

THÈSES D'ORSAY

PASCAL POULLET

1. Simulation de la turbulence par la méthode des grandes échelles. 2. Résolution d'équations de la mécanique des fluides par la méthode des inconnues incrémentales

Thèses d'Orsay, 1996

http://www.numdam.org/item?id=BJHTUP11_1996__0455__A1_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/>*

63791

ORSAY
n° 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

Pascal POULLET

Sujet: 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

Soutenue le : 15 Janvier 1996 devant la Commission d'examen

M.	BREZINSKI	
M.	BRUNEAU	Rapporteur
Mme	CANCELIER	
M.	LAMINIE	
M.	SCHEURER	Rapporteur
M.	TEMAM	Président

A Caroline,

A Suzie et Jacques,

Je remercie Monsieur Roger Temam, mon directeur de thèse, de m'avoir proposé un sujet aussi passionnant. Après avoir bien orienté mes premiers pas de jeune chercheur en me prodiguant ses bons conseils, il me laissa toute sa confiance pour mener à bien ce travail au Laboratoire d'Analyse Numérique d'Orsay.

Je remercie également Monsieur Bruno Scheurer, de m'avoir accepté dans son équipe au Centre d'Etudes de Limeil-Valenton pendant mon service militaire et d'avoir bien voulu rapporter cette thèse.

Je remercie aussi Monsieur Charles-Henri Bruneau, d'avoir donné son avis précieux sur ce travail et je suis très sensible à l'honneur qu'il me fait ainsi.

Merci à Monsieur Jacques Laminie, d'avoir bien voulu juger cette thèse. En plus de partager son savoir, j'ai souvent apprécié sa disponibilité et sa bonne humeur.

Je remercie Madame Claudy Cancelier et Monsieur Claude Brezinski de me faire l'honneur d'être membres de mon jury.

C'est avec un réel plaisir que je remercie mon fidèle compagnon du front numérique Olivier Goyon. Par son esprit opiniâtre et l'amitié qu'il a su me témoigner, il m'a épaulé tout au long de ces dernières années. Nous avons mené un excellent travail d'équipe et je compte bien sur les fruits de notre collaboration. J'espère que l'avenir nous réservera une récolte encore abondante.

Bien sûr, je n'oublie pas Marc Sancandi à qui je suis très heureux de témoigner ici de mon amitié et mon estime. Je garde un très bon souvenir de la fructueuse collaboration menée avec lui.

Frédéric Pascal mérite bien aussi que je ne l'oublie pas. Je lui dois une "fière chandelle" pour toutes ses recommandations et son savoir-faire. Il a su me porter toute son attention à maintes reprises, je l'en remercie. Ceci est aussi vrai pour Arnaud Debbusche.

Merci à Messieurs Claude Jouron et Jean-Claude Saut ainsi que tous les autres membres du Laboratoire d'Analyse Numérique d'Orsay, notamment Madame Danielle Le Meur pour sa gentillesse.

Je profite aussi de cette occasion pour saluer tous mes compagnons numériques et théoriciens, Caterina, Catherine, Azzedine, Cédric, Hervé et Seif, pour ne citer qu'eux. L'ambiance dans les salles du deuxième étage a souvent été formidable.

Enfin, je ne saurais oublier ma femme Caroline qui m'a patiemment soutenu tout au long de ces dernières années de dur labeur, je remercie également mon père et ma mère et je leur dédie à tous les trois cette thèse.

Turbulence simulation by the Large Eddy Scale simulation method

Resolution of fluid mechanics equations with the Incremental Unknowns method

ABSTRACT :

In this work, we explore two different areas : the development of new simulation method of incompressible viscous fluid flows and the study of the Large Eddy Simulation (L.E.S.) method for a compressible flow in an isotropic homogeneous turbulence state.

For our L.E.S. study, we filter the whole field of the three-dimensionnal Navier-Stokes equations, and model the action of the residual component onto the solving mean component. Moreover, from the modified equation concept, a theoretical study on the numerical scheme enables us to obtain the real influence of the discretization on the closing models.

In a mathematical context of the expansion of the nonlinear Galerkin methods, the Incremental Unknowns (I.U.) have been introduced in order to hierarchize the unknowns in finite differences. First of all, we verify the efficiency of the I.U. as a preconditioner for solving different elliptic problems (2D and 3D-Dirichlet, 2D-Neuman). Secondly, we find a unique solution in a Sobolev space for a steady-state Navier-Stokes like equations, develop several implicit solving methods and prove the efficiency of the I.U. in many situations. Thirdly, we propose to find numerical solutions of the Stokes and incompressible 2D-Navier-Stokes problems with an orthogonal projection algorithm and using staggered mesh (M.A.C.). Our method is valid for the lid driven cavity problem for a Reynolds number of 5000. And Finally, we start developing a new adaptive multilevel method by creating a new hierarchization of a M.A.C. mesh, introducing estimators of global dynamic of the flow. Thus, we propose new schemes which seem adapted for the dynamic of solenoidal flows for high Reynolds numbers.

KEY WORDS :

Poisson equations - Neuman problem - Navier-Stokes equations - finite volume - finite difference - multilevel discretization - modified equation - preconditioning systems - homogeneous isotropic turbulence.

Table des matières

INTRODUCTION	1
I SIMULATION DE LA TURBULENCE PAR LA METHODE DES GRANDES ECHELLES	5
Introduction	7
1 Problème physique et simulations numériques	9
1.1 Principe de la LES	9
1.2 Equations traitées	10
1.3 Modèles de Fermeture	12
1.4 Simulations numériques	13
1.4.1 Résultats du modèle de Smagorinsky	15
1.4.2 Résultats du modèle “Fonction de Structure”	21
1.5 Conclusions	23
2 Etude du schéma numérique	25
2.1 Schéma Numérique de Vreman <i>et al.</i>	26
2.2 Notion d’équation équivalente	29
2.3 Anisotropie théorique du schéma	32
2.4 Propriété régularisante du schéma	35
2.5 Schéma régularisant et schéma 4-isotrope	36
2.6 Comparaison théorie-simulations	38
2.7 Conclusions	47
3 Quelques remarques techniques	49
3.1 Discrétisation du tenseur de taux de déformation	49
3.2 Critique des modèles de fermeture utilisés	52
Conclusions et développements	53
Liste des Annexes	55

A	Analyse du comportement de ν en fonction de Re	55
B	Définition du système sans dimension	57
B.1	Définition des variables adimensionnées	57
B.2	Adimensionnement des équations	58
B.3	Choix des paramètres de calcul	61
B.4	Obtention des grandeurs caractéristiques physiques	62
C	Initialisation : génération aléatoire d'un champ turbulent	65
C.1	Le problème	65
C.2	La méthode	66
C.3	Expression du spectre initial	67
C.4	Méthode de calcul du spectre	67
D	Equation equivalente : schéma de Vremann	69
	Bibliographie	74

II LE PRECONDITIONNEMENT PAR LES INCONNUES INCREMENTALES (I.U.) **75**

	Introduction	77
1	Résolution de quelques problèmes elliptiques linéaires	79
1.1	Problèmes modèles et discrétisations	79
1.2	Présentation de deux méthodes multi-niveaux	80
1.2.1	Méthode Multigrilles	80
1.2.2	Méthode des Inconnues Incrémentales	84
1.2.3	Résultats numériques	92
1.3	Conclusions	100
2	Résolution d'un problème non linéaire stationnaire	101
2.1	Existence et unicité dans $H_0^1(\Omega)$ et système discret	102
2.2	Méthodes de Newton	108
2.2.1	Newton-BICGSTAB dans la base des I.U.	109
2.2.2	Résolution de systèmes linéaires non symétriques	110
2.2.3	Quelques Préconditionnements pour BI-CGSTAB	112
2.3	Méthode de GMRES non linéaire	114
2.3.1	Résolution de systèmes linéaires non symétriques par GMRES	115
2.3.2	Les I.U. pour preconditionner GMRES non linéaire	117
2.4	Résultats Numériques	120
2.4.1	Présentation des tests	120
2.4.2	Résultats pour les méthodes de Newton	123

2.4.3	Résultats pour les méthodes de GMRES	126
2.5	Conclusions	128
Liste des Annexes		130
A Discrétisation du problème de Neumann		131
Bibliographie		137
III RESOLUTION DES EQUATIONS DE NAVIER-STOKES BI-DIMENSIONNELLES		139
Introduction		141
1	Présentation du problème	143
1.1	Rappels théoriques et formulation	143
1.2	Problèmes modèles et Visualisation	145
1.2.1	La cavité carrée entraînée	145
1.2.2	Visualisation des résultats	147
2	Méthode de différences finies et méthode M.A.C.	149
2.1	Problème discret	149
2.1.1	Schéma numérique	150
2.1.2	Conditions aux limites	152
2.1.3	Stabilité linéaire	153
2.2	Algorithme de projection et préconditionnements	157
2.3	Résultats numériques	158
2.3.1	Problèmes de Stokes et de Stokes généralisé	159
2.3.2	Problème de Navier-Stokes	162
3	Développement d'une méthode multi-niveaux	173
3.1	Hiérarchisation des inconnues pour un maillage MAC à l'aide des <i>Staggered Incremental Unknowns</i> (S.I.U.)	174
3.1.1	S.I.U. pour la vitesse	175
3.1.2	S.I.U. pour la pression	178
3.2	Analyse "a posteriori" de la solution	182
3.3	Problème de Stokes généralisé sur la grille grossière	187
3.3.1	Analyse des opérateurs intervenant dans la base nodale	187
3.3.2	Restriction directe	189
3.3.3	Restriction sur la forme variationnelle	192
3.3.4	Résolution du problème grossier	195
3.4	Schéma de type Galerkin Non Linéaire	196

Conclusions et développements	199
Bibliographie	203

INTRODUCTION

De nombreuses activités, de nos jours, sont en relation avec la mécanique des fluides. Les deux domaines que nous allons explorer représentent deux approches différentes de deux phénomènes physiques très proches : la conception de nouvelles méthodes de simulation d'écoulements de fluides visqueux incompressibles et l'étude d'une plus ancienne méthode en état de turbulence idéale.

Il nous faut rappeler que le nombre de degrés de liberté permettant la simulation d'un état de **turbulence homogène isotrope** tridimensionnelle est fonction du nombre de Reynolds (constante caractéristique de l'écoulement), soit : $N^* = Re^{\frac{9}{4}}$ [McC92] : ce qui nous amène, très vite à la limite des moyens matériels informatiques dont nous disposons actuellement. Plusieurs méthodes ont été proposées afin d'y remédier ; elles envisagent toutes un découpage des inconnues. Pour surmonter cette difficulté, différents auteurs (Léonard, Smagorinsky,...) ont repris l'idée de Reynolds, de décomposer les différents champs turbulents en une composante qui est calculée explicitement, et une autre dont l'effet sur le champ total est *modélisé*.

Notre première étude s'appuie sur une décomposition statistique des inconnues, en partie moyenne et partie fluctuante aléatoire (analogue à la décomposition de Reynolds) par le biais d'un procédé de filtrage obtenu par convolution. En appliquant cette technique aux équations de Navier-Stokes, on obtient des équations similaires, excepté la présence du tenseur de sous-maille qui relie les composantes aléatoires de la vitesse, il est appelé le **tenseur de Reynolds**. La simulation des grandes échelles se ramène alors à résoudre la composante de grande échelle dès que les corrélations dans lesquelles intervient la composante résiduelle sont modélisées. Afin d'étudier cette modélisation, ils nous a paru intéressant de comparer deux expressions différentes de la viscosité turbulente intervenant dans les relations de fermeture des équations de Navier-Stokes compressibles filtrées, pour un état de turbulence homogène isotrope. Notre choix s'est porté sur le modèle **Fonction de structure** proposé par Métais et Lesieur [ML91], et celui de **Smagorinsky** plus classique [McC92].

La résolution de ces équations filtrées se fait par un schéma temporel de type Runge-Kutta d'ordre 4, et par une discrétisation en volumes finis pour l'espace, utilisée avec succès par Vreman *et al.* [VGKZ92]. Cette discrétisation spatiale est obtenue en approchant, d'une part, les opérateurs différentiels présents dans les termes diffusifs, par des schémas aux différences centrés classiques, d'ordre 2. D'autre part,

les autres opérateurs différentiels du premier ordre sont approchés par moyenne des mêmes schémas centrés, sur les voisins dans le plan perpendiculaire à la direction dans laquelle est faite les différences. L'initialisation se fait de manière aléatoire par une procédure introduite par Roy [ROY80]. Le résultat de nos diverses simulations, met en évidence quelques interrogations auxquelles il est difficile d'apporter des réponses précises. Nous avons donc entrepris une étude théorique sur notre schéma numérique, en utilisant le concept de **l'équation modifiée**. En s'intéressant au spectre de densité d'énergie, ce remarquable outil nous a permis de trouver une propriété de régularisation et un défaut d'isotropie de notre schéma numérique. Bien qu'il ne soit pas possible d'envisager une configuration où le choix des coefficients de pondération permette de corriger ce défaut tout en gardant le même écart-type, la corrélation très forte des distributions d'anisotropie théorique et numérique nous prouve que le schéma numérique devrait tenir compte de cette propriété.

Ce développement éclaire l'influence certaine de la discrétisation sur les modèles de fermeture et permettra de concevoir des schémas plus adaptés à l'état physique que l'on espère simuler.

L'autre étude, s'appuyant sur de récents concepts mathématiques, cherche à développer une méthode multi-résolution aussi bien dans l'espace physique que dans le temps. En effet, grâce à la jonction entre la théorie des systèmes dynamiques et la théorie phénoménologique de la turbulence, on a pu noter l'apparition de nouveaux objets, les **variétés inertielles exactes ou approchées** [FMT88], ainsi que de nouvelles méthodes numériques les **méthodes de Galerkin non linéaire** [MT90]. Ces méthodes ont permis la création d'algorithmes utilisant les discrétisations spectrales et pseudo-spectrales [DJT93] ainsi que les éléments finis hiérarchiques [YSE86], [LPT90], [PAS92], [CLT95]. Dans ce contexte, les Inconnues Incrémentales permettent de hiérarchiser les inconnues en différences finies [TEM90]. Dans la première partie de ce travail, nous vérifions l'aspect préconditionneur que revêt la méthode sur différents problèmes modèles : les problèmes bidimensionnel et tridimensionnels de Dirichlet, le problème bidimensionnel de Neuman, ainsi qu'un problème non linéaire stationnaire de type Navier-Stokes dépendant de deux paramètres. Comme les Inconnues Incrémentales (I.U.) apparaissent comme une méthode multi-niveaux récente, il nous a paru intéressant, dans un premier temps, de développer une méthode multi-grilles classique, afin de mener une étude comparative entre ces deux méthodes. Ensuite, après avoir montré que la méthode la plus performante reste la méthode classique, nous vérifions que le préconditionnement par les I.U. est extrêmement efficace sur calculateurs vectoriels, pour la résolution de problèmes elliptiques. Les autres préconditionneurs testés sont le préconditionnement par la factorisation de Choleski incomplète IC(0) et celui d'Evans SSOR. Pour le problème non-linéaire de type Navier-Stokes, après avoir mené une étude théorique préliminaire, nous appliquons le préconditionnement par les I.U. à plusieurs méthodes implicites de résolution afin de calculer les solutions numériques du problème discret. Nous prouvons que nos algorithmes *incrémentaux* sont performants

et efficaces pour de petites valeurs du paramètre voué à simuler le caractère évolutif du problème évolutif discret sous-jacent. Une autre étude portant sur la résolution du problème continu est récapitulée dans [GP95].

La deuxième partie consiste plus précisément à entreprendre la résolution des problèmes de Stokes et Navier-Stokes, et de développer de nouvelles méthodes permettant d'y arriver de façon efficace. Nous commençons par implémenter une méthode de résolution efficace pour résoudre les équations de Stokes et de Navier-Stokes. Pour le problème de Stokes, la méthode est basée sur un algorithme de projection orthogonale sur l'espace à divergence nulle [BRA92] et utilise la discrétisation en grilles décalées, identique à celle qu'utilise la méthode M.A.C. [PT83]. L'algorithme consiste alors, à appliquer une méthode de gradient conjugué sur le système projeté ayant la vitesse comme inconnue. Le projecteur étant défini par le biais d'un pseudo-inverse, chaque itération de gradient nécessite la résolution d'un problème où la pression est l'inconnue. Le problème de Navier-Stokes est discrétisé par un schéma temporel classique, semi-implicite (Cranck-Nicholson) pour l'opérateur linéaire de diffusion, explicite (Adams-Bashforth) pour le terme convectif. Dès que la condition de stabilité de type CFL est vérifiée, nous sommes donc amenés à résoudre un problème de Stokes généralisé à chaque itération. Pour réduire le coût de cette résolution, nous inversons la matrice intervenant sur la pression par un gradient conjugué préconditionné par l'inverse d'un laplacien discret. Cette inversion interne est faite par l'algorithme multi-grilles performant développé précédemment. Nous validons notre méthode de résolution pour les cas tests des cavités entraînées régularisées et non régularisées et des nombres de Reynolds allant jusqu'à 5000.

Une des caractéristiques du maillage utilisé est que le nombre d'inconnues dans chacune des directions n'est pas le même, et comme leur parité n'est pas la même ($N_x = N_y \pm 1$), il n'est pas possible de hiérarchiser un maillage M.A.C. à l'aide des Inconnues Incrémentales classiques. Nous proposons donc, une nouvelle hiérarchisation permettant de définir plusieurs grilles d'inconnues ayant toutes la même propriété du maillage M.A.C.. Pour une configuration à deux grilles, les inconnues de la grille fine privée des nœuds de la grille grossière sont appelées les **Staggered Incremental Unknowns (S.I.U.)**. Tout comme les (I.U.), les S.I.U. sont obtenues par les formules de Taylor et sont des incréments du second ordre du pas du maillage sur lequel ils sont définis. Grâce à cet outil, il est possible de décomposer le système discret directement hiérarchisé, et l'on s'aperçoit qu'une analyse globale de la solution met en évidence des estimateurs de dynamique non-linéaire globale de l'écoulement.

A partir du comportement asymptotique du fluide, il semble bien que dans les cas de solutions stationnaires, à certains instants de simulation il n'est pas nécessaire de calculer toutes les inconnues du système avec la même précision. C'est pour cette raison que nous envisageons de figer dans le temps (tout comme l'on fait d'autres auteurs pour une équation de Burgers [CLT95]), à certaines itérations de Navier-Stokes, le terme d'interaction entre les différentes structures ainsi que nos S.I.U.. Après projection du système discret sur la grille grossière, nous obtenons un

problème de Stokes généralisé non homogène sur la grille grossière, qu'il est toujours possible de résoudre avec un algorithme de projection sur l'espace à divergence nulle. Enfin, on développe une nouvelle classe de schémas numériques multi-niveaux qui nous semblent adaptés à la dynamique non linéaire présente dans l'écoulement de fluides incompressibles à grand nombre de Reynolds.

Il est possible d'interpréter ce travail comme la mise en œuvre numérique d'une modélisation des petites structures par le contrôle de leur évolution et leur interaction sur les grandes structures cohérentes présentes dans un écoulement instationnaire.

Partie I

SIMULATION DE LA TURBULENCE PAR LA METHODE DES GRANDES ECHELLES

Introduction

The Large Eddy Simulation (*LES*) can be interpreted as a compromise between the two concepts of simulation methods: Direct Numerical Simulation (*DNS*) and modeling by transport equations.

The basic idea is to compute directly the mean velocity of the flow with modeling the subgrid components correlation. We just tell that the decomposition sets in a **large-scale** component and a **small-scale** one, by a **filtering** step.

In the first chapter, we describe the method and give the equations which govern the flow. Then, we propose to study the comparison between simulations carried out two different closing models: the Smagorinsky subgrid model and the structure function of the velocity model. In that way, one can feel all the questions it is difficult to answer.

The second chapter is dedicated to a theoretical study on a reaction-diffusion equation. Thus, with the **modified equation** tool, we find a pretty linear analysis of our numerical scheme. And we verify the relationship between our analysis results and the numerical simulations which enables many clarifications.

In the last chapter, we deal with the influence of the discretization of the strain rate tensor over the closing model. We also make a critic analysis on the closing models tested mentioned in the first chapter.

We conclude and propose some developments.

This study has been made in conjunction with M. Sancandi during my military service in 1994 at CEA de Limeil-Valenton DAM/DMA/MCN, 94195 Villeneuve-Saint-Georges.

Chapitre 1

Problème physique et simulations numériques

Dans ce chapitre, nous allons comparer deux modèles de fermeture en vue de simuler l'état de Turbulence Homogène Isotrope (THI) pour un fluide quasi-incompressible. Il est difficile de définir précisément la notion de turbulence. D'après les expériences de laboratoire, un fluide est dit turbulent quand, pour deux configurations initiales très proches, il est difficile d'obtenir des solutions présentant les mêmes caractéristiques au même instant de simulation. Ce faisant, la formulation statistique a été souvent prise pour la présentation et l'analyse de ce phénomène, faisant appel à la théorie ergodique. L'état physique THI se caractérise de plus, par une homogénéité spatiale des propriétés statistiques des inconnues, ainsi qu'une invariance des moments de la vitesse par rapport au groupe des rotations.

Nous commençons par présenter le principe de la simulation des grandes échelles (LES). Ensuite, nous exposons les équations compressibles filtrées résolues par notre implémentation. La fermeture du problème nécessite une écriture de la viscosité turbulente : nous donnons les deux écritures, l'une selon Smagorinsky et l'autre, suivant Métails *et al.*.

Après avoir précisé les bases de la comparaison, nous donnons les résultats numériques de nos simulations et concluons brièvement.

1.1 Principe de la LES

Faisant suite à l'introduction, donnons une description plus détaillée du principe LES.

Il se compose de cinq étapes :

- *étape initiale*

Nous considérons que nous avons à résoudre une équation d'évolution non linéaire, avec de "bonnes conditions aux limites" et une "bonne condition

initiale”, de la forme suivante :

$$\frac{\partial u}{\partial t} + \mathcal{L}(u) + \mathcal{N}(u) = 0, \quad (1.1)$$

où $\mathcal{L}$ est un opérateur linéaire et $\mathcal{N}$ une fonctionnelle non linéaire.

- *décomposition en petites et grosses structures*

On définit un filtre $\mathcal{G}$ (gaussienne, fonction porte, ...) qui permet de définir un champ filtré $\bar{u}$:

$$\bar{u} = u * \mathcal{G}, \text{ et, on écrit } u = \bar{u} + u'.$$

- *convolution*

On applique le filtre $\mathcal{G}$ à l'équation (1.1) :

$$\mathcal{G} * \left(\frac{\partial u}{\partial t} + \mathcal{L}(u) + \mathcal{N}(u) \right) = 0.$$

- *justification, et couplage entre les termes de structures différentes*

On donne les hypothèses sous lesquelles on peut commuter les opérateurs de convolution et de dérivée temporelle ainsi que de dérivée spatiale ($\mathcal{L}$)

$$\mathcal{G} * \frac{\partial u}{\partial t} = \frac{\partial}{\partial t}(\mathcal{G} * u) \text{ et, } \mathcal{G} * \mathcal{L}(u) = \mathcal{L}(\mathcal{G} * u).$$

on obtient donc, $\frac{\partial \bar{u}}{\partial t} + \mathcal{L}(\bar{u}) + \overline{\mathcal{N}(u)} = 0$ que l'on écrit, par exemple,

$$\frac{\partial \bar{u}}{\partial t} + \mathcal{L}(\bar{u}) + \mathcal{N}(\bar{u}) + \mathcal{R}(\bar{u}, u') = 0.$$

- *fermeture algébrique du problème*

Les petites structures de l'écoulement n'étant pas résolues, il est nécessaire de connaître leur action sur les grosses structures de l'écoulement. Ainsi, on substitue le terme $\mathcal{R}(\bar{u}, u')$ par un terme dépendant uniquement de $\bar{u}$. Cette action s'appelle la modélisation :

$$\mathcal{R}(\bar{u}, u') = \mathcal{M}(\bar{u}).$$

1.2 Equations traitées

Il s'agit des équations de Navier-Stokes adimensionnées pour un fluide compressible, intervenant sur les grosses structures des différentes quantités, pour un

écoulement sans apport de forces extérieures. Le schéma numérique est donné au chapitre suivant.

$$\frac{\partial \bar{p}}{\partial t} + \sum_{j=1}^3 \frac{\partial \bar{\rho} u_j}{\partial x_j} = 0 \quad (1.2)$$

$$\frac{\partial}{\partial t}(\bar{\rho} u_i) + \sum_{j=1}^3 \frac{\partial}{\partial x_j} \left(\frac{\bar{\rho} u_i \cdot \bar{\rho} u_j}{\bar{\rho}} \right) - \sum_{j=1}^3 \frac{\partial \bar{\sigma}_{ij}}{\partial x_j} + \frac{\partial \bar{p}}{\partial x_i} = - \sum_{j=1}^3 \frac{\partial \tau_{ij}}{\partial x_j} \quad (1.3)$$

$$\begin{aligned} \frac{\partial \hat{E}}{\partial t} + \sum_{j=1}^3 \frac{\partial}{\partial x_j} \left((\hat{E} + \bar{p}) \frac{\bar{\rho} u_j}{\bar{\rho}} + \bar{Q}_j \right) - \sum_{i,j=1}^3 \frac{\partial}{\partial x_j} (\bar{\sigma}_{ij} u_i) = & - \sum_{i,j=1}^3 \frac{\partial}{\partial x_j} \left(\frac{\bar{\rho} u_i}{\bar{\rho}} \tau_{ij} \right) \\ & - \sum_{j=1}^3 \frac{\partial}{\partial x_j} (\pi_j + q_j) \end{aligned} \quad (1.4)$$

où $\hat{E}$ représente l'énergie totale des variables filtrées, et les trois tenseurs τ_{ij} , π_j , q_j proviennent du filtrage et de la non linéarité des termes convectifs. Nous détaillons les procédures d'adimensionnement et de filtrage dans l'annexe B.

Les quantités de sous-maille résultant de la non linéarité de σ_{ij} , Q_j et μ sont négligées (cf. [VGKZ92]).

$$\begin{aligned} \hat{E} &= \frac{\bar{p}}{\gamma - 1} + \frac{1}{2} \sum_{i=1}^3 \frac{\bar{\rho} u_i \cdot \bar{\rho} u_i}{\bar{\rho}}, \\ \bar{\sigma}_{ij} &= \frac{\mu}{Re} \bar{S}_{ij} = \frac{\mu}{Re} \left(\left(\frac{\partial \bar{u}_i}{\partial x_j} + \frac{\partial \bar{u}_j}{\partial x_i} \right) - \frac{2}{3} \delta_{ij} \sum_{k=1}^3 \frac{\partial \bar{u}_k}{\partial x_k} \right), \\ \bar{Q}_j &= - \frac{\mu}{(\gamma - 1) Re P_r M_r^2} \frac{\partial \bar{T}}{\partial x_j}, \end{aligned}$$

Re étant le nombre de Reynolds, M_r le nombre de Mach, P_r le nombre de Prandtl.

L'expression des trois termes τ_{ij} , π_j et q_j est la suivante:

$$\tau_{ij} = \bar{\rho} u_i u_j - \frac{\bar{\rho} u_i \cdot \bar{\rho} u_j}{\bar{\rho}}, \quad (1.5)$$

$$\pi_j = \bar{p} u_j - \bar{p} \cdot \frac{\bar{\rho} u_j}{\bar{\rho}}, \quad (1.6)$$

$$q_j = \frac{1}{(\gamma - 1) \gamma M_r^2} (\bar{\rho} T u_j - \frac{\bar{\rho} T \cdot \bar{\rho} u_j}{\bar{\rho}}). \quad (1.7)$$

Le tenseur τ_{ij} est modélisé par :

$$\tau_{ij} = -\bar{\rho} \nu_t \frac{\bar{\rho} S_{ij}}{\bar{\rho}} - \frac{1}{3} trace(\tau_{ij}),$$

et les deux autres termes π_j et q_j par la diffusion tourbillonnaire μ_t/P_{rt} .

1.3 Modèles de Fermeture

Nous avons testé deux modèles de fermeture, le modèle sous-maille de Smagorinsky, déjà mis en place par G. BAUDIER [BAU93] et le modèle sous-maille “Fonction de Structure”, proposé par Métais et Lesieur.

Modèle sous-maille de Smagorinsky

Nous rappellerons seulement, que la viscosité turbulente adimensionnée, ν_t^* , a pour expression :

$$\nu_t^* = \left(\frac{C_S}{N}\right)^2 \sqrt{\sum_{i=1}^3 \sum_{j=1}^3 \left(\frac{\partial \bar{u}_i}{\partial x_j}\right) \frac{\rho \bar{S}_{ij}}{\bar{\rho}}} \quad (1.8)$$

pour le modèle de Smagorinsky.

Modèle sous-maille Fonction de Structure

Il s’agit d’une modélisation proposée par Métais et Lesieur [ML91], dérivée de celle de Chollet-Lesieur [CL82], mais ayant la particularité de calculer la viscosité turbulente dans l’espace physique, contrairement au modèle de Chollet-Lesieur (C-L) qui la calcule dans l’espace spectral.

Dans (C-L), la viscosité turbulente ν_t (ici dimensionnée) est calculée en fonction de la densité spectrale d’énergie et du nombre d’onde de coupure k_c ,

$$\nu_t(k_c, t) = \frac{2}{3} C_K^{-\frac{3}{2}} \sqrt{\frac{E(k_c, t)}{k_c}}.$$

Kraichnan [LES73] définit la fonction de structure de la vitesse, dans l’espace physique, qui est une quantité locale, du second ordre, comme :

$$F_2(\vec{x}, \Delta, t) = \langle \|\vec{u}(\vec{x} + \vec{r}, t) - \vec{u}(\vec{x}, t)\|^2 \rangle_{\|\vec{r}\|=\Delta},$$

où, l’opérateur $\langle g(r) \rangle_{\|\vec{r}\|=\Delta}$ représente la moyenne de l’expression locale $g(r)$ sur la sphère de rayon Δ .

En utilisant la relation (1.9) ci-après, liant le spectre d’énergie et la fonction de structure pour une turbulence isotrope [McC92], et la loi de Kolmogorov en turbulence développée (1.10) (dans la zone inertielle),

$$F_2(r, t) = 4 \int_0^\infty E(k, t) \left[1 - \frac{\sin(kr)}{kr}\right] dk \quad (1.9)$$

$$E(k, t) = C_K \epsilon^{\frac{2}{3}} k^{-\frac{5}{3}} \quad (1.10)$$

on trouve [ML91] [NGML92],

$$F_2(r, t) = 4.82 C_K (\epsilon r)^{\frac{2}{3}}.$$

On supprime ensuite la dépendance en ϵ en évaluant (1.10) au nombre d'onde de coupure $k_c = \frac{\pi}{\Delta}$, donc :

$$\nu_t(\Delta, t) = 0.067 C_K^{-\frac{3}{2}} \Delta \sqrt{F_2(\Delta, t)}.$$

Après avoir tronqué la fonction de structure sur le champ des grosses structures et en utilisant (1.10) pour les structures de sous-maillages, on obtient, finalement [NGML92] :

$$\nu_t(\vec{x}, \Delta, t) = 0.104 C_K^{-\frac{3}{2}} \Delta \sqrt{\bar{F}_2(\vec{x}, \Delta, t)}$$

où, $\bar{F}_2(\vec{x}, \Delta, t) = \langle \|\vec{u}(\vec{x} + \vec{r}, t) - \vec{u}(\vec{x}, t)\|^2 \rangle_{\|\vec{r}\|=\Delta}$, est la fonction de structure filtrée, qui peut être évaluée à partir des équations filtrées (sur les grosses structures).

