	Conclusions	108
2	Le problème de Navier-Stokes en base hiérarchique	111
	Introduction	111
2.1	Décomposition sur la base hiérarchique	112
2.1.1	Stokes en base hiérarchique	113
2.1.2	Navier-Stokes en base hiérarchique	119
2.2	Analyse de la décomposition hiérarchique	120
2.2.1	Taille des différentes structures	121
2.2.2	Evaluation pratique des différents termes	123
2.2.3	Etude globale des différents termes	125
2.2.4	Conclusions relatives à l'étude globale	132
2.2.5	Etude par zones des différents termes	133
	Bibliographie	161
	Conclusions et développements	165

Partie I

Résolution du problème de BURGERS visqueux 2D

Chapitre 1

Dynamical multilevel schemes for
the solution of evolution equations
by hierarchical finite element
discretization

Dynamical Multilevel Schemes for the Solution of Evolution Equations by Hierarchical Finite Element Discretization ¹

C. CALGARO, J. LAMINIE and R. TEMAM ²

Laboratoire d'Analyse Numérique et Equations aux dérivées partielles
Université Paris-Sud et CNRS URA 760 Bâtiment 425, 91405 Orsay Cedex, France.

Abstract

The full numerical simulation of turbulent flows in the context of industrial applications remains a very challenging problem. One of the difficulties is the presence of a large number of interacting scales of different orders of magnitude ranging from the macroscopic scale to the Kolmogorov dissipation scale. In order to better understand and simulate these interactions, new algorithms of *incremental type* have been recently introduced, stemming from the theory of infinite dimensional dynamical systems (see e.g. [7], [11], [14]-[17], [22]). These algorithms are based on decompositions of the unknown functions into a large and a small scale components, one of the underlying ideas being to approximate the corresponding attractor. In the context of finite elements methods, the decomposition of solution into small and large scale components appears when we consider hierarchical bases [23], [24]. In the present article we derive several estimates of the size of the structures for the linear and nonlinear terms which correspond to interactions of different hierarchical components of the velocity field, and also their time variation. The one-step time variation of these terms can be smaller than the expected accuracy of the computation. Using this remark, we implement an adaptive spatial and temporal multilevel algorithm which treats differently the small and large scale components of the flow. We derive several *a priori* estimates in order to study the perturbation introduced into the approximated equations. All the interaction terms between small and large structures are frozen during several time steps. Finally we implement the multilevel method in order to simulate a bidimensional flow described by the Burgers equations. We perform a parametric study of the procedure and its efficiency. The gain on CPU time is significant and the trajectories computed by our multi-scale method remain close to the trajectories obtained with the classical Galerkin method.

AMS Codes : 35K55, 65G99, 65M55, 65M60, 76E30, 76M10

Key words : Hierarchical Finite Elements, 2D Burgers equations,
auto-adaptive and multi-scale solvers, long time integration.

¹To appear in *Applied Numerical Mathematics*.

²The numerical simulations of the work was supported by the Institut du Développement et des Ressources en Informatique Scientifique (IDRIS), Bât. 506, B.P. 167, 91403 Orsay Cedex, France.

Introduction

The integration of evolution equations on large intervals of time remains a difficult problem, despite the significant increase in the computing power. It is well known that the large time behaviour of solutions of dissipative evolution equations depends on certain non-dimensional parameters such as the Reynolds or Grashof numbers or other similar numbers. When the values of those physical parameters are small the system evolves, as $t \rightarrow \infty$, towards a stationary solution. On the contrary, for large values of such parameters the system exhibits nontrivial dynamics related to turbulence.

For dissipative evolution equations, such as the Navier-Stokes equations, it was shown that the long-time behaviour is encompassed by a compact attractor and that such an attractor has a finite fractal dimension (see C. Foias and R. Temam [8]). Hence, although the attractor can be complex, or even fractal, the long-time dynamics of the system are finite dimensional. In practice, Constantin *et al.* in [2] estimate the number of degrees of freedom, needed to insure that a numerical solution of the Navier-Stokes equations is at least qualitatively correct, to be of order of the fractal dimension of the global attractor.

To solve dissipative evolution equations in dynamically nontrivial situations, C. Foias, G. Sell and R. Temam have introduced in [9] the concept of Inertial Manifolds, which are finite dimensional manifolds containing the global attractor. The Approximate Inertial Manifolds, introduced by C. Foias, O. Manley and R. Temam in [11], are more flexible and a less restrictive idea. They are explicit smooth manifolds approximating the attractor up to a certain level of accuracy and they allow to build numerical schemes providing orbits lying on the approximate set.

From this M. Marion and R. Temam have subsequently developed new algorithms of incremental type, called nonlinear Galerkin methods (N.L.G. methods), in the context of spectral and finite element discretizations (cf. [16] and [17]) and incremental unknowns methods in the context of finite differences [22]. These innovative methods are based on a decomposition of the unknowns, such as the velocity field, into a small scale component and a large scale one. Essential in the N.L.G. methods is a differentiated treatment of small and large scales and in that respect the nonlinear Galerkin methods and their variations appear as dynamical multi-resolutions methods.

A multilevel scheme, which develops a self-adaptive procedure, is implemented for spectral discretizations by T. Dubois, F. Jauberteau and R. Temam (cf. [5] and the reference therein ; see also [6]) ; moreover A. Debussche, T. Dubois et R. Temam develop further the study of the numerical implementation of the N.L.G. method for turbulent flows (cf. [3]).

In the context of the finite elements discretizations, implementations of the N.L.G. method have been done by J. Laminie, F. Pascal and R. Temam [14], [15]. The classical Galerkin approximation of Burgers and Navier-Stokes problems have been studied using the hierarchical basis decomposition (with a coarse initial triangulation and a finer one obtained by subdividing any triangle into four similar sub-triangles). The approximations made for the N.L.G. method are justified when the mesh size h converges to zero. Although the quantities which represent the small scale components of the flow

are small, numerical experiments done by F. Pascal [18] show that they can not be neglected, especially in the case of higher Reynolds numbers.

So we introduce a new idea which is to consider the time variations of terms depending upon the small scale components of the flow, with the intention of freezing them when their time variation is small. In order to split the solution into its small and large scale components, we use the hierarchical finite element bases, defined by considering several levels of nested triangulations. The study is performed for the two-dimensional Burgers problem, considering the evolution of all linear and nonlinear terms. The first results show that the time variation, over one time step, of the coupled terms can be small. Therefore these terms can be frozen during a short period of time representing a certain number of time steps.

Based on this results, we propose a spatial and temporal multilevel self-adaptive method. The level of refinement, which defines the separation of the solution into its small and large scales, evolves in time as a multilevel V-cycle. Moreover the small components of the solution are frozen and, besides, we freeze the suitable terms in the equation during a suitable short time period. The structure of a series of V-cycles is dynamically determined using criteria which guarantee that all frozen quantities are smaller than a given parameter, representing the expected accuracy of the computation. Significant savings of computing time are obtained as compared to the classical finite element Galerkin implementation. The promises of the multi-scale algorithm are important : most often only one coarse grid nonlinear equation is solved instead of two nonlinear equations as in the previous implementations [14], [15].

The article is organized as follows. In Section 1.1 we introduce the nonlinear Burgers problem and its weak formulation. In Section 1.2 we recall briefly the hierarchical finite element basis and we split the solution into its components of different scales. The coarsest mesh gives the large structures of the flow, whereas on the finer grids one derives the corrections of the solution, which are smaller and smaller according to the level of refinement. In Section 1.3 we recall the classical Galerkin (F.E.M.) implementation : the approximate solution computed by the standard Galerkin method will be compared with that provided by our multilevel algorithm. Moreover this solution will be the object, in Section 1.4, of an *a posteriori* estimate performed on the different terms of the equations. These quantities and the time variation of the coupled terms are analyzed in order to emphasize the numerical justifications which are applied in the development of the algorithm proposed in the following Section 1.5. The level of refinement evolves in time undergoing successive multilevel V-cycles. Moreover, the corrections to the approximate solution are frozen (indeed only the nonlinear equation corresponding to the large structures is solved) and also we freeze the coupled interaction terms during a suitable short period of time. After a complete description of this multilevel procedure, we derive important estimates of the characteristic numbers which determine the structure of the V-cycle series. The last Section is devoted to the description and the analysis of numerical results performed with our multilevel method. This simulation is compared with the approximate solution determined by the classical Galerkin (F.E.M.) method.

1.1 Test problem and weak formulation

Let Ω be a bounded domain of $\mathbf{R}^2$ with boundary $\partial\Omega$. We consider, as a model of nonlinear evolution equations, the two-dimensional Burgers problem. The Burgers equations are a simplified formulation of the incompressible Navier-Stokes equations. They constitute a model of viscous flows which incorporates the diffusive viscous terms and the nonlinear convection processes but the pressure gradient terms and the incompressibility condition are removed. These equations encompass some qualitative aspects of the Navier-Stokes equations and they produce a suitable model for the development of algorithms for Navier-Stokes equations themselves.

In the two-dimensional Burgers problem the unknown function u maps $\Omega \times [0, T]$ into $\mathbf{R}^2$ and it satisfies the following equations :

$$\left\{ \begin{array}{ll} \frac{\partial u}{\partial t} - \frac{1}{Re} \Delta u + (u \cdot \nabla) u = f & \text{in } \Omega \times [0, T] \\ u = g & \text{on } \partial\Omega \\ u(x, 0) = u_0(x) & \forall x \in \Omega \end{array} \right. \quad \begin{array}{l} (a) \\ (b) \\ (c) \end{array} \quad (1.1)$$

We denote by $(.,.)$ the scalar product in $L^2(\Omega)$ and by $((.,.))$ that of $H_0^1(\Omega)$ defined by $((u, v)) = \int_{\Omega} \nabla u \cdot \nabla v$. We use the following function spaces :

$$\begin{aligned} \mathcal{H} &= L^2(\Omega) \times L^2(\Omega) \quad ; \quad \mathcal{V} = H^1(\Omega) \times H^1(\Omega) \quad ; \\ \mathcal{V}_0 &= H_0^1(\Omega) \times H_0^1(\Omega) \quad ; \quad \mathcal{V}_g = \{u \in \mathcal{V} \ ; \ u = g \text{ on } \partial\Omega\} . \end{aligned}$$

Let $B(.,.)$ be a bilinear operator from $\mathcal{V} \times \mathcal{V}$ into $\mathcal{V}'$. We denote by b the trilinear continuous form on $\mathcal{V}$ given by :

$$b(u, v, w) = \langle B(u, v), w \rangle_{\mathcal{V}', \mathcal{V}} , \quad \forall u, v, w \in \mathcal{V} .$$

For the Burgers equations

$$b(u, v, w) = \sum_{i,j=1}^2 \int_{\Omega} u_i \cdot \frac{\partial v_j}{\partial x_i} \cdot w_j \, dx \quad (1.2)$$

For the finite element treatment, we write (1.1) in weak form.

Let $f \in \mathcal{H}$ and $u_0 \in \mathcal{V}_g$; for all $t > 0$, the unknown function $u = u(t) \in L^2(0, T; \mathcal{V}_g) \cap \mathcal{C}(0, T; \mathcal{H})$ is solution of :

$$\left\{ \begin{array}{ll} \left(\frac{\partial u}{\partial t}, \tilde{u} \right) + \frac{1}{Re} ((u, \tilde{u})) + b(u, u, \tilde{u}) = (f, \tilde{u}) , & \forall \tilde{u} \in \mathcal{V}_0 , \\ (u(0), \tilde{u}) = (u_0, \tilde{u}) , & \forall \tilde{u} \in \mathcal{V}_0 . \end{array} \right. \quad (1.3)$$

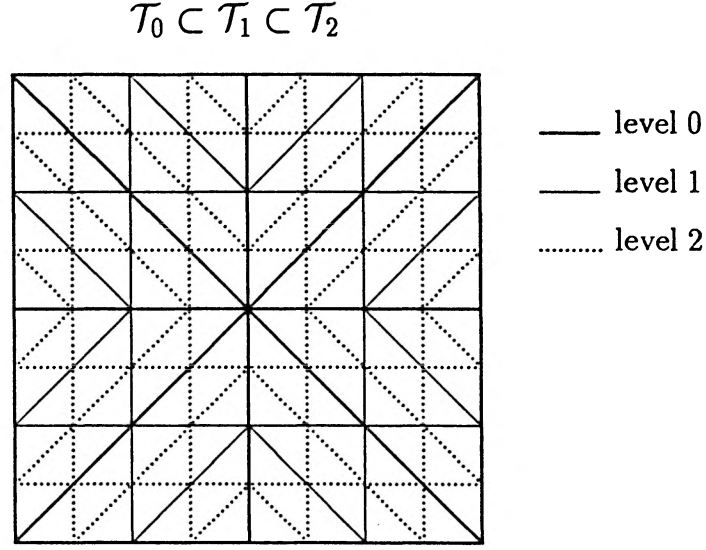


Figure 1.1: Hierarchical triangulation of domain Ω : $\mathcal{T}_0 \subset \mathcal{T}_1 \subset \mathcal{T}_2$.

1.2 The hierarchical finite element basis

In order to split the solution into its small and large scale components, we use the hierarchical finite element basis. This concept has been introduced by Zieniewicz *et al.* in [24] and has been developed by H. Yserentant [23] in the context of elliptic linear problems.

From an initial finite element mesh $\mathcal{T}_0$ of the computational domain Ω , we construct a family of nested meshes

$$\mathcal{T}_0 \subset \mathcal{T}_1 \subset \dots \mathcal{T}_d$$

by subdividing each triangle into four similar sub-triangles. Here, $\mathcal{T}_d$ represents the finest triangulation, called fine mesh or nodal mesh. On Figure 1.1 we represent three nested meshes ($d = 2$) ; for each level of triangulation, we have considered a uniform grid.

Let Σ_k be the set of vertices of triangulation $\mathcal{T}_k$ and let V_k be the finite element space which approximates $H^1(\Omega)$:

$$V_k = \{v_k \in C^0(\Omega) \mid \forall \tau \in \mathcal{T}_k, v_k|_{\tau} \in P_1(\tau)\},$$

where $P_1(\tau)$ is the set of polynomials of degree less than or equal to one on each triangle of $\mathcal{T}_k$. Obviously we have :

$$\Sigma_k \subseteq \Sigma_{k+1} \quad \text{and} \quad V_k \subseteq V_{k+1} \quad (0 \leq k \leq d-1).$$

The finite element space V_d associated with the fine grid $\mathcal{T}_d$ is split into :

$$V_d = V_{d-1} + W_d ,$$

and recursively :

$$V_d = V_0 + W_1 + \dots + W_d .$$

The usual nodal basis of V_k is denoted by $\phi_{k,i}$. It is the canonical basis consisting of functions which take the value 1 at the node i of $\mathcal{T}_k$ and 0 at all other nodal points. The functions in the hierarchical basis of W_k ($k = 1, \dots, d$) are denoted by $\psi_{k,j}$, where j is a nodal point of $\mathcal{T}_k \setminus \mathcal{T}_{k-1}$; $\psi_{k,j}$ is a function of the nodal basis of V_k which vanishes on the nodes of triangulation $\mathcal{T}_{k-1}$ (i.e. $\psi_{k,j} = \phi_{k,i}$ at the node j of $\mathcal{T}_k \setminus \mathcal{T}_{k-1}$). Then we denote :

$$\begin{aligned} V_k &= \text{Span} \{ \phi_{k,1}, \dots, \phi_{k,n_k} \}, & n_k &= \text{card}(\Sigma_k) = \dim(V_k), \\ W_k &= \text{Span} \{ \psi_{k,n_{k-1}+1}, \dots, \psi_{k,n_k} \}, & n_k - n_{k-1} &= \text{card}(\Sigma_k \setminus \Sigma_{k-1}) = \dim(W_k). \end{aligned}$$

If u_d is an element of V_d , then the corresponding decomposition reads :

$$u_d = y_{d-1} + z_d, \quad \text{where } y_{d-1} \in V_{d-1} \text{ and } z_d \in W_d.$$

Recursively we find :

$$u_d = y_0 + z_1 + \dots + z_d, \quad \text{where } y_0 \in V_0 \text{ and } z_k \in W_k \text{ } (k = 1, \dots, d).$$

For each level k , $1 \leq k \leq d$, we write :

$$y_k = y_{k-1} + z_k = \sum_{i=1}^{n_{k-1}} y_{k-1,i} \phi_{k-1,i} + \sum_{i=n_{k-1}+1}^{n_k} z_{k,i} \psi_{k,i}$$

The values of $y_{k-1,i}$ and $z_{k,i}$ are uniquely determined by :

$$\begin{aligned} \forall i \in [1, n_{k-1}], & \quad y_{k-1,i} = y_k(A_i), \\ \forall i \in [n_{k-1} + 1, n_k], & \quad z_{k,i} = y_k(B_i) - \frac{1}{2}(y_{k-1}(A_{i1}) + y_{k-1}(A_{i2})), \end{aligned} \tag{1.4}$$

where $A_i, A_{i1}, A_{i2} \in \mathcal{T}_{k-1}$ and $B_i \in \mathcal{T}_k \setminus \mathcal{T}_{k-1}$, B_i being the midpoint of the edge $[A_{i1}, A_{i2}]$ (with generalized notation, we set $y_d = u_d$).

The large structures defined by y_{k-1} are of the same order as u_d . From Taylor's formula, we easily infer the order of the corrections z_k over the grid $\mathcal{T}_k$, for each $k \leq d$. Let $A_{i1}, A_{i2} \in \mathcal{T}_{k-1}$ and $B_i \in \mathcal{T}_k \setminus \mathcal{T}_{k-1}$, B_i the midpoint of the edge $[A_{i1}, A_{i2}]$. We obtain :

$$\begin{aligned} u_d(A_{i2}) &= u_d(B_i) + \overrightarrow{B_i A_{i2}} \cdot u'_d(B_i) + \frac{\overrightarrow{B_i A_{i2}}^2}{2} \cdot \frac{u''_d(B_i)}{2} + \mathcal{O}(\|\overrightarrow{B_i A_{i2}}\|^3), \\ u_d(A_{i1}) &= u_d(B_i) - \overrightarrow{A_{i1} B_i} \cdot u'_d(B_i) + \frac{\overrightarrow{A_{i1} B_i}^2}{2} \cdot \frac{u''_d(B_i)}{2} + \mathcal{O}(\|\overrightarrow{A_{i1} B_i}\|^3). \end{aligned}$$

From (1.4) we find :

$$z_k(B_i) = u_d(B_i) - \frac{u_d(A_{i1}) + u_d(A_{i2})}{2} = - \frac{\overrightarrow{A_{i1} B_i}^2}{2} \cdot \frac{u''_d(B_i)}{2} + \mathcal{O}(\|\overrightarrow{A_{i1} B_i}\|^3)$$

Hence the incremental components z_k can be of order h_k^2 , where h_k is the mesh width of the triangulation $\mathcal{T}_k$ (cf. [22]) ; but in fact, if the solution has strong gradients, the size

of the corrections z_k can grow considerably. We will keep this difficulty in mind when we build our algorithm.

Our approach is based on differentiated treatments of the small scale and the large scale components. This approach utilizes the hierarchization of finite element spaces introduced above. We will consider the size of different structures and especially the time variation over one time step of the small scale components as well as the time variation of the linear and nonlinear terms which depend on the small scales. When these variations are small, we will be able to freeze the corresponding terms during a short period of time representing a few time steps. Then we will freeze the small corrections z_k computed over one to several finer grids, and we will approximate the solution only over the coarser grids. The solution computed by our procedure will be compared with that given by the classical Galerkin method applied to the nodal basis which we now present in the following section.

1.3 Discretization on the nodal basis

Let V_d be a finite element space which approximates $H^1(\Omega)$. We set :

$$\begin{aligned}\mathcal{V}_d &= V_d \times V_d , \\ \mathcal{V}_{0,d} &= (V_d \cap H_0^1(\Omega))^2 , \\ \mathcal{V}_{g,d} &= \{v_d \in \mathcal{V}_d , \quad v_d = g_d \text{ on } \partial\Omega\} ,\end{aligned}$$

with g_d an appropriated approximation of g , and we equip $\mathcal{V}_{g,d}$ with the usual nodal basis. The discret form of (1.3) consists in finding $u_d(t) \in \mathcal{V}_{g,d}$ such that

$$\begin{cases} \left(\frac{\partial u_d}{\partial t}, \tilde{u}_d \right) + \frac{1}{Re}((u_d, \tilde{u}_d)) + b(u_d, u_d, \tilde{u}_d) = (f, \tilde{u}_d) , & \forall \tilde{u}_d \in \mathcal{V}_{0,d} , \\ (u_d(0), \tilde{u}_d) = (u_{0,d}, \tilde{u}_d) , & \forall \tilde{u}_d \in \mathcal{V}_{0,d} . \end{cases} \quad (1.5)$$

The time derivative is approximated by a finite difference scheme. We use a second-order implicit Crank-Nicholson scheme for the linear terms and a first-order explicit Euler scheme for the nonlinear terms.

Once we know the solution u_d^n at time $n\Delta t$, then the unknown $u_d^{n+1} \in \mathcal{V}_{g,d}$ at time $(n+1)\Delta t$ is solution of the following linear system :

$$\left(\frac{u_d^{n+1} - u_d^n}{\Delta t}, \tilde{u}_d \right) + \frac{1}{Re} \left(\left(\frac{u_d^{n+1} + u_d^n}{2}, \tilde{u}_d \right) \right) = (f, \tilde{u}_d) - b(u_d^n, u_d^n, \tilde{u}_d) , \quad \forall \tilde{u}_d \in \mathcal{V}_{0,d} . \quad (1.6)$$

The scheme given by (1.6) is found to be conditionally stable : it requires the ratio $|u_d|_\infty \frac{\Delta t}{h_d}$ not to exceed a critical value (CFL condition). Moreover the linear system for u_d^{n+1} is well conditioned.

1.4 Analysis of the decomposition based on the hierarchical basis

Using the hierarchical basis decomposition, we split the solution into its small scale and its large scale components. In this section we give some numerical (empirical) estimates of the terms which appear in the equations projected on the different meshes. More precisely we will analyze the norm L^2 of all linear and nonlinear terms and also the time variation of the terms which depend on the small scale z . We aim at estimating the size of these terms and especially at checking if their variation over one time step can be small.

The hierarchical basis for the finite element spaces considered, gives us the decomposition :

$$\mathcal{V}_{g,d} = \mathcal{V}_{g,d-1} + \mathcal{W}_{g,d} ,$$

and recursively :

$$\mathcal{V}_{g,d-1} = \mathcal{V}_{g,d-2} + \mathcal{W}_{g,d-1} , \quad \dots , \quad \mathcal{V}_{g,1} = \mathcal{V}_{g,0} + \mathcal{W}_{g,1} .$$

From the variational formulation (1.5), we infer a spatial discretization corresponding to the decomposition of $\mathcal{V}_{g,d}$ into $\mathcal{V}_{g,d-1} + \mathcal{W}_{g,d}$. For all $t > 0$, we look for $y_{d-1}(t) \in \mathcal{V}_{g,d-1}$ and $z_d(t) \in \mathcal{W}_{g,d}$ solutions of the following system :

$$\begin{cases} \left(\frac{\partial y_{d-1}}{\partial t} + \frac{\partial z_d}{\partial t}, \tilde{y}_{d-1} \right) + \frac{1}{Re} ((y_{d-1} + z_d, \tilde{y}_{d-1})) + \\ b(y_{d-1}, y_{d-1}, \tilde{y}_{d-1}) + b_{int}(y_{d-1}, z_d, \tilde{y}_{d-1}) = (f, \tilde{y}_{d-1}) , \quad \forall \tilde{y}_{d-1} \in \mathcal{V}_{0,d-1} , \end{cases} \quad (1.7)$$

$$\begin{cases} \left(\frac{\partial y_{d-1}}{\partial t} + \frac{\partial z_d}{\partial t}, \tilde{z}_d \right) + \frac{1}{Re} ((y_{d-1} + z_d, \tilde{z}_d)) + \\ b(y_{d-1}, y_{d-1}, \tilde{z}_d) + b_{int}(y_{d-1}, z_d, \tilde{z}_d) = (f, \tilde{z}_d) , \quad \forall \tilde{z}_d \in \mathcal{W}_{0,d} . \end{cases} \quad (1.8)$$

We recall that b stands for the trilinear form defined in (1.2). Here, we have split the nonlinear term $b(u, u, w)$ into $b(y, y, w) + b_{int}(y, z, w)$ where

$$\begin{aligned} b_{int}(y, z, w) &= b(y, z, w) + b(z, y, w) + b(z, z, w) \\ &= b(u, u, w) - b(y, y, w) . \end{aligned}$$

This term represents *the interaction between the small and large structures and the interaction of the small structures among themselves*.

Due to the decomposition of $\mathcal{V}_{g,d-1}$ into $\mathcal{V}_{g,d-2} + \mathcal{W}_{g,d-1}$, we can replace (1.7) by another system equivalent to that consisting of equations (1.7)-(1.8) where we replace $d-1$ and d by $d-2$ and $d-1$. Recursively, the same decomposition is made for each level k replacing d by k into equations (1.7)-(1.8). The process can be iterated till we reach the decomposition $\mathcal{V}_{g,0} + \mathcal{W}_{g,1}$.

For the sake of generality, we rewrite the system (1.7)-(1.8) for the decomposition $\mathcal{V}_{g,k} = \mathcal{V}_{g,k-1} + \mathcal{W}_{g,k}$, for each level k :

$$\left\{ \begin{aligned} & \left(\frac{\partial y_{k-1}}{\partial t}, \tilde{y}_{k-1} \right) + \left(\frac{\partial z_k}{\partial t}, \tilde{y}_{k-1} \right) + \frac{1}{Re} ((y_{k-1}, \tilde{y}_{k-1})) + \frac{1}{Re} ((z_k, \tilde{y}_{k-1})) + \\ & \quad b(y_{k-1}, y_{k-1}, \tilde{y}_{k-1}) + b_{int}(y_{k-1}, z_k, \tilde{y}_{k-1}) = (f, \tilde{y}_{k-1}), \\ & \quad \forall \tilde{y}_{k-1} \in \mathcal{V}_{0,k-1}, \end{aligned} \right. \quad (1.9)$$

$$\left\{ \begin{aligned} & \left(\frac{\partial y_{k-1}}{\partial t}, \tilde{z}_k \right) + \left(\frac{\partial z_k}{\partial t}, \tilde{z}_k \right) + \frac{1}{Re} ((y_{k-1}, \tilde{z}_k)) + \frac{1}{Re} ((z_k, \tilde{z}_k)) + \\ & \quad b(y_{k-1}, y_{k-1}, \tilde{z}_k) + b_{int}(y_{k-1}, z_k, \tilde{z}_k) = (f, \tilde{z}_k), \\ & \quad \forall \tilde{z}_k \in \mathcal{W}_{0,k}. \end{aligned} \right. \quad (1.10)$$

During a simulation of the Burgers problem solved by the classical Galerkin method, we perform several tests in order to evaluate the terms which appear into equations (1.9)-(1.10), for each level k , $1 \leq k \leq d$. Moreover we estimate the time variation of the terms depending upon the corrections z_k into equations (1.9) relative to y_{k-1} . The time variation is given by the difference of values over two successive time steps. This means that, here, the hierarchical basis decomposition of the nodal solution is made *a posteriori*, just for evaluation purposes.

1.4.1 Small and large scale structures

In our simulation we will consider $d = 4$ hierarchical levels. The Figure 1.2 shows the evolution of the l^2 -norm and l^∞ -norm of the ratio $\frac{z_k}{y_{k-1}}$ (for $k = 1, \dots, d$) and also the norm of the nodal solution u_d . The l^2 -norm of a vector $v_k = (v_{k,1}, \dots, v_{k,N_k}) \in \mathcal{V}_k$ is given by the following relation :

$$|v_k|_{l^2} = \sqrt{\left(\sum_{i=1}^{N_k} v_{k,i}^2 \right) / N_k}.$$

This norm is of the same order as the L^2 -norm of v_k :

$$|v_k|_{L^2} = \sqrt{v_k^T M^{(k)} v_k},$$

where $M^{(k)}$ is the mass matrix at level k :

$$M_{i,j}^{(k)} = (\phi_{k,i}, \phi_{k,j}).$$

Similarly, the l^∞ -norm of v_k is given by :

$$|v_k|_{l^\infty} = \sup_{1 \leq i \leq N_k} |v_{k,i}|,$$

where $v_{k,i} = (v_{k,i}^1, v_{k,i}^2) \in \mathbf{R}^2$, and $|v_{k,i}|^2 = ((v_{k,i}^1)^2 + (v_{k,i}^2)^2)$.

We recall that h_k is the width of triangulation $\mathcal{T}_k$, and that $h_k = \frac{h_0}{2^k}$ (Figure 1.2 corresponds to a coarsest mesh of size $h_0 = \frac{1}{16}$).

We can see that the incremental values z_k are not of the order of h_k^2 . The estimation of the order of small structures was

$$|z_k|_{l^2} \sim h_k^2 |u_d''|_{l^2} ; \quad (1.11)$$

and this relation shows that the increments z_k are not small if the solution displays strong gradients. We know of course that these gradients grow up when the Reynolds number increases.

On Figure 1.3 we plotted the l^2 and l^∞ norms of the quantities :

$$z_k^n - z_k^{n-1} = \Delta t \frac{\partial z_k^n}{\partial t} = \Delta t \dot{z}_k^n$$

as well as the norms of $y_0^n - y_0^{n-1}$ and $u_d^n - u_d^{n-1}$. These quantities represent the time variation over one time step of the hierarchical and nodal solutions.

Let us now introduce ϵ , the desired accuracy of the computation. We assume that ϵ is a given parameter, which can be chosen bases on a number of practical considerations. When we freeze the small structures z_k during a few time iterations, we introduce in the solution an error which depends on the time variation of the corrections z_k . This error can be evaluated to be of order

$$c \Delta t \left| \frac{\partial z_k(t)}{\partial t} \right|_{l^2} \sim \epsilon$$

where c is a positive constant. We will justify the above estimate in Section 1.5 devoted to the description of the space-time multilevel algorithm.

1.4.2 The time derivative terms

On Figure 1.4 we can see the time evolution of the y and z components of the variational term $(\frac{\partial u_d}{\partial t}, \tilde{u}_d)$. For each time iteration n and for all hierarchical levels k , $1 \leq k \leq d$, we plotted the l^2 -norm of the following quantities :

$$\begin{aligned} \left(\frac{\partial y_{k-1}^n}{\partial t}, \phi_{k-1,j} \right) &= \left(\sum_{i=1}^{n_{k-1}} \frac{y_{k-1,i}^n - y_{k-1,i}^{n-1}}{\Delta t} \cdot \phi_{k-1,i}, \phi_{k-1,j} \right), \quad 1 \leq j \leq n_{k-1}, \\ \left(\frac{\partial z_k^n}{\partial t}, \phi_{k-1,j} \right) &= \left(\sum_{i=n_{k-1}+1}^{n_k} \frac{z_{k,i}^n - z_{k,i}^{n-1}}{\Delta t} \cdot \psi_{k,i}, \phi_{k-1,j} \right), \quad 1 \leq j \leq n_{k-1}, \\ \left(\frac{\partial y_{k-1}^n}{\partial t}, \psi_{k,j} \right) &= \left(\sum_{i=1}^{n_{k-1}} \frac{y_{k-1,i}^n - y_{k-1,i}^{n-1}}{\Delta t} \cdot \phi_{k-1,i}, \psi_{k,j} \right), \quad n_{k-1} + 1 \leq j \leq n_k, \\ \left(\frac{\partial z_k^n}{\partial t}, \psi_{k,j} \right) &= \left(\sum_{i=n_{k-1}+1}^{n_k} \frac{z_{k,i}^n - z_{k,i}^{n-1}}{\Delta t} \cdot \psi_{k,i}, \psi_{k,j} \right), \quad n_{k-1} + 1 \leq j \leq n_k; \end{aligned} \quad (1.12)$$

$\left(\frac{\partial y_{k-1}^n}{\partial t}, \phi_{k-1,j}\right)$ (resp. $\left(\frac{\partial y_{k-1}^n}{\partial t}, \psi_{k,j}\right)$) represents the time derivative of y_{k-1} projected on the coarse grid $k-1$ (resp. on the fine grid k), whereas $\left(\frac{\partial z_k^n}{\partial t}, \phi_{k-1,j}\right)$ (resp. $\left(\frac{\partial z_k^n}{\partial t}, \psi_{k,j}\right)$) is the time derivative of z_k projected on the grid $k-1$ (resp. grid k). Obviously, these quantities are vectors. Hereafter the subscripts j will be omitted.

Considering Figure 1.4, we observe that the quantities (1.12) verify the following inequalities :

$$\left|\left(\frac{\partial z_k}{\partial t}, \psi_{k,j}\right)\right|_{l^2} \leq \left|\left(\frac{\partial z_k}{\partial t}, \phi_{k-1,j}\right)\right|_{l^2} \leq \left|\left(\frac{\partial y_{k-1}}{\partial t}, \psi_{k,j}\right)\right|_{l^2} \leq \left|\left(\frac{\partial y_{k-1}}{\partial t}, \phi_{k-1,j}\right)\right|_{l^2}$$

Nevertheless, the difference of size between these terms is not significant and consequently, we do not take into consideration the idea of neglecting certain terms by comparison with others. Such approximations done by M. Marion and R. Temam in [17] were justified because the mesh width converged to zero and thus could be assumed to be much smaller than in actual computations with realistic mesh sizes. Similarly the first numerical implementations done by J. Laminie and F. Pascal (cf. [14], [15] and [18]), developed using only two hierarchical levels, showed that the contributions of evolutive term can not be neglected. Therefore the schemes introduced in [17] are not practical, especially when the velocity vector field presents steep gradients.

To overcome this difficulty we aim to define a multilevel procedure which, during a suitable short time period, solves the Burgers equations (1.9) on a grid coarser than the nodal grid $\mathcal{T}_d$. Hence we analyze the time variation of $\left(\frac{\partial z_k}{\partial t}, \phi_{k-1,j}\right)$ which gives a linear interaction between the small scale and the large scale structures in the y -equation on the grid $\mathcal{T}_{k-1}$. For each level $k \leq d$ and for each time iteration $n \geq 2$, we evaluate :

$$\left|\left(\frac{\partial z_k^n}{\partial t} - \frac{\partial z_k^{n-1}}{\partial t}, \phi_{k-1,j}\right)\right|_{l^2} = \Delta t \left|\left(\sum_{i=n_{k-1}+1}^{n_k} \frac{z_{k,i}^n - 2z_{k,i}^{n-1} + z_{k,i}^{n-2}}{\Delta t^2} \cdot \psi_{k,i}, \phi_{k-1,j}\right)\right|_{l^2}$$

On Figure 1.5 we can observe that the time variation of $\left(\frac{\partial z_k}{\partial t}, \phi_{k-1,j}\right)$ over one time step is much smaller than the term itself. More precisely the variation is locally small even though the term itself is not so small. The above result means that the time variation of $\left(\frac{\partial z_k}{\partial t}, \tilde{y}_{k-1}\right)$ is locally negligible, and this will allow us to freeze it during a few steps of the iterations.

1.4.3 The diffusive viscous terms

As for the time derivative terms, we now analyze the evolution during actual simulations of the Burgers problem of the different parts of $\frac{1}{Re}((u_d, \tilde{u}_d))$ corresponding to the splitting associated with the hierarchical basis. For each level $1 \leq k \leq d$ and for each time

iteration n , we display the l^2 -norm of the following terms :

$$\begin{aligned}
\frac{1}{Re} ((y_{k-1}^n, \phi_{k-1,j})) &= \frac{1}{Re} \left(\sum_{i=1}^{n_{k-1}} y_{k-1,i}^n \cdot \nabla \phi_{k-1,i}, \nabla \phi_{k-1,j} \right), \quad 1 \leq j \leq n_{k-1}, \\
\frac{1}{Re} ((z_k^n, \phi_{k-1,j})) &= \frac{1}{Re} \left(\sum_{i=n_{k-1}+1}^{n_k} z_{k,i}^n \cdot \nabla \psi_{k,i}, \nabla \phi_{k-1,j} \right), \quad 1 \leq j \leq n_{k-1}, \\
\frac{1}{Re} ((y_{k-1}^n, \psi_{k,j})) &= \frac{1}{Re} \left(\sum_{i=1}^{n_{k-1}} y_{k-1,i}^n \cdot \nabla \phi_{k-1,i}, \nabla \psi_{k,j} \right), \quad n_{k-1} + 1 \leq j \leq n_k, \\
\frac{1}{Re} ((z_k^n, \psi_{k,j})) &= \frac{1}{Re} \left(\sum_{i=n_{k-1}+1}^{n_k} z_{k,i}^n \cdot \nabla \psi_{k,i}, \nabla \psi_{k,j} \right), \quad n_{k-1} + 1 \leq j \leq n_k.
\end{aligned} \tag{1.13}$$

The different components of the Laplacian term are of the same order ; consequently we can not neglect certain terms by comparison with other, a *sharp difference with spectral methods* where some terms automatically vanish. Nevertheless, considering the projection of Laplacian term on the coarse grid, we obtain :

$$|((z_k, \phi_{k-1,j}))|_{l^2} \leq |((y_{k-1}, \phi_{k-1,j}))|_{l^2}.$$

Analyzing the interactions between y_{k-1} and z_k in the diffusive terms we see that their variation over one time step is very small. On Figure 1.7 we plotted

$$\begin{aligned}
\frac{\Delta t}{Re} \left| \left(\left(\frac{\partial z_k^n}{\partial t}, \phi_{k-1,j} \right) \right) \right|_{l^2} &= \frac{1}{Re} \left| ((z_k^n - z_k^{n-1}, \phi_{k-1,j})) \right|_{l^2} \\
&= \frac{1}{Re} \left| \left(\sum_{i=n_{k-1}+1}^{n_k} (z_{k,i}^n - z_{k,i}^{n-1}) \cdot \nabla \psi_{k,i}, \nabla \phi_{k-1,j} \right) \right|_{l^2}
\end{aligned}$$

The time variation of the vector $\frac{1}{Re}((z_k, \phi_{k-1,j}))$ is much smaller than that of vector $\left(\frac{\partial z_k}{\partial t}, \phi_{k-1,j} \right)$. Consequently, as for the time derivative term, we will propose to freeze, during a few time steps, the term $\frac{1}{Re}((z_k, \tilde{y}_{k-1}))$ in the y -equation (1.9).

1.4.4 The nonlinear terms

Concluding the analysis of the hierarchical terms, we consider for each level $1 \leq k \leq d$, the time evolution of the l^2 -norm of the following quantities :

$$\begin{aligned}
b(u_d^n, u_d^n, \phi_{k-1,j}) &= ((u_d^n \cdot \nabla) u_d^n, \phi_{k-1,j}), \quad 1 \leq j \leq n_{k-1}, \\
b(u_d^n, u_d^n, \psi_{k,j}) &= ((u_d^n \cdot \nabla) u_d^n, \psi_{k,j}), \quad n_{k-1} + 1 \leq j \leq n_k.
\end{aligned} \tag{1.14}$$

We can establish from the numerical grafics on top of Figure 1.8 the following relation :

$$|b(u_d^n, u_d^n, \phi_{k-1,j})|_{l^2} \geq |b(u_d^n, u_d^n, \psi_{k,j})|_{l^2}.$$

Hence we direct our attention towards the splitting of the nonlinear term on the coarse grid $k - 1$. The splitting consists, for $1 \leq j \leq n_{k-1}$, of :

$$\begin{aligned} & b(y_{k-1}^n, y_{k-1}^n, \phi_{k-1,j}), \\ \text{and} \\ & b_{int}(y_{k-1}^n, z_k^n, \phi_{k-1,j}) = b(u_d^n, u_d^n, \phi_{k-1,j}) - b(y_{k-1}^n, y_{k-1}^n, \phi_{k-1,j}) \\ & = b(y_{k-1}^n, z_k^n, \phi_{k-1,j}) + b(z_k^n, y_{k-1}^n, \phi_{k-1,j}) + b(z_k^n, z_k^n, \phi_{k-1,j}). \end{aligned}$$

We can observe that the l^2 -norm of the vectors $b(u_d, u_d, \phi_{k-1,j})$ and $b(y_{k-1}, y_{k-1}, \phi_{k-1,j})$ are almost the same (see the left of the Figure 1.8). On the other hand, we establish numerically the following inequality :

$$|b(y_{k-1}^n, y_{k-1}^n, \phi_{k-1,j})|_{l^2} \geq |b_{int}(y_{k-1}^n, z_k^n, \phi_{k-1,j})|_{l^2},$$

which confirms that the contribution of the nonlinear interaction terms b_{int} is smaller than the component of the nonlinear term involving only the large scale y_{k-1} (see the bottom of Figure 1.8).

However, the nonlinear interaction $b_{int}(y_{k-1}, z_k, \tilde{y}_{k-1})$ can not be neglected in the y -equation (1.9) : indeed their effects on long time computations can modify the behaviour of the large scale structures ³.

Their time variation, represented on Figure 1.9, is the following :

$$\Delta t \left| \frac{\partial b_{int}}{\partial t}(y_{k-1}^n, z_k^n, \phi_{k-1,j}) \right|_{l^2} = |b_{int}(y_{k-1}^n, z_k^n, \phi_{k-1,j}) - b_{int}(y_{k-1}^{n-1}, z_k^{n-1}, \phi_{k-1,j})|_{l^2}$$

As for the time derivative and viscous terms, we can consider to freeze the b_{int} -term during some parts of the time iterations when particular relations are verified (we will present several inequalities in Section 1.5.2, devoted to estimating some parameters of our space-time multilevel procedure).

Since the calculation of nonlinear terms are very expensive, we intend to take their effects into account in a simplified manner. Hence, we can expect a significant gain on CPU time.

³This is similar to the famous butterfly effect, well-known in meteorology.

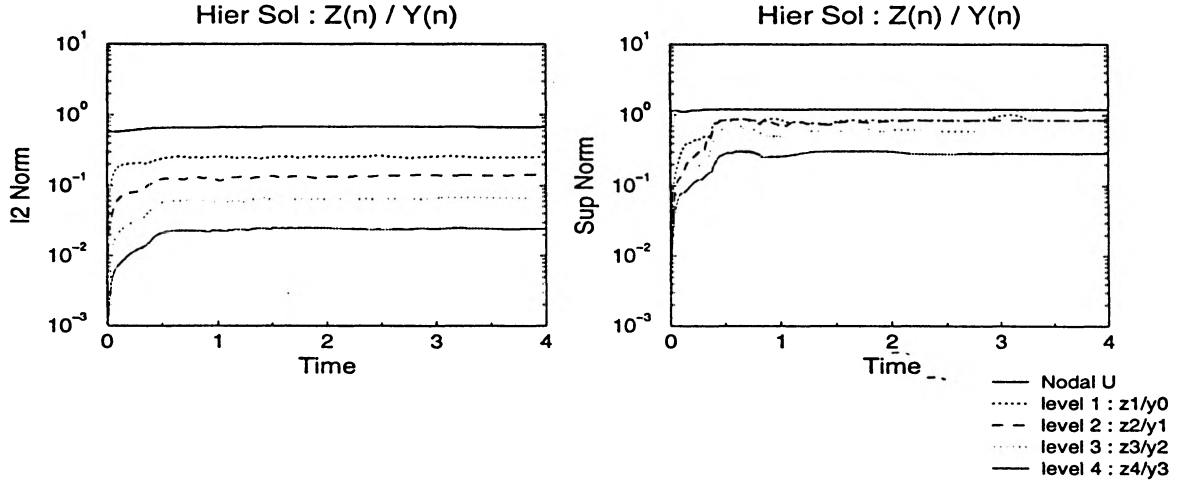


Figure 1.2: l^2 and l^∞ norms of u_d and $\frac{z_k}{y_{k-1}}$ ($k \leq d$).

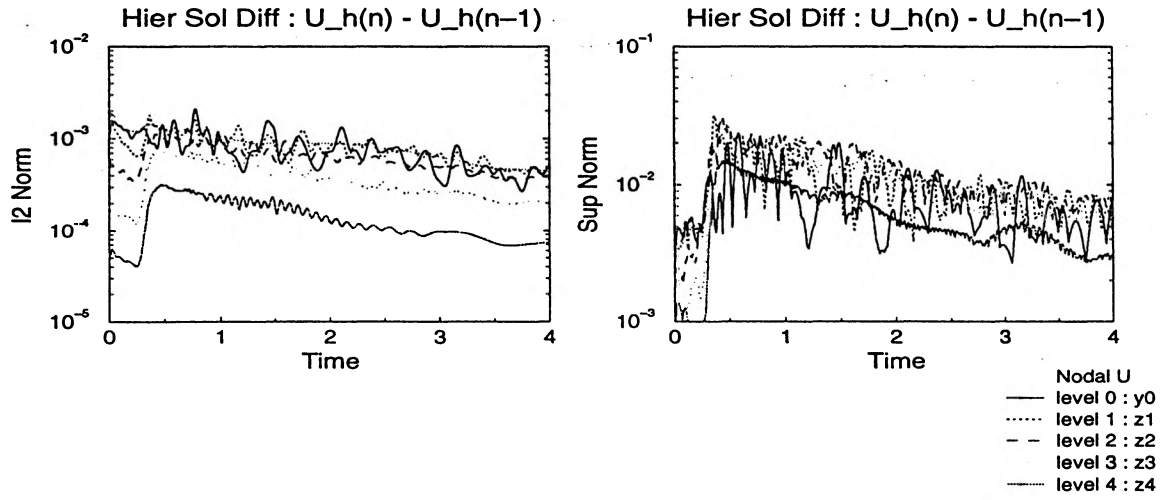


Figure 1.3: l^2 and l^∞ norms of $(u_d^n - u_d^{n-1})$, $(y_0^n - y_0^{n-1})$ and $(z_k^n - z_k^{n-1})$ ($k \leq d$).

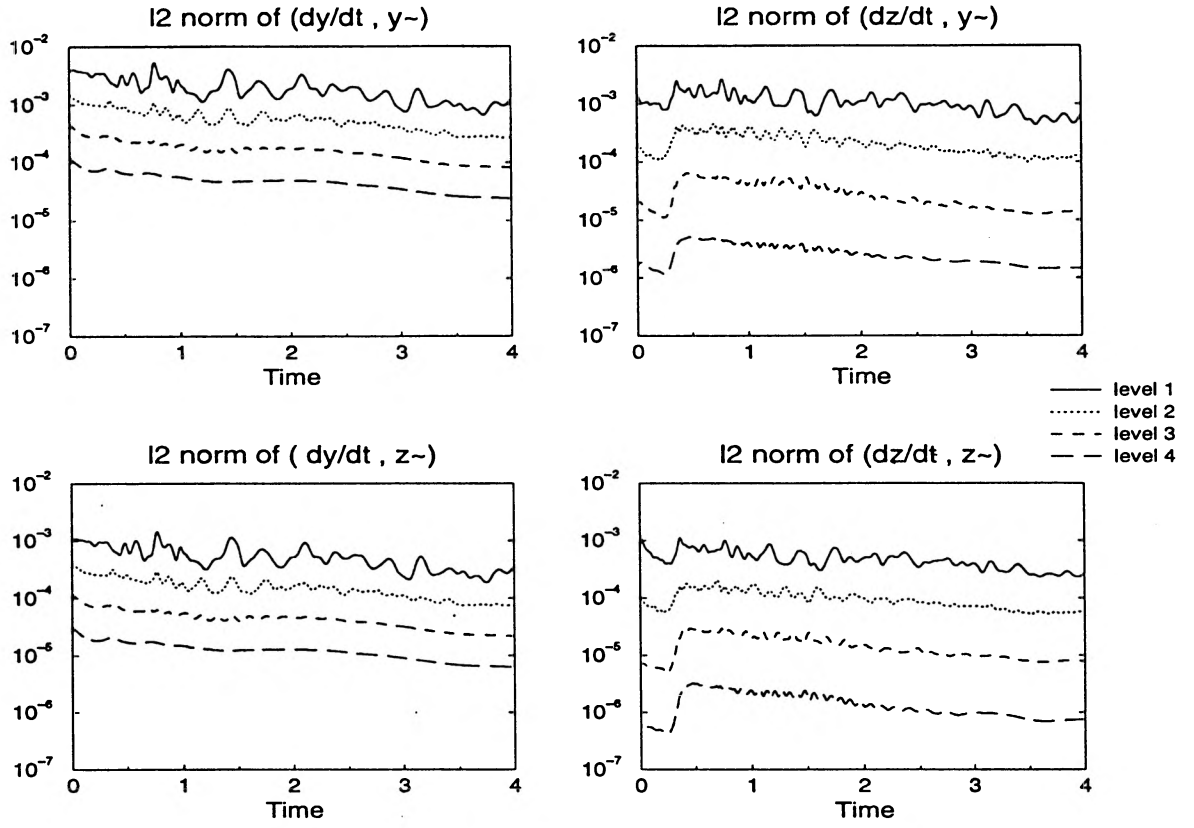


Figure 1.4: l^2 -norm of $\left(\frac{\partial y_{k-1}}{\partial t}, \phi_{k-1,j}\right)$, $\left(\frac{\partial z_k}{\partial t}, \phi_{k-1,j}\right)$, $\left(\frac{\partial y_{k-1}}{\partial t}, \psi_{k,j}\right)$ and $\left(\frac{\partial z_k}{\partial t}, \psi_{k,j}\right)$.

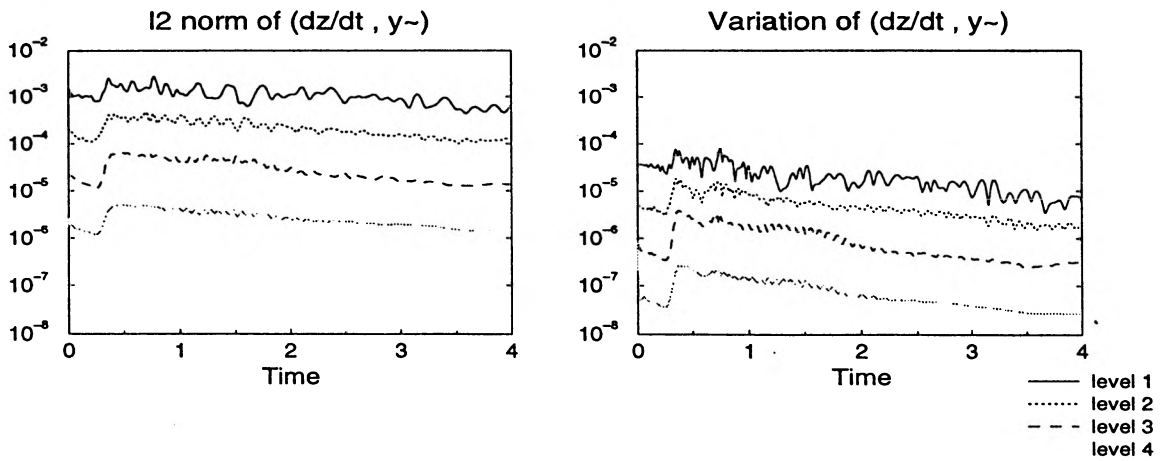


Figure 1.5: Time evolution of $\left(\frac{\partial z_k}{\partial t}, \phi_{k-1,j}\right)$ and $\Delta t \cdot \left(\frac{\partial^2 z_k}{\partial t^2}, \phi_{k-1,j}\right)$.

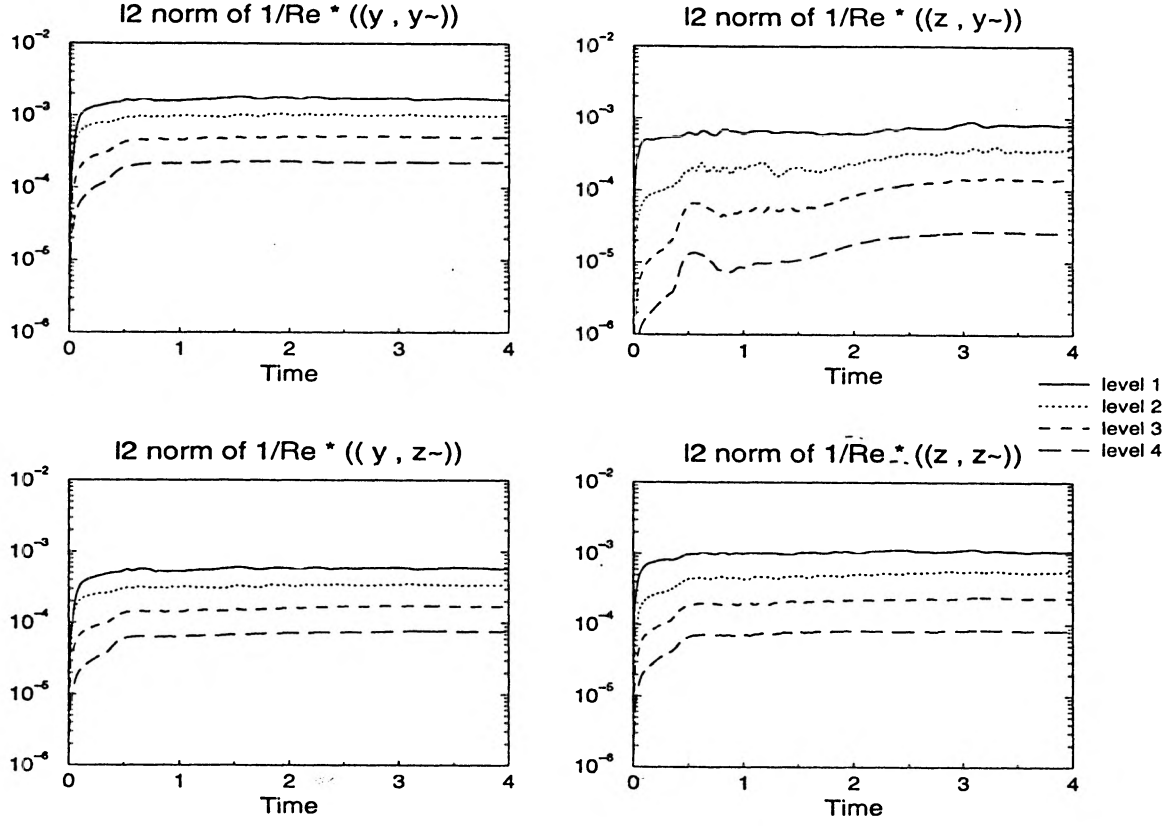


Figure 1.6: l^2 -norm of the viscous terms $\frac{1}{Re}((y_{k-1}, \phi_{k-1,j}))$, $\frac{1}{Re}((z_k, \phi_{k-1,j}))$, $\frac{1}{Re}((y_{k-1}, \psi_{k,j}))$ and $\frac{1}{Re}((z_k, \psi_{k,j}))$.

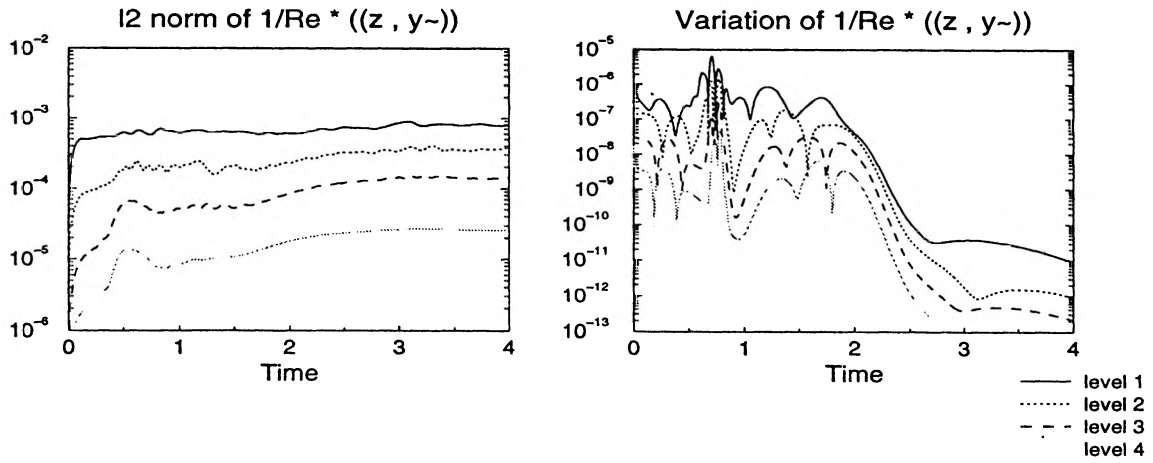


Figure 1.7: Time evolution of $\frac{1}{Re}((z_k, \phi_{k-1,j}))$ and $\Delta t \cdot \left(\left(\frac{\partial z_k}{\partial t}, \phi_{k-1,j} \right) \right)$.

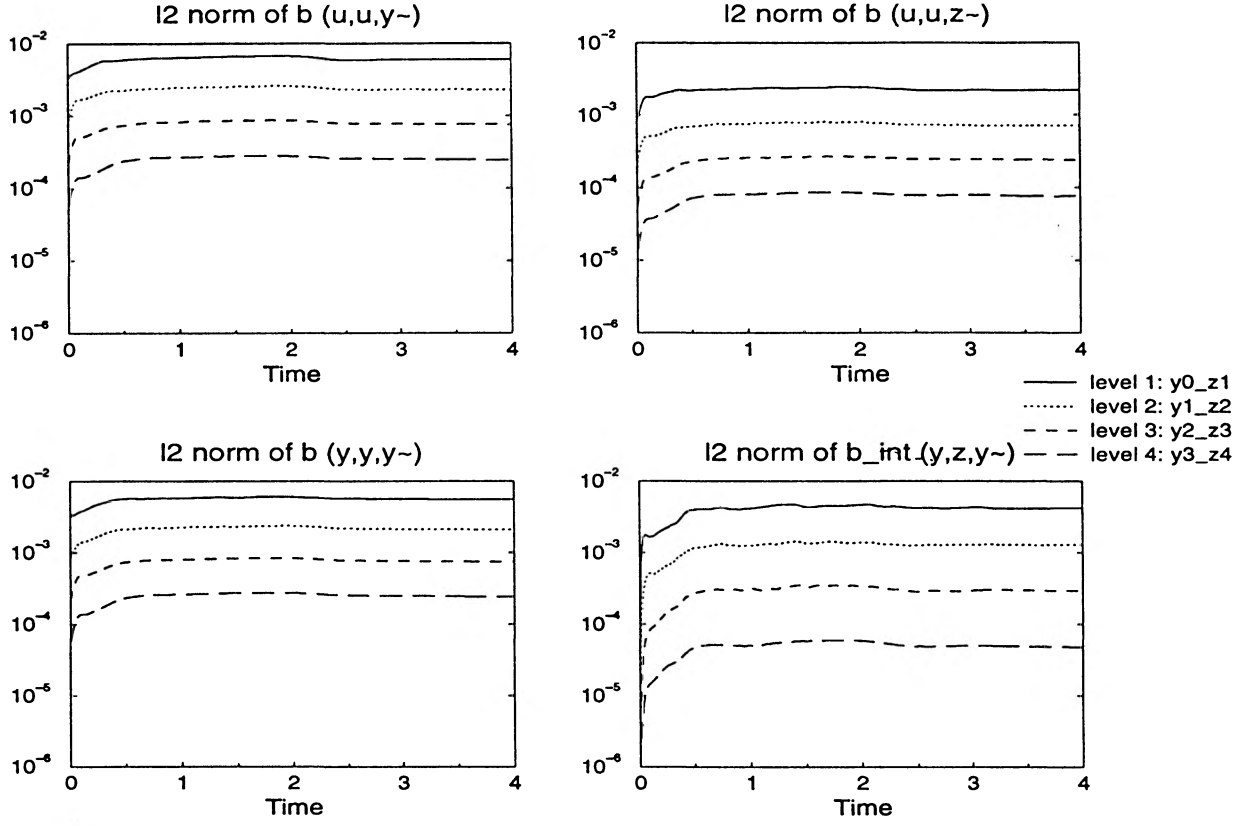


Figure 1.8: l^2 -norm of $b(u_d, u_d, \phi_{k-1,j})$, $b(u_d, u_d, \psi_{k,j})$, $b(y_{k-1}, y_{k-1}, \phi_{k-1,j})$ and $b_{int}(y_{k-1}, z_k, \phi_{k-1,j})$.

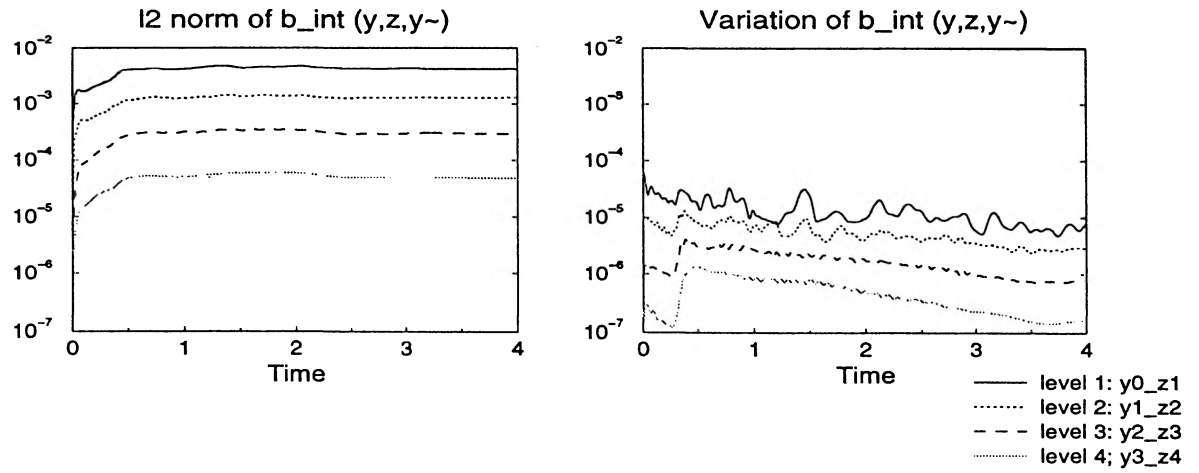


Figure 1.9: Time evolution of the nonlinear interaction term $b_{int}(y_{k-1}, z_k, \phi_{k-1,j})$ and its variation with time $\Delta t \cdot \frac{\partial b_{int}}{\partial t}(y_{k-1}, z_k, \phi_{k-1,j})$.

In the previous section we have numerically shown that the time variation of certain terms related to the small structures z_k can be considered as small during appropriate time intervals.

In this way we obtain a multilevel algorithm which realizes a completely self-adaptive procedure for solving our problem on $\mathcal{V}_d$. The procedure depends on only one parameter, namely ϵ , and it produces dynamical space-time V-cycles. The structure of a series of V-cycles is presented on Figure 1.10.



Let us now assume that the approximation $u_d(x, t)$ is known at the time t_0 ($t_0 = (n - 1)\Delta t$ on Figure 1.10). In order to characterize a series of V-cycles, we introduce the following numbers : p_V and n_V . These quantities are determined according to the tolerance ϵ :

- (i) p_V represents the depth of a V-cycle and it is determined by estimating the coarsest acceptable hierarchical level l ; consequently $p_V = d - l$. Till the coarsest level l , the time variation of the small scale structures z_k ($l < k \leq d$) is smaller than the tolerance ϵ .
- (ii) n_V represents the global number of V-cycles which can be made with the same frozen interaction terms. Let $\tau(t_0)$ be the length of the time interval during which we retain the depth p_V and we freeze the nonlinear terms $b_{int}(y_{k-1}(t), z_k(t), \phi_{k-1,j})$ at their last value, at time t_0 . After determination of p_V we can deduce the time interval t_V , during which we perform only one V-cycle

$$t_V = 2 \cdot \Delta t \cdot p_V ,$$

and the global interval $\tau(t_0)$:

$$\tau(t_0) = n_V \cdot t_V + \Delta t .$$

The nonlinear terms b_{int} can be frozen during the interval $\tau(t_0)$ without losing the order of approximation on the large scales.

The estimates for the depth p_V and the characteristic time $\tau(t_0)$, from which we deduce n_V , will be derived in Section 1.5.2. We remark that a series of V-cycles starts at the nodal level d (see Figure 1.10). Indeed, when we verify the estimate for the time variation of $b_{int}(y_{k-1}(t_0), z_k(t_0), \phi_{k-1,j})$, we must compute the value of $b(u_d(t_0), u_d(t_0), \phi_{d,j})$. After this estimation, it is not expensive to solve the equation on the nodal grid, at time $t_0 + \Delta t$.

Finally, at the end of a complete series of cycles, one to several time steps of integration are made on the nodal mesh $\mathcal{T}_d$. This procedure allows to compute the characteristic numbers for the next series. Certainly, it also plays an important role in our approach, like the projection on a Approximate Inertial Manifold in the original version of the nonlinear Galerkin method, which can be viewed as a filtering procedure ⁴.

We note here that the linear terms $\left(\frac{\partial z_k}{\partial t}(t), \phi_{k-1,j}\right)$ and $\frac{1}{Re}((z_k(t), \phi_{k-1,j}))$, unlike the nonlinear term b_{int} , are frozen at their last value during an interval less than or equal to t_V . Indeed, the small scales z_k , as well as the two linear z -terms, are updated during the upward phase of each V-cycle. It must be underlined that the fundamental role of the V-cycles is to allow a constant update of the z contributions. A choice obviously inappropriate is to realize a procedure which, during several time steps of integration, solves the equation on a mesh coarser than the nodal mesh. In this case we could produce a large accumulation of errors in the small structures.

⁴This question will be discussed elsewhere.

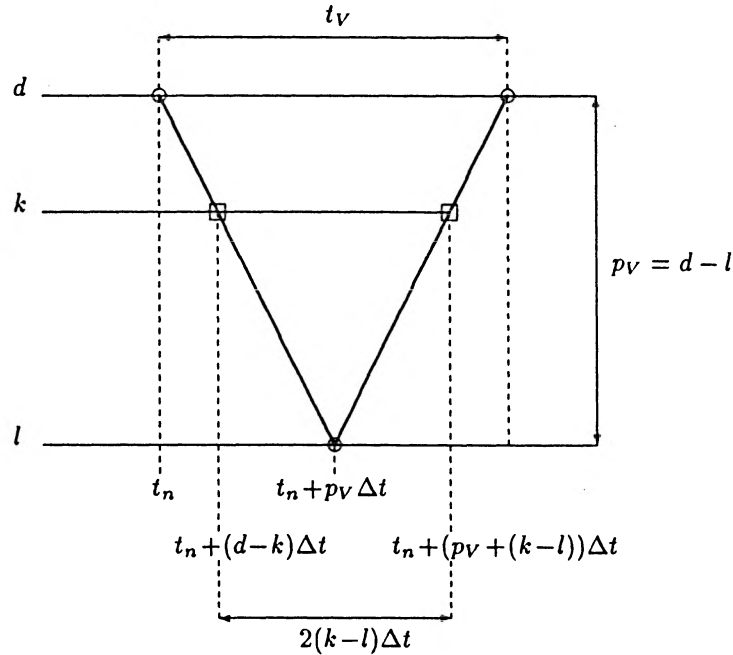


Figure 1.11: Description of a single V-cycle : estimate of the time interval for which the small structures z_k are frozen.

Let us consider in more details the integration process for the downward and the upward phases of a single V-cycle. Let us assume that the approximation $u_d(x, t)$ is known at time t_n , on the nodal mesh $\mathcal{T}_d$. For the downward phase, when we work on the mesh of level k , $l < k \leq d$, we determine at time $t_n + (d-k)\Delta t$ the approximation

$$y_k^{n+d-k} = y_{k-1}^{n+d-k} + z_k^{n+d-k},$$

as the solution of the equation projected on the grid $\mathcal{T}_k$ (see Figure 1.11). The incremental components z_k are frozen to their value computed at $t_n + (d-k)\Delta t$ during the process of the downward phase and for the upward phase till level k . Hence they are not updated on the time interval of length $2(k-l)\Delta t$ (if $k = d$, then $2(k-l)\Delta t = t_V$). As for the small structures z_k , the linear terms $\left(\frac{\partial z_k}{\partial t}(t), \phi_{k-1,j}\right)$ and $\frac{1}{Re}((z_k(t), \phi_{k-1,j}))$ are kept to the values assumed at time $t_n + (d-k)\Delta t$ and frozen during the time interval $2(k-l)\Delta t$. At the end of each V-cycle, we derive an *a posteriori* estimate of the quantities p_V and $\tau(t_0)$. Only if these estimates remain smaller than the tolerance ϵ , can we execute the next downward phase ; again the value of $b_{int}(y_{k-1}(t), z_k(t), \phi_{k-1,j})$ is replaced by that at time t_0 .

1.5.1 The multilevel procedure

We have determined $u_d(x, t)$ at the time $n\Delta t$ as the solution of the discretized equation corresponding to the nodal level d (see (1.6)). Let us consider now the split system at

the levels $d - 1$ and d . The second step of a V-cycle, starting at level d , consists in solving the y -equation on the mesh $\mathcal{T}_{d-1}$, putting each term which depends on z_d on the right hand side. We recall that the z -equation is not solved.

Let us enter in a more detailed description of the algorithm. The linear part of the y -equation is discretized with the Crank-Nicholson scheme and we use the explicit Euler method for the nonlinear part.

For $u_d^{n-1} = y_{d-1}^{n-1} + z_d^{n-1}$ and $u_d^n = y_{d-1}^n + z_d^n$ given, we compute at $t = (n + 1)\Delta t$, $u_d^{n+1} = y_{d-1}^{n+1} + z_d^{n+1}$ by the two following steps :

- **Step 1** z_d^{n+1} is frozen :

$$z_d^{n+1} = z_d^n . \quad (1.15)$$

- **Step 2** y_{d-1}^{n+1} is solution, for all $\tilde{y}_{d-1}$ in $\mathcal{V}_{0,d-1}$, of :

$$\begin{aligned} & \left(\frac{y_{d-1}^{n+1} - y_{d-1}^n}{\Delta t}, \tilde{y}_{d-1} \right) + \frac{1}{Re} \left(\left(\frac{y_{d-1}^{n+1} + y_{d-1}^n}{2}, \tilde{y}_{d-1} \right) \right) + b(y_{d-1}^n, y_{d-1}^n, \tilde{y}_{d-1}) \\ & = (f, \tilde{y}_{d-1}) - \left(\frac{z_d^n - z_d^{n-1}}{\Delta t}, \tilde{y}_{d-1} \right) - \frac{1}{Re} \left((z_d^n, \tilde{y}_{d-1}) \right) - b_{int}(y_{d-1}^{n-1}, z_d^{n-1}, \tilde{y}_{d-1}) . \end{aligned} \quad (1.16)$$

In equation (1.16) some approximations have been made for the z -terms. More precisely :

- for the time derivative term :

As the time variation of $\left(\frac{\partial z_d}{\partial t}(t), \tilde{y}_{d-1} \right)$ is small, it is approximated by :

$$\left(\frac{\partial z_d}{\partial t}(t), \tilde{y}_{d-1} \right) \cong \left(\frac{z_d^n - z_d^{n-1}}{\Delta t}, \tilde{y}_{d-1} \right) ; \quad (1.17)$$

- for the viscous term :

Due to step 1, the Crank-Nicholson scheme can be reduced to an explicit Euler scheme on the term : $\frac{1}{Re} \left((z_d(t), \tilde{y}_{d-1}) \right)$;

- for the nonlinear interaction :

$b_{int}(y_{k-1}(t), z_k(t), \tilde{y}_{d-1})$ is computed for $t = t_0 = (n - 1)\Delta t$ and stored before the beginning of the series of V-cycle for each mesh $k, l < k \leq d$.

Remark 1 Due to step 1, if the evolutive z -term is evaluated at $(n+1)\Delta t$, one obtains :

$$\left(\frac{\partial z_d^{n+1}}{\partial t}, \tilde{y}_{d-1} \right) \cong \left(\frac{z_d^{n+1} - z_d^n}{\Delta t}, \tilde{y}_{d-1} \right) = 0 . \quad (1.18)$$

The numerical schemes proposed by M. Marion and R. Temam in [17] neglected this evolutive term, by analogy with the schemes performed on the pseudo-spectral case (where the corresponding y_{k-1} and z_k are orthogonal for the L^2 and H^1 norms). The implementation of the nonlinear Galerkin methods, developed by J. Laminie and F. Pascal in the context of the finite elements discretization (cf. [14] and [15]), consisted in finding

an approximation u_h of the solution, decomposed into $y_{2h} + z_h$, where y_{2h} and z_h are solutions of a highly coupled system (the coarse grid has mesh of width $2h$ and h is the mesh size of the finer grid). The schemes performed in the finite element case neglected some terms of the coupled system in y_{2h} and z_h , especially the evolutive terms

$$\left(\frac{\partial z_h}{\partial t}(t), \tilde{y}_{2h} \right) \quad \text{and} \quad \left(\frac{\partial y_{2h}}{\partial t}(t), \tilde{z}_h \right)$$

(here, $\tilde{y}_{2h}$ and $\tilde{z}_h$ are functions test defined on the corresponding hierarchical spaces). Moreover, further approximations were performed on the z -equation which became a linear problem in the z_h unknown. The authors deduced that these algorithms are performing in the case of sufficiently small Reynolds numbers of else for very small mesh sizes. Otherwise, steep gradients involve some instability and the approximations made on the z -terms are not justified.

In the implementation of our multilevel procedure we do not enforce the approximation (1.18) which would introduce further significant error. For this reason the evolutive term is approximated at its last value, namely we enforce (1.17).