1.4 Simulations numériques

Avant de présenter les résultats obtenus, il est nécessaire de préciser la définition adoptée dans la suite pour le nombre de Reynolds Re . Nous détaillons la procédure d'initialisation en annexe C.

Ensuite, il nous faut aussi préciser à quel temps nous devons nous référer, et jusqu'où nous pouvons nous attendre à simuler dans les hypothèses de la LES.

Le critère de validation et le plan d'investigations précède les paragraphes récapitulant les résultats de chacun des modèles de fermeture.

Nombre de Reynolds

Puisque les écoulements considérés sont homogènes et isotropes, et ne possèdent donc pas de mouvement moyen, nous avons choisi de prendre comme vitesse caractéristique :

$$U_c = \left(\int_0^{+\infty} E(k, t) dk \right)^{1/2}$$

Cette vitesse, étroitement liée au spectre de densité d'énergie $E(k, t)$, nous amène alors à utiliser pour longueur caractéristique L_c , une grandeur liée elle aussi à ce spectre. Parmi les choix possibles, nous avons retenu l'expression suivante :

$$L_c = \frac{2\pi}{k_{peak}},$$

où k_{peak} représente le nombre d'onde pour lequel le spectre initial atteint sa valeur maximale. Le nombre de Reynolds s'écrit donc :

$$Re = \frac{U_c L_c}{\mu}$$

Les valeurs de k_{peak} , Re et μ étant fixées, on obtient donc une valeur pour U_c

$$\text{avec } l_{peak} = \frac{2\pi}{k_{peak}} \quad \text{et} \quad \bar{u} = c.M_r = \sqrt{\frac{\gamma p}{\rho}}.M_r$$

Puisque les équations traitées ne contiennent aucun terme source permettant d'entretenir la turbulence initiale, on doit s'attendre à observer une décroissance de l'énergie cinétique turbulente dans le temps, celle-ci étant progressivement transformée en énergie interne. On peut prédire dans ces conditions une diminution de la valeur du nombre de Reynolds global de l'écoulement au cours du temps. Cela signifie qu'un état turbulent initial compatible avec les hypothèses sous-jacentes à l'utilisation des modèles de fermeture étudiés (en particulier le fait que la "coupure" se situe dans la zone inertielle), peut conduire, après un certain temps, à un état incompatible avec ces mêmes hypothèses.

Temps significatif

D'après l'annexe C, l'état initial n'est pas un état THI car le spectre initial n'est pas en accord avec la théorie de Kolmogorov. En conséquence, si l'on suppose que l'on peut obtenir un spectre de Kolmogorov avec les modèles de fermeture utilisés, on peut penser que toutes les simulations présentent une phase transitoire correspondant au temps nécessaire au changement d'état.

Rien ne prouve que les durées de cette phase transitoire, et de la période pendant laquelle une simulation est cohérente avec les hypothèses de LES, sont indépendantes des conditions de simulation (on peut même penser au contraire que ces durées sont très liées aux valeurs du nombre de Reynolds).

Pour toutes ces raisons il est apparu nécessaire d'introduire un temps caractéristique pour, à la fois, permettre une comparaison entre des différentes simulations réalisées, et pouvoir rattacher par exemple la durée de la période de cohérence à quelque chose de significatif.

Le temps caractéristique que nous avons retenu ici est le "*Temps de retournement des grandes structures*" (*large eddy turnover time*, cf. [McC92]). Malheureusement plusieurs définitions de ce temps étant disponibles dans la littérature, nous avons donc été amené à effectuer un choix.

L'expression du temps de retournement adoptée est celle donnée dans [SAG], basée sur la valeur $L(t)$ de "*l'échelle intégrale*" et sur la "*vitesse moyenne*" de l'écoulement, et dont la définition est la suivante :

$$\Omega(t) = \frac{L(t)}{\sqrt{\langle \bar{u}^2 \rangle(t)}} \equiv \frac{L(t)}{\sqrt{\mathcal{E}(t)}} \quad (1.11)$$

où,

$$L(t) = \frac{3\pi}{4} \frac{\int_0^\infty k^{-1} E(k, t) dk}{\int_0^\infty E(k, t) dk}$$

et

$$\mathcal{E}(t) = \int_0^\infty E(k, t) dk$$

Critère de validation et objectifs

Le seul critère utilisé ici pour “valider” les simulations, est l’obtention d’une pente correcte pour le spectre (en accord avec la théorie de Kolmogorov). Une solution présentant un tel spectre est dite auto-semblable.

Il faut bien se rendre compte que l’une des difficultés principales des simulations par LES est de pouvoir assurer un transfert d’énergie des grandes échelles jusqu’aux plus petites simulées, en accord avec les expériences réalisées.

1.4.1 Résultats du modèle de Smagorinsky

Dans cette section, outre la présentation des résultats obtenus, nous essaierons de répondre à un certain nombre de questions.

Parmi celles-ci :

La constante de Smagorinsky C_s , est-elle une “vraie” constante, ou bien varie-t-elle, dans ce cas pourquoi et comment, en fonction des caractéristiques de l’écoulement simulé, et/ou des caractéristiques du schéma numérique adopté? En particulier, C_s dépend-elle du nombre de Reynolds et du nombre N de points de discrétisation par dimension d’espace?

Dans les conditions d’une turbulence non entretenue, peut-on garantir l’obtention d’un spectre de Kolmogorov jusqu’au temps maximum T_{max} pour lequel la simulation reste cohérente avec les hypothèses sous-jacentes aux modèles de fermeture?

Les écoulements simulés étant THI, rappelons que le domaine de calcul est constitué par un cube de dimension 3, maillé uniformément dans chaque direction d’espace. Les conditions aux limites sont périodiques, ce qui se programme naturellement sur un ordinateur massivement parallèle comme la CM-5. On peut trouver quelques détails sur l’implémentation dans [BAU93].

Rappelons l'expression de la viscosité turbulente utilisée dans le modèle de Smagorinsky :

$$\nu_t^* = C_s^2 \Delta^2 \sqrt{\sum_{i=1}^3 \sum_{j=1}^3 \left(\frac{\partial \overline{u_i}}{\partial x_j} \right) \frac{\overline{\rho S_{ij}}}{\overline{\rho}}}$$

La condition CFL choisie est de 0.6, elle est bien inférieure à celle correspondant à la limite théorique de stabilité linéaire de $2\sqrt{2}$, (cf. [HIR92]).

Aux figures suivantes (cf. Fig. 1.2), nous présentons l'allure du spectre de densité d'énergie cinétique turbulente $E(k, t)$ associé à la sphère de l'espace spectral de rayon $|k|$, à différents temps de simulation (à T_0 , temps initial où l'on initialise de manière aléatoire, ainsi qu'à $T = 0.05, 0.1, 0.2, \dots$ et ainsi de suite tous les 0.1 de temps de simulation). Si l'on considère que notre espace spectral est maillé de N modes dans chaque direction, le calcul du spectre s'effectue en associant à $E(K, t)$, la quantité :

$$4\pi K^2 \cdot \langle \hat{u}(k, t) \rangle_{\{k \in \{-N \dots N\}^3, |k|=K\}},$$

où l'expression $\langle g(k) \rangle_{F_k}$ désigne toujours la moyenne de $g(k)$ sur l'espace F_k . La méthode de calcul complète est donnée dans l'annexe C.

Nous utilisons une échelle logarithmique suivant les deux axes, ce qui nous donne un repère visuel de la pente (en $-\frac{5}{3}$) dans la zone inertielle d'une solution auto-semblable.

Nombre de Reynolds	1500
Discretisation	32^3
Fermeture	$C_s = 0,22$

Evolution de spectres d'énergie

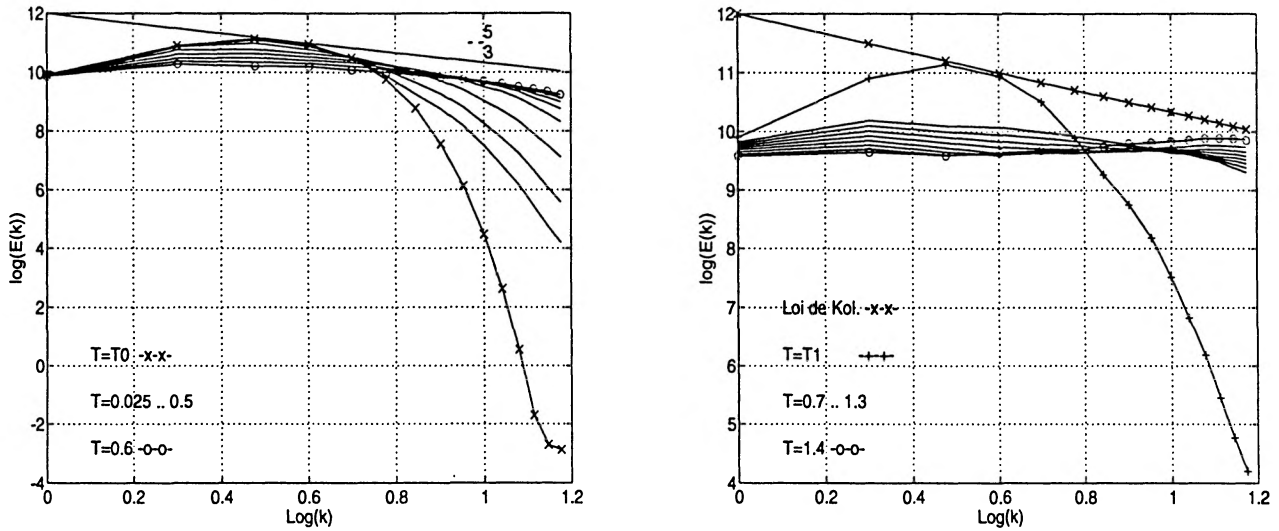
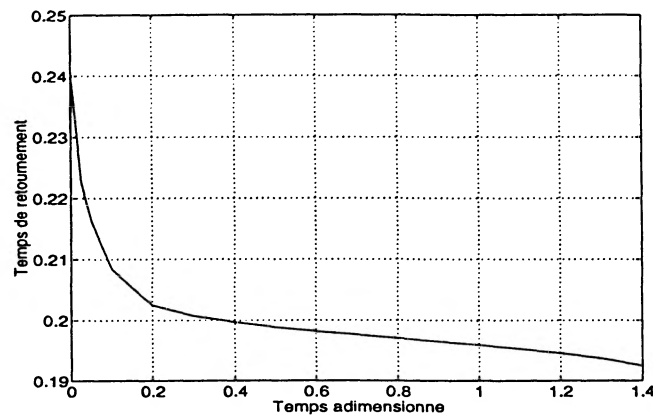


FIG. 1.2 - Simulation avec une condition CFL de 0,6

On notera que la valeur de 0.22 pour la constante de Smagorinsky, est en total accord avec les valeurs mentionnées dans la littérature ([McC92], [BER93], [SAG]).

La figure ci-après représente l'évolution temporelle du temps de retournement des grandes échelles défini par l'équation (1.11). On remarque, malgré une légère diminution de celle-ci, que la valeur de $\Omega(t)$ varie relativement peu dans le temps, ce qui nous conduit à adopter pour valeur caractéristique de ce temps la valeur 0.2.

Si l'on se base sur de simples considérations visuelles, on peut alors considérer *a priori*, que la simulation précédente fournit un spectre correct jusqu'à un temps de l'ordre de 0.7. En termes de temps de retournement, la simulation est donc *acceptable* jusqu'à environ 3.5 à 4 fois le temps de retournement moyen.

FIG. 1.3 - Temps de retournement $\Omega(t)$ des grandes échelles

Une simulation effectuée à une valeur du nombre de Reynolds double de la précédente, toutes les autres conditions demeurant inchangées, nous a conduit à corriger la valeur de la constante de Smagorinsky, pour obtenir une décroissance correcte du spectre. La valeur précédente ($C_s = 0.22$) conduisant à une trop forte dissipation (pente "plus raide" du spectre), des résultats *qualitativement* corrects nous ont imposé le choix $C_s = 0.17$ (cf. les figures ci-dessous).

Nombre de Reynolds	3000
Discretisation	32^3
Fermeture	$C_s = 0,17$

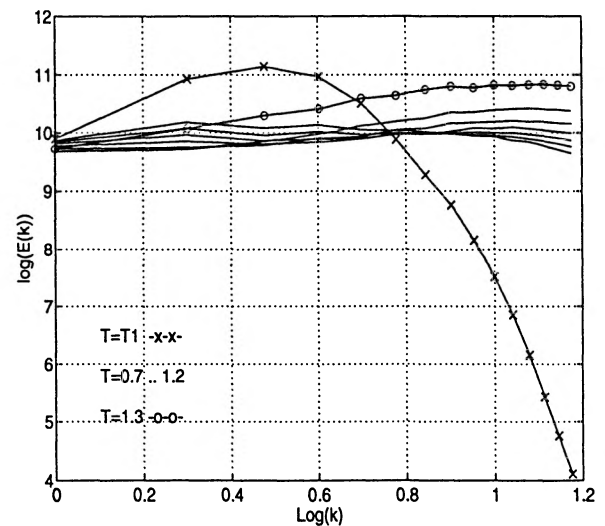
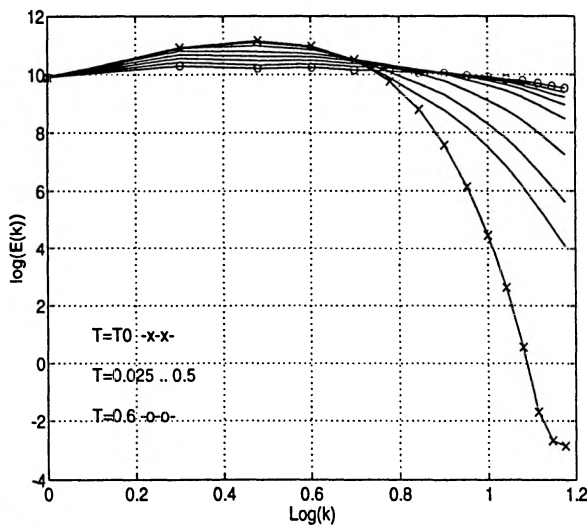


FIG. 1.4 - Evolution de spectres d'énergie (CFL de 0,6)

Ces deux premières simulations prouvent bien que la valeur de la *constante* de Smagorinsky dépend de la valeur du nombre de Reynolds Re .

Une explication est proposé en Annexe A.

Le problème principal de ces simulations (déjà rencontré dans [BAU93]), à savoir la remontée du spectre pour des temps de simulation *importants*, est ici bien visible (*cf.* 1.2 droite et 1.4 droite).

On observera d'ailleurs que la simulation à $Re = 3000$ donne, de ce point de vue, des résultats plutôt plus mauvais que la simulation à $Re = 1500$.

Indiquons enfin que pour ces deux simulations le pas de temps (adimensionné) était de l'ordre de 6.10^{-4} , et qu'une itération en temps correspondait à un temps de 0.3 seconde CPU.

Avec ces données, une simulation jusqu'au temps (adimensionné) 1 nécessite environ 500 secondes CPU pour un maillage 32^3 . Et logiquement, une simulation pour une discrétisation de 64^3 points, jusqu'à un même temps (de simulation) demande environ 16 fois plus de temps, soit 8000 secondes (le pas de temps étant divisé par 2 puisque le nombre de Courant reste fixé à 0.6).

Nombre de Reynolds	3000
Discrétisation	64^3
Fermeture	$C_s = 0,17$

Evolution de spectres d'énergie

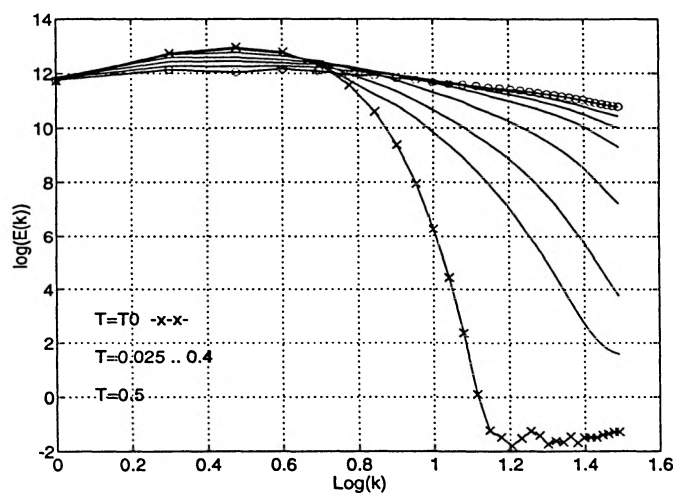


FIG. 1.5 - Simulation avec une condition CFL de 0,6

Nombre de Reynolds	1500
Discrétisation	64^3
Fermeture	$C_s = 0,22$

Evolution de spectres d'énergie

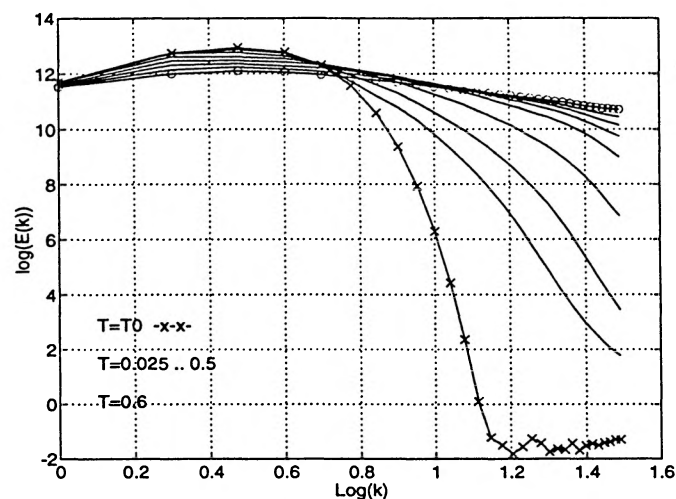


FIG. 1.6 - Simulation avec une condition CFL de 0,6

Le calcul du temps de retournement des grandes échelles pour la simulation considérée précédemment donne le résultat suivant :

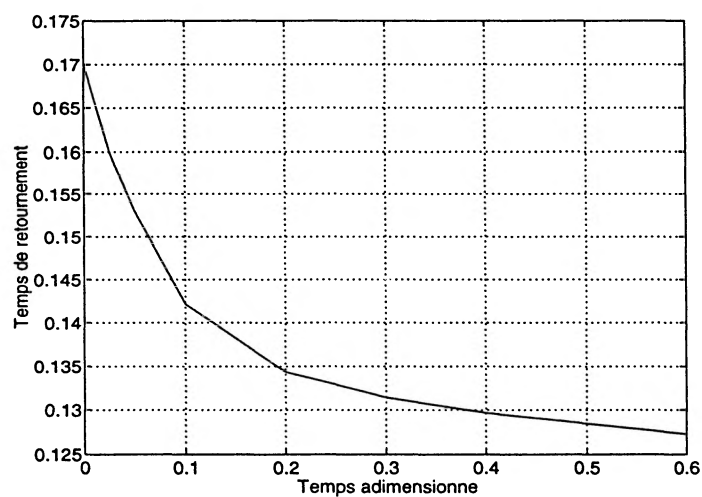


FIG. 1.7 - Temps de retournement des grandes échelles

On remarque que ce temps caractéristique évolue exactement de la même manière que celui de la simulation avec fermeture Smagorinsky pour une discrétisation de

32^3 points. La décroissance quasi-hyperbolique semble tout à fait logique et s'explique par le transfert d'énergie des petits nombres d'onde vers les grands nombres d'onde. En effet, à l'instant initial le spectre de densité d'énergie considéré a pour but de concentrer la majeure partie de l'énergie transmise au système sur les petits nombres d'ondes. Ensuite, la simulation va donner lieu à un échange d'énergie créé par le terme non linéaire de couplage de l'équation. En pratique, le taux de transfert sera élevé au début de la simulation (l'énergie étant initialement accumulée sur les petits nombres d'onde), et après quelques itérations temporelles le transfert va s'effectuer au taux prédit et donc va diminuer au fur et à mesure la taille des structures de l'écoulement. Par conséquent, les nouvelles grandes échelles se retournent plus rapidement que les grandes échelles de la phase initiale de l'écoulement.

1.4.2 Résultats du modèle "Fonction de Structure"

Reprenons les mêmes configurations de la section précédente pour l'étude comparative des deux modèles de fermeture.

A la Figure 1.8, nous obtenons le même comportement du spectre de densité d'énergie en utilisant le modèle de fermeture par la fonction de structure (d'ordre 2 de la vitesse). La pente reste correcte jusqu'au même temps adimensionné de simulation (0,7).

Nombre de Reynolds	1500
Discretisation	32^3
Fermeture	Fonct. de Struc.

Evolution des spectres d'énergie

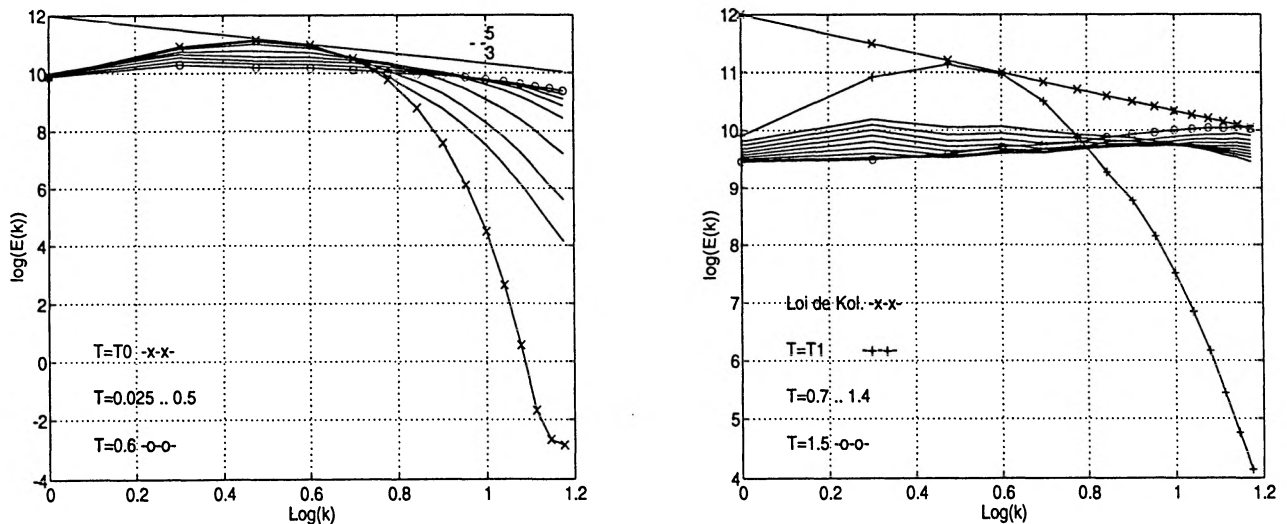


FIG. 1.8 - Simulation avec une condition CFL de 0,6

En augmentant le nombre de Reynolds, nous obtenons les résultats suivants (cf.

Figure 1.9):

Nombre de Reynolds	3000
Discrétisation	32^3
Fermeture	Fonct. de Struc.

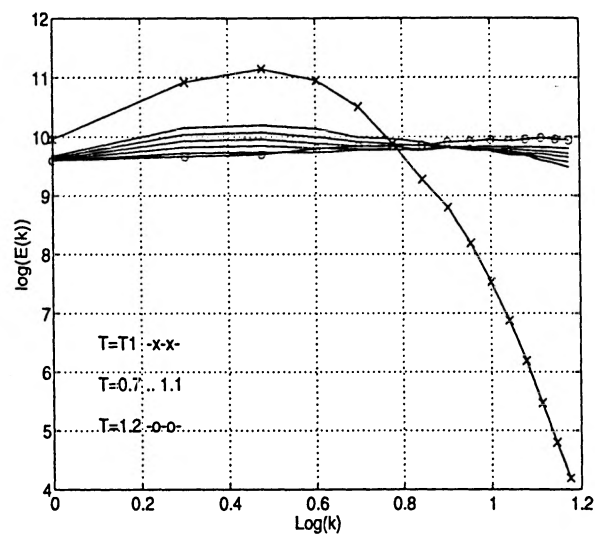
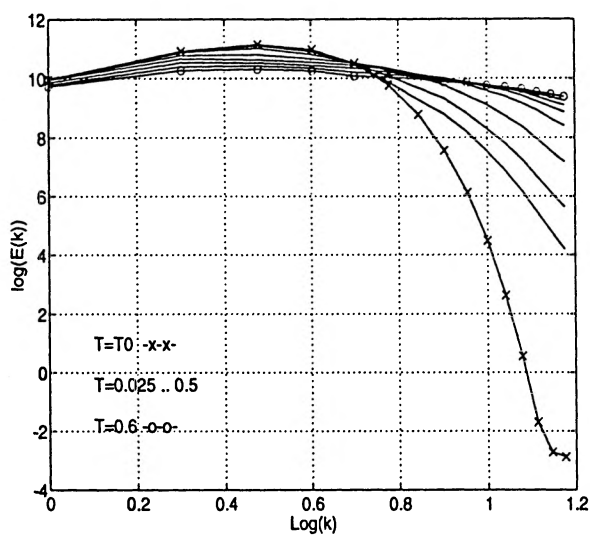
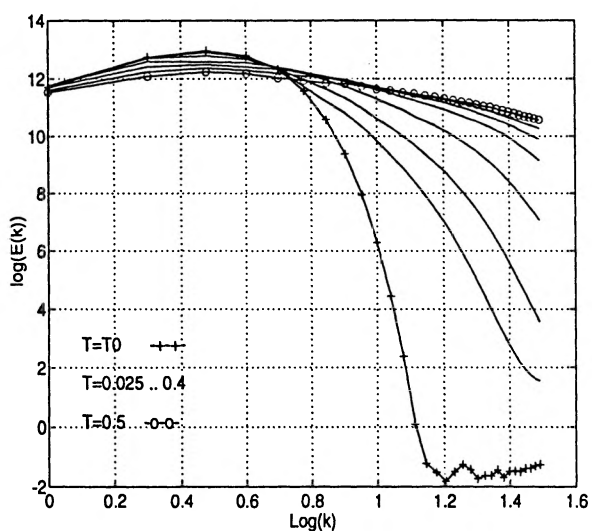


FIG. 1.9 - Evolution de spectres d'énergie (CFL de 0.6)

Nombre de Reynolds	1500
Discrétisation	64^3
Fermeture	Fonct. de Struc.



Nombre de Reynolds	3000
Discrétisation	64^3
Fermeture	Fonct. de Struc.

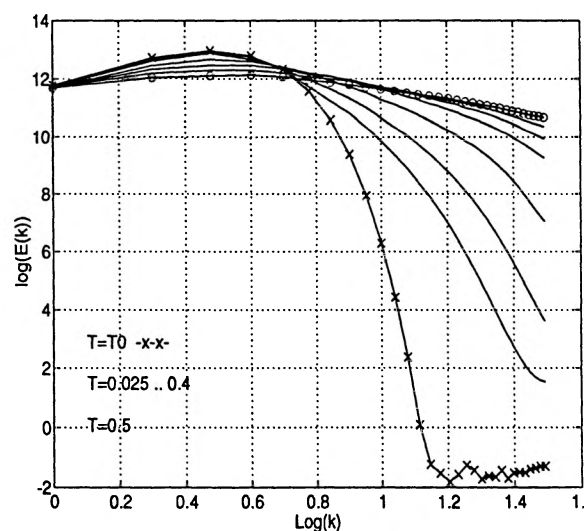


FIG. 1.10 - Evolution de spectres d'énergie (CFL de 0.6)

1.5 Conclusions

La comparaison objective de plusieurs simulations entre elles, qu'elles proviennent de modèles de fermeture différents ou d'une variation de conditions de simulation (Reynolds, N , ...) pour un même modèle, se heurte, comme nous allons le montrer, à de nombreuses difficultés.

Tout d'abord, nous devons faire un choix de configurations qui vont nous servir de plan d'expérimentations. Nous avons pensé faire varier le nombre de points de discrétisation et le nombre de Reynolds. Et cette dernière modification, nous mène à une ambiguïté du temps auquel la comparaison doit être effectuée (temps adimensionné ou temps de retournement et de quelles échelles).

De plus, un seul critère a été réellement pris en compte, il s'agit de la loi de Kolmogorov (1.10). Par conséquent, comment juger de la bonne pente dans une plage de l'espace spectral (zone inertielle) qui n'est pas définie explicitement?

D'après nos simulations, et le critère sur la pente du spectre de densité d'énergie, nous pouvons dire que les deux modèles de fermeture, globalement, donnent des résultats équivalents. Toutefois, la spécificité du modèle Fonction de Structure vient du fait qu'il ne demande pas de "réglage", contrairement à l'écriture de la viscosité turbulente de Smagorinsky, qui fait intervenir une *constante* dépendant de la simulation. Le temps calcul avec le modèle de Smagorinsky ou celui qui utilise la fonction de Structure reste le même. Et la programmation de l'un ou l'autre des modèles est tout aussi simple. Ce qui nous amène à conclure qu'aucun de ces deux modèles ne permet d'obtenir une solution auto-semblable avec un spectre correct. Cette étude nous a permis d'obtenir des renseignements en particulier sur le modèle de fermeture de Smagorinsky. Concernant l'étude comparative, nous avons vu que la Turbulence Homogène Isotrope n'est pas suffisamment restrictive pour permettre de départager ces deux modèles. Le chapitre suivant nous apportera des éclaircissements, et surtout nous permettra d'obtenir un lien étroit entre le schéma numérique utilisé et les résultats des simulations.

Chapitre 2

Etude du schéma numérique

L'objet de ce chapitre est l'étude de l'influence du schéma de discrétisation proposé par Vremann *et al.* (cf. [VGKZ92]) sur la simulation des grandes échelles. Pour cela, nous menons une étude théorique du schéma sur l'équation modèle suivante (convection-diffusion linéaire (2.1)) :

$$\frac{\partial u}{\partial t} + (v \cdot \nabla)u - r\Delta u = 0, \quad (2.1)$$

définie sur le domaine $\Omega \times [0, +\infty[$, u étant une fonction scalaire.

Après avoir présenté le schéma utilisé par Vremann *et al.*, nous introduisons à la deuxième section, la notion d'équation équivalente. Nous donnons ensuite un procédé permettant d'obtenir son expression à partir d'un développement en série formelle (cf. [BLC94]). Il semble alors que la solution calculée à partir du schéma numérique peut être bien approchée par l'équation équivalente tronquée à un ordre fixé (cf. [LEV90]).

Nous établissons, à la troisième section, une relation entre le spectre de densité d'énergie cinétique de la solution de cette équation équivalente et du spectre de la solution de l'équation modèle. Ce qui nous permet d'une part, de mettre en évidence un défaut d'isotropie du schéma de Vremann *et al.* à l'ordre 4, et d'autre part de donner une relation visant à corriger ce défaut.

La quatrième section fait état d'une nouvelle propriété liée au schéma cité ci-dessus. Nous voyons que le fait de moyenner dans les plans perpendiculaires à la direction de l'opérateur spatial induit une convolution discrète avec une gaussienne tronquée ayant donc un effet régularisant. Malheureusement, à la section suivante, nous obtenons une incompatibilité entre cette propriété de régularisation et l'isotropie à l'ordre 4.

Ensuite, nous rapprochons nos résultats de l'étude théorique de nos simulations numériques, ce qui nous permet d'élucider quelques points. Nous concluons et proposons quelques développements.

2.1 Schéma Numérique de Vremann *et al.*

Pour présenter le schéma numérique utilisé lors de nos simulations, nous formulons le problème (1.2-1.4) par le problème différentiel suivant (complété de bonnes conditions initiales).

$$\frac{\partial u(t)}{\partial t} = F(u(t)), \quad (2.2)$$

Nous décidons d'utiliser une méthode de type Runge-Kutta du quatrième ordre, semi-implicite (*cf.* [HIR92]) par exemple). Pour résoudre le problème différentiel (2.2),

$$\begin{aligned} \text{on écrit: } k_{n,1} &= F(u^n), \\ k_{n,2} &= F(u^n + \frac{\Delta t}{2} k_{n,1}), \\ k_{n,3} &= F(u^n + \frac{\Delta t}{2} k_{n,2}), \\ k_{n,4} &= F(u^n + \frac{\Delta t}{2} k_{n,3}), \\ u^{n+1} &= u^n + \frac{\Delta t}{6} (k_{n,1} + 2k_{n,2} + 2k_{n,3} + k_{n,4}). \end{aligned}$$

La discrétisation spatiale du terme $F(u(t))$ utilise deux opérateurs d'approximation notés D_σ et D_v . On désigne par $D_{\sigma,i}[u(\vec{M})]$ l'approximation de $\partial u / \partial x_i$ au point $\vec{M}$, et on pose :

$$D_{\sigma,i}[u(\vec{M})] = \frac{u(\vec{M} + \Delta x_i \vec{e}_i) - u(\vec{M} - \Delta x_i \vec{e}_i)}{2\Delta x_i},$$

où $\vec{e}_i$ représente le vecteur unitaire dans la direction i . D_σ réalise donc une approximation différence finie centrée d'ordre 2 de $\partial u / \partial x_i$. L'opérateur D_v a une expression plus complexe. Désignant par $D_{v,i}[u(\vec{M})]$ l'approximation de $\partial u / \partial x_i$ au point $\vec{M}$ et $\mathcal{E}$ l'ensemble $\{-1, 0, +1\} \times \{-1, 0, +1\}$ de toutes les valeurs possibles du couple (ϵ, η) , on pose :

$$D_{v,i}[u(\vec{M})] = \sum_{(\epsilon, \eta) \in \mathcal{E}} \phi_{(\epsilon, \eta)} D_{\sigma,i}[u(\vec{M} + \epsilon \vec{e}_2 + \eta \vec{e}_3)].$$

Par cette procédure nous obtenons une approximation centrée de $\partial u / \partial x_i$ au point $\vec{M}$, pondérée par les valeurs approchées de $\partial u / \partial x_i$ aux points $\vec{P}$ situés dans le plan passant par $\vec{M}$ et perpendiculaire au vecteur $\vec{e}_i$.

Plus précisément, et en prenant par exemple $i = 1$, on définit l'ensemble des voisins de $\vec{M}$ comme l'ensemble formé, en dimension 3, du point $\vec{M}$ lui même, et des 8

points $\vec{P}_n$ ($1 \leq n \leq 8$) les plus proches de $\vec{M}$ tels que :

$$\begin{cases} x_1(\vec{P}_n) = x_1(\vec{M}) & \forall n : 1 \leq n \leq 8 \\ x_2(\vec{P}_n) = x_2(\vec{M}) + \epsilon \Delta x_2 & \epsilon \in \{-1, 0, +1\} \\ x_3(\vec{P}_n) = x_3(\vec{M}) + \eta \Delta x_3 & \eta \in \{-1, 0, +1\} \end{cases}$$

(ϵ et η non simultanément nuls).

Un voisin de $\vec{M}$, y compris ce point, est donc caractérisé par un couple (ϵ, η) ($\vec{M}$ correspondant alors au couple $(0, 0)$).

Toute une série de schémas différents peuvent être construits suivant le choix effectué pour les coefficients de pondération ($\phi_{(-1,-1)} \cdots \phi_{(+1,+1)}$) utilisés.

Parmi ceux-ci nous avons adopté la solution suivante :

$$\begin{cases} (\epsilon, \eta) = (0, 0) & : \quad \phi_{(0,0)} = \frac{1}{2\Delta x} \frac{4}{9} \\ (\epsilon, \eta) = (0, \pm 1) & : \quad \phi_{(0,\pm 1)} = \frac{1}{2\Delta x} \frac{1}{9} \\ (\epsilon, \eta) = (\pm 1, 0) & : \quad \phi_{(\pm 1,0)} = \frac{1}{2\Delta x} \frac{1}{9} \\ (\epsilon, \eta) = (\pm 1, \pm 1) & : \quad \phi_{(\pm 1,\pm 1)} = \frac{1}{2\Delta x} \frac{1}{36} \end{cases}$$

Nous posons dans la suite :

$$\phi = \phi_{(0,0)} \quad ; \quad \alpha = \phi_{(0,\pm 1)} = \phi_{(\pm 1,0)} \quad ; \quad \beta = \phi_{(\pm 1,\pm 1)}$$

Vremann *et al.* ont montré (cf. [VGKZ92]) que ce choix était celui qui, parmi ceux qu'ils ont testés, apportait le moins de dissipation numérique.

Les termes convectifs et diffusifs sont alors approchés de la manière suivante :

– les termes convectifs de la forme :

$$\frac{\partial \overline{u_i}}{\partial x_j}(\vec{M})$$

sont approchés par :

$$D_{v,j}(\overline{u_i})(\vec{M})$$

– les termes diffusifs de la forme :

$$\frac{\partial \overline{S_{ij}}}{\partial x_j}(\vec{M}) = \frac{\partial}{\partial x_j} \left(\frac{\partial \overline{u_i}}{\partial x_j} + \frac{\partial \overline{u_j}}{\partial x_i} \right) (\vec{M})$$

sont approchés par :

$$D_{v,j}(D_{\sigma,j}(\bar{u}_i(\vec{M})) + D_{\sigma,i}(\bar{u}_j(\vec{M})))$$

Soit $\vec{M}$ le point de coordonnées :

$$x_1(\vec{M}) = x_1^i \quad ; \quad x_2(\vec{M}) = x_2^j \quad ; \quad x_3(\vec{M}) = x_3^k$$

et soit $u_{i,j,k}$ la valeur de u au point $\vec{M}$.

Avec les notations et définitions précédentes, un terme convectif tel que $(\partial u / \partial x)(\vec{M})$, est approché par l'expression :

$$\begin{aligned} \frac{\partial u}{\partial x}(\vec{M}) \rightarrow & \frac{1}{2\Delta x(4\alpha + 4\beta + \phi)} \left(\phi(u_{i+1,j,k} - u_{i-1,j,k}) \right. \\ & + \alpha \sum_{(\epsilon,\eta) \in \{-1,1\} \times \{-1,1\}} (u_{i+1,j+\epsilon,k+\eta} - u_{i-1,j+\epsilon,k+\eta}) \\ & \left. + \beta \sum_{(\epsilon,\eta) \in \mathcal{E} - \{-1,1\} \times \{-1,1\}} (u_{i+1,j+\epsilon,k+\eta} - u_{i-1,j+\epsilon,k+\eta}) \right). \end{aligned} \quad (2.3)$$

La Figure 2.2 représente la molécule de calcul associée à l'approximation d'un terme convectif. Dans le cas d'un terme diffusif, la molécule correspondante est symétrique par rapport à l'un quelconque des plans $[\alpha, \beta, \phi]$.

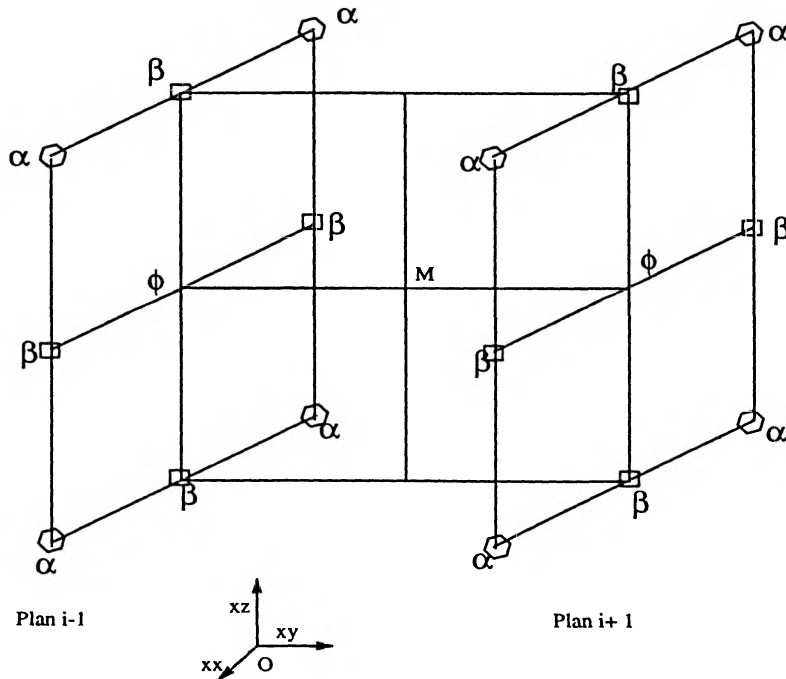


FIG. 2.2 - Molécule de calcul

2.2 Notion d'équation équivalente

Convenons de représenter l'équation (2.1) sous la forme $(T + A)u = 0$, où A est un opérateur différentiel linéaire en espace, et T représente l'opérateur de dérivation partielle par rapport à t . La discrétisation de cette équation à l'aide d'un schéma homogène en temps et en espace, consiste à remplacer celle-ci, sur un ensemble Ω_N de N points $M = (i, j, k) \in \Omega$, par une expression de la forme :

$$U_{i,j,k}^{n+1} \equiv U^{n+1}(M) = \sum_{n' \geq 0} \sum_{M' \in \mathcal{V}(M)} C_{M'-M}^{n'} U^{n-n'}(M') \quad (2.4)$$

où M' est le point de coordonnées $(i + i', j + j', k + k')$, i', j' et k' étant des entiers relatifs ; $\mathcal{V}(M)$ représente l'ensemble des "voisins" du point M , défini comme l'ensemble des points M' tels que $C_{M'-M}^{n'} \neq 0$; et les $C_{M'-M}^{n'}$ sont des nombres réels caractérisant le schéma utilisé.

La fonction discrète U qui vérifie la relation précédente en chacun des points M considérés, constitue la solution numérique du schéma représenté par cette relation. On définit alors une fonction $\tilde{U}$, "suffisamment régulière" (théoriquement de classe $\mathcal{C}^\infty$ mais, en pratique, de classe $\mathcal{C}^r$ avec r suffisamment grand), définie sur $\Omega \times [0, +\infty[$, et telle que $\tilde{U}$ et U coïncident sur Ω_N .

On pourra par exemple se reporter à ([VGKZ92]) pour les détails techniques de cette opération.

Les contraintes de régularité imposées à $\tilde{U}$, nous autorisent alors à développer les termes $U^{n-n'}(M' - M) = \tilde{U}(M' - M, t_{n-n'})$ en série de Taylor au voisinage du point (M, t_n) .

Convenons à cet effet de désigner simplement par T, X, Y et Z les opérateurs de dérivation partielle par rapport à t, x, y et z respectivement.

Avec ces symboles, on associe alors à l'équation discrète (2.4) une équation continue, d'ordre infini, que l'on écrit sous la forme suivante :

$$\left[\sum_{p,q,r,s} \mu_{p,q,r}^s T^s X^p Y^q Z^r \right] \tilde{U} = 0 \quad (2.5)$$

Cette équation représente l'équation équivalente **non résolue**, associée au schéma (2.4).

Remarquons que dans le cas d'un schéma explicite, celle-ci admet la forme simplifiée suivante :

$$\left[\sum_s \alpha^s T^s \right] \tilde{U} + \left[\sum_{p,q,r} \beta_{p,q,r} X^p Y^q Z^r \right] \tilde{U} = 0 \quad (2.6)$$

Ces équations étant de structure extrêmement complexe, en particulier à cause de la présence des opérateurs T^s pour, a priori toute valeur entière de s , il est quasiment impossible d'étudier le comportement de leur solution $\tilde{U}$.

Pour contourner cette difficulté, on va essayer, par un ensemble de règles de réécriture spécifiques, de ramener l'équation équivalente non résolue à une équation d'évolution "avec reste".

Plus précisément, on cherche à construire une équation de la forme :

$$T\tilde{U} + \left[\sum_{\substack{p,q,r \\ p+q+r \leq S}} \gamma_{p,q,r} X^p Y^q Z^r \right] \tilde{U} + \left[\sum_{\substack{p,q,r,s \\ p+q+r > S}} \tilde{\mu}_{p,q,r}^s T^s X^p Y^q Z^r \right] \tilde{U} = 0, \quad (2.7)$$

dans laquelle le second crochet représente un *reste* en $\mathcal{O}(S+1)$. Cette équation sera appelée "équation équivalente résolue à l'ordre S ".

Nous n'évoquerons pas ici les notions de consistance de cette équation avec l'équation $(T+A)u=0$ (bien qu'il soit immédiat de voir qu'une condition nécessaire est que $\mu_{0,0,0}^1 = 1$) et supposerons donc que cette propriété est vérifiée.

On voit donc que les règles de réécriture à utiliser, reviennent à éliminer successivement toutes les dérivées en temps d'ordre supérieur à 1 contenues dans l'équation équivalente non résolue, et ceci jusqu'à l'ordre (ici S) souhaité.

C'est en 1974 que Warming et Hyett (cf. [WH74]) ont proposé une méthode permettant de mener à bien cette élimination. Celle-ci est en fait un procédé itératif, qui génère une suite d'équations équivalentes partiellement résolues, d'ordre de plus en plus élevé.

$$T U_S + \left[\sum_{\substack{p,q,r \\ p+q+r \leq S}} \gamma_{p,q,r} X^p Y^q Z^r \right] U_S = 0 \quad (2.8)$$

Si le schéma (2.4) est consistant avec l'équation $(T+A)u=0$, alors l'équation précédente (dans laquelle on suppose implicitement que S est au moins égal à l'ordre $\mathcal{O}(A)$ de l'opérateur A) s'écrit aussi :

$$(T+A)U_S + B_S U_S = 0 \quad (2.9)$$

où B_S est un opérateur différentiel linéaire en espace, et dans lequel ne figurent que des termes $X^p Y^q Z^r$ tels que $\mathcal{O}(A) < p+q+r \leq S$.

L'intérêt de cette construction, réside dans la considération heuristique établie par Warming et Hyett, selon laquelle la solution continue $\tilde{U}$ de l'équation équivalente non résolue est aussi solution, en chaque point de discrétisation, de l'équation discrète (2.4).

Moyennant alors certaines conditions de convergence du procédé développé par ces auteurs, on peut espérer que les solutions U_S de (2.9) convergent, lorsque S tend vers l'infini, vers celle de (2.4) :

$$\lim_{S \rightarrow \infty} \|U_S - U\| = 0 \quad (2.10)$$

pour une norme $\|\cdot\|$ convenable.

Il semble bien que la fonction $\tilde{U}$ demeure “plus proche” de U que ne l’est la solution u de l’équation à résoudre (2.1), et les fonctions U_S jouissent de la même propriété :

$$\forall S \geq \mathcal{O}(A) : \lim_{S \rightarrow \infty} \|U_S - U\| \leq \|u - U\| \quad (2.11)$$

On comprend alors mieux pourquoi, malgré les objections qui ont pu lui être faite, l’équation équivalente (en pratique résolue jusqu’à un ordre S arbitraire), constitue un outil privilégié d’étude d’un schéma numérique.

La méthode proposée par Warming et Hyett souffre cependant d’être extrêmement lourde à mettre en œuvre, tout particulièrement dans le cas d’équations en dimension d’espace supérieure à 1, et de mal se généraliser au cas de schémas à pas en temps fractionnaires.

A tel point que son utilisation est quasiment tributaire de celle d’un logiciel de calcul formel, et que la construction à *la main* de l’équation équivalente est souvent une gageure.

Fort heureusement, Carpentier, De La Bourdonnaye et Larrouturou (*cf.* [BLC94]) ont proposé récemment une nouvelle procédure de construction de l’équation modifiée (ou équivalente), beaucoup plus simple à mettre en œuvre que celle de Warming et Hyett, et parfaitement adaptée au traitement de schémas de type Runge-Kutta.

Avec les notations introduites plus haut, considérons l’équation d’évolution linéaire suivante :

$$Tu = -Au = \left[\sum_{(p,q,r)} a_{p,q,r} X^p Y^q Z^r \right] u \quad (2.12)$$

Supposons pour simplifier que les pas du maillage Ω_N sont identiques dans les 3 directions d’espaces, et désignons par h leur valeur commune.

Supposons enfin que l’équation précédente est résolue à l’aide d’un schéma de Runge-Kutta à M niveaux, et que l’opérateur $-A$ est approché par un même opérateur discret $-A_h$ quel que soit le pas fractionnaire considéré (schéma homogène en temps).

A l’aide des résultats partiels démontrés dans [BLC94] : expressions de l’équation équivalente résolue pour un schéma explicite en temps en dimension 2 d’espace, et pour un schéma de Runge-Kutta en dimension 1 d’espace ; il est possible de déterminer l’expression de l’équation équivalente pour le schéma utilisé dans nos simulations (cas tridimensionnel, schéma de Runge-Kutta d’ordre 4). La proposition suivante en fait état :

Proposition 2.2.1

Si le schéma est consistant avec l'équation (2.1), son équation équivalente s'écrit :

$$T = \sum_{(i,j,k)} \alpha_{i,j,k}(\Delta t, h) X^i Y^j Z^k,$$

où $\sum_{(i,j,k)} \alpha_{i,j,k}(\Delta t, h) X^i Y^j Z^k$ est le développement en série formelle de la fonction :

$$\mathcal{F}(X, Y, Z) = \frac{\ln \left(1 + \sum_{K=1}^{K=M} \frac{[\Delta t A_h(X, Y, Z)]^K}{K!} \right)}{\Delta t}$$

2.3 Anisotropie théorique du schéma

Avant d'entreprendre l'étude proprement dite de ce schéma, commençons par établir un résultat intermédiaire concernant les spectres de densité d'énergie des solutions u et U_S des équations $(T + A)u = 0$ et (2.9).

Limitons nous dans un premier temps au cas d'une équation d'évolution linéaire, en dimension 1 d'espace :

$$Tu + \left[\sum_{m=1}^{m=M} a_m X^m \right] u = 0 \quad (2.13)$$

avec des conditions aux limites périodiques.

Pour toute fonction $f \equiv f(x, t)$, désignons, si elle existe, sa transformée de Fourier spatiale $\hat{f} \equiv \hat{f}(\xi, t)$ par la relation :

$$\hat{f}(\xi, t) = \int_{\mathbf{R}} f(x, t) e^{-ix\xi} dx$$

Moyennant des conditions convenables sur u , on a alors :

$$\widehat{X^m u}(\xi, t) = \int_{\mathbf{R}} X^m u(x, t) e^{-ix\xi} dx = (i\xi)^m \hat{u}(\xi, t)$$

$$\widehat{Tu}(\xi, t) = \frac{\partial}{\partial t} \left(\int_{\mathbf{R}} u(x, t) e^{-ix\xi} dx \right) = T\hat{u}(\xi, t).$$

Dans l'espace de Fourier, l'équation précédente s'écrit donc :

$$T\hat{u}(\xi, t) + \left[\sum_{m=1}^{m=M} a_m (i\xi)^m \right] \hat{u}(\xi, t) = 0, \quad (2.14)$$

équation différentielle ordinaire, dont la solution est :

$$\hat{u}(\xi, t) = \hat{u}(\xi, 0)e^{-P(\xi)t},$$

où P est le polynôme à valeurs complexes de la variable ξ égal à :

$$P(\xi) = \sum_{m=1}^{m=M} a_m (i\xi)^m.$$

Le spectre de densité d'énergie de u , ou plus simplement "*spectre*" de u , défini par $E_u = \hat{u}\bar{\hat{u}}$ ($\bar{\hat{u}}$ désignant le conjugué complexe de $\hat{u}$), a donc pour expression :

$$E_u(\xi, t) = E_u(\xi, 0)e^{-2\Re[P(\xi)]t}. \quad (2.15)$$

Cette expression montre que seuls les opérateurs de dérivation spatiale d'ordre pair interviennent dans l'expression du spectre de la solution de l'équation (2.13).

Ce résultat se généralise immédiatement en dimension d'espace supérieure. En effet, en posant :

$$\widehat{X_i u}(\xi_i, \dots, \xi_n, t) = i\xi_i \hat{u}(\xi_1, \dots, \xi_n, t),$$

on voit que l'on a la transformation suivante :

$$\left[X_1^{p_1} \dots X_n^{p_n} \right] u(x_1, \dots, x_n, t) \longrightarrow i^{p_1 + \dots + p_n} \xi_1^{p_1} \dots \xi_n^{p_n} \hat{u}(\xi_1, \dots, \xi_n, t)$$

Ainsi, seuls les opérateurs de degré total pair (égal ici à $p_1 + \dots + p_n$) interviennent dans l'expression du spectre de la solution.

Désignons par $P[A](\xi)$ le polynôme à valeurs complexes obtenu en appliquant les règles de transformation ci-dessus à l'opérateur A de l'équation $(T + A)u = 0$. L'expression (2.9) nous montre alors que le spectre de U_S s'écrit :

$$E_{U_S}(\xi, t) = E_{U_S}(\xi, 0)e^{-2\Re(P[A](\xi) + P[B_S](\xi))t} \quad (2.16)$$

Si les conditions initiales imposées à l'équation (2.9) et à l'équation $(T + A)u = 0$ sont identiques (ce qu'il semble naturel de considérer), les spectres de u et de U_S sont reliés par la relation :

$$E_{U_S}(\xi, t) = E_u(\xi, t)e^{-2\Re(P[B_S](\xi))t}. \quad (2.17)$$

Cette égalité établit donc la relation directe entre, le schéma numérique utilisé (2.4), et la perturbation qu'il induit sur le spectre de la solution exacte E_u .

Munis de ces résultats, et à l'aide de la proposition établie au paragraphe précédent, nous avons déterminé l'expression de l'équation équivalente à l'ordre 4, pour le schéma de Vremann appliqué au cas de l'équation scalaire d'advection-diffusion linéaire en dimension 3 d'espace (équation (2.1)).

Les calculs complets sont donnés dans l'Annexe D.

On a supposé pour ces calculs que le pas de maillage était constant et égal à h dans les 3 directions d'espace.

Il est commode, pour analyser ces résultats, de représenter le polynôme formel $A + B_S$ d'indéterminées X, Y et Z sous la forme :

$$(A + B_S)(X, Y, Z) = \sum_{n=0}^{n=S} Q_n(X, Y, Z)$$

où les Q_n sont des polynômes formels d'indéterminées X, Y et Z , homogènes de degré total égal à n .

Dans le cas considéré, on a :

$$\begin{aligned} Q_0(X, Y, Z) &= 0 \\ Q_1(X, Y, Z) &= v(X + Y + Z) \\ Q_2(X, Y, Z) &= -r(X^2 + Y^2 + Z^2) \end{aligned}$$

ce qui permet de vérifier a posteriori que le schéma considéré est consistant et d'ordre 2.

Le polynôme Q_4 (le premier intervenant dans la perturbation sur le spectre) a pour expression :

$$Q_4(X, Y, Z) = \frac{rh^2}{12} \left\{ \left[X^4 + Y^4 + Z^4 \right] + 24(a + 2b) \left[X^2Y^2 + Y^2Z^2 + Z^2X^2 \right] \right\}$$

Ce polynôme, dont on vérifie immédiatement l'invariance par permutation des indéterminées X, Y, Z , admet pour image, par transformation de Fourier spatiale, la fonction polynômiale $\widehat{Q}_4(\xi_X, \xi_Y, \xi_Z)$ suivante :

$$\widehat{Q}_4(\xi_X, \xi_Y, \xi_Z) = \frac{rh^2}{12} \left\{ \left[\xi_X^4 + \xi_Y^4 + \xi_Z^4 \right] + 24(a + 2b) \left[\xi_X^2 \xi_Y^2 + \xi_Y^2 \xi_Z^2 + \xi_Z^2 \xi_X^2 \right] \right\} \quad (2.18)$$

La fonction ("polynôme" dans la suite) $\widehat{Q}_4$ n'est, en général, pas isotrope au sens où elle n'est pas invariante par rotation arbitraire du repère ξ_X, ξ_Y, ξ_Z .

En effet, si tel était le cas, ce polynôme devrait admettre (à une constante multiplicative près) la représentation suivante :

$$\widehat{Q}_4(\xi_X, \xi_Y, \xi_Z) = \left[\xi_X^2 + \xi_Y^2 + \xi_Z^2 \right]^2 \quad (2.19)$$

Ceci n'est possible que si la relation suivante est vérifiée :

$$a + 2b = \frac{1}{3} \quad (2.20)$$

Compte tenu des valeurs de a et b pour le schéma de Vremann ($a = 1/9$, $b = 1/36$) on en déduit que ce schéma n'est pas isotrope à l'ordre 4 au sens précédent. Cette conclusion est a priori quelque peu gênante, dans la mesure où les écoulements simulés sont théoriquement isotropes, et que ce schéma est anisotrope. Par soucis de cohérence, il semble plus correct de corriger le schéma de Vremann pour le rendre isotrope à l'ordre 4.

La relation (2.20) montre cependant qu'il existe une infinité de tels schémas. Aussi, une modification du schéma de Vremann nécessite t-elle une étude plus approfondie de celui-ci afin de déterminer, parmi cette infinité de schémas, si certains d'entre eux ne sont pas "*plus proches*" du schéma de Vremann que d'autres ...

2.4 Propriété régularisante du schéma

La caractéristique essentielle de ce schéma réside dans l'approximation utilisée pour les dérivées spatiales. Rappelons en effet que dans ce schéma, la dérivée d'ordre 1 d'une fonction F quelconque est approchée par une formule de différences finies centrées, portant non pas sur F elle même, mais sur une fonction G déduite de F par "*moyenne pondérée*" dans le plan perpendiculaire à la direction de dérivation. Plus concrètement, si $\Pi(i_1; 1)$ désigne le plan perpendiculaire à la direction 1 et coupant celle-ci au point d'abscisse i_1 , la dérivée au point $M = (i_1, i_2, i_3)$ de la fonction F dans la direction 1 est approchée par la formule :

$$(X_1 F)(M) \sim \frac{G(i_1 + 1, i_2, i_3) - G(i_1 - 1, i_2, i_3)}{2h},$$

où

$$G(i_1 \pm 1, i_2, i_3) = \sum_{(i'_2, i'_3) \in \Pi(i_1 \pm 1; 1)} c_{i'_2 - i_2, i'_3 - i_3} F(i_1, i'_2, i'_3).$$

Puisque dans le schéma de Vremann les coefficients $c_{p,q}$ vérifient les relations :

$$c_{p,q} = c_{-p,q} = c_{-p,-q} = c_{p,-q} \quad \text{et} \quad c_{p,q} = c_{q,p},$$

la relation introduisant $G(i_1 \pm 1, i_2, i_3)$ s'écrit aussi :

$$G(i_1 \pm 1, i_2, i_3) = \sum_{(i'_2, i'_3) \in \Pi(i_1 \pm 1; 1)} c_{i_2 - i'_2, i_3 - i'_3} F(i_1, i'_2, i'_3) \quad (2.21)$$

Sous cette forme G apparait comme le produit de convolution discret de F par une fonction discrète C , dont les valeurs aux points de coordonnées $(i_2 - i'_2, i_3 - i'_3)$ sont égales à $c_{i_2 - i'_2, i_3 - i'_3}$.

Cherchons, pour le schéma considéré, une fonction $\Phi(x_2, x_3)$ vérifiant ces conditions. C étant invariante par symétrie et permutation de (x_2, x_3) , et de support limité aux 8 voisins les plus proches du point $M^\pm = (i_1 \pm 1, i_2, i_3)$ et de ce point lui-même, il est immédiat de voir que C “ne vérifie” en fait que 3 conditions :

$$c_{0,0} = c \quad ; \quad c_{0,\pm h} = c_{\pm h,0} = a \quad ; \quad c_{\pm h,\pm h} = b$$

Avec un peu d'intuition, on trouve alors que la restriction au pavé $[-h, h] \times [-h, h]$ d'une fonction répondant à ces contraintes, est la fonction **radiale** Φ définie par :

$$\Phi(r) = c \exp\left(-\frac{r^2}{\sigma^2}\right) \quad (2.22)$$

et pour le schéma de Vremann, $\sigma^2 = \frac{h^2}{4 \log 2}$.

Ce schéma consiste donc à effectuer, un filtrage bidimensionnel des différents champs turbulents à l'aide d'une gaussienne d'écart-type $\sigma = 0.601 h$ environ.

Les auteurs qui l'ont introduit, propose l'analyse suivante : ils quantifient la diffusion apporté par ce schéma en mesurant l'effet de l'opérateur d'approximation spatiale sur une harmonique simple.

Notre nouvelle interprétation suffit peut être à expliquer, grâce à l'effet régularisant bien connu de la convolution par une gaussienne, le relativement bon comportement de ce schéma vis à vis des problèmes d'explosion pour des temps importants.

2.5 Schéma régularisant et schéma 4-isotrope

Comme nous savons qu'il existe une infinité de schémas “4-isotropes” (isotropes à l'ordre 4, cf. section 2.3), recherchons maintenant quels sont parmi ceux-ci ceux qui jouissent de la “*propriété de régularisation*” du schéma de Vremann. De tels schémas seront dits “régularisants, 4-isotropes”.

Soit Φ , une gaussienne centrée, ces schémas vérifient les 5 conditions suivantes :

$$\begin{aligned} c + 4a + 4b &= 1 \\ a + 2b &= \frac{1}{3} \\ \Phi(0) &= c \\ \Phi(h) &= a \\ \Phi(h\sqrt{2}) &= b \end{aligned}$$

La première de ces 5 conditions correspond au fait que la moyenne utilisée est barycentrique, la seconde liée à la 4-isotropie, et les 3 dernières au caractère régularisant du schéma.

Un calcul simple montre alors que seules 3 conditions parmi les 5 imposées sont indépendantes, et qu'il existe **un seul schéma** qui les vérifie. Celui-ci correspond au choix des paramètres suivant :

$$a = b = c = \frac{1}{9} \quad (2.23)$$

On observe alors immédiatement que la fonction Φ est en fait la fonction caractéristique du pavé $[-h, h] \times [-h, h]$ (au coefficient $1/9$ près), et que la convolution de la relation (2.21) ne fait pas réellement intervenir une gaussienne. Ce résultat, qui peut sembler en contradiction avec ce qu'il vient d'être dit plus haut, provient simplement du fait qu'avec le choix (2.23), l'écart-type de la "gaussienne" Φ est infini !

On peut conclure en prétendant qu'il n'existe pas de schéma **réellement** régularisant qui soit aussi 4-isotrope sur le pavé $[-h, h] \times [-h, h]$.

Explicitement, les deux conditions provenant de l'aspect régularisant s'écrivent aussi :

$$\frac{h^2}{\sigma^2} = \text{Log} \frac{c}{a} \quad ; \quad \left(\frac{c}{a} \right)^2 = \frac{c}{b} \quad (2.24)$$

Avec la condition $c + 4a + 4b = 1$, il est alors facile de construire toute une classe de schémas "*plus ou moins*" régularisants, ne dépendant que d'un seul paramètre $\lambda = h/\sigma$.