The process of the downward phase stops on the coarsest level of the V-cycle, namely $l = d - p_V$. We determine y_l^{m+1} at time $t = (m + 1)\Delta t$, where $m = n + p_V - 1$, by an equation similar to (1.16) ; more precisely, for all $\tilde{y}_l \in \mathcal{V}_{0,l}$:

$$\begin{aligned} & \left(\frac{y_l^{m+1} - y_l^m}{\Delta t}, \tilde{y}_l \right) + \frac{1}{Re} \left(\left(\frac{y_l^{m+1} + y_l^m}{2}, \tilde{y}_l \right) \right) + b(y_l^m, y_l^m, \tilde{y}_l) \\ & = (f, \tilde{y}_l) - \left(\frac{z_{l+1}^m - z_{l+1}^{m-1}}{\Delta t}, \tilde{y}_l \right) - \frac{1}{Re} \left((z_{l+1}^m, \tilde{y}_l) \right) - b_{int}(y_l^{n-1}, z_{l+1}^{n-1}, \tilde{y}_l) - (F, \tilde{y}_l), \end{aligned} \quad (1.19)$$

where $(F, \tilde{y}_l)$ is the projection on the grid $\mathcal{T}_l$ of all the z -terms frozen on the previous levels $k = l + 2, \dots, d$. Certainly, this term was also in the y -equations related to the above levels, till level $d - 2$. If we fully neglect this contribution, the difference between our approximation and the classical solution increases ; moreover this term has been computed on the above levels and only a projection on the present level of solution must be made.

For the upward phase, an equation similar to (1.16) is solved in order to obtain y_{k-1}^{m+1} , for $l + 1 < k \leq d$ and $m = n + p_V + k - l - 1$. We note that during the upward process the evolutive z -term as well as the viscous z -term were frozen since $2(k - l)$ time steps. We conclude a V-cycle with the iteration on the nodal mesh $\mathcal{T}_d$, which determines $u_d(x, t)$, solution of (1.6), at $t = n\Delta t + t_V$. At the end of each V-cycle some inexpensive tests are made in order to verify the validity of the current V-cycle series. Hence, we take into account the evolution of the small scales as well as that of the nonlinear term (these tests verify the estimates of the time variation relatively to z_k and b_{int}). If all estimates are smaller than the tolerance ϵ , we can proceed and perform another V-cycle. Otherwise the procedure is stopped and, after one to several time steps of integration on the nodal (fine) mesh, we compute the new characteristic numbers p_V and n_V for the next series. The complete algorithm is described in the Table 1.1.

Table 1.1: Description of the multilevel algorithm

- Initialization :
 - ★ construct the family of nested meshes,
 - ★ assemble the mass and stiffness matrices and the source terms,
 - ★ initialize the solution,
 - ★ solve the equation on the nodal mesh $\mathcal{T}_d$ determining u_d ,
 - ★ set $vcycle = 0$.
- Time iterations :
 - if ($vcycle = 0$) then
 - ★ solve the equation on the nodal mesh $\mathcal{T}_d$ determining u_d ,
 - ★ compute the time variation of the z_k -terms,
 - ★ if ($variation < \epsilon$) then
 - evaluate the depth of a V-cycle : $p_V = d - l$,
 - evaluate the maximum number of V-cycles : n_V ,
 - set $vcycle = 1$;
 - else
 - ★ evaluate and store $b_{int}(y_{k-1}, z_k, \tilde{y}_{k-1})$ (for $k = d$ till $l + 1$),
 - ★ solve the equation on the nodal mesh $\mathcal{T}_d$ determining u_d ,
 - ★ **V-cycle series** : execute $2 \cdot n_V \cdot p_V$ time steps :
 - for $i_V = 1$ to n_V :
 - **downward phase** :
 - for $k = d$ down to $l + 1$
 - compute y_{k-1} solution of (1.9) on the grid $\mathcal{T}_{k-1}$,
 - **upward phase** :
 - for $k = l + 2$ to $d - 1$
 - compute y_{k-1} solution of (1.9) on the grid $\mathcal{T}_{k-1}$,
 - for $k = d$ determine u_d on the nodal mesh $\mathcal{T}_d$,
 - evaluate $\Delta z = 2\Delta t \left| \frac{\partial z_{l+1}}{\partial t} \right|_{l^2}$ and $\Delta b = \left| \frac{\partial b_{int}(y_l, z_{l+1}, \phi_{l,j})}{\partial t} \right|_{l^2}$
 - if ($\Delta z > \epsilon$) or ($\Delta b > \epsilon \cdot (\tau(t_0))^{-2}$) then
 - set $vcycle = 0$,
 - get out of the V-cycle series
 - next i_V
 - end if.

1.5.2 Time scale estimate for the small structures and for the linear and nonlinear interaction terms

After the description of the numerical scheme realizing the space-time V-cycles, we direct our attention towards the criteria determining the depth of a V-cycle as well as the *a priori* estimates establishing the global time length $\tau(t_0)$ during which a series of V-cycles can be executed.

A. Debussche, in a personal communication, provides us the fundamental idea for obtaining the criteria performing the characteristic quantities of a series of V-cycles. Their theoretical justification is lightly different from that presented in the pseudo-spectral case (see [3]).

1.5.2.1 Estimate of p_V , the depth of a V-cycle

In the case of a pseudo-spectral discretization, one chooses an integer N which represents the total number of modes retained, relatively to each space direction, in the truncation of the Fourier expansion of the solution. In this case one can define two levels of discretization $N_{i1}(t)$ and $N_{i2}(t)$, changing during the time, such that $N_{i1}(t) < N_{i2}(t) < N$. Then, a V-cycle consists of $2(i_2 - i_1) + 1$ time iterations between the levels N_{i1} and N_{i2} . These levels are chosen so that suitable estimates of the L^2 and H^1 norms of the ratio $\frac{z_{N_{ip}}}{y_{N_{ip}}}$ ($p = 1$ or 2) are verified. Mainly, one requires

$$\frac{|z_{N_{i2}}|_{L^2}}{|y_{N_{i2}}|_{L^2}} \leq \theta ,$$

where θ is a constant of the order of ϵ^2 .

In the finite elements case, [17], [15] and our numerical experiments show that the incremental terms z_k can be much larger than h_k^2 , where h_k is the width of the mesh $\mathcal{T}_k$. This is one of the most important differences with the spectral case : the ratio between the small scale and the large scale components is not as small as the spectral ratio.

During a V-cycle, we restrict the solution of the Burgers equation to the y_{k-l} -equation, where $l < k \leq d$ (see steps (1.15) and (1.16)). The corrections z_k are frozen during the time interval $2(k - l)\Delta t$, where l is the coarsest level of solution (as shown on Figure 1.11). Here, we impose that their time variation be smaller than the parameter ϵ .

For the upward phase, in order to get $y_k(t)$, $t = t_n + (p_V + k - l + 1)\Delta t$, we solve the y -equation on the grid $\mathcal{T}_k$, but we commit an error on the initial data. Actually we have replaced the initial data

$$y_k(t) = y_{k-1}(t) + z_k(t) , \quad \text{where } t = t_n + (p_V + k - l)\Delta t,$$

by

$$y_{k-1}(t_n + (p_V + k - l)\Delta t) + z_k(t_n + (d - k)\Delta t) .$$

From this, we infer an estimate of the error at the level k :

$$e(k) = \int_{t_n + (d-k)\Delta t}^{t_n + (p_V + k - l)\Delta t} [z_k(s) - z_k(t_n + (d - k)\Delta t)] ds .$$

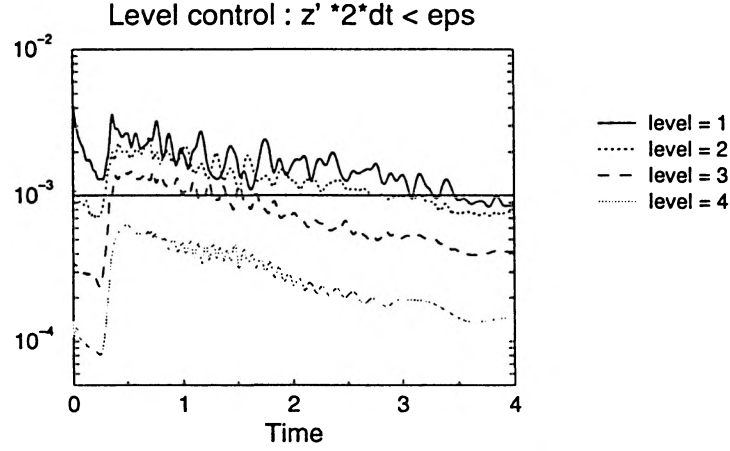


Figure 1.12: Evolution of the l^2 norm of $2\Delta t \cdot \frac{\partial z_k^m}{\partial t}$; here $\epsilon = 10^{-3}$.

Thus we obtain :

$$e(k) \leq 2(k-l)\Delta t \cdot \sup_t \left| \frac{\partial z_k}{\partial t}(t) \right|_{l^2} ,$$

with $t \in [t_n + (d-k)\Delta t, t_n + (p_V + k-l)\Delta t]$. Hence, we impose, for $k = l+1, \dots, d$:

$$e(k) \leq 2(k-l)\Delta t \cdot \sup_t \left| \frac{\partial z_k}{\partial t}(t) \right|_{l^2} < \epsilon .$$

We observe on Figure 1.12 that, for each time iteration m , the quantity $|z_k^m - z_k^{m-1}|_{l^2} = \Delta t \cdot \left| \frac{\partial z_k}{\partial t}(m\Delta t) \right|_{l^2}$ increases with decreasing values of the level k . In this case we choose the level l as the coarsest level for which the error $e(l+1)$ introduced on the initial data is smaller than ϵ . Then we obtain the following estimate :

$$2\Delta t \left| \frac{\partial z_{l+1}^m}{\partial t} \right|_{l^2} = 2|z_{l+1}^m - z_{l+1}^{m-1}|_{l^2} < \epsilon . \quad (1.20)$$

This test is not expensive and it inhibits the beginning of a series of V-cycles when $|z_d^m - z_d^{m-1}|_{l^2} \geq \frac{\epsilon}{2}$ (see Figure 1.12).

In order to control the time variation of z_k on the different levels, a sufficient condition is to estimate (1.20) on the lower level $l+1$, at the end of each V-cycle. This condition can be written in the form :

$$2\Delta t \left| \frac{\partial z_{l+1}}{\partial t}(t_0 + qt_V) \right|_{l^2} < \epsilon , \quad \text{for } q = 0, \dots, n_V - 1 . \quad (1.21)$$

If (1.21) is violated, then the V-cycle process is stopped.

1.5.2.2 Estimate of $\tau(t_0)$, the time interval for the transfer terms

Now we derive the estimate of the global time length $\tau(t_0)$ of the whole series of V-cycles on $[t_0, t_0 + \tau(t_0)]$ during which the coupled z -terms are frozen in the y -equation (1.9), namely :

$$\left(\frac{\partial z_k}{\partial t}(t), \tilde{y}_{k-1} \right), \quad \frac{1}{Re} ((z_k(t), \tilde{y}_{k-1})) \quad \text{et} \quad b_{int}(y_{k-1}(t), z_k(t), \tilde{y}_{k-1}) . \quad (1.22)$$

We recall that we do not freeze these terms in the same manner. Indeed, the transfer terms b_{int} are frozen during several V-cycles and kept at their value at time t_0 (independently of the level k) ; whereas the two linear terms are updated inside a V-cycle and they are frozen during the time interval $2(k-l)\Delta t$, less than or equal to t_V .

We start by considering the error introduced in a single V-cycle, at level $k-1$ ($l < k \leq d$), when the computation proceeds from $t_0 + m\Delta t$ to $t_0 + (m+1)\Delta t$, where $m = 0, \dots, 2p_V - 1$. We approximate the exact equation

$$\begin{cases} \left(\frac{\partial y_{k-1}}{\partial t}(t), \tilde{y}_{k-1} \right) + \left(\frac{\partial z_k}{\partial t}(t), \tilde{y}_{k-1} \right) + \frac{1}{Re} ((y_{k-1}(t), \tilde{y}_{k-1})) + \frac{1}{Re} ((z_k(t), \tilde{y}_{k-1})) \\ + b(y_{k-1}(t), y_{k-1}(t), \tilde{y}_{k-1}) + b_{int}(y_{k-1}(t), z_k(t), \tilde{y}_{k-1}) \\ = (f, \tilde{y}_{k-1}) \quad \forall \tilde{y}_{k-1} \in \mathcal{V}_{0,k-1} , \end{cases} \quad (1.23)$$

by the same equation where the coupled terms (1.22) are replaced by those frozen at their last value, namely

$$\left(\frac{\partial z_k}{\partial t}(t_i), \tilde{y}_{k-1} \right), \quad \frac{1}{Re} ((z_k(t_i), \tilde{y}_{k-1})) \quad \text{and} \quad b_{int}(y_{k-1}(t_0), z_k(t_0), \tilde{y}_{k-1}) . \quad (1.24)$$

We note t_0 the time corresponding to the beginning of the V-cycle series and

$$t_i = \begin{cases} t_0 + m\Delta t & \text{for the downward phase,} \\ t_0 + (m - 2(k-l) - 1)\Delta t & \text{for the upward phase.} \end{cases}$$

Hence, the error made on the equation used to compute $y_{k-1}(t)$ is exactly given by the integral, over one time step, of the difference between the exact equation (1.23) and its approximation, namely, for all $\tilde{y}_{k-1} \in \mathcal{V}_{0,k-1}$:

$$\begin{aligned} e(m) &= \int_{t_0+m\Delta t}^{t_0+(m+1)\Delta t} \left(\frac{\partial z_k}{\partial t}(s) - \frac{\partial z_k}{\partial t}(t_i), \tilde{y}_{k-1} \right) ds \\ &+ \frac{1}{Re} \int_{t_0+m\Delta t}^{t_0+(m+1)\Delta t} ((z_k(s) - z_k(t_i), \tilde{y}_{k-1})) ds \\ &+ \int_{t_0+m\Delta t}^{t_0+(m+1)\Delta t} [b_{int}(y_{k-1}(s), z_k(s), \tilde{y}_{k-1}) - b_{int}(y_{k-1}(t_0), z_k(t_0), \tilde{y}_{k-1})] ds . \end{aligned} \quad (1.25)$$

Remark 2 On the right hand side of the equation approximating (1.23), we have not considered the possible z -terms frozen on the previous levels and projected on the current level $k - 1$. This choice involves an incomplete estimation of the error. Nevertheless (1.25) and the successive estimates are sufficient to insure the accuracy of the algorithm.

Let $\tilde{y}_{k-1} = \phi_{k-1,j}$ ($1 \leq j \leq n_{k-1}$) ; by the mean value theorem :

$$\begin{aligned}
 e(m) \leq & \sup_{t \in [t_0+m\Delta t, t_0+(m+1)\Delta t]} \left| \left(\frac{\partial^2 z_k}{\partial t^2}(t), \phi_{k-1,j} \right) \right|_{l^2} \cdot \int_{t_0+m\Delta t}^{t_0+(m+1)\Delta t} (s - t_i) ds \\
 & + \sup_{t \in [t_0+m\Delta t, t_0+(m+1)\Delta t]} \frac{1}{Re} \left| \left(\left(\frac{\partial z_k}{\partial t}(t), \phi_{k-1,j} \right) \right) \right|_{l^2} \cdot \int_{t_0+m\Delta t}^{t_0+(m+1)\Delta t} (s - t_i) ds \quad (1.26) \\
 & + \sup_{t \in [t_0+m\Delta t, t_0+(m+1)\Delta t]} \left| \frac{\partial b_{int}}{\partial t}(y_{k-1}(t), z_k(t), \phi_{k-1,j}) \right|_{l^2} \cdot \int_{t_0+m\Delta t}^{t_0+(m+1)\Delta t} (s - t_0) ds
 \end{aligned}$$

Obviously

$$\int (s - t_i) ds \leq \int (s - t_0) ds ,$$

and observing that

$$\int_{t_0+m\Delta t}^{t_0+(m+1)\Delta t} (s - t_0) ds = \frac{((m+1)\Delta t)^2 - (m\Delta t)^2}{2} ,$$

we finally have :

$$e(m) \leq (M_1 + M_2 + M_3) \cdot \frac{((m+1)\Delta t)^2 - (m\Delta t)^2}{2} , \quad (1.27)$$

where

$$\begin{aligned}
 M_1 &= \sup_{t \in [t_0, t_0+t_V]} \left| \left(\frac{\partial^2 z_k}{\partial t^2}(t), \phi_{k-1,j} \right) \right|_{l^2} , \\
 M_2 &= \sup_{t \in [t_0, t_0+t_V]} \frac{1}{Re} \left| \left(\left(\frac{\partial z_k}{\partial t}(t), \phi_{k-1,j} \right) \right) \right|_{l^2} , \\
 M_3 &= \sup_{t \in [t_0, t_0+t_V]} \left| \frac{\partial b_{int}}{\partial t}(y_{k-1}(t), z_k(t), \phi_{k-1,j}) \right|_{l^2} .
 \end{aligned} \quad (1.28)$$

We find the sum of the inequality (1.27) from $m = 0$ to $2p_V - 1$. The sum of these errors, during a single V-cycle, is lower than

$$(M_1 + M_2 + M_3) \cdot \frac{t_V^2}{2} . \quad (1.29)$$

In conclusion, if we prescribe :

$$(M_1 + M_2 + M_3) < \epsilon \cdot t_V^{-2} \quad (1.30)$$

we introduce in the y_{k-1} -equation, during a single V-cycle, an error lower than $\frac{\epsilon}{2}$ (this is true for each hierarchical level between the coarsest level l and the nodal level d).

Henceforth, we consider the global error introduced when executing the whole series of V-cycles. The error over one time steps is given by inequality (1.25), but hence-forward m changes from 0 to $\frac{\tau(t_0)}{\Delta t} - 1$. We deduce the following inequality :

$$e(m) \leq (\overline{M}_1 + \overline{M}_2 + \overline{M}_3) \cdot \frac{((m+1)\Delta t)^2 - (m\Delta t)^2}{2}, \quad (1.31)$$

where we have set :

$$\begin{aligned} \overline{M}_1 &= \sup_{\substack{t \in [t_0 + qt_V, t_0 + (q+1)t_V] \\ 0 \leq q \leq n_V - 1, \, l < k \leq d}} \left| \left(\frac{\partial^2 z_k}{\partial t^2}(t), \phi_{k-1,j} \right) \right|_{l^2}, \\ \overline{M}_2 &= \sup_{\substack{t \in [t_0 + qt_V, t_0 + (q+1)t_V] \\ 0 \leq q \leq n_V - 1, \, l < k \leq d}} \frac{1}{Re} \left| \left(\left(\frac{\partial z_k}{\partial t}(t), \phi_{k-1,j} \right) \right) \right|_{l^2}, \\ \overline{M}_3 &= \sup_{\substack{t \in [t_0, t_0 + \tau(t_0)] \\ l < k \leq d}} \left| \frac{\partial b_{int}}{\partial t}(y_{k-1}(t), z_k(t), \phi_{k-1,j}) \right|_{l^2}. \end{aligned} \quad (1.32)$$

To find the sum of (1.31) for m varying from 0 to $\frac{\tau(t_0)}{\Delta t} - 1$, we split the right hand side of this estimate in two parts : the two linear contributions and the nonlinear one. For the two quantities deriving from linear terms, the error over the whole series is given by n_V times the error committed during one cycle (see the estimate (1.29)), whereas for the quantity which derives from the nonlinear interaction terms we compute

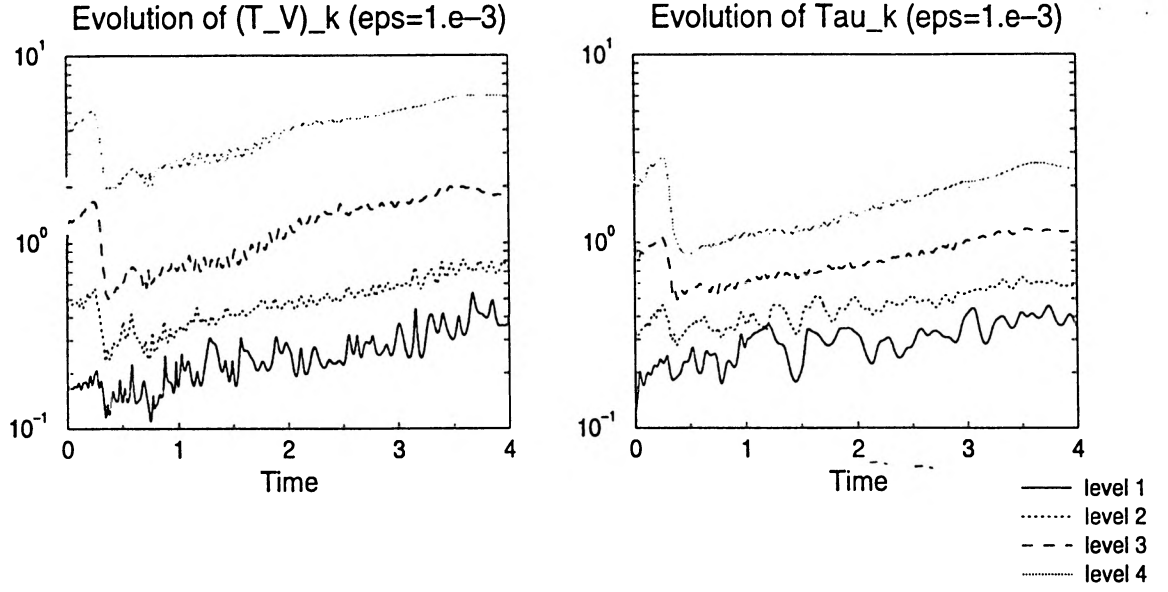
$$\overline{M}_3 \cdot \sum_{m=0}^{\frac{\tau(t_0)}{\Delta t} - 1} \frac{((m+1)\Delta t)^2 - (m\Delta t)^2}{2} = \overline{M}_3 \cdot \frac{(\tau(t_0))^2}{2}$$

Then, to insure that the global error introduced on the y -equation be lower than ϵ , we prescribe :

$$(\overline{M}_1 + \overline{M}_2) \cdot n_V \cdot \frac{t_V^2}{2} + \overline{M}_3 \cdot \frac{(\tau(t_0))^2}{2} < \epsilon. \quad (1.33)$$

Consequently, we will enforce the two following inequalities :

$$\begin{aligned} \sup_{\substack{t \in [t_0 + qt_V, t_0 + (q+1)t_V] \\ 0 \leq q \leq n_V - 1, \, l < k \leq d}} \left\{ \left| \left(\frac{\partial^2 z_k}{\partial t^2}(t), \phi_{k-1,j} \right) \right|_{l^2} \right. \\ \left. + \frac{1}{Re} \left| \left(\left(\frac{\partial z_k}{\partial t}(t), \phi_{k-1,j} \right) \right) \right|_{l^2} \right\} < \frac{\epsilon \cdot t_V^{-2}}{n_V}, \end{aligned} \quad (1.34)$$


 Figure 1.13: Estimate of $(t_V)_k$ and τ_k for $k = 0, \dots, d$.

and

$$\sup_{\substack{t \in [t_0, t_0 + \tau(t_0)] \\ l < k \leq d}} \left| \frac{\partial b_{int}}{\partial t} (y_{k-1}(t), z_k(t), \phi_{k-1,j}) \right|_{l^2} < \epsilon (\tau(t_0))^{-2}. \quad (1.35)$$

For the sake of simplicity, we note with the dot the differentiation with respect to t . We estimate numerically the two following quantities :

$$\tau_k = \left(\frac{\epsilon}{\left| \dot{b}_{int} (y_{k-1}(t), z_k(t), \phi_{k-1,j}) \right|_{l^2}} \right)^{\frac{1}{2}},$$

and

$$(t_V)_k = \left(\frac{\epsilon}{\left| (\ddot{z}_k(t), \phi_{k-1,j}) \right|_{l^2} + \frac{1}{Re} \left| (\dot{z}_k(t), \phi_{k-1,j}) \right|_{l^2}} \right)^{\frac{1}{2}}.$$

On Figure 1.13 we can see that τ_k , as well as $(t_V)_k$, decreases with decreasing values of the level k ($l < k \leq d$). Moreover we note that $(t_V)_k$ is of the order of τ_k . From (1.35) we deduce, obviously,

$$\tau(t_0) \leq \tau_k, \quad \text{for } k = l+1, \dots, d.$$

Instead, as the time $\tau(t_0)$ is adjusted to be a multiple of t_V , the period of a single V-cycle, we deduce from (1.34) :

$$\tau(t_0) \simeq n_V \cdot t_V = \sqrt{n_V} (\sqrt{n_V} t_V) < \sqrt{n_V} (t_V)_k, \quad \text{for } k = l+1, \dots, d.$$

In practice it appears that estimate (1.35) is generally more restrictive than estimate (1.34).

Hence-forward, the global time length of the whole series of V-cycles is written

$$\tau(t_0) = \left(\frac{\epsilon}{\left| \frac{\partial b_{int}}{\partial t}(y_l(t_0), z_{l+1}(t_0), \phi_{l,j}) \right|_{l^2}} \right)^{\frac{1}{2}}.$$

A sufficient condition is to verify the estimate (1.35) (but not necessarily (1.34)) at the end of each V-cycle, only on the coarsest level l . The estimate (1.35) can be written under the form :

$$\left| \frac{\partial b_{int}}{\partial t}(y_l(t_0 + qt_V), z_{l+1}(t_0 + qt_V), \phi_{l,j}) \right|_{l^2} < \epsilon \tau(t_0)^{-2}, \quad \text{for } q = 0, \dots, n_V - 1. \quad (1.36)$$

Once more, the series of V-cycles is stopped when (1.36) is violated.

Remark 3 *The time variation of the nonlinear interaction terms b_{int} is never computed inside a series of V-cycles, because its estimate is very expensive. In our numerical implementations, inequality (1.36) is exactly verified for $q = 0$, i.e. only before the beginning of the V-cycle series.*

In order to verify the constraint on the nonlinear term $\left| \frac{\partial b_{int}}{\partial t}(y_l, z_{l+1}, \phi_{l,j}) \right|_{l^2}$, we put some terms, easily calculable, in the place of the time variation of b_{int} . Naturally these terms must have the same size. The numerical experiments show, as in the pseudo-spectral case, a certain correlation between

$$\frac{\Delta t \cdot \left| \frac{\partial b_{int}}{\partial t}(y_l, z_{l+1}, \phi_{l,j}) \right|_{l^2}}{|b_{int}(y_l, z_{l+1}, \phi_{l,j})|_{l^2}} \quad \text{and} \quad \frac{\Delta t \cdot \left| \frac{\partial z_k}{\partial t} \right|_{l^2}}{|z_k|_{l^2}}. \quad (1.37)$$

This kind of correlation is accurate on the finer hierarchical levels (see Figure 1.14), but we lose the accuracy on the coarsest levels. In the near future our efforts will be concentrated on finding others correlations better than (1.37), especially in the case of the Navier-Stokes problem.

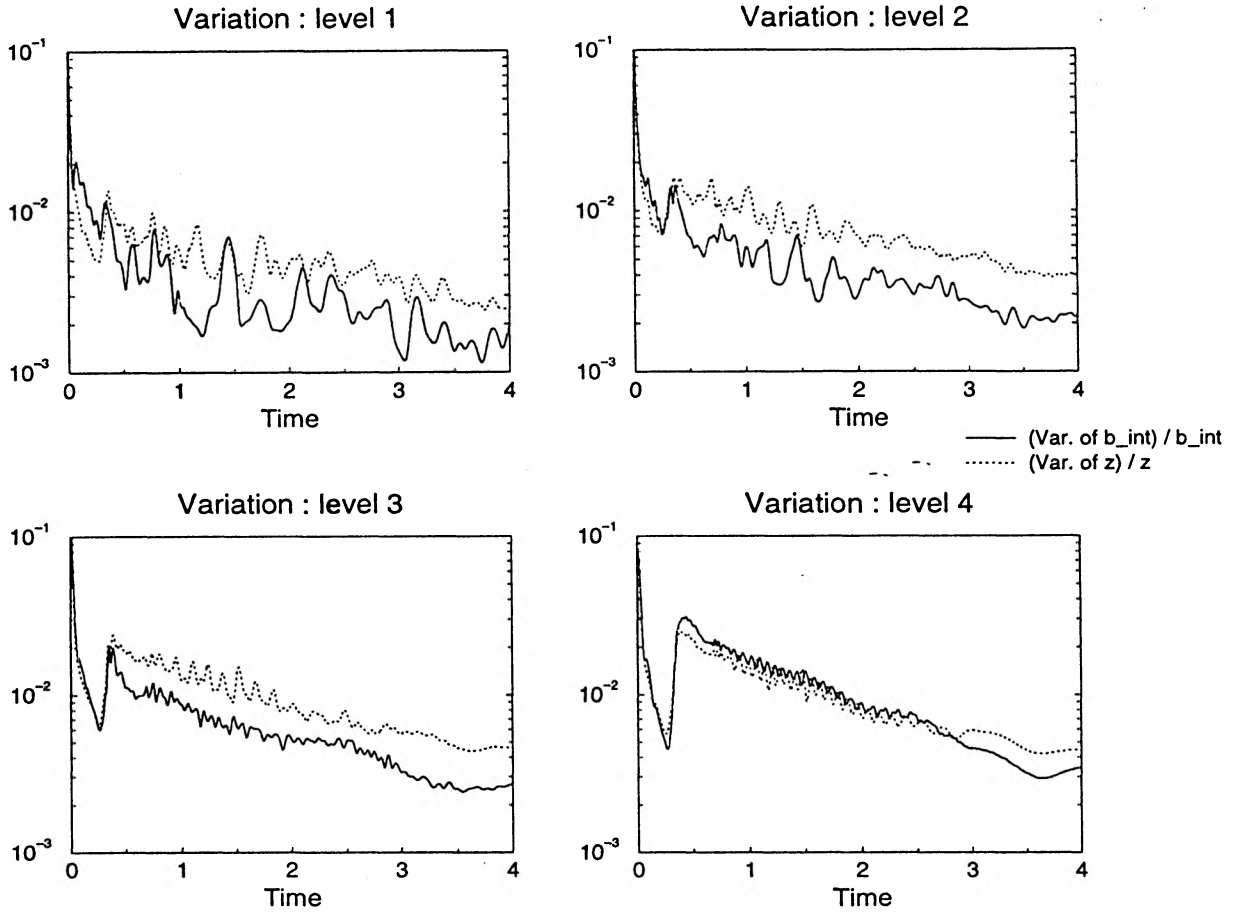


Figure 1.14: Correlation between $\frac{\Delta t \cdot |b_{int}^{\cdot}(y_l, z_{l+1}, \phi_{l,j})|_{l^2}}{|b_{int}(y_l, z_{l+1}, \phi_{l,j})|_{l^2}}$ and $\frac{\Delta t \cdot |\dot{z}_k|_{l^2}}{|z_k|_{l^2}}$.

1.6 Numerical results

In this section we report on the numerical results obtained by using the space-time multilevel algorithm presented in the previous sections, and we compare these results with those obtained by using the classical Galerkin method.

Let Ω be the square domain $[0, 1] \times [0, 1]$. We consider a nodal mesh of size $h_d = \frac{1}{256}$ and we follow the evolution of the Burgers flow from $T = 0$ to $T = 4$.

The initial condition is given by :

$$u(x_1, x_2, 0) = \begin{pmatrix} u_1(x_1, x_2, 0) \\ u_2(x_1, x_2, 0) \end{pmatrix} = \begin{pmatrix} \cos(\pi x_1) \cdot \cos(\pi x_2) \\ \frac{x_1 + x_2}{4} \end{pmatrix}.$$

The boundary conditions are the following non homogeneous Dirichlet conditions :

$$g(x_1, x_2) = \begin{pmatrix} g_1(x_1, x_2) \\ g_2(x_1, x_2) \end{pmatrix} = \begin{pmatrix} u_1(x_1, x_2, 0)|_{\partial\Omega} \\ u_2(x_1, x_2, 0)|_{\partial\Omega} \end{pmatrix},$$

and we force the Burgers flow by a time independent external source f term which acts on only some low frequency components of the velocity field :

$$f(x_1, x_2) = \begin{pmatrix} f_1(x_1, x_2) \\ f_2(x_1, x_2) \end{pmatrix} = \begin{pmatrix} \sin(3\pi x_1) \cdot \sin(4\pi x_2) + 3 \cdot \sin(16\pi x_1) \cdot \sin(8\pi x_2) \\ \sin(4\pi x_1) \cdot \sin(2\pi x_2) + 2 \cdot \sin(10\pi x_1) \cdot \sin(16\pi x_2) \end{pmatrix}.$$

The results are presented for the Reynolds number $Re = 500$, considering at most $d = 4$ hierarchical levels ($h_0 = \frac{1}{16}$ is the width of the coarsest mesh $\mathcal{T}_0$). We choose a time step Δt equal to 10^{-3} , so that the numerical scheme is found to verify the C.F.L. stability condition :

$$|u|_\infty \cdot \frac{\Delta t}{h_d} < c$$

where c is a constant of the order of the unity.

We recall that we actually use the Crank-Nicholson scheme for the linear terms and the Euler forward scheme for the nonlinear terms (this scheme is globally a first-order scheme). Finally we choose the parameter ϵ of the same order as the time step discretization, namely $\epsilon = 10^{-3}$.

Remark 4 *The nonlinear Galerkin method was introduced in the context of numerical simulations of turbulent flows on large intervals of time. Our spatial and temporal multilevel algorithm is implemented with success in the case of solutions whose long term behaviour is not merely the convergence to a steady state.*

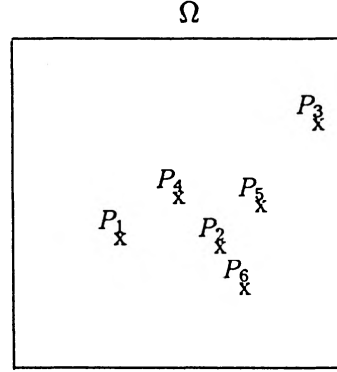
But the solution of the Burgers problem, when we take the right-hand side f equal to zero, converges very fast to a stationary solution depending on the initial condition $u_0(x_1, x_2)$ (see [1] and [12]). In order to obtain a Burgers flow sufficiently perturbed, we have imposed a time independent source f . The external force f puts sufficient energy to keep an almost turbulent flow, during a time interval long enough. This forcing is a substitute for the energy transfer phenomena, displayed into the energy spectrum, from the low frequencies towards the high frequencies. These phenomena are an important characteristic of the Navier-Stokes flow, but they are probably absent in the Burgers flow.

Figures 1.17 and 1.18 show the evolution of the velocity field $u(x_1, x_2, t)$ for $t = 1.0, \dots, 4.0$. We notice a deep gradient, relatively to the first component of the solution, into the lower part of the domain Ω . At $t = 4.0$ we stop the simulation ; the flow has not yet converged to the steady state, but it has almost totally lost its turbulent behaviour.

Figures 1.19 and 1.20 represent the time evolution of the characteristic numbers determining the structure of the V-cycles realized by our space-time multilevel algorithm. Figure 1.19 exhibits the time evolution of the levels in the V-cycles (we recall that the upward phase of a V-cycle go up until the nodal level that is $d = 4$ in our simulation). From Figure 1.20 we can determine the characteristic time length $\tau(t_0)$ of a V-cycle series. More precisely, the time iterations are represented on the x-axis and the depth p_V of a V-cycle on the y-axis. The quantity p_V is computed before the beginning of the series and it is the same during the whole interval $\tau(t_0)$, as far as the estimates (1.21)

P_1	(0.320, 0.394)	P_2	(0.605, 0.375)
P_3	(0.929, 0.745)	P_4	(0.488, 0.492)
P_5	(0.746, 0.488)	P_6	(0.683, 0.242)

Table 1.2: Coordinates of nodes retained in the computational domain.

Figure 1.15: Localization of the nodes $P_1, \dots, P_6$.

or (1.36) are verified. We can observe that, except during a transient time interval after the initial period, $\tau(t_0)$, and then n_V , are rather large. Hence this graph confirms that several V-cycles are realized with no update of the nonlinear terms b_{int} . Concluding, at the end of a whole series the process begins again on the nodal grid, with $p_V = 0$.

Remark 5 We described briefly in the previous section 1.5.2.1 the structure of a V-cycle in the case of a pseudo-spectral discretization. A V-cycle consists of $2(i_2 - i_1) + 1$ time iterations between two levels of discretization N_{i_1} and N_{i_2} , $N_{i_1} < N_{i_2} < N$. In this case one can choose several levels between N_{i_1} and N_{i_2} . This is an important difference with the finite elements case where, for a quasi-structured mesh, at each projection on a coarser mesh we neglect 3/4 of values of the solution.

To perform the comparison between the multilevel algorithm and the classical Galerkin method, we have retained six different points of the computational domain. Their coordinates are written in Table 1.2 and their position is plotted on Figure 1.15. They are vertices of the nodal mesh $\mathcal{T}_d$ (i.e. $P_i \in \Sigma_d$, for $i = 1, \dots, 6$) and some of them are close to the regions where the strong gradient of the solution appears. We have stored the value of the horizontal component of the velocity $u_1(x_1, x_2, t)$ at these points during the time evolution (see the top of Figures 1.21 and 1.22). At the bottom of the same figures we have plotted the quantity $u_1^{CG}(P_i) - u_1^{NLG}(P_i)$ which represents the difference between the first component of the velocity computed with the classical method (marked CG) and the component approximated with the multilevel method (marked NLG).

The difference between the trajectories obtained with the two different algorithms exceeds sometimes the tolerance ϵ . However, when we perform the series of V-cycles,

the error accumulated along the iterations is not so large as cause the blow up or the divergence of the approximate solution. Indeed the simulation of the flow obtained with the multi-scale process is satisfactory : the trajectories of the multilevel method remain close to those obtained with the classical Galerkin method.

Temps	$ u_1^{CG} - u_1^{NLG} _{l^2}$	$ u_2^{CG} - u_2^{NLG} _{l^2}$	$ u_1^{CG} - u_1^{NLG} _{l^\infty}$	$ u_2^{CG} - u_2^{NLG} _{l^\infty}$
1.0	0.00304	0.00187	0.07633	0.02263
2.0	0.00239	0.00159	0.05672	0.01489
3.0	0.00264	0.00158	0.05556	0.02255
4.0	0.00387	0.00299	0.09615	0.09304

Table 1.3: l^2 and l^∞ norm of the vector components of $u^{CG} - u^{NLG}$.

Table 1.3 gives, for each component of the velocity, the l^2 and l^∞ norms of the difference between the solutions computed with the CG and NLG methods. We can see that the vector field $u(x, t)$ is globally well approximated, indeed the l^2 norm of the difference is only lightly greater than ϵ . On the contrary the l^∞ norm of this difference is large.

Considering this result, we contemplate to study in the future the local behaviour of the linear and nonlinear terms that we would like to freeze during the time iterations. Up to now we have kept into account only the l^2 norm of these terms, but this norm is an average measure of these terms and it tends to mask the steep gradients localized upon a small size of the computational domain. We can take into account this local analysis by using some adaptive hierarchical meshes in order to sufficiently discretized the domain Ω in regions of steep gradients.

Finally, we show in Figure 1.16 the time evolution of the CPU time required by CG and NLG methods. To achieve the comparison between the different algorithms, we want to note the significant fact that the multilevel method requires nearly 30% less CPU time than the classical Galerkin method. The saving is given by the ratio

$$\frac{T_{CG} - T_{NLG}}{T_{CG}}$$

where T_{CG} is the CPU time required by the classical Galerkin method and T_{NLG} the CPU time required by the multilevel method.

Remark 6 *In order to justify the mesh size chosen for the nodal grid, we compare two approximate solutions of the Burgers problem determined by the classical Galerkin method ; the first solution is obtained using a space step $h_d = \frac{1}{256}$ and the second solution is computed using $h_d = \frac{1}{128}$ (obviously the Reynolds number and the time step remain unchanged). The approximate solution with mesh size $h_d = \frac{1}{128}$ is not totally satisfactory because of the presence of some fluctuations near the zone where the solution exhibits steep gradients. These fluctuations appear when we consider a too large mesh size and they disappear when we consider $h_d = \frac{1}{256}$.*

Moreover we compute the difference, at time $t = 2.0$ and $t = 4.0$, between $u_1(x_1, x_2, t)$ approximated with CG and a mesh size $h_d = \frac{1}{256}$ and $u_1(x_1, x_2, t)$ approximated with the same method but with a mesh size twice greater : $h_d = \frac{1}{128}$. We show this difference on the top of Figures 1.25 and 1.26. On the bottom of the same figures we represent the difference between u_1 computed with the CG method and u_1 computed with the multilevel method, again at $t = 2.0$ and $t = 4.0$, using the same space step $h_d = \frac{1}{256}$. We want to note the significant fact that the difference between u_1 computed with CG and NLG method is smaller than the difference between u_1 computed with the CG method and different mesh sizes. One must observe that these quantities are plotted with different extrema on the z-axis.

Finally, Figures 1.23 and 1.24 exhibit the evolution of the difference between u_1^{CG} ($h_d = \frac{1}{256}$) and u_1^{CG} ($h_d = \frac{1}{128}$), difference established at the points $P_1, \dots, P_6$ (see Table 1.2). These figures can be compared with the Figures 1.21 and 1.22, which give the difference $u_1^{CG}(P_i) - u_1^{NLG}(P_i)$ computed with the same mesh size $h_d = \frac{1}{256}$.

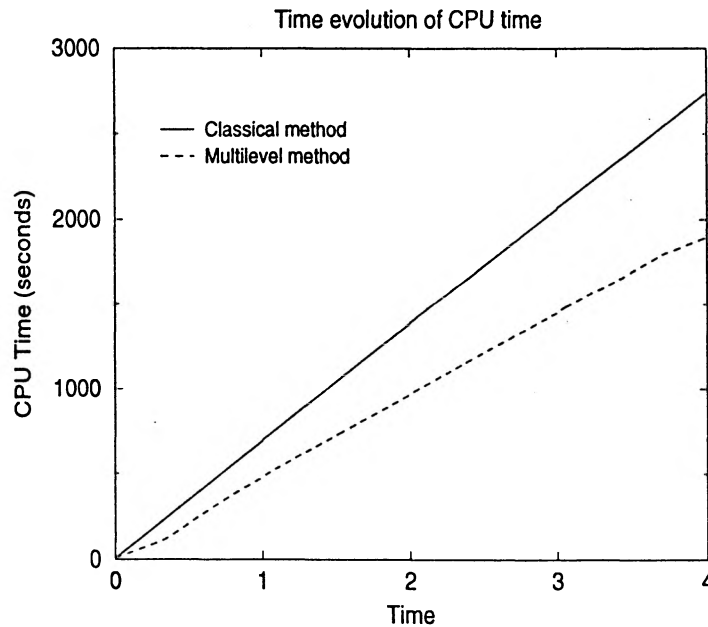


Figure 1.16: Time evolution of CPU time.

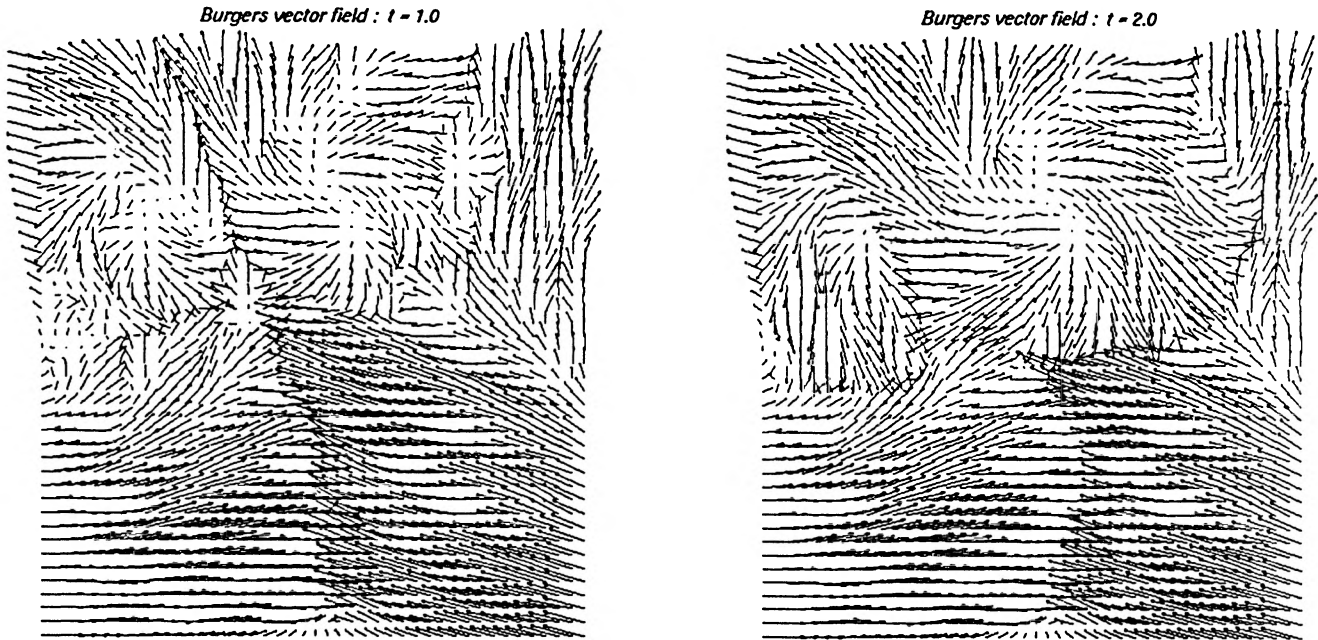


Figure 1.17: Velocity fields at $t = 1.0$ and $t = 2.0$; $\text{Re} = 500$ and $h_d = \frac{1}{256}$.

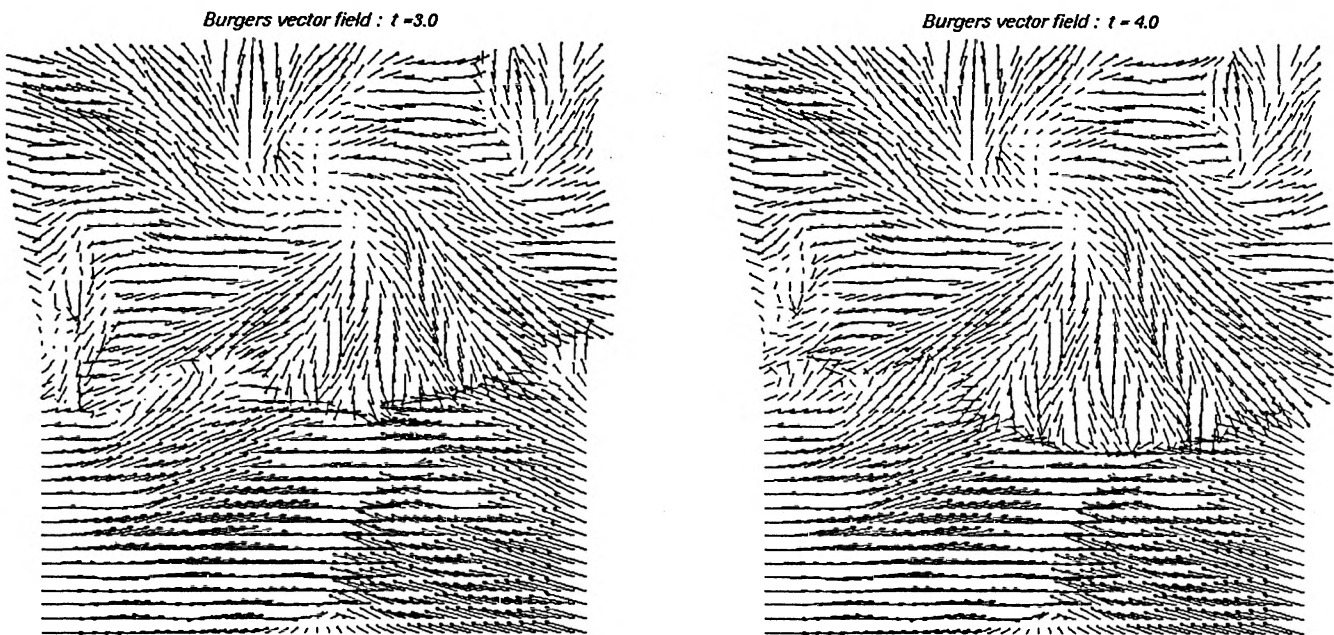


Figure 1.18: Velocity fields at $t = 3.0$ and $t = 4.0$; $\text{Re} = 500$ and $h_d = \frac{1}{256}$.

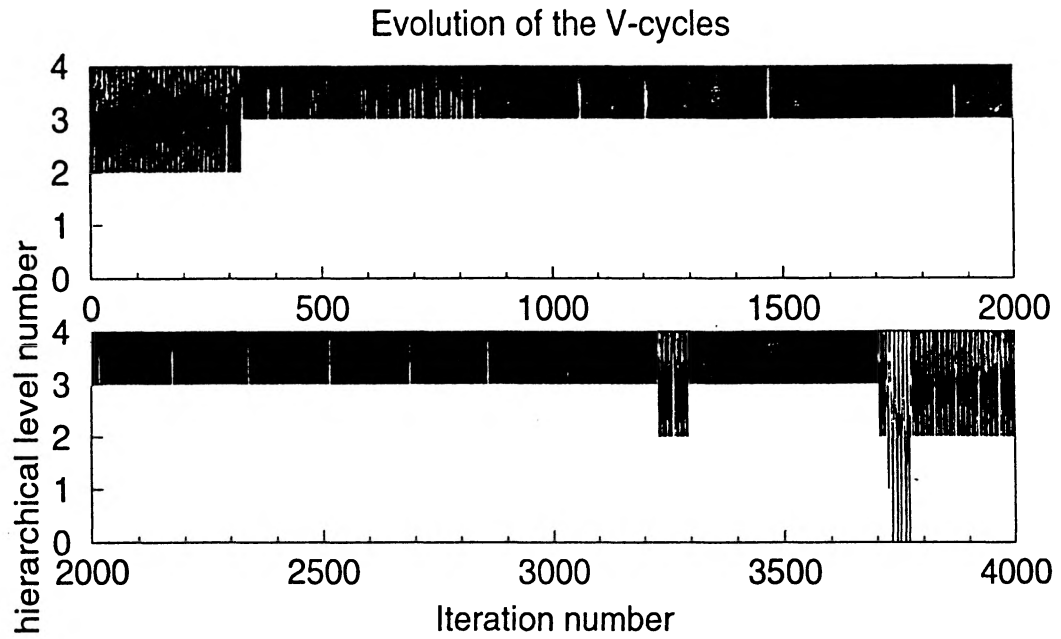
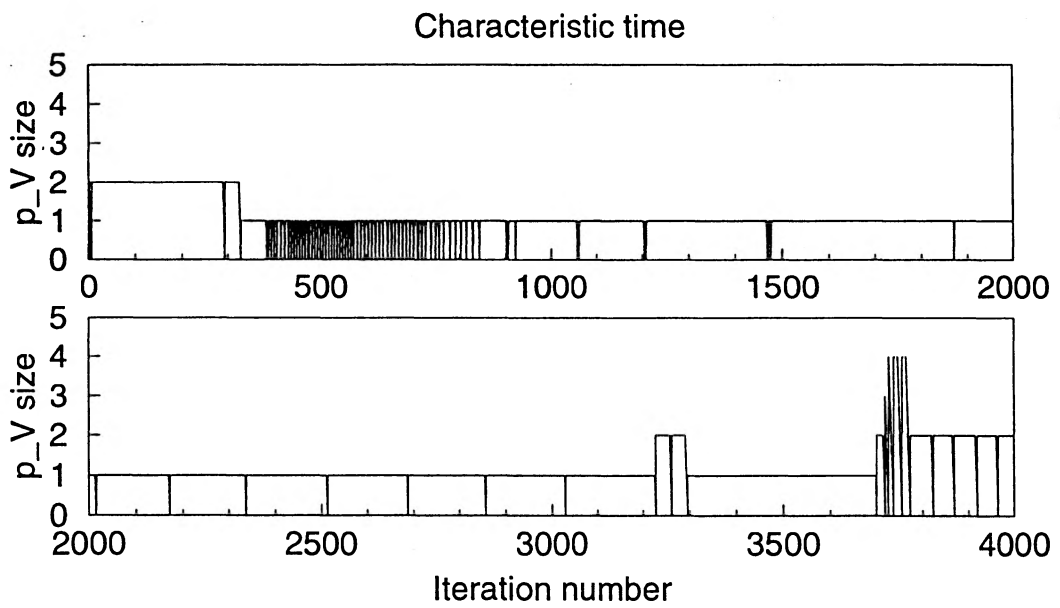


Figure 1.19: Time evolution of the levels in the V-cycles.


 Figure 1.20: Characteristic time $\tau(t_0)$ of the V-cycles.

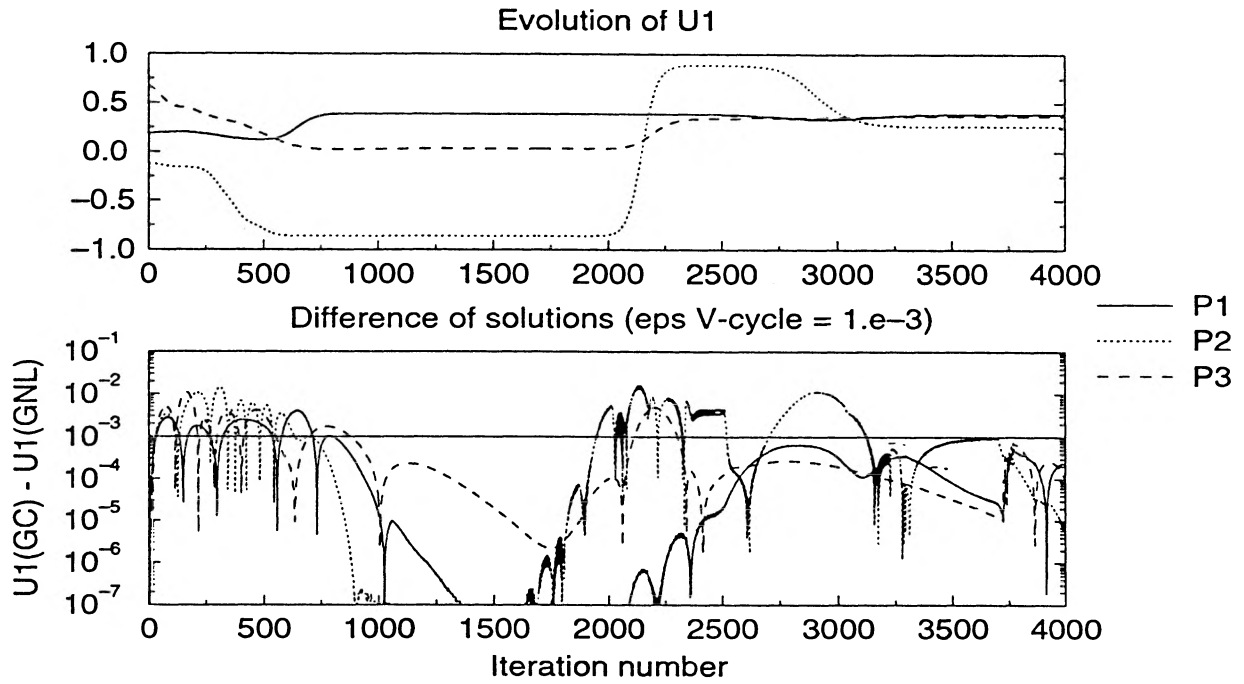


Figure 1.21: Time evolution of the difference $u_1^{CG} - u_1^{NLG}$ at nodes P_1, P_2, P_3 .

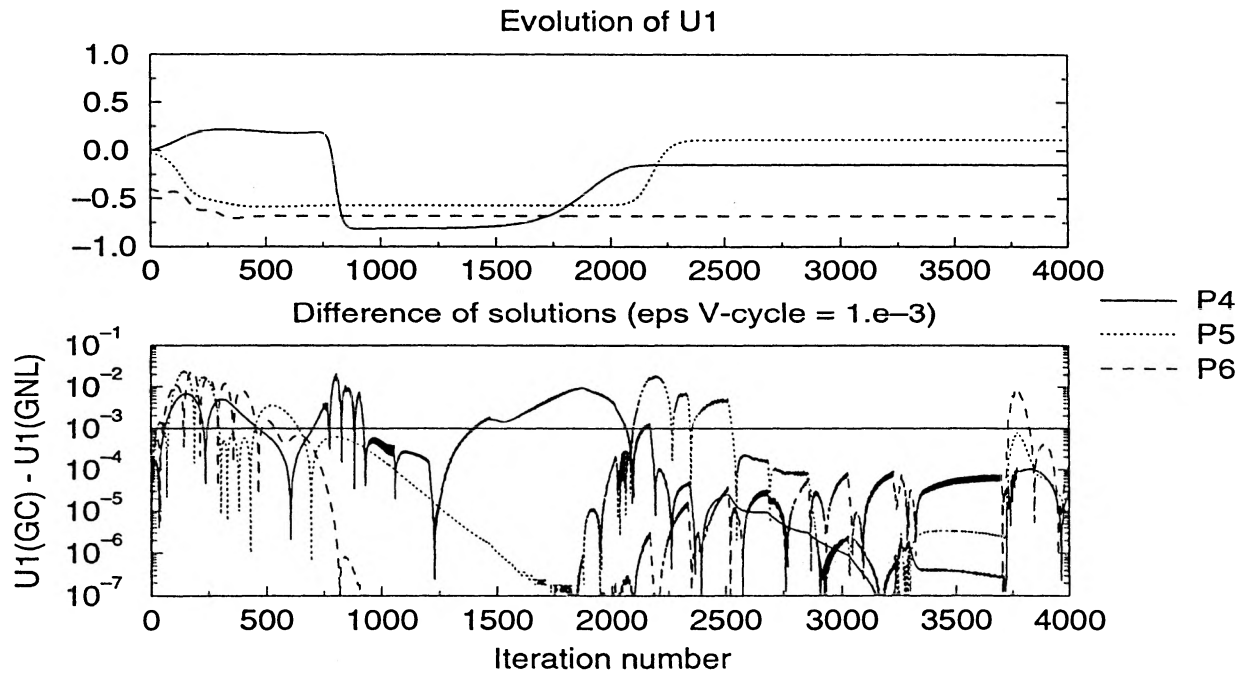


Figure 1.22: Time evolution of the difference $u_1^{CG} - u_1^{NLG}$ at nodes P_4, P_5, P_6 .

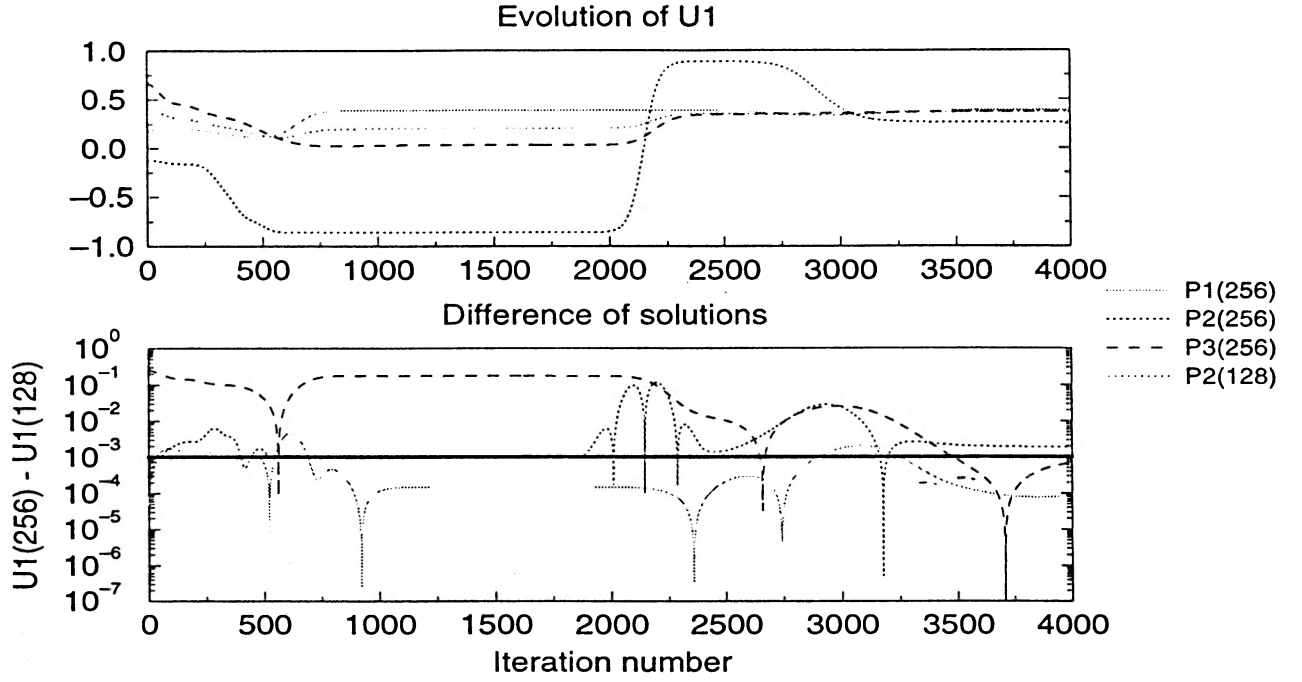


Figure 1.23: Time evolution of the difference between $u_1^{CG}(h_d = \frac{1}{256})$ and $u_1^{CG}(h_d = \frac{1}{128})$ at nodes P_1, P_2, P_3 .

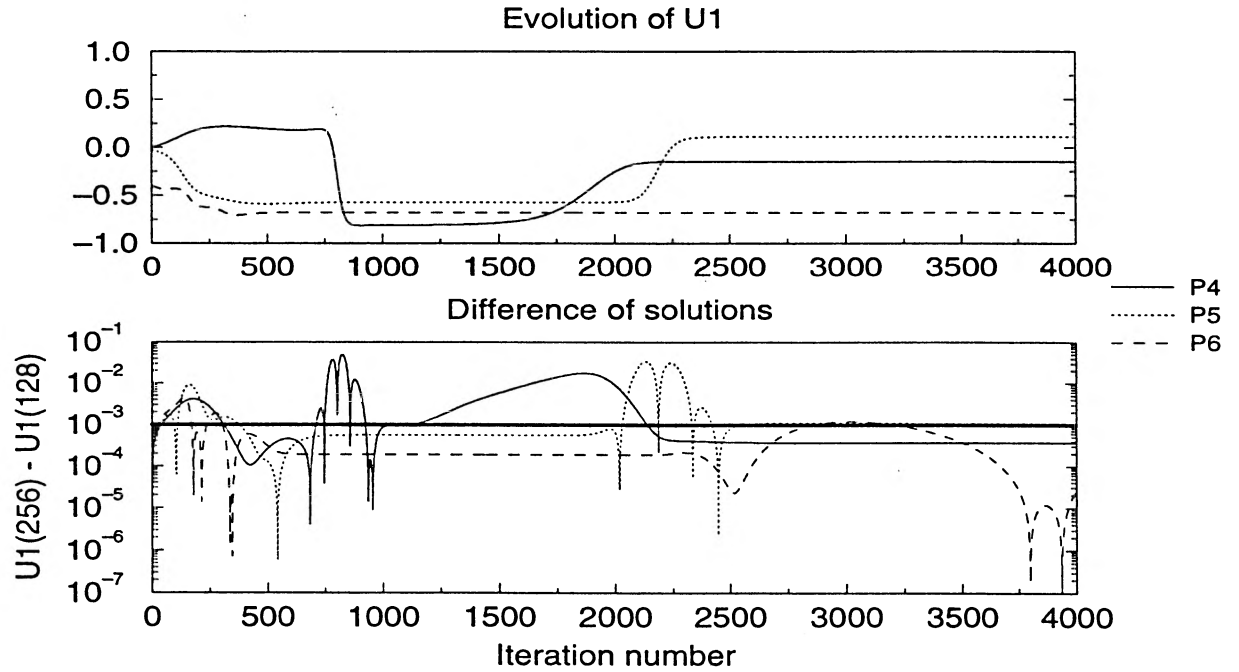
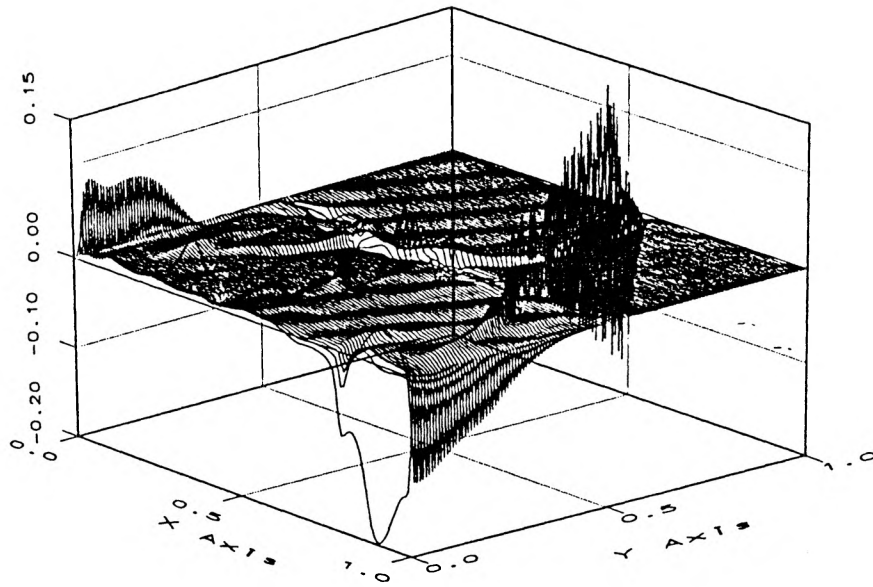


Figure 1.24: Time evolution of the difference between $u_1^{CG}(h_d = \frac{1}{256})$ and $u_1^{CG}(h_d = \frac{1}{128})$ at nodes P_4, P_5, P_6 .

Diff. Sol. GC(129) - GC(257) ; T = 2.0



Diff. Sol. GC(257) - GNL(257) ; T = 2.0

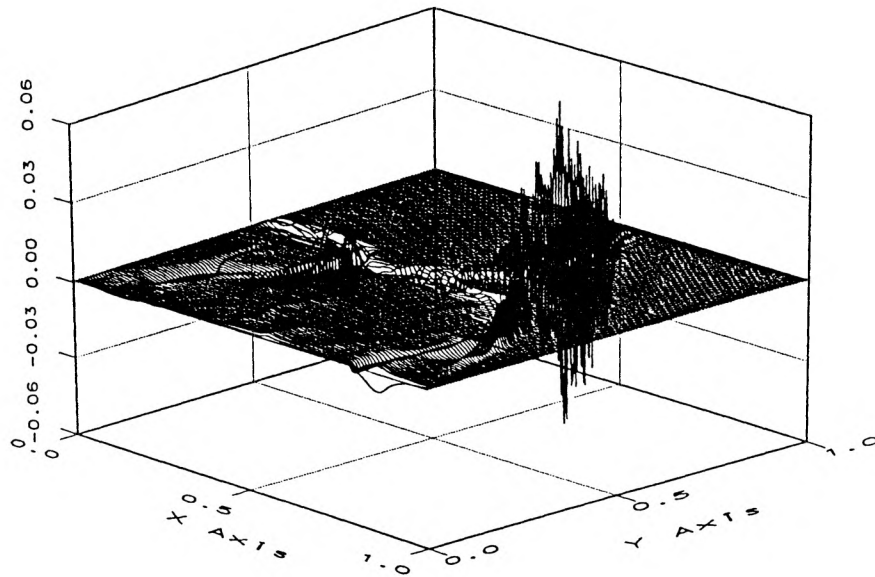
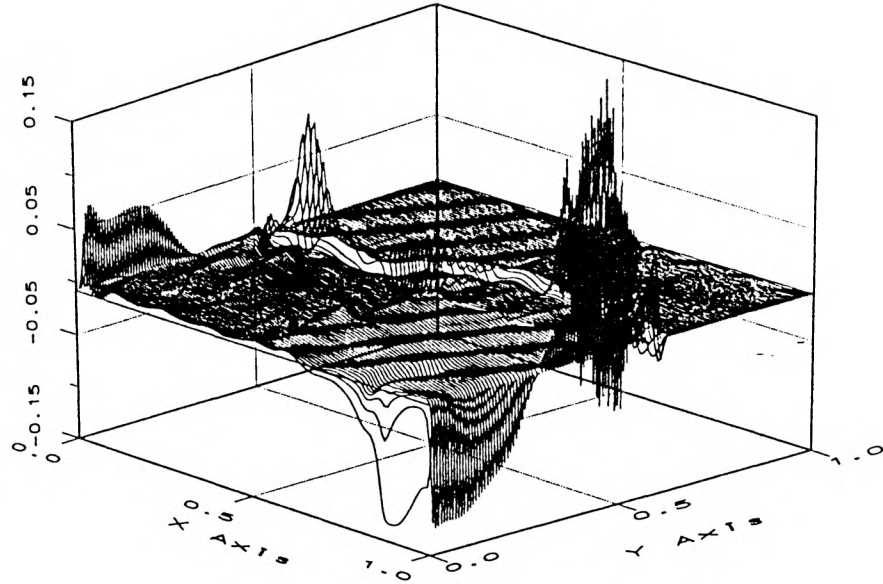


Figure 1.25: On top : difference between $u_1(t = 2)$ computed with CG method and $h_d = \frac{1}{256}$ or $h_d = \frac{1}{128}$. At the bottom : difference between $u_1(t = 2)$ computed with CG and NLG method ; $h_d = \frac{1}{256}$ is the same.

Diff. Sol. GC(129) - GC(257) ; T = 4.0



Diff. Sol. GC(257) - GNL(257) ; T = 4.0

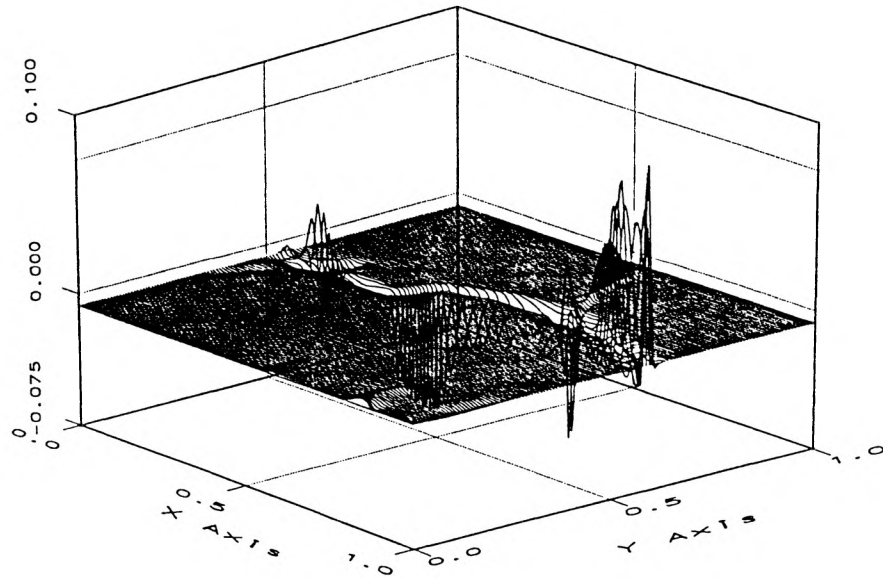


Figure 1.26: On top : difference between $u_1(t = 4)$ computed with CG method and $h_d = \frac{1}{256}$ or $h_d = \frac{1}{128}$. At the bottom : difference between $u_1(t = 4)$ computed with CG and NLG method ; $h_d = \frac{1}{256}$ is the same.

Concluding Remarks

In this article we have proposed a space-time multilevel method in the context of finite elements. This method treats differently the large and small components of the flow which is solution of the two-dimensional Burgers problem. This decomposition is induced by the utilization of the hierarchical bases. We have obtained with this multilevel algorithm the first results of an application of the nonlinear Galerkin method in the context of a finite element discretization using several hierarchical levels. In the previous numerical applications of the original schemes proposed by M. Marion and R. Temam, J. Laminie *et al.* [14], [15] observe that the approximations made in the y -equations and in the z -equation are not always justified from a numerical point of view. Although the size of certain coupled terms can be small, they can not be neglected, especially when the solution exhibits steep gradients in the computational domain. The new suggestion here is to consider instead the time variation of the linear and nonlinear terms. The subsequent multilevel algorithm, which freezes the coupled terms during suitable time intervals, converges with a good efficiency, producing a significant saving of computing time. Several numerical estimates of terms which appear in the equations have been performed before the conception and the realization of the multilevel algorithm. Moreover we have derived mathematical estimates for the parameters involved in this dynamical procedure. In the case of flows which present deep gradients, it will be possible to envisage a local analysis of suitably frozen terms ; this analysis can be obtained using adaptive hierarchical meshes. In the future our efforts will be concentrated on the criteria for determining the structure of the V-cycles as well as finding better correlations concerning the nonlinear interaction terms, in order to decrease the perturbations introduced by the multi-scale procedure. Other efforts will be concentrated on the implementation of a second-order scheme for the time discretization of the linear and nonlinear terms, and also on the implementation of our multilevel algorithm in the case of the Navier-Stokes problem.