Donnons 3 exemples de tels schémas :

$$\begin{array}{lll} c = \frac{100}{144} & a = \frac{10}{144} & b = \frac{1}{144} \\ c = \frac{16}{36} & a = \frac{4}{36} & b = \frac{1}{36} \\ c = \frac{4}{16} & a = \frac{2}{16} & b = \frac{1}{16} \end{array}$$

Le second exemple correspond, comme cela se vérifie immédiatement, au schéma de Vremann. Ces 3 schémas sont donnés par ordre décroissant de valeur de λ (respectivement $\sqrt{\text{Log } 10}$, $\sqrt{\text{Log } 4}$, $\sqrt{\text{Log } 2}$) et sont donc "*de plus en plus*" régularisants.

Remarque 1 Cette dernière affirmation doit être nuancée, car plus λ diminue (σ augmente donc !), et plus la fonction discrète définie sur $[-h, h] \times [-h, h]$, s'éloigne d'une véritable gaussienne et l'effet de ce support borné devient de plus en plus important (la véritable fonction est en fait le produit de la gaussienne considérée par la fonction caractéristique de $[-h, h] \times [-h, h]$).

2.6 Comparaison théorie-simulations

Sur la base d'une étude linéaire, nous avons démontré dans la section 2.3, que le schéma de Vremann n'était pas 4-isotrope : il est intéressant de regarder si ce "défaut" d'isotropie se retrouve dans les simulations que nous avons réalisées.

Pour répondre à cette question, nous nous sommes limités à l'étude de la répartition angulaire du spectre numérique pour le seul nombre d'onde $K = 10$ (nombre d'onde qui devrait se situer au moins un instant dans la zone inertielle, la coupure étant faite à $K = 16$).

Bien que l'anisotropie du schéma de Vremann soit indépendante de la valeur du nombre d'onde (cf. section 2.3), il peut en être autrement de l'anisotropie "numérique" à laquelle nous nous intéressons maintenant. De ce point de vue, l'étude de celle-ci pour le seul nombre d'onde $K = 10$ est bien évidemment restrictive.

Cependant, comme nous allons le voir, elle demeure suffisante pour tirer quelques conclusions fort intéressantes.

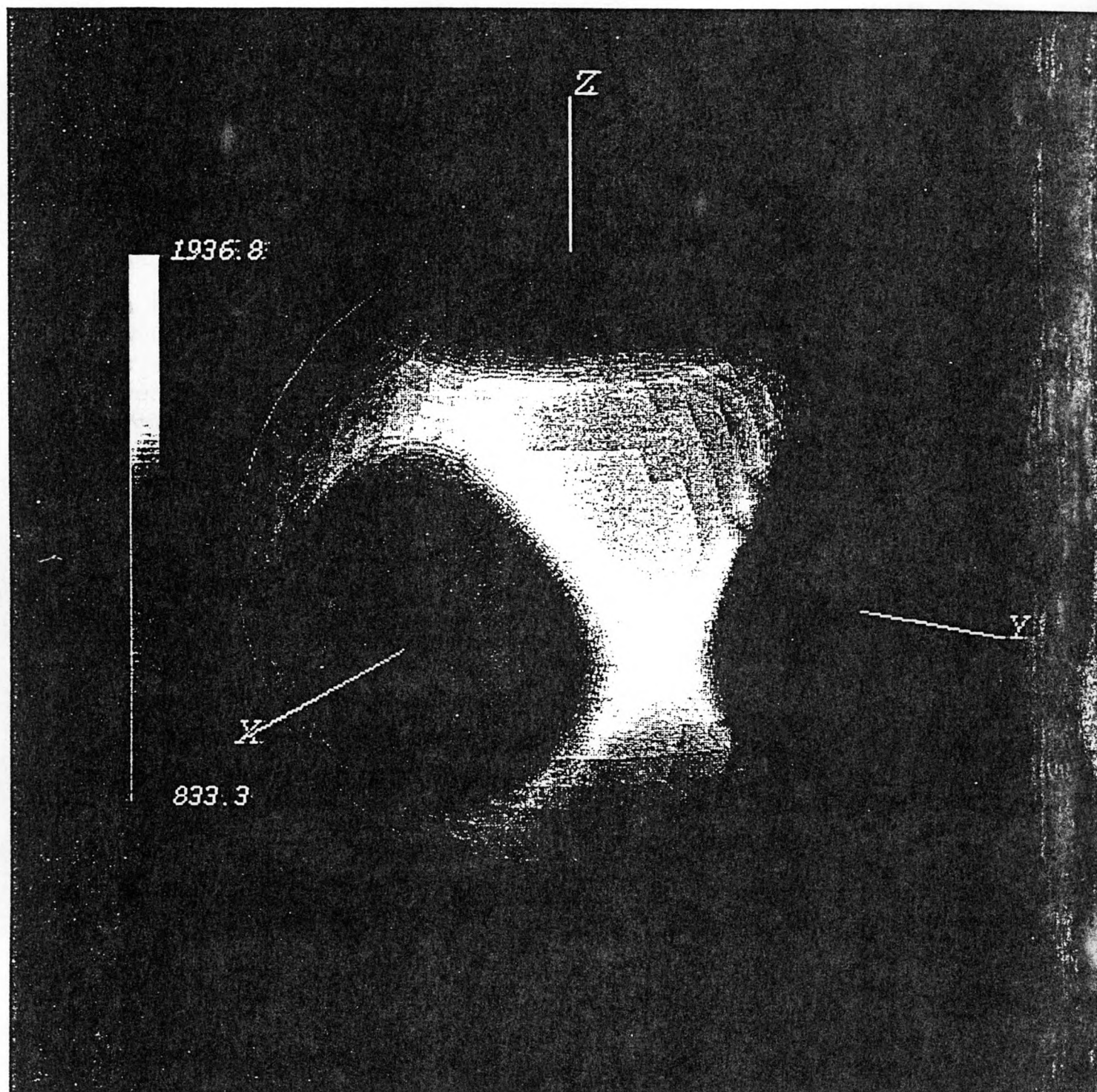
Afin d'avoir une carte d'anisotropie numérique suffisamment détaillée pour permettre une comparaison avec la carte d'anisotropie théorique, nous avons procédé comme suit :

- Construction, par Modulef¹, d'une sphère de rayon 10, pavée en triangles.
- Pour chacun des noeuds de cette triangulation, recherche, dans le cube de côté N défini par $[-N/2, N/2 - 1]^3$, du cube de côté unité contenant le noeud considéré.
- A l'intérieur de ce cube, aux sommets duquel est connue la valeur du spectre, approximation nodale du spectre pour en déterminer une valeur au noeud de la triangulation considéré.
- Coloration de la sphère de rayon 10 en fonction de l'intensité du spectre obtenue.

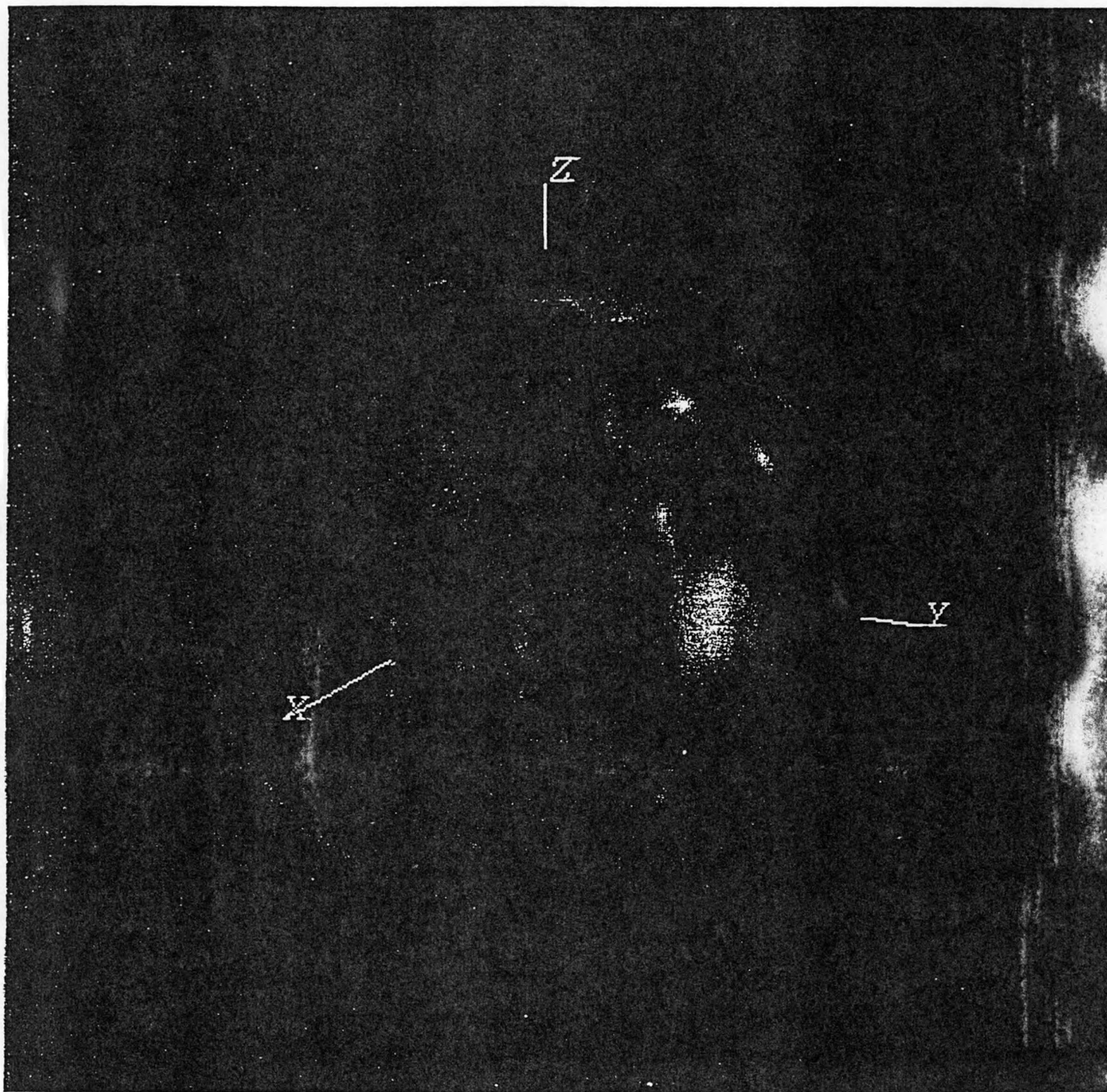
Les cartes C.2 à C.5 représentent les résultats obtenus aux temps $T = 0, 1.0, 1.8$ et $T = 2.4$ respectivement.

1. Logiciel de calcul scientifique SIMULOG-INRIA

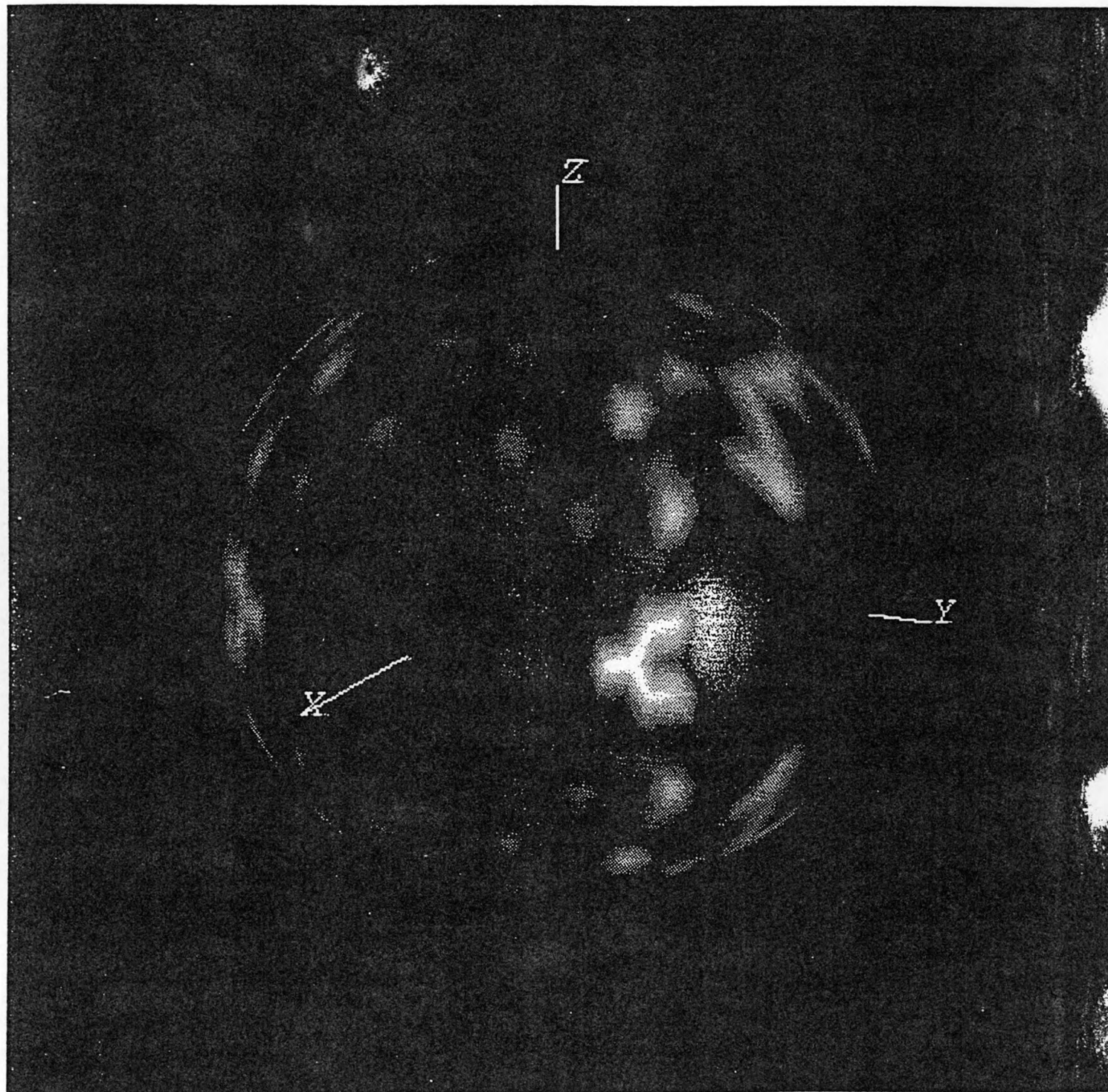
La comparaison avec la carte C.1 montre une “bonne” adéquation entre les variations angulaires des anisotropies théoriques et numériques.



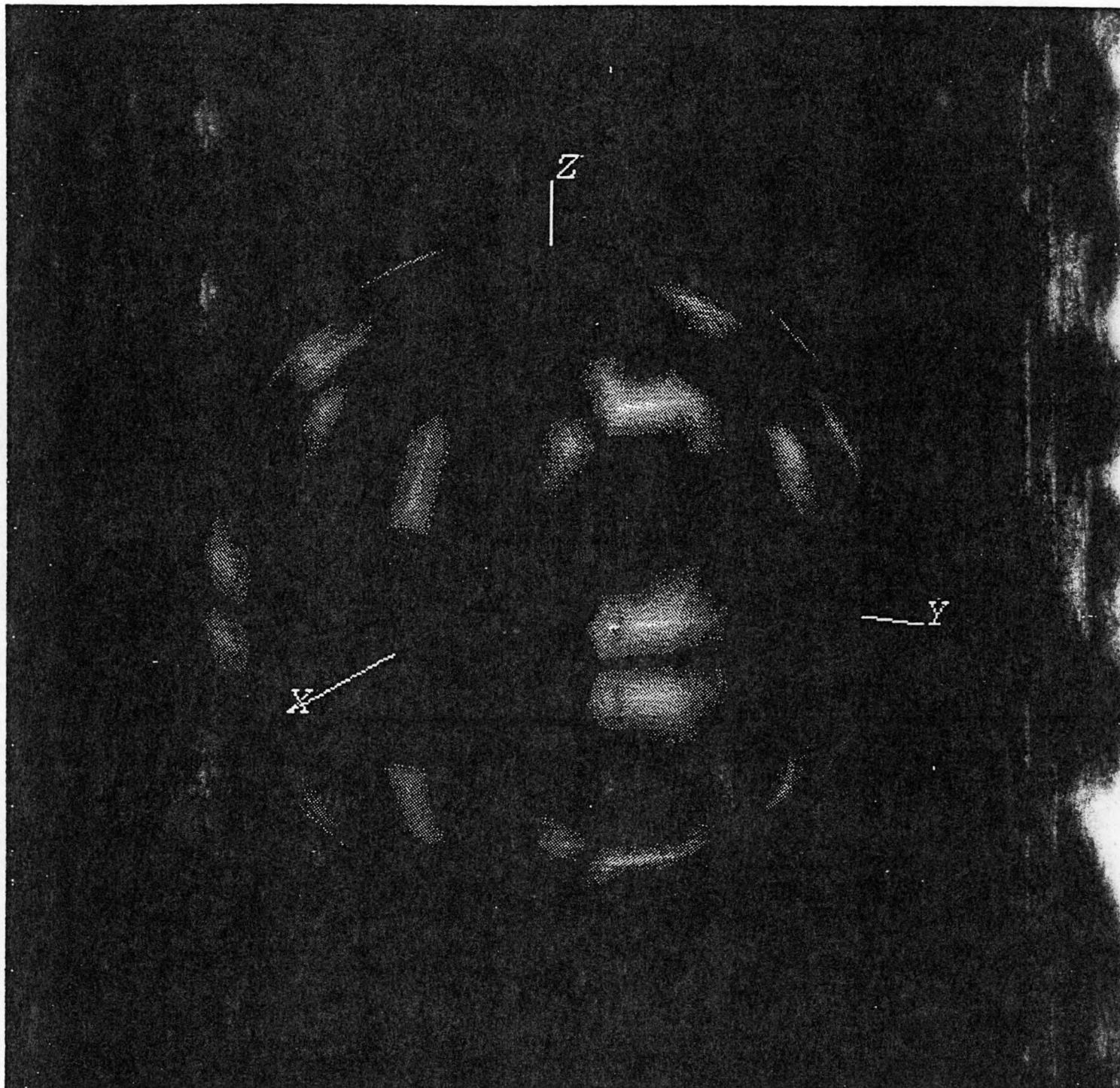
Carte C.1: Anisotropie théorique à l'ordre 4 du schéma de Vremann



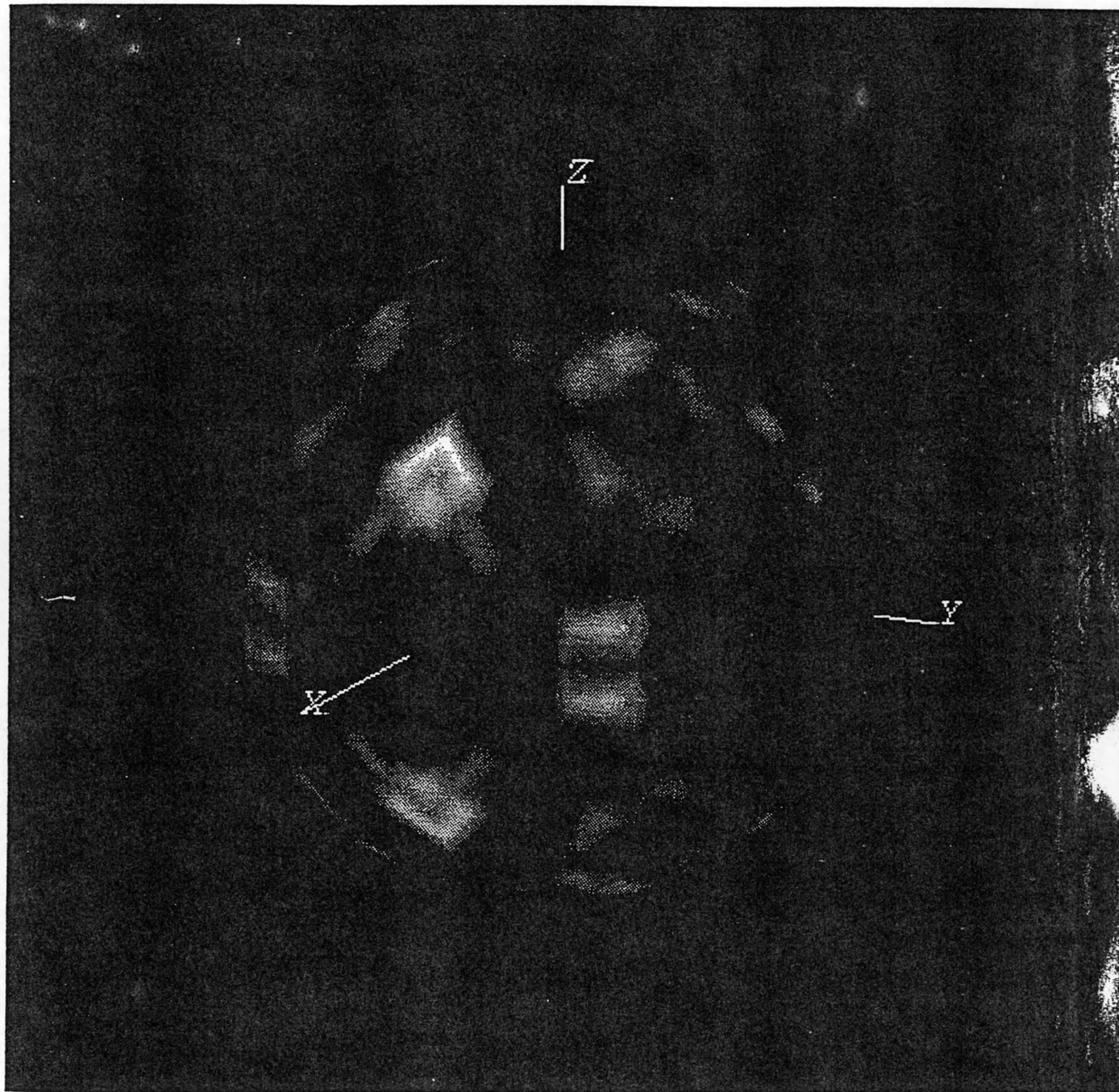
Carte C.2: Anisotropie numérique du schéma de Vremann, $t=0$



Carte C.3: Anisotropie numérique du schéma de Vremann, $t=1.0$



Carte C.4: Anisotropie numérique du schéma de Vreman, $t=1.8$



Carte C.5: Anisotropie numérique du schéma de Vremann, $t=2.4$

Afin de confirmer cette impression, nous avons calculé, toujours au nombre d'onde $K = 10$, le coefficient de corrélation linéaire entre les distributions théoriques et numériques de l'anisotropie.

En tout point i de la sphère construite par Modulef, on dispose en effet, et à tout instant, de deux valeurs T_i et $S_i(t)$, représentant respectivement l'anisotropie théorique et l'anisotropie au temps t obtenue par simulation.

N désignant le nombre de points considérés, le coefficient de corrélation linéaire au temps t a pour définition :

$$\kappa(t) = \frac{1}{N-1} \sum_{i=1}^{i=N} (T_i - E(T))(S_i(t) - E(S)) \quad (2.25)$$

L'évolution de $\kappa(t)$ en fonction du temps est donnée à la figure 2.3.

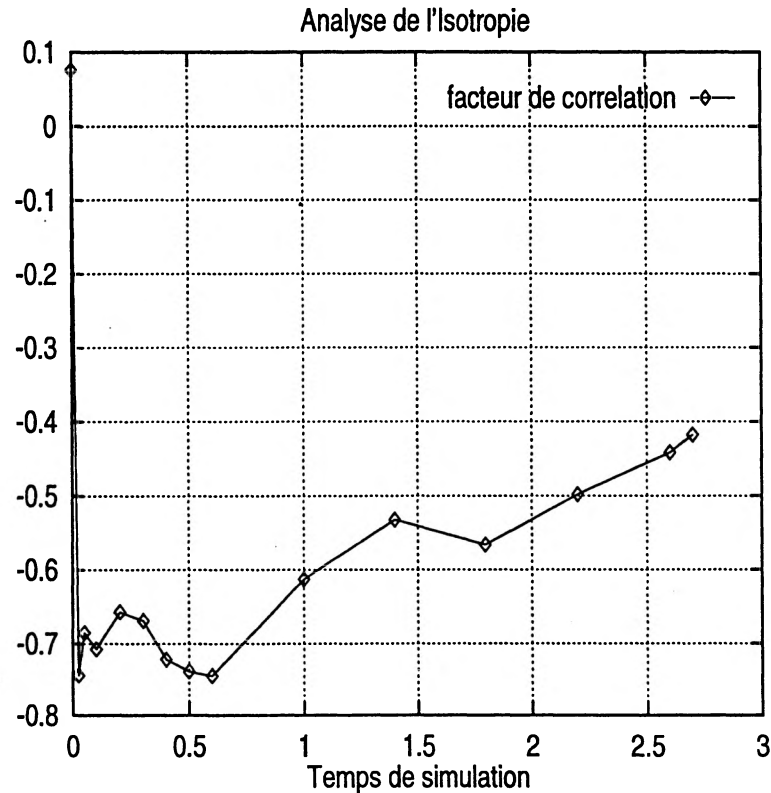


FIG. 2.3 - Propriétés statistiques de l'écoulement

La valeur très élevée de ce coefficient (N étant égal à 1178, une valeur de $|\kappa(t)|$ supérieure à 0.2 traduit déjà une forte probabilité de corrélation) montre que l'anisotropie de la distribution angulaire du spectre numérique, est quasiment entièrement imputable à l'anisotropie du schéma numérique.

Cette observation nous confirme dans notre idée première, qu'il serait préférable d'utiliser des schémas 4-isotropes pour simuler un écoulement THI!

Une étude plus approfondie de la figure 2.3 permet d'observer une lente décroissance de la valeur de $|\kappa(t)|$ lorsque t croît, celle-ci correspond à une progressive décorrélation entre les deux séries $\{T_i\}$ et $\{S_i(t)\}$.

Deux explications peuvent être apportées pour tenter d'expliquer ce phénomène.

- Le caractère aléatoire de l'écoulement (n'oublions pas que le champ de vitesse initial est aléatoire), ainsi que les inévitables perturbations numériques, peuvent entraîner une décorrélation progressive. En effet la série $\{S_i(t)\}$ peut être représentée sous la forme $S_i(t) = T_i + e_i(t)$ où $e_i(t)$ est un résidu aléatoire qui peut ne pas être stationnaire en temps : son rôle pourrait alors devenir dominant lorsque t augmente?
- La non linéarité du terme convectif de l'équation de conservation de l'impulsion, dont on sait qu'elle se traduit dans l'espace de Fourier par un mélange de modes (la "*fameuse*" interaction triadique) pourrait suffire à expliquer cette décorrélation.

Quelle que soit la part de vérité dans ces deux tentatives d'explication, elles traduisent en fait une sorte de "*retour à l'isotropie*" de l'écoulement (initialement isotrope) dont il serait intéressant d'approfondir la signification, ceci d'autant plus que ce phénomène a été mis en évidence dans de nombreuses expériences de laboratoire.

Dans la section 2.5 plusieurs schémas ont été proposés : un schéma 4-isotrope, mais non régularisant, et 2 schémas, non 4-isotropes, mais dont l'un d'entre eux était moins régularisant que le schéma de Vremann, et l'autre plus.

Commençons par étudier le comportement du schéma 4-isotrope.

La répartition angulaire du spectre sur la sphère de rayon 10, diffère assez sensiblement de la répartition obtenue pour le schéma de Vremann, mais n'est pas uniforme.

Ce qui s'explique par la répartition spatiale aléatoire du champ de vitesses et nous permet au passage d'avoir une idée de l'importance des fluctuations que celle-ci induit.

Comme nous pouvions le supposer compte tenu de son effet non régularisant ce schéma est moins "*robuste*" que le schéma de Vremann au sens où il ne permet pas de poursuivre des simulations aussi loin que celui-ci (respectivement $t = 0.64$ et $t = 1.6$).

Le caractère régularisant d'un schéma semblant donc être, pour les simulations entreprises ici, une propriété essentielle, nous avons décidé de ne tester que le troisième schéma construit à la section 2.5 ($c = 4/16$, $a = 2/16$, $b = 1/16$) dont

l'écart-type est $\sqrt{2}$ fois plus important que celui caractérisant le schéma de Vremann.

Avec ce schéma les simulations ont été poursuivies jusqu'au temps $t = 2.7$ (mais pas au delà pour des questions de disponibilité de la CM-5), ce qui confirme bien le rôle primordial joué par cet aspect régularisant sur la maîtrise des risques d'explosion des calculs.

La distribution angulaire d'anisotropie numérique, ne diffère guère de celle obtenue avec le schéma de Vremann, résultat prévisible dans la mesure où les anisotropies théoriques des deux schémas sont, à une homothétie près, très proches l'une de l'autre.

2.7 Conclusions

Tout d'abord, nous avons démontré tout l'intérêt qu'il y avait à entreprendre une étude fine du schéma utilisé, même si celle-ci se cantonnait au cas d'une équation scalaire linéaire.

En particulier, le concept d'équation équivalente s'est avéré être un outil précieux pour la compréhension du rôle du schéma numérique sur le spectre de la solution discrète.

Plus directement en relation avec les simulations effectuées, le rôle stabilisateur du schéma de Vremann a été clairement mis en évidence, et des variantes de ce schéma ont été construites et testées.

Nous avons ainsi pu montrer que l'une des propriétés essentielles que doivent posséder ces schémas, est la capacité à régulariser suffisamment les champs turbulents, pour prévenir tout arrêt accidentel des simulations (explosions ou pas de temps quasiment nuls). Il serait intéressant de changer d'état turbulent tout en gardant l'approche de discrétisation par volumes finis, pour sentir l'effet de ce type de schémas.

En outre, le fait que le schéma adopté soit ou non 4-isotrope, ne semble pas avoir une incidence importante sur la qualité des résultats obtenus.

Chapitre 3

Quelques remarques techniques

3.1 Discrétisation du tenseur de taux de déformation

Comme cela avait déjà été mis en évidence dans [BAU93], une des difficultés majeures auxquelles nous avons été confronté, a été le maintien d'une zone inertielle jusqu'à des temps élevés.

En d'autres termes, et en contradiction complète tant avec les prédictions théoriques (théorie de Kolmogorov [KOL41a] [KOL41b]) qu'avec les résultats expérimentaux, les simulations effectuées n'ont à aucun moment permis de mettre en évidence l'existence d'un état d'équilibre spectral, ni donc d'une solution auto-semblable.

Soulignons de plus que ce problème n'est pas propre à un type de modèle de sous-maille particulier, puisque aussi bien le modèle de Smagorinsky que celui de fonction de structure présentent le même défaut.

En première analyse, le phénomène de croissance du spectre aux nombres d'onde les plus élevés, nous incite à incriminer l'efficacité du modèle de sous-maille à assurer un transfert correct d'énergie des grandes vers les petites échelles.

Si l'on s'en tient exclusivement au modèle de Smagorinsky, il semble donc que celui-ci n'apporte pas un excès de dissipation suffisamment important.

Certes le schéma peut être en partie responsable de ce phénomène, mais rien dans l'analyse linéaire que nous en avons effectué au chapitre précédent, n'a permis de confirmer cette supposition.

C'est donc tout naturellement, que nous avons regardé si les différents éléments qui interviennent dans l'expression des modèles de fermeture testés, ne pouvaient, par une mauvaise estimation de leurs valeurs, permettre d'expliquer cette remontée du spectre.

Les deux modèles exhibant le même comportement, seul le modèle de Smago-

rinsky a été étudié dans la suite.

Pour celui-ci, la viscosité turbulente est donnée par la relation:

$$\nu_t = C^2 \Delta^2 \Pi^{1/2} \quad (3.1)$$

où Π est le second invariant du tenseur de taux de déformations ($\Pi = S_{ij}S_{ij}$).

Dans l'expression ci-dessus, Δ est une longueur caractéristique, censée caractériser la dimension des structures de sous-maille les plus énergétiques.

D'une définition peu précise, on tire au moins la conviction que Δ est au plus égale à une longueur caractéristique du maillage (pas h du maillage, diamètre des éléments de la triangulation, etc.).

Nous supposons donc avoir la relation suivante:

$$\Delta \leq h$$

Néanmoins, en faisant varier la constante C , qui intervient uniquement dans le groupement $(C\Delta)^2$, il est possible de "*corriger*" l'effet d'une plus ou moins grande dissipation liée à une sur ou sous évaluation de Δ .

Si l'on accepte l'inégalité ci-dessus, Π doit aussi représenter l'effet du tenseur de déformations de sous-maille. Or il est bien évident que ne disposant que de la connaissance d'informations concernant les grandes échelles, il ne sera pas possible d'obtenir une évaluation correcte de Π .

Comme les modes qui composent le champ de vitesses de sous-maille sont les plus grands, on comprend que ce champ est une fonction plus rapidement oscillante que le champ de vitesses de grandes échelles. De manière heuristique, on s'attend bien à observer que le champ de sous-maille soit plus oscillant que le champ de vitesses de grandes échelles. Et vu que l'évaluation de Π se fait par des composantes résolues, il y a malheureusement une sous-évaluation de cette réelle quantité.

Cette remarque nous pose le problème du procédé à utiliser pour calculer Π . De manière a priori fort logique, l'approximation des dérivées spatiales qui interviennent dans l'expression de Π , est effectuée en utilisant des différences finies centrées. Or cette approximation consiste à évaluer des gradients sur une base de calcul de l'ordre de $2h$, ce qui est en opposition complète avec le paragraphe ci-dessus!

La base de calcul minimale pour cette évaluation étant égale à h , et la valeur de Π demeurant encore sous-estimée dans ce cas, nous avons testé l'approximation suivante des dérivées spatiales (donnée ici uniquement pour la dérivée suivant x):

$$\frac{\partial u}{\partial x} = \text{Max} \left(\frac{u_{i,j,k} - u_{i-1,j,k}}{h}, \frac{u_{i+1,j,k} - u_{i,j,k}}{h} \right)$$

A la figure 3.2, nous montrons que les résultats obtenus avec cette expression sont "*qualitativement*" meilleurs que ceux obtenus, dans les mêmes conditions avec l'approximation centrée (à la Fig. 1.2, l'évolution des spectres obtenus a un comportement très proche pour la constante de Smagorinsky $C_s = 0.28$).

Comme nous pouvions nous y attendre, on observe dans ce cas une décroissance du spectre plus importante que dans le cas centré.

Avec cette approximation des dérivées, le positionnement sur la pente théorique de $-5/3$ nécessite, à un même instant qu'avec une approximation centrée, une valeur plus faible de la constante de Smagorinsky C (qui admet alors des valeurs plus conformes à celles disponibles dans la littérature).

Globalement ce n'est qu'une compensation, le fait que l'augmentation de la valeur de Π s'accompagnant d'une diminution de celle de C n'apporte aucune amélioration sensible dans la maîtrise des problèmes de remontée de spectre, ou d'arrêt accidentel des simulations.

En fin de compte, on peut douter de la capacité du modèle de Smagorinsky à permettre une simulation correcte d'un écoulement THI!

Comme le calcul de viscosité turbulente par le modèle "Fonction de structure" s'effectue autrement, on pourrait regarder aussi la discrétisation adoptée pour son écriture.

Nombre de Reynolds	1500
Discrétisation	32^3
Fermeture	$C_s = 0,28$

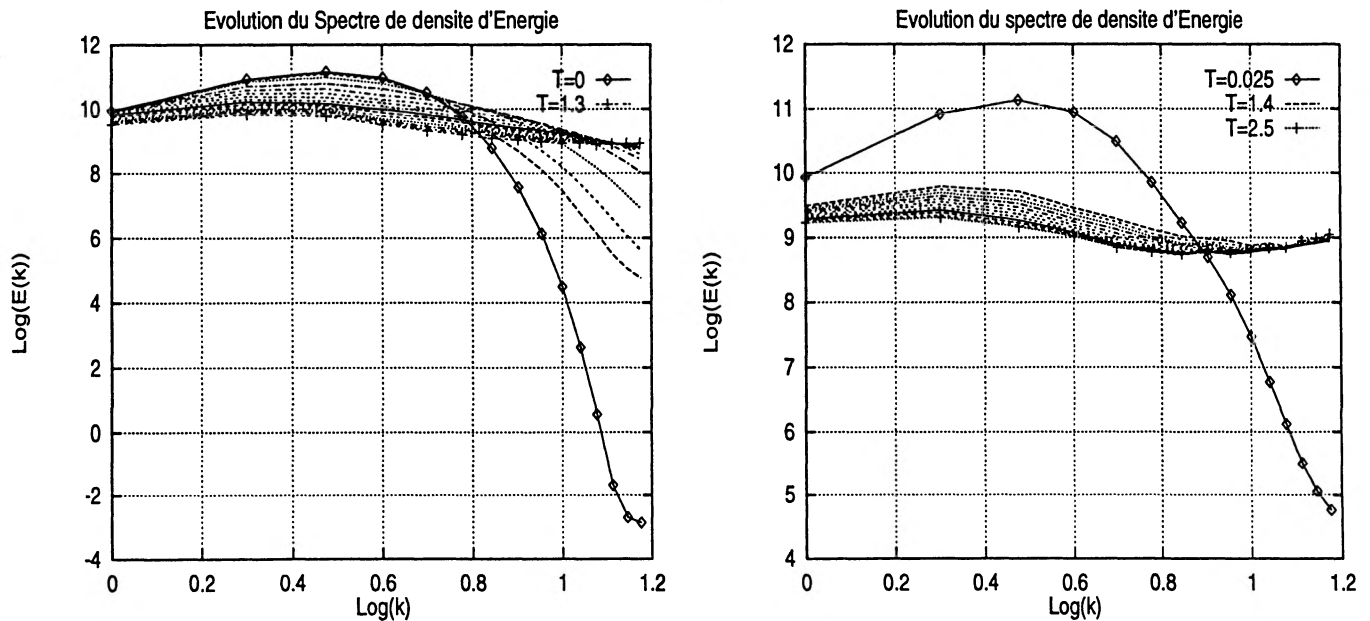


FIG. 3.2 - Evolution de spectres d'énergie (CFL de 0,6)

3.2 Critique des modèles de fermeture utilisés

Les limites du modèle de Smagorinsky sont bien connues et ont fait l'objet de multiples études (pour une synthèse de celles-ci, voir par exemple [SAG]).

Cependant ces restrictions d'utilisation concernent essentiellement les cas d'écoulements relativement complexes (cisailés, couches limites *etc.*), la capacité de ce modèle à fournir des résultats corrects en THI n'ayant, à notre connaissance, jamais été à controverses. Or il est évident, à la lumière des résultats présentés plus haut, que ce modèle, **dans les situations de simulations que nous nous sommes fixées**, ne nous a pas permis d'atteindre la qualité de résultats escomptée. A ce titre la question de la validité de ce modèle doit être envisagée.

Lorsqu'on examine l'expression de la viscosité turbulente introduite par ce modèle, un premier point gênant (déjà évoqué dans la section précédente), est que cette viscosité est fonction du groupement $C\Delta$, dans lequel d'une part Δ n'a pas de définition précise, et d'autre part C n'apparaît que comme une constante servant à ajuster la décroissance du spectre sur sa valeur théorique.

Le degré de liberté offert par cette constante est en fait bien pratique, car il permet pratiquement toujours d'obtenir, au moins dans une certaine plage de temps, une pente correcte du spectre. Cependant, comme nous l'avons montré plus haut, elle rend illusoire toute amélioration de l'évaluation de ν_t par une estimation plus précise des gradients de vitesse. Cette constante est donc en fait une arme à double tranchant !

Plus grave est l'impossibilité d'atteindre une solution auto-semblable, caractéristique d'ailleurs partagée avec le modèle de fonction de structure. Ce problème révèle une incapacité de ces deux modèles, à assurer un transfert correct d'énergie cinétique des grandes échelles vers les petites.

Nous devons cependant signaler, que des raffinements du modèle de fonction de structure ont été proposés (communication personnelle de O. Métais), qui permettent peut être de résoudre ce problème. Ceux-ci n'ayant cependant pas été pris en compte dans le cadre de cette étude, nous devons donc nuancer la conclusion précédente en disant que la "*version standard*" de ce modèle ne donne pas satisfaction.

Il n'en demeure pas moins que les simulations correctes, au sens où celles-ci peuvent être menées jusqu'à des temps arbitrairement longs, d'écoulements THI nous semble être un problème ouvert, du moins lorsqu'elles sont réalisées dans l'espace physique.

Afin d'essayer de pallier cette déficience des modèles utilisés, nous avons construit, et testé, un modèle de fermeture original que nous détaillons dans la référence [PS95].

Conclusions et développements

Nous avons pu mettre en évidence la difficulté, pour ne pas dire l'impossibilité intrinsèque au modèle de Smagorinsky à conduire à un état auto-semblable correct. Par ailleurs, l'étude d'autres caractéristiques du champ turbulent (*cf.* [BAU93]) (Skewness, micro-échelle de Taylor) a permis de voir qu'il est là aussi difficile d'obtenir, pour chacune d'entre elles, des résultats en accord avec les résultats expérimentaux disponibles .

Néanmoins, à la faveur de nombreuses communications personnelles (D. Laurence et V. Maupuis du Laboratoire National d'Hydraulique, J.P. Bertoglio de l'Ecole Centrale de Lyon en particulier), il est apparu clairement que les difficultés que nous avons rencontrées en Turbulence Homogène Isotrope, sont partagées par d'autres auteurs, et semblent même devoir être considérées comme le résultats de la majorité de simulations LES.

Dans la continuité de ces travaux préliminaires, il est donc normal d'entreprendre une seconde phase d'étude en axant celle-ci selon deux directions : d'une part étudier le comportement d'un autre modèle de fermeture "*assez différent*" du modèle de Smagorinsky, ceci afin de voir si les problèmes mentionnés plus haut sont inhérents au modèle ou à la méthode elle même ; d'autre part entreprendre une étude aussi fine que possible du rôle du schéma numérique.

Sur ce premier point, notre choix se porte sur le modèle de Fonction de Structure développé par Lesieur et Métais.

Nous avons cependant rapidement constaté, que les résultats auxquels il conduit ne sont pas significativement différents que ceux obtenus avec le modèle de Smagorinsky.

Bien que le modèle de Bardina -ou modèle dynamique- n'ait pas été testé, nous ne serions pas étonnés qu'aucun modèle de fermeture ne permette d'effectuer une simulation de type LES d'un écoulement THI **correcte**, au sens où le spectre obtenu présenterait une forme auto-semblable en accord avec la théorie de Kolmogorov.

A l'opposé de cette conclusion partielle négative, le second axe de recherche nous a permis d'obtenir des résultats très intéressants.

En effet, l'étude du schéma entreprise à partir de son équation équivalente, ainsi

que l'étude d'un certain nombre de points techniques faisant l'objet du chapitre 3, nous ont permis d'affiner notre compréhension de l'influence de la discrétisation sur les résultats obtenus en simulation. Nous avons pu montrer que le caractère régularisant d'un schéma, bien marqué dans le schéma de Vremann, avait un rôle essentiel pour limiter les problèmes de croissance du spectre dans les grands nombres d'ondes.

On voit donc, que le caractère régularisant d'un schéma peut compenser les déficiences intrinsèques des modèles de fermeture spatiaux.

Ce résultat, qui demande confirmation, est extrêmement important dans la mesure où il éclaire l'interaction schéma/modèle de fermeture, et pourrait permettre à terme de modifier des schémas existant ou d'en concevoir de nouveaux pour une utilisation spécifique LES.

Annexe A

Analyse du comportement de ν en fonction de Re

Nous montrons dans cette annexe pourquoi la constante de Smagorinsky diminue lorsque le nombre de Reynolds augmente.
Partons pour cela de l'expression de la viscosité turbulente effective introduite par Heisenberg (cf. [McC92], relation 2.154) :

$$\nu(k', t) = \int_{k'}^{\infty} A k^{-3/2} \sqrt{E(k, t)} dk$$

expression dans laquelle A est une constante strictement positive.
Supposons que la zone inertielle s'étende jusqu'au nombre d'onde K on a alors :

$$\nu(k', t) = \underbrace{A \int_{k'}^K k^{-3/2} \sqrt{\left(\alpha \epsilon^{2/3}(t) k^{-5/3}\right)} dk}_{J(k', K, t)} + A \int_K^{\infty} k^{-3/2} \sqrt{E(k, t)} dk$$

et, après intégration :

$$J(k', K, t) = -\frac{3}{4} AC(t) (K^{-\frac{4}{3}} - k'^{-\frac{4}{3}})$$

où $C(t)$ est une fonction à valeurs positives du temps t , donnée par la relation :

$$C(t) = \alpha^{1/2} \epsilon^{1/3}(t)$$

On a de plus :

$$\lim_{K \rightarrow \infty} J(k', K, t) = \frac{3}{4} C(t) k'^{-4/3} \stackrel{def}{=} \nu_{\infty}(k', t)$$

et donc :

$$\frac{\partial \nu_{\infty}(k', t)}{\partial k'} = -AC(t) k'^{-7/3}$$

Remarquons au passage que cette formule permet de prouver la positivité de A , puisque la valeur de la viscosité turbulente diminue, fort logiquement, lorsque le nombre de modes simulés (k') augmente.

Fixons maintenant la valeur de k' , et remarquons que $\partial J/\partial K$ est donnée par :

$$\frac{\partial J(k', K, t)}{\partial K} = AC(t)K^{-7/3} > 0.$$

k' étant fixé, posons $\lambda = K/k'$ (ce qui implique généralement $\lambda > 1$ dans le cas d'une simulation LES), et étudions la fonction $\tilde{J}$ définie par :

$$\tilde{J}(k', \lambda, t) = -\frac{3}{4}AC(t)k'^{-4/3}(\lambda^{-4/3} - 1)$$

K étant d'autant plus grand que le nombre de Reynolds est élevé [BER93], on en déduit les relations suivantes :

$$\frac{\partial J}{\partial Re} = \frac{\partial J}{\partial K} \frac{\partial K}{\partial Re} > 0$$

et,

$$\frac{\partial \nu}{\partial Re} = \frac{\partial J}{\partial K} \frac{\partial K}{\partial Re} + \frac{\partial}{\partial Re} \left(A \int_K^\infty \sqrt{\frac{E(k, t)}{k^3}} dk \right)$$

Le deuxième terme du membre de droite de l'équation précédente peut être négligé si K est très grand (cette intégrale est évidemment convergente puisque l'énergie E est finie, et que celle-ci décroît exponentiellement dans la zone de dissipation).

Si l'on utilise maintenant la relation existant entre K et la valeur Re du nombre de Reynolds (cf. [BER93], Annexe 5) :

$$K \sim Re^{\frac{3}{4}}$$

on trouve finalement :

$$\begin{aligned} \frac{\partial \nu}{\partial Re} &\sim K^{-7/3} Re^{-1/4} \\ &\sim Re^{-2} \end{aligned}$$

La viscosité turbulente varie donc en $1/Re$, ce qui confirme la décroissance de la valeur de la constante de Smagorinsky observée expérimentalement lorsque Re croît. On obtient en effet (cf. 1.8) la relation de proportionnalité suivante :

$$C_s \sim \frac{1}{Re}$$

Annexe B

Définition du système sans dimension

Il est très important d'adimensionner le système d'équations précédemment obtenu pour avoir en mémoire des variables du même ordre de grandeur.

B.1 Définition des variables adimensionnées

Pour chaque variable V , la grandeur de référence sera notée V_0 et la variable sans dimension V' (ces notations ne sont valables que pour ce chapitre).

- pour les longueurs: $x'_i = \frac{x_i}{L_0}$, L_0 étant la longueur de référence du domaine physique.
- pour les vitesses: $u'_i = \frac{u_i}{U_0}$, U_0 étant la norme moyenne des vitesses.
- pour le temps: $t' = \frac{U_0}{L_0} t$
- pour la densité: $\rho' = \frac{\rho}{\rho_0}$
- pour la température: $T' = \frac{T}{T_0}$ avec $T_0 = \frac{U_0^2}{\gamma R M_r^2}$
- pour la pression: $p' = \frac{p}{p_0}$ avec $p_0 = \rho_0 U_0^2$
- pour l'énergie: $E' = \frac{E}{E_0}$ avec $E_0 = \rho_0 U_0^2$
- pour la viscosité dynamique: $\mu' = \frac{\mu}{\mu_0}$
- pour la diffusivité dynamique: $\phi' = \frac{\phi}{\phi_0}$

B.2 Adimensionnement des équations

On part du système modélisant l'écoulement d'un fluide newtonien visqueux compressible. Ce problème, provenant de la physique, peut être régi par les trois lois de conservation suivantes : conservation de la masse (B.1), conservation de la quantité de mouvement (B.2) et la conservation de l'énergie totale (B.3).

$$\frac{\partial \rho}{\partial t} + \sum_{j=1}^3 \frac{\partial(\rho u_j)}{\partial x_j} = 0 \quad (\text{B.1})$$

$$\frac{\partial}{\partial t}(\rho u_i) + \sum_{j=1}^3 \frac{\partial}{\partial x_j}(\rho u_i u_j) - \sum_{j=1}^3 \frac{\partial \sigma_{ij}}{\partial x_j} + \frac{\partial p}{\partial x_i} = 0 \quad (\text{B.2})$$

$$\frac{\partial E}{\partial t} + \sum_{j=1}^3 \frac{\partial}{\partial x_j} \left((E + p) \rho u_j \right) - \sum_{i=1}^3 \sum_{j=1}^3 \frac{\partial}{\partial x_j} (\sigma_{ij} u_i) + \sum_{j=1}^3 \frac{\partial}{\partial x_j} Q_j = 0, \quad (\text{B.3})$$

Faisant suite aux changements de variables du paragraphe précédent, on obtient une relation entre dérivées spatiales et temporelles des variables physiques et celles des nouvelles variables adimensionnées :

$$\begin{aligned} \frac{\partial f}{\partial t}(t) &= \frac{U_0}{L_0} \frac{\partial f'}{\partial t'}(t') \\ \frac{\partial f}{\partial x}(x) &= \frac{1}{L_0} \frac{\partial f'}{\partial x'}(x') \end{aligned}$$

Il suffit de multiplier l'équation de masse (B.1) par $\frac{L_0}{\rho_0 U_0}$, et utiliser les formules précédentes pour obtenir la nouvelle équation de la masse :

$$\frac{L_0}{\rho_0 U_0} \cdot \left(\frac{\partial \rho'}{\partial t'} \cdot \frac{\rho_0 U_0}{L_0} + \sum_{j=1}^3 \frac{\partial(\rho' u'_j)}{\partial x'_j} \cdot \frac{\rho_0 U_0}{L_0} \right) = 0$$

soit :

$$\frac{\partial \rho'}{\partial t'} + \sum_{j=1}^3 \frac{\partial(\rho' u'_j)}{\partial x'_j} = 0 \quad (\text{B.4})$$

De manière analogue, on multiplie l'équation de la quantité de mouvement par $\frac{L_0}{\rho_0 U_0^2}$,

$$\begin{aligned} & \frac{L_0}{\rho_0 U_0^2} \cdot \left(\frac{U_0}{L_0} \frac{\partial}{\partial t'} (\rho_0 \rho' U_0 u'_i) + \frac{1}{L_0} \sum_{j=1}^3 \frac{\partial}{\partial x'_j} (U_0^2 u'_i u'_j) \right. \\ & \quad \left. - \frac{1}{L_0} \sum_{j=1}^3 \frac{\partial \sigma_{ij}}{\partial x'_j} + \frac{1}{L_0} \frac{\partial}{\partial x'_i} (p_0 p') \right) = 0, \end{aligned} \quad (\text{B.5})$$

or, $\sigma_{ij} = \frac{U_0 \mu_0}{L_0}$, d'où :

$$\begin{aligned} (\text{B.2}) \quad & \iff \frac{\partial}{\partial t'} (\rho' u'_i) + \sum_{j=1}^3 \frac{\partial}{\partial x'_j} \left(\frac{\rho_0 u'_i \cdot \rho_0 u'_j}{\rho_0} \right) \\ & - \frac{\mu_0}{\rho_0 L_0 U_0} \sum_{j=1}^3 \frac{\partial \sigma'_{ij}}{\partial x'_j} + \frac{\partial p'}{\partial x'_i} = 0, \end{aligned} \quad (\text{B.5}) \quad (\text{B.6})$$

ce qui fait apparaître le nombre sans dimension $\frac{\mu_0}{\rho_0 L_0 U_0}$ appelé le nombre de Reynolds et noté Re .

$$Re = \frac{\rho_0 U_0 L_0}{\mu_0}.$$

Pour obtenir l'équation de conservation de l'énergie totale adimensionnée, il suffit de multiplier (B.3) par $\frac{L_0}{\rho_0 U_0^3}$ et d'utiliser la technique similaire.

$$\begin{aligned} & \left(\frac{L_0}{\rho_0 U_0^3} \right) \cdot \left\{ \left(\frac{U_0}{L_0} \frac{\partial}{\partial t'} (E_0 E') + \frac{1}{L_0} \sum_{j=1}^3 \frac{\partial}{\partial x'_j} \left[E_0 U_0 (E' + p') \frac{\rho' u'_j}{\rho'} \right] \right. \right. \\ & \quad \left. \left. - \frac{1}{L_0} \sum_{i=1}^3 \sum_{j=1}^3 \frac{\partial}{\partial x'_j} \left[\frac{U_0^2 \mu_0}{L_0} \sigma'_{ij} u'_i \right] + \frac{1}{L_0} \sum_{j=1}^3 \frac{\partial}{\partial x_j} \left(\frac{\mu R \gamma}{(\gamma - 1) P_r} \frac{\partial T}{\partial x_j} \right) \right\} = 0, \end{aligned} \quad (\text{B.8})$$

le dernier terme du membre de gauche s'écrit encore :

$$\frac{\mu_0 T_0 \gamma R}{L_0 U_0^3 \rho_0 (\gamma - 1) P_r} \sum_{j=1}^3 \frac{\partial}{\partial x'_j} \left(\mu' \frac{\partial T'}{\partial x'_j} \right) = \frac{\gamma}{(\gamma - 1) P_r Re} \sum_{j=1}^3 \frac{\partial}{\partial x'_j} \left(\mu' \frac{\partial}{\partial x'_j} \left(\frac{p'}{\rho'} \right) \right)$$

car $T' = \gamma M_r^2 \frac{p'}{\rho'}$.

Finalement, nous obtenons :

$$\begin{aligned} & \frac{\partial E'}{\partial t'} + \sum_{j=1}^3 \frac{\partial}{\partial x'_j} \left[(E' + p') \frac{\rho' u'_j}{\rho'} \right] - \frac{1}{Re} \sum_{j=1}^3 \sum_{i=1}^3 \frac{\partial}{\partial x'_j} (\sigma'_{ij} u'_i) \\ & + \frac{\gamma}{(\gamma - 1) P_r Re} \sum_{j=1}^3 \frac{\partial}{\partial x'_j} \left(\mu' \frac{\partial}{\partial x'_j} \left(\frac{p'}{\rho'} \right) \right) = 0. \end{aligned} \quad (\text{B.9})$$

A partir de ce nouveau système adimensionné (B.4-B.9), nous donnons les points essentiels de l'étape de filtrage de la méthode LES. On fait les hypothèses que le filtre commute avec les dérivées temporelle et spatiale, ce qui pré-suppose que les variables doivent en partie être dérivables par rapport au temps. Cet abus est toléré pour ce type de problèmes.

Par contre, pour les termes non linéaires, il n'est pas possible de procéder directement car le filtrage ne peut pas se distribuer sur la multiplication. Pour l'équation de la quantité de mouvement (B.5), on fait apparaître ce défaut de distributivité, appelé le tenseur de Léonard, qu'il est nécessaire de modéliser car non résolue. Comme il s'exprime par des quantités non résolues, d'une échelle plus petite, ce tenseur fait partie des tenseurs de sous-maille.

De l'équation de la conservation d'énergie, on fait apparaître aussi deux autres tenseurs de sous-maille adimensionnés π'_j et q'_j qui sont modélisés par une expression faisant intervenir la diffusivité tourbillonnaire.

$$\frac{\partial \pi'_j}{\partial x'_j} = \frac{\partial}{\partial x'_j} (\overline{p' u'_j} - \overline{p'} \tilde{u}'_j) = \frac{\partial}{\partial x'_j} \left(-\frac{\phi'_t}{P'_{r,t}} \frac{\partial p'}{\partial x'_j} \right)$$

Pour des raisons pratiques, on introduit souvent la moyenne de Favre pour les écoulements compressibles, qui est définie par :

$$\widetilde{X} = \frac{\overline{\rho X}}{\overline{\rho}}$$

$$\begin{aligned} \frac{\partial q'_j}{\partial x'_j} &= \frac{\partial}{\partial x'_j} [\overline{\rho'} (\widetilde{T' u'_j} - \widetilde{T'} \widetilde{u}'_j)] \cdot \frac{1}{\gamma(\gamma-1)M_r^2} \\ &= \frac{\partial}{\partial x'_j} \left[-\frac{\phi'_t}{P'_{r,t}} \frac{\partial (\overline{\rho'} \widetilde{T'})}{\partial x'_j} \right] \cdot \frac{1}{\gamma(\gamma-1)M_r^2} \\ &= \frac{\partial}{\partial x'_j} \left[-\frac{\phi'_t}{P'_{r,t}} \frac{\partial \overline{p'}}{\partial x'_j} \right] \cdot \frac{1}{\gamma-1} \end{aligned}$$

de même, R étant la constante des gaz parfaits, on a :

$$\begin{aligned} \frac{\partial \overline{Q'_j}}{\partial x'_j} &= -\frac{\partial}{\partial x_j} \left[\frac{\mu \gamma R}{(\gamma-1)P_r} \cdot \frac{\partial \overline{T}}{\partial x_j} \right] \\ &= -\frac{\partial}{\partial x'_j} \left[\mu' \frac{\partial \overline{T'}}{\partial x'_j} \right] \cdot \frac{\gamma R \mu_0 T_0}{(\gamma-1)P_r L_0^2} \\ &= -\frac{\partial}{\partial x'_j} \left[\mu' \frac{\partial}{\partial x'_j} \left(\frac{\overline{p'}}{\overline{\rho'}} \right) \right] \cdot \frac{\gamma}{(\gamma-1)P_r Re} \end{aligned}$$

B.3 Choix des paramètres de calcul

L'adimensionnement précédent fait apparaître quatre nombres caractéristiques :

- Nombre de Reynolds Re (cf. B.1) :

$$Re = \frac{\rho_0 U_0 L_0}{\mu_0}$$

- Nombre de Prandtl P_r (ou $P_{r,t}$ si l'on considère sa valeur turbulente) :

$$P_r = \frac{\rho_0 U_0 L_0}{\phi_0} = \frac{\mu_0}{\phi_0} Re$$

- Nombre de Mach M_r :

$$M_r = \frac{U_0}{C} \quad \text{avec : } C = \sqrt{\frac{\gamma p_0}{\rho_0}}$$

C désignant ici la vitesse du son.

- Le coefficient γ .

La viscosité turbulente μ_t est adimensionnée, comme la viscosité moléculaire, par μ_0 : la viscosité turbulente adimensionnée μ'_t a alors pour définition :

$$\mu'_t = C_s^2 \frac{1}{N^2} \Pi'$$

où Π' est calculé à partir du tenseur de taux de déformations adimensionné $S'_{i,j}$ par la formule $\Pi = S'_{i,j} S'_{i,j}$.

Le calcul du spectre initial (cf. la relation 2.7) est effectué en variables adimensionnées, en particulier le nombre d'onde adimensionné k' est défini par :

$$k' = k/k_0 \quad \text{avec : } k_0 = \frac{2\pi}{L_0}$$

La valeur de L_0 est bien évidemment arbitraire, puisque elle ne peut être obtenue à partir des 4 nombres sans dimensions ci-dessus (il en est d'ailleurs de même de celle de U_0).

B.4 Obtention des grandeurs caractéristiques physiques

Dans tous les résultats de ce rapport, nous travaillons avec des variables adimensionnées. Donnons dans ce paragraphe les relations permettant de retrouver certaines grandeurs caractéristiques physiques.

Le nombre d'onde k que nous utilisons est en fait adimensionné. Les formules théoriques sont obtenues avec k dimensionné qui vaut :

$$k = \frac{2\pi}{L_0} k' \quad (\text{B.10})$$

Le spectre avec dimension $E(k)$ est donné par :

$$\int E(k) dk = \frac{2\pi}{L_0} \frac{1}{N^6} U_0^2 L_0 \int E'(k') dk' \quad (\text{B.11})$$

Le taux de dissipation ϵ est donné par :

$$E(k) = \alpha \epsilon^{\frac{2}{3}} k^{-\frac{5}{3}} \quad (\text{B.12})$$

$$\epsilon = (2\pi)^{\frac{5}{2}} \frac{u_0^3}{L_0} \epsilon' \quad (\text{B.13})$$

$$\epsilon' = \alpha^{-\frac{3}{2}} E'^{\frac{3}{2}} k'^{\frac{5}{2}} \quad (\text{B.14})$$

L'échelle de Kolmogorov vaut : (voir [BER])

$$\eta_K = \left(\frac{\nu^3}{\epsilon} \right)^{\frac{1}{4}} \quad (\text{B.15})$$

L'échelle de début de la zone est donnée par :

$$l = \frac{q^3}{\epsilon} \quad (\text{B.16})$$

$$q^2 = 2 \int E(k) dk \quad (\text{B.17})$$

L'échelle intégrale s'écrit :

$$L_i = \frac{\int k^{-1} E(k) dk}{\int E(k) dk} = \frac{L_0}{2\pi} L'_i \quad (\text{B.18})$$

et le Reynolds basé sur cette échelle :

$$Re_{L_i} = \frac{U_0 L_i}{\nu} \quad (\text{B.19})$$

$$\frac{L_i}{\eta_K} = Re_{L_i}^{\frac{3}{4}} \quad (\text{B.20})$$

Annexe C

Initialisation : génération aléatoire d'un champ turbulent

C.1 Le problème

Etant donné un spectre d'énergie $E(k)$, k module du vecteur d'onde $\vec{k}$, nous cherchons un champ de vitesse $\vec{u}$ satisfaisant les conditions suivantes :

1. $\vec{u}$ est un champ vectoriel réel.
2. $\vec{u}$ est un champ isotrope et homogène.
3. $\vec{u}$ est un champ solénoïdal ($\text{div} \vec{u} = 0$), c'est à dire qu'il existe une fonction $\vec{\psi}$ telle que $\vec{u} = \overrightarrow{\text{Rot}} \vec{\psi}$.
4. Le spectre d'énergie $E(k)$ de $\vec{u}$ est donné a priori (à un facteur multiplicatif près).

$$E(k) = \int_{\|\vec{k}\|=k} \|\overrightarrow{\hat{u}(\vec{k})}\|^2 dS.$$

En pratique, nous allons construire $\vec{\psi}$ dans l'espace de Fourier. Nous introduisons donc un opérateur défini par :

$$\begin{aligned} \overrightarrow{\widehat{\text{Rot}} \hat{\phi}} &= \widehat{\overrightarrow{\text{Rot}} \hat{\phi}} \\ &= 2\pi i(k_2 \hat{\phi}_3 - k_3 \hat{\phi}_2, k_3 \hat{\phi}_1 - k_1 \hat{\phi}_3, k_1 \hat{\phi}_2 - k_2 \hat{\phi}_1) \end{aligned}$$

Cet opérateur est linéaire :

$$\overrightarrow{\widehat{\text{Rot}}(a(\vec{k})\hat{\phi}_1(\vec{k}) + b(\vec{k})\hat{\phi}_2(\vec{k}))} = a(\vec{k})\overrightarrow{\widehat{\text{Rot}} \hat{\phi}_1(\vec{k})} + b(\vec{k})\overrightarrow{\widehat{\text{Rot}} \hat{\phi}_2(\vec{k})}$$

$\vec{\psi}$ obtenu par tirage aléatoire doit alors satisfaire les deux conditions :

$$\begin{aligned} \overrightarrow{\|\widehat{Rot}\vec{\psi}(\vec{k})\|^2} &= \frac{E(k)}{k^2} \\ \vec{\psi}(\vec{k}) &= \vec{\psi}(-\vec{k}) \end{aligned}$$

Pour satisfaire la deuxième condition, nous ne construisons $\vec{\psi}$ que dans le demi espace $k_3 \leq 0$.