Bibliography

- [1] C. CALGARO
Méthodes multi-résolution auto-adaptatives en éléments finis : application aux équations de la mécanique des fluides
(1996) *Thèse Univ. Paris-Orsay : to appear*
- [2] P. CONSTANTIN, C. FOIAS, O. MANLEY and R. TEMAM
Determining modes and fractal dimension of turbulent flows
(1985) *J. Fluid Mech.*, vol 150, pp 427-440
- [3] A. DEBUSSCHE, T. DUBOIS and R. TEMAM
The nonlinear Galerkin method: a multi-scale method applied to the simulation of homogeneous turbulent flow
(1995) *Theoret. and Comput. Fluid Dynamics*, vol 7, pp. 279-315
- [4] T. DUBOIS
Simulation numérique d'écoulements homogènes et non-homogènes par des méthodes multi-résolution
(1993) *Thèse Univ. Paris-Orsay*
- [5] T. DUBOIS, F. JAUBERTEAU and R. TEMAM
Solution of the incompressible Navier-Stokes equations by the Nonlinear Galerkin Method
(1993) *J. Scient. Computing*, vol 8, pp 167-194
- [6] T. DUBOIS, F. JAUBERTEAU and R. TEMAM
Dynamical multilevel methods in turbulence simulation
Computational Fluid Dynamics Review : to appear
- [7] C. FOIAS, M.S. JOLLY, I.G. KEVREKIDIS, G.R. SELL and E.S. TITI
On the computation of inertial manifolds
(1988) *Physics Letters*, vol A 131, pp 433-436
- [8] C. FOIAS and R. TEMAM
Some analytic and geometric properties of the solutions of the evolution Navier-Stokes equations
(1979) *J. Math. Pures et Appl.*, vol 58, pp 339-368

- [9] C. FOIAS, G.R. SELL and R. TEMAM
Inertial manifolds for nonlinear evolutionary equations
(1988) *J. Diff Equ.*, vol 73, pp 309-353
- [10] C. FOIAS, O. MANLEY and R. TEMAM
Sur l'interaction des petits et grands tourbillons dans les écoulements turbulents
(1987) *C.R. Acad. Sci. Paris Série I*, 305, pp 497-500
- [11] C. FOIAS, O. MANLEY and R. TEMAM
On the interaction of small and large eddies in two dimensional turbulent flows
(1988) *Math. Model. and Numer. Anal.*, vol 22, pp 93-114
- [12] P.C. JAIN and D. N. HOLLA
Numerical solutions of coupled Burgers' equations
(1978) *Int. J. Nonlinear Mech.*, vol 13, pp 213-222
- [13] F. JAUBERTEAU
Résolution numérique des équations de Navier-Stokes instationnaires par méthodes spectrales. Méthode de Galerkin non linéaire
(1990) *Thèse Univ. Paris-Orsay*
- [14] J. LAMINIE, F. PASCAL and R. TEMAM
Implementation of finite element nonlinear Galerkin methods using hierarchical bases
(1993) *J. Comp. Mech.*, vol 11, pp 384-407
- [15] J. LAMINIE, F. PASCAL and R. TEMAM
Implementation and numerical analysis of the nonlinear Galerkin methods with finite elements discretization
(1994) *Applied Numerical Mathematics*, vol 15, pp 219-246
- [16] M. MARION and R. TEMAM
Nonlinear Galerkin Methods
(1989) *SIAM J. Math. Anal.*, vol 26, pp 1139-1157
- [17] M. MARION and R. TEMAM
Nonlinear Galerkin Methods : the finite elements case
(1990) *Num. Math.*, vol 57, pp 205-226
- [18] F. PASCAL
Méthodes de Galerkin non linéaires en discrétisation par éléments finis et pseudo-spectrale. Application à la mécanique des fluides
(1992) *Thèse Univ. Paris-Orsay*
- [19] R. TEMAM
Sur l'approximation des équations de Navier-Stokes
(1966) *C.R. Acad. Sci. Paris, Ser. A*, 262, pp 219-221

-
- [20] R. TEMAM
Navier-Stokes equations
(1984) North Holland ; Amsterdam
- [21] R. TEMAM
Infinite dimensional dynamical systems in mechanics and physics
(1988) Applied Mathematical Sciences 68, Springer Verlag ; New-York
- [22] R. TEMAM
Inertial manifolds and multigrid methods,
(1990) *SIAM J. Math. Anal.*, vol 21, pp 154-178
- [23] H. YSERENTANT
On the multi-level splitting of finite element spaces
(1986) *Num. Math.*, vol 49, pp 379-412
- [24] O.C. ZIENCIEWICZ, D.W. KELLY, J. GAGO, I. BABUSKA
Hierarchical finite element approaches, error estimates and adaptive refinement
The mathematics of finite elements and applications IV (J.R. Whiteman ed.)
(1982) Uxbridge , London - Academics Press

Partie II

Résolution des équations de NAVIER-STOKES

Chapitre 1

Résolution des équations de Stokes et de Navier-Stokes

Introduction

Dans ce chapitre, après avoir introduit les équations de Stokes et de Navier-Stokes, nous précisons les motivations qui nous amènent à étudier ces problèmes sous leur forme pénalisée. En effet nous cherchons à simplifier le problème de départ, en évitant surtout les difficultés dues à la contrainte d'incompressibilité, afin de pouvoir étendre à ces équations la méthode multi-résolution de *type incrémental* mise au point dans la première partie de cette thèse.

Une étude très circonstanciée du problème de Stokes précède la résolution numérique des équations de Navier-Stokes. Cette étude débute par la présentation de la discrétisation en espace du problème de Stokes pénalisé, considéré sous sa forme variationnelle. Deux choix d'éléments finis mixtes, i.e. les éléments 4P1-P0 et 4P1-P1, seront exposés en parallèle, bien que dans nos applications l'un ait suivi l'autre, à la suite de difficultés rencontrés dans l'approximation des solutions du problème de Stokes généralisé.

La recherche d'une méthode performante, afin d'effectuer la résolution du système linéaire auquel nous parvenons, passe par différents essais cherchant à surmonter le défaut principal de la méthode de pénalisation, soit le très mauvais conditionnement de la matrice à inverser. Plusieurs préconditionnements classiques des méthodes itératives du type gradient conjugué ont été testés mais nous montrons, en les passant en revue, qu'aucun préconditionnement classique ne donne de résultats satisfaisants afin de déterminer un solveur efficace du problème linéaire considéré. Après une brève étude des valeurs propres de la matrice à inverser, nous optons pour une résolution directe du système linéaire. La méthode directe, performante en temps de calcul, limitera la taille de discrétisation du problème, à cause de l'encombrement mémoire dû au stockage de la matrice factorisée. Une réduction du stockage sera obtenue en réordonnant les coefficients de la matrice afin de diminuer sa largeur de bande.

Les résultats numériques pour les équations de Stokes sont ensuite exposés, en détaillant en même temps les problèmes révélés, qui se manifestent par des perturbations

dans le champ de vitesse calculé. L'étape finale consiste à résoudre les équations de Navier-Stokes pénalisées. Toutefois, nos simulations numériques se restreindront à des écoulements qui évoluent vers des solutions stationnaires. En conclusion, nous montrons que la correspondance entre nos résultats et ceux obtenus par d'autres auteurs est satisfaisante.

1.1 Les équations considérées

1.1.1 Les équations en variables primitives

Nous considérons un ouvert Ω de $\mathbf{R}^2$ de frontière $\partial\Omega$ régulière (par exemple Lipschitzienne). Dans le domaine Ω nous cherchons les caractéristiques de l'écoulement d'un fluide newtonien. Soient donc :

$u(x, t) = (u_1(x, t), u_2(x, t))$ la vitesse de l'écoulement au point $x \in \Omega$, à l'instant t ;

$p(x, t)$ la pression hydrostatique en x , à l'instant t ;

$\rho(x, t)$ la densité du fluide.

Les équations régissant l'écoulement sont les suivantes :

- *Conservation de la quantité de mouvement :*

$$\rho \left(\frac{\partial u_i}{\partial t} + u_j \frac{\partial u_i}{\partial x_j} \right) - \frac{\partial \sigma_{ij}}{\partial x_j} = \rho f_i \quad \text{pour } i = 1, 2, \quad (1.1)$$

la matrice σ étant le tenseur des contraintes et le vecteur $f(x, t) = (f_1(x, t), f_2(x, t))$ la densité des forces extérieures. Dans cette expression, comme dans les suivantes, quand un indice est doublé, cela signifie qu'il faut sommer par rapport à cet indice.

- *Conservation de la masse :*

$$\frac{\partial \rho}{\partial t} + \operatorname{div}(\rho u) = 0. \quad (1.2)$$

- *Loi de comportement des fluides newtoniens :*

$$\sigma_{ij} = -p \delta_{ij} + \mu \left(\frac{\partial u_i}{\partial x_j} + \frac{\partial u_j}{\partial x_i} \right), \quad (1.3)$$

μ étant la viscosité dynamique.

Si on considère un fluide visqueux incompressible, c'est-à-dire que l'on suppose que la densité ρ varie peu, l'équation (1.2) se réduit à celle d'incompressibilité

$$\operatorname{div} u = 0 \quad (1.4)$$

et le tenseur des contraintes σ vérifie :

$$\frac{\partial \sigma_{ij}}{\partial x_j} = -\frac{\partial p}{\partial x_i} + \mu \frac{\partial}{\partial x_j} \left(\frac{\partial u_i}{\partial x_j} \right) .$$

Dans ce cas on nomme viscosité cinématique la quantité $\nu = \mu/\rho$. La conservation de la quantité de mouvement se réécrit sous la forme classique suivante :

$$\frac{\partial u}{\partial t} + (u \cdot \nabla)u - \nu \Delta u + \nabla p = f , \quad (1.5)$$

où p dénote dorénavant la pression réduite, i.e. p/ρ .

On note L une longueur caractéristique du domaine (par exemple le diamètre de Ω), U une vitesse caractéristique de l'écoulement et T un temps caractéristique (*a priori* $T = L/U$). Si on pose :

$$u = \frac{u}{U}, \quad x = \frac{x}{L}, \quad t = \frac{t}{T}, \quad p = \frac{p}{U^2} \quad \text{et} \quad f = \frac{fL}{U^2} ,$$

alors le système (1.4)-(1.5) s'écrit sous la forme adimensionnée suivante :

$$\begin{cases} \frac{\partial u}{\partial t} - \frac{1}{Re} \Delta u + (u \cdot \nabla)u + \nabla p = f & \text{dans } \Omega \times \mathbf{R}^+ , \\ \nabla \cdot u = 0 & \text{dans } \Omega \times \mathbf{R}^+ , \end{cases} \quad (1.6)$$

où le nombre de Reynolds Re , qui joue le rôle de paramètre de bifurcation dans le comportement de la solution, s'écrit sous la forme

$$Re = \frac{UL}{\nu} .$$

Les équations de Navier-Stokes (1.6) sont complétées par

- une condition initiale

$$u(x, 0) = u_0(x) \quad \forall x \in \Omega , \quad (1.7)$$

la vitesse initiale u_0 étant à divergence nulle ;

- des conditions aux limites qui varient selon le problème que l'on veut résoudre. Dans notre étude numérique nous considérons seulement des conditions aux limites de type Dirichlet non homogènes

$$u(x, t) = g(x, t) \quad \text{sur } \partial\Omega , \quad (1.8)$$

g vérifiant la condition de flux nul imposée par l'incompressibilité :

$$\int_{\partial\Omega} g \cdot n \, d\gamma = 0 ,$$

n étant la normale extérieure à Ω .

Lorsqu'on envisage d'étudier des écoulements de fluides très visqueux, c'est-à-dire à nombre de Reynolds très faible ($Re \ll 1$), on peut approcher les équations de Navier-Stokes par celles de Stokes (stationnaires) :

$$\begin{cases} -\frac{1}{Re}\Delta u + \nabla p = f & \text{dans } \Omega \\ \nabla \cdot u = 0 & \text{dans } \Omega . \end{cases} \quad (1.9)$$

De façon plus générale, on parle des équations de Stokes généralisées lorsqu'on considère :

$$\begin{cases} \alpha u - \frac{1}{Re}\Delta u + \nabla p = f & \text{dans } \Omega \\ \nabla \cdot u = 0 & \text{dans } \Omega . \end{cases} \quad (1.10)$$

avec le paramètre $\alpha > 0$. De nombreux schémas en temps qui discrétisent les équations de Navier-Stokes font apparaître le problème (1.10) avec $\alpha = \frac{1}{\Delta t}$, Δt étant le pas de discrétisation en temps. Il s'agit par exemple des schémas classiques aux différences finies (Euler - Euler ou Crank-Nicholson - Adams-Bashforth), ou des schémas de splitting, appelés également θ -schémas, ou encore des méthodes à pas fractionnaires proposées par Temam (cf. [34] et [35]).

1.1.2 Cadre fonctionnel

On considère Ω un ouvert borné de $\mathbf{R}^2$ dont la frontière $\partial\Omega$ est localement Lipschitzienne. On suppose que $L^2(\Omega)$ est muni du produit scalaire usuel et de la norme usuelle :

$$(u, v) = \int_{\Omega} u(x)v(x)dx \quad \forall u, v \in L^2(\Omega) ,$$

$$|u| = (u, u)^{1/2} \quad \forall u \in L^2(\Omega) ,$$

et que l'espace de Sobolev $H^1(\Omega)$ est muni du produit scalaire et de la norme :

$$(u, v)_1 = (u, v) + (\nabla u, \nabla v) \quad \forall u, v \in H^1(\Omega) ,$$

$$|u|_1 = (u, u)_1^{1/2} \quad \forall u \in H^1(\Omega) .$$

Soit $H_0^1(\Omega)$ sous-espace de $H^1(\Omega)$ des fonctions qui s'annulent sur la frontière $\partial\Omega$. L'inégalité de Poincaré implique que $H_0^1(\Omega)$ est un espace d'Hilbert pour le produit scalaire :

$$((u, v)) = (\nabla u, \nabla v) \quad \forall u, v \in H_0^1(\Omega) .$$

La norme correspondante est notée :

$$\|u\| = ((u, u))^{1/2} \quad \forall u \in H_0^1(\Omega) ,$$

$\|\cdot\|$ étant une norme équivalente à la norme $|\cdot|_1$.

Lorsqu'on suppose les conditions aux limites de type Dirichlet non homogènes $u = g$ sur $\partial\Omega$, avec g fonction vérifiant la condition de flux nul, on pose $H_g^1(\Omega)$ le sous-espace de $H^1(\Omega)$ tel que :

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

On suppose enfin que $L^2(\Omega)/\mathbb{R}$, espace des fonctions de $L^2(\Omega)$ déterminées à une constante additive près, est muni de la norme quotient suivante :

$$\|v\|_{L^2(\Omega)/\mathbb{R}} = \inf_{c \in \mathbb{R}} |v + c| \quad \forall v \in L^2(\Omega)/\mathbb{R}.$$

1.1.3 Le problème de Stokes : formulations équivalentes

Afin de simplifier les notations, on se restreint pour le moment aux équations linéaires de Stokes stationnaire (1.9) ou généralisé (1.10) (le paramètre α peut être positif ou nul). Si on considère la condition de divergence nulle comme une contrainte linéaire sur la solution u , la pression p apparaît alors comme un multiplicateur de Lagrange.

Ecrivons avant tout la formulation variationnelle associée au problème de Stokes :

$$\left\{ \begin{array}{l} \text{Trouver } (u, p) \in (H_g^1(\Omega))^2 \times L^2(\Omega) \text{ tel que} \\ a(u, v) - b(v, p) = (f, v) \quad \forall v \in (H_0^1(\Omega))^2 \\ b(u, q) = 0 \quad \forall q \in L^2(\Omega), \end{array} \right. \quad (1.11)$$

où les formes bilinéaires continues a et b sont définies de la façon suivante :

$$\begin{aligned} a(u, v) &= \alpha(u, v) + \frac{1}{Re}((u, v)), \\ b(u, q) &= (\operatorname{div} u, q). \end{aligned} \quad (1.12)$$

Remarque 1 *La pression peut être déterminée seulement à une constante additive près i.e. $p \in L^2(\Omega)/\mathbb{R}$.*

Pour $f \in (L^2(\Omega))^2$ et $v \in (H_g^1(\Omega))^2$ considérons la fonctionnelle quadratique J définie par :

$$J(v) = \frac{1}{2}a(v, v) - (f, v).$$

Le problème de Stokes (1.11) est alors équivalent au problème de minimisation avec contraintes :

$$\left\{ \begin{array}{l} \text{Trouver } u \in V = \{v \in (H_g^1(\Omega))^2 \mid \operatorname{div} v = 0 \text{ dans } \Omega\} \text{ tel que} \\ J(u) \leq J(v) \quad \forall v \in V. \end{array} \right. \quad (1.13)$$

L'existence et l'unicité de la solution du problème de Stokes est une conséquence immédiate du théorème de Lax-Milgram, car la forme bilinéaire $a(.,.)$ est coercive sur le sous-espace V (cf. [13]).

Il est naturel d'imposer la contrainte linéaire au moyen d'un multiplicateur de Lagrange, transformant ainsi le problème (1.13) en un problème de point-selle. En définissant, pour $v \in (H_g^1(\Omega))^2$ et $q \in L^2(\Omega)$, le Lagrangien :

$$\mathcal{L}(v, q) = J(v) - b(v, q) = \frac{1}{2}a(v, v) - (f, v) - b(v, q) ,$$

on cherche un couple (u, p) point-selle de $\mathcal{L}$ sur $(H_g^1(\Omega))^2 \times L^2(\Omega)$ tel que

$$\begin{cases} \text{Trouver } (u, p) \in (H_g^1(\Omega))^2 \times L^2(\Omega) \text{ tel que} \\ \mathcal{L}(u, q) \leq \mathcal{L}(u, p) \leq \mathcal{L}(v, p) \quad \forall v \in (H_g^1(\Omega))^2 \text{ et } \forall q \in L^2(\Omega) . \end{cases} \quad (1.14)$$

Un critère d'existence de point-selle est donné par la proposition suivante :

Proposition 1 *La fonction $\mathcal{L}$ à valeurs réelles possède un point-selle (u, p) sur $(H_g^1(\Omega))^2 \times L^2(\Omega)$ si et seulement si*

$$\min_{v \in (H_g^1(\Omega))^2} \max_{q \in L^2(\Omega)} \mathcal{L}(v, q) = \max_{q \in L^2(\Omega)} \min_{v \in (H_g^1(\Omega))^2} \mathcal{L}(v, q)$$

et ce nombre est alors égal à $\mathcal{L}(u, p)$.

Suivant [14], on peut introduire le Lagrangien augmenté $\mathcal{L}_r$ défini, pour tout $r > 0$, par :

$$\begin{aligned} \mathcal{L}_r(v, q) &= \mathcal{L}(v, q) + \frac{r}{2} |\operatorname{div} v|^2 \\ &= J(v) - b(v, q) + \frac{r}{2} |\operatorname{div} v|^2 . \end{aligned} \quad (1.15)$$

Alors tout point-selle de $\mathcal{L}_r$ est point-selle de $\mathcal{L}$ et réciproquement (puisque $r |\operatorname{div} v|^2$ s'annule lorsque la contrainte $\operatorname{div} v = 0$ est satisfaite).

Remarque 2

1. Le fait d'ajouter le terme quadratique $\frac{r}{2} |\operatorname{div} v|^2$ au Lagrangien $\mathcal{L}$ améliore les propriétés de convergence des algorithmes de dualité employés pour le calcul du point-selle (voir [14]).
2. Pour $q = 0$ on a $\mathcal{L}_r(v, 0) = J(v) + \frac{r}{2} |\operatorname{div} v|^2$, soit la fonctionnelle pénalisée relative à la contrainte $\operatorname{div} v = 0$. Considérer le Lagrangien augmenté est donc équivalent, grâce aux arguments classiques de dualité, à l'introduction d'une fonction de pénalisation dans le problème primal (1.13).

1.1.4 Les problèmes de Stokes et Navier-Stokes pénalisés

Plusieurs auteurs, dans le but de discrétiser les problèmes de Stokes ou de Navier-Stokes tout en évitant les difficultés dues à la contrainte d'incompressibilité, ont proposé d'étudier des problèmes pénalisés qui ne vérifient plus exactement la condition $\operatorname{div} u = 0$. Cette méthode de pénalisation a fait l'objet de nombreux travaux ; on pourra consulter entre autre, dans l'ordre chronologique, [33], [36], [2], [3], [37], [19], [12] et [10].

Les raisons qui motivent notre choix pour la méthode de pénalisation résident essentiellement dans l'intention de simplifier le problème de départ, en le réduisant à une seule équation en vitesse. Il semble nécessaire d'effectuer une simplification du problème originel, en tenant compte du nombre limité de résultats obtenus jusqu'à présent avec les discrétisations par éléments finis ou par différences finies (cf. [27], [23], [17], [29] et [8]). Grâce à la pénalisation nous espérons éviter les difficultés rencontrées par certains auteurs cités ci-dessus, difficultés dues principalement au traitement de la contrainte d'incompressibilité et à la façon de hiérarchiser les inconnues du problème, surtout la pression. De plus, la méthode de pénalisation nous semble appropriée en prévision d'étendre aisément aux équations de Navier-Stokes la méthode multi-résolution en espace et en temps déjà proposée pour Burgers. En effet, l'équation en vitesse que l'on trouve après la pénalisation de la pression se rapproche assez des équations de Burgers généralisées. Donc, pour toutes ces raisons, la méthode de pénalisation se révèle une piste à suivre.

Toutefois la pénalisation ne permettra pas d'obtenir une résolution plus facile du système linéaire auquel on parvient après la discrétisation en espace (et éventuellement en temps) du problème considéré. En effet, étant donné que pour discrétiser notre problème il faut d'abord choisir une méthode d'éléments finis mixte en vitesse et en pression, d'autres méthodes largement utilisées exhibent des bonnes solutions numériques à coût raisonnable (voir par exemple la méthode du Lagrangien augmenté ou l'algorithme d'Uzawa). En revanche, la pénalisation demande des moyens de calcul puissants (surtout des machines qui possèdent un grand espace mémoire) car, comme nous le montrerons dans la suite, seulement une méthode de résolution directe permettra d'approcher efficacement le problème linéaire trouvé, la matrice à inverser possédant un conditionnement très mauvais.

Bien évidemment le problème de Navier-Stokes pénalisé (celui de Stokes étant seulement une phase préliminaire) sera la première étape de l'étude globale que nous nous proposons d'amener. Cette étape devrait nous permettre d'appliquer, sans trop de difficultés, les résultats obtenus précédemment (voir partie I). On anticipe déjà les perspectives futures qui consisteront dans la simulation numérique des vraies équations de Navier-Stokes en variables primitives grâce aux méthodes de *type incrémental* que nous cherchons à mettre au point avec le travail actuel.

Le problème de Stokes pénalisé

Le problème de Stokes généralisé (1.10) sera dorénavant approché par une famille à un paramètre ϵ ($\epsilon > 0$) de problèmes similaires, mais pour lesquels la condition d'incompressibilité $\operatorname{div} u = 0$ n'est plus vérifiée, soit :

$$\begin{cases} \alpha u_\epsilon - \frac{1}{Re} \Delta u_\epsilon + \nabla p_\epsilon = f & \text{dans } \Omega \\ \nabla \cdot u_\epsilon + \epsilon p_\epsilon = 0 & \text{dans } \Omega, \end{cases} \quad (1.16)$$

avec conditions aux limites suivantes :

$$u_\epsilon(x) = g(x) \quad \text{sur } \partial\Omega.$$

La formulation variationnelle associée à (1.16) est la suivante :

$$\begin{cases} \text{Trouver } (u_\epsilon, p_\epsilon) \in (H_g^1(\Omega))^2 \times L^2(\Omega) \text{ tel que} \\ a(u_\epsilon, v) - b(v, p_\epsilon) = (f, v) \quad \forall v \in (H_0^1(\Omega))^2 \\ b(u_\epsilon, q) + \epsilon(p_\epsilon, q) = 0 \quad \forall q \in L^2(\Omega), \end{cases} \quad (1.17)$$

avec les formes bilinéaires a et b définies par (1.12). On démontre que le problème (1.16) possède une solution unique et que, lorsque le paramètre ϵ tend vers 0, u_ϵ et p_ϵ tendent respectivement vers u et p solutions du problème (1.10) (voir [33] et [36]).

Rappelons la *condition Inf-Sup* due à Babuska et Brezzi :

il existe une constante $\beta > 0$ telle que

$$\inf_{q \in L^2(\Omega)} \sup_{v \in (H_0^1(\Omega))^2 \setminus \{0\}} \frac{b(v, q)}{\|v\| \|q\|} \geq \beta. \quad (1.18)$$

Cette condition est nécessaire et suffisante pour démontrer, dans un cadre général, l'existence et l'unicité de la solution du problème de Stokes ou de celui de Stokes pénalisé. En effet, afin de déterminer la variable p , il est nécessaire d'avoir des conditions supplémentaires sur l'opérateur associé à la forme bilinéaire b , ce qui revient à imposer (1.18) (voir [6] ou [19]). Si on fait l'hypothèse (1.18) et, en outre, s'il existe une constante $\gamma > 0$ telle que

$$a(v, v) + |\nabla \cdot v| \geq \gamma \|v\|^2 \quad \forall v \in (H_0^1(\Omega))^2,$$

alors pour $\epsilon < \epsilon_0$, ϵ_0 suffisamment petit, on a l'estimation d'erreur suivante :

$$|u_\epsilon - u|_1 + |p_\epsilon - p| \leq K \epsilon \quad (1.19)$$

avec K une constante qui dépend de la norme des fonctions f et g et des constantes β et γ seulement (cf. [2] ou [19] pour la démonstration de la majoration (1.19)).

Le problème de Navier-Stokes pénalisé

De la même manière que pour l'étude du problème de Stokes, le problème de Navier-Stokes (1.6) sera dorénavant approché par une famille à un paramètre ϵ positif de problèmes similaires pour lesquels la contrainte $\operatorname{div} u = 0$ est remplacée par une condition d'incompressibilité 'relaxée', soit :

$$\begin{cases} \frac{\partial u_\epsilon}{\partial t} - \frac{1}{Re} \Delta u_\epsilon + (u_\epsilon \cdot \nabla) u_\epsilon + \frac{1}{2} (\operatorname{div} u_\epsilon) \cdot u_\epsilon + \nabla p_\epsilon = f & \text{dans } \Omega \times [0, T] \\ \nabla \cdot u_\epsilon + \epsilon p_\epsilon = 0 & \text{dans } \Omega \times [0, T], \end{cases} \quad (1.20)$$

avec conditions initiales et aux limites suivantes :

$$\begin{aligned} u_\epsilon(x, 0) &= u_0(x) \quad \forall x \in \Omega, \\ u_\epsilon(x, t) &= g(x, t) \quad \forall x \in \partial\Omega \text{ et } \forall t \in [0, T]. \end{aligned}$$

Le terme $\frac{1}{2} (\operatorname{div} u_\epsilon) \cdot u_\epsilon$ a été ajouté dans l'équation de conservation de la quantité de mouvement pour que le problème (1.20) soit bien posé (voir remarque 3).

On démontre (en dimension d'espace égale à deux) que le problème (1.20) possède une solution unique et que, pour $\epsilon \rightarrow 0$, $u_\epsilon \rightarrow u$ et $p_\epsilon \rightarrow p$, u et p étant les solutions de (1.6) (voir [33], [3] et [19]). Le système (1.20), comme auparavant le système (1.16), est particulièrement intéressant car il permettra de se réduire à une seule équation en u_ϵ , obtenue en éliminant la pression p_ϵ . Cela permet non seulement une diminution importante du nombre d'inconnues du problème, mais en second lieu la possibilité de développer une méthode multi-résolution relativement proche de celle déjà exhibée pour résoudre les équations de Burgers.

La formulation variationnelle associée au problème (1.20) est la suivante :

$$\begin{cases} \text{Trouver } (u_\epsilon(t), p_\epsilon(t)) \in (H_g^1(\Omega))^2 \times L^2(\Omega) \text{ tel que} \\ \left(\frac{\partial u_\epsilon}{\partial t}, v \right) + \frac{1}{Re} ((u_\epsilon, v)) + \hat{b}(u_\epsilon, u_\epsilon, v) - (p_\epsilon, \operatorname{div} v) = (f, v) \quad \forall v \in (H_0^1(\Omega))^2 \\ (\operatorname{div} u_\epsilon, q) + \epsilon (p_\epsilon, q) = 0 \quad \forall q \in L^2(\Omega), \end{cases} \quad (1.21)$$

avec $\hat{b}$ la forme trilinéaire définie par :

$$\hat{b}(u, v, w) = ((u \cdot \nabla) v, w) + \frac{1}{2} ((\operatorname{div} u) \cdot v, w) \quad \forall u, v, w \in (H^1(\Omega))^2.$$

Remarque 3 La forme trilinéaire $\hat{b}$ satisfait la condition d'orthogonalité suivante :

$$\hat{b}(u, v, v) = 0 \quad \forall u, v \in (H^1(\Omega))^2,$$

cette hypothèse étant fondamentale pour démontrer l'existence et l'unicité de la solution du problème de Navier-Stokes (voir [36]).

L'estimation d'erreur (1.19) reste toujours valable pour le problème de Navier-Stokes pénalisé (cf. [3]). D'un point de vue numérique, le terme non linéaire ajouté dans (1.20), i.e. $\frac{1}{2} (\operatorname{div} u_\epsilon) \cdot u_\epsilon$, ne change pas le comportement qualitatif de la solution approchée des équations de Navier-Stokes (la vérification a été faite pour une solution stationnaire obtenue avec un faible nombre de Reynolds).

1.2 Discrétisation du problème de Stokes pénalisé

L'étape de pénalisation du problème de Stokes, considérée simplement comme une technique de résolution, suit la discrétisation du problème variationnel (1.11) ou, en d'autres termes, l'élimination de la pression dans le problème pénalisé (1.17) fait suite à la discrétisation en espace de ce problème. Nous avons choisi de poursuivre avec l'utilisation de la discrétisation par éléments finis, implémentée avec succès dans la partie I. En effet, la technique des éléments finis est la seule actuellement qui puisse s'appliquer à toutes les équations de la mécanique de fluides, quels que soient le domaine d'application, la complexité de la géométrie et les conditions aux limites.

Nous présentons dans cette section deux méthodes classiques d'éléments finis mixtes, employées pour approcher le problème variationnel considéré. La première méthode utilise l'élément 4P1-P0 et la seconde l'élément 4P1-P1 (appelé aussi P1 iso P2-P1). Nous avons opté dans un premier temps pour la méthode mixte 4P1-P0 plutôt que pour l'élément P2-P0, largement adopté dans la littérature. Ainsi nous réduisons l'encombrement mémoire puisque, avec un choix pareil, les matrices obtenues sont beaucoup plus creuses. En même temps nous continuons à approcher la pression de façon discontinue, ce qui permet de l'éliminer simplement lorsqu'on veut réduire le nombre d'équations à résoudre. Cependant, comme nous détaillerons dans la section 1.4.2, ce choix devient inadéquat quand nous cherchons à résoudre un problème de Stokes généralisé (étape fondamentale dans la résolution des équations de Navier-Stokes). Pour cela le deuxième choix a été celui du 4P1-P1, implémenté fructueusement déjà par de nombreux auteurs (cf. [27] et ses références).

1.2.1 Espaces d'approximation

Soit $P_l(\tau)$ l'ensemble des polynômes de degré inférieur ou égal à l sur chaque triangle τ (pour nos choix on se restreint aux valeurs $l = 0$ et 1). Comme pour les équations de Burgers, on considère la famille de grilles emboîtées

$$\mathcal{T}_0 \subset \mathcal{T}_1 \subset \dots \subset \mathcal{T}_d,$$

où chaque $\mathcal{T}_k$ ($k = 1, \dots, d$) est une triangulation régulière de Ω obtenue à partir de la triangulation $\mathcal{T}_{k-1}$ en divisant chaque triangle $\tau \in \mathcal{T}_{k-1}$ en quatre sous-triangles congruents. On note V_k l'espace des fonctions continues linéaires par morceaux de Ω dans $\mathbf{R}$:

$$V_k = \{v_k \in C^0(\Omega) \mid \forall \tau \in \mathcal{T}_k, v_k|_\tau \in P_1(\tau)\} \subset H^1(\Omega),$$

et Z_k l'espace des fonctions constantes sur chaque triangle de $\mathcal{T}_k$:

$$Z_k = \{z_k \in L^2(\Omega) \mid \forall \tau \in \mathcal{T}_k, z_k|_\tau \in P_0(\tau)\}.$$

De plus soient

$$V_{0,k} = V_k \cap H_0^1(\Omega)$$

et

$$V_{g,k} = \{v_k \in V_k \mid v_k = g_k^l \text{ sur } \partial\Omega\} \quad \text{pour } l = 1 \text{ ou } 2 ,$$

$g_k = (g_k^1, g_k^2) \in V_k^2$ étant une approximation de $g = (g_1, g_2)$ qui vérifie la condition de flux nul.

Les espaces d'éléments finis approchant $(H^1(\Omega))^2$, $(H_0^1(\Omega))^2$ et $(H_g^1(\Omega))^2$ sont alors respectivement (pour $k = 1, \dots, d$):

$$\begin{aligned} \mathcal{V}_k &= V_k \times V_k; \\ \mathcal{V}_{0,k} &= V_{0,k} \times V_{0,k}; \\ \mathcal{V}_{g,k} &= V_{g,k} \times V_{g,k}. \end{aligned}$$

De plus, les espaces d'éléments finis approchant $L^2(\Omega)$ sont les suivants (pour $k = 1, \dots, d$):

$$\begin{aligned} \text{élément } P1 : \mathcal{Q}_k &= V_{k-1}; \\ \text{élément } P0 : \mathcal{S}_k &= Z_{k-1}. \end{aligned}$$

Les degrés de liberté pour la vitesse sont alors les sommets A_i et les milieux de côtés B_i des triangles de $\mathcal{T}_{k-1}$ ($k = 1, \dots, d$) pour les deux méthodes mixtes. Pour ce qui concerne la pression les degrés de liberté sont les sommets des triangles de $\mathcal{T}_{k-1}$ dans le cas de l'élément 4P1-P1 ou le barycentre des triangles de $\mathcal{T}_{k-1}$ pour l'élément 4P1-P0.

1.2.2 Formulation discrète du problème

Cas du 4P1-P1

La version discrète du problème de Stokes (1.11), lorsque nous considérons les espaces d'éléments finis associés à la triangulation fine $\mathcal{T}_d$, s'écrit :

$$\left\{ \begin{array}{ll} \text{Trouver } (u_d, p_d) \in \mathcal{V}_{g,d} \times \mathcal{Q}_d \text{ tel que} & \\ a(u_d, v_d) - b(v_d, p_d) = (f, v_d) & \forall v_d \in \mathcal{V}_{0,d} \\ b(u_d, q_d) = 0 & \forall q_d \in \mathcal{Q}_d. \end{array} \right. \quad (1.22)$$

Remarque 4 *Il est bien connu que l'on ne peut choisir indépendamment les espaces d'approximations $\mathcal{V}_{0,d}$ et $\mathcal{Q}_d$ (de même pour les espaces $\mathcal{V}_{0,d}$ et $\mathcal{S}_d$ dans le cas de l'élément 4P1-P0). En fait la condition Inf-Sup, qui traduit la compatibilité entre la divergence discrète et la pression discrète, doit être vérifiée. En particulier on ne pourra obtenir une solution correcte du problème pénalisé que si le problème mixte est bien posé. La pénalisation est donc indissociable d'une méthode mixte en vitesse-pression (voir [14] ou [6]).*

L'approche correcte consiste à pénaliser le problème discret, plutôt que discrétiser le problème pénalisé où la pression a déjà été supprimée. En effet, ce dernier procédé nous conduirait à un choix inadéquat de l'espace d'approximation de $L^2(\Omega)$.

Dorénavant, pour alléger les notations, l'indice ϵ des solutions du problème pénalisé sera supprimé. On parlera encore de vitesse u et de pression p comme des solutions du problème (1.17).

De la même façon que pour (1.22), la version discrète du problème pénalisé (1.17) s'écrit :

$$\left\{ \begin{array}{l} \text{Trouver } (u_d, p_d) \in \mathcal{V}_{g,d} \times \mathcal{Q}_d \text{ tel que} \\ a(u_d, v_d) - b(v_d, p_d) = (f, v_d) \quad \forall v_d \in \mathcal{V}_{0,d} \\ b(u_d, q_d) + \epsilon(p_d, q_d) = 0 \quad \forall q_d \in \mathcal{Q}_d . \end{array} \right. \quad (1.23)$$

L'existence et l'unicité de la solution du problème (1.23) (de même que pour celle de (1.22)) sont toujours liées à la vérification de la condition Inf-Sup pour les espaces d'approximations $\mathcal{V}_{0,d}$ et $\mathcal{Q}_d$ de $(H_0^1(\Omega))^2$ et $L^2(\Omega)$:

il existe $\beta > 0$ indépendant du pas de discrétisation h_d tel que

$$\sup_{v_d \in \mathcal{V}_{0,d} \setminus \{0\}} \frac{b(v_d, q_d)}{\|v_d\|} \geq \beta |q_d| \quad \forall q_d \in \mathcal{Q}_d . \quad (1.24)$$

Ce résultat technique a été démontré par Bercovier et Pironneau (cf. [5]). La condition Inf-Sup permet également d'établir la convergence et l'ordre de la méthode (voir [5] pour la démonstration des relations suivantes) :

$$\begin{aligned} \|u_d - u\|_{(H^1(\Omega))^2} &\leq C_1 h (\|u\|_{(H^2(\Omega))^2} + \|p\|_{H^1(\Omega)/\mathbb{R}}) , \\ \|p_d - p\|_{L^2(\Omega)/\mathbb{R}} &\leq C_2 h (\|u\|_{(H^2(\Omega))^2} + \|p\|_{H^1(\Omega)/\mathbb{R}}) . \end{aligned}$$

Cas du 4P1-P0

Les relations (1.22), (1.23) et (1.24) se réécrivent de façon analogue lorsque les espaces d'approximation de $(H_0^1(\Omega))^2$ et $L^2(\Omega)$ sont respectivement $\mathcal{V}_{0,d}$ et $\mathcal{S}_d$ (il suffit de remplacer $\mathcal{Q}_d$ par $\mathcal{S}_d$ dans les relations mentionnées).

La condition Inf-Sup est vérifiée pour les espaces d'approximation $\mathcal{V}_{0,d}$ et $\mathcal{S}_d$ (cf. [28]). Ce résultat se démontre en construisant un opérateur de restriction r_h de W dans W_h , W étant le sous-espace des fonctions $(H_0^1(\Omega))^2$ à divergence nulle et W_h son approximation de Galerkin (cf. [36]). Cet opérateur r_h (généralement non linéaire) doit satisfaire

$$\lim_{h \rightarrow 0} \|r_h u - u\| = 0 , \quad (1.25)$$

sur tout W , mais il suffit de savoir le construire et démontrer la limite (1.25) seulement sur un sous-espace dense $\mathcal{W}$ de W ($\mathcal{W} = \{u \in (\mathcal{D}(\Omega))^2, \operatorname{div} u = 0\}$) pour en déduire la convergence sur tout W (voir [36]).

Choix du 4P1-P0 et du 4P1-P1

Le choix de l'élément 4P1-P1 a été postérieur à celui de l'élément 4P1-P0 et motivé par les perturbations détectées dans la solution approchée du problème de Stokes généralisé

avec paramètre α strictement positif (voir les sections 1.4.2 et 1.4.3). Dans un premier temps la pression a été approchée de façon discontinue, en utilisant les fonctions de base de Z_k ($k = 0, \dots, d-1$), pour pouvoir l'éliminer simplement triangle par triangle; ceci n'est plus le cas pour l'élément 4P1-P1.

De toute façon les raisons pour lesquelles nous avons choisi ces éléments sont résumées dans la suite :

1. Les éléments 4P1-P0 et 4P1-P1 vérifient la condition Inf-Sup, à l'inverse de l'élément P1-P0 par exemple.
2. Ces éléments sont tous les deux conformes en vitesse (à l'inverse de l'élément P1 non conforme - P0). L'élément 4P1-P1 est aussi conforme en pression; par contre la pression est approchée de façon discontinue dans le cas du 4P1-P0.
3. Par rapport aux éléments P2-P0 et P2-P1, les matrices sont plus creuses, en réduisant ainsi l'encombrement mémoire. De même la matrice de passage de la base nodale à la base hiérarchique sera triviale à calculer. Tout cela reste avantageux bien que l'ordre de convergence soit inférieur.
4. La construction de la base ainsi que l'assemblage des matrices du problème discret est facile (ce qui n'est plus le cas pour des bases à divergence nulle ou pour les éléments P1 bulle - P1 ou P2 bulle - P1 non conforme).

1.2.3 Forme matricielle du problème

Notons Σ_k l'ensemble des sommets de $\mathcal{T}_k$. Evidemment

$$\Sigma_{k-1} \subset \Sigma_k \quad (k = 1, \dots, d).$$

On pose

$$n_k = \text{Card}(\Sigma_k), \quad n_{0,k} = \text{Card}(\Sigma_k \cap \overset{\circ}{\Omega}),$$

et soit $t_{k-1} = \dim(Z_{k-1})$ le nombre de triangles dans la triangulation $\mathcal{T}_{k-1}$.

La base canonique dite nodale de V_k , pour $k = 0, \dots, d$, est notée $\{\phi_{k,i}\}_{i=1}^{n_k}$. Elle consiste en des fonctions $\phi_{k,i}$ prenant la valeur 1 au nœud i de $\mathcal{T}_k$ et 0 aux autres nœuds de la même triangulation.

Afin de garder une notation uniforme, indépendamment de l'élément choisi, on indique par $\{\xi_{k,i}\}_{i=1}^{m_k}$ la base nodale de l'espace d'éléments finis approchant $L^2(\Omega)$. Si ce dernier est approché par des éléments P1, alors on a $m_k = n_{k-1}$ et $\xi_{k,i}$ coïncide avec la fonction $\phi_{k-1,i}$ de la base de V_{k-1} , prenant donc la valeur 1 au nœud i et 0 aux autres nœuds de la triangulation $\mathcal{T}_{k-1}$. En revanche dans le cas des éléments P0, $m_k = t_{k-1}$ et $\xi_{k,i}$ est la fonction constante prenant la valeur 1 sur tout le triangle τ_i de $\mathcal{T}_{k-1}$ et 0 ailleurs.

On peut définir sur chaque niveau k ($k = 1, \dots, d$) les matrices nodales suivantes :

$$\begin{aligned} (\overline{\mathcal{A}}_k)_{ij} &= a(\phi_{k,i}, \phi_{k,j}) \quad 1 \leq i, j \leq n_{0,k} \\ (\mathcal{B}_{l,k})_{ij} &= -(\xi_{k,i}, \frac{\partial \phi_{k,j}}{\partial x_l}) \quad 1 \leq i \leq m_k, \quad 1 \leq j \leq n_{0,k} \quad \text{et } l = 1 \text{ ou } 2, \end{aligned}$$

et les matrices blocs

$$A_k = \begin{pmatrix} \bar{\mathcal{A}}_k & 0 \\ 0 & \bar{\mathcal{A}}_k \end{pmatrix} \quad \text{et} \quad B_k = (B_{1,k}, B_{2,k}).$$

De plus on définit $U_{l,k}$ et $F_{l,k}$ ($l = 1$ ou 2) les vecteurs colonnes formés des composantes de la vitesse et des forces extérieures exprimées dans la base nodale de $V_{0,k}$ et P_k le vecteur colonne de la pression exprimée dans la base nodale de V_{k-1} dans le cas P1 (ou Z_{k-1} dans le cas de l'élément P0) :

$$\begin{aligned} U_{l,k} &= (u_{l,1}, \dots, u_{l,n_{0,k}})^T, \\ F_{l,k} &= (f_{l,1}, \dots, f_{l,n_{0,k}})^T, \\ P_k &= (p_1, \dots, p_{m_k})^T. \end{aligned}$$

Alors vitesse et pression dans la base nodale s'écrivent :

$$\begin{aligned} u_k = (u_k^1, u_k^2) \in \mathcal{V}_{g,k}, \quad \text{étant} \quad u_k^l &= \sum_{i=1}^{n_{0,k}} u_{k,i}^l \phi_{k,i} + \sum_{i=n_{0,k}+1}^{n_k} g_k^l(M_i) \phi_{k,i} \\ & \quad (\text{avec } l = 1, 2 \text{ et } M_i \in \Sigma_k \cap \partial\Omega), \\ p_k = \sum_{i=1}^{m_k} p_{k,i} \xi_{k,i} &\in \mathcal{Q}_k \quad (\text{ou } \mathcal{S}_k). \end{aligned}$$

Nous allons écrire la forme matricielle du problème de Stokes discrétisé relative à la grille nodale $k = d$ (pour les grilles plus grossières la forme matricielle peut s'écrire de la même façon). De plus, pour alléger les notations, on négligera l'indice de grille. Le problème de Stokes discrétisé (1.22), écrit pour les fonctions de base de $\mathcal{V}_{0,d}$ et $\mathcal{Q}_d$ (resp. $\mathcal{S}_d$) est équivalent au problème suivant :

$$\begin{cases} \text{Trouver } U = (U_1, U_2)^T \in \mathbb{R}^{2 \times n_{0,d}} \text{ et } P \in \mathbb{R}^{m_d} \text{ solutions de :} \\ \begin{pmatrix} \mathcal{A} & B^T \\ B & 0 \end{pmatrix} \begin{pmatrix} U \\ P \end{pmatrix} = \begin{pmatrix} F \\ 0 \end{pmatrix}. \end{cases} \quad (1.26)$$

De la même façon, la forme matricielle du problème de Stokes pénalisé (1.23) s'écrit :

$$\begin{cases} \text{Trouver } U \in \mathbb{R}^{2 \times n_{0,d}} \text{ et } P \in \mathbb{R}^{m_d} \text{ solutions du système linéaire :} \\ \begin{pmatrix} \mathcal{A} & B^T \\ B & -\epsilon C \end{pmatrix} \begin{pmatrix} U \\ P \end{pmatrix} = \begin{pmatrix} F \\ 0 \end{pmatrix} \end{cases} \quad (1.27)$$

avec C une matrice $m_d \times m_d$ symétrique définie positive.

En général il existe deux possibilités. L'une consiste à pénaliser le problème discret (1.26), en considérant simplement $C = Id$, la matrice identité ; l'autre réside dans l'introduction de la matrice de masse C relative à la pression, avec $(C)_{ij} = (\xi_{d,i}, \xi_{d,j})$. Nous

verrons dans la suite que dans le premier cas il faut choisir ϵ assez petit, de l'ordre de Kh_d^2 , où K est une constante de l'ordre de 10^{-2} (ceci est en accord avec les résultats proposés dans la littérature; cf. [3]). Dans nos applications nous choisissons plutôt la seconde approche qui considère la matrice de masse $\mathcal{C}$ (ce choix est meilleur que le premier, surtout en prévision d'un raffinement local du maillage). La matrice $\mathcal{C}$ reste diagonale dans le cas de l'élément 4P1-P0 ($(\mathcal{C})_{ij} = \text{Aire}(\tau_i) \cdot \delta_{ij}$ où $\tau_i \in \mathcal{T}_{d-1}$); pour l'élément 4P1-P1 on peut encore choisir $\mathcal{C}$ diagonale en calculant par exemple le *mass lumping*. Le fait de considérer $\mathcal{C}$ diagonale implique un coût minimal pour son inversion dans la résolution du système (1.28). Ayant choisi le problème de Stokes pénalisé (1.23), il suffit d'un coefficient ϵ de l'ordre de 10^{-2} ou 10^{-3} .

A partir du système (1.27) nous pouvons éliminer la variable P . Le système linéaire à résoudre se réduit au suivant :

$$\begin{cases} \text{Trouver } U \in \mathbb{R}^{2 \times n_{0,d}} \text{ solutions de :} \\ \left(\mathcal{A} + \frac{1}{\epsilon} \mathcal{B}^T \mathcal{C}^{-1} \mathcal{B} \right) U = F . \end{cases} \quad (1.28)$$

La matrice $\mathcal{A} + \frac{1}{\epsilon} \mathcal{B}^T \mathcal{C}^{-1} \mathcal{B}$ est symétrique, définie positive, donc inversible.

Remarque 5 *Si on utilise une méthode itérative pour la résolution du système (1.28), le produit de la partie $\mathcal{B}^T \mathcal{C}^{-1} \mathcal{B}$ de la matrice (1.28) est obtenu en faisant consécutivement deux ou trois produits matrice vecteur (on rappelle que $\mathcal{C}$ est l'identité ou une matrice diagonale). En revanche si nous employons une méthode directe pour la résolution de (1.28), nous sommes obligés d'assembler la matrice $\mathcal{B}^T \mathcal{C}^{-1} \mathcal{B}$. Celle-ci est obtenue en calculant le produit des matrices facteurs assemblées précédemment. On souligne que la largeur de bande de $\mathcal{B}^T \mathcal{C}^{-1} \mathcal{B}$ est considérablement plus petite dans le cas du 4P1-P0 que dans le cas du 4P1-P1 (voir figure 1.1).*

On remarque enfin que l'on ne se sert pas d'une intégration numérique "réduite" pour approcher l'opérateur $\frac{1}{\epsilon} \int_{\tau \in \mathcal{T}_{d-1}} (\text{div } \phi_d)^2 d\tau$ comme il était proposé par Bercovier (cf. [2] et [3]) et repris par d'autres auteurs (cf. [12]). Une intégration de cette forme doit être utilisée lorsqu'on discrétise en espace le problème déjà pénalisé, ce qui n'est pas notre approche.

1.3 Résolution du système linéaire issu du problème de Stokes

L'objet de ce chapitre est l'étude de l'implémentation de plusieurs méthodes, itératives ou directes, employées pour la résolution du système (1.28). Au fur et à mesure de leur présentation, nous détaillerons les multiples difficultés rencontrées. En effet, il est bien connu que, indépendamment de la méthode d'éléments finis mixtes, la matrice $\mathcal{A} + \frac{1}{\epsilon} \mathcal{B}^T \mathcal{C}^{-1} \mathcal{B}$ est mal conditionnée. Ce défaut est très pénalisant, d'autant plus que

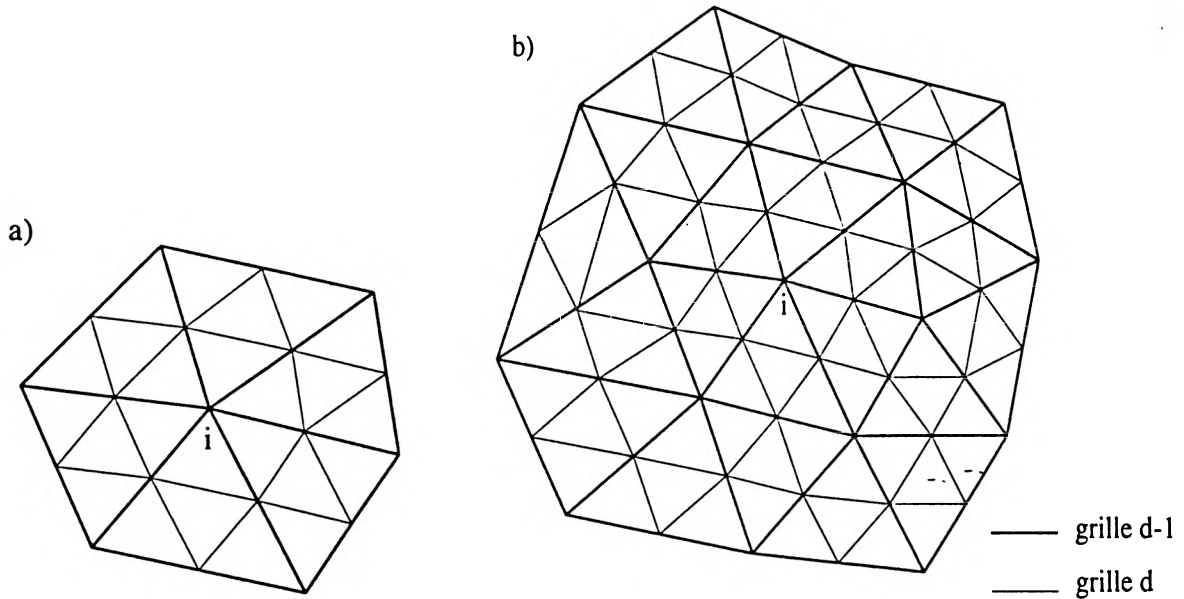


FIG. 1.1 – Nœuds qui donnent un coefficient a_{ij} non nul relativement à la ligne i de la matrice $B^T C^{-1} B$ (dans cette figure $i \in \Sigma_{d-1} \subset \Sigma_d$). a) cas du 4P1-P0; b) cas du 4P1-P1.

nous raffinons la discrétisation en espace afin de simuler des écoulements à nombres de Reynolds élevés. Effectivement on pourra constater dans la suite que le conditionnement de la matrice (1.28), défini comme le rapport entre la plus grande et la plus petite valeur propre, est de l'ordre de $\epsilon^{-1} h_d^4$, h_d étant le pas de discrétisation spatiale (l'estimation ci-dessus a été obtenue pour le problème de Stokes, i.e. avec le paramètre $\alpha = 0$; consulter [10] pour sa démonstration).

Quand on souhaite employer des méthodes itératives du type gradient conjugué pour résoudre le système linéaire (1.28), il faut retenir que la vitesse de convergence de ces méthodes est directement liée au conditionnement de la matrice à inverser (cf. [20]). Dans notre cas il est donc absolument nécessaire d'opter pour un préconditionnement du système linéaire, afin de réduire le nombre d'itérations pour atteindre la convergence de la solution. Malheureusement on verra que les préconditionnements classiques ont une influence négative sur la vitesse de convergence de la méthode itérative ou que, de toute manière, ils sont trop coûteux en temps de calcul (la performance d'un préconditionnement se mesure en nombre d'itérations pour la convergence mais aussi en temps CPU utilisé pour les calculs).

En revanche les méthodes de résolutions directes seront beaucoup plus performantes en temps de calcul que les méthodes itératives, leur utilisation devenant intéressante surtout lorsqu'on résoudra les équations évolutives de Navier-Stokes (la factorisation de la matrice à inverser sera calculée une fois pour toutes au début des itérations en temps). Cependant ces méthodes demandent le stockage en mémoire de la matrice factorisée complète, ce qui limitera fortement la taille de discrétisation du problème.

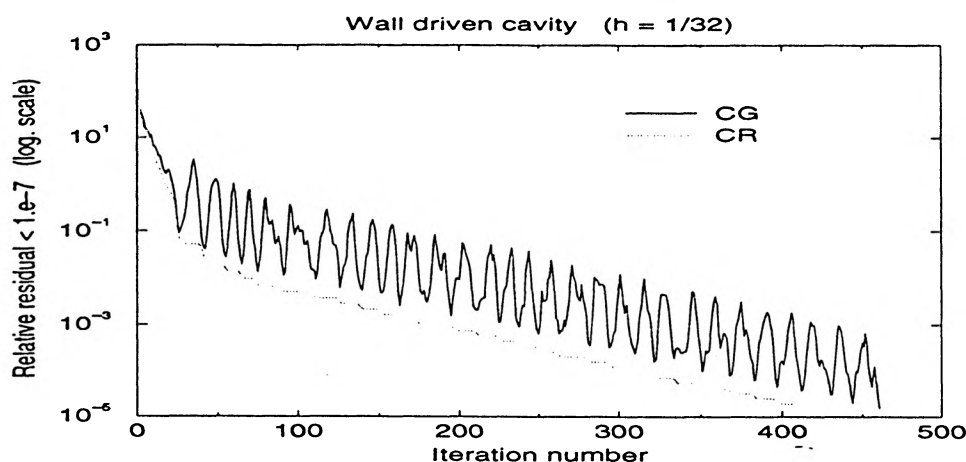


FIG. 1.2 – Résidu relatif des méthodes CG et CR employées pour résoudre le problème de Stokes dans la cavité entraînée discrétisé par 4P1-P1 ($\alpha = 1$, $Re = 100$ et $\epsilon = 10^{-3}$).

1.3.1 La méthode itérative : une étude de différents préconditionnements

Etant donné que la matrice à inverser est symétrique, définie positive, nous l'avons inversée par la méthode du gradient conjugué, notée CG, et par celle du résidu minimal (*conjugate residual method*), notée CR. En utilisant la méthode CG on observe que des oscillations sont présentes dans le résidu. Ces oscillations vont disparaître lorsqu'on se sert de la méthode CR. Dans le cas du résidu minimal en fait, au lieu de minimiser une fonctionnelle mesurant l'erreur absolue, on choisit de minimiser une fonctionnelle qui mesure le résidu. Cependant, dans le cas de CR les oscillations dans le résidu ont été remplacées par des stagnations dans sa descente (voir figure 1.2). A cause de cela, les deux méthodes restent à peu près comparables en nombre d'itérations pour atteindre la convergence du processus itératif. Toutefois la méthode CR converge en moins d'itérations ; on prendra alors celle-ci comme méthode de référence, bien que souvent les deux soient testées.

Obligés de préconditionner le système linéaire (1.28), nous avons testé, pour le choix initial de discrétisation par l'élément 4P1-P0, plusieurs méthodes de préconditionnement classiques qui se sont révélées inefficaces. On résume dans la suite les méthodes employées et les raisons pour lesquelles elles se manifestent inadéquates au préconditionnement.

1. Le préconditionnement par la méthode de Jacobi, en général moins efficace mais triviale à programmer par rapport à la méthode SSOR, réduit le nombre d'itérations pour la convergence de CR seulement si on considère la partie de (1.28) relative à l'opérateur de la chaleur (soit la seule matrice $\mathcal{A}$). Pour l'opérateur de Stokes pénalisé la méthode diverge. On remarque de plus que la matrice (1.28) n'est pas à diagonale dominante.

2. L'utilisation des éléments finis - bases hiérarchiques (pour leur présentation consulter la première partie de cette thèse) permet de construire un solveur efficace de l'équation de Poisson bidimensionnelle (cf. [38] et [1]). Cette méthode a un comportement très similaire à celui d'une méthode multigrille classique. Le préconditionnement par la base hiérarchique perd son intérêt lorsqu'on veut résoudre l'équation de la chaleur, car la matrice associée à cet opérateur admet un nombre de conditionnement assez petit. Enfin ce solveur devient extrêmement inefficace quand on cherche à l'employer comme méthode de préconditionnement pour la matrice (1.28). On trouve une confirmation de ce résultat négatif dans [10].
3. Nous avons déjà souligné que la matrice $\frac{1}{\epsilon} B^T C^{-1} B$ n'est pas assemblée quand on résout le système (1.28) par une méthode itérative. Or la matrice $\mathcal{M}$ de préconditionnement de la méthode CR doit être d'une part une matrice facile à inverser et d'autre part la matrice inverse $\mathcal{M}^{-1}$ doit être proche de la matrice (1.28), c'est-à-dire qu'elle doit vérifier

$$\mathcal{M}^{-1} \left(\mathcal{A} + \frac{1}{\epsilon} B^T C^{-1} B \right) \simeq Id$$

où Id est la matrice identité.

Notre choix se porte sur la factorisation incomplète de Cholesky : IC(0). Pour celle-ci il faut prévoir le stockage partiel de la matrice $\mathcal{M}$ et la résolution séquentielle de deux systèmes triangulaires. Puisqu'on cherche seulement une approximation de la partie $B^T C^{-1} B$ de la matrice (1.28), sans la calculer explicitement, on considère les matrices $\mathcal{M}$ suivantes :

a)

$$\mathcal{M} = \begin{pmatrix} \overline{\mathcal{M}} & 0 \\ 0 & \overline{\mathcal{M}} \end{pmatrix}$$

avec

$$(\overline{\mathcal{M}})_{ij} = \alpha(\phi_{d,i}, \phi_{d,j}) + \left(\frac{1}{Re} + \frac{1}{\epsilon} \right) (\nabla \phi_{d,i}, \nabla \phi_{d,j}) \quad 1 \leq i, j \leq n_{0,d}.$$

b)

$$\mathcal{M} = \begin{pmatrix} \overline{\mathcal{M}}_1 & 0 \\ 0 & \overline{\mathcal{M}}_2 \end{pmatrix}$$

où, pour $l = 1, 2$ et $1 \leq i, j \leq n_{0,d}$:

$$(\overline{\mathcal{M}}_l)_{ij} = \alpha(\phi_{d,i}, \phi_{d,j}) + \frac{1}{Re} (\nabla \phi_{d,i}, \nabla \phi_{d,j}) + \frac{1}{\epsilon} \left(\frac{\partial \phi_{d,i}}{\partial x_l}, \frac{\partial \phi_{d,j}}{\partial x_l} \right).$$

Remarque 6 Une variante est obtenue en prenant $\frac{1}{\epsilon S_\tau}$ à la place du coefficient $\frac{1}{\epsilon}$, où $S_\tau = \text{Aire}(\tau)$, $\tau \in \mathcal{T}_{d-1}$. S_τ est constante $\forall \tau \in \mathcal{T}_{d-1}$ car les triangulations $\mathcal{T}_k$ pour $k = 0, \dots, d$ sont uniformes.

Dans ce cas, donc, on préconditionne CR par une factorisation IC(0) et on approche la partie $B^T C^{-1} B$ de la matrice (1.28) par la matrice du Laplacien (cas a)) ou un découpage de celle-ci (cas b)). Pour cela cette méthode sera notée LICCR. Les résultats numériques (cf. tableau 1.1) montrent que dans les deux cas le préconditionnement par la matrice $\mathcal{M}$ réduit le nombre d'itérations pour atteindre la convergence d'environ la moitié par rapport à la méthode CR non préconditionnée, mais seulement lorsqu'on résout le problème de Stokes (le nombre d'itérations pour la convergence est donné dans le tableau 1.1). Les résultats dans ce tableau sont relatifs au test qui considère la matrice $\mathcal{M}$ du cas b) avec coefficient égal à $\frac{1}{\epsilon S_\tau}$. En revanche, le préconditionnement n'agit plus pour un Stokes généralisé ($\alpha > 0$), ou plutôt son action devient même néfaste. Or, même pour le problème de Stokes le préconditionnement est inintéressant du point de vue du temps de calcul, car l'assemblage et surtout l'inversion de $\mathcal{M}$ sont trop onéreuses par rapport au gain obtenu dans l'accélération de la convergence, ce qui explique, en conclusion, une perte d'efficacité.

Une condition nécessaire et suffisante pour la convergence de CR préconditionné est que la matrice de préconditionnement $\mathcal{M}$ soit symétrique définie positive. Or si on regarde les valeurs de la factorisation IC(0) de $\mathcal{M}$ (vérification effectuée dans le cas b)) on observe que beaucoup de valeurs sur la diagonale sont très proches de zero et parfois même négatives. Cela implique que $\mathcal{M}$ n'est pas définie positive.

4. Nous cherchons à résoudre les problèmes de *breakdown* ou de non convergence de la méthode LICCR, entraînés par le fait que $\mathcal{M}$ n'est pas définie positive. Comme suggéré par Manteuffel (cf. [25] et [26]), on essaie de rendre la factorisation de Cholesky stable. La stabilité de la factorisation incomplète implique la convergence de la méthode itérative préconditionnée. Une condition suffisante pour la stabilité est que la matrice soit définie positive et à diagonale dominante.

Cette méthode, appelée *shifted incomplete Cholesky factorization* et notée SICCR, demande avant le calcul de la factorisation les deux étapes suivantes :

- i) Faire un *scaling* de $\mathcal{M}$, soit calculer $\widehat{\mathcal{M}}$ où

$$\widehat{\mathcal{M}} = D^{-1/2} \mathcal{M} D^{-1/2} \quad \text{avec } D = \text{diag}(\dots \mathcal{M}_{ii} \dots).$$

Alors on a $\widehat{\mathcal{M}} = I - B$ avec I la matrice identité et B matrice à diagonale nulle i.e. $B_{ii} = 0$.

- ii) Considérer

$$\widehat{\mathcal{M}}(\sigma) = I - \frac{1}{1 + \sigma} B \quad \text{avec } \sigma \in \mathbf{R}^+.$$

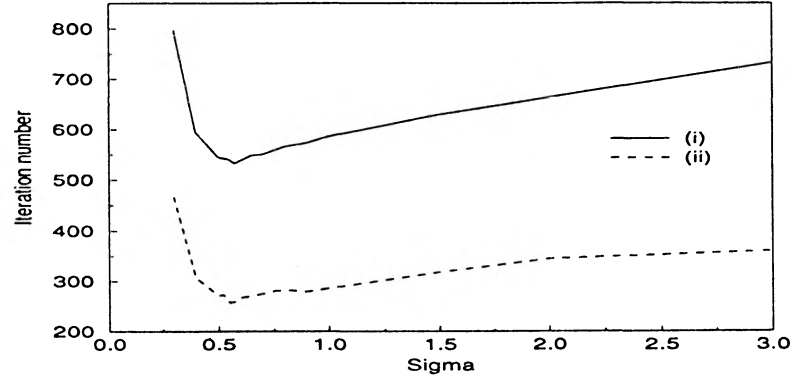


FIG. 1.3 – Nombre d'itérations de SICCR en fonction du coefficient $\bar{\sigma}$. Deux choix des paramètres: (i) $\alpha = 1, Re = 100, h_d = 1/32$; (ii) $\alpha = 10, Re = 1000, h_d = 1/64$.

Pour σ suffisamment grand la matrice $\widehat{\mathcal{M}}(\sigma)$ est à diagonale dominante. On cherche donc des σ pour lesquels la factorisation incomplète IC(0) de $\widehat{\mathcal{M}}(\sigma)$ soit stable. Les tests numériques montrent que $\widehat{\mathcal{M}}(\sigma)$ est à diagonale dominante pour des σ très grands, de l'ordre de $\frac{1}{\epsilon}$. On note

$$\tilde{A}(\sigma) = \left(\widehat{\mathcal{M}}(\sigma) \right)^{-1} \left(A + \frac{1}{\epsilon} B^T C^{-1} B \right),$$

la matrice que la méthode CR préconditionnée doit inverser. Manteuffel montre que le conditionnement de cette matrice, noté $\kappa(\tilde{A}(\sigma))$, atteint un minimum pour σ_0 assez petit, $\sigma_0 > \sigma_c$ (où σ_c est tel que, si $\sigma \leq \sigma_c$, la factorisation n'est plus définie positive) et que pour $\sigma_0 < \sigma_1 < \sigma_2$ on a $\kappa(\tilde{A}(\sigma_1)) < \kappa(\tilde{A}(\sigma_2))$. Puisque " $\widehat{\mathcal{M}}(\sigma)$ à diagonale dominante" est une condition suffisante mais non nécessaire pour la stabilité de la factorisation, il nous reste à trouver des coefficients $\sigma \ll \frac{1}{\epsilon}$ pour lesquels la diagonale de $\widehat{\mathcal{M}}(\sigma)$ ne contient aucun terme nul ou négatif. En figure 1.3 on trace le nombre d'itérations pour la convergence de SICCR en fonction du coefficient σ . Ces résultats correspondent à une résolution du problème de Stokes généralisé dans la cavité entraînée, avec deux choix des paramètres :

- (i) $\alpha = 1, Re = 100, \epsilon = 10^{-3}$ et $h_d = 1/32$;
- (ii) $\alpha = 10, Re = 1000, \epsilon = 10^{-3}$ et $h_d = 1/64$.

La convergence de la méthode itérative est atteinte quand la norme du résidu devient inférieure à 10^{-6} . On remarque que le choix de σ pour obtenir le minimum d'itérations pour la convergence n'est pas évident *a priori*. Enfin les tableaux 1.1 établissent une comparaison entre les nombres d'itérations pour atteindre la convergence des différentes méthodes employées jusqu'ici dans la résolution du système (1.28) (de même ou pourra comparer le temps CPU de ces méthodes). On constate que, malgré une certaine accélération dans la convergence de SICCR par rapport à

α	CR	LICCR	SICCR - $\tilde{\mathcal{M}}(0.55)$	SICCR - $\tilde{\mathcal{M}}(0.575)$
0.0	1189 (12.8)	662 (16.7)	726 (41.9)	723 (41)
1.0	878 (9.6)	509 (13)	540 (31)	532 (30)
2.0	698 (7.6)	429 (11)	437 (25)	438 (25)
10.0	347 (4)	392 (10)	239 (14.2)	232 (13)

(A)

α	CR	LICCR	SICCR - $\tilde{\mathcal{M}}(0.6)$	SICCR - $\tilde{\mathcal{M}}(1.0)$
0.0	4186 (185)	2337 (240)	2300 (532)	2555 (588)
1.0	1203 (54)	834 (87)	680 (159)	758 (177)
2.0	887 (40)	763 (80)	513 (121)	562 (131)
10.0	410 (19)	778 (81)	267 (64)	286 (68)
100.0	345 (16)	962 (99)	258 (62)	267 (64)

(B)

TAB. 1.1 – Résolution du problème de Stokes généralisé :

(A) $Re = 100, \epsilon = 10^{-3}, h_d = 1/32$; (B) $Re = 1000, \epsilon = 10^{-3}, h_d = 1/64$.

Comparaison, en fonction du coefficient α , des méthodes CR, LICCR et SICCR (pour deux choix du paramètre σ) : nombre d'itérations pour la convergence et, entre parenthèses, temps CPU.

CR, ce dernier type de préconditionnement reste également inintéressant du point de vue du temps de calcul.

1.3.2 Les valeurs propres de la matrice à inverser

Les difficultés rencontrées pour définir un solveur performant du système (1.28) motivent l'étude des valeurs propres de la matrice à inverser. Une deuxième motivation porte sur la recherche d'un préconditionnement efficace des méthodes itératives employées afin de résoudre le système linéaire (1.28). Cette étude se révèle intéressante en elle même car elle nous montrera une séparation inattendue des ces valeurs propres. Néanmoins, anticipons dès maintenant que l'étude présentée ici n'aboutira à aucun nouveau préconditionneur de CR, surtout à cause des difficultés théoriques rencontrées pour le définir. Le choix qui en résulte sera l'emploi des méthodes directes pour inverser la matrice (1.28).

La détermination des valeurs propres de la matrice a été réalisée à l'aide de la bibliothèque de calcul EISPACK qui réduit toute matrice symétrique, donnée avec un stockage bande, à une matrice tridiagonale et ensuite, grâce à une transformation QL, évalue ses valeurs propres (et éventuellement les vecteurs propres associés).

Représentons toutes les valeurs propres λ_i de la matrice $\mathcal{A} + \frac{1}{\epsilon} \mathcal{B}^T \mathcal{B}$ en fonction de la constante K (dans ce cas, $\mathcal{C}$ étant la matrice identité, on pose $\epsilon = Kh_d^2$ pour la discrétisation à l'aide du 4P1-P1 et $\epsilon = K \frac{h_d^2 - 1}{2} = \frac{K}{8} h_d^2$ dans le cas du 4P1-P0). Nous découvrons une nette séparation des valeurs propres de la matrice en deux zones, ceci

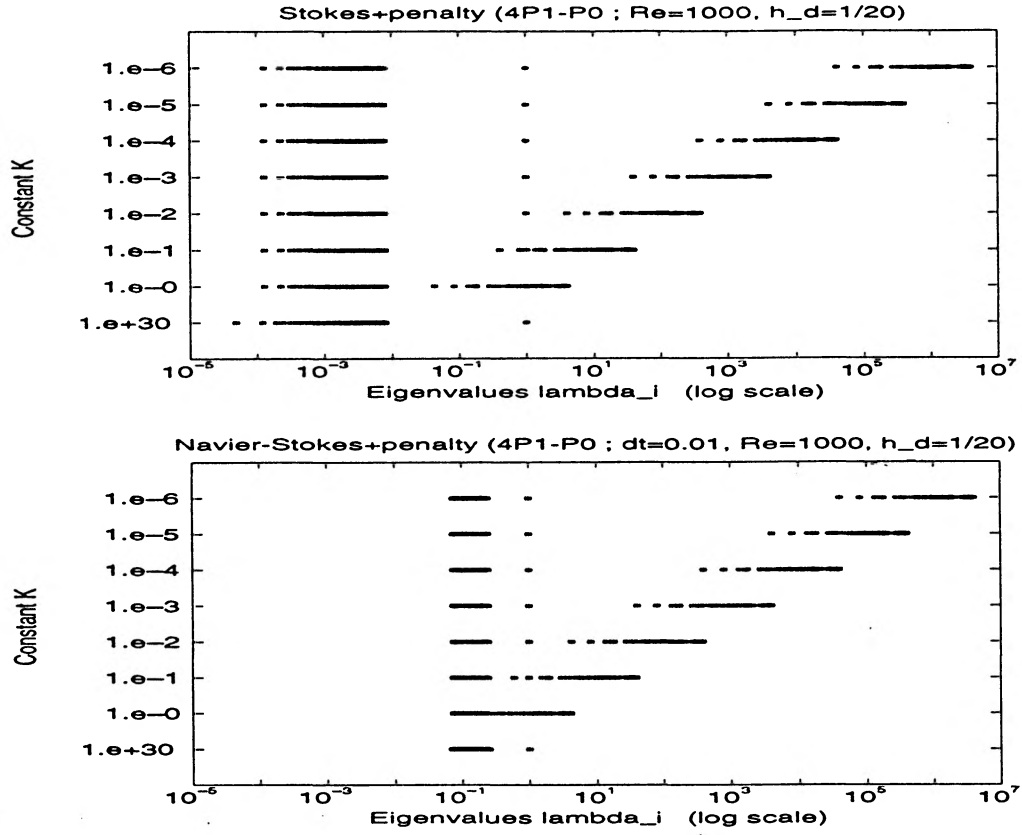


FIG. 1.4 – Valeurs propres de la matrice relative au problème de Stokes et de celle relative au problème de Stokes généralisé ($\epsilon = \frac{K}{8}h_d^2$).

indépendamment de la valeur du coefficient ϵ , de l'élément fini choisi pour discrétiser le problème et même du pas de discrétisation spatiale h_d (cf. par exemple la figure 1.5).

Relativement au premier choix de l'élément fini, i.e. l'élément 4P1-P0, nous avons tracé (voir la figure 1.4) les valeurs propres de la matrice issue de la discrétisation du problème de Stokes en les comparant avec celles relatives à la matrice du problème de Stokes généralisé. Remarquons que la valeur propre unité correspond aux lignes de la matrice pour lesquelles nous avons imposé les conditions aux limites. Dans cet exemple, le choix des paramètres est le suivant : pas de discrétisation spatiale $h_d = \frac{1}{20}$ (rappelons que la triangulation $\mathcal{T}_d$ est uniforme), nombre de Reynolds $Re = 1000$ et $\alpha = 100$ pour le problème de Stokes généralisé. La figure 1.4 permet aussi de vérifier l'estimation relative au conditionnement de la matrice de Stokes, qui théoriquement est de l'ordre de $\epsilon^{-1}h_d^2$, avec $\epsilon = \frac{K}{8}h_d^2$ (cf. [10]). Il est toutefois difficile, à partir de la répartition des valeurs propres de ces deux problèmes, d'en déduire les causes pour lesquelles la méthode de résolution LICCR du système (1.28) réduit le nombre d'itérations pour attendre la convergence relativement au problème de Stokes et échoue totalement dans la résolution du problème de Stokes généralisé.

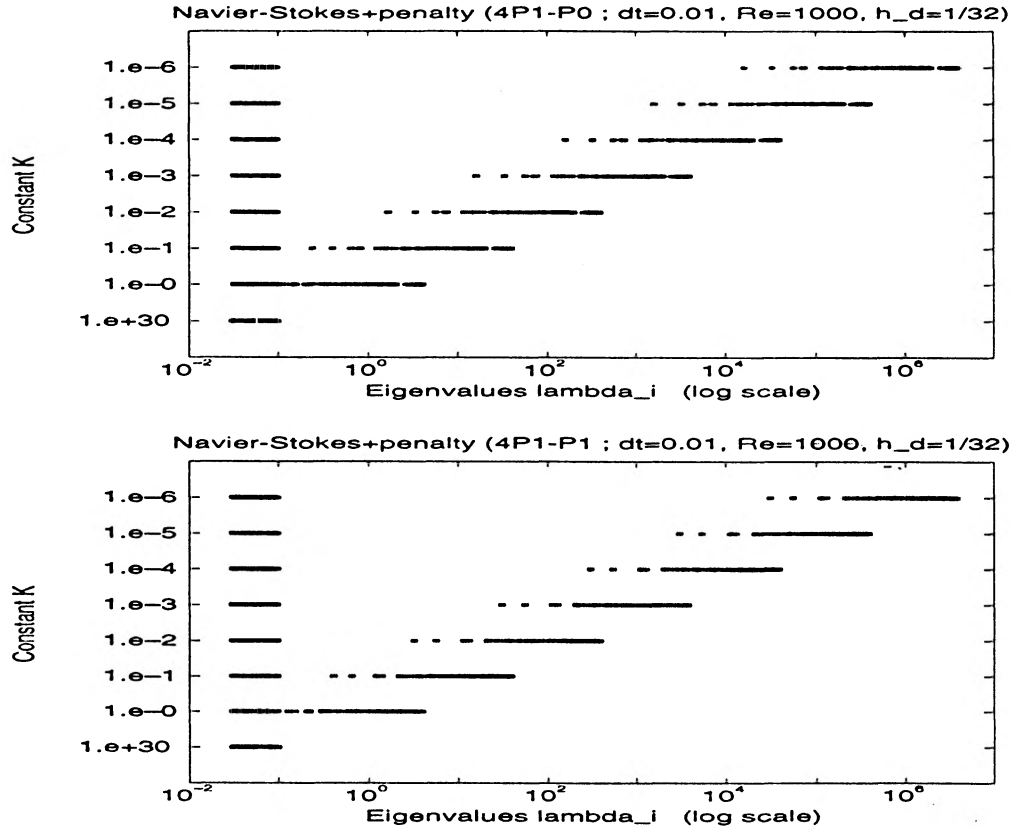


FIG. 1.5 – Valeurs propres de la matrice relative au problème de Stokes généralisé : discrétisation par l'élément 4P1-P0 ($\epsilon = \frac{K}{8} h_d^2$) et par l'élément 4P1-P1 ($\epsilon = K h_d^2$).

En figure 1.5 on représente les valeurs propres λ_i pour un choix des paramètres qui est le suivant : $\alpha = 100$, $Re = 1000$ et $h_d = 32$. Les résultats de l'élément 4P1-P0 sont comparables à ceux de l'élément 4P1-P1.

Si l'on trace seulement les valeurs propres de la matrice $\frac{1}{\epsilon} \mathcal{B}^T \mathcal{B}$ (voir la figure 1.6) on s'aperçoit qu'un grand nombre de λ_i sont nulles (le calcul numérique donne des valeurs très proches du zéro machine).

Détaillons ce dernier résultat qui ne surprend pas, sachant que le noyau de la matrice $\mathcal{B}$, noté $\text{Ker}(\mathcal{B})$, ne se réduit pas à l'ensemble $\{0\}$. On rappelle le théorème qui établit la décomposition en valeurs singulières d'une matrice A qu'on note usuellement SVD (*singular value decomposition*) (cf. [20]).

Proposition 2 Toute matrice $A \in \mathbb{R}^{m \times n}$ ($m \geq n$) peut être factorisée sous la forme

$$A = W \Sigma V^T$$

avec $W \in \mathbb{R}^{m \times m}$ et $V \in \mathbb{R}^{n \times n}$ matrices orthogonales. Σ est alors une matrice de $\mathbb{R}^{m \times n}$

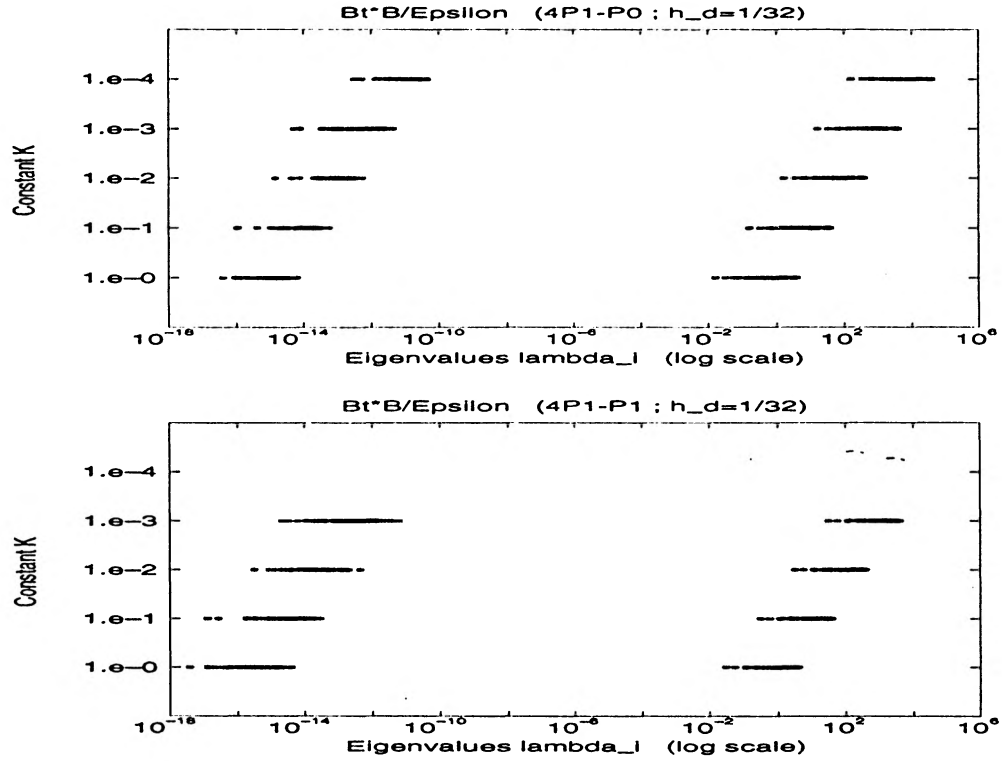


FIG. 1.6 – Valeurs propres de la matrice $\frac{1}{\epsilon} B^T B$: discrétisation par l'élément 4P1-P0 ($\epsilon = \frac{K}{8} h_d^2$) et par l'élément 4P1-P1 ($\epsilon = K h_d^2$).

de la forme :

$$\Sigma = \begin{pmatrix} \sigma_1 & & 0 \\ & \ddots & \\ 0 & & \sigma_n \\ 0 & \dots & 0 \end{pmatrix},$$

où les σ_i sont les valeurs singulières de A .

Remarque 7 Si on admet la Proposition 2 on peut écrire :

$$A^T A = (V \Sigma^T W^T) (W \Sigma V^T) = V \Sigma^T \Sigma V^T = V \begin{pmatrix} \sigma_1^2 & & 0 \\ & \ddots & \\ 0 & & \sigma_n^2 \end{pmatrix} V^T.$$

Les σ_i^2 sont bien les valeurs propres de $A^T A$ et V est formée des vecteurs propres de $A^T A$ (bien évidemment certaines σ_i^2 peuvent être nulles, selon la dimension du noyau de A).

On rappelle de plus la relation d'Euler pour les triangles :

$$S - C + T = 1,$$

où S est le nombre de sommets, C le nombre de côtés et T le nombre de triangles dans la triangulation $\mathcal{T}_{d-1}$. Asymptotiquement on a les relations suivantes :

$$C \simeq \frac{3T}{2} \quad \text{et} \quad S \simeq C - T \simeq \frac{T}{2} .$$

On en déduit la dimension des espaces de discrétisation :

élément 4P1-P0 :

$$\begin{aligned} \dim(\mathcal{V}_d) &\simeq (2S + 2C) \simeq 4T \\ \dim(\mathcal{S}_d) &\simeq T \end{aligned}$$

élément 4P1-P1 :

$$\begin{aligned} \dim(\mathcal{V}_d) &\simeq (2S + 2C) \simeq 4T \\ \dim(\mathcal{Q}_d) &\simeq S \simeq \frac{T}{2} \end{aligned}$$

On a donc plus de contraintes sur la pression lorsqu'on la discrétise par P0 que lorsqu'on la discrétise par P1. On en déduit que la dimension de $\text{Ker}(\mathcal{B})$ est plus petite dans le cas de l'élément 4P1-P0 que dans le cas du 4P1-P1 (l'estimation asymptotique est confirmée par les valeurs du tableau 1.2).

Revenons à la matrice $\mathcal{B}$ associée à l'opérateur qui discrétise la contrainte $\text{div } u = 0$. Soit $\mathcal{B} \in \mathbf{R}^{m_d \times (2 \times n_{0,d})}$ avec $m_d < (2 \times n_{0,d})$. Grâce à la Proposition 2, la matrice $\mathcal{B}$ peut être factorisée sous la forme

$$\mathcal{B} = V \Gamma W^T$$

avec $V \in \mathbf{R}^{m_d \times m_d}$ et $W \in \mathbf{R}^{(2 \times n_{0,d}) \times (2 \times n_{0,d})}$ matrices orthogonales et $\Gamma \in \mathbf{R}^{m_d \times (2 \times n_{0,d})}$ matrice de la forme :

$$\Gamma = \begin{pmatrix} \gamma_1 & & 0 & 0 \\ & \ddots & & \vdots \\ 0 & & \gamma_n & 0 \end{pmatrix} .$$

Nous pouvons finalement écrire :

$$\mathcal{B}^T \mathcal{B} = (W \Gamma^T V^T) (V \Gamma W^T) = W \Gamma^T \Gamma W^T$$

avec

$$\Gamma^T \Gamma = \begin{pmatrix} \gamma_1 & & 0 \\ & \ddots & \\ 0 & & \gamma_n \\ 0 & \dots & 0 \end{pmatrix} \begin{pmatrix} \gamma_1 & & 0 & 0 \\ & \ddots & & \vdots \\ 0 & & \gamma_n & 0 \end{pmatrix} = \begin{pmatrix} \gamma_1^2 & & 0 & 0 \\ & \ddots & & \vdots \\ 0 & & \gamma_n^2 & 0 \\ 0 & \dots & 0 & 0 \end{pmatrix} .$$

Les γ_i^2 sont les valeurs propres de $\mathcal{B}^T \mathcal{B}$, positives ou nulles, de multiplicité supérieure ou égale à un. Comptons alors le nombre de γ_i^2 strictement positives en déduisant le nombre de valeurs propres nulles (écrites entre parenthèses dans le tableau 1.2). Lorsqu'on divise le pas h_d par 2 le nombre de $\gamma_i^2 > 0$ se multiplie exactement par 4 dans le cas du 4P1-P0. Malheureusement cette règle ne semble plus être respectée dans le cas du 4P1-P1.

$1/h_d$	4P1-P0	4P1-P1
16	127 (323)	80 (370)
32	511 (1411)	288 (1634)
64	2047 (5891)	1088 (6850)

TAB. 1.2 – Nombre des valeurs propres strictement positives et nulles (les dernières entre parenthèses) de $\frac{1}{\epsilon} \mathcal{B}^T \mathcal{B}$ en fonction du pas de discrétisation du maillage $\mathcal{T}_d$.

Nous concluons ce paragraphe relatif à l'étude des valeurs propres, en donnant seulement une esquisse de l'idée sous-jacente au fait d'avoir compté les γ_i^2 positives en fonction du pas de discrétisation h_d , idée liée à la recherche d'un préconditionnement performant de la méthode itérative CR. Pour l'élément 4P1-P0, la manière dans laquelle les valeurs propres de la matrice $\mathcal{A}_k + \frac{1}{\epsilon} \mathcal{B}_k^T \mathcal{B}_k$ se séparent est totalement indépendante de la grille de discrétisation k . L'idée serait alors de ramener la résolution du système nodal (1.28) à celle du système restreint sur une grille plus grossière, pour lequel une factorisation complète de la matrice est peu coûteuse en encombrement mémoire. Le passage d'une grille à l'autre serait réalisé à chaque itération de CR par une matrice de restriction et une de prolongement, semblables aux matrices de passage de la base nodale à la base hiérarchique et vice versa. Toutefois, il n'est pas facile de montrer théoriquement que ce procédé puisse amener à une réduction du nombre d'itérations pour la convergence de la méthode itérative. De plus, nous verrons dans la suite que la discrétisation par l'élément 4P1-P0 sera abandonnée pour maintenir seulement celle par l'élément 4P1-P1. Mais pour cet élément la séparation des valeurs propres de la matrice $\mathcal{A}_k + \frac{1}{\epsilon} \mathcal{B}_k^T \mathcal{B}_k$ ne semble plus être complètement indépendante de la grille k . En fait le nombre de valeurs propres $\gamma_i^2 > 0$ ne se multiplie plus par 4 lorsqu'on raffine le pas d'espace h_d en le divisant par 2. Il en résulte un mélange entre les valeurs propres de la partie $\frac{1}{\epsilon} \mathcal{B}_k^T \mathcal{B}_k$ et celles de la partie $\mathcal{A}_k$ de notre matrice, mélange qui dépend du niveau de discrétisation k . En conclusion, face aux difficultés théoriques, nous avons décidé d'arrêter le développement et la mise au point d'une nouvelle méthode de préconditionnement liée à la hiérarchisation de la matrice à inverser, pour nous orienter définitivement vers l'emploi des méthodes directes afin de résoudre le système linéaire (1.28).

1.3.3 La méthode directe : mise en œuvre de la factorisation complète

Suivant les notations de la section 1.2.3, on a :

$$\Sigma_{k-1} \subset \Sigma_k \quad \text{pour } k = 1, \dots, d.$$

En plus de la numérotation lexicographique des nœuds de chaque triangulation, la prise en compte de la famille de grilles emboîtées implique une numérotation hiérarchique, sur au moins deux niveaux, de chaque ensemble de nœuds Σ_k . On note $\hat{A}^{(d)}$ la matrice

nodale du système linéaire (1.28) relative à la grille fine :

$$\hat{A}^{(d)} = \mathcal{A}_d + \frac{1}{\epsilon} \mathcal{B}_d^T \mathcal{C}_d^{-1} \mathcal{B}_d .$$

Cette matrice, grâce à la numérotation hiérarchique des sommets, se met sous la forme d'une matrice par blocs :

$$\hat{A}^{(d)} = \begin{pmatrix} \hat{A}_{gg}^{(d)} & \hat{A}_{gf}^{(d)} \\ \hat{A}_{fg}^{(d)} & \hat{A}_{ff}^{(d)} \end{pmatrix} , \quad (1.29)$$

où :

- $\hat{A}_{fg}^{(d)} = \left(\hat{A}_{gf}^{(d)} \right)^T$;
- $\hat{A}_{ff}^{(d)}$ est une matrice carrée, symétrique, définie positive, d'ordre $(n_d - n_{d-1})^2$ (ou $(n_{0,d} - n_{0,d-1})^2$ si l'on élimine les points du bord) ;
- $\hat{A}_{gg}^{(d)}$ est une matrice carrée, symétrique, définie positive, d'ordre n_{d-1}^2 (où resp. $n_{0,d-1}^2$). Cette matrice se réduit à une matrice diagonale seulement lorsque $\hat{A}^{(d)} = \mathcal{A}_d$, c'est-à-dire si on considère seulement la partie de (1.28) relative à l'opérateur de la chaleur (car les n_{d-1}^2 premières fonctions de la base nodale, avec numérotation issue de la base hiérarchique, ont des supports deux à deux disjoints).

Naturellement cette structure par blocs peut être réitérée de façon récursive pour les matrices nodales de niveau plus grossier $\hat{A}^{(k)}$, $k = 1, \dots, d-1$.

On reconnaît la structure par blocs (1.29) dans les deux représentations en figure 1.7 des matrices issues respectivement de la discrétisation par l'élément 4P1-P0 et par l'élément 4P1-P1. En effet on a représenté la position des blocs 2×2 non nuls $\left(\hat{A}^{(d)} \right)_{ij}$, avec $j \neq i$. On rappelle l'observation déjà faite dans la remarque 5 sur les nœuds donnant un coefficient non nul de la matrice : ces coefficients sont plus nombreux dans le cas de l'élément 4P1-P1 que dans le cas du 4P1-P0.

Il est impossible de réaliser une factorisation complète de Cholesky des matrices creuses avec cette structure par blocs à cause de l'énorme taille mémoire nécessaire pour stocker la matrice factorisée, car cette opération implique toujours un remplissage du profil de la matrice. Il est indispensable donc de réduire le plus possible la largeur de bande de la matrice que l'on veut factoriser.

Récursivement sur chaque niveau hiérarchique k ($k = 1, \dots, d$), nous cherchons une permutation σ_k des nœuds Σ_k telle que, si on note P_{σ_k} la matrice de permutation associée à σ_k , la matrice $P_{\sigma_k}^T \hat{A}^{(k)} P_{\sigma_k}$ ait une largeur de bande inférieure à celle de $\hat{A}^{(k)}$.

L'algorithme de Cuthill-McKee (cf. [11]), qui emploie la notion de graphe non orienté associé à une matrice creuse symétrique, permet d'obtenir une telle permutation σ_k dans le cas d'une matrice irréductible. En connaissant pour chaque nœud les nœuds voisins et après avoir défini un degré sur les sommets (qui peut être le nombre de voisins), on part d'un nœud élu à point de départ et on ordonne les sommets selon leur degré en

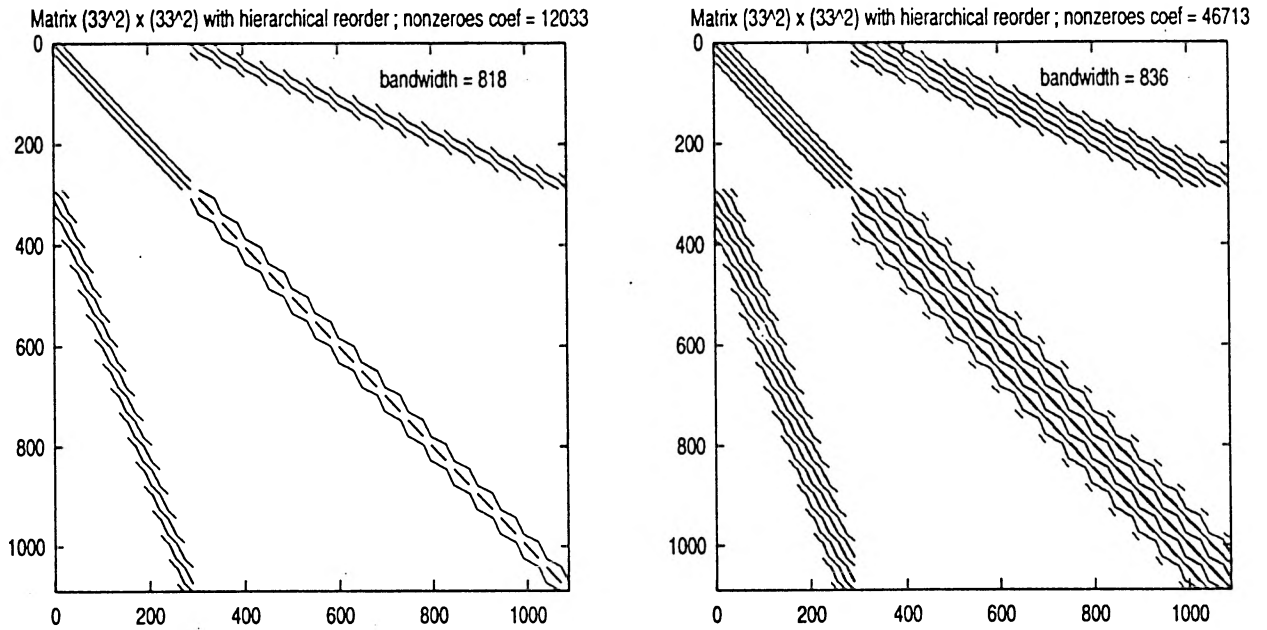


FIG. 1.7 – Structure par blocs, issue de la numérotation hiérarchique sur deux niveaux, des matrices de taille $33^2 \times 33^2$: à gauche matrice $\hat{A}^{(d)}$ issue de la discrétisation par 4P1-P0, à droite matrice $\hat{A}^{(d)}$ issue de la discrétisation par 4P1-P1.

couche de nœuds voisins. Cet algorithme n'est pas optimal mais il donne des résultats de bonne qualité pour un coût raisonnable.

Détaillons un peu plus la méthode de Cuthill-McKee qui consiste à faire une renumérotation des nœuds en couches successives. Pour cela, on dispose de deux tableaux d'entiers : le premier donne pour chaque nœud i , interne au domaine, le nombre de nœuds voisins (grâce à une liste simplement chaînée), le second fournit les nœuds voisins, ce qui revient à donner l'indice de colonne j des coefficients non nuls $(\hat{A}^{(k)})_{ij}$ de la matrice, avec $j \neq i$. Les sommets sur $\partial\Omega$ sont rangés (en ordre croissant ou décroissant) à la fin du tableau de permutation σ_k , car les lignes correspondantes de la matrice, ayant imposé les conditions aux limites de Dirichlet, ont seulement le coefficient sur la diagonale qui est non nul. Ceci permet aussi d'obtenir une matrice irréductible. Il est également nécessaire de définir un critère pour choisir un sommet de départ. Pour cela le point de départ sera un nœud d'un coin du domaine ou un nœud de degré minimal. Ensuite, pour une couche donnée, on construit la liste des nœuds non encore renumérotés voisins d'un nœud de cette couche, en les triant par insertion en fonction de leur degré croissant. Enfin, il est bien connu qu'en général l'algorithme de Cuthill-McKee inverse réduit ultérieurement la largeur de bande de la matrice considérée.

La permutation ainsi déterminée nous fournit des matrices symétriques définies positives qui n'ont plus une structure par blocs mais une disposition des coefficients non nuls qui est celle représentée en figure 1.8 (respectivement pour l'élément 4P1-P0 et pour le 4P1-P1). On remarque que les indices de colonne ont été rangés en ordre croissant pour nous permettre ensuite d'extraire facilement la partie triangulaire inférieure de la

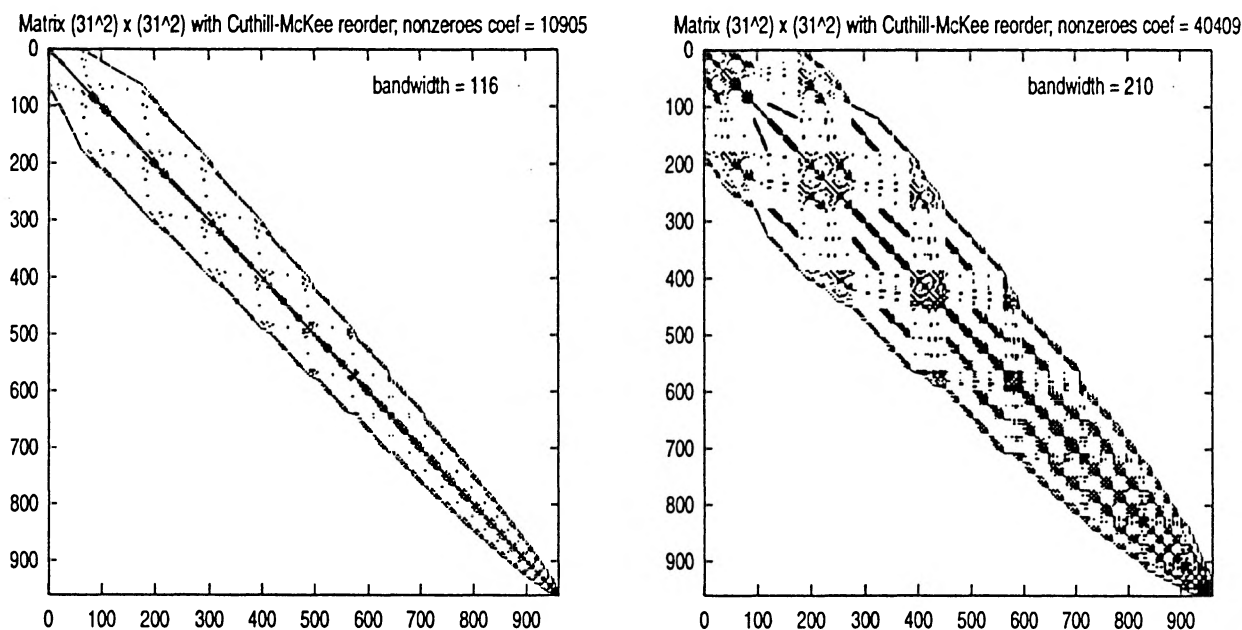


FIG. 1.8 – *Disposition des coefficients non nuls, après la renumérotation de Cuthill-McKee, des matrices de taille $31^2 \times 31^2$: à gauche matrice issue de la discrétisation par 4P1-P0, à droite matrice issue de la discrétisation par 4P1-P1.*

matrice.

La factorisation complète de la partie triangulaire inférieure des matrices avec nouvelle numérotation requiert un stockage qui reste malheureusement de taille mémoire encore très grande. On peut voir en figure 1.9 le remplissage du profil des matrices dans le cas des deux différentes discrétisations par éléments finis : les coefficients non nuls sont presque trois fois plus nombreux pour la matrice relative à la discrétisation 4P1-P1 par rapport à celle relative à la discrétisation 4P1-P0. Ceci nous limitera fortement dans la réduction du pas de discrétisation spatiale h_d du domaine Ω .

La méthode de Cuthill-McKee a été choisie car elle est assez facile à programmer bien qu'elle ne donne pas la permutation σ_k optimale. Une autre méthode qui semble être plus intéressante, surtout lorsque l'on cherche à inverser des matrices de taille supérieure à celle que nous considérons actuellement, est celle du *degré minimal* (que pour le moment nous n'avons pas implémentée). Cette méthode cherche à garder dans la matrice factorisée la structure creuse de la matrice de départ (cf. références dans [16]).

Des tests effectués par C. Jouron (cf. [21]) sur les matrices qui proviennent de la discrétisation par l'élément 4P1-P1 donnent les résultats suivants : pour une matrice de départ de dimension $31^2 \times 31^2$, la matrice factorisée issue de la méthode de Cuthill-McKee possède environ 100000 coefficients non nuls (voir en figure 1.9) par rapport à environ 73000 coefficients non nuls de la matrice factorisée issue de la méthode du degré minimal. Ce gain en taille mémoire devient encore plus intéressant pour une matrice de départ de dimension $63^2 \times 63^2$. En effet, appliquant Cuthill-McKee on trouve une matrice factorisée qui possède un peu plus de 900000 coefficients non nuls, par rapport

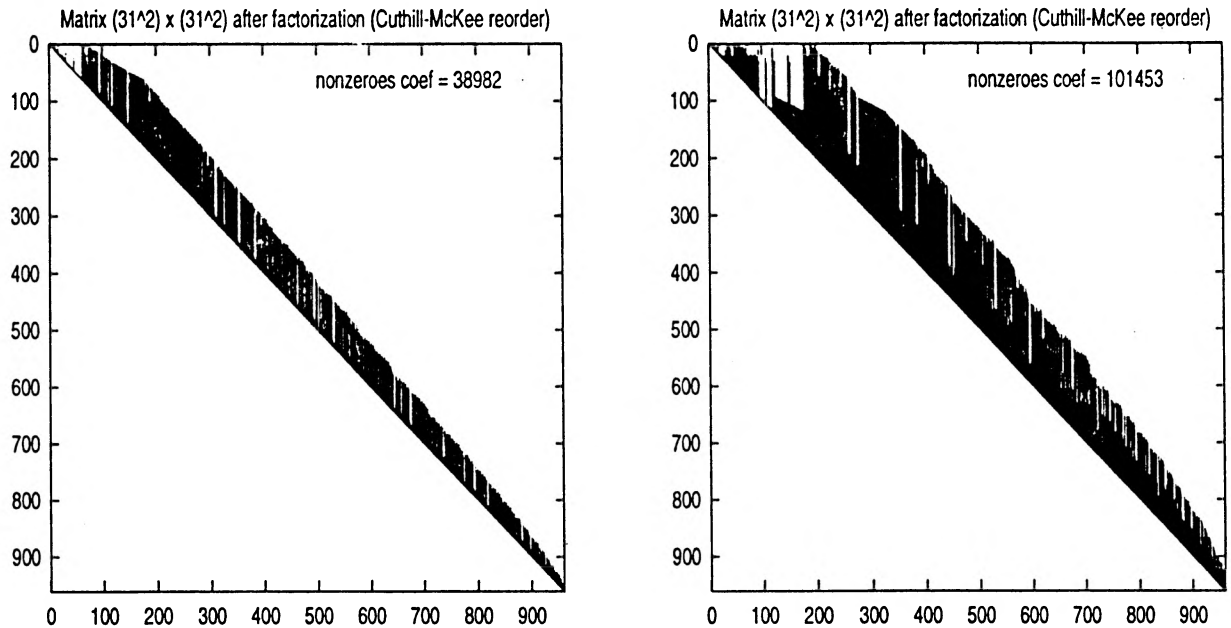


FIG. 1.9 – *Partie triangulaire inférieure des matrices factorisées (de taille $31^2 \times 31^2$) après la renumérotation de Cuthill-McKee : à gauche matrice issue de la discrétisation par $4P1-P0$, à droite matrice issue de la discrétisation par $4P1-P1$.*

à environ 540000 coefficients pour celle déterminée grâce au degré minimal.

La factorisation complète de la matrice qui provient de la permutation de Cuthill-McKee permet d'appliquer une méthode directe pour la résolution du système linéaire (1.28). Ayant calculé et stocké la matrice factorisée au début du programme, il suffit simplement de permuter la numérotation hiérarchique du second membre de (1.28) en fonction de celle de Cuthill-McKee (grâce à la permutation inverse de σ_k), d'effectuer l'inversion de deux matrices triangulaires, et de repasser à la numérotation hiérarchique du vecteur solution grâce à σ_k .

Un dernier test porte sur la nécessité d'envisager cette méthode de résolution directe par rapport à une méthode itérative du type CG ou CR, préconditionnée par une factorisation incomplète de la matrice (mais différente de IC(0)). Au fur et à mesure de la construction ligne par ligne de la matrice factorisée, on choisit d'éliminer les coefficients, en valeur absolue, les plus petits. Ce procédé permet de garder seulement les coefficients les plus significatifs dans la factorisation incomplète, le but étant celui de réduire la taille mémoire de son stockage.

Dans les tableaux 1.3 on exhibe, en fonction d'un paramètre de tolérance donné (appelé *drop-fill coefficient* et noté d_{fill}), le nombre d'itérations pour la convergence de CR préconditionnée par cette factorisation de Cholesky incomplète (notée ICCR), le nombre de coefficients non nuls que l'on stocke et le nombre de ceux éliminés car inférieurs, en valeur absolue, au *drop-fill coefficient*. Le premier tableau 1.3 concerne la résolution du problème de Stokes et le second tableau 1.3 celle de la première itération en temps du problème de Navier-Stokes, ce qui revient à résoudre un problème de Stokes

d_{fill}	Nb itér ICCR	Coef stockés	Coef éliminés
10^{-10}	1	919741	0
10^{-5}	1	919741	0
10^{-4}	4	919421	320
$2 \cdot 10^{-4}$	5	913256	6168
$3 \cdot 10^{-4}$	7	893850	25891
$5 \cdot 10^{-4}$	145	873149	46592
$8 \cdot 10^{-4}$	NC	859440	59981
$5 \cdot 10^{-3}$	NaN	490456	429285

(A)

d_{fill}	Nb itér ICCR	Coef stockés	Coef éliminés
10^{-10}	1	919741	0
10^{-5}	1	919741	0
10^{-4}	3	918836	305
$2 \cdot 10^{-4}$	3	913101	6640
$3 \cdot 10^{-4}$	4	908312	11429
$5 \cdot 10^{-4}$	7	904182	15559
$8 \cdot 10^{-4}$	18	902770	16971

(B)

TAB. 1.3 – En fonction du paramètre d_{fill} : nombre d'itérations pour la convergence de ICCR, nombre de coefficients non nuls stockés et nombre de coefficients non nuls éliminés, car inférieurs en valeur absolue à d_{fill} . (A) : problème de Stokes ; (B) : problème de Navier-Stokes.

généralisé. La discrétisation en espace est celle obtenue par l'élément 4P1-P1 et les valeurs des paramètres des deux problèmes sont les suivantes : $h_d = 1/64$, $Re = 1000$, $\epsilon = 10^{-2}$ et, pour le seul problème de Navier-Stokes, $\Delta t = 0.005$. Le test de convergence de la méthode itérative exige que la norme du résidu relatif soit inférieure à 10^{-8} .

A partir des résultats affichés dans les tableaux 1.3, ainsi que d'autres résultats obtenus avec des choix différents des paramètres du problème de départ, on observe que la procédure d'élimination même d'un petit nombre de coefficients de la matrice factorisée risque de faire échouer la méthode itérative employée pour la résolution du système linéaire (1.28). En effet, le nombre d'itérations pour atteindre la convergence de la méthode ICCR augmente considérablement en fonction du paramètre d_{fill} croissant. Ce phénomène peut s'étendre jusqu'à la non convergence (notée NC) de la méthode itérative, après un nombre maximum d'itérations fixé, ou même jusqu'aux problèmes de *breakdown* (notés NaN) de ICCR. Tout cela peut se produire malgré une très faible réduction de la taille mémoire du stockage de la matrice obtenue par la factorisation incomplète décrite ci-dessus. Ce dernier test permet aussi de justifier, dans une certaine mesure, les résultats négatifs présentés en section 1.3.1.

1.4 Les résultats numériques pour Stokes

Dans cette section nous exposons les résultats des tests préliminaires effectués sur le problème de Stokes et sur celui de Stokes généralisé, en tenant compte de la pénalisation de la contrainte d'incompressibilité. D'abord nous mettons l'accent sur l'influence de la valeur du paramètre ϵ , ceci afin de nous permettre une comparaison avec les conclusions données dans la littérature. Ensuite nous montrons les perturbations de la solution obtenues en discrétisant le problème par l'élément 4P1-P0, instabilités numériques qui vont disparaître lorsqu'on utilise l'élément 4P1-P1. Enfin nous effectuons une brève étude numérique des perturbations détectées, le but étant celui de comprendre les causes de ces instabilités.

1.4.1 Influence de la valeur de ϵ

Pour examiner l'effet du paramètre de pénalisation ϵ sur le champ de vitesse solution du problème de Stokes généralisé, considérons dans le domaine carré unitaire la solution exacte suivante :

$$\begin{aligned} u_1^{ex}(x, y) &= \pi \sin^2(\pi x) \sin(2\pi y) , \\ u_2^{ex}(x, y) &= -\pi \sin(2\pi x) \sin^2(\pi y) , \\ p^{ex}(x, y) &= x^3 y^3 , \end{aligned} \tag{1.30}$$

$u^{ex} = (u_1^{ex}, u_2^{ex})$ étant à divergence nulle.

Le second membre F du système linéaire (1.28) est alors déterminé à partir de la solution exacte donnée, lorsqu'on la remplace dans le système non pénalisé (1.26). On résout le système discret ainsi caractérisé, en employant les discrétisations 4P1-P0 et 4P1-P1 et un choix des paramètres qui est le suivant : $Re = 1000$, $\alpha = 100$ et le pas d'espace

$$h_d = \frac{1}{64} \simeq 0.0156 \quad \text{ou} \quad h_d = \frac{1}{128} \simeq 0.0078 .$$

Afin d'évaluer l'influence de ϵ sur la résolution du problème, nous présentons plusieurs tableaux qui, en fonction du paramètre de pénalisation, donnent le nombre d'itérations pour atteindre la convergence de la méthode CR, l'erreur relative commise et la norme l^∞ de la divergence de la vitesse calculée. Rappelons que l'erreur relative de u , notée e_r , est la suivante :

$$e_r = \frac{|u - u^{ex}|_{l^2}}{|u^{ex}|_{l^2}} ,$$

où

$$|v|_{l^2}^2 = \left(\sum_{i=1}^N |v_i|^2 \right) / N$$

est proportionnelle à la norme l^2 du vecteur v de longueur N , où $v_i = (v_i^1, v_i^2) \in \mathbf{R}^2$ et $|v_i|^2 = ((v_i^1)^2 + (v_i^2)^2)$. De façon similaire la norme l^∞ est définie par :

$$|v|_\infty = \sup_{1 \leq i \leq N} |v_i| .$$

	$h_d = 1/64 \quad (h_d^2 \simeq 2.4 \times 10^{-4})$			$h_d = 1/128 \quad (h_d^2 \simeq 6.1 \times 10^{-5})$		
ϵ	Nb itér CR	e_r	$ div u _\infty$	Nb itér CR	e_r	$ div u _\infty$
1	109	0.0037499	0.23×10^{-3}	150	0.00274	0.56×10^{-4}
10^{-1}	234	0.0012586	0.36×10^{-4}	354	0.000903	0.93×10^{-5}
10^{-2}	280	0.0007418	0.40×10^{-5}	495	0.000224	$\simeq 10^{-6}$
10^{-3}	387	0.0007301	0.41×10^{-6}	526	0.000186	$\simeq 10^{-7}$
10^{-4}	443	0.0007299	0.41×10^{-7}	715	0.000186	$\simeq 10^{-8}$
10^{-5}	531	0.0007299	0.41×10^{-8}	850	0.000186	$\simeq 10^{-9}$
10^{-6}	600	0.0007299	0.41×10^{-9}	1001	0.000186	$\simeq 10^{-10}$

TAB. 1.4 – *Discrétisation par 4P1-P0: nombre d'itérations pour la convergence de CR non préconditionnée (norme l^2 du résidu inférieure à 10^{-7}), erreur relative e_r et divergence de la vitesse calculée, tout en fonction du paramètre ϵ .*

Les résultats des tableaux 1.4, 1.5 et 1.6 ont été établis pour le CR non préconditionné, afin de comparer également le nombre d'itérations pour atteindre la convergence de cette méthode itérative en fonction du paramètre ϵ . Quand on utilise la méthode de résolution directe, les variations de l'erreur relative et de la divergence discrète sont comparables à celles que nous allons présenter pour le CR non préconditionné. Bien évidemment la méthode CR préconditionnée par la factorisation complète converge en une seule itération, avec un résidu relatif de norme très petite, mais qui dépend de l' ϵ choisi. Par exemple, dans le cas du 4P1-P1 avec $\mathcal{C}$ approchée par son *mass lumping*, le résidu relatif est d'ordre 10^{-13} pour $\epsilon = 10^{-1}$, d'ordre 10^{-11} pour $\epsilon = 10^{-3}$ mais il croît jusqu'à 10^{-5} pour $\epsilon = 10^{-6}$ (dans ce dernier cas il faut deux itérations de CR préconditionné par la factorisation complète pour obtenir un résidu relatif d'ordre 10^{-9}). La croissance de la norme l^2 du résidu relatif est due au conditionnement de plus en plus mauvais de la matrice à inverser.

Le tableau 1.4 concerne les résultats numériques relatifs à la discrétisation par l'élément 4P1-P0: les trois premières colonnes sont relatives au pas d'espace uniforme $h_d = \frac{1}{64}$, les autres au pas $h_d = \frac{1}{128}$. Indépendamment du pas d'espace, il suffit que le paramètre ϵ soit de l'ordre de 10^{-2} pour obtenir une erreur relative satisfaisante et une solution ayant une divergence très petite. Remarquons que le paramètre h_d est intrinsèque aux valeurs de la matrice diagonale $\mathcal{C}$.

Dans le cas de la discrétisation par l'élément 4P1-P1 avec un pas d'espace uniforme $h_d = \frac{1}{64}$, on donne dans le tableau 1.5 le nombre d'itérations pour la convergence de CR, l'erreur relative commise et la norme l^∞ de la divergence de la solution, tout cela en suivant deux approches différentes. Dans l'une des simulations on discrétise le problème de Stokes pénalisé (voir (1.27)), donc on considère le *mass lumping* de la matrice $\mathcal{C}$; dans l'autre simulation on pénalise le problème discret (1.26) ($\mathcal{C} = Id$, matrice identité). On remarque la différence de valeurs du paramètre ϵ en suivant les deux manières d'approcher la matrice $\mathcal{C}$. Dans le premier cas il suffit de prendre un ϵ de l'ordre de 10^{-2} , comme déjà constaté pour la discrétisation par 4P1-P0, car les valeurs

	$\mathcal{C}$: <i>mass lumping</i>			$\mathcal{C} = Id$		
ϵ	Nb itér CR	e_r	$ \operatorname{div} u _\infty$	Nb itér CR	e_r	$ \operatorname{div} u _\infty$
10^{-1}	114	0.001041	$\simeq 10^{-4}$			
10^{-2}	161	0.000382	$\simeq 10^{-5}$	14	0.00342	
10^{-3}	187	0.000348	$\simeq 10^{-6}$	40	0.00263	$\simeq 10^{-3}$
10^{-4}	250	0.000348	$\simeq 10^{-7}$	119	0.00117	$\simeq 10^{-4}$
10^{-5}	286	0.000348	$\simeq 10^{-8}$	173	0.000395	$\simeq 10^{-5}$
10^{-6}	336	0.000348	$\simeq 10^{-9}$	202	0.000348	$\simeq 10^{-6}$
10^{-7}				271	0.000348	$\simeq 10^{-7}$
10^{-8}				309	0.000348	$\simeq 10^{-8}$

TAB. 1.5 – *Discrétisation par 4P1-P1 : nombre d'itérations pour la convergence de CR non préconditionnée (norme l^2 du résidu relatif inférieure à 10^{-6}), erreur relative e_r et divergence de la vitesse calculée, tout en fonction du paramètre ϵ .*

de la matrice diagonale obtenue par le *mass lumping* sont de l'ordre de h_d^2 . En revanche, dans le deuxième cas il faudra choisir ϵ de l'ordre de $10^{-2}h_d^2$; ceci est totalement en accord avec les résultats numériques que l'on trouve dans la littérature (cf. par exemple [2] et [3], où d'abord on pénalisait le problème continu et ensuite on le discrétisait).

Dans la suite nous serons plutôt intéressés par des maillages qui présentent un resserrement des mailles au voisinage des bords. Ce sera le cas des simulations à nombre de Reynolds élevé. En effet l'encombrement mémoire, dû à la factorisation complète de la matrice à inverser, nous empêchera de considérer des pas d'espaces trop petits. Pour cela, afin d'obtenir une approximation numérique relativement correcte, nous raffinerons le maillage près des bords du domaine, là où le phénomène des couches limites va apparaître. A la suite de ces motivations, la discrétisation dans les deux directions spatiales est obtenue par exemple à l'aide de la fonction $f(\zeta_i) = (1 - \cos \zeta_i)$, où $\zeta_i = \frac{\pi i}{2N}$, $N = 64$ étant le nombre de segments dans chaque direction. Même pour ce type de maillages raffinés nous pouvons tirer des conclusions analogues à celles données précédemment, obtenues pour le 4P1-P1 à partir des maillages uniformes (voir tableau 1.6). Cette fois-ci le choix des paramètres du problème de Stokes généralisé est le suivant : $Re = 5000$ et $\alpha = 200$. Pour des maillages pareils, il faudra plutôt choisir un paramètre de pénalisation ϵ de l'ordre de 10^{-3} (ou $\frac{10^{-3}}{N^2}$ si $\mathcal{C}$ est l'identité). De toute manière l'approximation de la matrice de masse $\mathcal{C}$ par son *mass lumping* est préférable à l'identification $\mathcal{C} = Id$, ceci à cause des variations très importantes du pas de discrétisation spatiale.

1.4.2 Perturbations dans le champ de vitesse

Examinons à présent la discrétisation du problème de Stokes pénalisé (voir (1.27)) qui nous amène au système linéaire (1.28). Le problème modèle est celui de la cavité entraînée (en anglais *wall driven forced cavity*). Le domaine est un carré ($\Omega = [0, 1]^2$) et les forces extérieures f sont nulles; le mouvement du fluide est dû à la distribution

	C : mass lumping		$C = Id$	
ϵ	Nb itér CR	e_r	Nb itér CR	e_r
10^{-1}	346	0.0009652		
10^{-2}	653	0.0003369	42	0.0033355
10^{-3}	934	0.0002854	110	0.0030091
10^{-4}	1193	0.0002851	258	0.0020189
10^{-5}	1457	0.0002851	639	0.0005986
10^{-6}	1717	0.0002851	1147	0.0002916
10^{-7}			1672	0.0002851
10^{-8}			2120	0.0002851

TAB. 1.6 – *Discrétisation par 4P1-P1 avec maillage raffiné près des bords: nombre d'itérations pour la convergence de CR (norme l^2 du résidu relatif inférieure à 10^{-6}) et erreur relative e_r en fonction du paramètre ϵ .*

de la vitesse sur le bord supérieur de Ω . Les conditions aux limites sont celles de la cavité entraînée non régularisée, notée CANR: $u(x, t) = 0$ sur tout $\partial\Omega$ sauf sur le bord supérieur $x_2 = 1$ où

$$u(x, t) = \begin{pmatrix} u_1(x, t) \\ u_2(x, t) \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \end{pmatrix}.$$

Nous envisageons aussi les conditions aux limites de la cavité entraînée régularisée, notée CARE: $u(x, t) = 0$ sur tout $\partial\Omega$ sauf sur le bord $x_2 = 1$ où

$$u(x, t) = \begin{pmatrix} u_1(x, t) \\ u_2(x, t) \end{pmatrix} = \begin{pmatrix} (1 - (2x_1 - 1)^2)^2 \\ 0 \end{pmatrix}.$$

Le problème de la cavité entraînée n'a aucun sens physique mais il est intéressant du point de vue numérique à cause de la difficulté d'approcher correctement le fort gradient de vitesse qui est présent dans les couches supérieures du maillage, lorsque le paramètre α devient grand. Le choix de CARE au lieu de CANR permet de régulariser l'écoulement dans la cavité en éliminant les singularités dans les deux coins supérieurs de Ω dues à la discontinuité des conditions au bord.

Nous traçons les champs de vitesse solutions du problème de Stokes généralisé dans la cavité CANR, pour les discrétisations 4P1-P0 et 4P1-P1. Les paramètres Re et ϵ sont fixés respectivement à 1000 et 10^{-2} et le pas d'espace est uniforme ($h_d = \frac{1}{64}$). L'influence du paramètre α sur le champ de vitesse est illustrée en figure (1.10) pour l'élément 4P1-P0 et en figure (1.11) pour le 4P1-P1. Le zoom dans le coin supérieur gauche du domaine permet d'observer aisément les perturbations dans le champ de vitesse pour l'élément 4P1-P0, celles-ci paraissant même lorsque le paramètre α est très petit. Le phénomène des instabilités numériques se révèle par l'orientation des flèches représentant le champ de vitesse calculé: celles-ci ont tendance à se croiser lorsqu'on regarde la solution près de la zone perturbée. En revanche dans le cas d'une discrétisation spatiale par l'élément

4P1-P1 ces perturbations ne sont présentes que si α devient trop grand (voir les quelques petites oscillations près du bord supérieur dans le champ de vitesse tracé pour $\alpha = 50$).

De toute manière les singularités aux deux coins supérieurs de Ω ne jouent aucun rôle sur la présence des perturbations existantes dans la discrétisation 4P1-P0. En fait ces oscillations existent même lorsque nous simulons l'écoulement plus régulier de la cavité CARE, écoulement qui ne présente plus les discontinuités des conditions aux limites. C'est donc bien la méthode de discrétisation qui est en cause.

Remarque 8 Lorsque l'on regarde les résultats numériques de la littérature (voir [3] par exemple), on s'aperçoit que le champs de vitesse solution de Navier-Stokes présente des oscillations près des zones où le gradient de la vitesse est important. Cependant il est difficile de détecter les causes qui produisent ces perturbations, la raison la plus probable étant le pas de discrétisation du maillage qui est trop grossier.

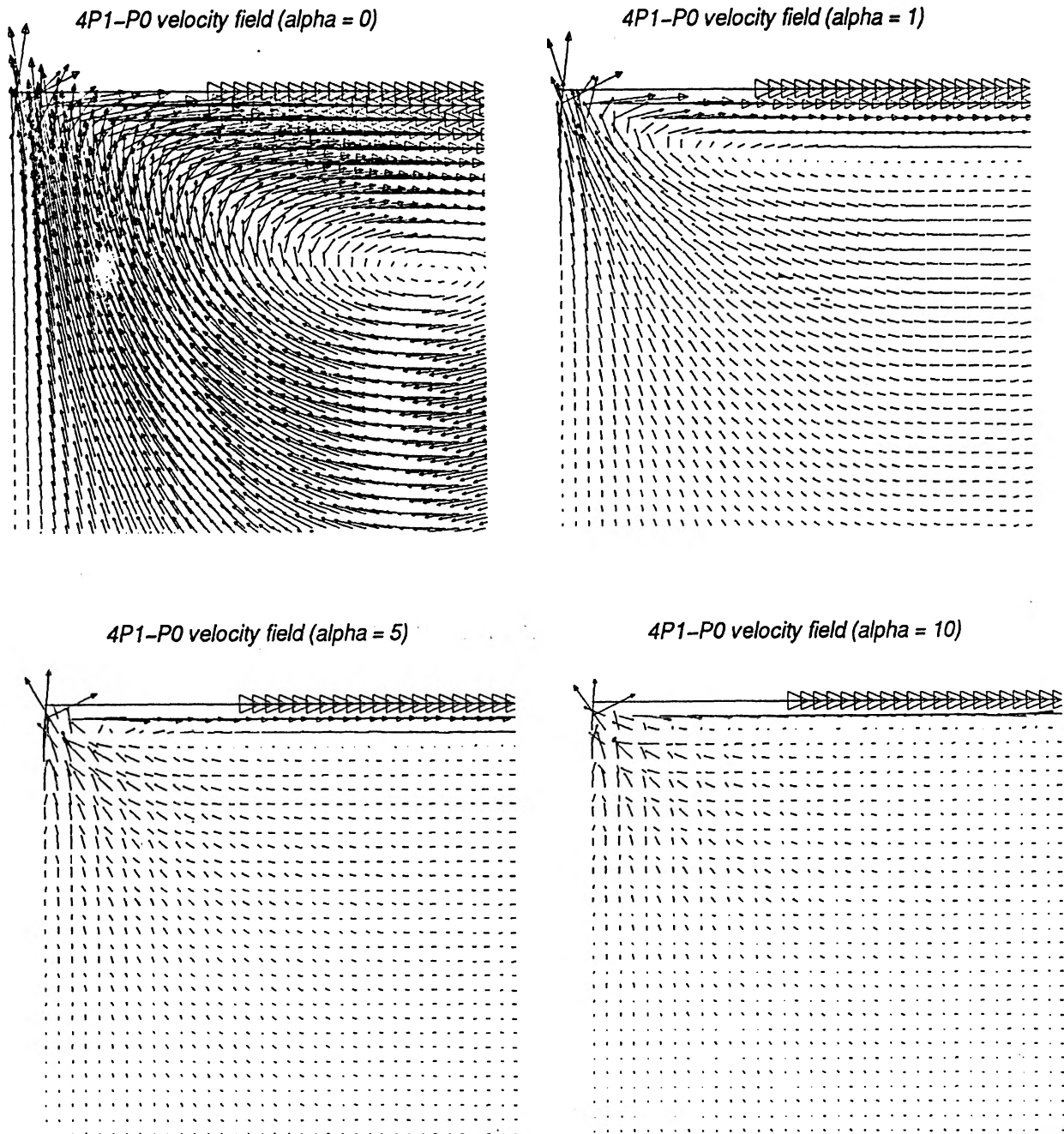


FIG. 1.10 – Problème de Stokes généralisé dans la cavité CANR: discrétisation par l'élément 4P1-P0. Champ de vitesse pour $Re = 1000$ et $\alpha = 0.0, 1.0, 5.0, 10.0$.

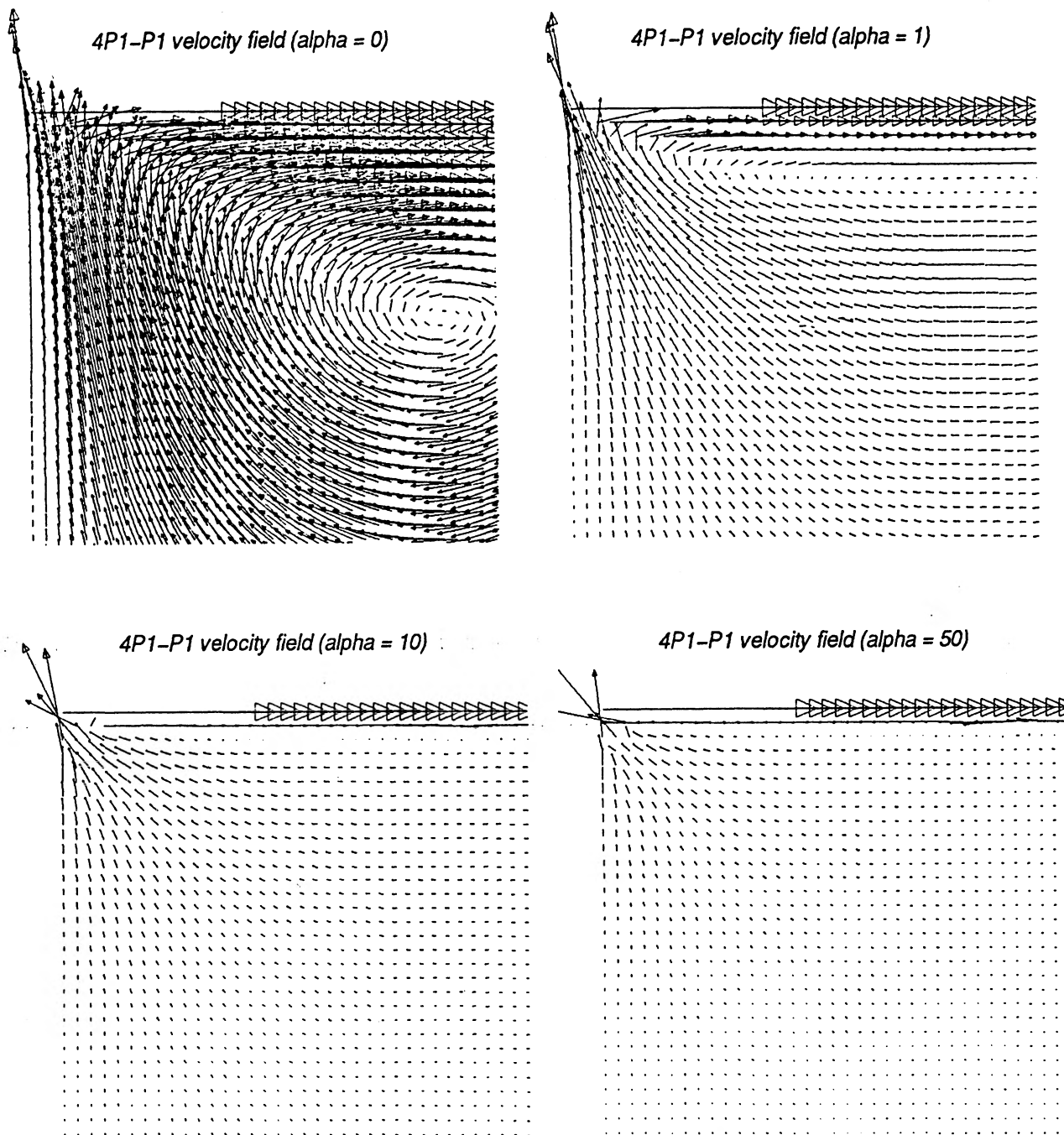


FIG. 1.11 – Problème de Stokes généralisé dans la cavité CANR : discrétisation par l'élément 4P1-P1. Champ de vitesse pour $Re = 1000$ et $\alpha = 0.0, 1.0, 10.0, 50.0$.

	$\Omega_1 = [0, 1]^2$		$\Omega_2 = [0.5, 1.5]^2$	
$\alpha \backslash e_r$	4P1-P0	4P1-P1	4P1-P0	4P1-P1
0	0.03122	0.004188	0.3594	0.05006
1	0.02820	0.004870	0.2973	0.004682
5	0.02227	0.007524	0.2298	0.004125
10	0.01803	0.01005	0.1804	0.004148
50	0.01178	0.01991	0.06732	0.004513
10^2	0.01259	0.02429	0.03800	0.004869
10^3	0.01643	0.03152	0.006964	0.005622
10^4	0.01725	0.03260	0.005801	0.005714

TAB. 1.7 – Erreur relative en fonction du paramètre α : discrétisation par l'élément 4P1-P0 et par le 4P1-P1.

1.4.3 Etude numérique des perturbations

L'objet de ce paragraphe est la compréhension des causes qui engendrent les perturbations décelées auparavant dans le champ de vitesse. En effet, nous voulons analyser si leur origine est liée à la technique de pénalisation, au choix des 'éléments finis mixtes employés pour discrétiser en espace le problème de Stokes ou encore au choix des différents paramètres du problème considéré, soit α , Re et h_d .

La série de tests suivante examine l'approximation du problème de Stokes (1.16) ayant pour second membre la fonction analytique f calculée à partir de la solution exacte (1.30). Nous ferons varier les différents paramètres afin d'exhiber dans les tableaux 1.7, 1.8 et 1.9 l'erreur relative e_r , commise en discrétisant le problème respectivement par 4P1-P0 et 4P1-P1 (le paramètre ϵ est fixé à la valeur 10^{-2}). En résumé, on observe une apparition ou une amplification des oscillations dans la solution calculée pour des valeurs croissantes de Re ou α et pour des maillages de plus en plus grossiers. De toute façon les instabilités présentes dans le champ de vitesse sont largement plus intenses pour l'élément 4P1-P0 que pour le 4P1-P1. Nous pouvons déjà en déduire que ces perturbations ne proviennent pas de la technique de pénalisation employée pour résoudre le problème analytique.

En revanche, elles semblent liées au choix de l'élément fini mixte. Effectivement des résultats comparables aux nôtres ont été obtenus par F. Pascal (cf. [28]) qui utilise l'algorithme d'Uzawa comme méthode de résolution. Les discrétisations considérées sont également celles des éléments 4P1-P0 et 4P1-P1, mais des vérifications supplémentaires se basaient sur les éléments P2-P0 et P2-P1.

Le tableau 1.7 établit la variation de l'erreur relative e_r en fonction du paramètre α , lorsque $Re = 1000$ et $h_d = \frac{1}{64}$. Les deux premières colonnes donnent l'erreur relative du problème analytique dans le domaine $\Omega = [0, 1]^2$ (dans ce domaine $u^x = 0$ sur $\partial\Omega$) et les deux dernières l'erreur dans le domaine $\Omega = [0.5, 1.5]^2$ (dans ce cas la solution exacte est non nulle sur le bord et s'annule le long des médianes). Dans le carré unitaire,

$\frac{e_r}{h_d}$	4P1-P0	4P1-P1
1/16	0.088869	0.10467
1/32	0.021367	0.021701
1/64	0.005554	0.004866
1/128	0.001462	0.001183

TAB. 1.8 – Erreur relative en fonction du pas d'espace h_d : discrétisation par l'élément 4P1-P0 et par le 4P1-P1.

les perturbations se manifestent, pour le choix du 4P1-P0, dans toute la zone proche du coin supérieur droit, même pour des α très petits ou nuls. De plus, quand α devient grand on peut observer, en regardant le champ de vitesse normalisé, que les oscillations apparaissent près des bords et aux quatre coins, sans distinction pour les deux discrétisations. Dans le domaine $[0.5, 1.5]^2$ la solution exacte vérifie des conditions aux limites de Dirichlet non homogènes. Lorsque α est très grand, l'erreur relative est petite pour les deux discrétisations. En contrepartie, les perturbations se manifestent de plus en plus avec α décroissant (dans ce cas, les oscillations dans le champ de vitesse se localisent plutôt le long des médianes, où la vitesse est nulle). Pour l'élément 4P1-P0, plus α est petit, plus la vitesse calculée est différente de la solution exacte. Cependant on décelez une grande erreur relative même pour 4P1-P1 si on choisit $\alpha = 0$.

Dorénavant le domaine Ω sera le carré unitaire. Il est bien connu que la différence entre la solution exacte et la solution calculée est, en norme L^2 , de l'ordre de h_d^2 . Dans le tableau 1.8 on donne l'erreur relative de u en fonction du pas de discrétisation spatiale pour les deux choix d'éléments finis, les valeurs des paramètres étant les suivantes : $\alpha = 10$ et $Re = 100$. Pour l'élément 4P1-P1 la vitesse est calculée correctement à partir de $h_d = \frac{1}{64}$ (voir le champ normalisé et non), des légères oscillations restant visibles dans les quatre coins du champ normalisé relatif au pas $h_d = \frac{1}{32}$. En revanche, si on considère l'élément 4P1-P0, bien que l'erreur relative ne le montre pas, les perturbations dans le champ de vitesse calculé restent évidentes dans toute la zone proche du coin supérieur droit. Elles sont très visibles pour $h_d = \frac{1}{32}$ (dans ce cas on peut observer des oscillations aussi dans les autres coins du domaine, même pour le champ non normalisé). Ce phénomène reste présent dans les champs normalisés relatifs aux pas $h_d = \frac{1}{64}$ et $h_d = \frac{1}{128}$. La figure (1.12) permet d'apprécier ces oscillations et de comparer le champ obtenu pour le 4P1-P0 à celui du 4P1-P1.

Enfin dans le tableau 1.9 on montre l'erreur relative du problème de Stokes ($\alpha = 0$ et $h_d = \frac{1}{64}$) en fonction du nombre de Reynolds Re . L'apparition des perturbations est liée à la croissance du paramètre Re . Notamment pour l'élément 4P1-P0 les instabilités se manifestent déjà dans le champ normalisé relatif à $Re = 100$ et, de manière plus faible, à $Re = 10$. Plus le nombre de Reynolds augmente plus ces oscillations se propagent vers l'intérieur du domaine et plus elles s'intensifient, devenant observables même dans le champ non normalisé. Dans la discrétisation par 4P1-P1 ce phénomène ressort pour des Re beaucoup plus grands (légères perturbations, visibles seulement dans le champ

$\frac{e_r}{Re}$	4P1-P0	4P1-P1
1	0.004396	0.004157
10	0.004421	0.004158
100	0.005428	0.004159
500	0.01611	0.004165
1000	0.03122	0.004185
5000		0.004773
10000		0.006267

TAB. 1.9 – Erreur relative en fonction du nombre de Reynolds Re : discrétisation par l'élément 4P1-P0 et par le 4P1-P1.

normalisé, à partir de $Re = 5000$).

Remarque 9

1. Pour éliminer les perturbations qui se manifestent dans le champ de vitesse lorsque α est petit et Re est grand, il semblerait plus adéquat de prendre en compte des conditions de type $u \cdot n = 0$ sur $\partial\Omega$ car, pour un pareil choix des paramètres, le problème examiné approche de plus en plus un problème d'Euler incompressible (qui est bien posé si on tient compte des conditions aux limites de glissement). Cependant, des tests supplémentaires effectués par F. Pascal ne montrent aucune amélioration de l'erreur relative même lorsqu'on résout le problème de type Euler avec $u \cdot n = 0$ sur le bord.
2. Le problème approché est celui de Stokes généralisé sur le carré unitaire, avec par exemple $\alpha = 1$ et $Re = 5000$. Les instabilités de la vitesse, présentes principalement pour le champ discrétisé par 4P1-P0 (elles sont pratiquement absentes pour le 4P1-P1), semblent plutôt liées à la pression exacte considérée. Effectivement, pour u^{ex} et p^{ex} données par (1.30), l'erreur relative est de l'ordre de 10^{-1} ; en revanche si on prend p^{ex} nulle, la valeur de e_r descend à peu près à 0.006. A partir des tests effectués par F. Pascal (cf. [28]), aucun lien peut être établi entre l'apparition des perturbations et le problème des conditions aux limites sur la pression (rappelons qu'en utilisant l'algorithme d'Uzawa on se ramène à résoudre un système linéaire en pression). De plus, aucune condition aux limites ne nécessite d'être imposée sur la pression, lorsqu'on résout notre problème par la méthode de pénalisation. Même le gradient de la pression ne semble avoir qu'un rôle marginal dans l'apparition et dans l'intensité des perturbations détectées.

Enfin, les résultats numériques obtenus par F. Pascal (cf. [28]) établissent que les perturbations sont dues au choix conjoint des paramètres Re et α . En effet, pour des Re grands et des α de l'ordre de l'unité, l'erreur sur la vitesse devient de plus en plus visible, même si la pression continue à être approchée correctement. Ceci

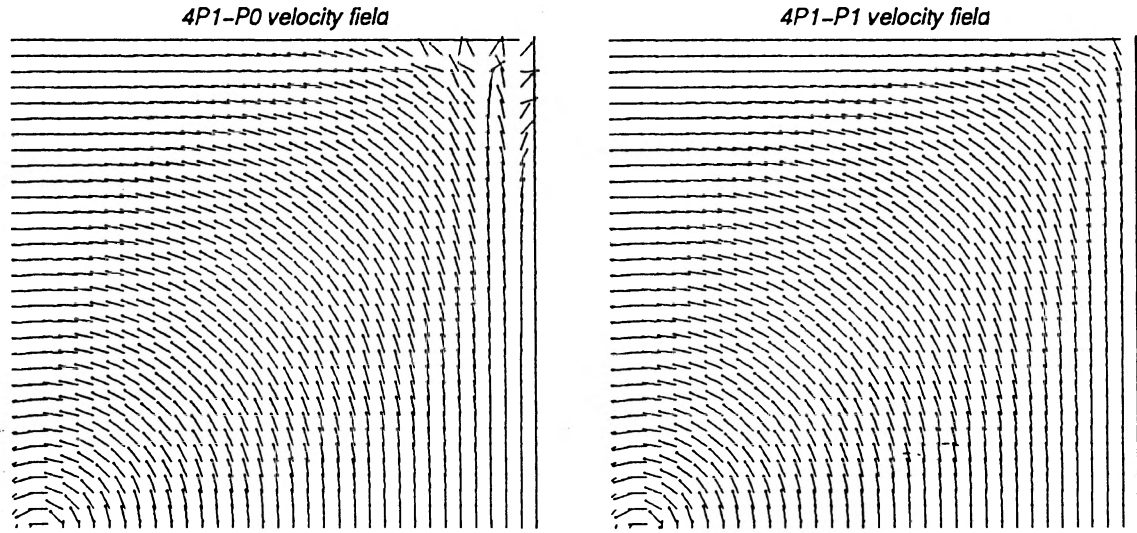


FIG. 1.12 – Problème de Stokes généralisé avec second membre analytique. Champ de vitesse normalisée pour $\alpha = 10$, $Re = 100$ et $h_d = 1/64$. A gauche discrétisation par 4P1-P0 et à droite par 4P1-P1.

est dû à la matrice que l'on inverse pour calculer le champ de vitesse, matrice qui s'approche de plus en plus de la matrice de masse. Les résultats numériques obtenus sont confirmés par des explications théoriques nous données par F. Brezzi, dans une communication privée. Les perturbations présentes dans le champ de vitesse du problème de Stokes généralisé sont engendrées par la perte du caractère elliptique de la forme bilinéaire $a_h(u, v)$ sur le $\text{Ker}(B_h)$, car la forme bilinéaire ressemble de plus en plus au produit scalaire dans $L^2(\Omega)$ plutôt qu'à celui dans $H_0^1(\Omega)$.

Afin de résoudre le problème de Navier-Stokes, nous approchons dans un premier temps les équations de Stokes (avec $\alpha = 0$) dans la cavité entraînée. L'élément 4P1-P0, comme aussi le 4P1-P1, donne une approximation suffisamment correcte de la solution de Stokes. Les résultats obtenus pour les deux discrétisations sont comparables. Ces résultats concernent la position du tourbillon principale et du contre-tourbillon en bas à gauche, les extrema de la vitesse le long des médianes, les extrema de la vorticit   ou encore la sym  trie de la pression par rapport    la m  diane verticale. La solution du probl  me de Stokes est choisie comme solution initiale du probl  me de Navier-Stokes, m  me si d'autres initialisations sont possibles. En effet, le probl  me de Navier-Stokes p  nalis   pourrait   tre initialis   simplement par la solution nulle    l'int  rieur de Ω et $u_0(x) = g(x)$ sur $\partial\Omega$. Cependant, une telle initialisation demande un nombre d'it  rations en temps bien plus important pour atteindre la convergence vers la solution stationnaire. De plus, tous les probl  mes observ  s en section 1.4.2 vont r  appara  tre. A partir de la solution de Stokes et seulement apr  s quelques it  rations en temps pour Navier-Stokes, on note une   volution de l'  coulement qui pr  sente de plus en plus de perturbations dans le champ de vitesse, ceci seulement quand on consid  re la discr  tisation par 4P1-P0. L'apparition des instabilit  s dans le champ de vitesse se situe en proximit   du coin

supérieur droit. Elles se propagent ensuite vers l'intérieur du domaine, suivant l'évolution des lignes d'isovorticité.

C'est ainsi, à la suite de ces dernières remarques et des analyses précédentes, relatives à l'approximation de la solution de Stokes, que nous avons décidé d'abandonner de manière définitive l'élément 4P1-P0 pour garder seulement le choix du 4P1-P1.

1.5 Résolution du problème de Navier-Stokes

1.5.1 Discrétisation en espace et en temps

Formulation discrète du problème

La version discrète du problème de Navier-Stokes pénalisé (1.21), issue de l'approximation par éléments finis mixtes 4P1-P1, est la suivante :

$$\left\{ \begin{array}{l} \text{Trouver } (u_d(t), p_d(t)) \in \mathcal{V}_{g,d} \times \mathcal{Q}_d \text{ tel que} \\ \left(\frac{\partial u_d}{\partial t}, v_d \right) + \frac{1}{Re} ((u_d, v_d)) + \hat{b}(u_d, u_d, v_d) - (p_d, \text{div } v_d) = (f, v_d) \quad \forall v_d \in \mathcal{V}_{0,d} \\ (\text{div } u_d, q_d) + \epsilon(p_d, q_d) = 0 \quad \forall q_d \in \mathcal{Q}_d . \end{array} \right. \quad (1.31)$$

Schéma en temps

La dérivée temporelle est approchée par un schéma aux différences finies. La discrétisation en temps relative à (1.31) est composée :

- d'un schéma d'Euler (implicite du premier ordre) ou de Crank-Nicholson (semi-implicite du second ordre) pour les termes linéaires ;
- d'un schéma d'Euler (explicite d'ordre un) ou d'Adams-Bashforth (explicite d'ordre deux) pour les termes non linéaires.