C.2 La méthode

La génération de $\vec{u}$ se fait en cinq étapes (cf. [ROY80]) :

1. tirage aléatoire de 2 champs de vecteurs réels $\vec{\psi}_1(\vec{k})$ et $\vec{\psi}_2(\vec{k})$ dans le demi espace $k_3 \leq 0$.
2. tirage aléatoire de 2 fonctions a et b réelles telles que :

$$\overrightarrow{\|\widehat{Rot}(a(\vec{k})\vec{\phi}_1(\vec{k}) + ib(\vec{k})\vec{\phi}_2(\vec{k}))\|^2} = \frac{E(k)}{k^2}$$

3. extension de $\vec{\psi}(\vec{k}) = a\vec{\psi}_1(\vec{k}) + ib\vec{\psi}_2(\vec{k})$ à tout l'espace spectral.
4. calcul de $\hat{\vec{u}} = \overrightarrow{\widehat{Rot}\vec{\psi}}$.
5. retour dans l'espace physique pour obtenir $\vec{u}$, que l'on norme par le module moyen car la variable est adimensionnée.

Pour générer $\vec{\psi}_1$ et $\vec{\psi}_2$, nous cherchons un champ de vecteurs répartis avec une probabilité constante sur la sphère de rayon 1. Sur cette sphère l'élément de surface dS vaut $\sin(\theta)d\phi d\theta = d\phi d(-\cos(\theta))$, $\phi \in [0, 2\pi[$, $\theta \in [0, \pi[$. ϕ et $\cos(\theta)$ suivent des lois uniformes. La densité de probabilité du vecteur $\vec{x} = (\sin(\theta)\cos(\phi), \sin(\theta)\sin(\phi), \cos(\theta))$ est alors constante. On écrit :

$$\vec{\psi}(\vec{k}) = \left(\frac{E(k)}{k^2}\right)^{\frac{1}{2}} \left[\cos(\lambda(\vec{k})) \frac{\vec{\psi}_1(\vec{k})}{\|\widehat{Rot}\vec{\psi}_1(\vec{k})\|} + i \sin(\lambda(\vec{k})) \frac{\vec{\psi}_2(\vec{k})}{\|\widehat{Rot}\vec{\psi}_2(\vec{k})\|} \right].$$

C.3 Expression du spectre initial

L'expression du spectre initial retenue est celle proposée par plusieurs auteurs (*cf.* [VGKZ92], [DEL85]) :

$$E_{init}(k) = A.k^4.exp(\frac{-2k^2}{k_{peak}^2}). \quad (C.1)$$

Dans cette expression, k_{peak} représente le nombre d'onde pour lequel le spectre initial atteint sa valeur maximale, et A est une constante positive choisie de telle sorte que le champ de vitesse initial soit de moyenne RMS égale à 1. Pour les calculs effectués, k_{peak} a été pris égal à 3 [VGKZ92], la valeur de A étant par ailleurs fournie par la formule suivante :

$$\int E_{init}(k)dk = N^6 < u^2 > = A.\frac{3\sqrt{\pi}}{32\sqrt{2}}.k_{peak}^5$$

C.4 Méthode de calcul du spectre

Le spectre (sous entendu "le spectre d'énergie cinétique turbulente") est obtenu facilement à partir du champ de vitesse (*cf.* [McC92] par exemple). Néanmoins, son calcul pratique nécessitant un certain nombre d'opérations non triviales, il nous a paru intéressant de détailler ici comment celui-ci était effectué.

- Afin de réduire le coût de calcul de la transformée de Fourier du champ de vitesse, on transforme ce champ réel en un champ complexe. En dimension 1, cela revient par exemple à associer à un champ réel u défini sur $N = 2^n$ points, un champ complexe $V = V_r + iV_i$ défini sur $N/2$ points en posant :

$$\forall j : 1 \leq j \leq 2^{n-1} : \begin{cases} V_r(j) = u(2j-1) \\ V_i(j) = u(2j) \end{cases}$$

Le coût de la transformée de Fourier $\hat{V}$ de V est alors de l'ordre de $N \log_2(N/2)$, alors que celui de $\hat{u}$ est de l'ordre de $N \log_2 N$.

- Après avoir reconstruit $\hat{u}$ à partir de $\hat{V}$, on calcule le carré de la norme de ce vecteur pour tous les triplets $\{k_1, k_2, k_3\} \in [-N/2, N/2 - 1]^3$.
- Pour chaque nombre d'onde $|k|$ compris entre $-N\sqrt{d}/2$ et $(N/2 - 1)\sqrt{d}$ (où $d = 3$ est la dimension de l'écoulement considéré), on calcule la quantité $E(|k|)$ correspondant à l'énergie cinétique turbulente totale associée à $|k|$. Cette quantité est approchée par :

$$E(|k|) \sim \frac{4\pi |k|^2}{M} \sum_{k_1} \sum_{k_2} \sum_{k_3} \hat{u}(k_1, k_2, k_3)$$

où k_1, k_2 et k_3 sont tels que $(k_1^2 + k_2^2 + k_3^2)^{1/2} \in [|k| - 1/2, |k| + 1/2[$ et M est le nombre de triplets $\{k_1, k_2, k_3\}$ vérifiant cette condition.

- La fonction $E(|k|)$ est alors représentée en coordonnées log-log.

Annexe D

Equation equivalente : schéma de Vremann

La construction de l'équation équivalente au schéma de Vremann, pour une équation de diffusion-convection en dimension 3 d'espace, a été réalisée avec l'aide du logiciel de calcul formel MAPLE.

Bien qu'il eût été possible de programmer cette construction en MAPLE, nous avons néanmoins choisi d'utiliser ce logiciel de façon interactive (la session correspondante est donnée dans les pages qui suivent).

Pour en faciliter la lecture, précisons les notations qui y sont employées :

- Tout d'abord, l'équation traitée est la suivante :

$$\frac{\partial u}{\partial t} + (v \cdot \nabla) \cdot u - r \Delta u = 0$$

- Les variables d, e, f représentent les opérateurs de décalages suivant x, y, z respectivement, définis par :

$$\begin{aligned} d^n[u(x, y, z)] &= u(x + n.h, y, z) \\ e^n[u(x, y, z)] &= u(x, y + n.h, z) \\ f^n[u(x, y, z)] &= u(x, y, z + n.h) \end{aligned}$$

(on suppose que le pas du maillage est un réseau tridimensionnel de pas h).

- Les opérateurs cx, cy, cz représentent des approximations centrées des dérivées premières suivant x, y, z respectivement, obtenues sur des molécules de base $[-h, h]$ (cf. paragraphe 2.1).
- Les opérateurs ex, ey, ez représentent des approximations centrées des dérivées secondes suivant x, y, z respectivement, obtenues sur des molécules de base $[-2h, 2h]$ (cf. paragraphe 2.1).

- Les opérateurs gx, gy, gz représentent les opérateurs de filtrage dans les plans perpendiculaires aux directions x, y, z respectivement (cf. paragraphe 2.1).
- F désigne le flux numérique discret du schéma de Vremann, et G en est la représentation “continue” obtenue en remplaçant les opérateurs de décalage par leur développement de Taylor.
- Le résultat final est représenté par la variable `solution`, $P(i, j, k)$ représente le coefficient du terme $X^i Y^j Z^k$ dans cette expression, et $Q4$ correspond au terme d’ordre total égal à 4.
- $Q4V$ représente l’expression de $Q4$ obtenue pour les coefficients du schéma de Vremann ($a = 1/9$, $b = 1/36$, $c = 4/9$).

```

> cx:=1/(2*h)*(d-1/d);
> cy:=1/(2*h)*(e-1/e);
> cz:=1/(2*h)*(f-1/f);
> ex:=1/(4*h^2)*(d^2-2+1/d^2);
> ey:=1/(4*h^2)*(e^2-2+1/e^2);
> ez:=1/(4*h^2)*(f^2-2+1/f^2);
> gx:=a*(e+1/e+f+1/f)+b*(e+1/e)*(f+1/f)+c;
> gy:=a*(d+1/d+f+1/f)+b*(d+1/d)*(f+1/f)+c;
> gz:=a*(d+1/d+e+1/e)+b*(d+1/d)*(e+1/e)+c;
> F:=(gx*(v*cx-r*ex)+gy*(v*cy-r*ey)+gz*(v*cz-r*ez));
> G:=(X,Y,Z)->subs(f=exp(h*Z),subs(e=exp(h*Y),subs(d=exp(h*X),F)));
> G(X,Y,Z);
> resultX:=series(G(X,Y,Z),X=0,5);
> resultY:=series(convert(resultX,polynomial),Y=0,5);
> resultZ:=series(convert(resultY,polynomial),Z=0,5);
> solution:=convert(resultZ,polynomial);
> solution:=expand(solution);solution:=collect(solution,X);
> P:=(i,j,k)->coeff(coeff(coeff(solution,X,i),Y,j),Z,k);
> c:=1-4*(a+b);
> Q4:=simplify(P(4,0,0))*(X^4+Y^4+Z^4)+factor(P(2,2,0))*(X^2*Y^2+Y^2*Z^2+X^2*Z^2);
> Q4V:=subs(a=1/9,b=1/36,Q4);
>

```


Bibliographie

- [BAU93] G. BAUDIER. Simulation de la turbulence par la méthode des grandes echelles. Technical report, Commissariat à l'Energie Atomique CEA/DMA/MCN, 07/1993.
- [BER93] J.P. BERTOGLIO. Simulation numérique d'écoulements turbulents. Technical report, Ecole de Printemps à Carcans-Maubuisson, 05/1993.
- [BLC94] A. DE LA BOURDONNAYE, B. LARROUTUROU, and R. CARPENTIER. On the derivation of the modified equation for the analysis of linear numerical methods. Technical Report 94-26, Centre d'Enseignement et de Recherche en Modélisation, Informatique et Calcul Scientifique, 01/1994.
- [CL82] J.P. CHOLLET and M. LESIEUR. Parametrization of small scales of threedimensional isotropic turbulence utilizing spectral closures. *Journal of Atmosphere Science*, (38):2747–2757, 1982.
- [DEL85] P. DELORME. *Simulation Numérique en Dynamique de la Turbulence Homogène Compressible*. PhD thesis, Université de Poitiers, 10/1985.
- [HIR92] C. HIRSCH. *Fundamentals of Numerical Discretization*, volume 1 of *Numerical Computation of Internal and External Flows*. John Wiley & Sons, 1992.
- [KOL41a] A.N. KOLMOGOROV. The local structure of turbulence in incompressible viscous fluid for very large Reynolds numbers. *C.R. Acad. Sci. URSS*, 30, 1941.
- [KOL41b] A.N. KOLMOGOROV. On degeneration of isotropic turbulence in an incompressible viscous fluid. *C.R. Acad. Sci. URSS*, 32, 1941.
- [LES73] D.C. LESLIE. *Developments in the theory of turbulence*. Oxford Science Publications, 1973.
- [LEV90] R.J. LEVEQUE. *Numerical Methods for conservation laws*. Birkhäuser, 1990.

- [McC92] W.D. McCOMB. *The Physics of Fluid Turbulence*. Oxford Science Publications, 1992.
- [ML91] O. METAIS and M. LESIEUR. Spectral Large-Eddy Simulation of isotropic and stably-stratified turbulence. submitted to *Journal of Fluid Mechanics*, 1991.
- [NGML92] A. SILVEIRA NETO, D. GRAND, O. METAIS, and M. LESIEUR. A numerical investigation of the coherent structures of turbulence behind a backward-facing step. *Journal of Fluid Mechanics*, 1992.
- [PS95] P. POULLET and M. SANCANDI. Simulation de la turbulence par la méthode des grandes echelles. Technical Report 2796, Commissariat a l'Energie Atomique, CEL-V/DMA/MCN, 94195 Villeneuve St Georges, 1995.
- [ROY80] P. ROY. Résolution des équations de Navier-Stokes par un schéma de haute précision en espace et en temps. *La Recherche Aéronautique*, 6:373-385, 1980.
- [SAG] P. SAGAUT. Modèles de viscosité turbulente. Communication Personnelle, ONERA CERT.
- [VGKZ92] A.W. VREMAN, B.J. GEURTS, J.G.M. KUERTEN, and P.J. ZANDBERGEN. A finite volume approach to Large Eddy Simulation of compressible, homogeneous, isotropic, decaying turbulence. *Int. J. Numerical Methods in Fluids*, 15:799-816, 1992.
- [WH74] R.F. WARMING and B.J. HYETT. The modified equation approach to the stability and accuracy analysis of finite-difference methods. *Journal of Computational Physics*, 14:159-179, 1974.

Partie II

LE PRECONDITIONNEMENT PAR LES INCONNUES INCREMENTALES (I.U.)

Introduction

This part is divided into two chapters.

First, we are going to apply Incremental Unknowns (I.U.) as a multilevel preconditioner in order to solve several elliptic problems. Then, after reviewing briefly the theoretical formalism linked with the I.U. and the mathematical results obtained by M. CHEN and R. TEMAM in [CT91], we propose to give a numerical verification and extend some results concerning the condition number of the Laplacian matrix in the I.U. basis. Our experience has led us to study and develop an efficient classical multigrid algorithm for which, a comparison has been made with the I.U. method's implementation, on vectorial computers, for solving the two-dimensional Dirichlet problem. We also prove that it is interesting to solve in using the I.U., an elliptic problem with homogeneous Neuman boundary conditions obtained by a variational approach [CEA64]. For the three-dimensional Dirichlet case, even if we have a smaller creasing of the matrix condition number in the I.U. basis, one can find only good assumptions with a small number of I.U.'s grid levels.

The second chapter is devoted to a theoretical and numerical resolution of a class of two-dimensional nonlinear steady-state problems which depends on two positive parameters. One of them, ρ , enhances the linear operator character of the evolutive associate problem, and the other ν has the viscosity function. One can show that the solution exists in $H_0^1(\Omega)$ and is unique. According to the skew-symmetric property of the nonlinear operator, the proof is classically based on a Galerkin method and a fixed point theorem. Following this step, we give a discretized system with centered finite differences schemes. So, we decide to solve this problem with two different implicit class of methods: Newton-BiCGSTAB and nonlinear GMRES. For each one, we introduce the I.U. as a preconditioner and compare this algorithm with other preconditioned implementations. This study has been made with several ρ and ν values, and proves that the I.U. preconditioner becomes efficient when ρ is small. The second interesting point which is noted, asks to focus more in the preconditioner choice than the algorithm one.

Chapitre 1

Résolution de quelques problèmes elliptiques linéaires

Nous nous proposons dans ce chapitre de présenter une méthode de discrétisation multi-niveaux : les **Inconnues Incrémentales**. Elle se base sur une discrétisation des inconnues en différences finies, et, a la particularité de définir plusieurs grilles uniformes d'inconnues hiérarchiques. Nous allons l'appliquer dans le but de résoudre quelques problèmes elliptiques présentés dans la section 1.1, ensuite nous développerons une méthode multigrilles classique très efficace pour ce type de problèmes au paragraphe 1.2.1, et la comparerons avec l'implémentation incluant les Inconnues Incrémentales.

Ce travail peut être considéré comme une extension des travaux de [CT91], [CMT94] et aussi comme une étape préliminaire à la constitution d'une méthode multi-niveaux de résolution des équations de Navier-Stokes utilisant les Inconnues Incrémentales.

1.1 Problèmes modèles et discrétisations

Considérons deux modèles naturels d'équations elliptiques en dimension 2 (ou 3) :

- le problème de Dirichlet dans un rectangle (ou un parallélépipède rectangle) :

$$\left\{ \begin{array}{ll} - \Delta u = f & \text{dans } \Omega \\ u = 0 & \text{sur } \partial\Omega \end{array} \right. \quad (1.1)$$

- le problème de Neumann homogène dans un carré (ou un rectangle) :

$$\left\{ \begin{array}{ll} - \Delta u + u = f & \text{dans } \Omega \\ \frac{\partial u}{\partial n} = 0 & \text{sur } \partial\Omega \end{array} \right. \quad (1.2)$$

Dans le cas où Ω est un ouvert régulier borné, ces deux problèmes possèdent tous les deux une solution unique régulière. En appliquant le lemme de Lax-Milgram, on trouve la solution du premier problème dans $H_0^1(\Omega)$ et la deuxième dans $H^1(\Omega)$.

La méthode des différences finies est la plus vieille méthode de discrétisation spatiale basée sur des propriétés liées aux développements en série de Taylor et la définition de la dérivée.

La discrétisation du problème (1.1) peut se concevoir en une approche directe de l'opérateur de dérivation par un schéma centré aux différences. Ce qui donne l'opérateur discret comme la matrice penta-diagonale usuelle. Toutefois, elle peut aussi s'obtenir en utilisant une formulation variationnelle : pour le problème (1.2), la description détaillée de la méthode est donnée dans l'annexe A.

1.2 Présentation de deux méthodes multi-niveaux

1.2.1 Méthode Multigrilles

Descriptif

Depuis 1965, date à laquelle ont émergé ces méthodes, leur domaine d'application s'est vu grandir. Nous nous limiterons à la résolution d'équations aux dérivées partielles linéaires, cadre dans lequel leur grande efficacité n'est pas à mettre en doute. Dans ce paragraphe, nous présenterons un algorithme classique s'appliquant aux problèmes linéaires *Correction Storage* utilisant la stratégie de cycles V.

Pour simplifier les notations, décrivons le système discret de (1.1) pour $\Omega = [0, 1]^2$: Soient N un entier positif pair, on pose $h = \frac{1}{N}$, et le domaine Ω est approché par une grille uniforme de pas h :

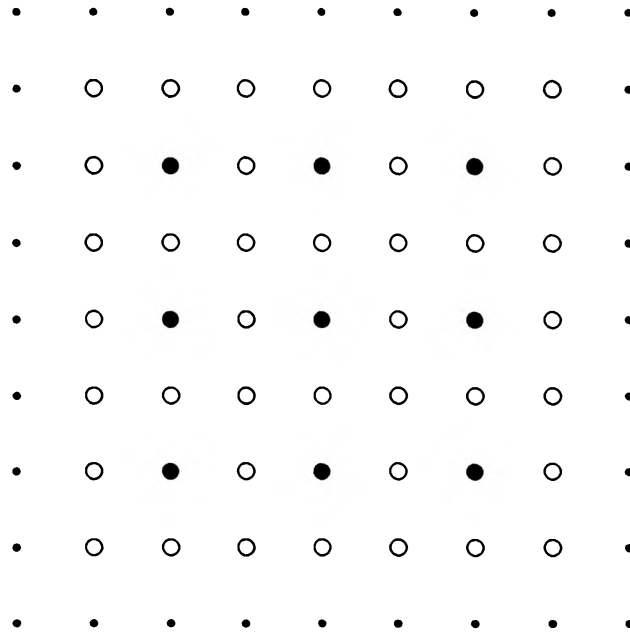
$$\begin{aligned}\Omega_h &= \{(i.h, j.h) \text{ où } 1 \leq i, j \leq N-1\}, \\ \partial\Omega_h &= \{(i.h, j.h) \text{ où } 0 \leq i \leq N \text{ et } j = 0 \text{ ou } N\} \\ &\quad \cup \{(i.h, j.h) \text{ où } 0 \leq j \leq N \text{ et } i = 0 \text{ ou } N\},\end{aligned}$$

u^h est la solution discrète cherchée définie sur $\Omega_h \cup \partial\Omega_h$,

$f_{i,j}^h$ représente $f(i.h, j.h)$,

et Δ_h , l'opérateur discret approchant le laplacien, soit :

$$-\Delta_h u_{i,j}^h = \frac{1}{h^2} (4u_{i,j}^h - u_{i+1,j}^h - u_{i-1,j}^h - u_{i,j+1}^h - u_{i,j-1}^h)$$

FIG. 1.1 - *Maillage constitué de deux grilles*

le système discret s'écrit :

$$\begin{cases} -\Delta_h u_{i,j}^h = f_{i,j}^h & 1 \leq i, j \leq N-1 \\ u_{i,j}^h = 0 & \forall i, j \text{ tels que } (i.h, j.h) \in \partial\Omega_h \end{cases} \quad (1.3)$$

En considérant deux grilles de discrétisation :

- Une grille grossière G_g de pas $(2h)$ constituée des $(\bullet)$
- Une grille fine G_f de pas (h) constituée des $(\bullet, \circ)$ (Fig. 1.1),

donnons la description de la méthode à deux grilles :

- 1) On effectue quelques itérations à l'aide d'une méthode itérative de base :

On obtient une certaine fonction u^h .

Soit u_{exacte}^h la solution exacte du système (1.3), si l'on pose $fe^h = u_{exacte}^h - u^h$ la fonction d'erreur, alors fe^h vérifie :

$$\begin{cases} -\Delta_h fe^h = f^h + \Delta_h u^h & \text{dans } \Omega_h, \\ fe^h = 0 & \text{sur } \partial\Omega_h. \end{cases}$$

2) On transfère le résidu sur G_g à l'aide de l'opérateur de restriction r_h^{2h} :

sur G_g , le résidu est en fait : $r_h^{2h}(f^h + \Delta_h u^h)$.

3) On résout le système (1.4) obtenu sur G_g pour calculer la fonction d'erreur.

$$\begin{cases} -\Delta_{2h} f e^{2h} = r_h^{2h}(f^h + \Delta_h u^h) & \text{dans } \Omega_h, \\ f e^{2h} = 0 & \text{sur } \partial\Omega_h. \end{cases} \quad (1.4)$$

4) On interpole cette fonction d'erreur obtenue sur G_g à G_f pour corriger la solution approchée sur G_f à l'aide de l'opérateur de prolongement p_{2h}^h . Cela s'effectue en calculant la nouvelle solution approchée :

$$u^h + p_{2h}^h(f e^{2h})$$

5) On effectue de nouveau quelques itérations à l'aide de la méthode itérative de base.

L'efficacité de ce type de méthodes vient du fait, qu'il n'est pas nécessaire de faire de nombreuses itérations de relaxation (en pratique moins de cinq itérations suffisent à chaque fois) et que le coût de ces itérations de relaxation du système (1.4) est très largement inférieur à celui nécessaire pour la résolution de (1.3).

Pour rechercher une méthode bien optimisée sur calculateurs vectoriels, nous avons utilisé la méthode itérative de base la plus courante : la méthode itérative de **Gauss-Seidel** avec un ordre **rouge-noir** (ou blanc-noir, ou damier) (cf. [HIR92]). C'est une technique de coloriage particulièrement efficace, ayant l'avantage d'avoir une implémentation très simple. Elle consiste à partager la grille fine en deux familles de points, les noirs et les blancs à la manière d'un damier, et ainsi, il suffit d'appliquer l'algorithme de Gauss-Seidel classique d'abord sur les noirs et ensuite sur les blancs (ou vice versa), pour rendre complètement vectorisable l'algorithme de Gauss-Seidel. En décomposant G_f comme un damier, on voit bien que pour (1.3), un point de la famille noire n'aura que des voisins blancs et vice versa, ce qui lève la dépendance informatique sur chacune des itérations pour une famille et donne le moyen d'augmenter le facteur de vectorisation du programme (cf. 1.2.3).

Pour le choix des opérateurs de prolongement et de restriction, la seule condition à assurer intervient sur l'ordre de ces opérateurs [HAC85]. En effet, le problème (1.1)

est d'ordre 2, et si l'on choisit un prolongement p d'ordre m_p et une restriction r d'ordre m_r , on doit vérifier que :

$$m_p + m_r > 2$$

En ce qui concerne l'opérateur de restriction, nous aurions pu choisir une injection simple ou une interpolation bi-linéaire (la liste est loin d'être exhaustive) mais, par souci d'efficacité nous avons utilisé l'opérateur "half-weighting" introduit par Trottenberg défini par le symbole suivant [HAC85]:

$$\text{sym}(r_h^{2h}) = \begin{bmatrix} 0 & \frac{1}{8} & 0 \\ \frac{1}{8} & \frac{1}{4} & \frac{1}{8} \\ 0 & \frac{1}{4} & 0 \end{bmatrix}$$

Le prolongement a été fait par l'opérateur "prolongement par neuf points" introduit par Wesseling (*cf.* [WES82]) défini par le symbole :

$$\text{sym}(p_{2h}^h) = \begin{bmatrix} \frac{1}{4} & \frac{1}{2} & \frac{1}{4} \\ \frac{1}{2} & 1 & \frac{1}{2} \\ \frac{1}{4} & \frac{1}{2} & \frac{1}{4} \end{bmatrix}$$

Ces deux opérateurs sont d'ordre 2.

On remarque de plus que le système (1.4) à résoudre sur la grille grossière G_g a exactement la même forme que celui sur G_f , ce qui nous autorise à appliquer la même stratégie sur G_g en utilisant (pour N entier positif multiple de 4) une grille encore plus grossière.

De manière générale, on pose $N = 2^k$ où $k \in \mathbb{N}^*$, et on peut définir jusqu'à k grilles d'inconnues de plus en plus grossière, le pas de discrétisation de la grille étant multiplié par 2 chaque fois.

Algorithme Correction Storage (CS)

Cas d'une itération de l'algorithme CS utilisant un cycle V.

étape 1 initialisation :

$$u_h = u_h^0$$

étape 2 on se place sur la grille G_h :

$$l = h,$$

étape 3 n_1 pré-relaxations sur G_l :

$$u_l \leftarrow Relax_l^{n_1}(u_l, f_l)$$

étape 4 on transfère le résidu de G_l à G_{2l} :

$$\begin{cases} f_{2l} \leftarrow r_l^{2l}(f_l - \Delta_l u_l) \\ u_{2l} \leftarrow 0 \\ l \leftarrow 2 \cdot l \end{cases}$$

étape 5 Si $l < \frac{1}{2}$ alors, on va à l'étape 3

étape 6 Résolution exacte sur la grille la plus grossière $G_{1/2}$:

$$u_{1/2} \leftarrow -\Delta_{1/2}^{-1} f_{1/2}$$

étape 7 Prolongement sur la grille immédiatement plus fine :

$$\begin{cases} u_{l/2} \leftarrow u_{1/2} + p_l^{l/2}(u_l) \\ l \leftarrow \frac{l}{2} \end{cases}$$

étape 8 n_2 post-relaxations sur G_l :

$$u_l \leftarrow Relax_l^{n_2}(u_l, f_l)$$

étape 9 Si $l > \frac{1}{N}$ alors, on va à l'étape 7

1.2.2 Méthode des Inconnues Incrémentales

Les Inconnues Incrémentales (connus sous l'acronyme anglais I.U.) ont été introduites en 1990 par R. TEMAM (cf. [TEM90]) dans le but initial d'approcher, en régime turbulent, des équations d'évolutions dissipatives sur de longs temps d'intégration.

La première étape de ce développement a consisté à s'intéresser aux équations elliptiques linéaires issues des problèmes de Dirichlet et problèmes de Neumann. Dans un souci de ne pas trop s'éloigner de la motivation principale, nous considérerons des discrétisations très fines ce qui nous amènera des résolutions de problèmes de grandes tailles.

Dans un premier temps, M. CHEN et R. TEMAM ont montré que l'implémentation de cette méthode était facile, comparable à celle de méthodes rapides de type multigrilles (cf. [CT91]). A l'aide d'un formalisme introduit dans [CT93] permettant de considérer la formulation variationnelle et les espaces appropriés, les auteurs ont justifié théoriquement le gain apporté par la nouvelle méthode pour les problèmes

bi-dimensionnels (cf. [CT93]). Ensuite, récemment dans [GOY94] et [CMT94], on peut voir l'extension au problème tri-dimensionnel de Dirichlet.

L'objectif de ce paragraphe est de présenter brièvement la méthode des Inconnues Incrémentales pour les problèmes bi-dimensionnels elliptiques et pour le problème tri-dimensionnel de Dirichlet, et de discuter des résultats numériques obtenus dans [CT91] et [CMT94].

Les I.U. pour les problèmes elliptiques en dimension 2

Il nous paraît nécessaire de rappeler que le conditionnement d'une matrice est le rapport entre la plus grande et la plus petite valeur propre de la matrice. De plus, l'ellipticité de nos problèmes nous garantit que toutes les valeurs propres des matrices issues des problèmes discrets sont des réels strictement positifs.

Considérons le système discret (1.3) issu du problème de Dirichlet (1.1). On rappelle que la résolution d'un système linéaire de grande taille nécessite généralement l'utilisation d'une méthode itérative. Citons une des plus répandue: *la méthode du Gradient Conjugué*, qui converge pour toute matrice dont la partie symétrique¹ est définie positive (cadre dans lequel nous nous situons pour nos problèmes discrets). Pour obtenir la convergence, elle demande un nombre d'opérations fonction du conditionnement κ de la matrice à inverser. Si l'on appelle A la matrice intervenant dans (1.3), alors on connaît bien ses valeurs propres et :

$$n_{\text{opérations}} \propto \sqrt{\kappa(A)}, \quad \kappa(A) \in \left[\frac{4}{\pi^2 h^2}, \frac{1}{h^2} \right] \quad (1.5)$$

Nous allons restreindre notre présentation au cas d'une grille d'I.U. pour des raisons de clarté.

Nous fixons un pas uniforme $h = \frac{1}{2N}$ ($N \in \mathbb{N}$), et nous décomposons notre maillage en une grille grossière d'inconnues aux noeuds $(2ih, 2jh)$ de pas $2h$ (notée $\mathcal{G}_g$) et la grille complémentaire d'inconnues appelée grille fine (notée $\mathcal{G}_f$). On peut donc réordonner le système discret (1.3) précédent, pour obtenir le système équivalent suivant (la matrice étant renotée A) :

$$A u = b \quad \text{où } u = \begin{pmatrix} u_g \\ u_f \end{pmatrix} \text{ et } b = \begin{pmatrix} b_g \\ b_f \end{pmatrix}, \quad (1.6)$$

avec u_g les inconnues aux points de la grille grossière, u_f les inconnues aux points de la grille complémentaire.

Nous appellerons $y = u_g$ les inconnues de grille grossière et, par une transformation particulière, nous considérerons les inconnues incrémentales (I.U.) z sur la grille complémentaire. Les I.U. seront les inconnues sur cette grille complémentaire auxquelles on aura oté la moyenne des inconnues voisines de la grille grossière. En fait ce procédé revient à un changement de variables linéaire.