1. Schéma Euler - Euler (noté E-E)

Si l'on connaît u_d^n au temps $n\Delta t$ ($n \geq 0$), u_d^{n+1} et p_d^{n+1} sont les approximations au temps $(n+1)\Delta t$ de u_d et p_d définies par :

$$\left\{ \begin{array}{l} \left(\frac{u_d^{n+1} - u_d^n}{\Delta t}, v_d \right) + \frac{1}{Re} ((u_d^{n+1}, v_d)) - (p_d^{n+1}, \text{div } v_d) = \\ \quad (f^{n+1}, v_d) - \hat{b}(u_d^n, u_d^n, v_d) \quad \forall v_d \in \mathcal{V}_{0,d} \\ (\text{div } u_d^{n+1}, q_d) + \epsilon(p_d^{n+1}, q_d) = 0 \quad \forall q_d \in \mathcal{Q}_d . \end{array} \right. \quad (1.32)$$

2. Schéma Crank-Nicholson - Adams-Bashforth (noté CN-AB)

Après une initialisation effectuée par une méthode d'Euler (voir schéma E-E), u_d^{n+1} et p_d^{n+1} ($n \geq 1$) sont définies à partir de u_d^n et u_d^{n-1} par :

$$\left\{ \begin{array}{l} \left(\frac{u_d^{n+1} - u_d^n}{\Delta t}, v_d \right) + \frac{1}{2Re} ((u_d^{n+1} + u_d^n, v_d)) - (p_d^{n+1}, \operatorname{div} v_d) = \\ \quad (f^{n+1}, v_d) - \left[\frac{3}{2} \hat{b}(u_d^n, u_d^n, v_d) - \frac{1}{2} \hat{b}(u_d^{n-1}, u_d^{n-1}, v_d) \right] \quad \forall v_d \in \mathcal{V}_{0,d} \\ \\ (\operatorname{div} u_d^{n+1}, q_d) + \epsilon (p_d^{n+1}, q_d) = 0 \quad \forall q_d \in \mathcal{Q}_d. \end{array} \right. \quad (1.33)$$

Remarque 10 Le terme $(p_d^{n+1}, \operatorname{div} v_d)$ dans le schéma CN-AB peut-être remplacé par $(\frac{p_d^{n+1} + p_d^n}{2}, \operatorname{div} v_d)$. Cela revient à discrétiser aussi la pression par le schéma de Crank-Nicholson, comme déjà fait pour les termes linéaires en vitesse. Les résultats numériques obtenus avec le schéma CN-AB modifié sont équivalents à ceux du schéma (1.33). De toute manière cette modification est peu intéressante car à chaque itération en temps on ajoute le calcul d'un produit matrice-vecteur supplémentaire.

Ces deux schémas, respectivement du premier et du second ordre, ont l'inconvénient d'être conditionnellement stables. Il est nécessaire qu'une condition CFL du type :

$$|u_d|_\infty \cdot \frac{\Delta t}{h_d} \leq cst$$

soit vérifiée, car le terme non linéaire de transport est traité explicitement. Cela rend ces schémas coûteux en temps de calcul dès que le nombre de Reynolds est grand car, dans ce cas, on doit choisir h_d petit. Cependant on envisage ces schémas temporels, plutôt que des schémas implicites ou de splitting, en vue d'une méthode multi-niveaux en espace et en temps similaire à celle déjà appliquée pour résoudre les équations non linéaires de Burgers.

Il est enfin à noter que chacune des étapes de ces deux algorithmes se ramène à un problème de Stokes généralisé. La pression peut ainsi être éliminée par le procédé de pénalisation qui permet d'obtenir un système en vitesse comme unique inconnue. La nécessité d'avoir un solveur performant de Stokes, tel qu'une méthode de résolution directe, est donc évidente.

1.5.2 Résultats numériques

Nous comparons nos résultats avec ceux obtenus par d'autres auteurs (cf. [18], [4], [27], [17] et [29] pour la cavité entraînée non régularisée (i.e. CANR) et [30], [27] et [29] pour la cavité régularisée (i.e. CARE)).

La figure 1.13 présente la configuration de référence ainsi que la nomenclature employée afin d'exposer nos résultats numériques. Comme pour les autres auteurs, l'étude

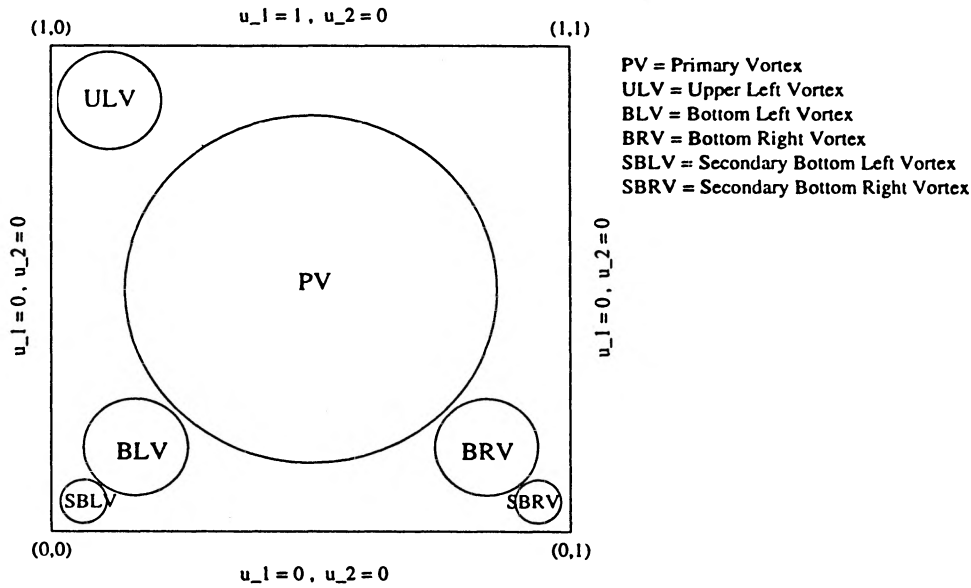


FIG. 1.13 – Configuration de la cavité entraînée.

qualitative et quantitative du problème modèle se base sur différentes quantités et en particulier sur les outils suivant :

- Le profil de la vitesse horizontale le long de la médiane verticale $u_1(0.5, .)$, le profil médian horizontal de la vitesse verticale $u_2(., 0.5)$ et les extrema de $u_1(0.5, .)$ et de $u_2(., 0.5)$ (notés $\min(u_1)$, $\min(u_2)$ et $\max(u_2)$) avec leur positions respectives ($y_{\min}$, $x_{\min}$ et $x_{\max}$).
- Les isobares qui permettent de relever des singularités ; celles-ci sont tracées après avoir calculé la pression grâce à l'équation de continuité pénalisée.
- L'isovorticité, soit les lignes de niveaux de la fonction $\omega(x, t)$, le rotationnel de la vitesse. Pour un écoulement bidimensionnel la vorticité est parallèle à l'axe Ox_3 ; elle est définie par :

$$\omega(x, t) = \frac{\partial u_2}{\partial x_1}(x, t) - \frac{\partial u_1}{\partial x_2}(x, t)$$

et calculée par l'inversion d'une matrice de masse.

- Les isovaleurs de la fonction de courant $\psi(x, t)$, qui est une fonction scalaire définie par :

$$u_1(x, t) = \frac{\partial \psi}{\partial x_2}(x, t) \quad \text{et} \quad u_2(x, t) = -\frac{\partial \psi}{\partial x_1}(x, t)$$

et calculée en résolvant un problème de Poisson avec conditions aux limites de type Dirichlet homogène.

Les calculs ont été réalisés à $Re = 1000$ et 5000 pour la cavité entraînée non régularisée comme aussi pour celle régularisée. Pour de tels nombres de Reynolds, la solution du

problème de Navier-Stokes dans la cavité entraînée évolue vers une solution stationnaire. Le critère de convergence est le suivant :

$$\frac{|u_d^{l,n+1} - u_d^{l,n}|_{l^2}}{|u_d^{l,n+1}|_{l^2}} \leq 10^{-5} \quad \text{pour } l = 1 \text{ et } 2 .$$

Les résultats ont été obtenus pour des maillages plus resserrés près des bords (présentés à la page 86) et un nombre de points de discrétisation équivalent à 65 dans chaque direction spatiale.

Les schémas en temps E-E ou CN-AB donnent des résultats comparables. Le schéma du premier ordre E-E permet de capter correctement la solution stationnaire, bien qu'avec le schéma CN-AB d'ordre deux la convergence vers cette solution semble être plus rapide. Cependant, dans nos simulations numériques nous appliquons le schéma E-E, surtout en prévision d'une future implémentation de la méthode multi-résolution en espace et en temps (en effet on avait déjà remarqué, pendant la résolution des équations de Burgers, qu'un schéma multi-pas en temps engendre, lors des V-cycles, trop de perturbations dans la solution calculée).

Les calculs ont été effectués sur des machines séquentielles (stations de travail SUN Sparc 20 ou SUN Ultra Sparc ou encore sur un nœud de la machine parallèle IBM-SP2). La résolution directe du problème de Stokes généralisé, qui apparaît à chaque itération en temps, permet d'avoir des temps de calcul moyens de 0.71 - 0.75 seconde par itération (à ces temps de calcul il faut ajouter environ 0.25 seconde par itération lorsqu'on réalise simultanément les estimations *a posteriori*, faites pour la seule cavité entraînée régularisée et qu'on présentera dans le prochain chapitre).

Les principaux résultats concernant la localisation des tourbillons et des extrema de la vitesse sont résumés dans les tableaux 1.10 et 1.11 pour la cavité CANR et dans les tableaux 1.12 et 1.13 pour la cavité CARE. On peut constater une correspondance satisfaisante entre nos résultats et ceux précédemment publiés. Les quelques différences sont dues aux différents maillages utilisés, aux critères de convergence considérés et à la variété des méthodes numériques employées. Bien évidemment un raffinement du maillage améliore les résultats et une meilleure précision sur le test de convergence influe sur l'intensité ainsi que sur la position des extrema de la fonction de courant, i.e. sur la localisation du centre des tourbillons.

Pour mieux comparer nos résultats à ceux des autres auteurs, nous traçons 20 lignes de courant significatives, aux intensités fixées, dont les valeurs sont écrites ci-dessous :

-0.1175	-0.09	-0.01	10-6	0.0005
-0.115	-0.07	-0.001	10-5	0.001
-0.11	-0.05	-0.0001	5 10-5	0.0015
-0.1	-0.03	0	0.0001	0.003

Pour l'isovorticité, les isobares et les lignes de niveau de l'énergie cinétique, on ne donne que les lignes équidistantes les plus significatives.

Solution pour $Re = 1000$

Les simulations pour $Re = 1000$ sont faites sur un maillage resserré près des bords, car les solutions obtenues serviront à initialiser les problèmes de Navier-Stokes à $Re = 5000$. Toutefois un maillage uniforme de 65^2 nœuds aurait été suffisant pour obtenir des simulations correctes pour un tel nombre de Reynolds. La solution stationnaire est atteinte aux temps adimensionnés $T = 25.3$ pour CANR et $T = 29.7$ pour CARE. Dans les tableaux 1.10 et 1.12 sont présentés les résultats concernant respectivement la cavité CANR et la cavité CARE. Enfin les figures 1.14 et 1.15 donnent les différentes caractéristiques de la solution dans la cavité entraînée non régularisée.

Solution pour $Re = 5000$

Un maillage uniforme de 65^2 nœuds est insuffisant pour obtenir une simulation correcte pour $Re = 5000$. En effet, dans le cas d'un maillage uniforme il y a trop d'oscillations, visibles au cours de toute l'évolution en temps, dans les lignes de niveau de la vorticit . En m me temps, l'encombrement m moire d    la factorisation compl te de la matrice   inverser nous emp che de consid rer des maillages de taille 129^2 , qui donneraient des r sultats corrects m me pour une grille uniforme. C'est pour cette raison que nous optons pour un maillage de 65^2 n uds, plus raffin  pr s des bords, qui permet de capter le ph nom ne des couches limites, d' paisseur de l'ordre de $O(Re^{-1/2})$.

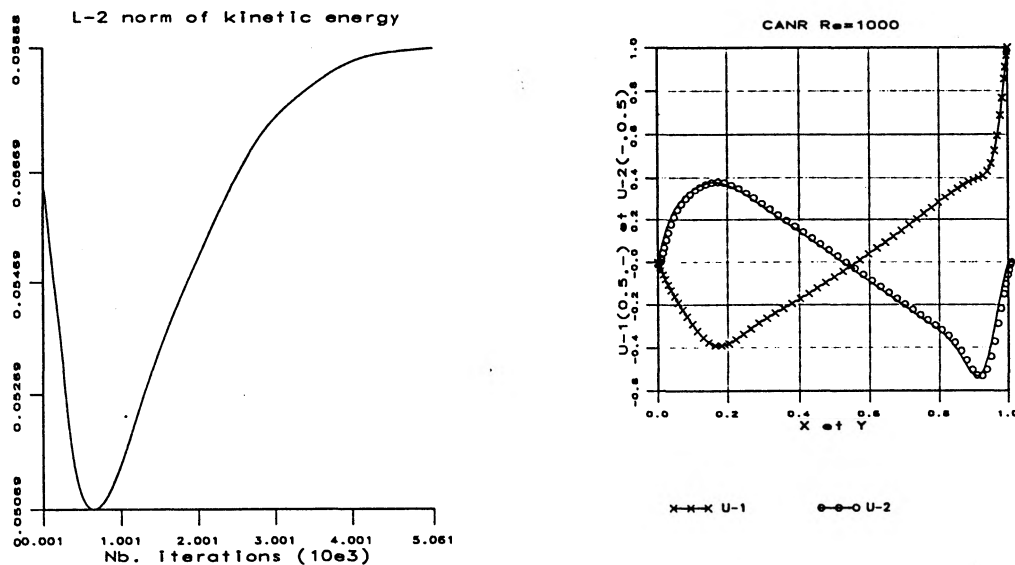
Le probl me de Navier-Stokes est initialis , pour les deuxavit s, par la solution stationnaire calcul e   $Re = 1000$. Cela permet d' viter les perturbations dans le champ de la vitesse qui sont pr sentes, par contre, lorsqu'on part de la solution du probl me de Stokes   $Re = 5000$. En fait, quand la solution initiale est celle de Stokes, des perturbations dans la vitesse se d veloppent   proximit  du coin sup rieur droit du domaine (ces instabilit s sont bien visibles dans le champ normalis  comme aussi en examinant l'isovorticit ). Ces perturbations sont essentiellement dues au maillage qui reste encore trop grossier. En revanche, avec un maillage uniforme de 129^2 sommets, les oscillations sont presque inexistantes (dans les lignes de niveau de la vorticit  on peut encore remarquer quelques petites oscillations pr s du coin sup rieur droit, pr sentes   $T = 1$, mais qui ont pratiquement disparu   $T = 2$).

Nous avons suivi l' volution du ph nom ne des perturbations pendant un certain nombre d'it rations en temps, en partant de la solution de Stokes   $Re = 5000$ et en consid rant le maillage raffin  de 65^2 n uds. Les instabilit s dans le champ de vitesse, qui naissent   proximit  du coin sup rieur droit et se propagent ensuite vers l'int rieur, l  o  le tourbillon principale s'instaure, semblent  tre essentiellement li es aux sch mas de discr tisation en temps. En fait, afin d'exclure toute cause li e   la m thode de p nalisation, nous avons compar  nos r sultats   ceux obtenus par F. Pascal, lorsqu'il discr tise en temps par les m mes sch mas et ensuite r sout le probl me par l'algorithme d'Uzawa. Les oscillations apparaissant dans le champ de vitesse sont comparables, bien que d'intensit  l g rement inf rieure   la n tre (elles sont tr s visibles lorsqu'on trace l'isovorticit ). En revanche, aucune perturbation n' tait observable si l'on discr tisait en temps par un sch ma de splitting   trois pas (voir [27]). Le ph nom ne semble provenir

du terme non linéaire. Même si la condition CFL est bien respectée, il est nécessaire d'avoir un schéma suffisamment dissipatif sur le terme non linéaire (comme des schémas décentrés en différences finies, des méthodes type SUPG ou des éléments P1 bulle en éléments finis ou encore des schémas en temps semi-implicites ou de splitting avec des pas Δt suffisamment grands). L'apparition de perturbations dans la solution est visible également avec des méthodes de projections, lorsqu'on prend en compte des Δt trop petits (cf [28]). Pour conclure, le schéma E-E ne donne pas des résultats meilleurs que ceux du schéma multi-pas CN-AB, moins dissipatif.

A partir de la solution de Navier-Stokes à $Re = 1000$, la solution stationnaire est atteinte aux temps adimensionnés $T = 40.554$ pour CANR et $T = 19.362$ pour CARE (pour CARE le test de sortie est moins précis et il est fixé à $5 \cdot 10^{-5}$). Au fur et à mesure que Re augmente, les recirculations deviennent de plus en plus importantes, les couches limites de plus en plus minces et l'écoulement stationnaire de plus en plus complexe. Dans les tableaux 1.11 et 1.13 sont présentés les résultats concernant respectivement la cavité CANR et la cavité CARE. Enfin les figures 1.16 et 1.17, relatives à la cavité CANR, et les figures 1.18 et 1.19, relatives à la cavité CARE, donnent les différentes caractéristiques de la solution. Notons qu'en figure 1.19 deux isovorticités ont été tracées : celle obtenue par la discrétisation resserrée près des bords de 65^2 sommets et celle calculée avec une discrétisation uniforme de 129^2 nœuds, relative au temps $T = 18.5$. On observe aisément que les oscillations ont totalement disparu dans le cas de la discrétisation plus fine. Toutefois, pour les deux cavités entraînées, bien que la simulation avec 65^2 points montre des isovorticités non satisfaisantes, les autres caractéristiques de la solution stationnaire sont captées correctement (on remarque, dans les tableaux 1.11 et 1.13, la correspondance satisfaisante entre nos résultats et ceux des autres auteurs).

$Re = 1000$						
	Calgaro	Poullet [29]	Pascal [27]	Bruneau [4]	Goyon [17]	Ghia [18]
Formulation	U, P	U, P	U, P	U, P	ω, ψ	ω, ψ
Problème	évolutif	évolutif	évolutif	stationnaire	évolutif	stationnaire
Discrétisation	E.F.	D.F.	E.F.	D.F.	D.F.	D.F.
Maillage	65^2	65^2	65^2	129^2	129^2	129^2
Δt	$5 \cdot 10^{-3}$	$8 \cdot 10^{-3}$	10^{-1}		$2 \cdot 10^{-2}$	
P.V.						
(x, y)	(0.524, 0.573)	(0.539, 0.571)	(0.531, 0.562)	(0.531, 0.559)	(0.531, 0.562)	(0.531, 0.562)
$\psi(x, y)$	-0.122	-0.112	-0.123	-0.116	-0.116	-0.118
$\omega(x, y)$	-2.298	-1.99	-0.022			2.05
B.L.V.						
(x, y)	(0.084, 0.072)	(0.079, 0.079)	(0.078, 0.078)	(0.086, 0.082)	(0.086, 0.078)	(0.086, 0.078)
$\psi(x, y)$	$2.32 \cdot 10^{-4}$	$1.83 \cdot 10^{-4}$	$2.28 \cdot 10^{-4}$	$3.25 \cdot 10^{-4}$	$2.11 \cdot 10^{-4}$	$2.31 \cdot 10^{-4}$
$\omega(x, y)$	0.297	0.295	0.324			-0.362
B.R.V.						
(x, y)	(0.870, 0.113)	(0.857, 0.111)	(0.859, 0.109)	(0.871, 0.109)	(0.867, 0.117)	(0.859, 0.109)
$\psi(x, y)$	$1.78 \cdot 10^{-3}$	$1.71 \cdot 10^{-3}$	$1.69 \cdot 10^{-3}$	$1.91 \cdot 10^{-3}$	$1.63 \cdot 10^{-3}$	$1.75 \cdot 10^{-3}$
$\omega(x, y)$	1.051		1.22			-1.155
$\min(u_1)$	-0.391		-0.399	-0.376		-0.383
$y_{\min}$	0.183		0.172	0.160		0.172
$\min(u_2)$	-0.531		-0.536	-0.521		-0.516
$x_{\min}$	0.916		0.906	0.910		0.906
$\max(u_2)$	0.380		0.386	0.366		0.371
$x_{\max}$	0.165		0.156	0.152		0.156

TAB. 1.10 – Résultats pour la cavité CANR pour $Re = 1000$.FIG. 1.14 – Evolution de la norme L^2 de l'énergie cinétique et profils médians de la vitesse pour CANR à $Re = 1000$ et 65^2 points.

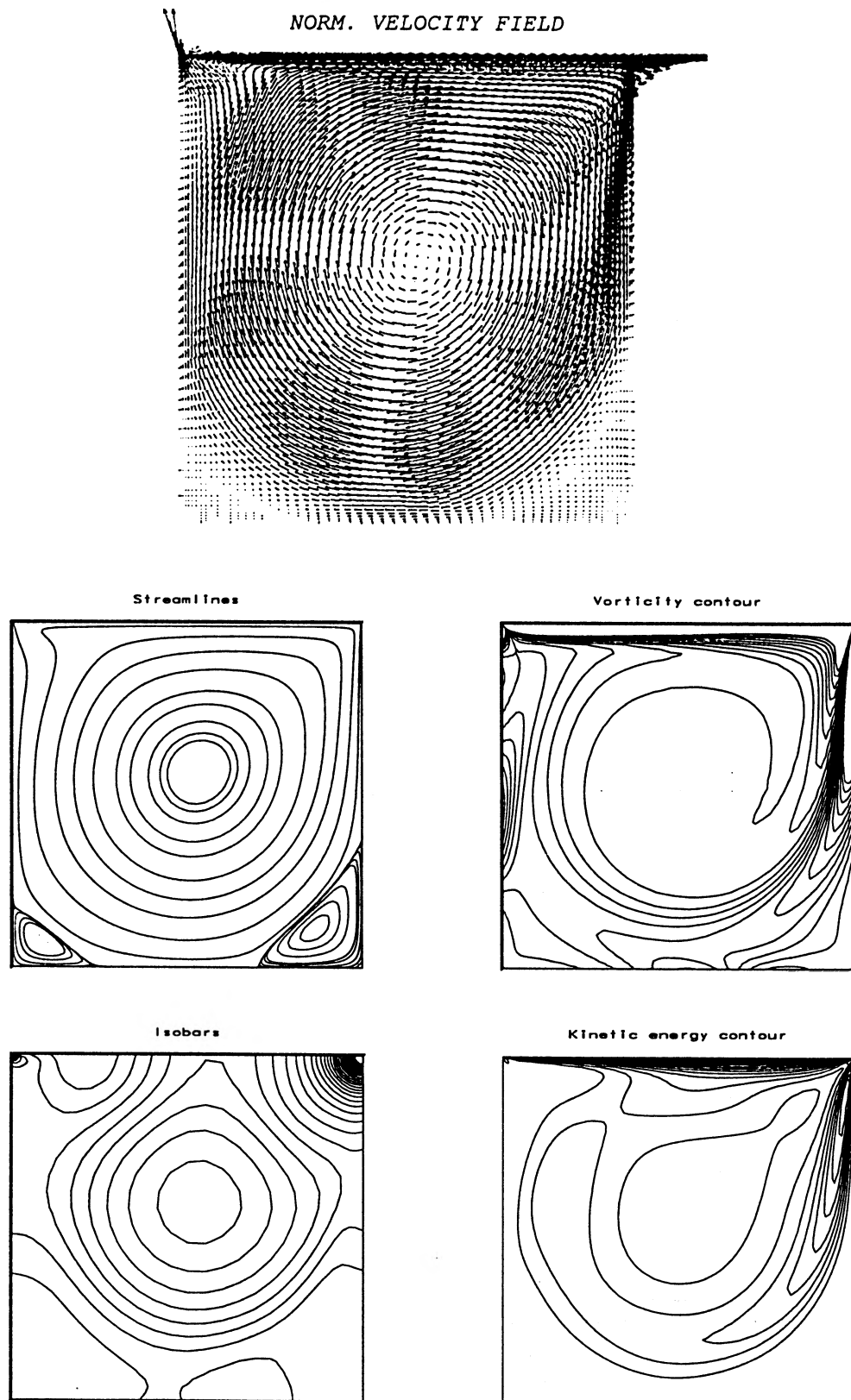


FIG. 1.15 - CANR à $Re = 1000$: champ de vitesse, lignes de courant, isovorticité, isobares et énergie cinétique pour 65^2 points.

$Re = 5000$				
	Calgaro	Bruneau [4]	Goyon [17]	Ghia [18]
Formulation	U, P	U, P	ω, ψ	ω, ψ
Problème	évolutif	stationnaire	évolutif	évolutif
Discrétisation	E.F.	D.F.	D.F.	D.F.
Maillage	65^2	513^2	257^2	257^2
Δt	10^{-3}		$5 \cdot 10^{-3}$	$> 10^{-1}$
P.V.				--
(x, y)	(0.524, 0.524)	(0.516, 0.531)	(0.516, 0.539)	(0.512, 0.535)
$\psi(x, y)$	-0.127	-0.114	-0.112	-0.119
$\omega(x, y)$	-2.271			1.86
U.L.V.				
(x, y)	(0.059, 0.916)	(0.062, 0.910)	(0.062, 0.910)	(0.062, 0.910)
$\psi(x, y)$	$1.39 \cdot 10^{-3}$	$1.75 \cdot 10^{-3}$	$1.30 \cdot 10^{-3}$	$1.46 \cdot 10^{-3}$
$\omega(x, y)$	2.309			-2.088
B.L.V.				
(x, y)	(0.084, 0.130)	(0.066, 0.148)	(0.074, 0.137)	(0.070, 0.137)
$\psi(x, y)$	$1.28 \cdot 10^{-3}$	$2.22 \cdot 10^{-3}$	$1.35 \cdot 10^{-3}$	$1.36 \cdot 10^{-3}$
$\omega(x, y)$	1.29			-1.53
B.R.V.				
(x, y)	(0.797, 0.072)	(0.830, 0.070)	(0.809, 0.074)	(0.809, 0.074)
$\psi(x, y)$	$3.60 \cdot 10^{-3}$	$4.65 \cdot 10^{-3}$	$2.99 \cdot 10^{-3}$	$3.08 \cdot 10^{-3}$
$\omega(x, y)$	3.146			-2.66
S.B.L.V.				
(x, y)	(0.002, 0.002)	(0.012, 0.009)	(0.008, 0.008)	(0.012, 0.008)
$\psi(x, y)$	$-1.53 \cdot 10^{-9}$	$-2.33 \cdot 10^{-7}$	$-6.57 \cdot 10^{-8}$	$-7.09 \cdot 10^{-8}$
$\omega(x, y)$	$-1.73 \cdot 10^{-4}$			$1.88 \cdot 10^{-2}$
S.B.R.V.				
(x, y)	(0.970, 0.021)	(0.967, 0.029)	(0.980, 0.016)	(0.980, 0.019)
$\psi(x, y)$	$-3.36 \cdot 10^{-6}$	$-2.47 \cdot 10^{-5}$	$-9.60 \cdot 10^{-7}$	$-1.43 \cdot 10^{-6}$
$\omega(x, y)$	$-5.10 \cdot 10^{-2}$			$3.19 \cdot 10^{-2}$
$\min(u_1)$	-0.449	-0.436		-0.436
$ymin$	0.072	0.066		0.070
$\min(u_2)$	-0.581	-0.567		-0.554
$xmin$	0.962	0.959		0.953
$\max(u_2)$	0.452	0.426		0.436
$xmax$	0.084	0.076		0.078

TAB. 1.11 – Résultats pour la cavité CANR pour $Re = 5000$.

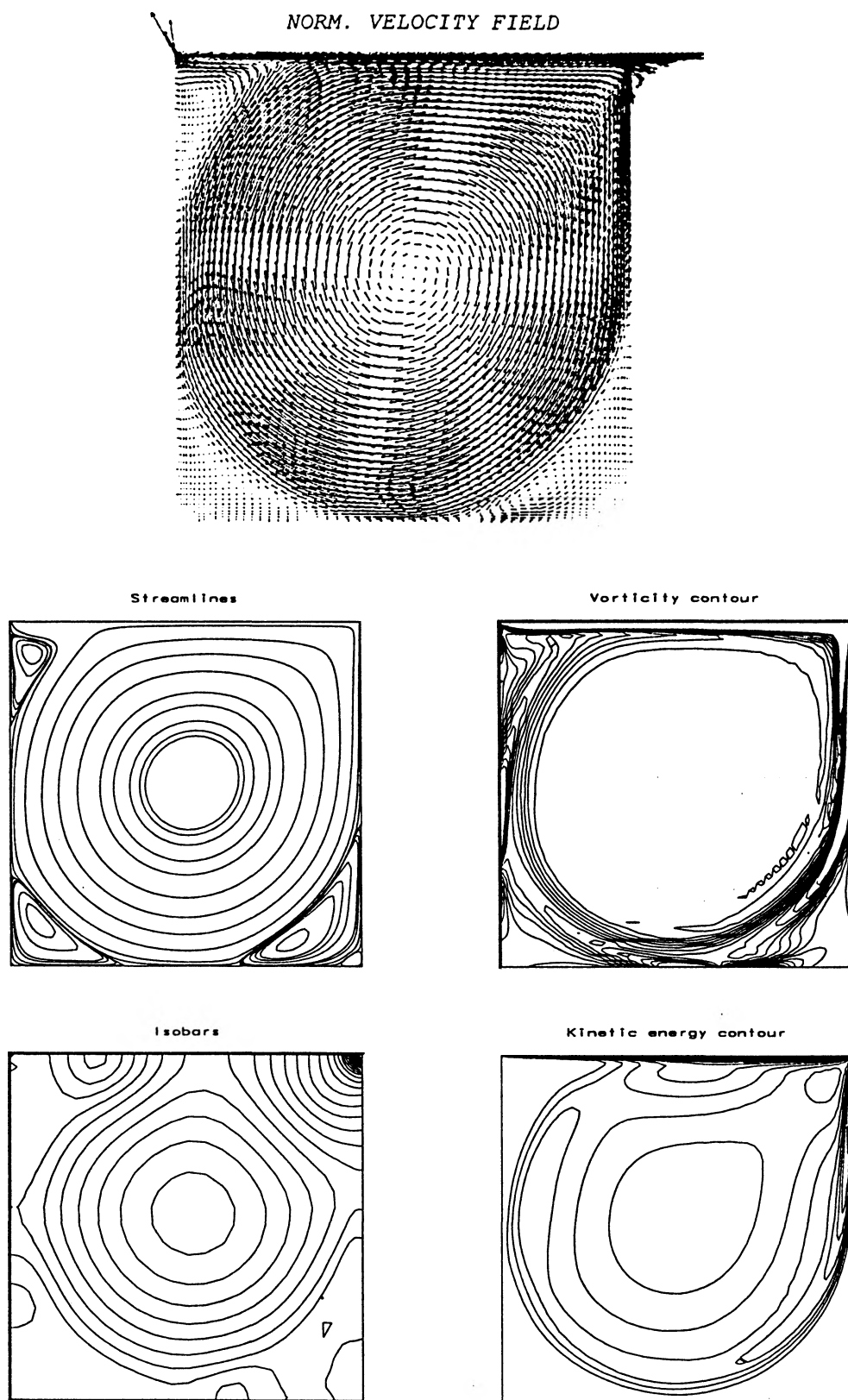


FIG. 1.16 – CANR à $Re = 5000$: champ de vitesse, lignes de courant, isovorticité, isobares et énergie cinétique pour 65^2 points.

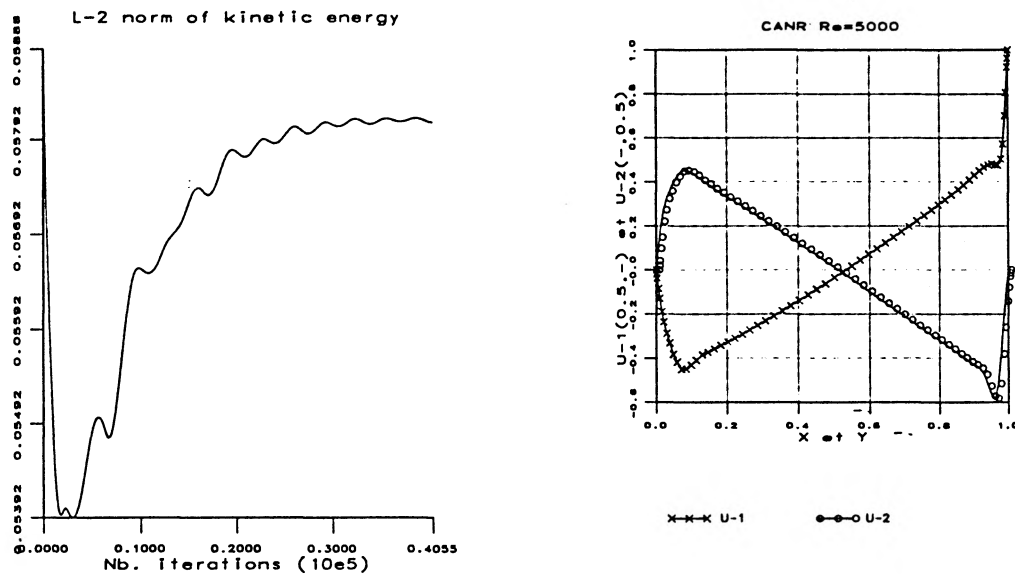


FIG. 1.17 – Evolution de la norme L^2 de l'énergie cinétique et profils médians de la vitesse pour CANR à $Re = 5000$ et 65^2 points.

$Re = 1000$				
	Calgaro	Pascal [27]	Pascal [27]	Shen [30]
Formulation	U, P	U, P	U, P	U, P
Problème	évolutif	évolutif	évolutif	évolutif
Discrétisation	E.F.	E.F.	E.F.	Spectral-Ch.
Maillage	65^2	65^2	129^2	25^2
Δt	$5 \cdot 10^{-3}$	10^{-1}	$5 \cdot 10^{-2}$	$1.5 \cdot 10^{-1}$
P.V.				
(x, y)	(0.548, 0.572)	(0.547, 0.578)	(0.555, 0.586)	(0.547, 0.578)
$\psi(x, y)$	-0.089	-0.092	-0.090	-0.087
$\omega(x, y)$	-1.728	-1.944	-2.168	
B.L.V.				
(x, y)	(0.079, 0.065)	(0.062, 0.062)	(0.055, 0.055)	(0.078, 0.063)
$\psi(x, y)$	$8.09 \cdot 10^{-5}$	$3.15 \cdot 10^{-5}$	$1.27 \cdot 10^{-5}$	$8.28 \cdot 10^{-5}$
$\omega(x, y)$	0.140	0.076	0.041	
B.R.V.				
(x, y)	(0.867, 0.119)	(0.875, 0.125)	(0.867, 0.125)	(0.922, 0.094)
$\psi(x, y)$	$1.01 \cdot 10^{-3}$	$9.22 \cdot 10^{-4}$	$8.85 \cdot 10^{-4}$	$5.68 \cdot 10^{-4}$
$\omega(x, y)$	0.704	0.607	0.510	
$\min(u_1)$	-0.279	-0.279	-0.275	
$ymin$	0.207	0.219	0.242	
$\min(u_2)$	-0.382	-0.376	-0.366	
$xmin$	0.894	0.891	0.891	
$\max(u_2)$	0.263	0.260	0.252	
$xmax$	0.187	0.219	0.242	

TAB. 1.12 – Résultats pour la cavité CARE pour $Re = 1000$.

$Re = 5000$				
	Calgaro	Poullet [29]	Pascal [27]	Shen [30]
Formulation	U, P	U, P	U, P	U, P
Problème	évolutif	évolutif	évolutif	évolutif
Discrétisation	E.F.	D.F.	E.F.	Spectral-Ch.
Maillage	65^2	257^2	129^2	33^2
Δt	$2 \cdot 10^{-3}$	$2 \cdot 10^{-3}$	$5 \cdot 10^{-2}$	$3 \cdot 10^{-2}$
P.V.				
(x, y)	(0.524, 0.548)	(0.522, 0.569)	(0.539, 0.531)	(0.516, 0.531)
$\psi(x, y)$	-0.094	-0.095	-0.097	-0.088
$\omega(x, y)$	-1.711	-2.198	-2.169	
U.L.V.				
(x, y)	(0.092, 0.908)	(0.082, 0.910)	(0.078, 0.906)	(0.078, 0.92)
$\psi(x, y)$	$1.10 \cdot 10^{-3}$	$5.74 \cdot 10^{-4}$	$7.86 \cdot 10^{-4}$	$6.78 \cdot 10^{-4}$
$\omega(x, y)$	1.346	0.965	1.159	
B.L.V.				
(x, y)	(0.092, 0.119)	(0.071, 0.153)	(0.086, 0.117)	(0.094, 0.094)
$\psi(x, y)$	$1.15 \cdot 10^{-3}$	$1.37 \cdot 10^{-3}$	$6.72 \cdot 10^{-4}$	$7.53 \cdot 10^{-4}$
$\omega(x, y)$	0.879	1.461	0.731	
B.R.V.				
(x, y)	(0.813, 0.079)	(0.816, 0.082)	(0.805, 0.078)	(0.922, 0.094)
$\psi(x, y)$	$2.18 \cdot 10^{-3}$	$2.04 \cdot 10^{-3}$	$2.42 \cdot 10^{-3}$	$7.75 \cdot 10^{-4}$
$\omega(x, y)$	1.575	1.648	2.009	
S.B.R.V.				
(x, y)	(0.990, 0.009)	(0.992, 0.012)	(0.984, 0.016)	
$\psi(x, y)$	$-2.72 \cdot 10^{-8}$	$-1.23 \cdot 10^{-7}$	$-2.39 \cdot 10^{-7}$	
$\omega(x, y)$	$-6.91 \cdot 10^{-3}$	$-8.26 \cdot 10^{-3}$	$-2.34 \cdot 10^{-2}$	
$\min(u_1)$	-0.304		-0.318	
$ymin$	0.106		0.094	
$\min(u_2)$	-0.404		-0.414	
$xmin$	0.948		0.945	
$\max(u_2)$	0.303		0.310	
$xmax$	0.106		0.102	

TAB. 1.13 – Résultats pour la cavité CARE pour $Re = 5000$.

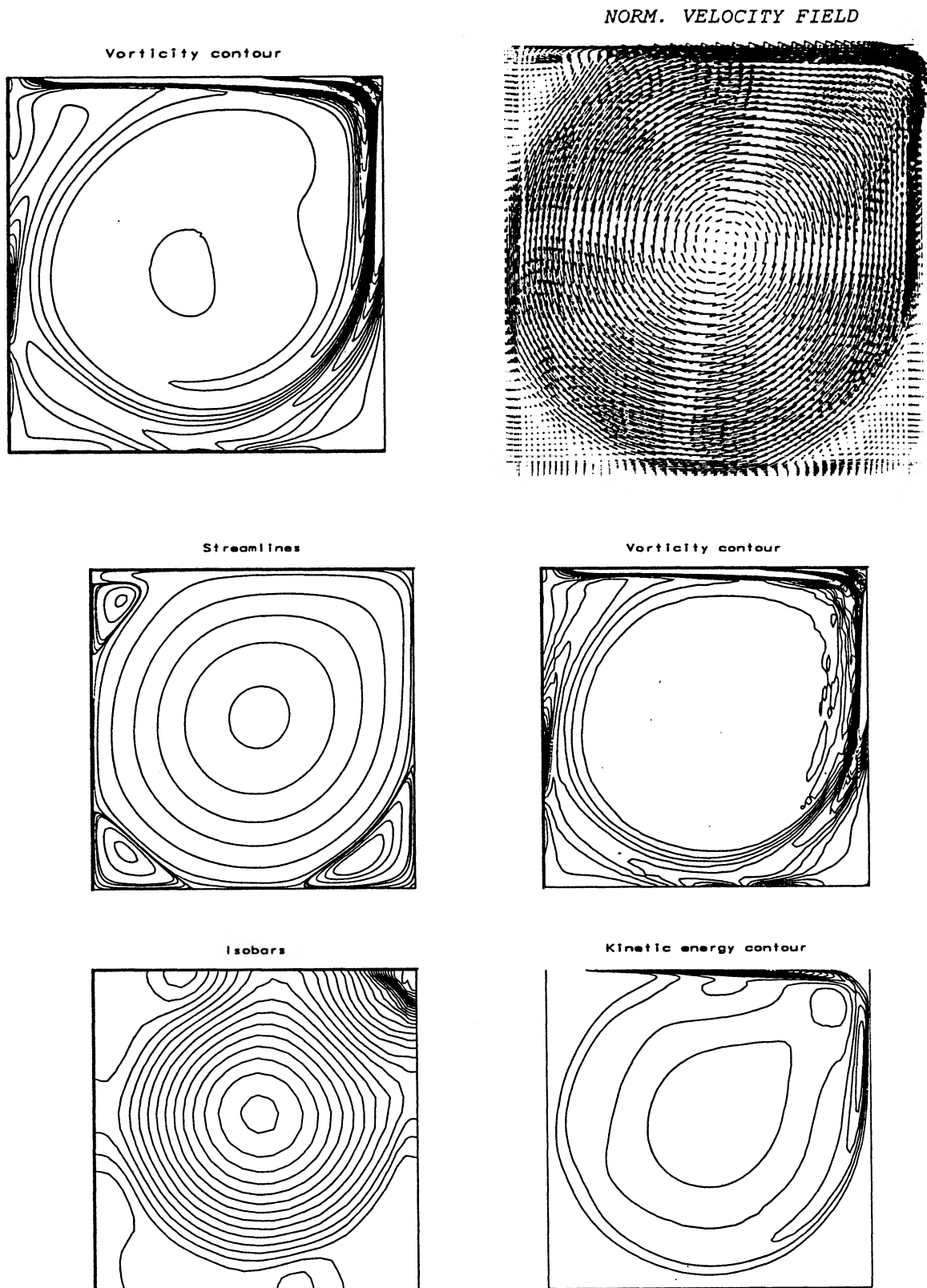


FIG. 1.18 – CARE à $Re = 5000$: isovorticité pour 129^2 points et champ de vitesse, lignes de courant, isovorticité, isobares et énergie cinétique pour 65^2 points.

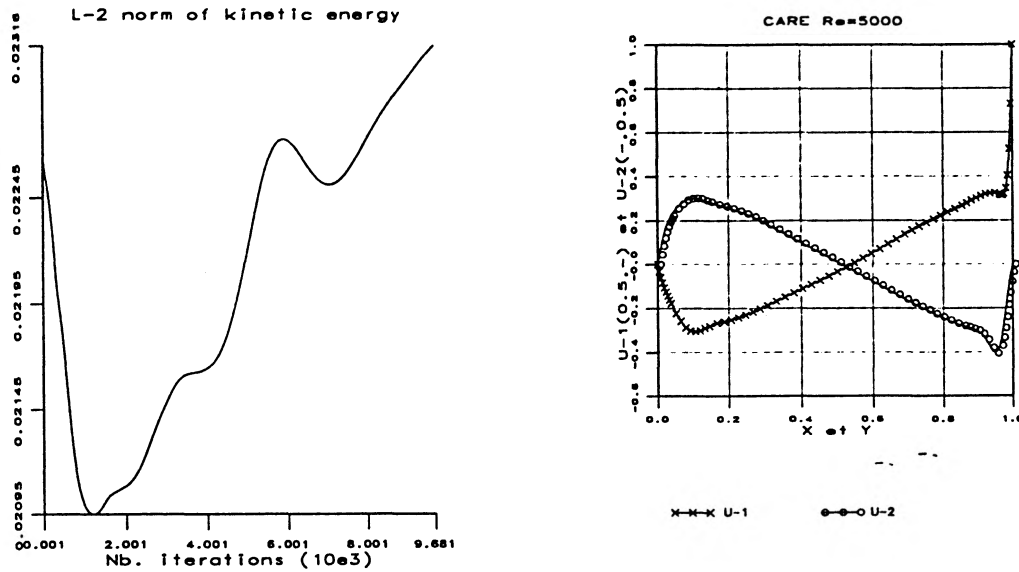


FIG. 1.19 – Evolution de la norme L^2 de l'énergie cinétique et profils médians de la vitesse pour CARE à $Re = 5000$ et 65^2 points.

1.6 Conclusions

Nous avons précisé, au début du chapitre, les raisons qui nous ont motivés à choisir la méthode de pénalisation plutôt que d'autres méthodes employées plus fréquemment dans la littérature afin d'approcher numériquement la solution du problème de Navier-Stokes. Actuellement notre but principal ne réside pas dans la recherche d'une méthode plus performante que celles existantes, mais plutôt dans l'exploitation d'une méthode qui nous permette de simplifier le problème examiné. Dans ces termes, la pénalisation devrait nous permettre d'éviter les difficultés qui se présentent en perspective, difficultés liées au traitement hiérarchique de l'incompressibilité. Grâce à la pénalisation on se ramène théoriquement et numériquement à une seule équation en vitesse pour laquelle nous souhaitons appliquer des idées analogues à celles déjà utilisées dans la première partie de la thèse pour résoudre les équations bidimensionnelles de Burgers.

Toutefois, nous avons dû affronter de nombreux problèmes provenant de la pénalisation. Il s'agissait surtout de problèmes liés au très mauvais conditionnement de la matrice à laquelle nous parvenons après l'élimination de la pression dans les équations étudiées. En pratique, nous avons testé plusieurs préconditionnements classiques des méthodes itératives employées, du type gradient conjugué. Aucun de ces préconditionnements n'est apparu adéquat pour la détermination d'un solveur performant du système linéaire à résoudre. En fait la méthode préconditionnée échoue, soit à cause du grand nombre d'itérations pour atteindre la convergence, soit, même si le nombre d'itérations est faible, à cause de l'augmentation du coût en temps de calcul. L'étude des valeurs propres de la matrice à inverser n'a abouti à aucun nouveau préconditionneur en raison des difficultés théoriques rencontrées pour le définir. Seule une méthode directe s'est

avérée efficace comme solveur du problème linéaire. Néanmoins, le stockage de la factorisation complète de la matrice à inverser nous restreint à considérer des discrétisations spatiales relativement grossières, et donc à simuler des écoulements avec des nombres de Reynolds pas trop grand. Ceci nous a limité à examiner des écoulements convergeant vers des solutions stationnaires.

Nous nous sommes aperçus qu'une solution aux problèmes de conditionnement dus à la pénalisation peut être donnée par les méthodes de projections (développées à présent par F. Pascal). Une autre possibilité vient de la prise en compte d'une deuxième forme de pénalisation, obtenue lorsqu'on remplace l'équation de continuité par

$$\operatorname{div} u + \epsilon \frac{\partial p}{\partial t} = 0 ,$$

(cf. par exemple les travaux récents de J. Shen [31] et [32]). Sous cette nouvelle forme, l'élimination de la pression s'effectue seulement numériquement. En revanche, elle nous permet de considérer un coefficient de pénalisation beaucoup plus grand car, à la place de la quantité ϵ , on prend en compte la quantité $\frac{\epsilon}{\Delta t}$. Cela donne un meilleur conditionnement de la matrice à inverser. On peut donc espérer qu'avec cette approche l'utilisation d'une méthode itérative soit suffisante pour une inversion efficace du système linéaire. Dans ce cas on éviterait les problèmes d'encombrement mémoire liés au stockage de la factorisation et donc on pourrait raffiner le pas en espace, afin de simuler des écoulements à des nombres de Reynolds beaucoup plus importants. Tout cela fait partie des développements que nous envisageons dans la résolution des équations de Navier-Stokes.

Relativement à la discrétisation par les éléments finis mixtes, le premier choix qui portait sur l'élément 4P1-P0 s'est révélé inapproprié. L'approximation de la pression par des fonctions constantes sur chaque triangle du maillage permettait une élimination simple de cette inconnue dans les équations de Stokes et de Navier-Stokes. Les matrices assemblées étaient plus creuses et possédaient une largeur de bande inférieure à leurs analogues, calculées à partir de l'élément 4P1-P1. Cela amenait aussi à une factorisation complète moins encombrante en espace mémoire. Toutefois, la solution obtenue avec le choix du 4P1-P0 avait des perturbations trop importantes, numériquement non négligeables. En revanche, les instabilités disparaissent en discrétisant le problème à l'aide du 4P1-P1. Nous avons donc opté pour ce dernier élément. Les simulations numériques effectuées donnent une solution correcte des équations de Navier-Stokes pénalisées. Les résultats, pour des écoulement jusqu'à $Re = 5000$ sont satisfaisants, comparés à ceux obtenus par d'autres auteurs.

Dans le chapitre suivant nous présenterons l'étude de la décomposition hiérarchique de la solution des équations de Navier-Stokes, obtenue à $Re = 5000$. Toute l'analyse sera réalisée *a posteriori*, pendant la résolution numérique des équations discrétisées par la méthode de Galerkin classique. C'est pour cette raison qu'il était indispensable d'avoir avant tout un code numérique donnant des résultats satisfaisants.

Chapitre 2

Le problème de Navier-Stokes en base hiérarchique

Introduction

Dans ce chapitre, après un bref rappel de la hiérarchisation de la vitesse et de la pression, nous présentons la décomposition sur la base hiérarchique du problème linéaire de Stokes, celle des équations évolutives de Navier-Stokes étant une simple généralisation. D'abord, nous présentons la hiérarchisation sur deux niveaux, qui ne pose aucune difficulté. Ensuite, nous proposons deux approches permettant de réitérer la hiérarchisation des nos équations. La première approche, considérant une hiérarchisation de la pression équivalente à celle de la vitesse, nous amène à des équations trop compliquées. En revanche, la deuxième décomposition de la pression hiérarchise simplement le problème de Stokes ou Navier-Stokes, mais elle n'apporte aucun résultat satisfaisant en vue de l'application de la méthode multi-niveaux, efficacement mise en place auparavant dans le cas des équations de Burgers.

Une analyse globale est ensuite détaillée : celle-ci nous permet d'évaluer les structures des différentes échelles de la vitesse et de la pression, ainsi que tous les termes linéaires et non linéaires présents dans les équations projetées. La hiérarchisation fait paraître les grandes structures des deux variables physiques et leurs corrections, de taille inférieure. Afin de décider si ces corrections peuvent être figées pendant un ou plusieurs pas de temps, il faut évaluer et comparer les différents termes de nos équations projetées sur la grille grossière, pour voir si les termes dépendant des incréments peuvent être négligés ou figés durant quelques itérations en temps. D'abord nous exposons l'étude globale des différents termes linéaires ou non linéaires, en vitesse ou en pression. L'analyse est réalisée en estimant ces termes sur le domaine Ω tout entier. Ce qui nous intéresse le plus est d'observer les variations en temps de ces termes, afin d'estimer si elles sont suffisamment petites pour être négligées sur certains intervalles de temps. Ensuite, puisqu'une séparation d'échelle globale semble être inadéquate pour implémenter efficacement notre méthode multi-résolution, nous localisons cette étude grâce à une décomposition de Ω en plusieurs sous-domaines. L'analyse des différents termes présents dans les équations

projetées, analyse qui a été réalisée dans chaque sous-domaine, montre de manière incontestable une évolution locale de ces termes-ci. Ce comportement local est entièrement noyé dans l'analyse globale. L'évolution par zone de tous les termes montrera que ceux qui dépendent des corrections peuvent être figés pendant quelques pas de temps et, en particulier, que les variations du terme d'interaction non linéaire peuvent être négligées sur de longs intervalles de temps.

Anticipons ici que les résultats positifs de l'étude locale permettront d'envisager le développement d'un algorithme multi-niveaux couplé avec une méthode de décomposition de domaine.

2.1 Décomposition sur la base hiérarchique

Les résultats concernant les algorithmes multi-niveaux de *type incrémental* appliqués au problème de Burgers (voir la première partie de la thèse) permettent d'envisager l'utilisation de ces méthodes pour le problème de Navier-Stokes évolutif. La décomposition en grandes structures et incréments de celles-ci, obtenue grâce à l'utilisation de la base hiérarchique, est appliquée au champ de vitesse et à la pression.

Les notations sont identiques à celles décrites en section 1.2 et restent cohérentes avec les notations introduites dans la partie I. Rappelons que l'espace d'éléments finis V_d associé à la grille fine $\mathcal{T}_d$ se décompose en

$$V_d = V_{d-1} + W_d$$

et, de façon récurrente sur les grilles plus grossières :

$$V_{d-1} = V_{d-2} + W_{d-1} \quad , \quad \dots \quad , \quad V_1 = V_0 + W_1 .$$

Ces sous-espaces sont munis des bases hiérarchiques respectives :

$$\begin{aligned} V_k &= \text{Span} \{ \phi_{k,1}, \dots, \phi_{k,n_k} \}, & n_k &= \dim(V_k) & (\text{pour } k = 0, \dots, d) , \\ W_k &= \text{Span} \{ \psi_{k,n_{k-1}+1}, \dots, \psi_{k,n_k} \}, & n_k - n_{k-1} &= \dim(W_k) & (\text{pour } k = 1, \dots, d) . \end{aligned}$$

Enfin pour le choix de l'élément fini mixte 4P1-P1 nous avons posé :

$$\mathcal{Q}_k = V_{k-1} \quad \text{pour } k = 1, \dots, d ;$$

ainsi l'espace complémentaire de la pression sera :

$$\mathcal{R}_k = W_{k-1} \quad \text{pour } k = 2, \dots, d .$$

Afin de garder une notation en même temps indépendante de l'élément fini mixte choisi et cohérente avec celles introduites précédemment, on munit les sous-espaces relatifs à la pression des bases hiérarchiques respectives :

$$\begin{aligned} \mathcal{Q}_k &= \text{Span} \{ \xi_{k,1}, \dots, \xi_{k,m_k} \}, & m_k &= \dim(\mathcal{Q}_k) & (\text{pour } k = 1, \dots, d) , \\ \mathcal{R}_k &= \text{Span} \{ \eta_{k,m_{k-1}+1}, \dots, \eta_{k,m_k} \}, & m_k - m_{k-1} &= \dim(\mathcal{R}_k) & (\text{pour } k = 2, \dots, d) . \end{aligned}$$

De la même manière qu'en section 1.4 de la partie I, nous décomposons de façon récurrente les espaces de discrétisation de la vitesse et de la pression (respectivement $u_d(t) \in \mathcal{V}_{g,d}$ et $p_d(t) \in \mathcal{Q}_d$). La décomposition est analogue pour l'espace $\mathcal{V}_{0,d} = \mathcal{V}_d \cap (H_0^1(\Omega))^2$. Nous obtenons :

$$\begin{aligned} \mathcal{V}_{g,d} &= \mathcal{V}_{g,d-1} + \mathcal{W}_{g,d} & \text{et} & & \mathcal{Q}_d &= \mathcal{Q}_{d-1} + \mathcal{R}_d \\ \mathcal{V}_{g,d-1} &= \mathcal{V}_{g,d-2} + \mathcal{W}_{g,d-1} & \text{et} & & \mathcal{Q}_{d-1} &= \mathcal{Q}_{d-2} + \mathcal{R}_{d-1} \\ &\vdots & \text{et} & & \vdots & \\ \mathcal{V}_{g,2} &= \mathcal{V}_{g,1} + \mathcal{W}_{g,2} & \text{et} & & \mathcal{Q}_2 &= \mathcal{Q}_1 + \mathcal{R}_2 \end{aligned}$$

Par conséquent, après généralisation des notations en posant $y_d = u_d$, la vitesse et la pression en base hiérarchique s'écrivent, pour tout $k = 2, \dots, d$:

- $y_k = y_{k-1} + z_k$ avec

$$\begin{aligned} y_{k-1} &= (y_{k-1}^1, y_{k-1}^2) \in \mathcal{V}_{g,k-1} \quad \text{et} \quad y_{k-1}^l = \sum_{i=1}^{n_{k-1}} y_{k-1,i}^l \phi_{k-1,i} \quad (l = 1, 2), \\ z_k &= (z_k^1, z_k^2) \in \mathcal{W}_{g,k} \quad \text{et} \quad z_k^l = \sum_{i=n_{k-1}+1}^{n_k} z_{k,i}^l \psi_{k,i} \quad (l = 1, 2); \end{aligned} \quad (2.1)$$

- $p_k = p_{k-1} + \pi_k$ avec

$$\begin{aligned} p_{k-1} &= \sum_{i=1}^{m_{k-1}} p_{k-1,i} \xi_{k-1,i} \in \mathcal{Q}_{k-1}, \\ \pi_k &= \sum_{i=m_{k-1}+1}^{m_k} \pi_{k,i} \eta_{k,i} \in \mathcal{R}_k. \end{aligned} \quad (2.2)$$

2.1.1 Stokes en base hiérarchique

Afin de faciliter la présentation de l'étude concernant la base hiérarchique, examinons dans un premier temps les équations linéaires de Stokes. Cette simplification est motivée par le fait que le terme en temps des équations évolutives de Navier-Stokes se décompose de la même manière que le laplacien, tandis que le terme non linéaire $\hat{b}$ peut être inclus dans le terme source, $\hat{b}$ étant traité de façon explicite.

Hiérarchisation sur deux niveaux

A partir de la formulation discrète dans la base nodale du problème de Stokes pénalisé (1.23), écrivons la discrétisation correspondant à la décomposition $\mathcal{V}_{g,d} = \mathcal{V}_{g,d-1} + \mathcal{W}_{g,d}$ et $\mathcal{Q}_d = \mathcal{Q}_{d-1} + \mathcal{R}_d$.

Cherchons $(y_{d-1}, p_{d-1}) \in \mathcal{V}_{g,d-1} \times \mathcal{Q}_{d-1}$ et $(z_d, \pi_d) \in \mathcal{W}_{g,d} \times \mathcal{R}_d$ solutions des systèmes suivants :

$$\begin{cases} \frac{1}{Re} ((y_{d-1} + z_d, \tilde{y}_{d-1})) - (p_{d-1} + \pi_d, \text{div } \tilde{y}_{d-1}) = (f, \tilde{y}_{d-1}) & \forall \tilde{y}_{d-1} \in \mathcal{V}_{0,d-1} \\ (\text{div } y_{d-1} + \text{div } z_d, \tilde{p}_{d-1}) + \epsilon (p_{d-1} + \pi_d, \tilde{p}_{d-1}) = 0 & \forall \tilde{p}_{d-1} \in \mathcal{Q}_{d-1}, \end{cases} \quad (2.3)$$

$$\begin{cases} \frac{1}{Re}((y_{d-1} + z_d, \tilde{z}_d)) - (p_{d-1} + \pi_d, \operatorname{div} \tilde{z}_d) = (f, \tilde{z}_d) & \forall \tilde{z}_d \in \mathcal{W}_{0,d} \\ (\operatorname{div} y_{d-1} + \operatorname{div} z_d, \tilde{\pi}_d) + \epsilon(p_{d-1} + \pi_d, \tilde{\pi}_d) = 0 & \forall \tilde{\pi}_d \in \mathcal{R}_d. \end{cases} \quad (2.4)$$

Les systèmes (2.3) et (2.4) sont équivalents aux projections respectivement sur la grille *grossière* et sur la grille *fine complémentaire* des équations de Stokes pénalisées. Le problème en y et p (2.3) et celui en z et π (2.4) sont bien posés (cf. [27]).

L'idée à la base des méthodes multi-niveaux en espace et en temps est de figer pendant de courts intervalles de temps les corrections z et π . Par conséquent, lorsqu'on approchera les équations de Navier-Stokes le système équivalent à (2.4) ne sera jamais résolu. De plus il faut écrire la forme matricielle associée au problème pénalisé (2.3) afin d'éliminer la pression et de nous réduire à un problème dans la seule inconnue y_{d-1} , vitesse correspondante à la grille grossière $\mathcal{T}_{d-1}$. Naturellement tous les termes en z_d et π_d apparaissant dans les équations évolutives seront traités explicitement.

Dorénavant, en accord avec la notation introduite au début de la section 1.3.3, on note avec un chapeau les matrices et les vecteurs exprimés dans la base nodale, et sans chapeau leurs équivalents dans la base hiérarchique. De plus, pour $k = 1, \dots, d$, les matrices blocs $\hat{\mathcal{A}}_k$ sont scindées en deux matrices, celle de masse notée $\widehat{\mathcal{M}}_k$ et celle du laplacien que l'on persiste à noter $\hat{\mathcal{A}}_k$. Ces matrices sont ainsi définies :

$$\widehat{\mathcal{M}}_k = \begin{pmatrix} \overline{\mathcal{M}}_k & 0 \\ 0 & \mathcal{M}_k \end{pmatrix} \quad \text{et} \quad \hat{\mathcal{A}}_k = \begin{pmatrix} \overline{\mathcal{A}}_k & 0 \\ 0 & \mathcal{A}_k \end{pmatrix},$$

avec

$$(\overline{\mathcal{M}}_k)_{ij} = \frac{1}{\Delta t} (\phi_{k,i}, \phi_{k,j}) \quad \text{et} \quad (\overline{\mathcal{A}}_k)_{ij} = \frac{1}{Re} ((\phi_{k,i}, \phi_{k,j})) \quad \text{pour } 1 \leq i, j \leq n_{0,k}.$$

Rappelons que les matrices hiérarchiques $\mathcal{A}_k, \mathcal{B}_k, \mathcal{C}_k$ et $\mathcal{M}_k$ ($k = 1, \dots, d$) sont des matrices par blocs qui se mettent sous la forme :

$$\mathcal{A}_k = \begin{pmatrix} \mathcal{A}_{gg}^{(k)} & \mathcal{A}_{gf}^{(k)} \\ \mathcal{A}_{fg}^{(k)} & \mathcal{A}_{ff}^{(k)} \end{pmatrix} \quad (2.5)$$

et de même pour $\mathcal{B}_k, \mathcal{C}_k$ et $\mathcal{M}_k$. Les matrices $\mathcal{A}_k, \mathcal{C}_k$ et $\mathcal{M}_k$ sont symétriques, donc on a $\mathcal{A}_{fg}^{(k)} = (\mathcal{A}_{gf}^{(k)})^T$, ceci étant pareil pour les blocs $\mathcal{C}_{fg}^{(k)}$ et $\mathcal{M}_{fg}^{(k)}$.

Ces matrices hiérarchiques ne sont jamais assemblées. Les calculs se réalisent à l'aide des matrices de passage de la base nodale à la base hiérarchique : S la matrice relative à la vitesse et T à la pression. Les relations entre les matrices nodales et hiérarchiques sont classiques :

$$\begin{aligned} \mathcal{A}_k &= S^T \hat{\mathcal{A}}_k S, \\ \mathcal{B}_k &= T^T \hat{\mathcal{B}}_k S, \\ \mathcal{C}_k &= T^T \hat{\mathcal{C}}_k T, \\ \mathcal{M}_k &= S^T \widehat{\mathcal{M}}_k S. \end{aligned}$$

Le système (2.3) peut alors s'écrire sous la forme matricielle suivante :

$$\begin{cases} \mathcal{A}_{gg}^{(d)} Y_{d-1} + \mathcal{A}_{gf}^{(d)} Z_d + {}^T\mathcal{B}_{gg}^{(d)} P_{d-1} + {}^T\mathcal{B}_{fg}^{(d)} \Pi_d = F_{d-1} \\ \mathcal{B}_{gg}^{(d)} Y_{d-1} + \mathcal{B}_{gf}^{(d)} Z_d - \epsilon \mathcal{C}_{gg}^{(d)} P_{d-1} - \epsilon \mathcal{C}_{gf}^{(d)} \Pi_d = 0 \end{cases} \quad (2.6)$$

avec $Y_{d-1} = (Y_{d-1}^1, Y_{d-1}^2)$ et $F_{d-1} = (F_{d-1}^1, F_{d-1}^2)$ les vecteurs colonnes formés des composantes de la vitesse et des forces extérieures exprimées dans la base hiérarchique de $V_{0,d-1}$, $Z_d = (Z_d^1, Z_d^2)$ le vecteur colonne formé des composantes des corrections de la vitesse exprimées dans la base hiérarchique de $W_{0,d}$ et enfin P_{d-1} et Π_d les vecteurs colonnes de la pression et de sa correction exprimées respectivement dans la base hiérarchique de Q_{d-1} et dans celle de $\mathcal{R}_d$.

La seconde équation du système (2.6) nous fournit la valeur de la pression sur la grille grossière :

$$P_{d-1} = \frac{1}{\epsilon} \left(\mathcal{C}_{gg}^{(d)} \right)^{-1} \left(\mathcal{B}_{gg}^{(d)} Y_{d-1} + \mathcal{B}_{gf}^{(d)} Z_d - \epsilon \mathcal{C}_{gf}^{(d)} \Pi_d \right). \quad (2.7)$$

Remplaçant la valeur de P_{d-1} donnée par (2.7) dans la première équation de (2.6), on trouve le système suivant dont la seule inconnue est la vitesse sur la grille grossière Y_{d-1} :

$$\begin{aligned} \left(\mathcal{A}_{gg}^{(d)} + \frac{1}{\epsilon} {}^T\mathcal{B}_{gg}^{(d)} \left(\mathcal{C}_{gg}^{(d)} \right)^{-1} \mathcal{B}_{gg}^{(d)} \right) Y_{d-1} = \\ F_{d-1} - \mathcal{A}_{gf}^{(d)} Z_d - {}^T\mathcal{B}_{fg}^{(d)} \Pi_d - \frac{1}{\epsilon} {}^T\mathcal{B}_{gg}^{(d)} \left(\mathcal{C}_{gg}^{(d)} \right)^{-1} \mathcal{B}_{gf}^{(d)} Z_d + {}^T\mathcal{B}_{gg}^{(d)} \left(\mathcal{C}_{gg}^{(d)} \right)^{-1} \mathcal{C}_{gf}^{(d)} \Pi_d. \end{aligned} \quad (2.8)$$

Remarque 11 Une variante du système (2.6) est celle qui consiste à négliger le terme de pénalisation $\epsilon \mathcal{C}_{gf}^{(d)} \Pi_d$ dans la seconde équation, car sa contribution est petite par rapport aux autres termes de la même équation. En effet les termes en ϵ sont très petits par rapport aux termes de divergence et de plus la norme L^2 de la correction de la pression π_d est inférieure à la norme de la pression sur la grille grossière p_{d-1} . Ceci permet de simplifier légèrement l'expression du second membre de (2.8) sans introduire a priori une erreur trop grande dans l'approximation des équations considérées.

Remarque 12 Il faut avant tout hiérarchiser le problème de Stokes pénalisé et seulement dans un deuxième temps éliminer la pression. En revanche, si l'on cherche à hiérarchiser le système nodal dans lequel la pression a déjà été éliminée,

$$\left(\hat{\mathcal{A}}_d + \frac{1}{\epsilon} \hat{\mathcal{B}}_d^T \hat{\mathcal{C}}_d^{-1} \hat{\mathcal{B}}_d \right) \hat{U}_d = \hat{F}_d,$$

on découvre aussitôt deux difficultés insurmontables. D'abord on remarque que le choix de remplacer la matrice $\hat{\mathcal{C}}_d$ par la matrice identité n'est pas adéquat. En effet, en réécrivant le système ci-dessus en base hiérarchique, on trouverait la matrice $T^T \cdot T$ à la place de la matrice identité. La deuxième difficulté réside dans l'impossibilité de hiérarchiser la matrice produit $\hat{\mathcal{B}}_d^T \hat{\mathcal{C}}_d^{-1} \hat{\mathcal{B}}_d$ à cause de l'inversion de $\hat{\mathcal{C}}_d$, matrice de masse relative à la pression, $\hat{\mathcal{C}}_d$ étant une matrice non diagonale.

Hiérarchisation sur plusieurs niveaux

Nous avons présenté ci-dessus la hiérarchisation sur deux niveaux des équations de Stokes. Le but est maintenant celui de hiérarchiser le système d'équations considérées sur plusieurs niveaux.

Pour cela notons $\bar{F}_{d-1}$ le second membre complet de (2.8), composé du terme source projeté sur la grille grossière et de tous les termes en Z_d et Π_d qui seront figés. Soit de plus $\tilde{F}_{d-1}$ une partie du second membre de (2.8), dans laquelle les termes provenant de la pénalisation sont exclus, i.e. $\tilde{F}_{d-1} = F_{d-1} - \mathcal{A}_{gf}^{(d)} Z_d - \mathcal{B}_{fg}^{(d)} \Pi_d$.

Rappelons enfin que tout bloc hiérarchique du type $\mathcal{A}_{gg}^{(k)}$ ($k = 1, \dots, d$) est, par sa définition, équivalent à la matrice nodale sur la grille plus grossière, i.e. $\mathcal{A}_{gg}^{(k)} = \hat{\mathcal{A}}_{k-1}$; de même pour $\mathcal{B}_{gg}^{(k)}$ et $\mathcal{C}_{gg}^{(k)}$, ($k = 2, \dots, d$). Le système (2.8) peut alors être considéré comme un système de Stokes pénalisé, nodal sur la grille $d-1$, ayant simplement un terme source corrigé.

Nous avons étudié deux approches afin de réitérer la hiérarchisation du problème. La première consiste à hiérarchiser un problème de Stokes modifié, relatif à la grille de niveau $d-1$, trouvé à partir du système (2.6). La seconde approche implique l'introduction d'une nouvelle pression qui n'a aucun lien avec la pression physique, mais qui permet une meilleure hiérarchisation du problème de Stokes, cela à partir du système (2.8).

Développons à présent la première approche. Pour hiérarchiser le problème de Stokes relatif à la grille de niveau $d-1$ (voir le système (2.3) ou (2.6)), écrivons la discrétisation de (2.3) correspondant à la décomposition $\mathcal{V}_{g,d-1} = \mathcal{V}_{g,d-2} + \mathcal{W}_{g,d-1}$ et $\mathcal{Q}_{d-1} = \mathcal{Q}_{d-2} + \mathcal{R}_{d-1}$. D'abord notons $(\omega, \tilde{y}_{d-1})$ le terme source modifié par les termes en z_d et π_d :

$$(\omega, \tilde{y}_{d-1}) = (f, \tilde{y}_{d-1}) - \frac{1}{Re}((z_d, \tilde{y}_{d-1})) + (\pi_d, \text{div } \tilde{y}_{d-1}) ,$$

et $(\vartheta, \tilde{p}_{d-1})$ le second membre qui va apparaître dans l'équation de continuité pénalisée:

$$(\vartheta, \tilde{p}_{d-1}) = -(\text{div } z_d, \tilde{p}_{d-1}) - \epsilon(\pi_d, \tilde{p}_{d-1}) .$$

Ainsi le système (2.3) se réécrit:

$$\left\{ \begin{array}{l} \text{Trouver } (y_{d-1}, p_{d-1}) \in \mathcal{V}_{g,d-1} \times \mathcal{Q}_{d-1} \text{ tel que} \\ \frac{1}{Re}((y_{d-1}, \tilde{y}_{d-1})) - (p_{d-1}, \text{div } \tilde{y}_{d-1}) = (\omega, \tilde{y}_{d-1}) \quad \forall \tilde{y}_{d-1} \in \mathcal{V}_{0,d-1} , \\ (\text{div } y_{d-1}, \tilde{p}_{d-1}) + \epsilon(p_{d-1}, \tilde{p}_{d-1}) = (\vartheta, \tilde{p}_{d-1}) \quad \forall \tilde{p}_{d-1} \in \mathcal{Q}_{d-1} . \end{array} \right. \quad (2.9)$$

Omettons l'écriture du système en (z_{d-1}, π_{d-1}) qui ne sera pas résolu car les corrections z_{d-1} et π_{d-1} seront figées pendant quelques itérations en temps, donc traitées de façon explicite dans le système suivant (2.10). Il nous reste à trouver $(y_{d-2}, p_{d-2}) \in \mathcal{V}_{g,d-2} \times \mathcal{Q}_{d-2}$ solution de:

$$\left\{ \begin{array}{l} \frac{1}{Re}((y_{d-2} + z_{d-1}, \tilde{y}_{d-2})) - (p_{d-2} + \pi_{d-1}, \text{div } \tilde{y}_{d-2}) = (\omega, \tilde{y}_{d-2}) \\ (\text{div } y_{d-2} + \text{div } z_{d-1}, \tilde{p}_{d-2}) + \epsilon(p_{d-2} + \pi_{d-1}, \tilde{p}_{d-2}) = (\vartheta, \tilde{p}_{d-2}) \\ \forall \tilde{y}_{d-2} \in \mathcal{V}_{0,d-2} \text{ et } \forall \tilde{p}_{d-2} \in \mathcal{Q}_{d-2} . \end{array} \right. \quad (2.10)$$

Le système (2.10) s'écrit alors sous la forme matricielle suivante :

$$\begin{cases} \mathcal{A}_{gg}^{(d-1)} Y_{d-2} + \mathcal{A}_{gf}^{(d-1)} Z_{d-1} + {}^T\mathcal{B}_{gg}^{(d-1)} P_{d-2} + {}^T\mathcal{B}_{fg}^{(d-1)} \Pi_{d-1} = \\ \quad \Pi_{d-2} (\tilde{F}_{d-1}) \\ \mathcal{B}_{gg}^{(d-1)} Y_{d-2} + \mathcal{B}_{gf}^{(d-1)} Z_{d-1} - \epsilon \mathcal{C}_{gg}^{(d-1)} P_{d-2} - \epsilon \mathcal{C}_{gf}^{(d-1)} \Pi_{d-1} = \\ \quad \bar{\Pi}_{d-2} (-\mathcal{B}_{gf}^{(d)} Z_d + \epsilon \mathcal{C}_{gf}^{(d)} \Pi_d) \end{cases} \quad (2.11)$$

Π_{d-2} et $\bar{\Pi}_{d-2}$ étant deux opérateurs discrets de projection sur la grille de niveau $d-2$:

$$\Pi_{d-2} : \mathbb{R}^{2 \times n_{0,d-1}} \rightarrow \mathbb{R}^{2 \times n_{0,d-2}} \quad \text{et} \quad \bar{\Pi}_{d-2} : \mathbb{R}^{m_{d-1}} \rightarrow \mathbb{R}^{m_{d-2}} .$$

Naturellement ces deux opérateurs extraient les composantes hiérarchiques relatives à la grille grossière $d-2$ à partir des vecteurs associés à la grille de niveau $d-1$.