1. la partie symétrique d'une matrice A est la matrice $\frac{1}{2}(A + {}^t A)$

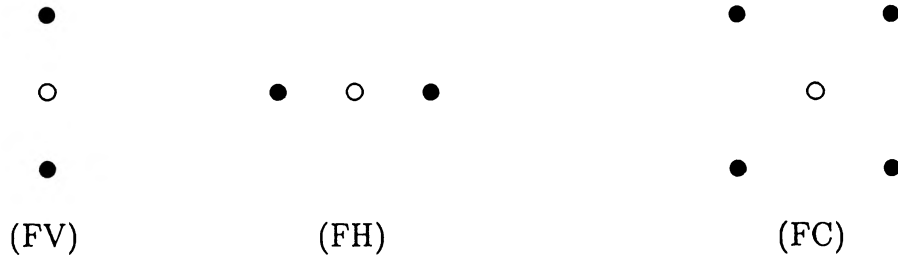


FIG. 1.2 - Familles d'I.U. en dim 2

Si l'on regarde le maillage de la Fig. 1.1, la structure multigrilles des I.U. correspond à une grille grossière constituée de ($\bullet$), et une grille fine constituée de ($\circ$).

Ce qui nous amène à regrouper les I.U. en trois familles : la famille verticale (FV), la famille horizontale (FH) et la famille centrale (FC) (voir Fig. 1.2). Ci-dessous, nous donnons leur définition :

(FV) constituée des points $(2ih, (2j+1)h)$, pour $1 \leq i \leq N-1$, $0 \leq j \leq N-1$

$$z_{2i,2j+1} = u_{2i,2j+1} - \frac{1}{2}(u_{2i,2j} + u_{2i+2,2j})$$

(FH) constituée des points $((2i+1)h, 2jh)$, pour $0 \leq i \leq N-1$, $1 \leq j \leq N-1$

$$z_{2i+1,2j} = u_{2i+1,2j} - \frac{1}{2}(u_{2i,2j} + u_{2i+2,2j})$$

(FC) constituée des points $((2i+1)h, (2j+1)h)$, pour $0 \leq i \leq N-1$, $0 \leq j \leq N-1$

$$z_{2i+1,2j+1} = u_{2i+1,2j+1} - \frac{1}{4}(u_{2i,2j} + u_{2i,2j+2} + u_{2i+2,2j} + u_{2i+2,2j+2})$$

En notant $\bar{u}_1 = \begin{pmatrix} y \\ z_1 \end{pmatrix}$, avec $z_1 = \begin{pmatrix} z_1^{(FH)} \\ z_1^{(FV)} \\ z_1^{(FC)} \end{pmatrix}$, on définit la matrice S_1 de changement de base (l'indice 1 désigne la seule grille d'I.U.) de la manière suivante :

$$u = \begin{pmatrix} u_g \\ u_f \end{pmatrix} = S_1 \bar{u}_1,$$

ce qui nous permet d'obtenir un nouveau système discret équivalent au précédent (cf. [CT93], [CEA64]):

$$\bar{A}_1 \bar{u}_1 = \bar{b}_1 \quad \text{avec} \quad \bar{A}_1 = {}^t S_1 A S_1 \quad \text{et} \quad \bar{b}_1 = {}^t S_1 b. \quad (1.7)$$

Cette construction peut s'étendre par récurrence à un nombre quelconque l de niveaux de grilles d'I.U.

En posant $h = h_l = \frac{1}{N^{2(l+1)}}$ (où N correspond au nombre d'inconnues sur la grille la plus grossière), on nomme, $z_1, z_2, \dots, z_l$ les inconnues de chacune des grilles d'I.U. ($\mathcal{G}_1, \mathcal{G}_2, \dots, \mathcal{G}_l$) auxquelles on a oté les points formant les grilles plus grossières, soit :

$$\mathcal{G}_g = G_{h_0}, \mathcal{G}_1 = G_{h_1} \setminus G_{h_0}, \dots, \mathcal{G}_l = G_{h_l} \setminus G_{h_{l-1}},$$

on obtient à nouveau, le système discret suivant, équivalent à (1.4):

$$\bar{A}_l \bar{u}_l = \bar{b}_l \quad \text{avec} \quad \bar{A}_l = {}^t S_l A S_l \quad \text{et} \quad \bar{b}_l = {}^t S_l b. \quad (1.8)$$

Formalisme mathématique

Nous verrons dans ce paragraphe qu'il est possible d'utiliser le formalisme des éléments finis hiérarchiques Q_1 pour poser le problème mathématique de la hiérarchisation par les I.U. (cf. [TEM90], [CT93], et pour les éléments finis: [YSE86]).

Soit $h_0 (\in \mathbb{R}^+)$ le pas grossier uniforme de discrétisation, $h_l = \frac{h_0}{2^l}$ (l représentant le nombre de grilles d'I.U.), on définit $\forall 0 \leq k \leq l$:

$$K_k^{i,j} = [i.h_k, (i+1).h_k] \times [j.h_k, (j+1).h_k]$$

$$\mathcal{U}_k \text{ l'ensemble des points } (ih_k, jh_k) \in \bar{\Omega}$$

V_k est l'espace de polynômes réels Q_1 (de degré 1 dans chacune des directions), dans chacun des rectangles $K_k^{i,j} \subset \Omega$, définis uniquement par les valeurs aux nœux de $\mathcal{U}_k$.

On observe sans peine que les espaces V_k sont emboîtés :

$$V_0 \subset V_1 \subset \dots \subset V_l$$

Suite à la relation d'emboîtement, on introduit l'espace W_l comme supplémentaire de V_{l-1} dans V_l :

$$V_l = V_{l-1} \oplus W_l, \quad (1.9)$$

et ainsi, on définit W_l comme l'espace de fonctions z_l s'annulant sur $\mathcal{U}_{l-1}$, autrement dit,

$$\forall u \in V_l, \exists ! (y, z_l) \in V_{l-1} \times W_l \text{ t.q. } u = y + z_l.$$

Donc, en réitérant la décomposition, on obtient :

$$V_l = V_0 \oplus W_1 \oplus W_2 \oplus \cdots \oplus W_l. \quad (1.10)$$

Après avoir remarqué qu'un rectangle K_l est obtenu en découpant en 4 un rectangle K_{l-1} , on retrouve les mêmes formules de définition des I.U. que précédemment.

En introduisant l'opérateur d'interpolation linéaire r_l associant à chaque $u \in \mathcal{C}(\bar{\Omega})$, $r_l u \in V_l$ définie par :

$$r_l u(M) = u(M), \quad \forall M \in \mathcal{U}_l,$$

on peut décomposer chaque fonction de $\mathcal{C}(\bar{\Omega})$ dans les espaces d'I.U. fidèlement à (1.10) :

$$u = r_l u = r_0 u + \sum_{k=1}^l (r_k u - r_{k-1} u).$$

Une équivalence de normes nous donne une estimation du conditionnement de la matrice issue de (1.8) [CT93].

Théorème 1.2.1

Si A est la matrice penta-diagonale usuelle de discrétisation du laplacien pour le problème de Dirichlet dans un rectangle, $\bar{A}_l = {}^t S_l A S_l$ celle issue du système (1.8) pour l grilles d'I.U. et un point sur $\mathcal{G}_g$ alors,

$$\exists C \in \mathbb{R}^+ \text{ t.q. } \forall l \in \mathbb{N}^* \quad \kappa(\bar{A}_l) \leq C.(l+1)^2 \quad (1.11)$$

La preuve de ce résultat est donnée dans [CT93].

D'après (1.5) et le théorème 1.2.1, et en comparant les croissances des fonctions $x \mapsto x^2$ et $x \mapsto \log^2(x)$, on espère bien trouver une discrétisation fine à partir de laquelle, la méthode de gradient conjugué appliquée au système (1.8) converge en moins d'itérations que celle appliquée au système d'origine (1.1) ou (1.6).

Il est important de noter que l'intérêt premier de la structure multigrilles réside dans une hiérarchisation des inconnues. En effet, pour une fonction suffisamment régulière on peut facilement montrer en utilisant les formules de Taylor, que les I.U. sont petites, d'ordre 2 du pas de la grille dont elles font partie. De ce fait, des auteurs ont défini d'autres I.U. ayant l'avantage d'être beaucoup plus petites, i.e. d'ordre plus élevé, mais souffrant d'une lourdeur dans leur définition.

Il convient aussi de mentionner qu'il existe d'autres I.U. appelées Inconnues Incrémentales Oscillantes introduites par [CT91] qui ont l'avantage d'être définies par une somme directe orthogonale (au lieu d'une somme directe comme en (1.9)) [MED95], ce qui nous montre bien que les I.U. ne peuvent faire partie inhérente de la même classe de méthodes que les bases hiérarchiques.

Les I.U. pour les problèmes elliptiques en dimension 3

Nous allons dans ce paragraphe décrire l'algèbre associée à l'extension des I.U. en dimension trois.

Le problème modèle traité est le problème de Dirichlet homogène dans le cube unité (1.1).

On discrétise l'opérateur discret par le schéma centré usuel à 7 points suivant :

$$\begin{aligned} -\Delta_h u = & \frac{1}{h_x^2}(2u_{i,j,k} - u_{i-1,j,k} - u_{i+1,j,k}) \\ & + \frac{1}{h_y^2}(2u_{i,j,k} - u_{i,j-1,k} - u_{i,j+1,k}) \\ & + \frac{1}{h_z^2}(2u_{i,j,k} - u_{i,j,k-1} - u_{i,j,k+1}) \end{aligned}$$

avec h_x le pas dans la direction des x , h_y celui dans la direction des y et h_z celui dans la direction des z .

Pour plus de clarté, nous cantonnerons notre présentation au cas d'une grille d'I.U.

Les points formant la grille grossière ($\mathcal{G}_g$) sont constitués des noeuds du maillage de la forme $(2i, 2j, 2k)$, auxquelles on associe les inconnues nodales $y_{2i,2j,2k}$. Sur la grille fine ($\mathcal{G}_f$) on définit sept types d'I.U., suivant la disposition des noeuds dans le cube (voir Fig. 1.3) :

- Type IC : les noeuds de la forme $(2i + 1, 2j, 2k)$,
- Type JC : les noeuds de la forme $(2i, 2j + 1, 2k)$,
- Type KC : les noeuds de la forme $(2i, 2j, 2k + 1)$,
- Type IF : les noeuds de la forme $(2i, 2j + 1, 2k + 1)$,
- Type JF : les noeuds de la forme $(2i + 1, 2j, 2k + 1)$,
- Type KF : les noeuds de la forme $(2i + 1, 2j + 1, 2k)$,
- Type CC : les noeuds de la forme $(2i + 1, 2j + 1, 2k + 1)$.

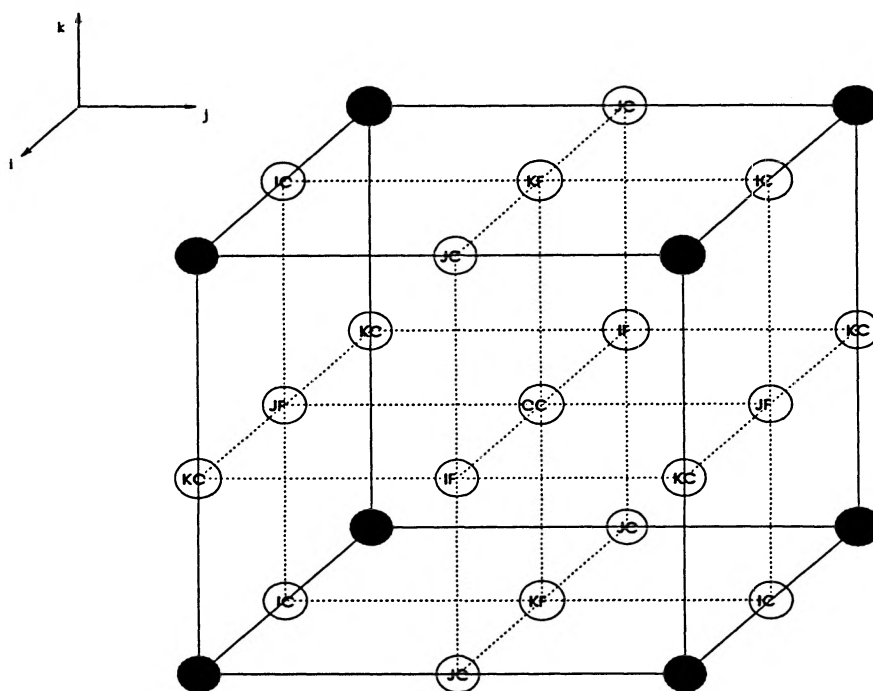


FIG. 1.3 - Grille grossière et grille fine

Les points du type IC sont les milieux des arêtes, parallèles à l'axe des x, reliant des noeuds de (G_g) , la quantité incrémentale associée, est :

$$z_{2i+1,2j,2k} = u_{2i+1,2j,2k} - \frac{1}{2}(u_{2i,2j,2k} + u_{2i+2,2j,2k})$$

Les points du type JC sont les milieux des arêtes, parallèles à l'axe des y, reliant des noeuds de (G_g) , la quantité incrémentale associée, est :

$$z_{2i,2j+1,2k} = u_{2i,2j+1,2k} - \frac{1}{2}(u_{2i,2j,2k} + u_{2i,2j+2,2k})$$

Les points du type KC sont les milieux des arêtes, parallèles à l'axe des z, reliant des noeuds de (G_g) , la quantité incrémentale associée, est :

$$z_{2i,2j,2k+1} = u_{2i,2j,2k+1} - \frac{1}{2}(u_{2i,2j,2k} + u_{2i,2j,2k+2})$$

Les points du type IF sont les centres des faces, parallèles au plan (y,z), reliant des noeuds de (G_g) , la quantité incrémentale associée, est :

$$z_{2i,2j+1,2k+1} = u_{2i,2j+1,2k+1} - \frac{1}{4}(u_{2i,2j,2k} + u_{2i,2j+2,2k} + u_{2i,2j,2k+2} + u_{2i,2j+2,2k+2})$$

Les points du type JF sont les centres des faces, parallèles au plan (x,z), reliant des noeuds de (G_g) , la quantité incrémentale associée, est :

$$z_{2i+1,2j,2k+1} = u_{2i+1,2j,2k+1} - \frac{1}{4}(u_{2i,2j,2k} + u_{2i+2,2j,2k} + u_{2i,2j,2k+2} + u_{2i+2,2j,2k+2})$$

Les points du type KF sont les centres des faces, parallèles au plan (x,y), reliant des noeuds de (G_g) , la quantité incrémentale associée, est :

$$z_{2i+1,2j+1,2k} = u_{2i+1,2j+1,2k} - \frac{1}{4}(u_{2i,2j,2k} + u_{2i+2,2j,2k} + u_{2i,2j+2,2k} + u_{2i+2,2j+2,2k})$$

Les points du type CC sont les centres des cubes reliant les noeuds de (G_g) , la quantité incrémentale associée, est :

$$z_{2i+1,2j+1,2k+1} = u_{2i+1,2j+1,2k+1} - \frac{1}{8}(u_{2i,2j,2k} + u_{2i+2,2j,2k} + u_{2i,2j+2,2k} + u_{2i+2,2j+2,2k} + u_{2i,2j,2k+2} + u_{2i+2,2j,2k+2} + u_{2i,2j+2,2k+2} + u_{2i+2,2j+2,2k+2})$$

Comme en dimension 2, on peut définir par récurrence, plusieurs grilles d'I.U. De même, en gardant les mêmes notations matricielles que précédemment, résoudre le problème (1.1), c'est résoudre le système suivant:

$$\bar{A}_l \bar{u}_l = \bar{b}_l \quad \text{avec} \quad \bar{A}_l = {}^t S_l A S_l \quad \text{et} \quad \bar{b}_l = {}^t S_l b. \quad (1.12)$$

où le vecteur $\bar{u}_l$ est constitué les inconnues y et des inconnues z_k rangées dans l'ordre suivant : d'abord les y dans l'ordre lexicographique, ensuite les z_k , dans l'ordre croissant des grilles et par type pour chacune des grilles k .

En étendant le formalisme mathématique du paragraphe précédent à la dimension 3, nous avons une estimation du conditionnement de la matrice issue de (1.12).

Théorème 1.2.2

Si A est la matrice hepta-diagonale usuelle de discrétisation du laplacien pour le problème de Dirichlet dans un parallélépipède rectangle, $\bar{A}_l = {}^t S_l A S_l$ celle issue du système (1.12) pour l grilles d'I.U. alors,

$$\exists C \in \mathbb{R}^+ \text{ t.q. } \forall l \in \mathbb{N}^* \quad \kappa(\bar{A}_l) \leq C \cdot \frac{l^4}{h_l} \quad (1.13)$$

La preuve de ce résultat est donnée dans [?].

1.2.3 Résultats numériques**Problème bi-dimensionnel de Dirichlet**

Nous présentons les résultats numériques obtenus pour le premier problème modèle (Dirichlet 2D). Nous comparerons plusieurs méthodes de résolution, dont les méthodes utilisant les I.U., la méthode multigrilles classique CS et deux autres pré-conditionnements de gradient conjugué plus anciens mais couramment utilisés, SSOR d'Evans et Choleski incomplet IC(0).

Ce travail a été réalisé en collaboration avec O.GOYON [GOY94] au Laboratoire d'Analyse Numérique d'Orsay.

L'étude comparative a été menée sur calculateurs vectoriels. Nous avons utilisé le TITAN de *Stardent* du L.A.N.O.R.S.² pour la mise au point, et les résultats qui suivent ont été obtenus sur le CRAY YMP/EL du C.R.I.³ d'Orsay.

Tout d'abord nous avons effectué le même test que celui proposé par M. CHEN et R. TEMAM dans [CT91]. Il s'agit d'une résolution du problème discret, dans le carré unité, ayant pour solution 0 et une condition initiale égale à la fonction : $u_0(x, y) = \sin(10^6 \pi xy(x-1)(y-1))$. Cette stratégie nous permet de tester la robustesse de la méthode. En effet, nous pouvons considérer un test de convergence très précis sur le résidu relatif du gradient : $\frac{\|r\|}{\|r_0\|} \leq 10^{-10}$.

Prenons une discrétisation fine de 513 points dans chaque direction, et comparons la décroissance du résidu relatif des deux méthodes multi-niveaux (Fig. 1.4).

2. Laboratoire d'Analyse Numérique d'Orsay

3. Centre de Ressources en Informatique

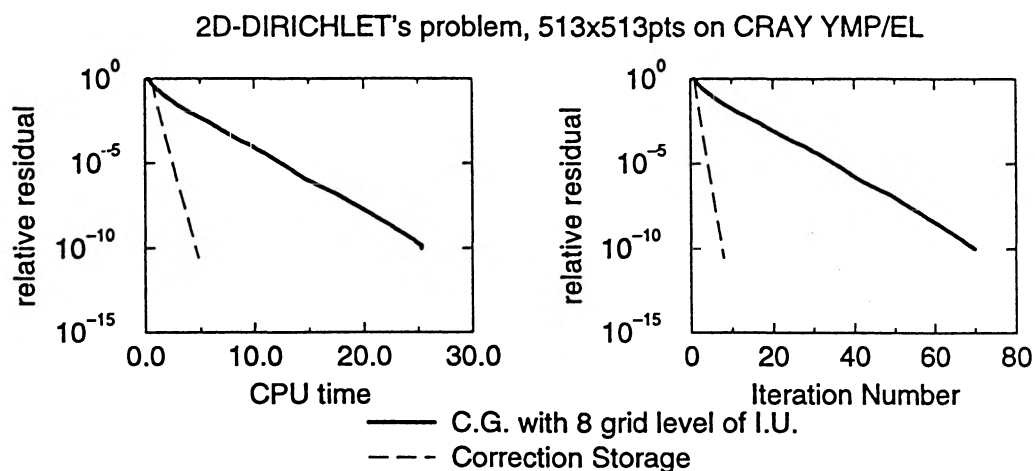


FIG. 1.4-Comparaison des performances de deux méthodes multi-niveaux

Ces courbes de convergence attestent bien une décroissance régulière du résidu en fonction des itérations. Il est évident tout de même que la méthode multi-niveaux la plus efficace est multigrilles Correction Storage. Les réalisations informatiques ont été vectorisées totalement. Pour le montrer, après avoir défini le *facteur de vectorisation* de chacun des programmes (cf. [PAS92]) comme le rapport du temps CPU du programme en scalaire sur le temps CPU du code de calcul vectorisé, nous le comparons pour chacune des méthodes et en faisons une analyse détaillée par la suite.

	Conj. Gradient	I.U. + C.G.	C.S. Multigrid
CPU time in s. with Vector.	177	25	5
CPU time in s. without Vector.	1770	181	65
It. number	1016	70	8
Vectorization factor	10	7.2	13
CPU/iter	0.17	0.36	0.62

Ce tableau nous donne plusieurs informations très intéressantes.

La première concerne, bien sur, l'efficacité de l'implémentation de notre algorithme multigrilles Correction Storage. Il s'agit d'une méthode extrêmement performante ayant le plus grand facteur de vectorisation; ce qui est dû au choix des opérateurs de restriction et de prolongement. Notre méthode des I.U. améliore tout de même très nettement la complexité et le temps calcul de la résolution. Mais elle demeure nettement moins performante que la méthode multigrilles C.S.

La deuxième concerne la très bonne vectorisation de toutes les implémentations (les facteurs de vectorisation sont tous très grands). On remarque que le fait d'introduire les I.U. dans la méthode de gradient conjugué diminue l'efficacité de la vectorisation. Ce qui s'explique facilement par l'implémentation du produit $(^tS_l A S_l).x$ pour un nombre l de grilles maximal (i.e. avec un seul point sur la grille grossière). En effet, notre réalisation a été guidée par deux objectifs un peu opposés : pouvoir traiter des discrétisations très fines et obtenir le programme le plus performant sur super-calculateurs vectoriels. Le fait d'inverser de grands systèmes nous oblige à simuler nos produits matrice-vecteur, i.e ne nous permet pas de stocker la matrice $\bar{A}_l$. Il suffit de voir que quelque soit le type de stockage utilisé, pour 513^2 points de discrétisation, il est pour le moment impossible de stocker cette matrice symétrique dont la structure est beaucoup trop complexe. Par conséquent, faire successivement, à chaque itération de gradient, les produits $S, ^tS$ par un vecteur nous fait perdre un temps précieux, causé par une vectorisation non efficace sur les premiers niveaux de grilles d'I.U. Avec un point sur la grille grossière, on a 8 z_1 , 40 z_2 , 176 z_3 , l'architecture des super-calculateurs CRAY utilise des registres vectoriels de longueur 64, d'où une efficacité croissante pour des tailles de vecteurs seulement supérieures à la centaine (on parle du temps d'amorçage du pipeline [GL90]).

Regardons aussi les méthodes de gradient conjugué utilisant des préconditionnements plus classiques.

	Conj. Gradient	I.U. + C.G.	SSOR + C.G.	IC(0) + C.G.
CPU time in s. with Vector.	177	25	293	864
CPU time in s. without Vector.	1770	181	546	1487
Iter. number	1016	70	207	425
Vectorization rate	10	7.2	1.8	1.7
CPU/iter	0.17	0.36	1.4	2.0

La factorisation incomplète de Choleski ainsi que la sur-relaxation SSOR d'Evans souffre du même problème. Elles sont toutes deux pénalisées par une inversion intermédiaire réalisée par réduction en deux systèmes triangulaires. Ceci se fait par un calcul récursif donc difficilement vectorisable. Cependant, il existe dans la littérature d'autres formes de factorisation de Choleski, comme celle proposée par VAN DER VORST dans [VOR82], ou encore celles obtenues par décomposition de domaines dont nous nous sommes écarté. Bien que nous n'ayons pas fait l'éventail de tous les préconditionnements de la méthode de gradient conjugué, nous pouvons tout de même classer nos I.U. comme bon préconditionneurs du gradient conjugué sur calculateurs vectoriels.

Continuons notre étude en nous interrogeant sur l'intérêt de considérer un nombre maximal de niveaux de grilles d'I.U. Il suffit de faire varier le nombre de points sur

Dans le tableau qui suit, nous présentons le nombre d'itérations de gradient conjugué nécessaire à la convergence.

Nbre de grilles	1	2	3	4	5	6	7	8	9
Nbre d'itérations	1016	464	240	128	84	71	68	70	70

Nous obtenons les résultats escomptés, le nombre d'itérations varie bien de la même manière que le conditionnement de $\tilde{A}_k$, avec un nombre d'itérations minimal pour un nombre de grilles quasiment le plus grand.

Problème bi-dimensionnel de Neuman

Pour le problème modèle (Neuman 2D), notre étude s'est limitée juste à la vérification du bon comportement lié à l'introduction des I.U. escomptée dans [CT93].

2D-NEUMANN's problem, 257x257pts on CRAY YMP/EL

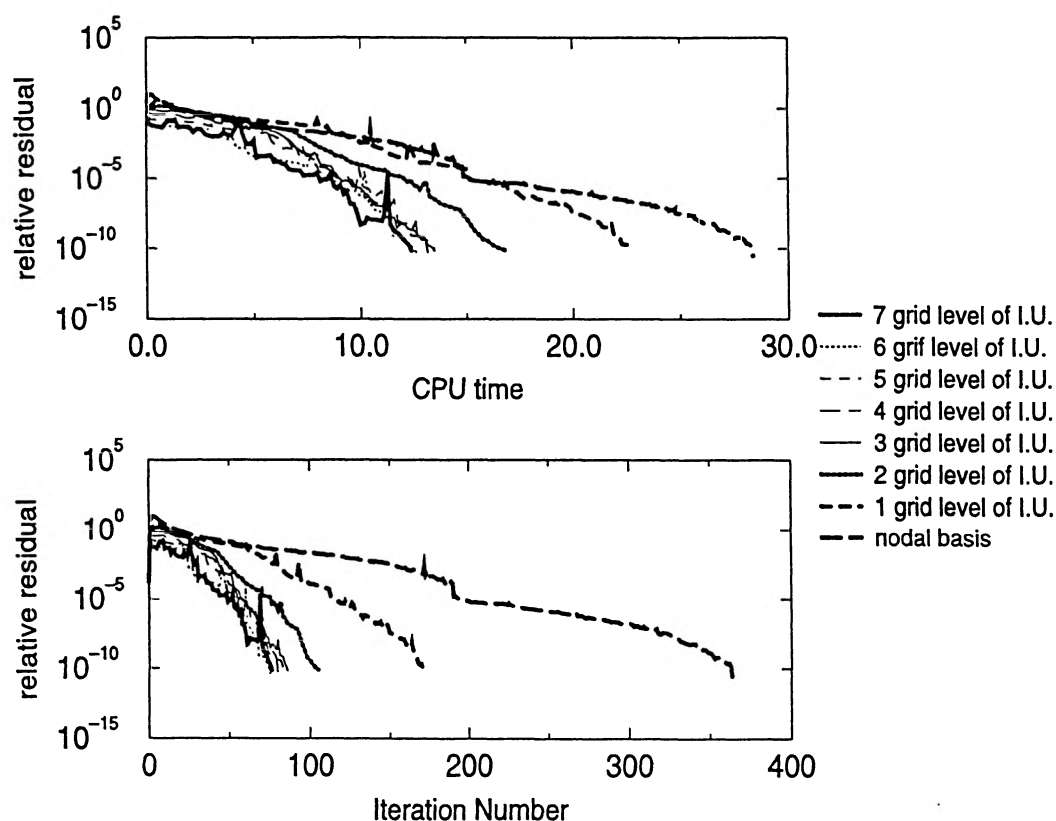


FIG. 1.6-Performance de la résolution en I.U. pour différentes grilles grossières

Nous arrivons à réduire d'un facteur 2 le temps de calcul et d'un facteur 5 le nombre d'itérations. Ces premiers résultats sont encourageants et méritent d'être un peu

approfondis. A mesure de l'augmentation du nombre de grilles d'I.U., l'efficacité de la méthode augmente de la même manière que pour le problème de Dirichlet. On observe tout de même beaucoup d'oscillations dans la courbe de décroissance du résidu de notre méthode multi-niveaux. Pour compléter cette étude, il serait intéressant d'utiliser d'autres méthodes de gradient conjugué comme MINRES ou BI-CGSTAB en vue de diminuer les oscillations [VOR92].

Problème tri-dimensionnel de Dirichlet

Les résultats numériques présentés dans ce paragraphe ont été réalisés sur le CRAY C98 d'I.D.R.I.S.⁴

Pour la résolution du système, notre choix a porté sur une méthode de gradient : BI-CGSTAB méthode que nous utiliserons plus tard pour sa robustesse et que nous détaillerons par la suite (Voir II.2.2.2). Nous présentons des résultats numériques obtenus pour une discrétisation de 257^3 points, soit 17 millions d'inconnues.

Pour déterminer les performances effectives d'un code de calcul, CRAY a développé un outil d'analyse : Jump Trace (simulateur de l'H.P.M.⁵). On obtient les performances directement en Megaflops⁶.

La puissance maximale théorique par processeur du C98 est de 1000Mflops, et l'on considère qu'un programme est bien optimisé s'il dépasse 200Mflops.

Dans le cas d'un maillage de 257^3 points, les performances des sous-routines sont de:

800 Mflops pour le sous-programme implémentant BI-CGSTAB dans les deux cas, 580 Mflops pour le produit matrice A par un vecteur et 370 Mflops pour l'autre produit $\tilde{A}_l$ par un vecteur.

A la lecture des performances, on remarque que notre sous-programme implémentant BI-CGSTAB est parfaitement optimisé. Comme nous l'avons remarqué en dimension 2, de part la nature récursive grille par grille des opérations en I.U., le produit faisant intervenir $\tilde{A}_l$ est moins efficace que celui dans la base nodale. De plus, par l'introduction des matrices tS et S , une itération classique de BI-CGSTAB coûte en moyenne 1.31s., alors qu'une itération de BI-CGSTAB en I.U. coûte entre 3.12s. pour 1 niveau d'I.U. et 3.6s. pour 7 niveaux d'I.U.

On présente aux Fig. 1.7, 1.8 les résultats numériques pour des maillages de 129^3 et 257^3 points, obtenus avec la version classique et celle faisant intervenir les I.U.

Si l'on regarde la Fig. 1.7, on observe que l'introduction des I.U. améliore le nombre d'itérations et que c'est pour deux niveaux d'I.U. que le meilleur résultat est obtenu.

Comme nous l'avons déjà précisé, le coût d'une itération augmente avec le nombre de

4. Institut de Développement et des Ressources en Informatique Scientifique

5. Hardware Performance Monitor

6. unité d'optimisation de programme informatique: Million floating operation per second

niveaux de grilles d'I.U., ce qui peut être interprété à la lecture des deux graphiques de la Fig. 1.7. On remarque de nouveau que le meilleur temps CPU est obtenu avec deux niveaux d'I.U.