Nous soulignons l'intérêt du système (2.11) qui considère la vraie pression sur la grille $d-2$, pression qui a été hiérarchisée grâce aux formules (2.2). L'expression de P_{d-2} peut être retrouvée en considérant la seconde équation de (2.11), i.e. :

$$P_{d-2} = \frac{1}{\epsilon} \left(\mathcal{C}_{gg}^{(d-1)} \right)^{-1} \left[\mathcal{B}_{gg}^{(d-1)} Y_{d-2} + \mathcal{B}_{gf}^{(d-1)} Z_{d-1} - \epsilon \mathcal{C}_{gf}^{(d-1)} \Pi_{d-1} \right. \\ \left. - \bar{\Pi}_{d-2} (-\mathcal{B}_{gf}^{(d)} Z_d + \epsilon \mathcal{C}_{gf}^{(d)} \Pi_d) \right] . \quad (2.12)$$

Malheureusement cette expression est devenue encore plus compliquée que celle de P_{d-1} (voir la relation (2.7)). Lorsque nous remplaçons la valeur de P_{d-2} donnée par (2.12) dans la première équation de (2.11), nous parvenons à un système dont la seule inconnue est Y_{d-2} , soit :

$$\begin{aligned} \left(\mathcal{A}_{gg}^{(d-1)} + \frac{1}{\epsilon} {}^T\mathcal{B}_{gg}^{(d-1)} \left(\mathcal{C}_{gg}^{(d-1)} \right)^{-1} \mathcal{B}_{gg}^{(d-1)} \right) Y_{d-2} = \\ \Pi_{d-2} (\tilde{F}_{d-1}) - \mathcal{A}_{gf}^{(d-1)} Z_{d-1} - {}^T\mathcal{B}_{fg}^{(d-1)} \Pi_{d-1} \\ - \frac{1}{\epsilon} {}^T\mathcal{B}_{gg}^{(d-1)} \left(\mathcal{C}_{gg}^{(d-1)} \right)^{-1} \mathcal{B}_{gf}^{(d-1)} Z_{d-1} + {}^T\mathcal{B}_{gg}^{(d-1)} \left(\mathcal{C}_{gg}^{(d-1)} \right)^{-1} \mathcal{C}_{gf}^{(d-1)} \Pi_{d-1} \\ + \frac{1}{\epsilon} {}^T\mathcal{B}_{gg}^{(d-1)} \left(\mathcal{C}_{gg}^{(d-1)} \right)^{-1} \bar{\Pi}_{d-2} (-\mathcal{B}_{gf}^{(d)} Z_d + \epsilon \mathcal{C}_{gf}^{(d)} \Pi_d) . \end{aligned} \quad (2.13)$$

Le système (2.13) reste semblable à (2.8) relativement au premier membre ; toutefois ce nouveau système possède un nombre bien supérieur de termes au second membre. De façon récurrente, à chaque étape de hiérarchisation du problème de Stokes il paraîtra un certain nombre de termes supplémentaires. Ceci rend l'écriture de la formulation matricielle compliquée et l'application numérique ardue (il faut noter que le second membre de l'équation relative à la divergence de la vitesse sera différent selon le niveau hiérarchique).

La seconde approche réside dans l'introduction d'une nouvelle pression afin de simplifier l'écriture de la hiérarchisation des équations. Notons immédiatement que cette

nouvelle pression ne conserve aucun lien avec la pression physique ; nous l'appelons pression *mathématique* afin qu'elle soit distinguée de la vraie pression.

D'après la remarque 12, il est impossible de hiérarchiser à nouveau le système (2.8), qui est nodal sur la grille $d - 1$. Afin de hiérarchiser de façon récurrente le problème de Stokes, nous cherchons à regarder le système (2.8) comme un nouveau système pénalisé, projeté sur la grille grossière. Pour cela introduisons une pression auxiliaire, notée $\bar{P}_{d-1}$, différente de P_{d-1} . Le vecteur colonne $\bar{P}_{d-1}$ a pour composantes les valeurs $\bar{p}_{d-1,1}, \dots, \bar{p}_{d-1,m_{d-1}}$ et on note

$$\bar{P}_{d-1} = \sum_{i=1}^{m_{d-1}} \bar{p}_{d-1,i} \xi_{d-1,i}$$

la nouvelle pression appartenant au sous-espace $\mathcal{Q}_{d-1}$. La nouvelle inconnue $\bar{p}_{d-1}$ n'est plus la projection de la pression nodale sur la grille grossière $d - 1$, mais une pression qui correspond à l'artifice de pénaliser l'équation de continuité discrète relative à la grille de niveau $d - 1$, c'est-à-dire qu'elle vérifie :

$$(\operatorname{div} y_{d-1}, \tilde{p}_{d-1}) + \epsilon(\bar{p}_{d-1}, \tilde{p}_{d-1}) = 0 \quad \forall \tilde{p}_{d-1} \in \mathcal{Q}_{d-1}. \quad (2.14)$$

On note de plus, en regardant le système (2.3), que le terme $\epsilon(\bar{p}_{d-1}, \tilde{p}_{d-1})$ contient la contribution venant aussi de $(\operatorname{div} z_d, \tilde{p}_{d-1})$ et $\epsilon(\pi_d, \tilde{p}_{d-1})$, le premier de ces deux termes étant du même ordre que $(\operatorname{div} y_{d-1}, \tilde{p}_{d-1})$.

Considérant cette pression *mathématique*, nous réécrivons (2.8) sous la forme matricielle suivante :

$$\begin{cases} \mathcal{A}_{gg}^{(d)} Y_{d-1} + {}^T \mathcal{B}_{gg}^{(d)} \bar{P}_{d-1} = \bar{F}_{d-1} \\ \mathcal{B}_{gg}^{(d)} Y_{d-1} - \epsilon \mathcal{C}_{gg}^{(d)} \bar{P}_{d-1} = 0 \end{cases} \quad (2.15)$$

Désormais nous pouvons hiérarchiser le système (2.15) afin d'obtenir deux systèmes qui correspondent à la partition des espaces de discrétisation suivante : $\mathcal{V}_{g,d-1} = \mathcal{V}_{g,d-2} + \mathcal{W}_{g,d-1}$ et $\mathcal{Q}_{d-1} = \mathcal{Q}_{d-2} + \mathcal{R}_{d-1}$. En effet la nouvelle pression peut être ainsi hiérarchisée : $\bar{p}_{d-1} = \bar{p}_{d-2} + \bar{\pi}_{d-1}$ avec $\bar{p}_{d-2} \in \mathcal{Q}_{d-2}$ et $\bar{\pi}_{d-1} \in \mathcal{R}_{d-1}$. Rappelons que les vecteurs Z_{d-1} et $\bar{\Pi}_{d-1}$ seront traités explicitement et que le système dans ces deux inconnues ne sera pas résolu. Le système restant s'écrit :

$$\begin{cases} \mathcal{A}_{gg}^{(d-1)} Y_{d-2} + \mathcal{A}_{gf}^{(d-1)} Z_{d-1} + {}^T \mathcal{B}_{gg}^{(d-1)} \bar{P}_{d-2} + {}^T \mathcal{B}_{gf}^{(d-1)} \bar{\Pi}_{d-1} = \Pi_{d-2} (\bar{F}_{d-1}) \\ \mathcal{B}_{gg}^{(d-1)} Y_{d-2} + \mathcal{B}_{gf}^{(d-1)} Z_{d-1} - \epsilon \mathcal{C}_{gg}^{(d-1)} \bar{P}_{d-2} - \epsilon \mathcal{C}_{gf}^{(d-1)} \bar{\Pi}_{d-1} = 0. \end{cases} \quad (2.16)$$

A partir de la seconde équation de (2.16) on trouve

$$\bar{P}_{d-2} = \frac{1}{\epsilon} \left(\mathcal{C}_{gg}^{(d-1)} \right)^{-1} \left(\mathcal{B}_{gg}^{(d-1)} Y_{d-2} + \mathcal{B}_{gf}^{(d-1)} Z_{d-1} - \epsilon \mathcal{C}_{gf}^{(d-1)} \bar{\Pi}_{d-1} \right) \quad (2.17)$$

qui, remplacée dans la première équation de (2.16), nous donne le système dans la seule

inconnue Y_{d-2} :

$$\begin{aligned} \left(\mathcal{A}_{gg}^{(d-1)} + \frac{1}{\epsilon} \mathcal{T} \mathcal{B}_{gg}^{(d-1)} \left(\mathcal{C}_{gg}^{(d-1)} \right)^{-1} \mathcal{B}_{gg}^{(d-1)} \right) Y_{d-2} = \\ \Pi_{d-2} \left(\bar{F}_{d-1} \right) - \mathcal{A}_{gf}^{(d-1)} Z_{d-1} - \mathcal{T} \mathcal{B}_{fg}^{(d-1)} \bar{\Pi}_{d-1} \\ - \frac{1}{\epsilon} \mathcal{T} \mathcal{B}_{gg}^{(d-1)} \left(\mathcal{C}_{gg}^{(d-1)} \right)^{-1} \mathcal{B}_{gf}^{(d-1)} Z_{d-1} + \mathcal{T} \mathcal{B}_{gg}^{(d-1)} \left(\mathcal{C}_{gg}^{(d-1)} \right)^{-1} \mathcal{C}_{gf}^{(d-1)} \bar{\Pi}_{d-1} . \end{aligned} \quad (2.18)$$

Naturellement ce procédé peut être réitéré sur les niveaux encore plus grossiers, de la même manière que celle décrite auparavant. A chaque étape de hiérarchisation une nouvelle pression doit être introduite, celle-ci étant calculée par une formule analogue à (2.14).

2.1.2 Navier-Stokes en base hiérarchique

De la même manière que pour le problème de Stokes en base hiérarchique, écrivons la discrétisation spatiale correspondante à la décomposition des espaces $\mathcal{V}_{g,d} = \mathcal{V}_{g,d-1} + \mathcal{W}_{g,d}$ et $\mathcal{Q}_d = \mathcal{Q}_{d-1} + \mathcal{R}_d$ pour le problème de Navier-Stokes pénalisé (1.31).

Pour tout $t > 0$, cherchons $(y_{d-1}(t), p_{d-1}(t)) \in \mathcal{V}_{g,d-1} \times \mathcal{Q}_{d-1}$ et $(z_d(t), \pi_d(t)) \in \mathcal{W}_{g,d} \times \mathcal{R}_d$ solutions des systèmes suivants :

$$\left\{ \begin{array}{l} \left(\frac{\partial y_{d-1}}{\partial t} + \frac{\partial z_d}{\partial t}, \tilde{y}_{d-1} \right) + \frac{1}{Re} ((y_{d-1} + z_d, \tilde{y}_{d-1})) + \\ \hat{b}(y_{d-1}, y_{d-1}, \tilde{y}_{d-1}) + \hat{b}_{int}(y_{d-1}, z_d, \tilde{y}_{d-1}) - (p_{d-1} + \pi_d, \text{div } \tilde{y}_{d-1}) = (f, \tilde{y}_{d-1}) \\ (\text{div } y_{d-1} + \text{div } z_d, \tilde{p}_{d-1}) + \epsilon (p_{d-1} + \pi_d, \tilde{p}_{d-1}) = 0 \\ \forall \tilde{y}_{d-1} \in \mathcal{V}_{0,d-1} \quad \text{et} \quad \forall \tilde{p}_{d-1} \in \mathcal{Q}_{d-1} , \end{array} \right. \quad (2.19)$$

$$\left\{ \begin{array}{l} \left(\frac{\partial y_{d-1}}{\partial t} + \frac{\partial z_d}{\partial t}, \tilde{z}_d \right) + \frac{1}{Re} ((y_{d-1} + z_d, \tilde{z}_d)) + \\ \hat{b}(y_{d-1}, y_{d-1}, \tilde{z}_d) + \hat{b}_{int}(y_{d-1}, z_d, \tilde{z}_d) - (p_{d-1} + \pi_d, \text{div } \tilde{z}_d) = (f, \tilde{z}_d) \\ (\text{div } y_{d-1} + \text{div } z_d, \tilde{\pi}_d) + \epsilon (p_{d-1} + \pi_d, \tilde{\pi}_d) = 0 \\ \forall \tilde{z}_d \in \mathcal{W}_{0,d} \quad \text{et} \quad \forall \tilde{\pi}_d \in \mathcal{R}_d . \end{array} \right. \quad (2.20)$$

Le terme $\hat{b}_{int}$ représente toujours les interactions entre les grandes structures et leurs corrections ainsi que les interactions au sein des corrections (voir sa définition exacte en section 1.4 de la partie I).

Comme déjà précisé auparavant, le système (2.20) ne sera jamais résolu, les corrections z_d et π_d étant figées pendant quelques itérations en temps. De plus, les termes en z_d et π_d apparaissant dans l'équation (2.19) sont traités de façon explicite, ainsi que les deux parties du terme non linéaire, ceci d'après les schémas en temps proposés en section 1.5.1 (dans un premier temps on considérera seulement le schéma Euler - Euler).

A chaque instant $t = (n + 1)\Delta t$, avec $n > 0$, le système (2.19) s'écrit sous la forme matricielle suivante :

$$\left\{ \begin{array}{l} (\mathcal{M}_{gg}^{(d)} + \mathcal{A}_{gg}^{(d)}) Y_{d-1}^{n+1} + (\mathcal{M}_{gf}^{(d)} + \mathcal{A}_{gf}^{(d)}) Z_d^n + \mathcal{B}_{gg}^{(d)} P_{d-1}^{n+1} + \mathcal{B}_{fg}^{(d)} \Pi_d^n = \\ F_{d-1} + \mathcal{M}_{gg}^{(d)} Y_{d-1}^n + \mathcal{M}_{gf}^{(d)} Z_d^{n-1} - \hat{b}_g(y_{d-1}^n) - \hat{b}_{int,g}(y_{d-1}^n, z_d^n) \\ \mathcal{B}_{gg}^{(d)} Y_{d-1}^{n+1} + \mathcal{B}_{gf}^{(d)} Z_d^n - \epsilon \mathcal{C}_{gg}^{(d)} P_{d-1}^{n+1} - \epsilon \mathcal{C}_{gf}^{(d)} \Pi_d^n = 0, \end{array} \right. \quad (2.21)$$

où $\hat{b}_g(y_{d-1}^n) = \hat{b}_g(y_{d-1}^n, y_{d-1}^n)$ est le vecteur du terme non linéaire $\hat{b}(y_{d-1}, y_{d-1}, \tilde{y}_{d-1})$ et $\hat{b}_{int,g}(y_{d-1}^n, z_d^n)$ le vecteur du terme non linéaire d'interaction $\hat{b}_{int}(y_{d-1}, z_d, \tilde{y}_{d-1})$, tous les deux projetés sur la grille grossière de niveau $d - 1$.

Comme pour les équations de Stokes, la seconde équation de (2.21) nous offre la valeur de la pression P_{d-1}^{n+1} . Celle-ci peut être remplacée dans la première équation de (2.21); on obtient ainsi, à chaque itération en temps, un système dont la seule inconnue est la vitesse Y_{d-1}^{n+1} :

$$\begin{aligned} & \left(\mathcal{M}_{gg}^{(d)} + \mathcal{A}_{gg}^{(d)} + \frac{1}{\epsilon} \mathcal{B}_{gg}^{(d)} (\mathcal{C}_{gg}^{(d)})^{-1} \mathcal{B}_{gg}^{(d)} \right) Y_{d-1}^{n+1} = \\ & F_{d-1} - \hat{b}_g(y_{d-1}^n) - \hat{b}_{int,g}(y_{d-1}^n, z_d^n) \\ & + \mathcal{M}_{gg}^{(d)} Y_{d-1}^n - (\mathcal{M}_{gf}^{(d)} + \mathcal{A}_{gf}^{(d)}) Z_d^n + \mathcal{M}_{gf}^{(d)} Z_d^{n-1} \\ & - \mathcal{B}_{fg}^{(d)} \Pi_d^n - \frac{1}{\epsilon} \mathcal{B}_{gg}^{(d)} (\mathcal{C}_{gg}^{(d)})^{-1} \mathcal{B}_{gf}^{(d)} Z_d^n + \mathcal{B}_{gg}^{(d)} (\mathcal{C}_{gg}^{(d)})^{-1} \mathcal{C}_{gf}^{(d)} \Pi_d^n. \end{aligned} \quad (2.22)$$

Le découplage récursif des équations de Navier-Stokes, donnant leur hiérarchisation sur plusieurs niveaux, s'effectuera de la même manière que dans le cas du problème de Stokes. Il nous reste toujours à choisir parmi les deux procédures présentées auparavant, i.e. soit la hiérarchisation de la pression physique, soit l'introduction d'une nouvelle pression qui n'a qu'un sens purement mathématique.

2.2 Analyse de la décomposition hiérarchique

La hiérarchisation de la vitesse et de la pression fait apparaître les grandes structures de ces deux quantités physiques et leurs corrections, de taille inférieure. Nous séparons ces structures de différentes échelles afin de négliger, sur des courts intervalles de temps, la variation temporelle des corrections de vitesse et pression en les figeant à la dernière valeur calculée. Pour cela il est nécessaire, dans un premier temps, d'évaluer toutes les quantités qui dépendent des incréments z et π et de les comparer avec les mêmes quantités dépendant des grandes structures (resp. y et p).

L'objet de cette section est donc d'analyser le comportement des différents termes des systèmes (2.19) et (2.20) issus de la projection des équations de Navier-Stokes sur la base hiérarchique relative à deux niveaux de discrétisation, ainsi que des termes apparaissant dans les systèmes obtenus lorsqu'on projette les équations considérées sur la base hiérarchique relative à plusieurs niveaux. Nous concentrons notre attention surtout sur les

termes du système issu de la projection des équations sur la grille grossière (voir (2.19)) car, comme déjà dit plusieurs fois, dans l'algorithme multi-niveaux que nous proposons le système (2.20) ne sera jamais résolu.

Pour effectuer cette analyse, la décomposition de la vitesse et de la pression sur la base hiérarchique est réalisée *a posteriori*, pendant la résolution numérique des équations de Navier-Stokes discrétisées par la méthode de Galerkin classique. La solution considérée dans ces premiers tests évolue assez rapidement vers la solution stationnaire (la figure 2.2 montre la décroissance de l'erreur relative des composantes de la vitesse U). Cette solution est obtenue pour $Re = 5000$, avec un pas de discrétisation spatiale raffiné près des bords ($N_d = 64$ étant le nombre de segments dans chaque direction spatiale), la condition initiale étant la solution du problème de Navier-Stokes pour $Re = 1000$.

2.2.1 Taille des différentes structures

Dans la série de tests suivante considérons $d = 3$ niveaux hiérarchiques. Le maillage nodal $\mathcal{T}_d$ correspond à une discrétisation ayant 65^2 sommets, tandis que la grille la plus grossière $\mathcal{T}_0$ possède seulement 9^2 sommets.

Hiérarchisation de la vitesse

En premier lieu traçons la norme $L^2(\Omega)$ de y_{k-1} et z_k ($k = 2, 3$) ainsi que celle de u_d (voir partie supérieure de la figure 2.3). En particulier on peut observer un facteur d'ordre 3.0-3.5 entre la norme de z_2 et celle de z_3 ; on attend un facteur 4.0 d'après les résultats théoriques attestant une dépendance quadratique en h ($|z_k|_{L^2} = \mathcal{O}(h_k^2)$ et $h_2 = 2h_3$). Cette norme L^2 sera tout le temps calculée en utilisant la matrice diagonale obtenue par le *mass lumping* (noté M.L.) de $\widehat{M}_k$, la matrice de masse relative au niveau k :

$$|v_k|_{L^2} = \left(\sum_{i=1}^{n_k-1} \left(\sum_{j=1}^{n_k-1} \widehat{M}_{i,j}^{(k)} \right) v_{k,i}^2 \right)^{1/2} \quad (2.23)$$

avec $v_k = (v_k^1, v_k^2) \in \mathbf{R}^{2 \times n_k}$ et $\widehat{M}_k = \Delta t \widehat{M}_k$, $\widehat{M}_k$ étant définie à la page 114. Vérifions que cette norme est du même ordre que la vraie norme $L^2(\Omega)$ du vecteur v_k (voir sa définition exacte en section 1.4.1 de la première partie de la thèse, et ses valeurs dans la partie inférieure de la figure 2.3). L'intérêt d'utiliser le *mass lumping* porte sur le coût du calcul, inférieur par rapport au produit matrice-vecteur. Toutefois le calcul avec le *mass lumping* amène à une erreur sur le résultat qui dépend du carré du pas de discrétisation h_k , cette erreur devenant visible dès que le maillage est trop grossier (comparons les graphiques à gauche dans la figure 2.3: la norme des grosses structures doit décroître si le maillage devient grossier).

La norme L^2 des dérivées première et seconde en temps de u_d , y_{k-1} et z_k , tracée en figure 2.4, révèle l'évolution temporelle de l'écoulement et sa convergence vers la solution stationnaire, le régime s'établissant à partir de $T = 5$.

Enfin la norme H^1 des grandes structures et de leurs corrections (cf. figure 2.5) indique que le gradient des z peut être important et presque du même ordre que celui des y , cela

étant vrai surtout dans les zones du domaine où la vitesse présente des forts gradients (consulter [27] pour l'analyse sur la décomposition de la solution stationnaire selon le nombre de Reynolds choisi).

Hiérarchisation de la pression

Rappelons que la pression physique ne joue plus le rôle d'inconnue dans le problème nodal de Navier-Stokes pénalisé. Les valeurs des degrés de liberté de cette quantité se retrouvent à partir des valeurs de la vitesse $\hat{U}_d^n$ grâce au système :

$$\hat{P}_d^n = \frac{1}{\epsilon} \hat{C}_d^{-1} \hat{B}_d \hat{U}_d^n ,$$

forme matricielle correspondante à l'équation de continuité pénalisée. La pression nodale p_d^n peut ensuite être hiérarchisée par (2.2).

En figure 2.6 est représentée, le long de la simulation numérique, la norme L^2 des grandes structures p_{k-1} et des incréments π_k de la pression physique, pour $k = 2, \dots, d$, ainsi que la pression relative à la grille nodale p_d . Le rapport entre π_2 et π_3 est, comme pour les z_k , d'ordre 3.0-3.5 (la hiérarchisation de la pression est analogue à celle de la vitesse). L'ordre des π_k respecte encore la dépendance quadratique en h ($|\pi_k|_{L^2} = \mathcal{O}(h_{k-1}^2)$, $h_{k-1} = 2h_k$ étant le pas de discrétisation de la pression lorsqu'on utilise l'élément fini 4P1-P1). La norme des grandes structures p_{k-1} est correcte en ordre de grandeur, toutefois on peut remarquer l'erreur causée par le mass lumping, le maillage de la pression devenant trop grossier : on ne retrouve plus la décroissance de la norme des p_{k-1} pour h_{k-1} croissant.

Considérons maintenant la pression $\bar{p}_{k-1}$, dite *mathématique*, introduite afin de réitérer facilement la hiérarchisation des équations de Stokes ou Navier-Stokes sur plusieurs niveaux. A la première étape de hiérarchisation de la pression nodale p_d nous obtenons la pression physique $p_d = p_{d-1} + \pi_d$ (p_{d-1} et π_d en trait pointillé en figure 2.7). A partir de la grille de niveau $d-1$, la pression *mathématique* $\bar{p}_{d-1}$ est introduite grâce à la relation (2.14). La norme L^2 de $\bar{p}_{d-1}$ est représentée par le trait épais en figure 2.7 : sa valeur est très grande, variant entre 20 et 40 au cours du temps. La hiérarchisation de $\bar{p}_{d-1}$ nous fournit les grandes structures $\bar{p}_{d-2}$ et les corrections $\bar{\pi}_{d-1}$, elles aussi ayant une norme très grande (lignes en trait continu en figure 2.7).

Anticipons déjà une remarque négative relative à cette pression *mathématique* : bien qu'elle hiérarchise simplement le problème de Navier-Stokes, elle ne semble pas adaptée aux objectifs de notre analyse *a posteriori*. En effet, notre but est de trouver des critères aptes à déterminer à quel moment les incréments de la vitesse et de la pression peuvent être figés. Ces critères doivent se baser sur la valeur de tous les termes apparaissant dans le système de Navier-Stokes hiérarchique et doivent être indépendants du niveau k de discrétisation. Nous verrons dans la suite que la norme des termes dépendant de la pression *mathématique* diffère largement de celle des termes linéaires ou non linéaires en vitesse. Alors, en anticipant sur les conclusions auxquelles nous parviendrons, il nous semble improbable de trouver, dans le cas de cette pression, des critères qui puissent déterminer la profondeur et la longueur de la série des V-cycles.

2.2.2 Evaluation pratique des différents termes

Nous voulons décider si certains termes de l'équation projetée sur la grille grossière peuvent être négligés ou figés sur certains intervalles de temps. Nous avons donc besoin de définir une norme permettant d'évaluer ces termes. Il est important de remarquer que cette norme doit être calculée à partir des composantes des vecteurs de la forme matricielle de notre équation (voir le système (2.22)), ces vecteurs étant les quantités dont on dispose au cours du calcul. De plus, on veut que cette norme ne dépende pas du pas de discrétisation, ce qui nous permettra de faire notre étude sur plusieurs niveaux simultanément.

Afin de voir quelle norme répond à ces contraintes, on peut remarquer que les équations (2.19) ou (2.22) sont la formulation variationnelle de

$$\begin{aligned} \frac{\partial y_{d-1}}{\partial t} + \mathbf{P}_{d-1} \frac{\partial z_d}{\partial t} + \frac{1}{Re} \tilde{\Delta}_{d-1} y_{d-1} + \frac{1}{Re} \mathbf{P}_{d-1} \tilde{\Delta}_d z_d + \\ \mathbf{P}_{d-1} \hat{b}(y_{d-1}) + \mathbf{P}_{d-1} \hat{b}_{int}(y_{d-1}, z_d) + \tilde{\nabla}_{d-1} p_{d-1} + \mathbf{P}_{d-1} \tilde{\nabla}_d \pi_d = \mathbf{P}_{d-1} f, \end{aligned} \quad (2.24)$$

où $\tilde{\Delta}_k$ (resp. $\tilde{\nabla}_k$) est le laplacien (resp. le gradient) discret sur la grille de niveau k et $\mathbf{P}_{d-1}$ l'opérateur de projection orthogonale sur l'espace $\mathcal{V}_{d-1}$. La formulation forte de nos équations sur une grille encore plus grossière (de niveau $k-1$) est obtenue en remplaçant l'indice d par k dans (2.24).

Il semble naturel de choisir la norme $L^2(\Omega)$ pour estimer un de ces termes. Notons T_{k-1} l'un d'entre eux, avec $k = 2, \dots, d$. Etant donné que T_{k-1} est dans l'espace $\mathcal{V}_{k-1}$, il peut s'écrire de la façon suivante :

$$T_{k-1} = \sum_{i=1}^{n_{k-1}} \tau_i \phi_{k-1,i}.$$

Si on note $\tau = (\tau_1, \dots, \tau_{n_{k-1}})^t$, on a alors la relation suivante :

$$\begin{aligned} |T_{k-1}|_{L^2(\Omega)}^2 &= \sum_{i,j=1}^{n_{k-1}} \tau_i \tau_j (\phi_{k-1,i}, \phi_{k-1,j}) \\ &= {}^t \widehat{M}_{k-1} \tau, \end{aligned}$$

où $\widehat{M}_{k-1} = \Delta t \widehat{\mathcal{M}}_{k-1} = \Delta t \mathcal{M}_{gg}^{(k)}$ (voir la définition de $\widehat{\mathcal{M}}_k$ à la page 114); $\widehat{M}_{k-1}$ est la matrice de masse nodale associée à la grille $k-1$ (matrice qui est identique au bloc grossier-grossier de la matrice de masse hiérarchique de niveau k).

Lorsque nous considérons le système (2.22), le terme T_{d-1} y est présent sous la forme d'un vecteur $L_{d-1} = (L_1, \dots, L_{n_{d-1}})^t$ vérifiant la relation suivante :

$$L_{d-1,j} = (T_{d-1}, \phi_{d-1,j}) \quad \text{pour } j = 1, \dots, n_{d-1}.$$

Par exemple, si T_{d-1} est le terme $\mathbf{P}_{d-1} \hat{b}_{int}(y_{d-1}, z_d)$, le vecteur L_{d-1} est alors $\hat{b}_{int,g}(y_{d-1}, z_d)$.

Or, nous voulons exprimer la quantité $|T_{d-1}|_{L^2(\Omega)}^2$ en fonction des composantes $L_1, \dots, L_{n_{d-1}}$ du vecteur L_{d-1} dont nous disposons. On a alors la relation suivante :

$$L_{d-1,j} = (T_{d-1}, \phi_{d-1,j}) = \sum_{i=1}^{n_{d-1}} \tau_i (\phi_{d-1,i}, \phi_{d-1,j}) ,$$

ce qui implique

$$L_{d-1} = \widehat{M}_{d-1} \tau .$$

Comme $\widehat{M}_k$, quel qu'il soit le niveau k , est une matrice symétrique, on en déduit :

$$|T_{d-1}|_{L^2(\Omega)}^2 = L_{d-1}^t \widehat{M}_{d-1}^{-1} L_{d-1} .$$

Afin de minimiser le coût du calcul, nous remplaçons la matrice $\widehat{M}_{d-1}$ par la matrice diagonale obtenue par le *mass lumping* de $\widehat{M}_{d-1}$. On obtient ainsi :

$$|T_{d-1}|_{L^2(\Omega)}^2 = \sum_{i=1}^{n_{d-1}} L_{d-1,i}^2 / \left(\sum_{j=1}^{n_{d-1}} \widehat{M}_{i,j}^{(d-1)} \right) .$$

La norme $L^2(\Omega)$ des termes T_{k-1} relatifs aux grilles plus grossières, s'exprime toujours à partir des composantes des vecteurs L_{k-1} dont on dispose. Pour les termes de divergence venant de la projection de l'équation de continuité sur la grille $k-1$ ($k = 2, \dots, d$), on remplace simplement la matrice $\widehat{M}_{k-1}$ par $\widehat{C}_{k-1} = C_{gg}^{(k)}$ (bloc grossier-grossier de la matrice de masse hiérarchique de niveau k relative à la pression). Enfin, si l'on veut évaluer la norme $L^2(\Omega)$ des termes appartenant aux équations de Navier-Stokes projetées sur la grille fine complémentaire, on remplace $\widehat{M}_{k-1}$ par $\Delta t \mathcal{M}_{ff}^{(k)}$ (bloc fin-fin de la matrice de masse hiérarchique de niveau k relative aux incréments de la vitesse) ou encore par $C_{ff}^{(k)}$ (bloc fin-fin de la matrice de masse hiérarchique de niveau k relative aux incréments de la pression).

Remarque 13 Dans le cas d'un maillage uniforme, tous les coefficients de la matrice diagonale issue du *mass lumping* sont de l'ordre du carré du pas de discrétisation, i.e. $h_{k-1}^2 \simeq \frac{1}{n_{k-1}}$, où $n_{k-1} = \text{card}(\Sigma_{k-1})$. D'après la définition de la norme l^2 d'un vecteur (voir à la page 84), on trouve la relation suivante :

$$|T_{k-1}|_{L^2} \simeq n_{k-1} |L_{k-1}|_{l^2} \quad \text{pour } k = 2, \dots, d . \quad (2.25)$$

Pour ce qui concerne l'analyse *a posteriori* des équations de Burgers, présentée dans la section 1.4 de la première partie, il suffit d'appliquer la relation (2.25) afin d'opérer un "scaling" de toutes les quantités étudiées. Ce "scaling" permet d'obtenir des estimations indépendantes de la grille sur laquelle les équations ont été projetées.

Dorénavant, chaque fois que nous examinerons la norme $L^2(\Omega)$ d'un terme appartenant à nos équations projetées (ou plutôt du vecteur qui lui correspond), nous considérerons la norme L^2 définie ci-dessus.

2.2.3 Etude globale des différents termes

Les termes d'évolution

La figure 2.8 montre l'évolution en temps des composantes y et z du terme variationnel d'évolution $(\frac{\partial u_d}{\partial t}, \tilde{u}_d)$. A chaque itération en temps n et pour chaque niveau hiérarchique $k = 2, \dots, d$, nous avons tracé la norme L^2 des vecteurs suivants :

$$\begin{aligned} \left(\frac{\partial y_{k-1}^n}{\partial t}, \phi_{k-1,j} \right) \quad \text{et} \quad \left(\frac{\partial z_k^n}{\partial t}, \phi_{k-1,j} \right) & \quad \text{pour } 1 \leq j \leq n_{k-1}, \\ \left(\frac{\partial y_{k-1}^n}{\partial t}, \psi_{k,j} \right) \quad \text{et} \quad \left(\frac{\partial z_k^n}{\partial t}, \psi_{k,j} \right) & \quad \text{pour } n_{k-1} + 1 \leq j \leq n_k. \end{aligned} \quad (2.26)$$

Comme déjà remarqué pour les équations de Burgers, la différence entre les tailles des termes (2.26) n'est pas significative, bien que la norme des termes évolutifs en z_k (projetés sur la grille grossière ou fine) soit légèrement inférieure à celle des termes évolutifs en y_{k-1} (projetés sur les mêmes grilles).

La norme L^2 des deux premiers vecteurs de (2.26) est représentée en figure 2.9 et est comparée à la norme des variations sur un pas de temps de ces termes (il s'agit des termes projetés sur la grille grossière). Rappelons que les vecteurs donnant les variations sur un pas Δt sont définis, pour $1 \leq j \leq n_{k-1}$, de la façon suivante :

$$\left(\frac{\partial y_{k-1}^n}{\partial t} - \frac{\partial y_{k-1}^{n-1}}{\partial t}, \phi_{k-1,j} \right) = \Delta t \left(\sum_{i=1}^{n_{k-1}} \frac{y_{k-1,i}^n - 2y_{k-1,i}^{n-1} + y_{k-1,i}^{n-2}}{\Delta t^2} \cdot \phi_{k-1,i}, \phi_{k-1,j} \right), \quad (2.27)$$

$$\left(\frac{\partial z_k^n}{\partial t} - \frac{\partial z_k^{n-1}}{\partial t}, \phi_{k-1,j} \right) = \Delta t \left(\sum_{i=n_{k-1}+1}^{n_k} \frac{z_{k,i}^n - 2z_{k,i}^{n-1} + z_{k,i}^{n-2}}{\Delta t^2} \cdot \psi_{k,i}, \phi_{k-1,j} \right). \quad (2.28)$$

On observe que les variations en temps des termes considérés sont nettement inférieures aux termes mêmes. De plus la norme de (2.28), variation du terme que l'on espère figer, est plus petite que la norme de (2.27), variation en temps du terme que nous appellerons dominant (traité implicitement dans le système relatif à la projection des équations de Navier-Stokes sur la grille grossière de niveau $k-1$). Comme déjà noté en section 2.2.1, la convergence de la solution vers une solution stationnaire se traduit dans une décroissance de la norme et un lissage de la variation en temps des termes évolutifs, phénomènes qui se manifestent à peu près dès $T = 5$.

Remarque 14 *Le terme évolutif en y (resp. en z) projeté sur la grille grossière possède une norme comparable au même terme en y (resp. en z) projeté sur la grille fine complémentaire. Ceci est dû au choix de la norme L^2 pour évaluer la taille des différents termes apparaissant dans les équations de Navier-Stokes en base hiérarchique. En effet, cette norme ne dépend plus du pas de discrétisation, ce qui revient à dire qu'elle est indépendante de la grille sur laquelle nous projetons nos équations.*

Le même comportement est présent dans tous les autres termes, lorsqu'on les projette sur la grille grossière ou sur la grille fine complémentaire. Nous observons cette équivalence d'ordre de grandeur en figure 2.14 pour les termes de diffusion, dans la partie

supérieure de la figure 2.20 pour les termes non linéaires, en figure 2.26 pour les termes de divergence et enfin en figure 2.32 (resp. 2.38) pour le gradient de la pression physique (resp. mathématique).

Les termes de diffusion

L'évolution de la valeur des composantes y et z du terme variationnel de diffusion $\frac{1}{Re}((u_d, \tilde{u}_d))$ est tracée en figure 2.14. A chaque instant $t = n\Delta t$ et pour chaque niveau hiérarchique $k = 2, \dots, d$, on a représenté la norme L^2 des vecteurs suivants :

$$\begin{aligned} \frac{1}{Re}((y_{k-1}^n, \phi_{k-1,j})) \quad \text{et} \quad \frac{1}{Re}((z_k^n, \phi_{k-1,j})) & \quad \text{pour } 1 \leq j \leq n_{k-1}, \\ \frac{1}{Re}((y_{k-1}^n, \psi_{k,j})) \quad \text{et} \quad \frac{1}{Re}((z_k^n, \psi_{k,j})) & \quad \text{pour } n_{k-1} + 1 \leq j \leq n_k. \end{aligned} \quad (2.29)$$

Tous ces termes possèdent une norme presque constante au cours de la simulation numérique et de plus on constate que la différence entre les normes est très réduite, ce qui empêche d'en négliger certains par rapport aux autres.

En figure 2.15 nous montrons les termes de diffusion projetés sur la grille grossière et à côté leurs variations sur un pas de temps, définies par :

$$\frac{1}{Re}((y_{k-1}^n - y_{k-1}^{n-1}, \phi_{k-1,j})) = \frac{1}{Re} \left(\sum_{i=1}^{n_{k-1}} (y_{k-1,i}^n - y_{k-1,i}^{n-1}) \cdot \nabla \phi_{k-1,i}, \nabla \phi_{k-1,j} \right), \quad (2.30)$$

$$\frac{1}{Re}((z_k^n - z_k^{n-1}, \phi_{k-1,j})) = \frac{1}{Re} \left(\sum_{i=n_{k-1}+1}^{n_k} (z_{k,i}^n - z_{k,i}^{n-1}) \cdot \nabla \psi_{k,i}, \nabla \phi_{k-1,j} \right), \quad (2.31)$$

pour $1 \leq j \leq n_{k-1}$. Les tracés de la norme de ces variations en temps, beaucoup plus petites que les termes mêmes, montrent des légères oscillations à peu près jusqu'à $T = 5$. Ensuite, même les variations en temps se stabilisent, puisque le tourbillon principal et les contre-tourbillons sont désormais en place. Comme pour les termes d'évolution, la norme de (2.31) reste inférieure à celle de (2.30), variation en temps du terme dominant.

Les termes non linéaires

En figure 2.20 est représentée, pour $k = 2, \dots, d$, l'évolution au cours du temps de la norme L^2 des vecteurs suivants :

en haut :

$$\hat{b}(u_d^n, u_d^n, \phi_{k-1,j}) \quad (1 \leq j \leq n_{k-1}) \quad \text{et} \quad \hat{b}(u_d^n, u_d^n, \psi_{k,j}) \quad (n_{k-1} + 1 \leq j \leq n_k); \quad (2.32)$$

en bas :

$$\hat{b}(y_{k-1}^n, y_{k-1}^n, \phi_{k-1,j}) \quad \text{et} \quad \hat{b}_{int}(y_{k-1}^n, z_k^n, \phi_{k-1,j}) \quad \text{pour } 1 \leq j \leq n_{k-1}. \quad (2.33)$$

Il s'agit, pour (2.32), de la projection du terme non linéaire de Navier-Stokes respectivement sur la grille grossière et sur la grille fine complémentaire; ces deux quantités sont

du même ordre (voir remarque 14). Décomposons le terme non linéaire projeté sur $\mathcal{V}_{k-1}$ dans le terme non linéaire dépendant des grandes structures et dans celui d'interaction (voir (2.33)). Il est à noter que $\hat{b}(y_{k-1}^n, y_{k-1}^n, \phi_{k-1,j})$ est du même ordre de grandeur que $\hat{b}(u_d^n, u_d^n, \phi_{k-1,j})$ et de norme presque constante, tandis que $\hat{b}_{int}(y_{k-1}^n, z_k^n, \phi_{k-1,j})$ possède une norme un peu plus petite, qui oscille au début de la simulation numérique.

La variation en temps des termes (2.33), comparée à la norme L^2 des termes mêmes, est montrée en figure 2.21. Il s'agit des vecteurs :

$$\hat{b}(y_{k-1}^n, y_{k-1}^n, \phi_{k-1,j}) - \hat{b}(y_{k-1}^{n-1}, y_{k-1}^{n-1}, \phi_{k-1,j}) \quad \text{et} \quad (2.34)$$

$$\hat{b}_{int}(y_{k-1}^n, z_k^n, \phi_{k-1,j}) - \hat{b}_{int}(y_{k-1}^{n-1}, z_k^{n-1}, \phi_{k-1,j}), \quad (2.35)$$

pour $1 \leq j \leq n_{k-1}$. Il faut noter que des oscillations dans les termes (2.34) et (2.35) apparaissent jusqu'à $T = 6$, le regime stationnaire s'établissant par la suite. De plus la norme de la variation en temps de $\hat{b}_{int}$ reste légèrement inférieure à celle de (2.34). Donc il semble possible de figer le terme non linéaire d'interaction pendant quelques itérations en temps.

Les termes de divergence

Pendant cette analyse *a posteriori*, considérons l'équation de continuité non pénalisée ($\text{div } u_d^n, \tilde{p}_d$) = 0, la pénalisation étant seulement une technique de résolution. La figure 2.26 prouve qu'à chaque itération en temps n et sur chaque grille $k = 2, \dots, d$ la divergence des grandes structures y_{k-1} et celle des incréments z_k projetées sur la grille grossière (i.e. sur $\mathcal{Q}_{k-1}$) ou sur la grille fine (i.e. $\mathcal{R}_k$), sont équivalentes en norme. En effet, à partir de l'équation de continuité sous forme variationnelle, on déduit que :

$$\begin{aligned} (\text{div } y_{k-1}^n, \xi_{k-1,j}) &= -(\text{div } z_k^n, \xi_{k-1,j}) && \text{pour } 1 \leq j \leq m_{k-1}, \\ (\text{div } y_{d-1}^n, \eta_{k,j}) &= -(\text{div } z_k^n, \eta_{k,j}) && \text{pour } m_{k-1} + 1 \leq j \leq m_k. \end{aligned} \quad (2.36)$$

De ce résultat on note l'impossibilité de négliger le terme $(\text{div } z_k^n, \tilde{p}_{k-1})$ ($k = 2, \dots, d$) dans l'équation pénalisée. On remarque de plus l'ordre de grandeur non négligeable de chacun des termes (2.36) qui, à regime stationnaire établi, deviennent constants. Cette valeur élevée s'explique par le choix de la norme L^2 fait en section 2.2.2. La norme l^2 des vecteurs (2.36) donnerait chacun de ces termes relativement petit (de l'ordre de 10^{-4} pour le niveau $k = 3$ et de $5 \cdot 10^{-4}$ pour $k = 2$). Après chaque résolution du système linéaire issu de la discrétisation en temps des équations de Navier-Stokes pénalisées, nous avons aussi vérifié la norme l^2 du vecteur $(\text{div } u_d^n, \xi_{d,j})$ ($1 \leq j \leq m_d$). Celle-ci étant de l'ordre de 10^{-7} durant toute la simulation, le résultat nous permet d'affirmer qu'une bonne approximation numérique de la condition d'incompressibilité a été obtenue.

Les variations en temps de la divergence de y_{k-1} et de z_k , projetées sur $\mathcal{Q}_{k-1}$, sont déterminées par les vecteurs suivants :

$$(\text{div } (y_{k-1}^n - y_{k-1}^{n-1}), \xi_{k-1,j}) = \left(\sum_{l=1}^2 \frac{\partial}{\partial x_l} \left[\sum_{i=1}^{n_{k-1}} (y_{k-1,i}^{l,n} - y_{k-1,i}^{l,n-1}) \cdot \phi_{k-1,i} \right], \xi_{k-1,j} \right), \quad (2.37)$$

$$(\operatorname{div}(z_k^n - z_k^{n-1}), \xi_{k-1,j}) = \left(\sum_{l=1}^2 \frac{\partial}{\partial x_l} \left[\sum_{i=n_{k-1}+1}^{n_k} (z_{k,i}^{l,n} - z_{k,i}^{l,n-1}) \cdot \psi_{k,i} \right], \xi_{k-1,j} \right), \quad (2.38)$$

avec $1 \leq j \leq m_{k-1}$. En figure 2.27 sont représentées les normes de (2.37) et (2.38), ces variations sur un pas Δt étant comparées aux termes mêmes. Les variations en temps ont une taille nettement inférieure aux termes respectifs et oscillent beaucoup jusqu'à l'établissement du regime stationnaire.

Rappelons (cf. la section 2.1.1) que notre intention est de résoudre principalement l'équation en y où les termes contenant z et π seront figés. En particulier la deuxième équation de (2.19) sera remplacée par

$$(\operatorname{div} y_{d-1}, \tilde{p}_{d-1}) + \epsilon(p_{d-1}, \tilde{p}_{d-1}) = -(\operatorname{div} z_d, \tilde{p}_{d-1}) - \epsilon(\pi_d, \tilde{p}_{d-1}) \quad \forall \tilde{p}_{d-1} \in \mathcal{Q}_{d-1}, \quad (2.39)$$

où le membre de droite sera figé. Malheureusement, pour tout $k = \bar{2}, \dots, d$, les variations de la divergence des grosses structures y_{k-1} et de la divergence des incréments z_k sont du même ordre. Il n'est donc pas évident a priori que l'on puisse figer la divergence des z_k sans introduire une erreur supplémentaire dans l'approximation de la solution.

Le problème lié aux termes de divergence reste encore ouvert. La technique de pénalisation, introduite dans l'espoir de résoudre les difficultés liées au traitement de la condition d'incompressibilité, n'aboutit effectivement à aucun résultat satisfaisant. Nous montrerons dans la suite comment les problèmes liées aux termes de divergence se retrouvent dans les termes de pénalisation. De plus, pour réitérer la hiérarchisation il faut se ramener toujours à un problème en vitesse et en pression (cf. la section 2.1.1). Les difficultés associées aux termes de divergence se manifestent aussi dans la hiérarchisation de l'équation de continuité et par conséquent dans la hiérarchisation de la pression physique ou *mathématique*.

Les termes du gradient de pression

On distingue l'analyse de ces termes selon le type de pression considérée.

Gradient de pression physique

La figure 2.32 montre l'évolution en temps du gradient de p_{k-1} et de celui de π_k , projetés sur la grille grossière et sur la grille fine complémentaire. Il s'agit, pour $k = 2, \dots, d$ et pour chaque itération en temps n , de la norme L^2 des vecteurs :

$$\begin{aligned} (\nabla p_{k-1}^n, \phi_{k-1,j}) \text{ et } (\nabla \pi_k^n, \phi_{k-1,j}) & \quad \text{pour } 1 \leq j \leq n_{k-1}, \\ (\nabla p_{k-1}^n, \psi_{k,j}) \text{ et } (\nabla \pi_k^n, \psi_{k,j}) & \quad \text{pour } n_{k-1} + 1 \leq j \leq n_k. \end{aligned} \quad (2.40)$$

On évalue un rapport de l'ordre de 10 entre la norme du gradient des grandes structures de la pression et celle du gradient des incréments π_k . La norme L^2 de ces quantités, excepté de légères oscillations au début de la simulation, est pratiquement constante. Toutefois, la différence d'ordre entre les gradients n'est pas suffisante pour nous permettre d'éliminer les termes en π_k dans les équations projetées.

Les variations sur un pas Δt des termes ∇p_{k-1} et $\nabla \pi_k$, projetés sur $\mathcal{V}_{k-1}$, sont ainsi définies :

$$(\nabla(p_{k-1}^n - p_{k-1}^{n-1}), \phi_{k-1,j}) = \left(\sum_{i=1}^{m_{k-1}} (p_{k-1,i}^n - p_{k-1,i}^{n-1}) \cdot \nabla \xi_{k-1,i}, \phi_{k-1,j} \right), \quad (2.41)$$

$$(\nabla(\pi_k^n - \pi_k^{n-1}), \phi_{k-1,j}) = \left(\sum_{i=m_{k-1}+1}^{m_k} (\pi_{k,i}^n - \pi_{k,i}^{n-1}) \cdot \nabla \eta_{k,i}, \phi_{k-1,j} \right), \quad (2.42)$$

avec $1 \leq j \leq n_{k-1}$. La norme L^2 des variations (2.41) et (2.42), comparée aux termes mêmes, est tracée en figure 2.33. On observe également pour les termes gradient de pression que les variations sur un pas de temps sont nettement inférieures à ces mêmes termes. De plus, la norme de (2.42), variation du gradient des corrections que nous espérons figer pendant quelques itérations en temps, reste inférieure à la norme de (2.41), variation en temps du terme dominant. L'évolution des variations en temps du gradient de la pression physique est donc analogue à celle des autres termes appartenant à l'équation de conservation de la quantité de mouvement.

Gradient de pression *mathématique*

En figure 2.38 est représentée la norme des termes (2.40) pour $k = 3$ et de leurs analogues obtenus en remplaçant p_{k-1}^n par $\bar{p}_{k-1}^n$ et π_k^n par $\bar{\pi}_k^n$ lorsque $k = 2$. La norme du gradient de la pression *mathématique* diffère énormément selon le niveau de discrétisation k . Ceci était déjà visible lorsqu'on évaluait la norme de la pression *mathématique* (voir figure 2.7). L'énorme différence sur l'ordre de grandeur se retrouve en figure 2.39 qui donne la norme des variations en temps du gradient de la pression *mathématique* à côté de la norme des termes mêmes (pour $k = 3$ il s'agit des vecteurs (2.41) et (2.42) et pour $k = 2$ de leurs correspondants obtenus remplaçant p_{k-1} et π_k respectivement par $\bar{p}_{k-1}$ et $\bar{\pi}_k$). Cependant, notons que le gradient des corrections, ainsi que sa variation sur un pas Δt , restent de norme inférieure aux mêmes termes relatifs aux grandes structures de la pression *mathématique*.

Pour $k = 3$ la valeur du gradient de pression, comme aussi la norme L^2 de sa variation sur un pas de temps, restent d'ordre comparable aux autres termes en vitesse. Il s'agit, nous le rappelons, de la décomposition sur deux niveaux hiérarchiques de la pression physique. En revanche, relativement au niveau $k = 2$, le gradient de la pression *mathématique* (c'est-à-dire des grandes structures $\bar{p}_{k-1}$ et des incréments $\bar{\pi}_k$), comme aussi la variation en temps de ces quantités, possède des valeurs nettement supérieures aux autres termes en vitesse projetés sur la grille de même niveau.

Ces résultats numériques ne nous conviennent pas car, à partir de ces quantités, il est impossible de trouver des critères définissant la profondeur et la longueur des V-cycles en espace et en temps. En pratique, les critères numériques que nous définirons doivent être valables pour le niveau $k = 3$ comme aussi pour $k = 2$. La difficulté ici réside dans l'impossibilité que ces critères considèrent en même temps les valeurs des termes en vitesse et en pression, ces termes étant trop différents en ordre de grandeur.

En conclusion, l'introduction de la pression *mathématique* permet la hiérarchisation sur plusieurs niveaux des équations de Navier-Stokes mais ne nous convient pas car les termes introduits n'ont plus des variations en temps négligeables. C'est pour cette raison que dorénavant nous arrêtons de la considérer comme une alternative à la pression physique. Désormais, chaque fois qu'on nomme la pression, ou les quantités qui en dépendent, il sera question de la variable physique.

Face à cette difficulté notre premier objectif sera la mise au point d'une méthode multi-résolution auto-adaptative sur seulement deux niveaux hiérarchiques. L'extension de la méthode à plusieurs niveaux fera partie de développements futurs ; à ce moment-là une autre décomposition sera probablement nécessaire.

Les termes de pénalisation

Lorsque nous éliminons la pression dans la formulation matricielle des équations de Navier-Stokes pénalisées (2.21), nous obtenons, à chaque itération en temps n , les termes suivants :

$$\frac{1}{\epsilon} T_{\mathcal{B}_{gg}^{(d)}} \left(\mathcal{C}_{gg}^{(d)} \right)^{-1} \mathcal{B}_{gg}^{(d)} Y_{d-1}^n, \quad \frac{1}{\epsilon} T_{\mathcal{B}_{gg}^{(d)}} \left(\mathcal{C}_{gg}^{(d)} \right)^{-1} \mathcal{B}_{gf}^{(d)} Z_d^n \quad \text{et} \quad T_{\mathcal{B}_{gg}^{(d)}} \left(\mathcal{C}_{gg}^{(d)} \right)^{-1} \mathcal{C}_{gf}^{(d)} \Pi_d^n \quad (2.43)$$

(voir le système (2.22) correspondant aux équations projetées sur la grille grossière $d-1$). Il s'agit respectivement du terme de pénalisation en y , de celui de pénalisation en z et d'un terme de correction dépendant de π . Nous présentons aussi, pour les avoir calculés en même temps, les deux termes de pénalisation et celui de correction relatifs aux grilles plus grossières (il suffit de remplacer dans (2.43) le niveau d par k pour $k = 2, \dots, d-1$). Les trois termes (2.43) apparaissent chaque fois qu'on hiérarchise récursivement les équations de Navier-Stokes. Cependant, sauf pour une hiérarchisation sur seulement deux niveaux, ces termes ne sont pas les seuls introduits par la pénalisation dans l'étape d'élimination de la pression physique. Déjà relativement au niveau $d-2$ (voir (2.13)) il nous reste à évaluer la taille d'un terme de la forme

$$\frac{1}{\epsilon} T_{\mathcal{B}_{gg}^{(d-1)}} \left(\mathcal{C}_{gg}^{(d-1)} \right)^{-1} \overline{\Pi}_{d-2} \left(-\mathcal{B}_{gf}^{(d)} Z_d + \epsilon \mathcal{C}_{gf}^{(d)} \Pi_d \right),$$

quantité probablement du même ordre que le terme de pénalisation en z .

Les variations sur un pas de temps des termes (2.43) sont respectivement

$$\frac{1}{\epsilon} T_{\mathcal{B}_{gg}^{(d)}} \left(\mathcal{C}_{gg}^{(d)} \right)^{-1} \mathcal{B}_{gg}^{(d)} \left(Y_{d-1}^n - Y_{d-1}^{n-1} \right), \quad (2.44)$$

$$\frac{1}{\epsilon} T_{\mathcal{B}_{gg}^{(d)}} \left(\mathcal{C}_{gg}^{(d)} \right)^{-1} \mathcal{B}_{gf}^{(d)} \left(Z_d^n - Z_d^{n-1} \right) \quad (2.45)$$

et

$$T_{\mathcal{B}_{gg}^{(d)}} \left(\mathcal{C}_{gg}^{(d)} \right)^{-1} \mathcal{C}_{gf}^{(d)} \left(\Pi_d^n - \Pi_d^{n-1} \right). \quad (2.46)$$

La norme L^2 des premiers deux termes de (2.43), ainsi que celle de (2.44) et (2.45), est tracée en figure 2.40. Notons immédiatement la différence d'ordre de grandeur entre

les termes de pénalisation et tous les autres termes en vitesse ou en pression évalués auparavant. Cette différence remarquable est due au facteur $\frac{1}{\epsilon}$ présent dans les termes de pénalisation. De la même manière que pour les autres termes, on observe que les variations sur un pas de temps sont d'ordre de grandeur inférieur aux termes de pénalisation mêmes. Cependant, on souligne que les termes de pénalisation en y et en z sont exactement de la même taille, ceci étant vrai aussi pour les variations respectives sur un pas de temps. Les termes de pénalisation sont en effet strictement liés aux termes de divergence.

Auparavant nous avons observé que la divergence des grandes structures y_{k-1} et celle des incréments z_k , projetées sur $\mathcal{Q}_{k-1}$, sont équivalentes en norme. Ce résultat se reflète dans les termes de pénalisation. Les deux problèmes sont donc inséparables : la technique de pénalisation a simplement transféré le problème d'un terme à l'autre, sans le résoudre.

De plus, dans le cas des équations pénalisées, le traitement multi-niveaux semble encore plus compliqué, car les termes introduits sont d'un ordre de grandeur beaucoup plus élevé que les autres termes en vitesse. Probablement, pour ce qui concerne ce dernier problème, il ne faudra trouver qu'un *scaling* adéquat permettant de ramener la norme de ces termes au même ordre que celle des autres termes en vitesse.

L'analyse du rôle joué par les termes provenant de la pénalisation nous conduit à une nouvelle façon de concevoir le problème. Les termes de pénalisation ont été introduits afin de résoudre numériquement les équations de Navier-Stokes. Par conséquent, on peut probablement négliger ces termes dans la mise au point des critères permettant de définir la profondeur et le nombre de V-cycles à réaliser lorsqu'on appliquera notre méthode multi-résolution. Dans cette optique, les critères devront prendre en compte seulement les termes à figer, termes qui dépendent des corrections de la vitesse et de la pression (c'est-à-dire de z_k et de π_k). Durant la réalisation de la série de V-cycles en espace et en temps, à l'étape correspondant à la résolution des équations relatives à la grille grossière de niveau $d - 1$, nous résoudrons un système pour lequel la divergence de la solution n'est plus nulle, mais égale à une valeur donnée. Il faudra donc considérer le système de Navier-Stokes (2.19) où la deuxième équation est remplacée par l'équation (2.39).

Pour conclure, la figure 2.41 confirme la remarque 11. Le terme de correction en π (troisième terme de (2.43)) est négligeable par rapport au terme de pénalisation en z . En fait, la norme du terme de correction en π (resp. sa variation en temps, soit (2.46)) est bien inférieure à celle du terme de pénalisation en z (resp. au terme (2.45)). Dans la partie inférieure de la figure 2.41 est représentée la norme L^2 de la somme du terme de pénalisation en z et du terme de correction en π . La norme de cette somme est du même ordre que celle du seul terme de pénalisation (confronter la partie inférieure des figures 2.40 et 2.41).

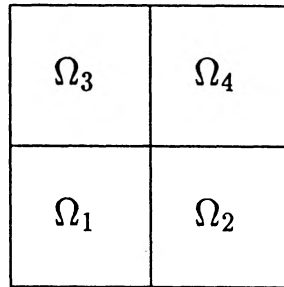
Le terme de correction (resp. sa variation en temps) possède une norme L^2 comparable aux autres termes linéaires et non linéaires en vitesse comme aussi au gradient de la pression (resp. aux variations en temps de ces termes). Négliger ce terme, ce qui revient à négliger le terme de pénalisation $\epsilon C_{gf}^{(d)} \Pi_d$ dans l'équation de la divergence, pourrait néanmoins introduire une erreur supplémentaire dans l'approximation des équations de Navier-Stokes. En revanche, il est certainement possible de le figer sur quelques pas de

temps, car la variation en temps est beaucoup plus petite que le terme de correction même.

2.2.4 Conclusions relatives à l'étude globale

Nous avons présenté dans la section précédente l'évolution au cours de toute la simulation numérique des différents termes qui apparaissent dans les équations de Navier-Stokes décomposées sur la base hiérarchique (et plus en particulier des termes présents dans la projection sur la grille grossière de ces équations). La norme L^2 de tous ces termes a été comparée à celle de leurs variations sur un pas de temps. En résumé, les termes évolutifs, de diffusion et non linéaires, comme aussi le gradient de la pression physique, ont presque le même comportement. Les variations en temps de ces termes sont toutes nettement inférieures aux termes mêmes. De plus, la variation sur un pas de temps de chaque terme en z_k est toujours plus petite de la variation du terme correspondant en y_{k-1} (de même pour le $\nabla\pi_k$ par rapport au ∇p_{k-1}). A partir de ces résultats nous pouvons déduire la possibilité de figer les termes qui dépendent des corrections z_k et π_k durant un court intervalle de temps. En revanche, nous avons vu que les termes de divergence ou de pénalisation n'ont pas un comportement aussi satisfaisant. Nous avons noté que la divergence des grosses structures y_{k-1} est équivalente en norme à la divergence des corrections z_k et surtout que les variations respectives sont du même ordre et évoluent de la même manière. Le même comportement se retrouve dans les termes de pénalisation en y et en z et dans leurs variations respectives, car les termes de pénalisation remplacent ceux de divergence. Le comportement des termes de divergence ou de pénalisation pourrait bloquer notre démarche. Cependant, il convient de relativiser ce problème. En effet, rappelons que seul le terme d'interaction non linéaire est destiné à être figé sur un nombre important de V-cycles; or les variations de celui-ci sont négligeables. Les autres termes peuvent être réévalués de façon économique à la fin de chaque cycle; des variations plus importantes peuvent alors être acceptées pour ceux-ci sans introduire une trop grosse erreur. Pour confirmer ceci, on pourra noter que pour les calculs effectués sur les équations de Burgers (voir la partie I) la condition (1.34) est toujours moins forte que (1.35). Une analyse adéquate nous fournira l'analogue de ces deux critères (1.34) et (1.35) et, vue l'étude faite en section 2.2.3, nous pensons que le même phénomène se produira.

Cette étude nous montre que la décomposition hiérarchique (au moins sur deux niveaux) permet de séparer les échelles et d'obtenir des phénomènes similaires à ceux qui ont été observés en discrétisation spectrale, notamment une variation plus lente des termes contenant les corrections z ou π . Toutefois, en examinant les variations du terme non linéaire d'interaction (voir figure 2.21), on observe que celles-ci ne seront négligeables que sur de très courts intervalles de temps. Sur la base de ces observations nous prévoyons l'impossibilité de mettre en place efficacement la méthode multi-résolution dans le domaine Ω tout entier. Si on cherche à restreindre la résolution de nos équations sur les niveaux associés à des triangulations de Ω plus grossières, on risque de ne plus capter correctement l'évolution et la position des divers contre-tourbillons, comme aussi

FIG. 2.1 – Partition du domaine Ω en quatre zones

d'introduire trop de perturbations dans la solution calculée et, par conséquent, trop d'oscillations dans la vorticité correspondante.

Pour remédier à ce problème, rappelons nous que les éléments finis donnent une discrétisation locale de notre équation et que le même type d'étude effectué de façon locale nous permettra peut-être de déceler d'autres comportements physiques non observables par une étude globale. On peut s'attendre à ce que des phénomènes physiques locaux soient noyés dans une séparation d'échelle globale, comme celle effectuée auparavant. En conséquence, nous tenterons de développer une approche originale, tenant compte de cette localisation, que nous détaillerons au paragraphe suivant. Nous avons choisi de refléter cette localisation par une décomposition de Ω en plusieurs sous-domaines, ce qui nous permettra dans le futur de coupler aisément notre méthode multi-niveaux avec une méthode de décomposition de domaine. En anticipant sur les résultats ci-dessous, nous pouvons envisager le développement d'une méthode multi-niveaux sur chaque sous-domaine.

Notons tout de même qu'il nous faudra être attentifs à la variation en temps des termes de pénalisation. Enfin une autre difficulté a été soulevée, celle de la décomposition de nos équations sur plus de deux niveaux. Les décompositions introduites en section 2.1.1 sont compliquées ou n'apportent pas de résultats satisfaisants. Une autre décomposition sera probablement introduite au moment d'étendre récursivement la hiérarchisation, ceci étant pour le moment un but secondaire.

2.2.5 Etude par zones des différents termes

L'idée de réaliser de façon locale le même type d'étude effectué en section 2.2.3 est née principalement des possibilités offertes par la discrétisation par éléments finis et de la manière dont évolue l'écoulement à l'intérieur de la cavité entraînée. Par leur nature même, les éléments finis permettent de travailler localement dans le domaine Ω . Nous pouvons alors décomposer de façon simple le domaine Ω en plusieurs sous-domaines sans recouvrement. De plus, l'emplacement des contre-tourbillons aux coins du domaine considéré, ainsi que le fait de constater que leur évolution se réalise de manière différente, nous suggère une première décomposition consistant à découper Ω en quatre sous-domaines (voir figure 2.1). Cette décomposition est mise en place de manière triviale et pourra facilement être généralisée dans l'avenir à un nombre supérieur

de sous-domaines qui pourront être aussi de forme différente.

Nous considérons de nouveau tous les termes appartenant aux équations de Navier-Stokes projetées sur la grille grossière. Suivant les notations introduites en section 2.2.2, soit T_{k-1} l'un de termes de (2.24) (formulation forte associée aux équations (2.19) ou (2.22)), ou de l'équation analogue relative à la grille $k-1$. Rappelons encore que chaque terme T_{k-1} est présent dans le système (2.22) (ou le système analogue relatif à la grille $k-1$) sous la forme d'un vecteur L_{k-1} de composantes

$$L_{k-1,j} = (T_{k-1}, \phi_{k-1,j}) \quad \text{pour } j = 1, \dots, n_{k-1},$$

ceci pour chaque niveau $k = 2, \dots, d$.

Soit $l = 1, \dots, 4$; l'indice l indique le sous-domaine Ω_l auquel on restreint l'étude *a posteriori*. Nous évaluons, pour chaque terme grossier, la norme L^2 du vecteur de composantes

$$L_{k-1,j}^l = (T_{k-1}, \phi_{k-1,j}^l),$$

avec j prenant les valeurs correspondant aux fonctions de base qui appartiennent au domaine Ω_l . Pour les termes relatifs à la décomposition de la divergence, le vecteur à évaluer aura pour composantes $(T_{k-1}, \xi_{k-1,j}^l)$. Bien évidemment la norme L^2 est toujours celle introduite en section 2.2.2. Nous effectuons de la même manière les calculs par zones de la variation en temps de tous les termes apparaissant dans les équations projetées sur la grille grossière.

Il serait trop long de détailler l'étude de chacun de ces termes comme nous l'avons fait dans la section 2.2.3. Pour cette raison, les estimations numériques obtenues dans le cadre de l'étude locale sont présentées grâce à une séparation des différents termes en deux groupes: au premier groupe appartiennent les termes évolutifs, de diffusion, non linéaires et le gradient de pression, tandis que le deuxième groupe contient les termes de divergence et ceux de pénalisation. Cette partition est due au comportement assez similaire des termes appartenant à chacun des deux groupes.

Voici brièvement la correspondance entre les termes évalués dans les quatre sous-domaines et les figures présentées à la fin de ce chapitre. Pour chacun des termes analysés, les figures qui montrent les résultats de l'étude locale suivent celles relatives à l'étude globale, ceci afin de comparer aisément les courbes pour permettre une meilleure compréhension du fait que les phénomènes physiques locaux sont noyés dans la séparation d'échelle globale.

Etablissons d'abord la liste des termes réunis dans le premier groupe. Les figures 2.10 et 2.11, 2.16 et 2.17, 2.22 et 2.23 et enfin 2.34 et 2.35 représentent respectivement l'évolution dans les quatre zones de la norme L^2 des vecteurs associés aux termes évolutifs grossiers (voir les deux premiers termes de (2.26)), aux termes de diffusions grossiers (cf. les deux premiers termes de (2.29)), au terme non linéaire grossier et au terme d'interaction non linéaire projeté sur la grille grossière (voir les termes (2.33)) et enfin aux gradients de la pression physique projetés sur la grille grossière (voir les deux premiers termes de (2.40)). Ensuite, les figures 2.12 et 2.13, 2.18 et 2.19, 2.24 et 2.25 et pour conclure 2.36 et 2.37 montrent comment les variations en temps de tous les termes énumérés ci-dessus évoluent dans chacun des quatre sous-domaines (il s'agit respectivement

du découplage des termes (2.27) et (2.28), (2.30) et (2.31), (2.34) et (2.35) et enfin (2.41) et (2.42)).

Mis à part les termes évolutifs, tous les autres termes en y (resp. en p) ou en z (resp. en π), projetés sur la grille grossière, possédaient une norme L^2 globale qui était pratiquement constante. Or, lorsqu'on observe l'évolution de ces mêmes termes dans chacun des sous-domaines, on remarque que la norme oscille, ces variations étant plus significatives dans les sous-domaines Ω_1 et Ω_2 . On retrouve dans chaque zone le fait que les valeurs des termes en z_k (ou en π_k) restent inférieures à celles des termes correspondant en y_{k-1} (ou en p_{k-1}). Enfin, la décroissance de la norme L^2 des termes évolutifs grossiers n'est plus aussi évidente car, même si le régime stationnaire s'est établi à peu près dès $T = 5$, les contre-tourbillons (surtout ceux dans les deux coins inférieurs) continuent à évoluer légèrement ; l'évolution concerne principalement la position du contre-tourbillon et l'intensité de la fonction de courant.

Les variations en temps des termes grossiers en y (resp. en p) ou en z (resp. en π) montrent de façon encore plus incontestable un comportement local qui ne se manifestait pas dans l'analyse globale. Des oscillations très importantes sont visibles dans les courbes représentant ces variations. Les minima et les maxima de ces oscillations ne se manifestent pas en même temps, ceci confirmant le fait que des phénomènes physiques se mettent en place à des instants différents dans les diverses zones du domaine. Une diminution de la valeur de la norme dans une zone est compensée par une croissance de cette valeur dans une autre zone du domaine, ceci étant vrai pour chacun des termes évalués. Ce comportement peut expliquer le fait que la séparation d'échelle globale noyait les phénomènes physiques locaux que maintenant nous arrivons à détecter. Enfin, on observe toujours une variation plus lente des termes en z (ou en π) par rapport aux mêmes termes en y (ou en p), ceci s'avérant dans chaque sous-domaine.

En ce qui concerne les termes restant, les figures 2.28 et 2.29 et les figures 2.42 et 2.43 donnent l'évolution par zones de la norme L^2 des vecteurs associés aux termes de divergence grossiers (il s'agit des deux premiers termes de (2.36)) et aux termes de pénalisation (voir les deux premiers termes de (2.43)), tandis que dans les figures 2.30 et 2.31 et enfin 2.44 et 2.45 sont tracées les variations en temps de ces mêmes termes dans chaque sous-domaine (voir respectivement les termes (2.37) et (2.38) et les termes (2.44) et (2.45)).

Ce dernier groupe de termes présente une évolution locale ayant la même allure que les autres termes appartenant à nos équations projetées. Néanmoins, pour ces termes, les difficultés à surmonter restent toujours les mêmes. Les termes de divergence en y et en z (resp. les variations en temps des termes de divergence grossiers) sont exactement du même ordre, ce résultat se retrouvant dans les termes de pénalisations (resp. dans les variations des termes de pénalisation). Toutefois, pour les mêmes raisons décrites au paragraphe 2.2.4, on devrait pouvoir appliquer notre méthode multi-résolution localement dans chaque sous-domaine. En effet, des variations plus importantes peuvent être admises pour ces termes figés seulement à l'intérieur de chaque cycle.

En regardant en particulier la figure 2.25, qui est relative à l'évolution par zones de la variation du terme non linéaire d'interaction, on observe que ces variations peuvent être

négligeables sur de très longs intervalles de temps. De plus, dans chaque zone l'évolution de cette variation est indépendante des autres zones. Sur la base de ce comportement on envisage le développement d'une méthode multi-niveaux qui exécute des V-cycles sur chaque sous-domaine. Relativement à chaque domaine Ω_l , l'approche multi-niveaux ne sera qu'une généralisation immédiate de l'algorithme mis au point pour les équations de Burgers. Donc, la méthode multi-résolution auto-adaptative sera couplée avec une méthode de décomposition de domaine qui permettra de travailler localement dans chaque sous-domaine. Les formules définissant la profondeur des V-cycles et la longueur totale de chaque série seront déterminées, de façon indépendante, dans chacun des sous-domaines. Une analyse adéquate de ces critères, analogues à ceux établis pour le problème de Burgers (voir les majorations (1.21) et (1.36) dans la partie I), suivra leur mise au point numérique.

La structure relative à la série de V-cycles dans chaque sous-domaine est déterminée simplement lorsqu'on décompose nos équations sur seulement deux niveaux hiérarchiques. Lors d'une généralisation de la décomposition sur plus de deux niveaux, il faudra être beaucoup plus attentifs afin de réaliser une synchronisation des cycles sur le niveau nodal. La synchronisation est particulièrement importante lorsqu'on termine ou on arrête une série de V-cycles dans l'un des sous-domaines (la fin étant due à l'épuisement de l'intervalle de temps global pendant lequel le terme d'interaction non linéaire est figé, alors que l'arrêt est dû à l'un des critères qui ne sont plus respectés). Il sera nécessaire, en effet, de réévaluer le terme non linéaire d'interaction relatif au sous-domaine Ω_l en connaissant la solution $y_d^n = u_d^n$ sur le domaine Ω tout entier, un problème de phase pouvant intervenir autrement. Rappelons que la j -ème composante du vecteur correspondant au terme non linéaire d'interaction relatif au sous-domaine Ω_l est, pour le niveau $d-1$, calculée de la façon suivante :

$$\hat{b}_{int}(y_{d-1}^n, z_d^n, \phi_{d-1,j}^l) = \hat{b}(y_d^n, y_d^n, \phi_{d-1,j}^l) - \hat{b}(y_{d-1}^n, y_{d-1}^n, \phi_{d-1,j}^l), \quad (2.47)$$

pour j qui varie entre 1 et n_{d-1} , j prenant les valeurs correspondant aux nœuds appartenant au domaine Ω_l . Cette mise à jour du terme non linéaire d'interaction se répète de manière récursive pour les niveaux plus grossiers, en remplaçant l'indice d par k dans (2.47).

Enfin, nous voulons souligner encore un autre aspect positif qui se manifeste dans l'étude locale des différents termes (précisément ceux appartenant au premier groupe présenté ci-dessus). On peut observer, en effet, une très bonne corrélation entre les différents termes, celle-ci nous permettant de deviner l'évolution du terme non linéaire d'interaction dans chaque sous-domaine à partir de termes plus simples et disponibles au cours de toute la simulation numérique. Il est à noter que la corrélation utilisée dans le problème de Burgers était peu satisfaisante, malgré les résultats obtenus (voir la remarque 3 dans la partie I).

Voici les prochaines étapes de notre travail. Avant de développer la stratégie multi-résolution auto-adaptative dans chaque sous-domaine, nous voulons voir si les comportements présentés dans l'étude effectuée par zones se maintiennent pour de plus grands

nombres de Reynolds. Le but à long terme est de simuler numériquement des écoulements non seulement stationnaires mais aussi périodiques. Ainsi, pour la cavité entraînée régularisée, l'écoulement obtenu à $Re = 10000$ converge encore vers une solution stationnaire. Toutefois, le temps d'obtention de cette solution est très long, ce qui nous permet d'avoir une longue période transitoire pendant laquelle la dynamique de l'écoulement s'installe. En considérant la taille des matrices utilisées et la nécessité de raffiner ultérieurement notre maillage, nous mettrons en place d'abord une méthode de décomposition de domaine. Celle-ci nous permettra de paralléliser le code numérique, donc d'augmenter la taille mémoire de notre programme. En effet, la décomposition de domaine facilitera la répartition du maillage et par conséquent la répartition de la factorisation complète de notre matrice dans la mémoire de plusieurs processeurs (on songe à utiliser des machines à mémoire distribuée comme l'IBM SP-2 ou plusieurs stations de travail liées en réseau). D'ailleurs, la décomposition de domaine cadre très bien avec notre stratégie multi-niveaux, les raisons étant suffisamment expliquées ci-dessus. Pour toutes ces motivations, le travail actuellement en cours de développement concerne la parallélisation du code numérique implémentant la méthode de Galerkin classique. Cette parallélisation est associée au traitement hiérarchique de la solution, traitement réalisé toujours *a posteriori*.

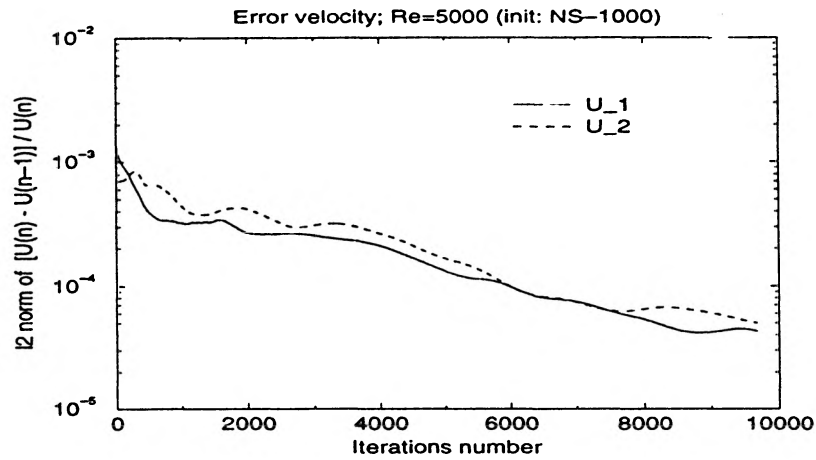


FIG. 2.2 - Erreur relative de la vitesse calculée pour $Re = 5000$.

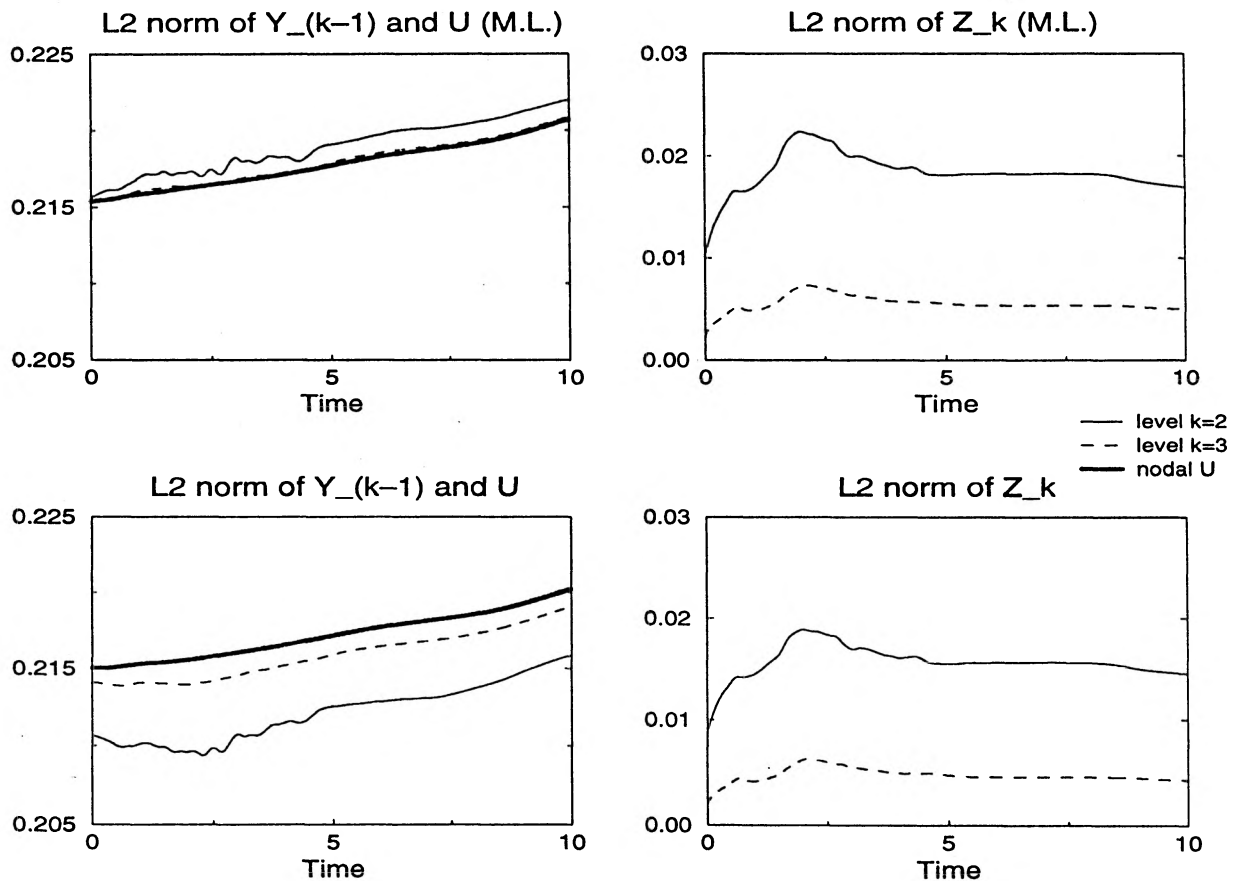


FIG. 2.3 - Norme L^2 (calculée grâce au mass lumping (M.L.) (en haut) ou au produit matrice-vecteur (en bas)) de la vitesse nodale, des grandes structures y_{k-1} et des corrections z_k ($k = 2, 3$).

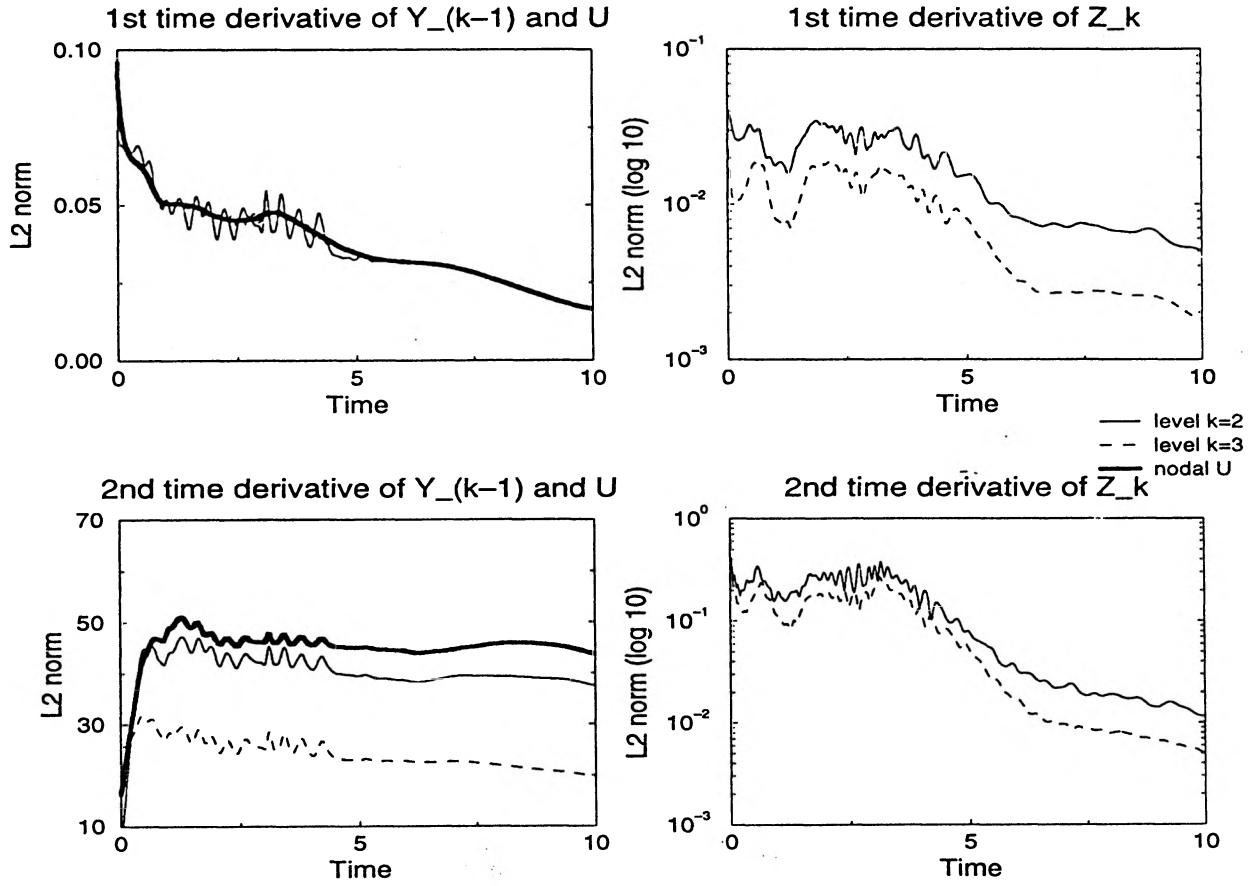


FIG. 2.4 – Norme L^2 des dérivées première et seconde en temps de la vitesse nodale, des grandes structures y_{k-1} et des corrections z_k ($k = 2, 3$).

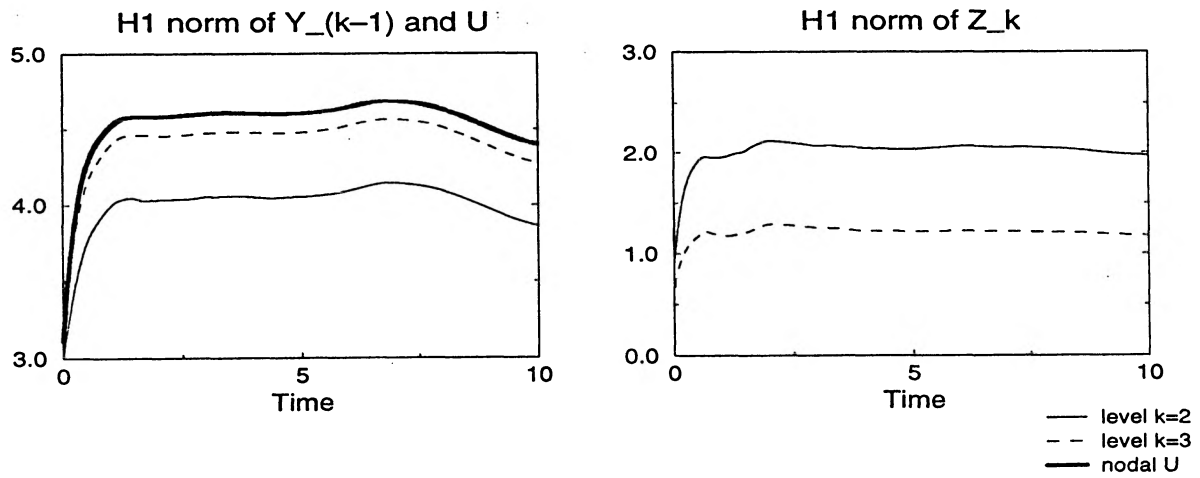


FIG. 2.5 – Norme H^1 de la vitesse nodale, des grandes structures y_{k-1} et des corrections z_k ($k = 2, 3$).

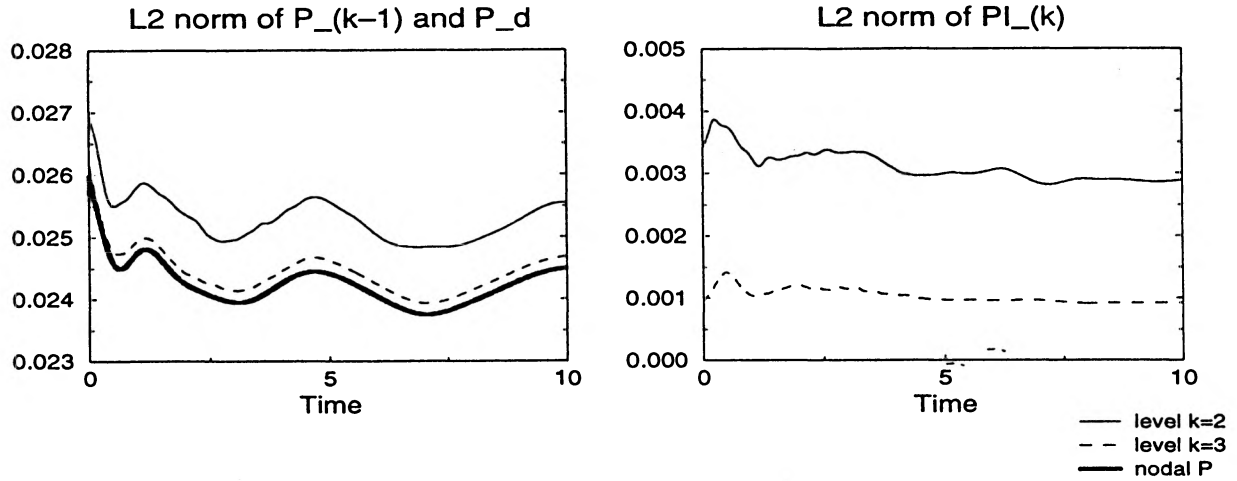


FIG. 2.6 – *Pression physique : norme L^2 de la pression nodale, des grandes structures p_{k-1} et des corrections π_k ($k = 2, 3$).*

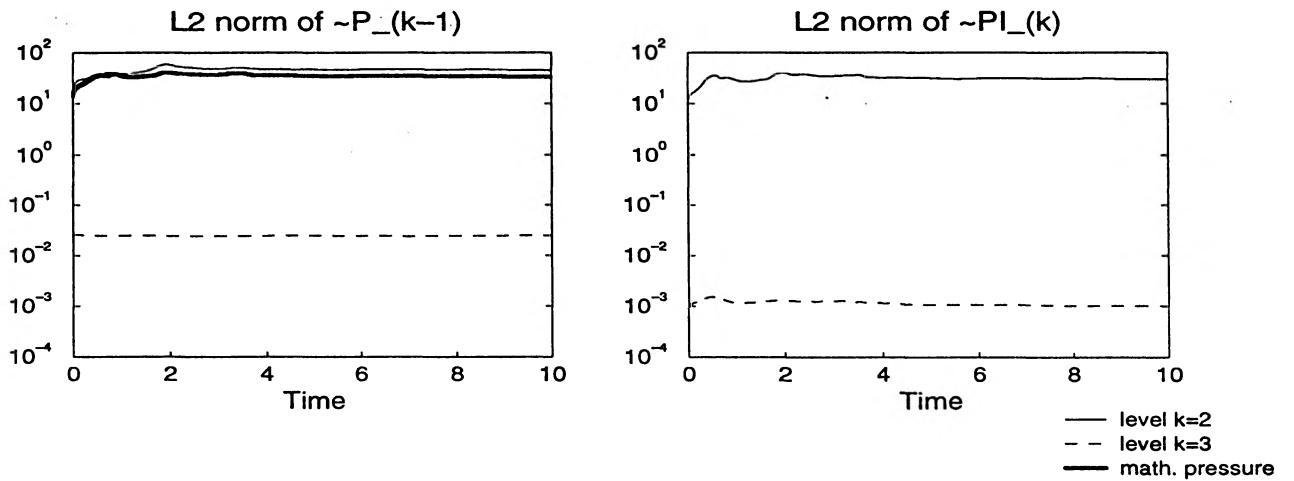


FIG. 2.7 – *Pression mathématique : norme L^2 des grandes structures $\bar{p}_{k-1}$ et des incréments $\bar{\pi}_k$ ($k = 2, 3$).*

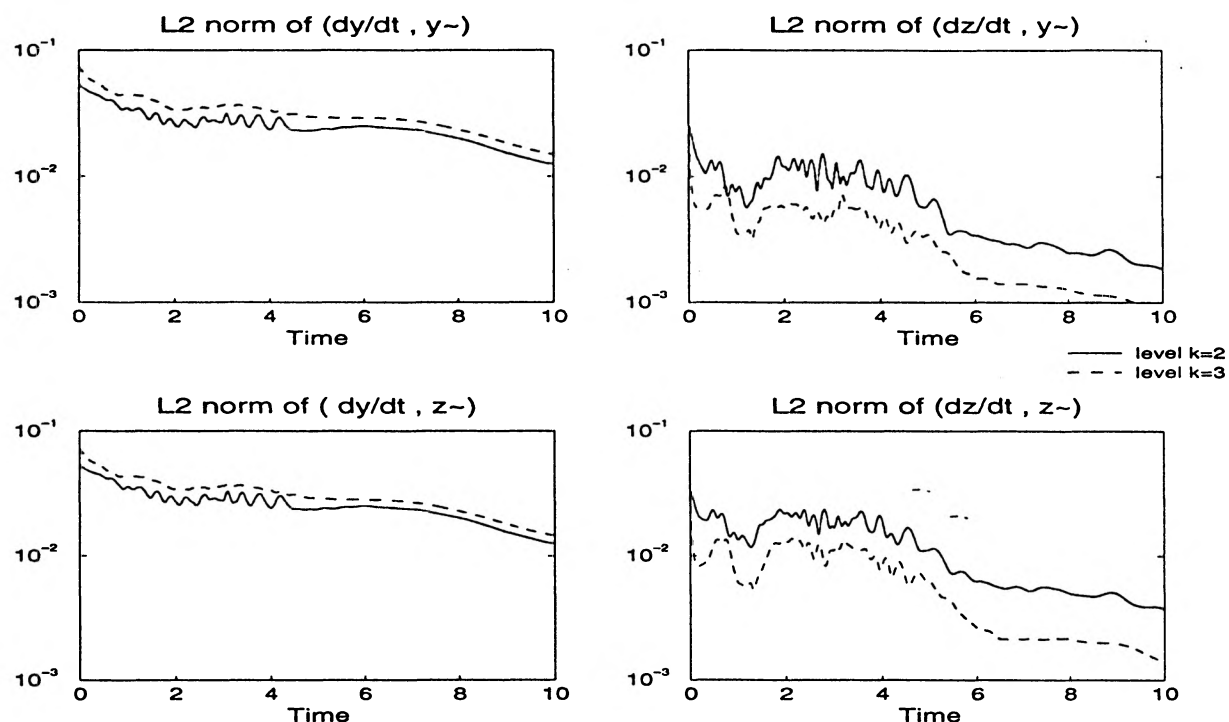


FIG. 2.8 – Evolution de la norme L^2 de la dérivée en temps des y et des z projetée sur la grille grossière et fine.

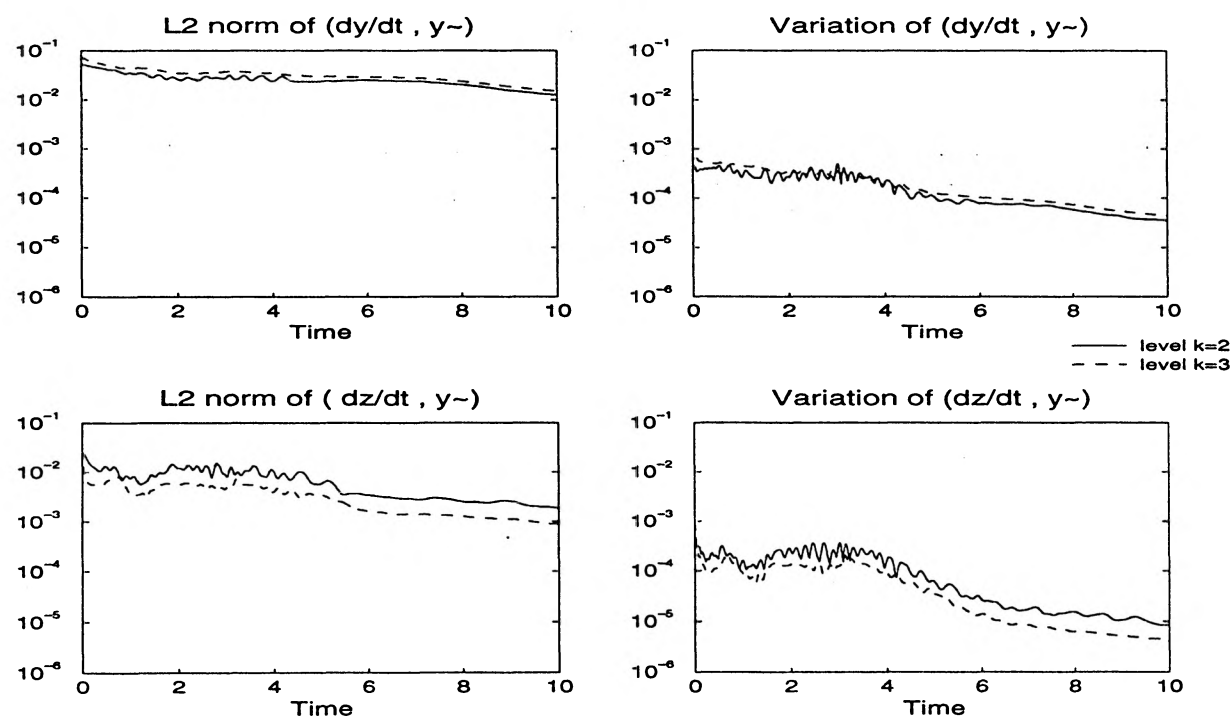


FIG. 2.9 – Evolution de la norme L^2 de $(\frac{\partial y_{k-1}}{\partial t}, \tilde{y}_{k-1})$ et $(\frac{\partial z_k}{\partial t}, \tilde{y}_{k-1})$ et de leurs variations respectives sur un pas de temps.

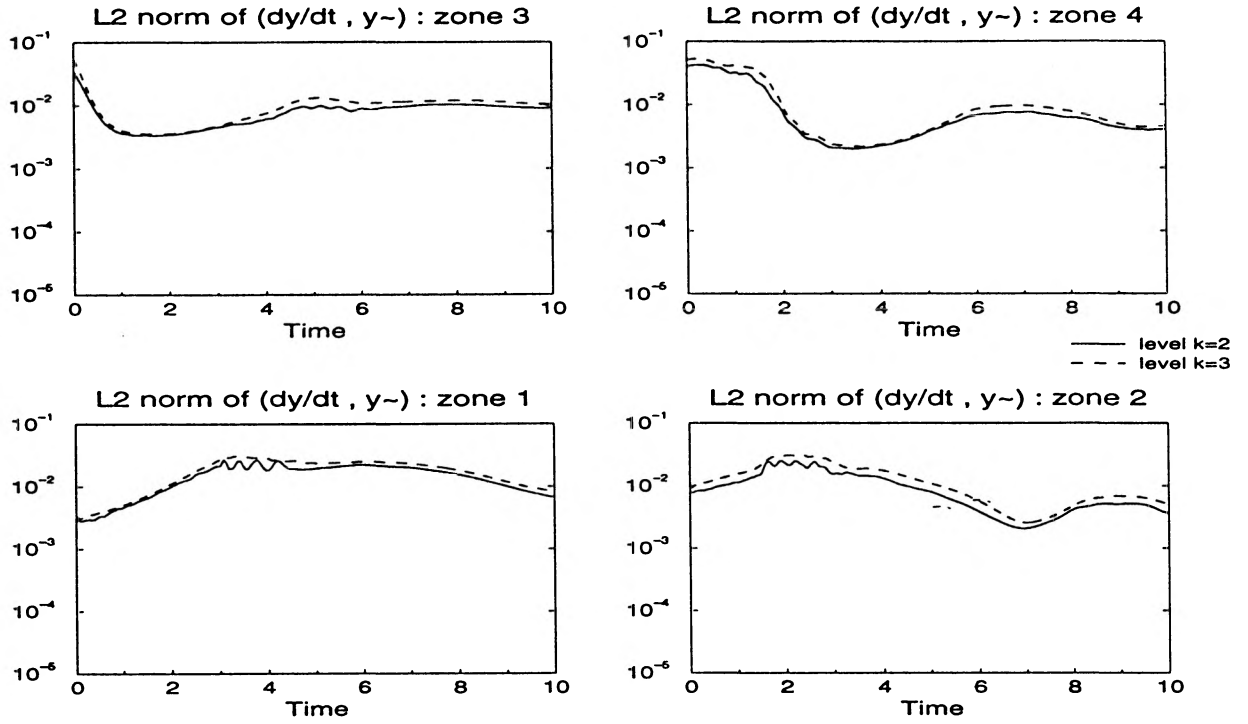


FIG. 2.10 – Evolution par zone de la norme L^2 du terme $(\frac{\partial y_{k-1}}{\partial t}, \tilde{y}_{k-1})$.

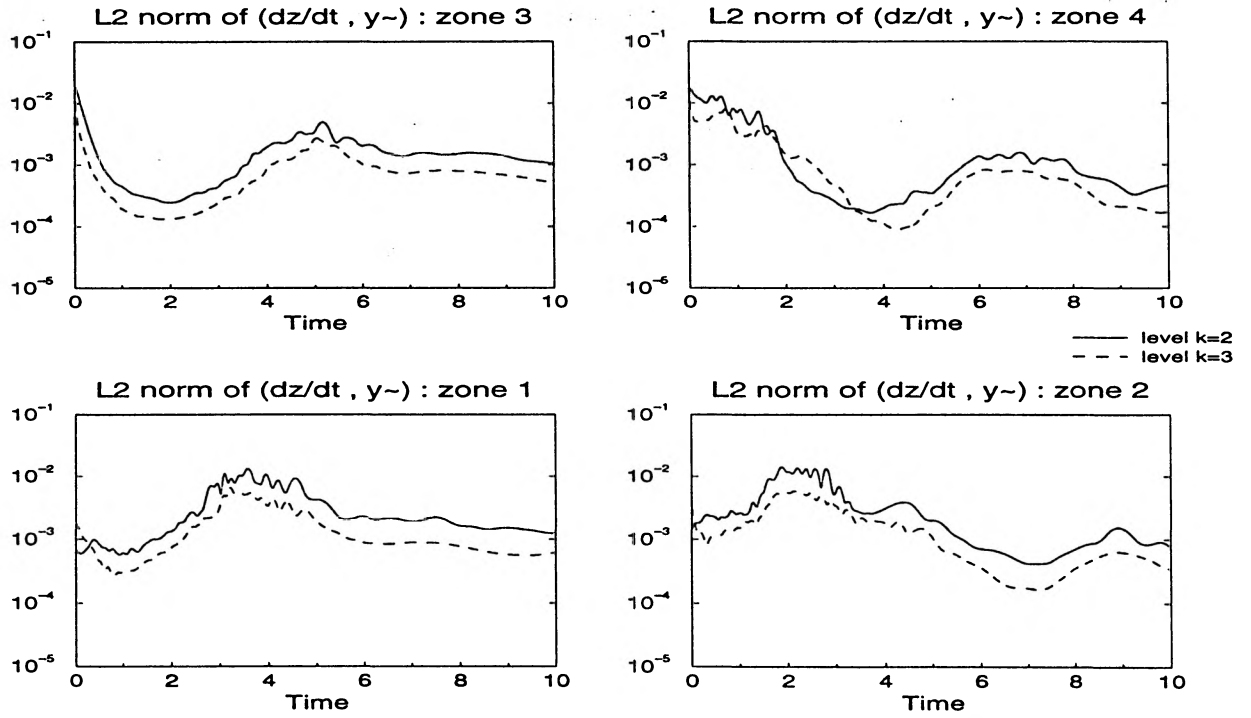


FIG. 2.11 – Evolution par zone de la norme L^2 du terme $(\frac{\partial z_k}{\partial t}, \tilde{y}_{k-1})$.

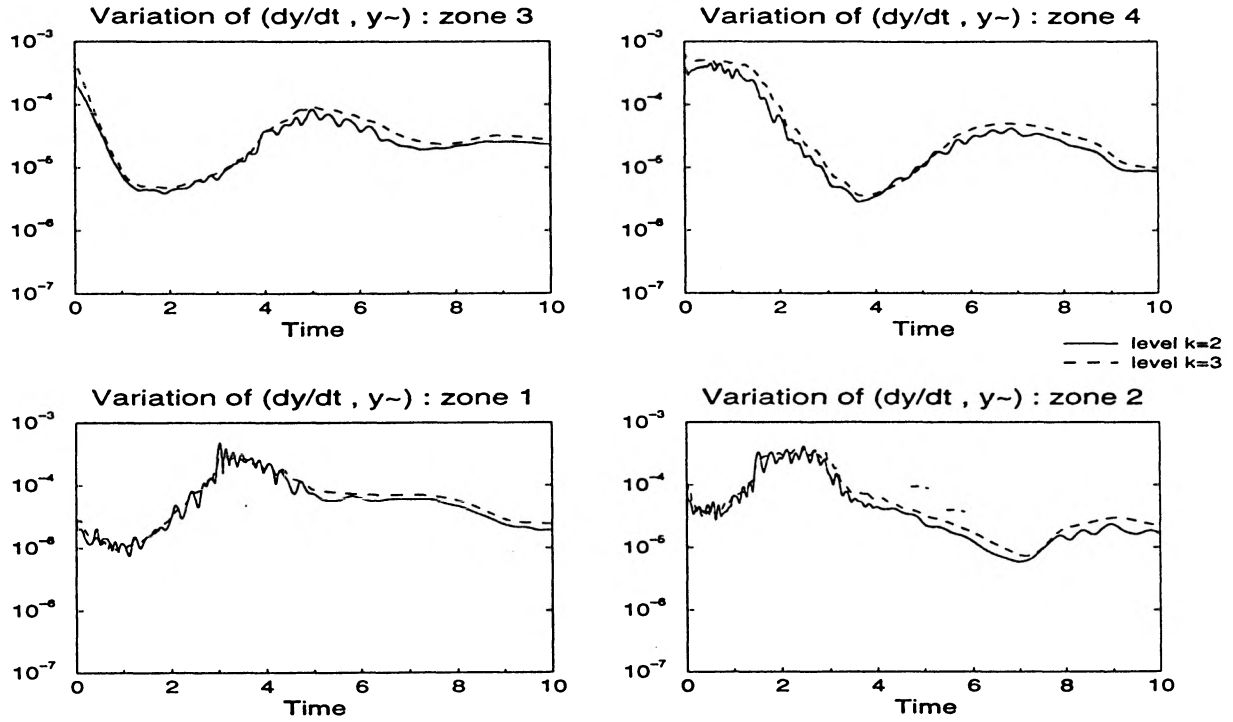


FIG. 2.12 – Evolution par zone de la norme L^2 de la variation sur un pas de temps du terme $(\frac{\partial y_{k-1}}{\partial t}, \tilde{y}_{k-1})$.

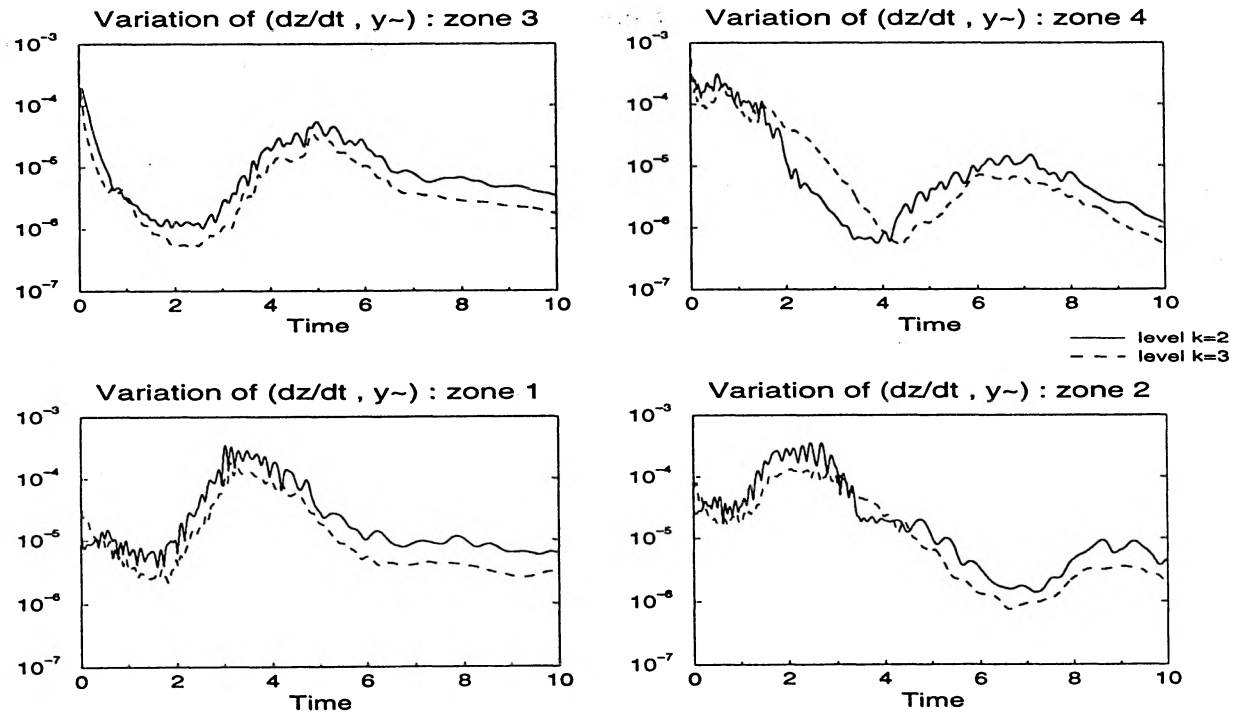


FIG. 2.13 – Evolution par zone de la norme L^2 de la variation sur un pas de temps du terme $(\frac{\partial z_k}{\partial t}, \tilde{y}_{k-1})$.

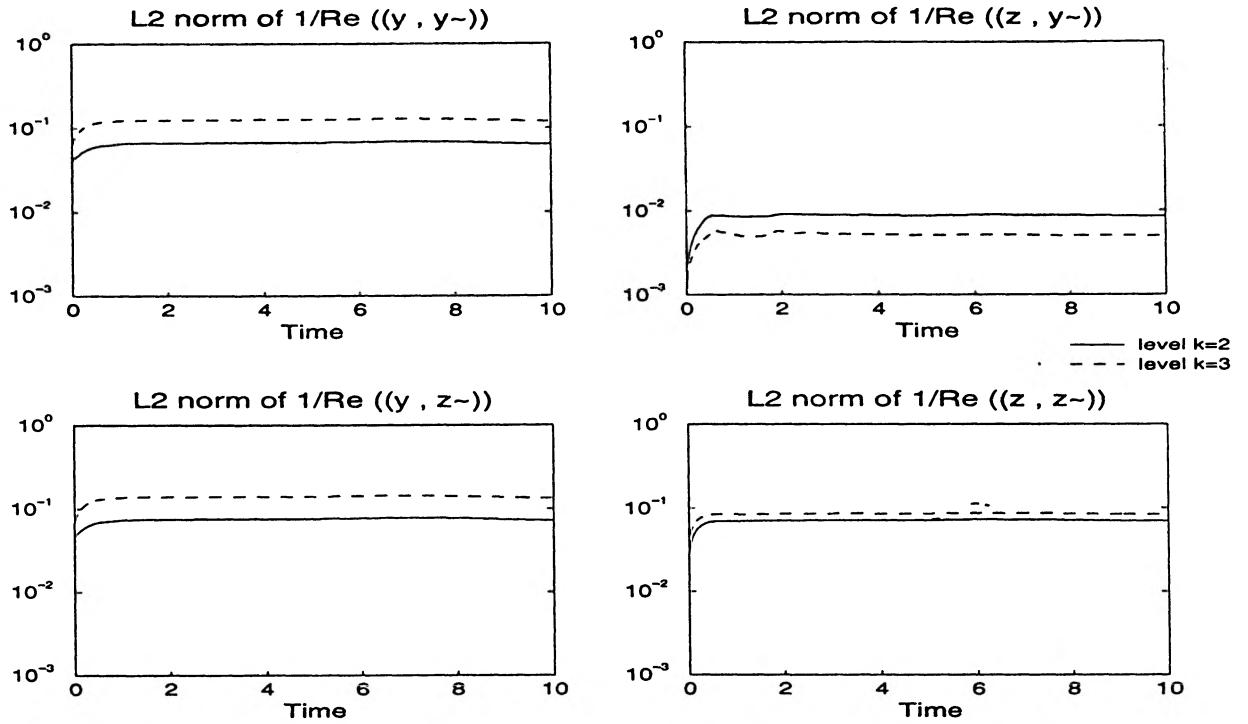


FIG. 2.14 – Evolution de la norme L^2 des termes de diffusions en y et en z projetés sur la grille grossière et fine.

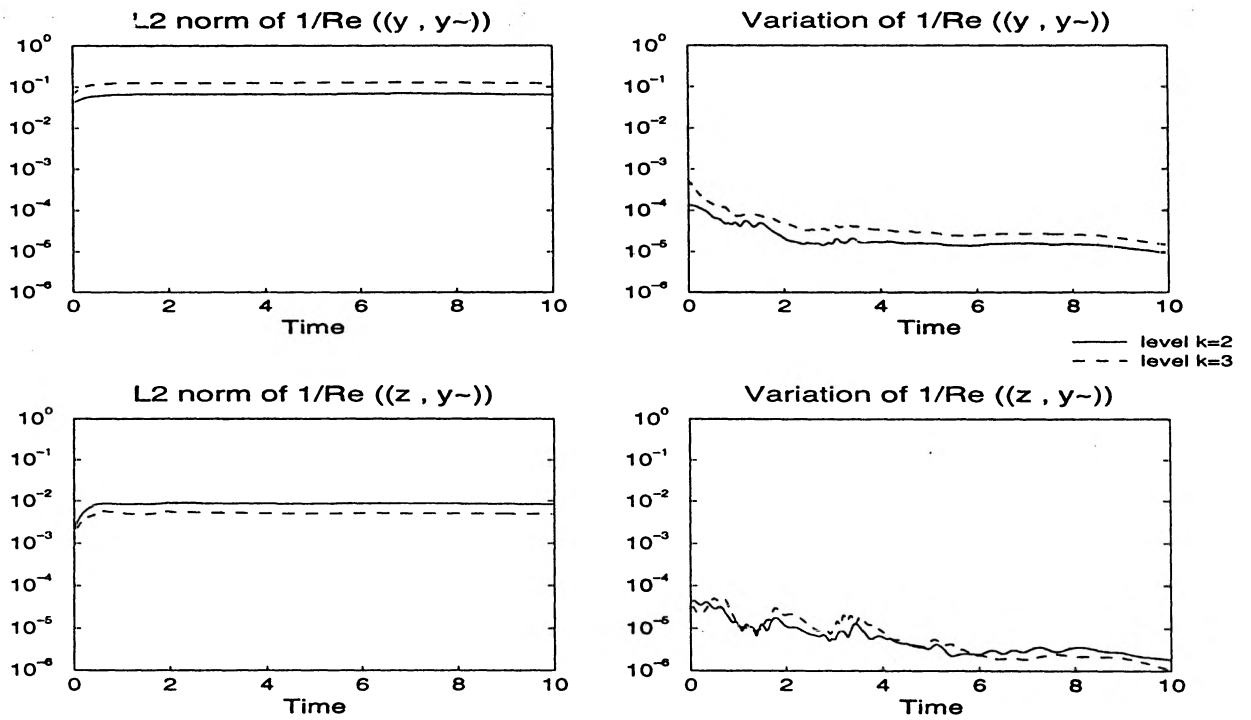


FIG. 2.15 – Evolution de la norme L^2 de $\frac{1}{Re}((y_{k-1}, \tilde{y}_{k-1}))$ et $\frac{1}{Re}((z_k, \tilde{y}_{k-1}))$ et de leurs variations respectives sur un pas de temps.

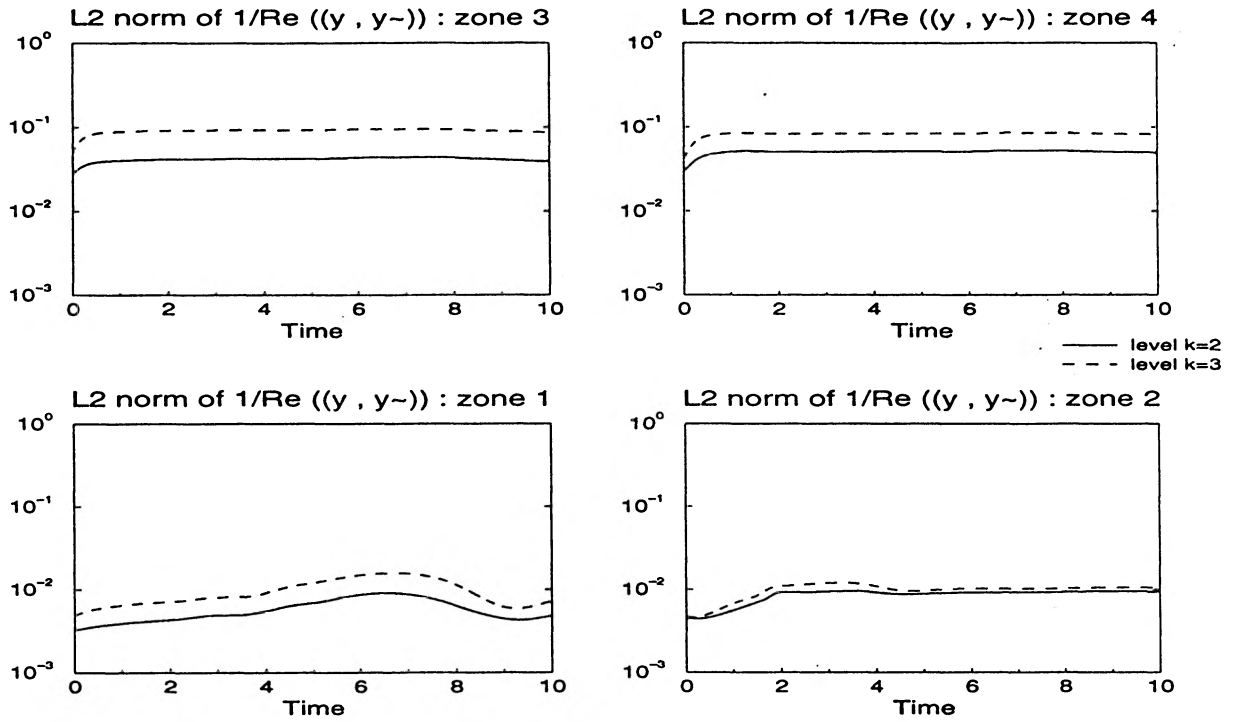


FIG. 2.16 – Evolution par zone de la norme L^2 du terme $\frac{1}{\text{Re}}((y_{k-1}, \tilde{y}_{k-1}))$.

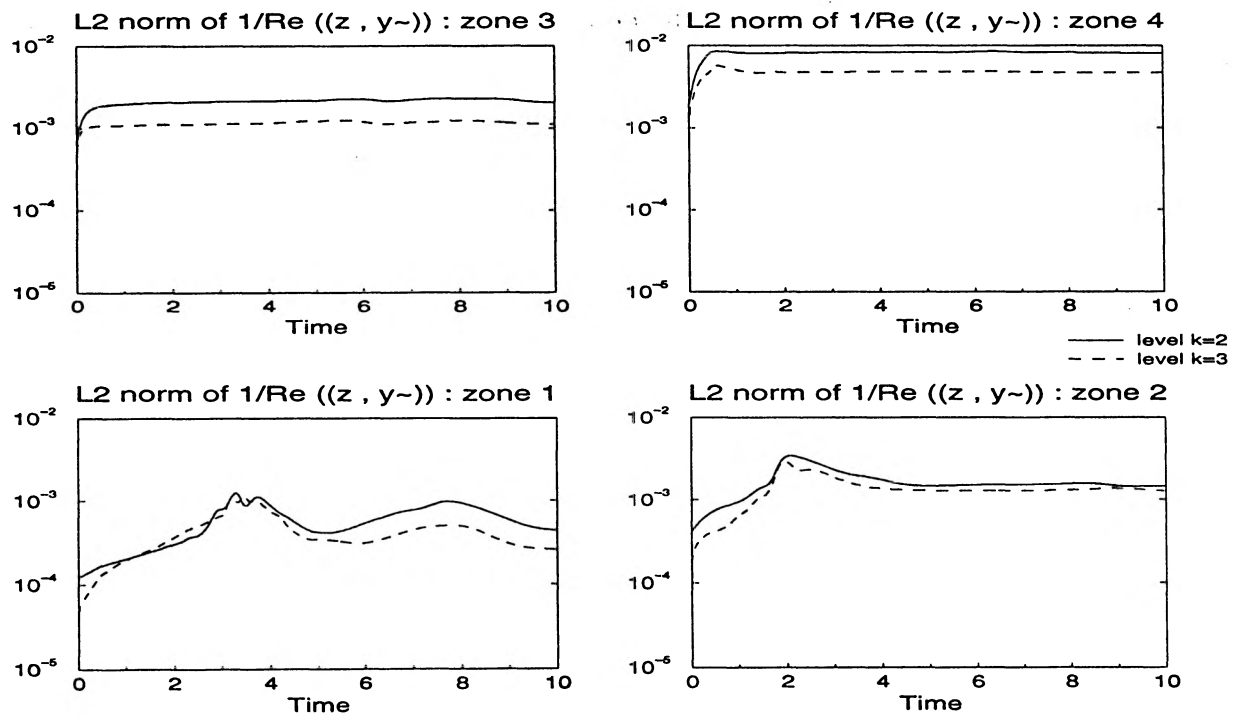


FIG. 2.17 – Evolution par zone de la norme L^2 du terme $\frac{1}{\text{Re}}((z_k, \tilde{y}_{k-1}))$.

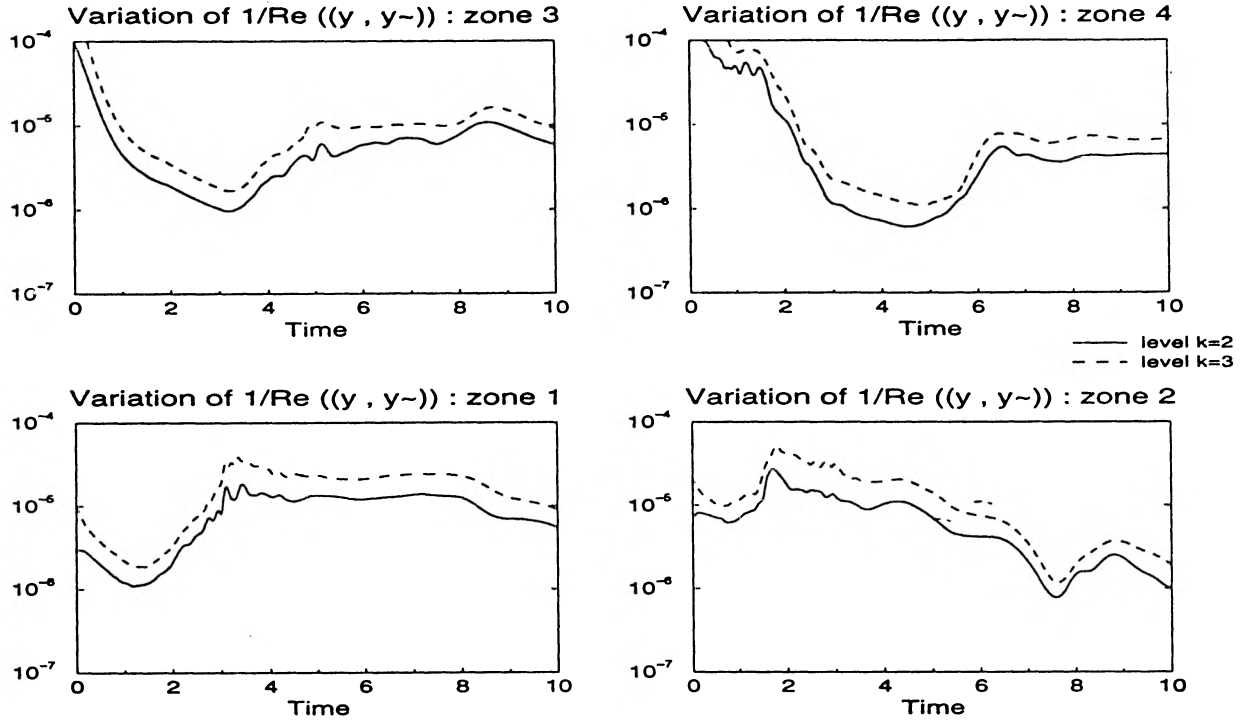


FIG. 2.18 – Evolution par zone de la norme L^2 de la variation sur un pas de temps du terme $\frac{1}{Re}((y_{k-1}, \tilde{y}_{k-1}))$.

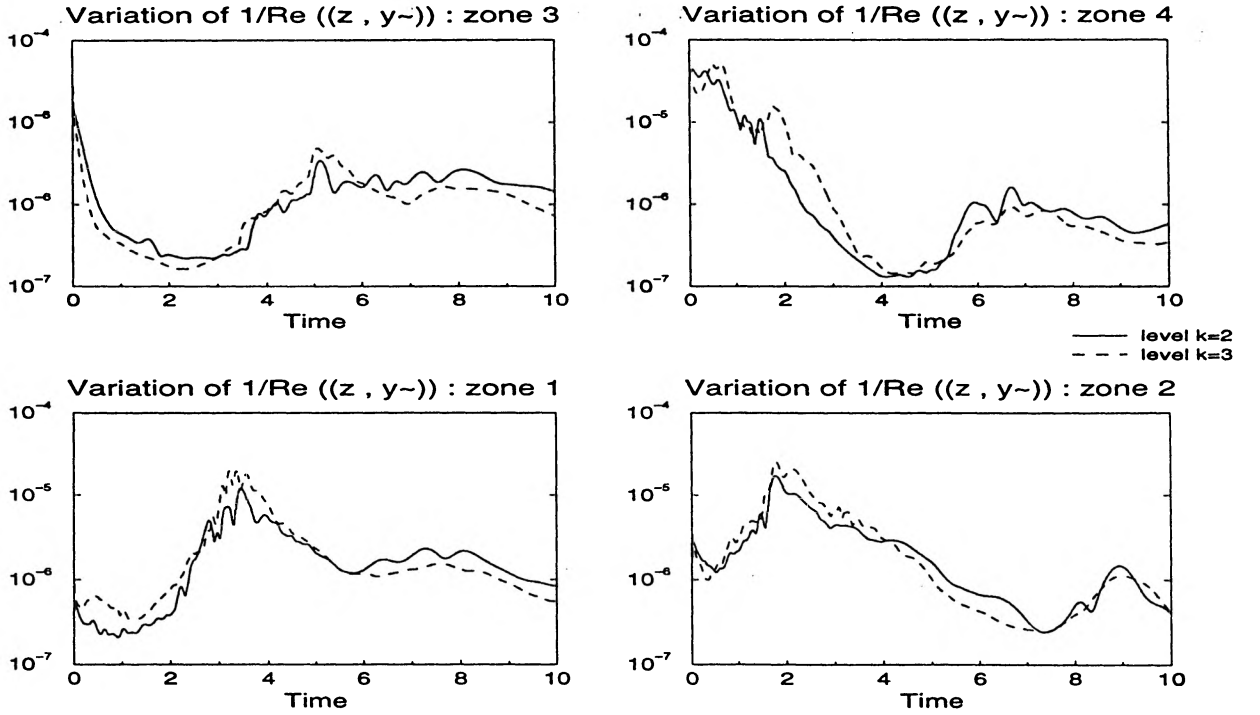


FIG. 2.19 – Evolution par zone de la norme L^2 de la variation sur un pas de temps du terme $\frac{1}{Re}((z_k, \tilde{y}_{k-1}))$.

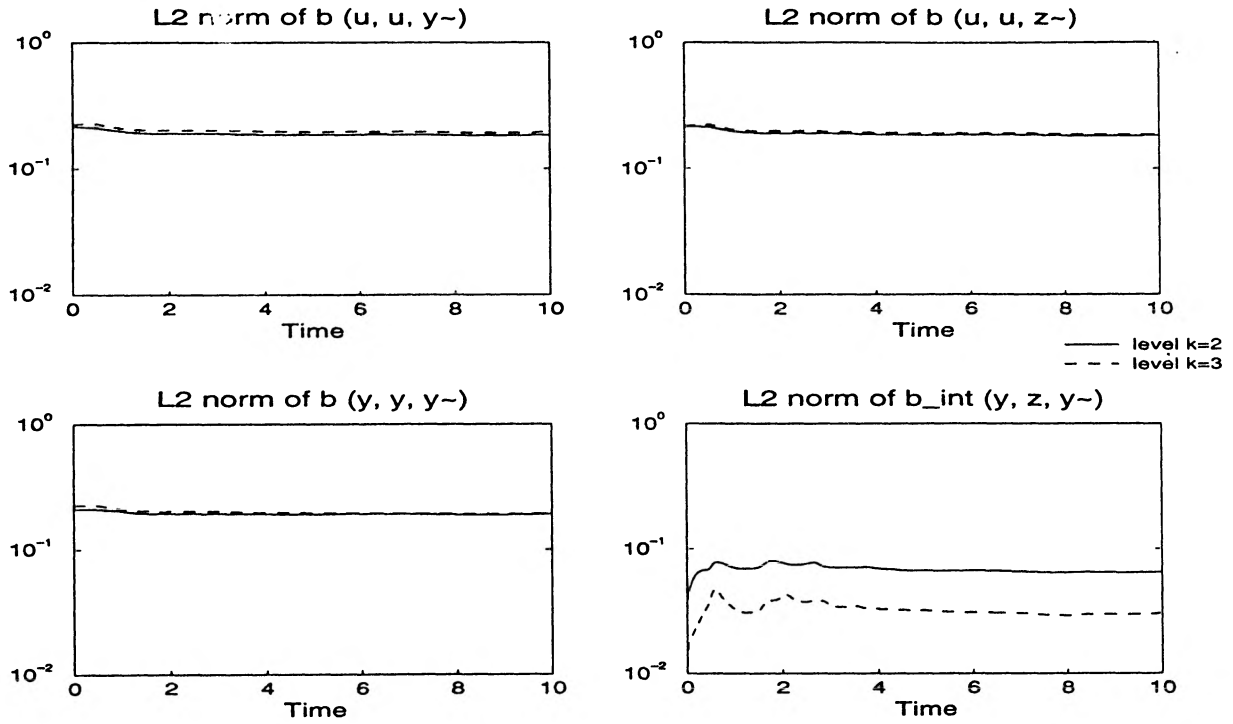


FIG. 2.20 – Evolution de la norme L^2 du terme non linéaire projeté sur la grille grossière et fine (en haut) et de la décomposition du terme non linéaire en $\hat{b}_g$ et $\hat{b}_{int,g}$ (en bas).

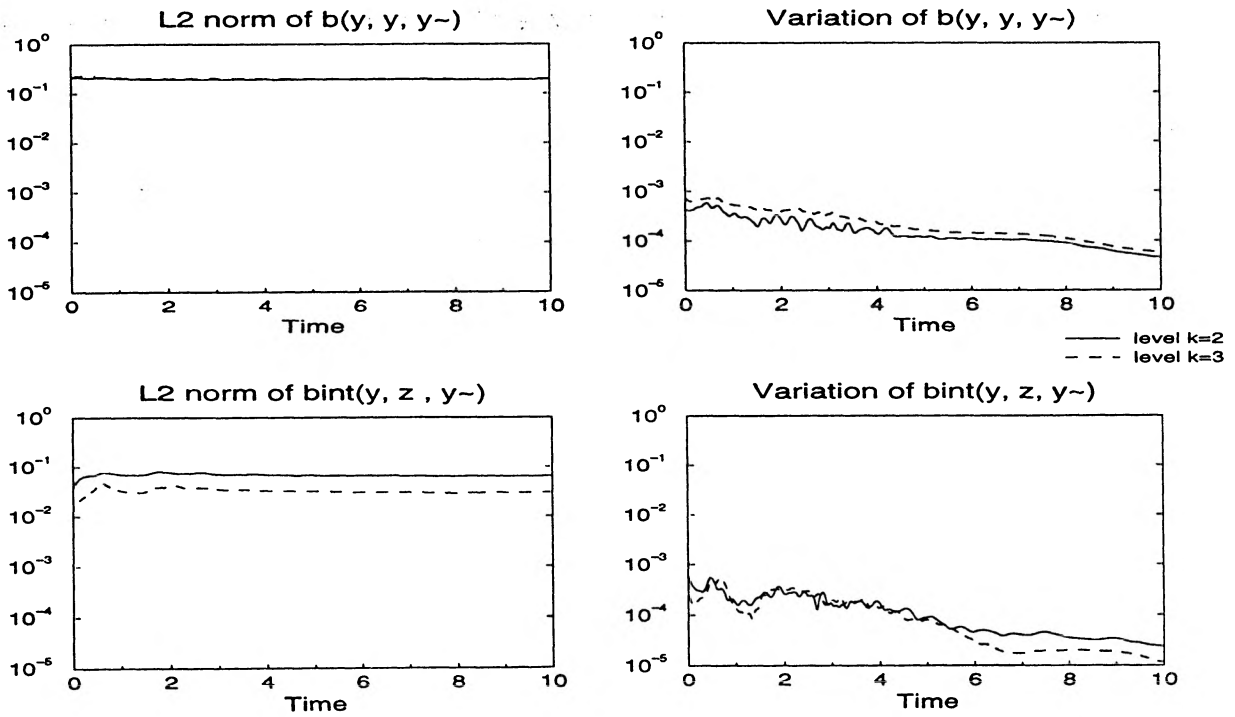


FIG. 2.21 – Evolution de la norme L^2 de $\hat{b}(y_{k-1}, y_{k-1}, \tilde{y}_{k-1})$ et $\hat{b}_{int}(y_{k-1}, z_k, \tilde{y}_{k-1})$ et de leurs variations respectives sur un pas de temps.

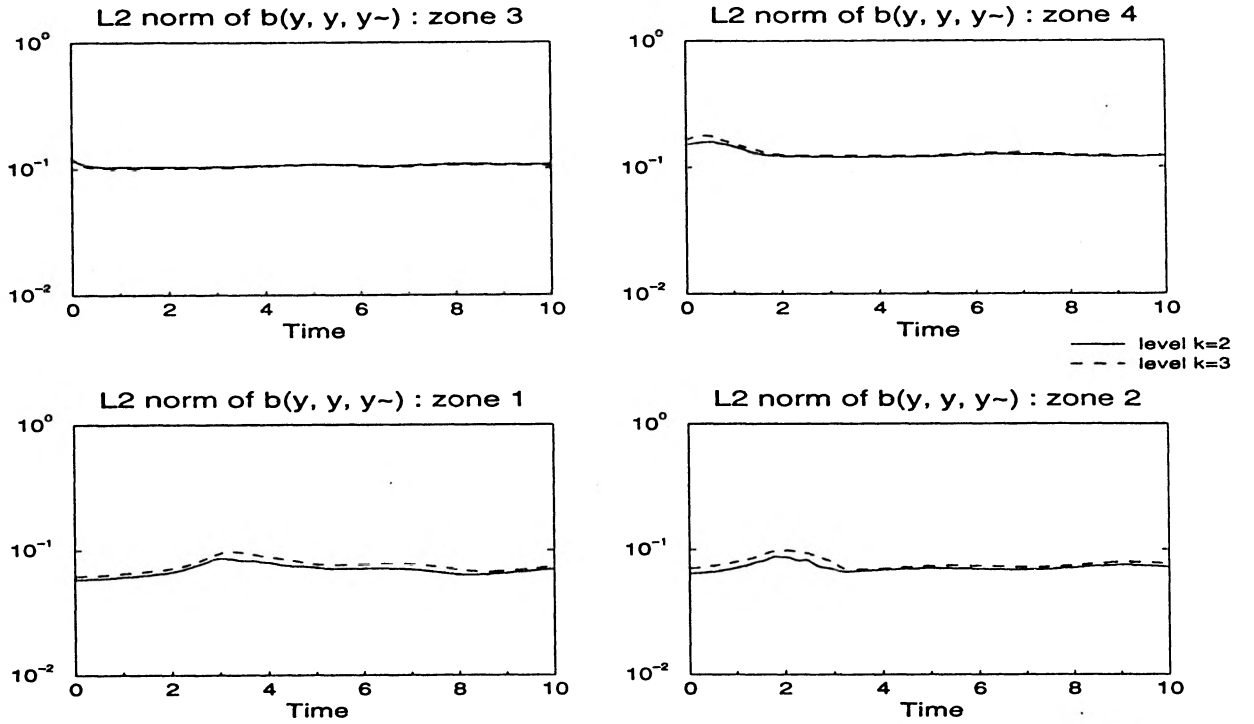


FIG. 2.22 – Evolution par zone de la norme L^2 du terme $\hat{b}(y_{k-1}, y_{k-1}, \tilde{y}_{k-1})$.

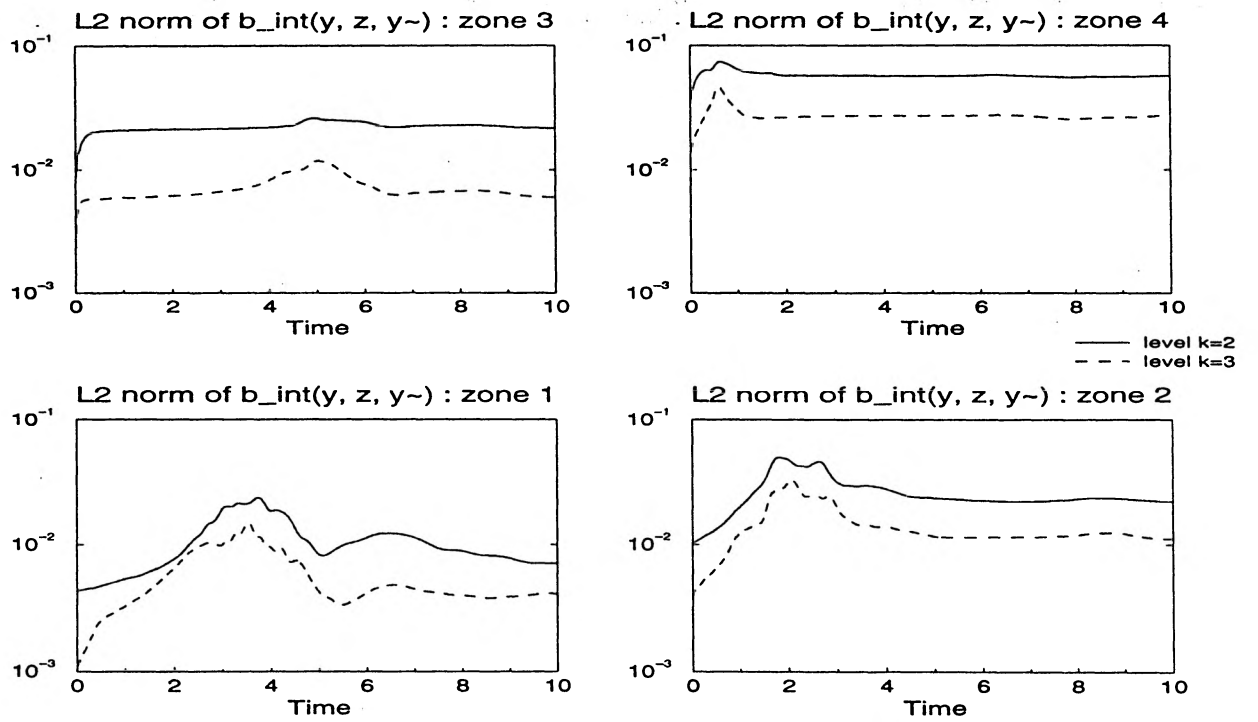


FIG. 2.23 – Evolution par zone de la norme L^2 du terme $\hat{b}_{\text{int}}(y_{k-1}, z_k, \tilde{y}_{k-1})$.

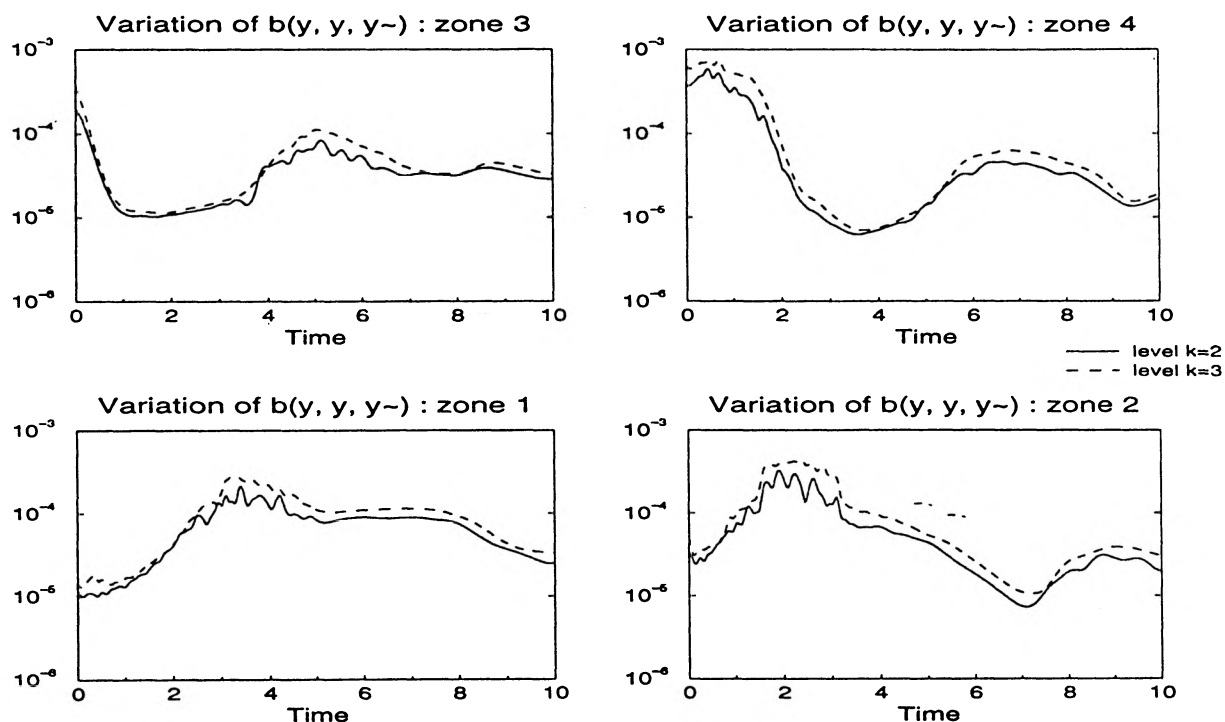


FIG. 2.24 – Evolution par zone de la norme L^2 de la variation sur un pas de temps du terme $\hat{b}(y_{k-1}, y_{k-1}, \tilde{y}_{k-1})$.

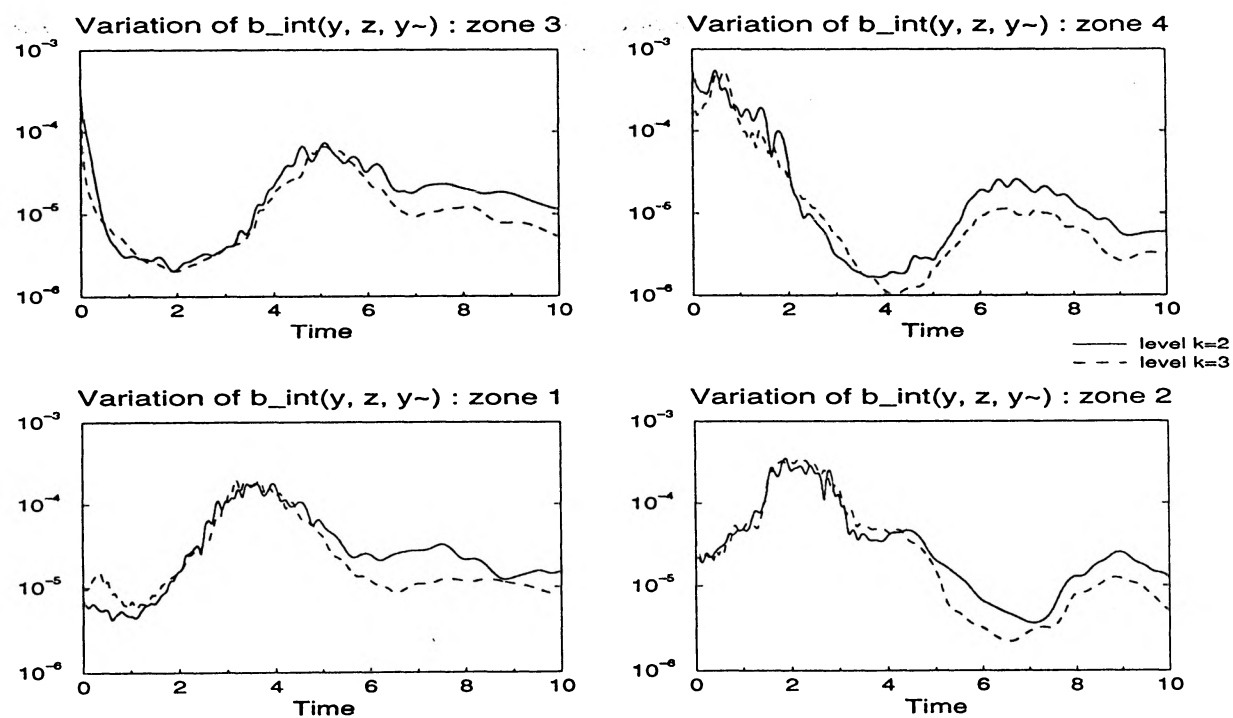


FIG. 2.25 – Evolution par zone de la norme L^2 de la variation sur un pas de temps du terme $\hat{b}_{\text{int}}(y_{k-1}, z_k, \tilde{y}_{k-1})$.

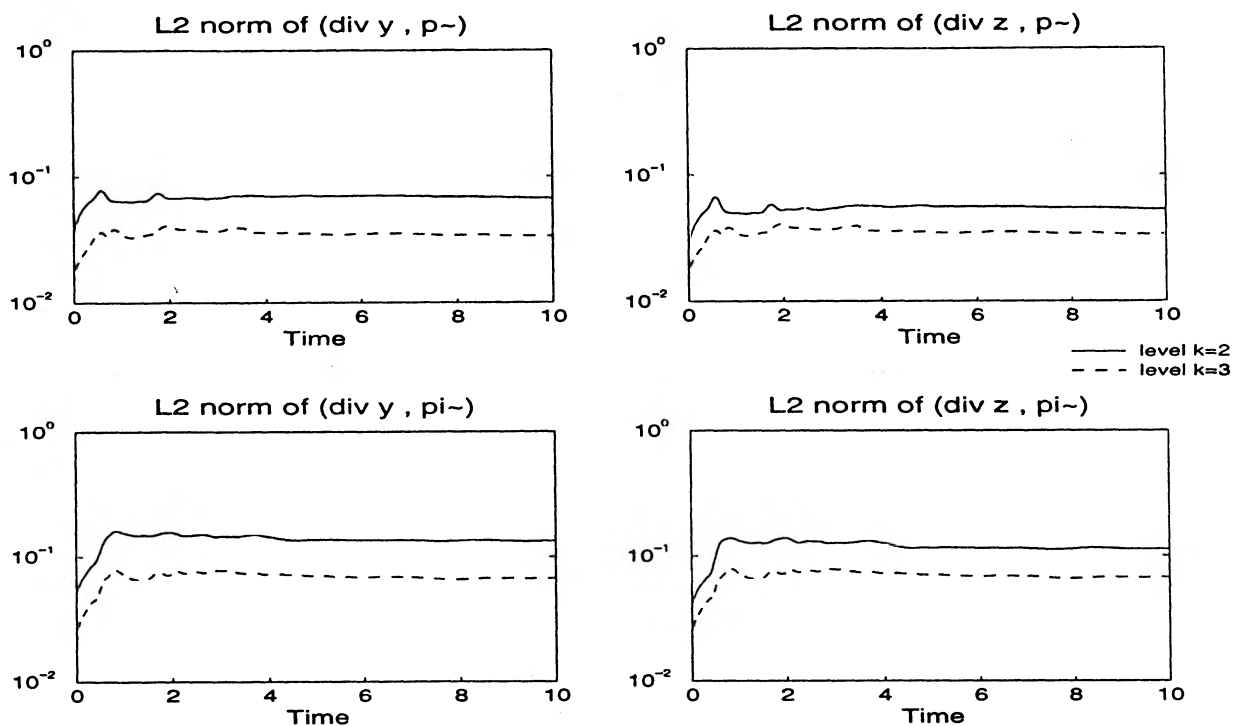


FIG. 2.26 – Evolution de la norme L^2 des termes de divergence en y et en z projetés sur la grille grossière et fine.

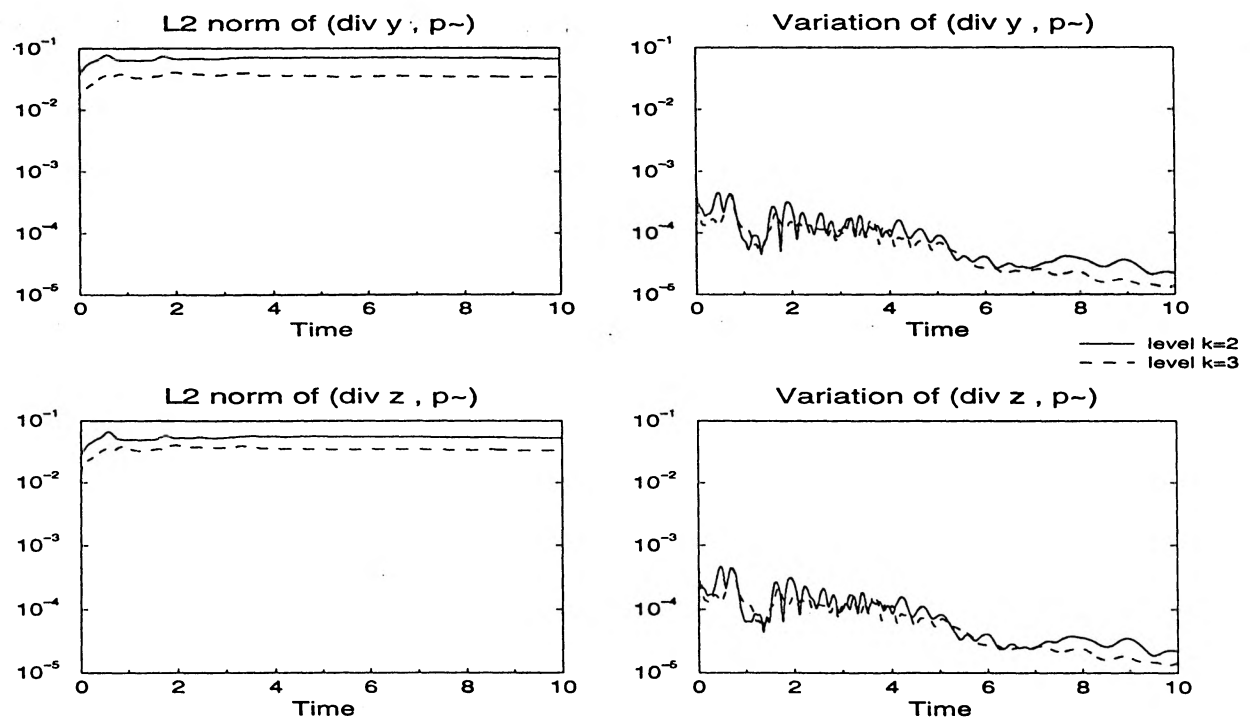
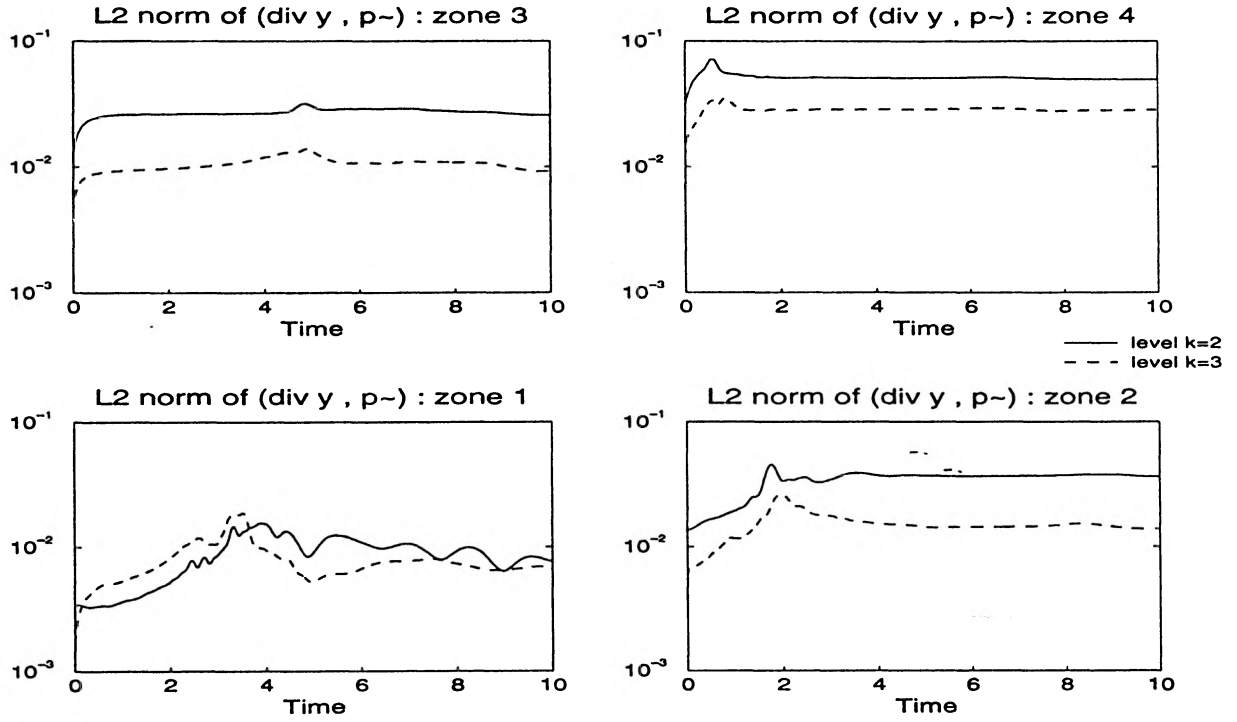
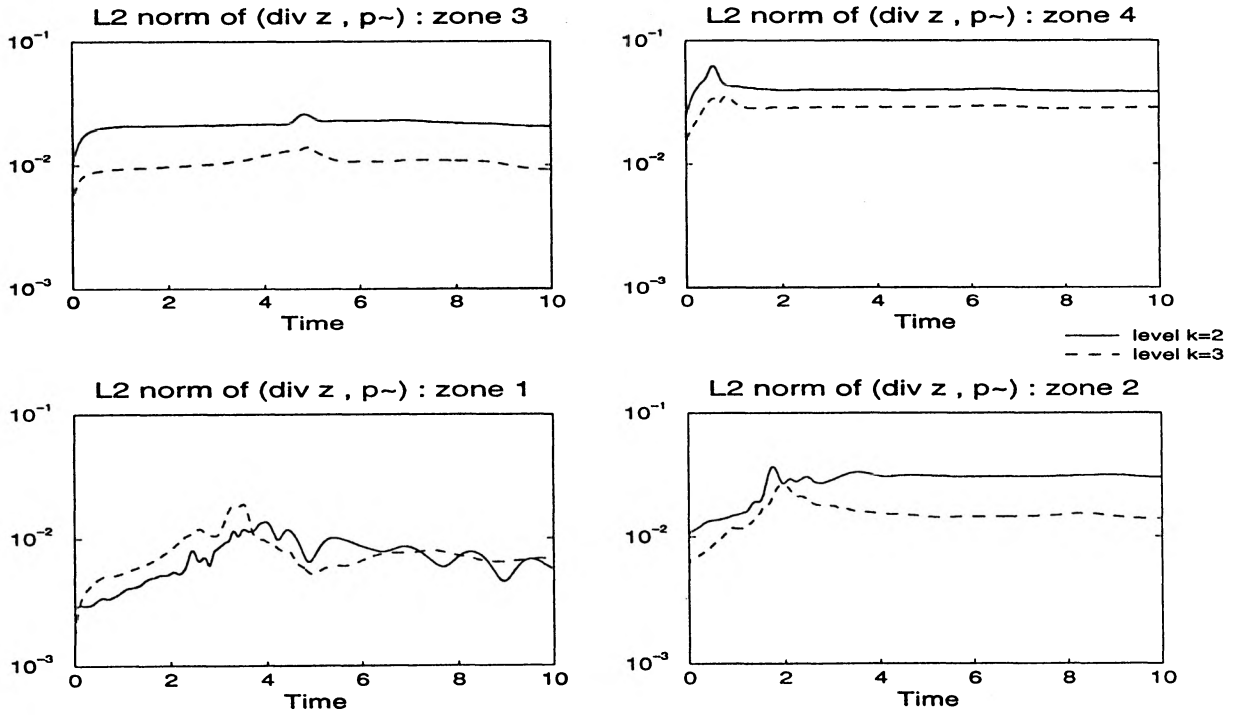


FIG. 2.27 – Evolution de la norme L^2 de $(div y_{k-1}, \tilde{p}_{k-1})$ et $(div z_k, \tilde{p}_{k-1})$ et de leurs variations respectives sur un pas de temps.

FIG. 2.28 – Evolution par zone de la norme L^2 du terme $(\text{div } y_{k-1}, \tilde{p}_{k-1})$.FIG. 2.29 – Evolution par zone de la norme L^2 du terme $(\text{div } z_k, \tilde{p}_{k-1})$.

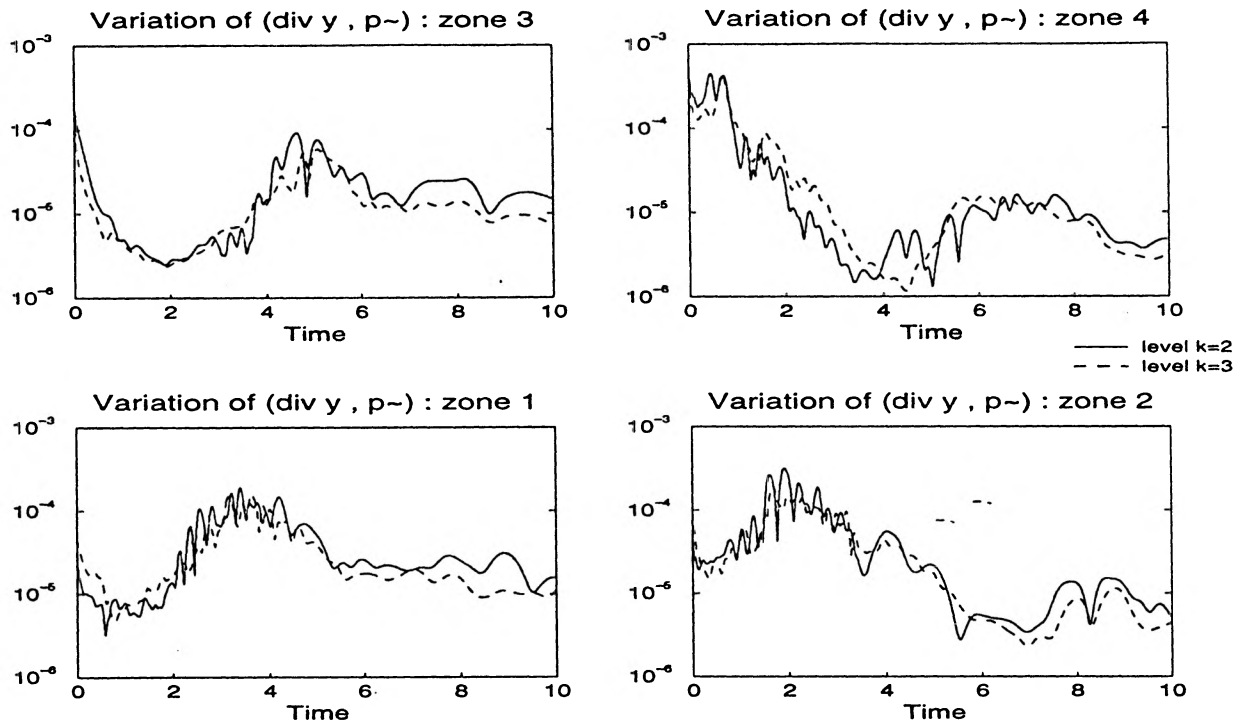


FIG. 2.30 – Evolution par zone de la norme L^2 de la variation sur un pas de temps du terme $(\text{div } y_{k-1}, \tilde{p}_{k-1})$.

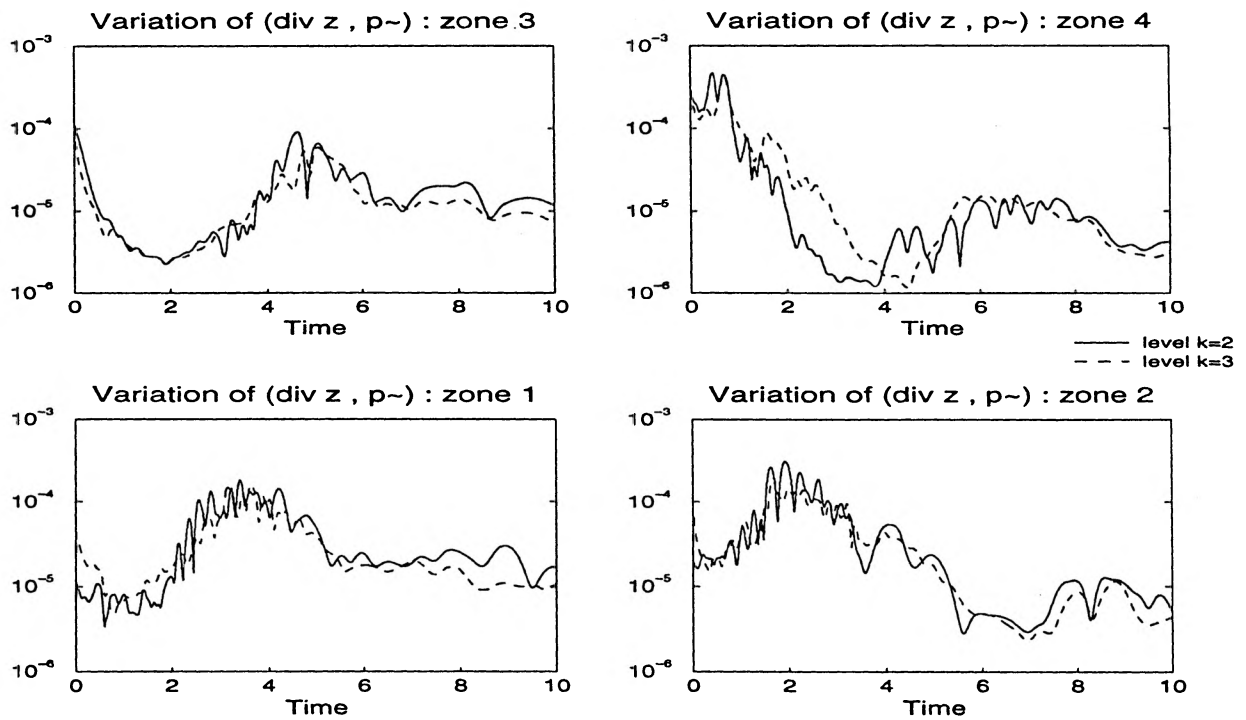


FIG. 2.31 – Evolution par zone de la norme L^2 de la variation sur un pas de temps du terme $(\text{div } z_k, \tilde{p}_{k-1})$.

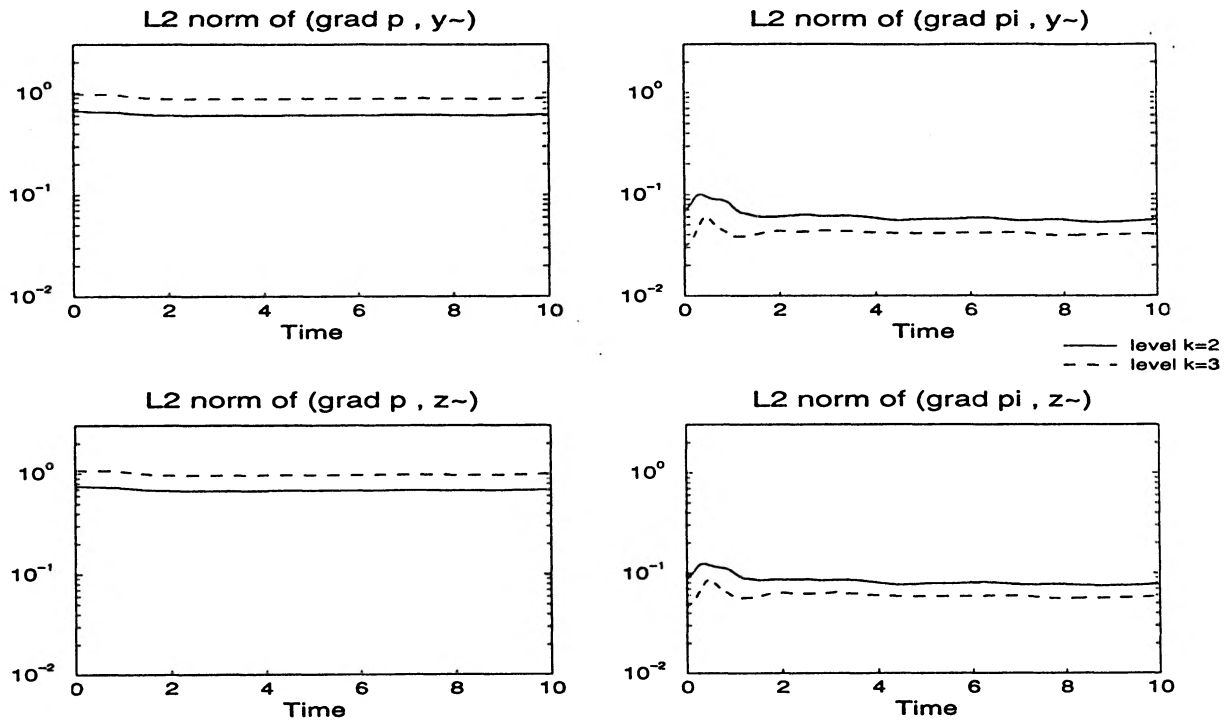


FIG. 2.32 – Evolution de la norme L^2 du gradient de la pression physique: termes en ∇p et en $\nabla \pi$ projetés sur la grille grossière et fine.

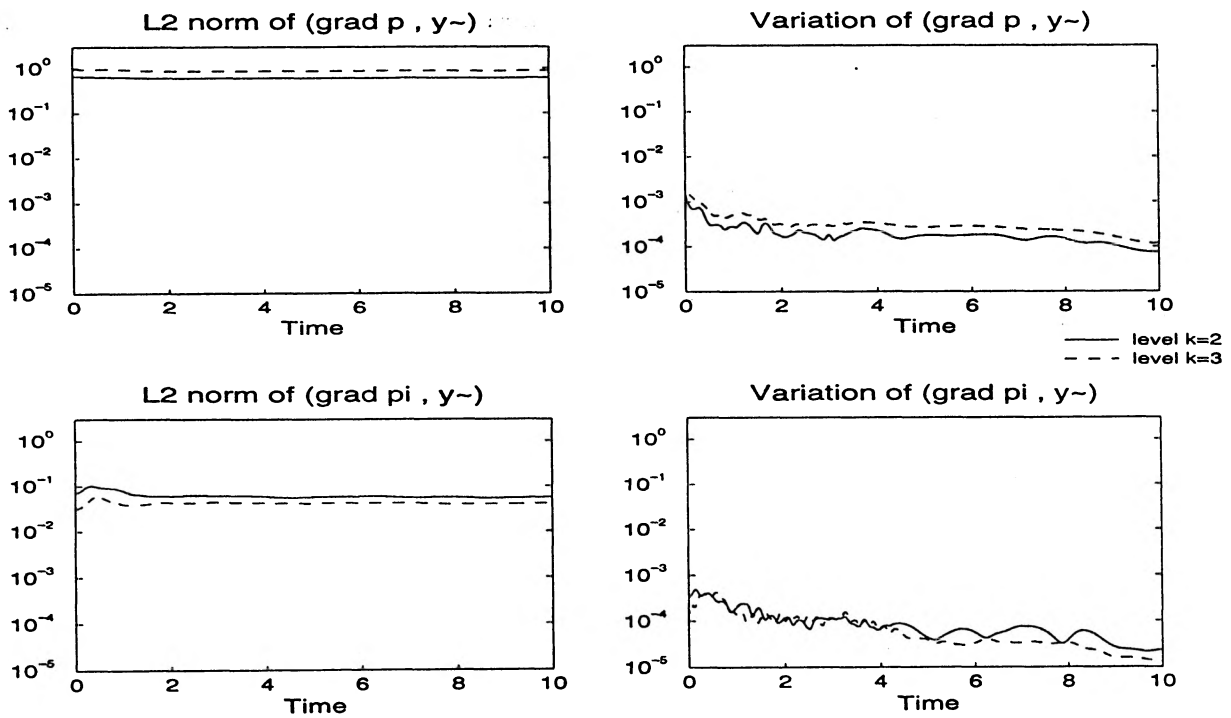
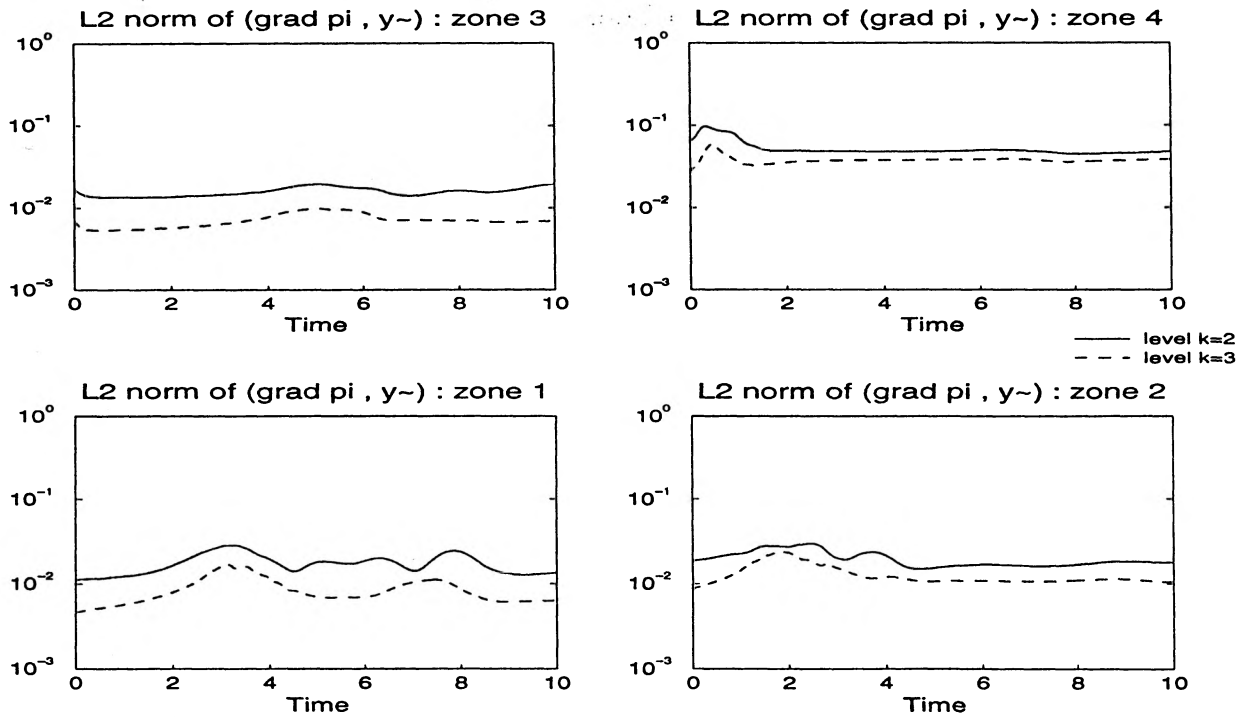
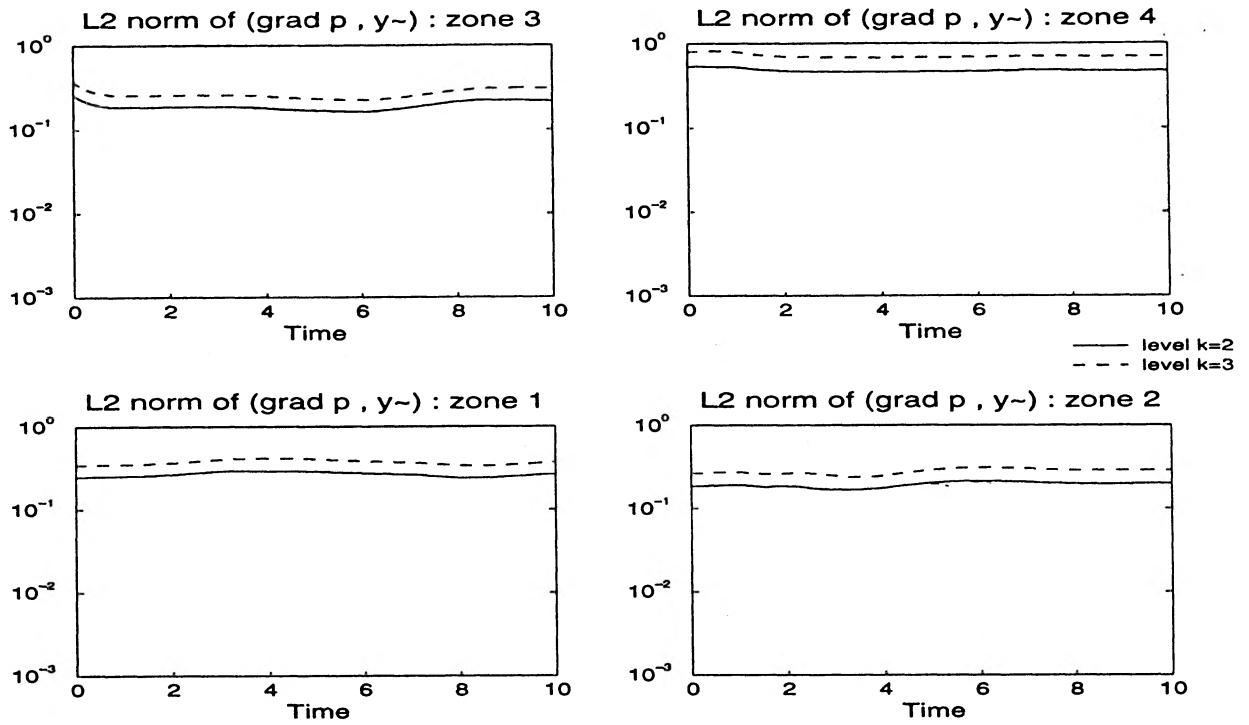


FIG. 2.33 – Evolution de la norme L^2 de $(\nabla p_{k-1}, \tilde{y}_{k-1})$ et $(\nabla \pi_k, \tilde{y}_{k-1})$ et de leurs variations respectives sur un pas de temps.



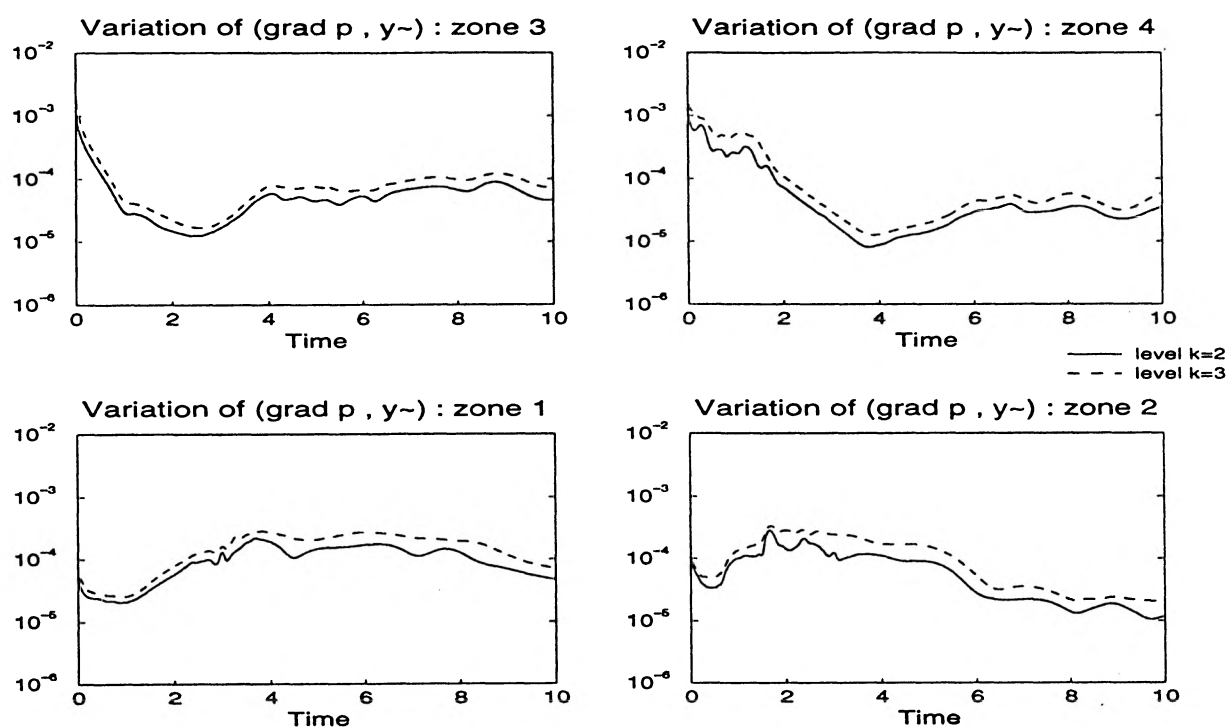


FIG. 2.36 – Evolution par zone de la norme L^2 de la variation sur un pas de temps du terme $(\nabla p_{k-1}, \tilde{y}_{k-1})$.

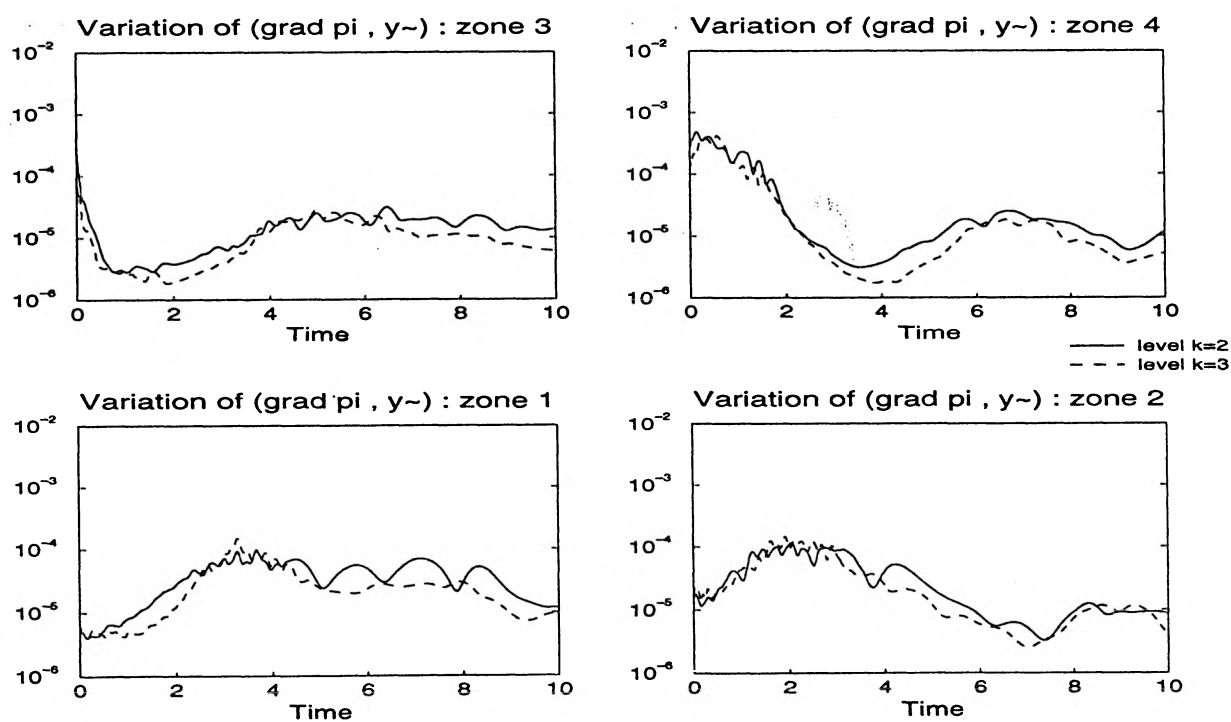


FIG. 2.37 – Evolution par zone de la norme L^2 de la variation sur un pas de temps du terme $(\nabla \pi_k, \tilde{y}_{k-1})$.

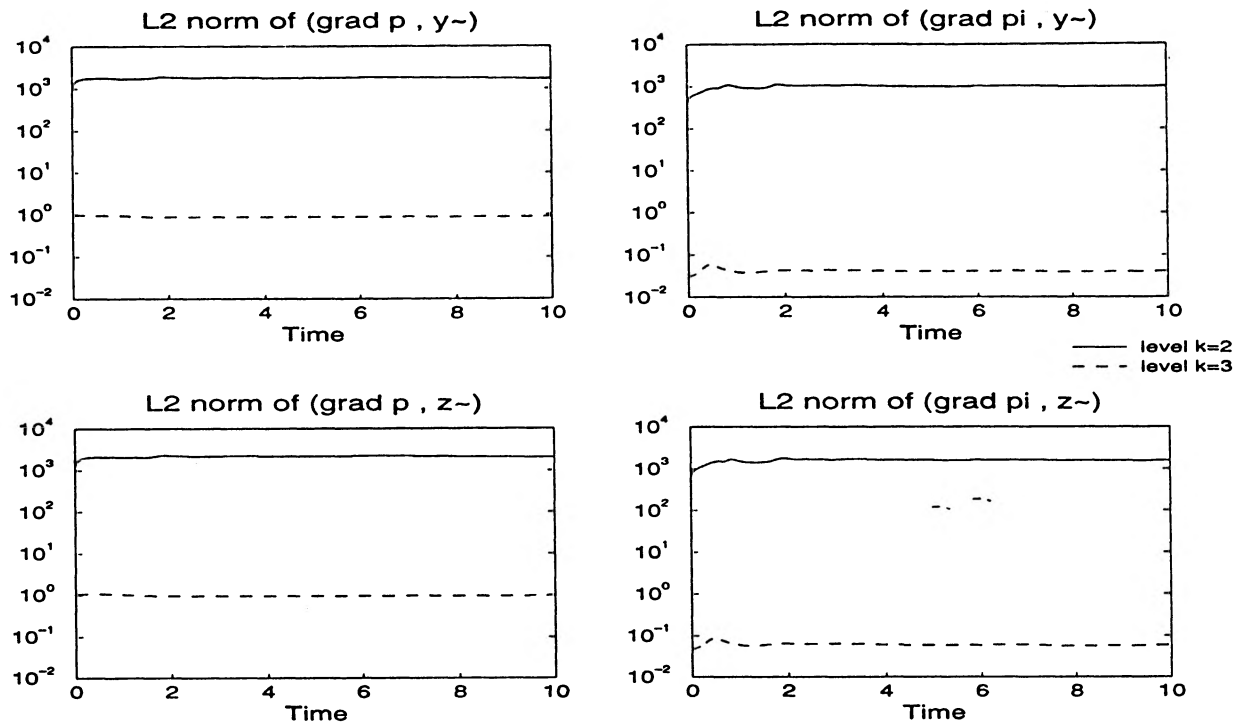


FIG. 2.38 – Evolution de la norme L^2 du gradient de la pression mathématique : termes en $\nabla \bar{p}$ et en $\nabla \bar{\pi}$ projetés sur la grille grossière et fine.

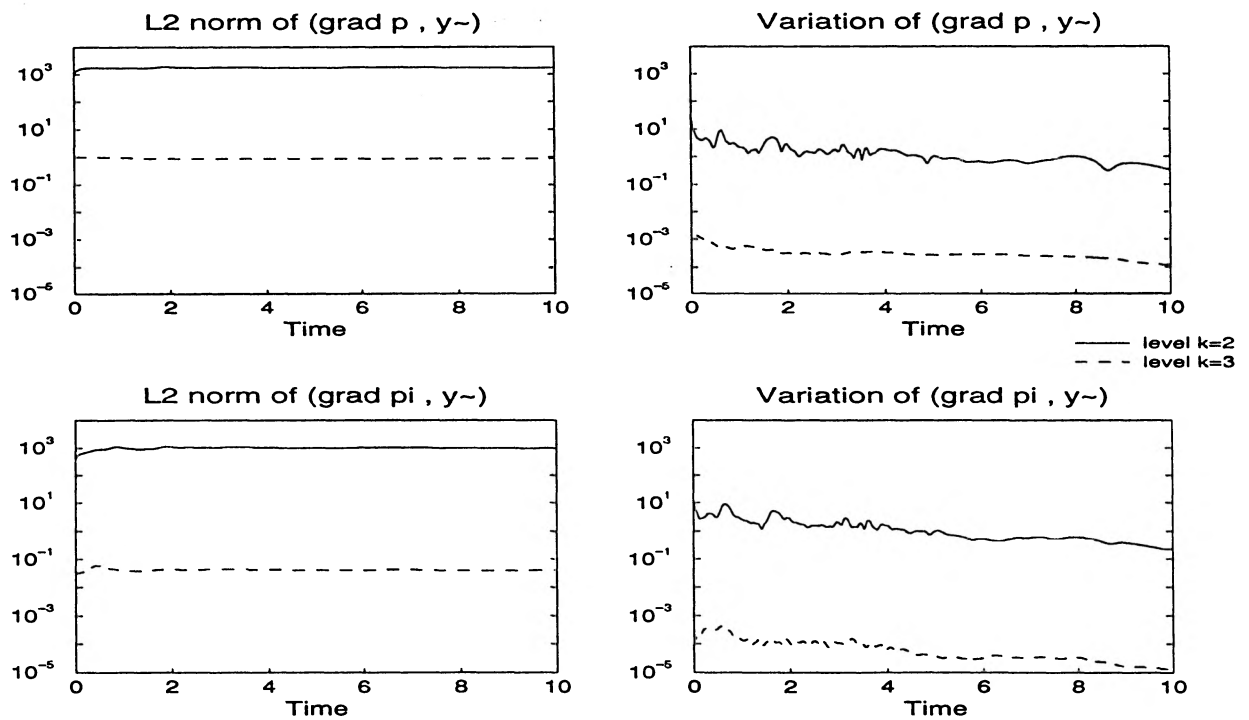


FIG. 2.39 – Evolution de la norme L^2 de $(\nabla \bar{p}_{k-1}, \tilde{y}_{k-1})$ et $(\nabla \bar{\pi}_k, \tilde{y}_{k-1})$ et de leurs variations respectives sur un pas de temps.

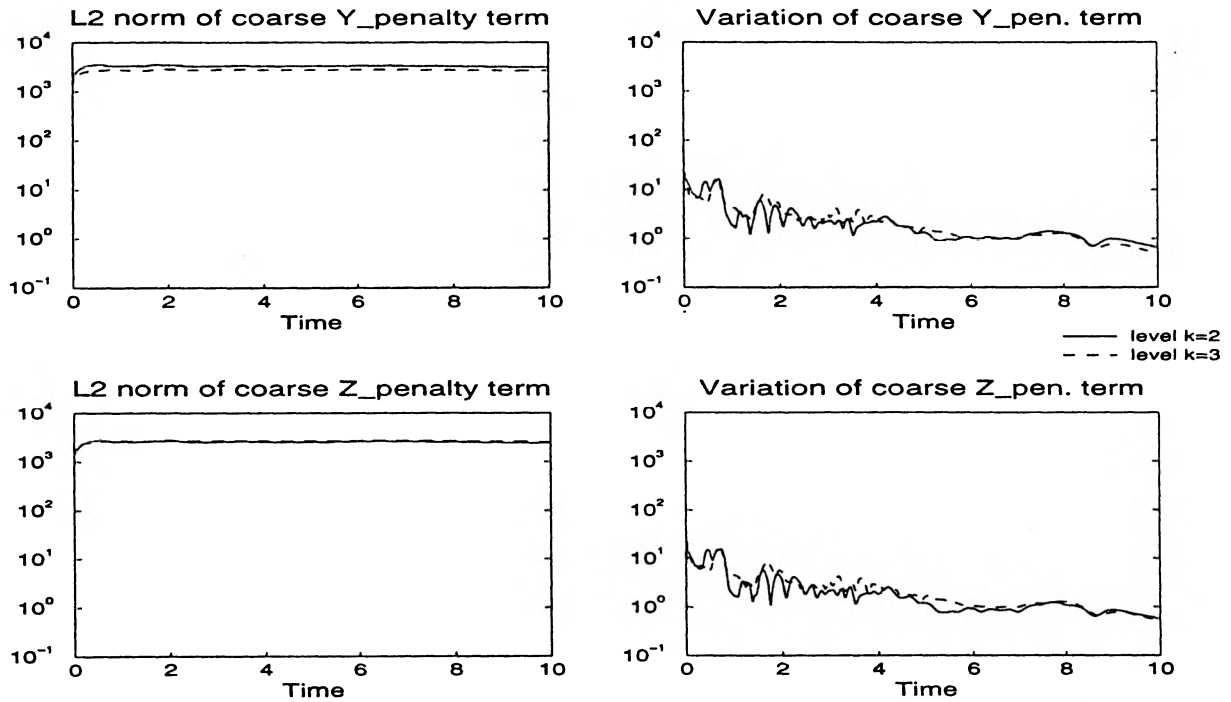


FIG. 2.40 – Evolution de la norme L^2 des termes de pénalisation $\frac{1}{\epsilon} T\mathcal{B}_{gg}^{(d)} (C_{gg}^{(d)})^{-1} \mathcal{B}_{gg}^{(d)} Y_{d-1}$ et $\frac{1}{\epsilon} T\mathcal{B}_{gg}^{(d)} (C_{gg}^{(d)})^{-1} \mathcal{B}_{gf}^{(d)} Z_d$ et de leurs variations respectives sur un pas de temps.

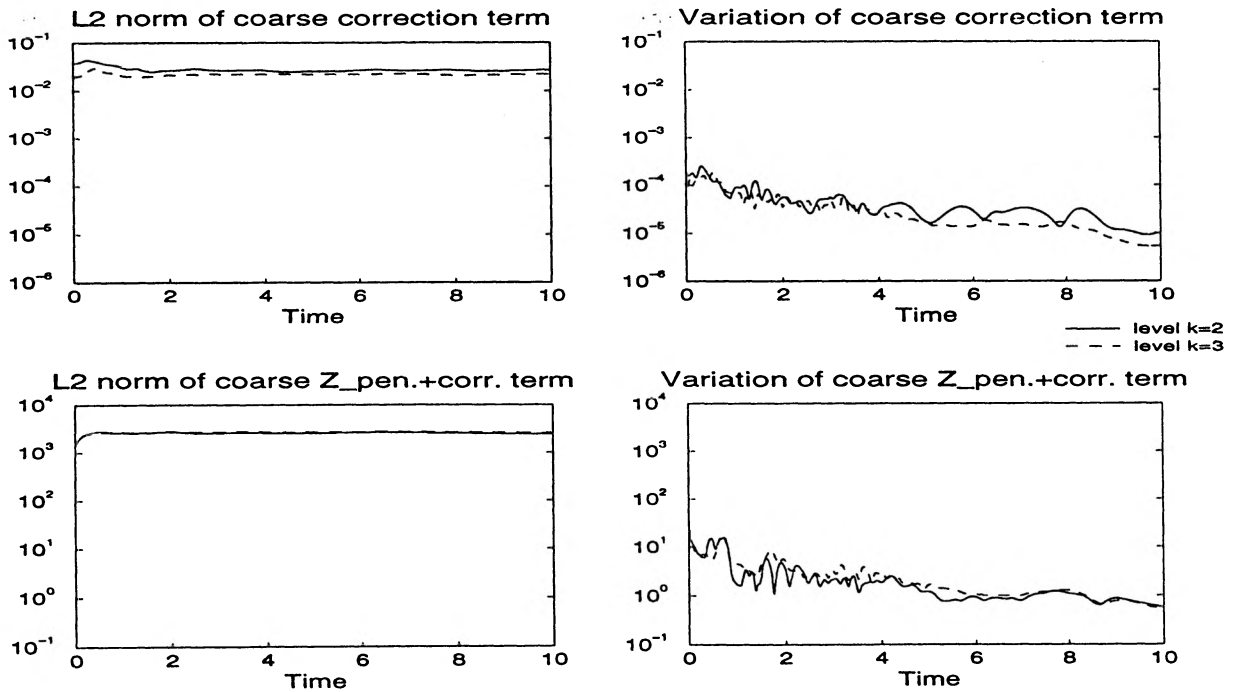


FIG. 2.41 – Evolution de la norme L^2 du terme de correction $T\mathcal{B}_{gg}^{(d)} (C_{gg}^{(d)})^{-1} C_{gf}^{(d)} \Pi_d$ et du terme somme $-\frac{1}{\epsilon} T\mathcal{B}_{gg}^{(d)} (C_{gg}^{(d)})^{-1} \mathcal{B}_{gf}^{(d)} Z_d + T\mathcal{B}_{gg}^{(d)} (C_{gg}^{(d)})^{-1} C_{gf}^{(d)} \Pi_d$ ainsi que de leurs variations respectives sur un pas de temps.

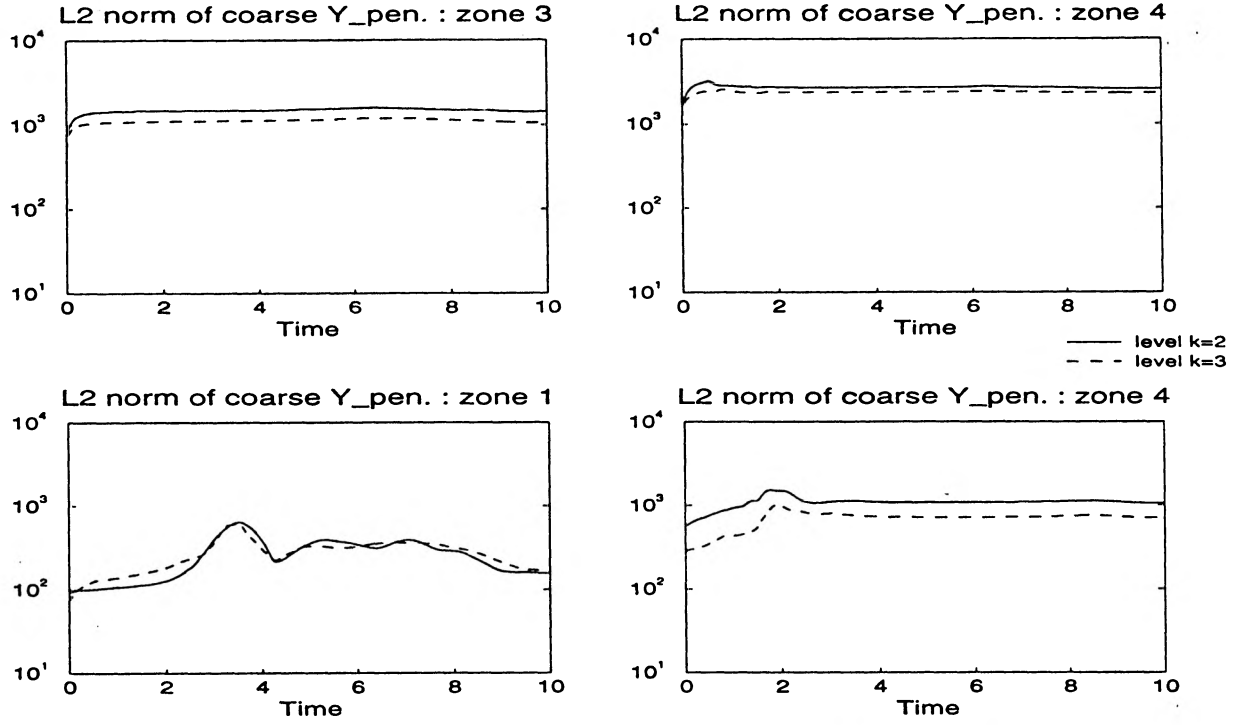


FIG. 2.42 – Evolution par zone de la norme L^2 du terme $\frac{1}{\epsilon} T \mathcal{B}_{gg}^{(d)} (C_{gg}^{(d)})^{-1} \mathcal{B}_{gg}^{(d)} Y_{d-1}$.

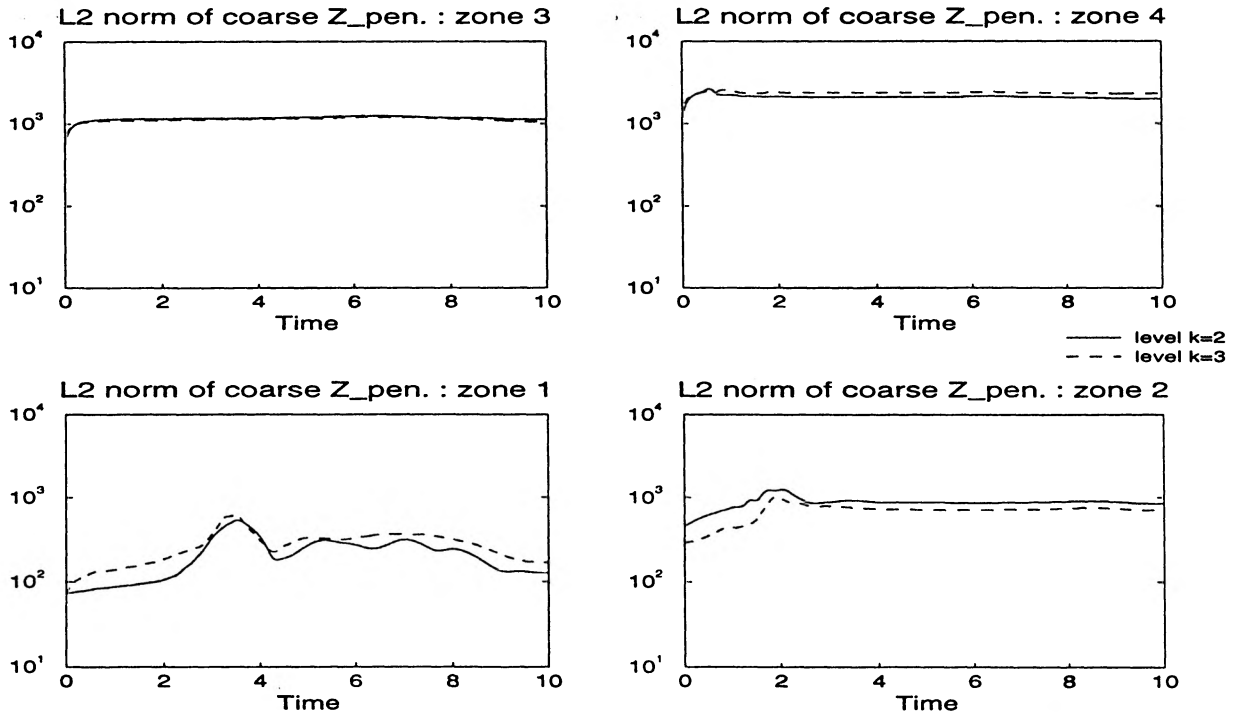


FIG. 2.43 – Evolution par zone de la norme L^2 du terme $\frac{1}{\epsilon} T \mathcal{B}_{gg}^{(d)} (C_{gg}^{(d)})^{-1} \mathcal{B}_{gf}^{(d)} Z_d$.

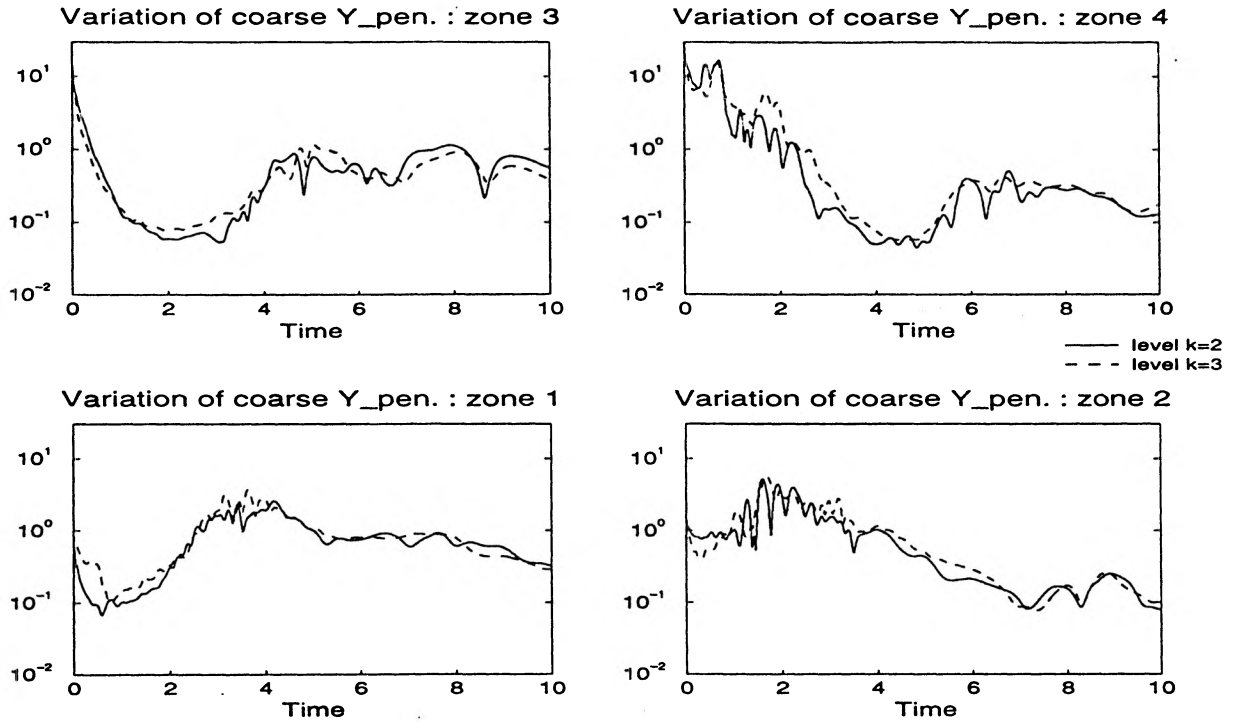


FIG. 2.44 – Evolution par zone de la norme L^2 de la variation sur un pas de temps du terme $\frac{1}{\epsilon} T_{\mathcal{B}_{gg}}^{(d)} (C_{gg}^{(d)})^{-1} \mathcal{B}_{gg}^{(d)} Y_{d-1}$.

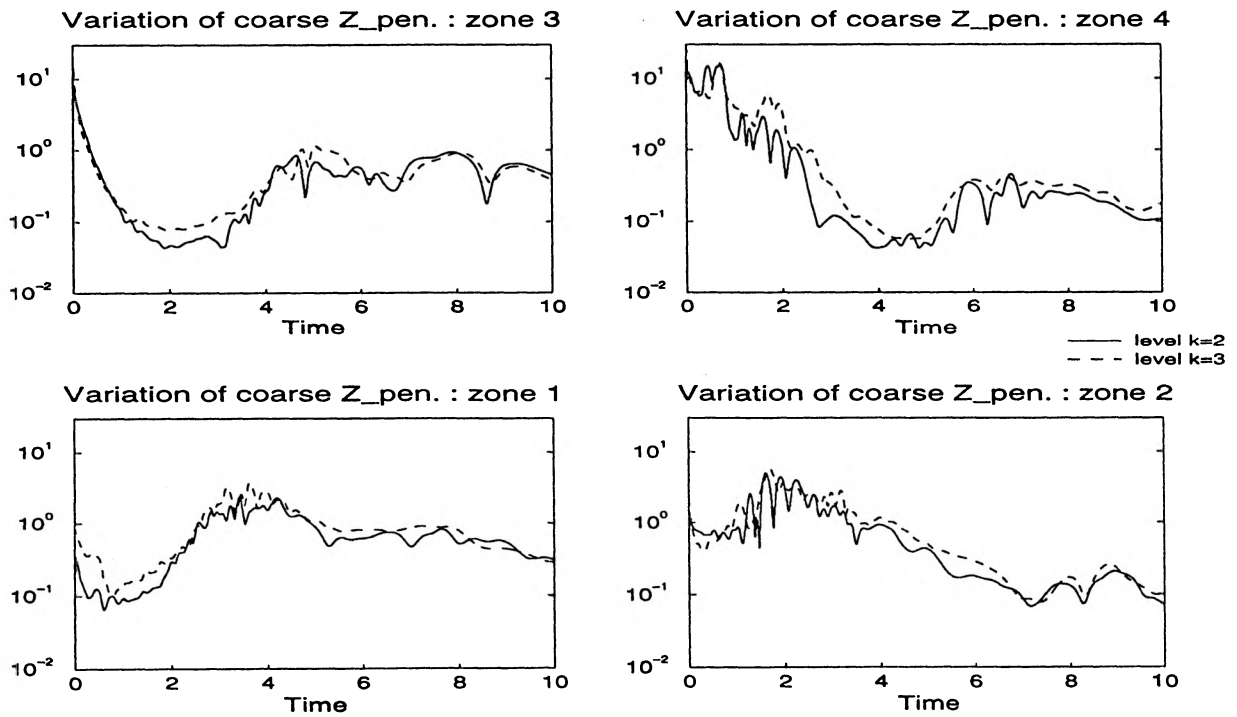


FIG. 2.45 – Evolution par zone de la norme L^2 de la variation sur un pas de temps du terme $\frac{1}{\epsilon} T_{\mathcal{B}_{gg}}^{(d)} (C_{gg}^{(d)})^{-1} \mathcal{B}_{gf}^{(d)} Z_d$.

Bibliographie

- [1] R.E. BANK, T.E. DUPONT, H. YSERENTANT
The hierarchical basis multigrid method
(1988) *Numer. Math.*, vol 52, pp 427-458
- [2] M. BERCOVIER
Perturbation of mixed variational problems. Application to mixed finite element methods
(1978) *R.A.I.R.O. Num. Analysis*, vol 12, n 3, pp 211-236
- [3] M. BERCOVIER, M. ENGELMAN
A finite element for the numerical solution of viscous incompressible flows
(1979) *Journal of Computational Physics*, vol 30, pp 181-201
- [4] C.H. BRUNEAU, C. JOURON
An efficient scheme for solving steady incompressible Navier-Stokes equations
(1990) *Journal of Computational Physics*, vol 89, pp 389-413
- [5] M. BERCOVIER, O. PIRONNEAU
Error estimates for finite element method solution of the Stokes problem in the primitive variables
(1973) *Numer. Math.*, vol 33, pp 211-224
- [6] F. BREZZI, M. FORTIN
Mixed and hybrid finite element methods
(1991) Springer-Verlag; New York
- [7] J. CAHOUE
On some difficulties occurring in the simulation of the incompressible fluid flows by Domain Decomposition methods
(1987) *Proceedings of the 1st International Symposium on Domain Decomposition Methods for P.D.E. (Paris)*; SIAM; pp 313-332
- [8] J.P. CHEHAB
Solution of generalized Stokes problems using hierarchical methods and Incremental Unknowns
(1996) *Applied Numerical Mathematics*, vol 21, pp 9-42

- [9] A.J. CHORIN
On the convergence of discrete approximations to the Navier-Stokes equations
(1969) *Math. Comp.*, vol 23, pp 341-353
- [10] Z. CAI, C.I. GOLDSTEIN, J.E. PASCIAK
Multilevel iteration for mixed finite element systems with penalty
(1993) *SIAM J. Sci. Comput.*, vol 14, n 5, pp 1072-1088
- [11] E. CUTHILL, J. MCKEE
Reducing the bandwidth of sparse symmetric matrices
(1969) *Proceedings of the 24th National Conference of the ACM*, pp 157-172
- [12] G.F. CAREY, K.C. WANG, W.D. JOUBERT
Performance of iterative methods for newtonian and generalized newtonian flows
(1989) *Int. J. of Num. Meth. in Fluids*, vol 9, pp 127-150
- [13] I. EKELAND, R. TEMAM
Analyse convexe et problèmes variationnels
(1974) Dunod; Paris
- [14] M. FORTIN, R. GLOWINSKI
Méthodes de lagrangien augmenté
(1982) Dunod; Paris
- [15] C. FARHAT, F.X. ROUX
Implicit parallel processing in structural mechanics
(1994) *Comput. Mech. Advances; North-Holland*, vol 2, pp 1-124
- [16] A. GEORGE, J.R. GILBERT, J.W.H. LIU Ed.
Graph theory and sparse matrix computation
(1993) Springer-Verlag; New York
- [17] O. GOYON
High-Reynolds number solutions of Navier-Stokes equations using Incremental Unknowns
(1996) *Computer Methods in Appl. Mech. and Eng.*, vol 130, pp 319-335
- [18] U. GHIA, K.N. GHIA, C.T. SHIN
High-Re solutions for incompressible flow using the Navier-Stokes equations and a multigrid method
(1982) *Journal of Computational Physics*, vol 48, pp 387-411
- [19] V. GIRAULT, P.A. RAVIART
Finite element methods for Navier-Stokes equations; theory and algorithms
(1986) Springer-Verlag; Berlin

- [20] G. GOLUB, C. VAN LOAN
Matrix computations
(1983) North Oxford Academic, Oxford
- [21] C. JOURON
(1996) *Communication privée*
- [22] J. LAMINIE
Quelques aspects de la résolution numérique de problèmes de la mécanique des fluides : discrétisation par éléments finis, mise en œuvre sur calculateurs parallèles
(1995) *Thèse d'état - Orsay - France*
- [23] J. LAMINIE, F. PASCAL and R. TEMAM
Implementation and numerical analysis of the nonlinear Galerkin methods with finite elements discretization
(1994) *Applied Numerical Mathematics*, vol 15, pp 219-246
- [24] P. LE TALLEC
Domain decomposition methods in computational mechanics
(1994) *Comput. Mech. Advances ; North-Holland*, vol 1, pp 121-220
- [25] T.A. MANTEUFFEL
Shifted incomplete Cholesky factorization
(1978) *Sparse Matrix Proceedings, SIAM Philadelphia*, pp 41-61
- [26] T.A. MANTEUFFEL
An incomplete factorization technique for positive definite linear systems
(1980) *Mathematics of Computation*, vol 34, n 150, pp 473-497
- [27] F. PASCAL
Méthodes de Galerkin non linéaires en discrétisation par éléments finis et pseudo-spectrale. Application à la mécanique des fluides
(1992) *Thèse Univ. Paris-Orsay*
- [28] F. PASCAL
(1996) *Communication privée*
- [29] P. POULLET
1) Simulation de la turbulence par la méthode des Grandes Echelles 2) Résolution d'équations de la mécanique des fluides par la méthode des Inconnues Incrémentales
(1996) *Thèse Univ. Paris-Orsay*
- [30] J. SHEN
Hopf bifurcation of the unsteady regularized driven cavity flow
(1991) *Journal of Computational Physics*, vol 95, pp 228-245

- [31] J. SHEN
On error estimates of the penalty method for unsteady Navier-Stokes equations
(1995) *SIAM J. Numer. Anal.*, vol 32, n. 2, pp 386-403
- [32] J. SHEN
A new pseudocompressibility method for the incompressible Navier-Stokes equations
(1996) *Applied Numerical Mathematics*, vol. 21, pp 71-90
- [33] R. TEMAM
Une méthode d'approximation de la solution des équations de Navier-Stokes
(1968) *Bull. Soc. Math. France*, vol 96, pp 115-152
- [34] R. TEMAM
Sur l'approximation de la solution des équations de Navier-Stokes par la méthode des pas fractionnaires I
(1969) *Arch. Rational Mech. Anal.*, vol 32, pp 135-153
- [35] R. TEMAM
Sur l'approximation de la solution des équations de Navier-Stokes par la méthode des pas fractionnaires II
(1969) *Arch. Rational Mech. Anal.*, vol 33, pp 377-385
- [36] R. TEMAM
Navier-Stokes equations
(1984) North Holland; Amsterdam
- [37] F. THOMASSET
Implementation of finite element methods for Navier-Stokes equations
(1981) Springer-Verlag; New York
- [38] H. YSERENTANT
On the multi-level splitting of finite element spaces
(1986) *Num. Math.*, vol 49, pp 379-412

Conclusions et développements

Ce travail a été consacré à la mise au point d'une nouvelle méthode, de *type incrémental*, permettant d'approcher la solution d'équations non linéaires évolutives, en dimension deux d'espace, issues de la mécanique des fluides (i.e. les équations de Burgers visqueuses et de Navier-Stokes incompressibles). Grâce à l'utilisation des bases hiérarchiques, dans le cadre d'une discrétisation par éléments finis, nous avons obtenu un découpage de la solution en ses grandes structures et dans leurs corrections. Cette décomposition est à la base de notre méthode multi-niveaux qui porte sur le traitement différent des diverses échelles de la solution.

Dans la première partie nous avons introduit une approche nouvelle consistant à négliger les variations en temps des termes d'interaction entre les grandes structures et leurs incréments. L'application numérique a été faite sur l'équation de Burgers, pour une hiérarchisation de la vitesse sur plusieurs niveaux de grilles. Une analyse détaillée des ordres de grandeur des différents termes appartenant aux équations projetées sur la grille grossière, ainsi que des variations sur un pas de temps des termes de couplage, a été mise en place afin de montrer numériquement que les termes d'interaction évoluent faiblement au cours de la simulation en temps. Le différent traitement des composantes hiérarchiques du champ de vitesse est alors réalisé grâce à la mise au point d'un algorithme multi-niveaux en espace et en temps. La structure en V-cycles de cet algorithme a permis de ne pas résoudre le problème étudié sur tous les niveaux à chaque pas de temps. Il en résulte un gain significatif en temps de calcul, pendant que la solution continue à rester assez proche de celle obtenue par une méthode classique, ceci grâce à la déduction de plusieurs estimations *a priori* qui ont permis d'effectuer un contrôle numérique sur l'erreur introduite dans les équations approchées.

Dans la deuxième partie nous avons généralisé cette méthode à la résolution numérique des équations de Navier-Stokes, en vitesse-pression, en simplifiant les équations de départ avec une méthode de pénalisation. Dans le premier chapitre, après avoir expliqué les motivations qui nous ont conduit à résoudre les équations sous la forme pénalisée et après avoir testé plusieurs éléments finis ainsi que plusieurs préconditionnements des méthodes itératives, nous nous sommes orientés vers une formulation éléments finis mixtes 4P1-P1 et une résolution directe, à chaque pas de temps, du système linéaire issu de la pénalisation de la pression. Le solveur mis au point, qui a permis d'obtenir des simulations numériques satisfaisantes pour des écoulements convergeant vers un état stationnaire, est devenu l'outil nécessaire aux développements présentés dans le dernier chapitre, consacré à l'extension des idées et des techniques de la première partie. Après

avoir constaté que la hiérarchisation des inconnues du problème ne s'est révélée aisée que sur deux niveaux, l'introduction d'une pression auxiliaire nous a permis d'étendre le processus à un nombre quelconque de niveaux. Malheureusement, les termes dépendant de cette pression ont présenté des variations en temps trop importantes par rapport aux autres termes en vitesse. Une analyse, analogue à celle effectuée pour Burgers, a été faite pour tous les termes des équations projetées ainsi que pour les variations en temps de ces termes. Il en résulte que la stratégie multi-niveaux ne semble pas être totalement adaptée aux équations de Navier-Stokes, car les variations des termes de couplage ne sont négligeables que sur des intervalles de temps trop courts.

Une nouvelle approche a été ensuite proposée. L'idée porte sur un couplage de la hiérarchisation avec la méthode de décomposition de domaine, ce qui permettra de localiser le traitement des interactions entre les différentes composantes hiérarchiques de la solution. L'analyse, faite zone par zone et analogue à celle réalisée sur le domaine tout entier, a permis de déceler des comportements physiques locaux qui n'étaient pas visibles dans la séparation d'échelle relative au domaine tout entier. Ainsi, localement, on retrouve le comportement attendu des variations des termes d'interaction : ces variations sont négligeables sur des longs intervalles de temps.

Bien que la généralisation pour les équations de Navier-Stokes de la méthode mise en place pour résoudre le problème de Burgers ne soit pas encore terminée, nous entrevoyons la possibilité d'un développement à court terme d'une résolution des équations de Navier-Stokes par une méthode multi-niveaux auto-adaptative couplée avec une méthode de décomposition de domaine. Ce couplage entre les deux méthodes doit passer par plusieurs étapes. La première porte sur la parallélisation du code séquentiel et est déjà en cours d'élaboration.

Il est bien connu que l'utilisation de la méthode de décomposition de domaine (que nous avons choisi sans recouvrement) augmente le potentiel de parallélisme du calcul en éléments finis. En considérant notre but, nous orientons nos choix vers l'utilisation d'une méthode de type duale. La méthode directe, dite du *complément de Schur* (cf. [22] et [24]), certainement plus facile à mettre en place, n'est pourtant pas adaptée au choix d'une pression P1 continue. Il en résulte l'impossibilité d'écrire une décomposition de Schur afin de ramener le problème dans le domaine Ω à un problème sur l'interface des sous-domaines. C'est pour cette raison que nous nous sommes orientés vers une formulation duale du problème (voir par exemple [7] et [15]). Dans cette méthode, l'égalité de la solution sur l'interface apparaît comme une contrainte du problème, contrainte qui est imposée sous sa forme faible. Si Γ_{ij} dénote l'interface entre les sous-domaines Ω_i et Ω_j , la méthode duale impose la continuité du vecteur des contraintes λ qui, le long de Γ_{ij} , est défini par la relation suivante :

$$\lambda_{ij} = \nu \frac{\partial u_i}{\partial n} - p_i n \quad ,$$

avec $\lambda_{ij} \in H^{-1/2}(\Gamma_{ij})$, $u_i = u|_{\Omega_i}$, $p_i = p|_{\Omega_i}$ et n la normale unitaire externe à Ω_i . Le vecteur des contraintes λ , que nous retrouvons facilement en écrivant la formulation variationnelle associée au problème de Stokes généralisé, apparaît comme la variable duale

de la vitesse sur l'interface. Le problème d'origine est donc équivalent à un problème de point-selle qui sera résolu sur l'interface grâce à la détermination du vecteur λ . Il nous reste à résoudre, à l'intérieur de chaque sous-domaine, un problème de Stokes généralisé avec des conditions de Neumann sur le(s) bord(s) appartenant à l'interface. La parallélisation du code nous permettra également d'accéder à des maillages plus fins et donc d'envisager la détermination de solutions hautement instationnaires (à grand nombre de Reynolds).

Des études préliminaires ont déjà été effectuées afin de montrer que cette approche duale se hiérarchise correctement dans chaque sous-domaine et qu'elle est compatible avec la stratégie de V-cycles qu'il faudra développer dans l'étape suivante. Le couplage de la méthode de décomposition de domaine avec notre stratégie multi-niveaux dans chaque zone demande seulement la mise au point de données (matrices et vecteurs relatifs à l'interface) ayant une structure multi-échelles.

Il conviendra ensuite de programmer les V-cycles sur le code parallèle obtenu. En même temps il faudra préciser, numériquement et théoriquement, les diverses estimations *a priori* pour contrôler l'erreur introduite dans les équations restreintes aux grilles plus grossières. De plus, il faudra proposer une autre décomposition hiérarchique des inconnues afin de réitérer le processus sur plusieurs niveaux.

Les perspectives à plus long terme concernent la mise au point d'un autre solveur de Navier-Stokes. Dans un premier temps, une pénalisation seulement numérique, obtenue en ajoutant le terme $\epsilon \frac{\partial p}{\partial t}$ dans l'équation de continuité, devrait permettre de résoudre les difficultés liées au mauvais conditionnement de la matrice à inverser à chaque pas de temps, en pouvant employer alors des méthodes itératives. Puisque la pénalisation n'a pas résolu les problèmes liés à la hiérarchisation de la pression ni à la contrainte d'incompressibilité, il conviendra ensuite de s'orienter vers d'autres méthodes de résolution, comme par exemple les méthodes de projection introduites par Chorin et Temam (cf. [9], [34] et [35]). En conclusion, d'autres écoulements physiquement plus représentatifs, comme par exemple l'écoulement autour d'un obstacle ou le problème de la marche, devront être testés, ceci afin de valider complètement et définitivement notre stratégie multi-résolution. Le choix d'autres éléments finis mixtes pourra être également envisagé, la condition fondamentale étant celle de considérer toujours des éléments aisément hiérarchisables.