3D-DIRICHLET'S PROBLEM with $129 \times 129 \times 129$ points

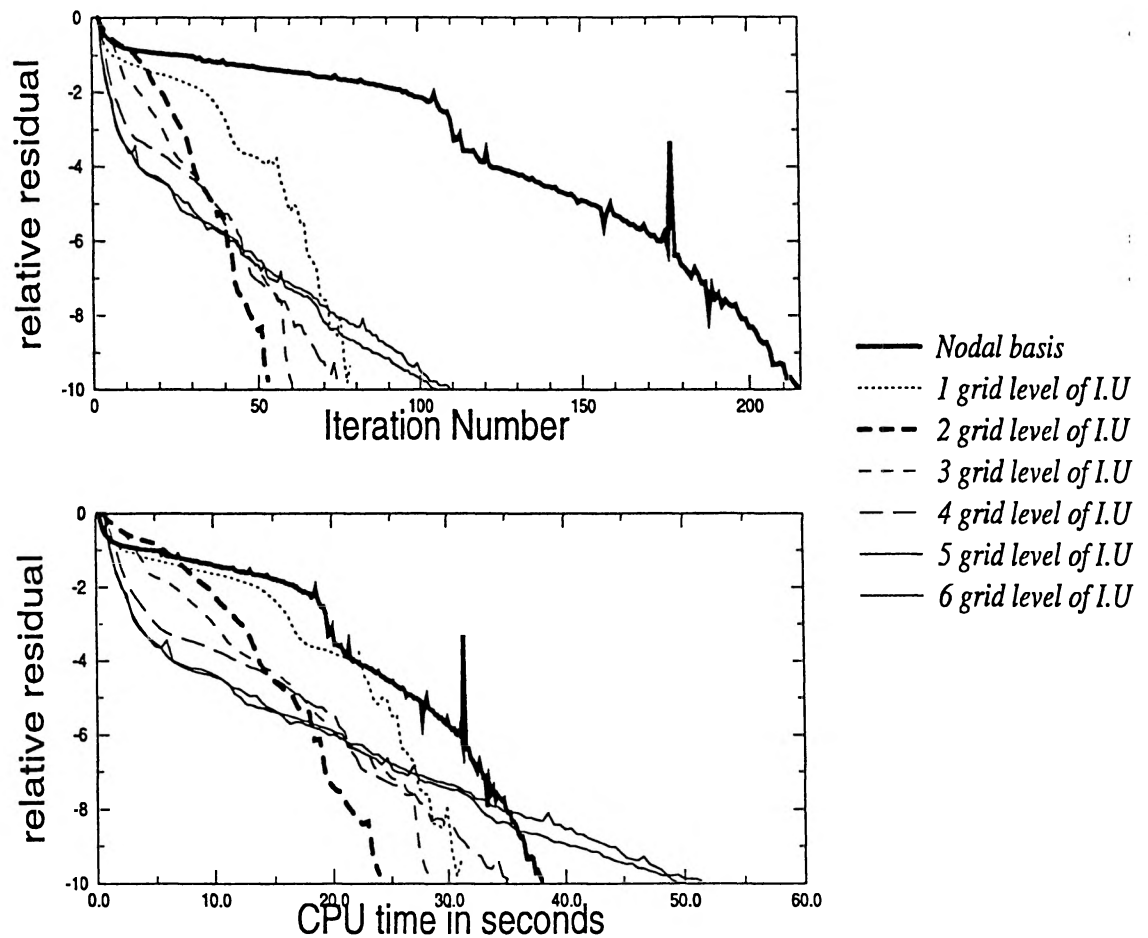


FIG. 1.7 - Résultats avec BI-CGSTAB sur un maillage 129^3 points

Pour la résolution avec une discrétisation de 257^3 points, la version de l'implémentation la plus efficace est obtenue pour trois niveaux d'I.U.

3D-DIRICHLET'S PROBLEM with 257^3 points

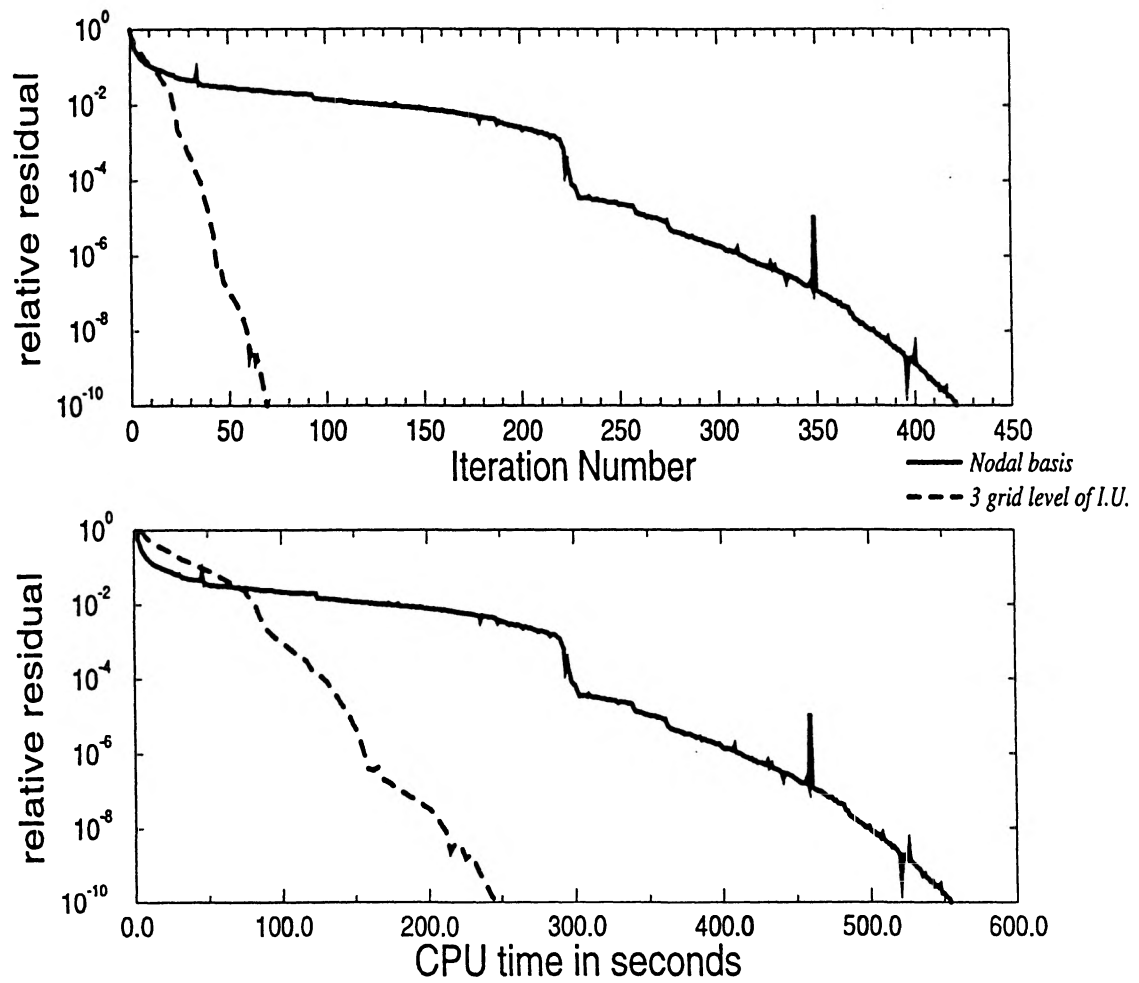


FIG. 1.8 - Résultats avec BI-CGSTAB sur un maillage 257^3 points

1.3 Conclusions

Nous avons montré que les Inconnues Incrémentales classiques étaient un bon candidat pour préconditionner les matrices intervenant dans la résolution de problèmes elliptiques. Toutefois, nous avons remarqué qu'elles n'y parvenaient pas avec la même efficacité qu'une méthode multigrilles classique. En effet, une méthode multigrilles combine l'utilisation de deux interpolateurs et d'une méthode de relaxation ayant la caractéristique de capter très rapidement les grandes structures (coefficients des petits modes de Fourier) de la solution cherchée. Le fait de restreindre le nombre de points de discrétisation atténue les coefficients des grands modes de Fourier et, améliore la précision de la nouvelle solution pour un même nombre d'itérations de relaxation. Ce qui explique qu'il suffise de peu d'itérations de relaxation sur chacune des grilles et assure l'efficacité de la méthode. Tandis que nos I.U., faisant intervenir deux interpolateurs symbolisées par S et $'S$, font appel à une méthode de gradient conjugué et donc permettent d'avoir une résolution uniquement rythmée par la valeur du conditionnement de la matrice à inverser.

Ce résultat comparatif est en accord avec les études menées sur les préconditionneurs multi-niveaux [BPX90]. Nous obtenons aussi des conditionnements pour le problème de Dirichlet-2D du même ordre que ceux obtenus pour les matrices intervenant lors de la discrétisation par éléments finis hiérarchiques [YSE86].

Une étude assez fine portant sur la vectorisation des codes de calcul met en évidence qu'il n'est pas nécessaire de restreindre la grille grossière à un seul point de discrétisation pour avoir la plus grande efficacité. Toutefois, la version maximisant le nombre de grilles d'I.U. permet de s'approcher de l'implémentation ayant le plus faible coût.

L'extension des I.U. aux problèmes tri-dimensionnels confirme leur aspect préconditionneur, mais à cause du grand nombre d'I.U., la complexité du changement de bases pénalise la version utilisant un grand nombre de grilles d'I.U. Nous nous sommes rendus compte que pour 129 points de discrétisation dans une direction, deux grilles d'I.U. représente déjà plus de 98,5% des inconnues du système, ce qui prouve qu'augmenter le nombre de niveaux risque surtout d'alourdir le coût d'une itération (en augmentant la récursivité) et pas fondamentalement d'améliorer l'efficacité de la résolution.

Chapitre 2

Résolution d'un problème non linéaire stationnaire

Cette étude permet d'appliquer le préconditionnement par les Inconnues Incrémentales à la résolution d'un problème non linéaire stationnaire.

Dans le premier paragraphe, après avoir présenté l'équation de type Navier-Stokes que nous considérons, nous donnons un résultat d'existence et d'unicité, et terminons en donnant le système discret.

Dans les deux paragraphes qui suivent, nous présentons les deux méthodes implicites de résolution que nous utilisons. Le premier constitue une étude comparative de plusieurs préconditionnements pour la méthode de Newton-BICGSTAB. Le deuxième reprend l'étude pour *Generalized Minimum RESidual method*¹. Pour chacune de ces méthodes, nous avons testé un nouveau préconditionnement : les Inconnues Incrémentales (IU) (cf. Chap. 1).

Nous avons effectué des comparaisons de la décroissance du résidu de Newton (en fonction du temps CPU nécessaire) pour quatre implémentations de la méthode de Newton-BICGSTAB :

- le préconditionnement par les Inconnues Incrémentales $\rightarrow IUNEXT$
- le préconditionnement par Multigrilles Correction Storage $\rightarrow MGNEXT$
- le préconditionnement par l'approximation polynômiale de l'inverse de la Jacobienne $\rightarrow P^nNEXT$
- Newton sans Préconditionnement $\rightarrow NEXT$

Pour la méthode de GMRES, trois versions ont été testé :

- le préconditionnement par les Inconnues Incrémentales $\rightarrow IUGMRES$
- le préconditionnement par Multigrilles Correction Storage $\rightarrow MGGMRES$

1. méthode du résidu minimal généralisé, plus connu sous son acronyme anglais (GMRES)

– GMRES sans Préconditionnement $\rightarrow$ *GMRES*

Dans les deux cas, nous avons considéré aussi le Résidu en fonction du Nombre d'itérations (ou nombre d'évaluations de la Fonctionnelle) nécessaires pour y parvenir.

Nous terminons cette étude en donnant des résultats numériques et une conclusion.

Les résultats de ce chapitre dans le cas particulier $\rho = 0$, ainsi que d'autres extensions de schémas multi-niveaux utilisant les Inconnues Incrémentales ont fait l'objet d'une prépublication commune avec O. GOYON [GP95].

2.1 Existence et unicité dans $H_0^1(\Omega)$ et système discret

Il s'agit d'un problème introduit sous sa forme évolutive par R. TEMAM [TEM66] dans le but de donner des équations ayant numériquement des difficultés comparables aux équations de Navier-Stokes en formulation vitesse-pression. Nous nous chargeons de résoudre théoriquement et numériquement le problème stationnaire sous-jacent suivant :

$$\begin{cases} \rho \mathbf{u} - \nu \Delta \mathbf{u} + (\mathbf{u} \nabla) \mathbf{u} + \frac{1}{2}(\operatorname{div} \mathbf{u}) \mathbf{u} &= \mathbf{f} \quad \text{dans } \Omega =]0, 1[^2 \\ \mathbf{u} &= 0 \quad \text{sur } \partial \Omega \end{cases} \quad (2.1)$$

où $\mathbf{f} \in L^2(\Omega)$

Théorème 2.1.1

Pour $\mathbf{f} \in L^2(\Omega)$ et ν tels que $\frac{|\mathbf{f}|}{\nu^2}$ soit suffisamment petit, il existe une unique solution

$$\mathbf{u} \in H_0^1(\Omega) \text{ vérifiant (2.1).}$$

La preuve se fait de manière classique par une méthode de Galerkin et par l'utilisation d'un théorème de point fixe.

On se place dans le cadre fonctionnel suivant :

$$V = H_0^1(\Omega), \quad H = L^2(\Omega), \quad \text{et } \mathcal{V} = \mathcal{D}(\Omega),$$

avec les normes habituelles, $|\mathbf{u}| = \left(\int_{\Omega} \mathbf{u}^2 dx \right)^{1/2}$, $\|\mathbf{u}\| = \left(\int_{\Omega} \nabla \mathbf{u}^2 dx \right)^{1/2}$.

H muni de $|\cdot|$ est un espace de Hilbert, et Ω étant borné, V muni de la norme $\|\cdot\|$

est aussi un espace de Hilbert (lemme de Poincaré).

Soient $f, \mathbf{u}$ deux fonctions régulières satisfaisant (2.1), en faisant le produit scalaire avec $\mathbf{v} \in \mathcal{V}$ on obtient :

$$(\rho \cdot \mathbf{u} - \nu \Delta \mathbf{u} + (\mathbf{u} \cdot \nabla) \mathbf{u} + \frac{1}{2}(\operatorname{div} \mathbf{u}) \mathbf{u}, \mathbf{v}) = (f, \mathbf{v})$$

après intégration par parties, on a :

$$\rho \cdot (\mathbf{u}, \mathbf{v}) + \nu((\mathbf{u}, \mathbf{v})) + b(\mathbf{u}, \mathbf{u}, \mathbf{v}) = (f, \mathbf{v}), \quad \forall \mathbf{v} \in \mathcal{V}$$

$$\text{où } b(\mathbf{u}, \mathbf{v}, \mathbf{w}) = \sum_{i,j=1}^2 \int_{\Omega} u_i \frac{\partial v_j}{\partial x_i} w_j dx + \frac{1}{2} \sum_{i=1}^2 \int_{\Omega} (\operatorname{div} \mathbf{u}) v_i w_i dx$$

Cette forme trilinéaire est antisymétrique et continue dans $H_0^1(\Omega)$:

- $b(\mathbf{u}, \mathbf{v}, \mathbf{v}) = 0 \quad \forall \mathbf{u}, \mathbf{v} \in H_0^1(\Omega)$
en effet, en appliquant la formule de Green,

$$\begin{aligned} \frac{1}{2} \sum_{i,j=1}^2 \int_{\Omega} \operatorname{div} \mathbf{u} \cdot (v_j)^2 dx &= \sum_{i,j=1}^2 \int_{\Omega} \frac{\partial u_i}{\partial x_i} \left(\frac{v_j^2}{2} \right) dx \\ &= - \sum_{i,j=1}^2 \int_{\Omega} u_i \frac{\partial}{\partial x_i} \left(\frac{v_j^2}{2} \right) dx \end{aligned}$$

donc,

$$b(\mathbf{u}, \mathbf{v}, \mathbf{v}) = \sum_{i,j=1}^2 \int_{\Omega} u_i \frac{\partial v_j}{\partial x_i} v_j dx - \sum_{i,j=1}^2 \int_{\Omega} u_i \frac{\partial}{\partial x_i} \left(\frac{v_j^2}{2} \right) dx = 0$$

- b est continu dans $(H_0^1(\Omega))^3$,
 $\forall \mathbf{u}, \mathbf{v}, \mathbf{w} \in H_0^1(\Omega)$, en appliquant l'inégalité de Hölder, on obtient :

$$\begin{aligned} |b(\mathbf{u}, \mathbf{v}, \mathbf{w})| &\leq \sum_{i,j=1}^2 |u_i|_{L^4(\Omega)} \cdot \left| \frac{\partial v_j}{\partial x_i} \right|_{L^2(\Omega)} \cdot |w_j|_{L^4(\Omega)} \\ &\quad + \sum_{i,j=1}^2 \left| \int_{\Omega} \frac{\partial u_i}{\partial x_j} v_i w_j dx \right| \\ &\leq \sum_{i,j=1}^2 (|u_i|_{L^4(\Omega)} \cdot |D_i v_j|_{L^2(\Omega)} \cdot |w_j|_{L^4(\Omega)} \\ &\quad + |D_j u_i|_{L^2(\Omega)} \cdot |v_i|_{L^4(\Omega)} \cdot |w|_{L^4(\Omega)}) \end{aligned}$$

de plus, pour Ω ouvert borné de $\mathbb{R}^2$,

$$|\mathbf{u}|_{L^4(\Omega)} \leq C(\Omega) \cdot |\mathbf{u}|_{H_0^1(\Omega)}$$

d'où,

$$|b(\mathbf{u}, \mathbf{v}, \mathbf{w})| \leq M_{\Omega} \cdot \|\mathbf{u}\| \cdot \|\mathbf{v}\| \cdot \|\mathbf{w}\|$$

Notons, pour $u, v \in V$, $B(u, v)$, la forme linéaire continue sur V par :

$$\langle B(u, v), w \rangle = b(u, v, w), \quad \forall w \in V$$

Toujours par le théorème de représentation de Riesz, comme l'application, pour $u \in V$ fixé :

$$v \in V \rightarrow \rho(u, v) + \nu((u, v))$$

est linéaire, continue sur V alors,

$$\exists! Au \in V' \text{ tel que } \rho(u, v) + \nu((u, v)) = \langle Au, v \rangle_{V', V}$$

et $u \rightarrow Au$ est un isomorphisme continu de V dans V' .

$$((u, v) = \langle u, v \rangle, \quad \forall v \in V, \quad \forall u \in H)$$

Ce qui nous amène à la formulation variationnelle :

$$\left\{ \begin{array}{l} \text{Pour } f \in V', \\ \text{Trouver } u \in V \text{ tel que} \\ Au + B(u, u) = f \text{ dans } V' \end{array} \right. \quad (2.2)$$

Pour $f \in H$, il est clair que (2.1) $\Leftrightarrow$ (2.2).

l'existence de u est prouvée par une méthode de Galerkin basée sur la séparabilité de V :

V est séparable donc, il existe $\{w_1, \dots, w_m\}$ libre et totale dans V

$$\left\{ \begin{array}{l} \forall m \geq 1, \\ \rho(u_m, w_i) + \nu((u_m, w_i)) + b(u_m, u_m, w_i) = \langle f, w_i \rangle, \quad i = 1, \dots, m \end{array} \right. \quad (2.3)$$

Lemme 2.1.2

Soit V_m un Hilbert de dimension finie et $F : V_m \rightarrow V_m$ une fonction continue telle que :

$$(F(x), x) > 0 \text{ pour } \|x\| = k > 0$$

alors,

$$\exists x \in V_m, \|x\| \leq k \text{ tel que } F(x) = 0$$

Pour la preuve, on pourra consulter [TEM84].

Soit $V_m = \text{Vect}(\{w_1, \dots, w_m\})$ et,

$$\forall u, v \in V, \quad ((F(u), v)) = \rho(u, v) + \nu((u, v)) + b(u, u, v) - (f, v)$$

F continue sur V_m et, comme $b(u, v, v) = 0 \forall u, v \in H_0^1(\Omega)$ alors,

$$((F(u), u)) = \rho \|u\|^2 + \nu \|u\|^2 - \langle f, u \rangle$$

d'où,

$$((F(u), u)) \geq \nu \|u\|^2 - \|f\|_{V'} \|u\|$$

et,

$$((F(u), u)) \geq \|u\| \cdot (\nu \|u\| - \|f\|_{V'})$$

donc, $((F(u), u)) > 0$ pour $\|u\| = k$ avec k vérifiant :

$$k > \frac{\|f\|_{V'}}{\nu}$$

Les hypothèses du lemme (2.1) sont vérifiées et il existe u_m solution de (2.3).

Le passage à la limite nécessite une estimation a priori :

en multipliant α_i^m par les équations du système (2.3) et additionnant ces dernières pour $i = 1, \dots, m$, on obtient :

$$\rho \|u_m\|^2 + \nu \|u_m\|^2 + b(u_m, u_m, u_m) = \langle f, u_m \rangle$$

et,

$$\begin{aligned} \nu \|u_m\|^2 &\leq \rho \|u_m\|^2 + \nu \|u_m\|^2 = \langle f, u_m \rangle \\ &\leq \|f\|_{V'} \|u_m\| \end{aligned}$$

d'où,

$$\|u_m\| \leq \frac{\|f\|_{V'}}{\nu}$$

donc $\{u_m\}_{m \geq 1}$ bornée dans V , or V réflexif,

d'où $u_m \rightharpoonup u$ dans V et comme $V \hookrightarrow H$ compacte, alors $u_m \rightarrow u$ dans H

Lemme 2.1.3

Si $u_\mu \rightharpoonup u$ dans V et $u_\mu \rightarrow u$ dans H , alors ,

$$b(u_\mu, u_\mu, v) \rightarrow b(u, u, v), \forall v \in V$$

En utilisant la trilinearité de b , on trouve que :

$$b(u, v, w) = -b(u, w, v), \forall u, v, w \in V$$

d'où, $\forall v \in \mathcal{V}$, $b(u_\mu, u_\mu, v) = - \sum_{i,j=1}^2 \int_{\Omega} u_{\mu_i} \frac{\partial v_j}{\partial x_i} u_{\mu_j} dx - \sum_{i=1}^2 \int_{\Omega} (\operatorname{div} u_\mu) u_{\mu_i} v_i dx$

Comme $\frac{\partial v_j}{\partial x_i} \in L^\infty(\Omega)$ alors,

$$\int_{\Omega} u_{\mu_i} u_{\mu_j} \frac{\partial v_j}{\partial x_i} dx \rightarrow \int_{\Omega} u_i u_j \frac{\partial v_j}{\partial x_i} dx \quad (u_{\mu_i} \rightarrow u_i \text{ dans } H)$$

$$\begin{aligned} \text{et, } & \left| \int_{\Omega} \frac{\partial u_{\mu_j}}{\partial x_j} v_i u_{\mu_i} dx - \int_{\Omega} \frac{\partial u_j}{\partial x_j} v_i u_i dx \right| \\ & \leq \left| \int_{\Omega} \left(\frac{\partial u_{\mu_j}}{\partial x_j} - \frac{\partial u_j}{\partial x_j} \right) v_i u_{\mu_i} dx \right| + \left| \int_{\Omega} \frac{\partial u_j}{\partial x_j} v_i (u_{\mu_i} - u_i) dx \right| \\ & \leq \|u_{\mu_j} - u_j\| \cdot |v_i u_{\mu_i} - v_i u_i| + \left| \int_{\Omega} \left(\frac{\partial u_{\mu_j}}{\partial x_j} - \frac{\partial u_j}{\partial x_j} \right) v_i u_i dx \right| \\ & + \left| \frac{\partial u_j}{\partial x_j} v_i \right| \cdot |u_{\mu_i} - u_i| \leq \varepsilon \quad (v_i u_i \in L^2(\Omega)) \end{aligned}$$

Par passage à la limite, et densité, on conclut que :

$$\rho(u, v) + \nu((u, v)) + b(u, u, v) = \langle f, v \rangle, \quad \forall v \in V$$

De plus, on obtient une unique solution de (2.1) pour :

$$\|f\|_{V'} < \nu^2 \cdot C \quad \text{où } C \text{ constante}$$

On pourra consulter [TEM84] pour la preuve.

Remarque 2 Théoriquement, il a été nécessaire de rajouter le terme " $\frac{1}{2}(\operatorname{div} u) \cdot u$ " à l'équation du problème (2.1) ressemblant à l'équation de la conservation de la quantité de mouvement, car n'ayant pas la relation $\operatorname{div}(u) = 0$, on ne peut pas avoir directement l'antisymétrie du terme non linéaire.

Remarque 3 Le terme " ρu " a été rajouté aussi, pour pouvoir obtenir le comportement de l'opérateur intervenant dans l'équation évolutive ci-dessous, complétée correctement de conditions aux limites et d'une condition initiale :

$$\frac{\partial u}{\partial t} - \nu \Delta u + (u \cdot \nabla) u + \frac{1}{2}(\operatorname{div} u) u = f \text{ dans } \Omega =]0, 1[^2$$

En effet, si l'on considère une discrétisation par différences finies avec le schéma Euler-Implicite pour cette équation évolutive, on obtient :

$$\begin{aligned} & \frac{u^{n+1} - u^n}{\Delta t} + A \cdot u^{n+1} + B(u^{n+1}, u^{n+1}) = F \\ \Leftrightarrow & \left(\frac{1}{\Delta t} \cdot I + A \right) u^{n+1} + B(u^{n+1}, u^{n+1}) = F + \frac{1}{\Delta t} \cdot u^n \end{aligned}$$

d'où, l'interprétation du ρ comme représentant $\frac{1}{\Delta t}$.

Système discret

Nous nous proposons d'approcher les opérateurs différentiels du système (2.1) par des schémas aux différences du second ordre. Ce qui nous permet d'obtenir le système discret (2.4) par différences finies classiques.

Posons les notations suivantes: soient $N, M \in \mathbb{N}^*$ le nombre de points dans chacune des directions,

$$h = \frac{1}{N+1}, \quad k = \frac{1}{M+1} \text{ les pas de discrétisation}$$

$$\Omega_{h,k} = \{(i,j) \in \mathbb{N}^2 / 1 \leq i \leq N \text{ et } 1 \leq j \leq M\}$$

$$\begin{aligned} \partial\Omega_{h,k} = & \{(i,0), 0 \leq i \leq N+1\} \cup \{(i,M+1), 0 \leq i \leq N+1\} \\ & \cup \{(0,j), 0 \leq j \leq M+1\} \cup \{(N+1,j), 0 \leq j \leq M+1\} \end{aligned}$$

Les composantes discrétisées de u et de f sont $\begin{pmatrix} u_{i,j} \\ v_{i,j} \end{pmatrix}$ et $\begin{pmatrix} f_{i,j} \\ g_{i,j} \end{pmatrix}$

$$(u(i.h, j.k) = \begin{pmatrix} u_{i,j} \\ v_{i,j} \end{pmatrix}, \text{ de même pour } f).$$

$$\left\{ \begin{array}{l} \rho u_{i,j} + \nu \left(\frac{2}{h^2} + \frac{2}{k^2} \right) u_{i,j} - \frac{\nu}{h^2} (u_{i-1,j} + u_{i+1,j}) \\ - \frac{\nu}{k^2} (u_{i,j-1} + u_{i,j+1}) + \frac{3}{4h} u_{i,j} \cdot (u_{i+1,j} - u_{i-1,j}) \\ + \frac{1}{4k} u_{i,j} \cdot (v_{i,j+1} - v_{i,j-1}) + \frac{1}{2k} v_{i,j} \cdot (u_{i,j+1} - u_{i,j-1}) = f_{i,j} \quad (i,j) \in \Omega_{h,k} \\ \rho v_{i,j} + \nu \left(\frac{2}{h^2} + \frac{2}{k^2} \right) v_{i,j} - \frac{\nu}{h^2} (v_{i-1,j} + v_{i+1,j}) \\ - \frac{\nu}{k^2} (v_{i,j-1} + v_{i,j+1}) + \frac{3}{4h} v_{i,j} \cdot (v_{i,j+1} - v_{i,j-1}) \\ + \frac{1}{4k} v_{i,j} \cdot (u_{i+1,j} - u_{i-1,j}) + \frac{1}{2k} u_{i,j} \cdot (v_{i+1,j} - v_{i-1,j}) = g_{i,j} \quad (i,j) \in \Omega_{h,k} \\ u_{i,j} = v_{i,j} = 0 \text{ pour } (i,j) \in \partial\Omega_{h,k} \end{array} \right. \quad (2.4)$$

2.2 Méthodes de Newton

Il s'agit d'une méthode itérative très robuste. De par sa convergence quadratique quand l'on se situe dans un voisinage (proche) de la solution, elle est souvent utilisée pour résoudre des problèmes de calculs d'écoulements stationnaires (cf. [FLE88]). Néanmoins, elle demeure plus lente que les méthodes de type Quasi-Newton souvent utilisées pour les problèmes d'optimisation. Nous citons un théorème qui assure la convergence (quadratique ou) géométrique de la méthode. Auparavant, précisons quelques notations :

$\mathcal{N}(x, r) = \{\bar{x} \in \mathbb{R}^n \text{ tel que } \|\bar{x} - x\| < r\}$ est le voisinage de x de rayon r , et, une fonction $g \in \text{Lip}_\gamma(X)$ si, $\forall x, y \in X \mid g(x) - g(y) \mid \leq \gamma \mid x - y \mid$. Pour $F : \mathbb{R}^n \rightarrow \mathbb{R}^n$ fonctionnelle continûment différentiable, nous noterons $J(x) = DF_x \in \mathcal{L}(\mathbb{R}^n)$.

Théorème 2.2.1

Soit $F : \mathbb{R}^n \rightarrow \mathbb{R}^n$ une fonctionnelle continûment différentiable sur un ouvert convexe $D \subset \mathbb{R}^n$.

Supposons qu'il existe $x_ \in \mathbb{R}^n$ et $r, \beta > 0$ tels que $\mathcal{N}(x_*, r) \subset D$, $F(x_*) = 0$ et que $J(x_*)^{-1}$ existe, avec $\|J(x_*)^{-1}\| \leq \beta$ et $J \in \text{Lip}_\gamma(\mathcal{N}(x_*, r))$ alors,*

il existe $\varepsilon > 0$ tel que : $\forall x_0 \in \mathcal{N}(x_, r)$, la suite $x_1, x_2, \dots$ générée par*

$$x_{k+1} = x_k - J(x_k)^{-1} F(x_k), k = 0, 1, \dots$$

est bien définie, et converge vers x_ et vérifie :*

$$\|x_{k+1} - x_k\| \leq \beta \gamma \|x_k - x_*\|^2, k = 0, 1, \dots$$

Pour la preuve du théorème, on pourra consulter [DS83b].

Bien que ce résultat soit très puissant, l'inversion du Jacobien $J(x_k)$ peut s'avérer difficile en pratique.

Un des atouts de la discrétisation par différences finies est que les matrices intervenant pour les problèmes issus de la mécanique des fluides peuvent se calculer aisément. En effet, aucun stockage n'est nécessaire dans notre cas pour la matrice Jacobienne, ce qui nous permet de pouvoir traiter des problèmes de grandes tailles - un point prometteur si l'on espère simuler des phénomènes nécessitant des discrétisations très fines avec une équation de Navier- Stokes pour des fluides incompressibles - . Dans l'optique de garder un stockage minimal, nous nous sommes écartés de toutes les méthodes de type Newton ou Quasi-Newton faisant intervenir des factorisations complètes de Gauss (LU) ou des préconditionnements nécessitant des factorisations incomplètes (ILU). Ces méthodes, très utilisées dans des implémentations sur l'algèbre linéaire, requièrent tout de même un stockage important pour

notre problème. Citons [GRMM92] où les auteurs effectuent des comparaisons pour ce type d'algorithmes.

2.2.1 Newton-BICGSTAB dans la base des I.U.

Par la méthode classique (NEWT), on peut décrire l'algorithme de la façon suivante :

Pour résoudre $F(\mathbf{u}_{i,j}) = 0$,

- On initialise $\mathbf{u}_{i,j}$ à $\mathbf{u}_{i,j}^{(0)}$
- A l'itération k , on résout :

$$\left| \begin{array}{l} [J(\mathbf{u}_{i,j}^{(k)})] \cdot \delta \mathbf{u} = F(\mathbf{u}_{i,j}^{(k)}) \\ \text{et,} \quad \mathbf{u}_{i,j}^{(k+1)} = \mathbf{u}_{i,j}^{(k)} - \delta \mathbf{u} \end{array} \right. \quad (2.5)$$

- Le test d'Arrêt s'effectue sur la norme de $F(\mathbf{u}_{i,j}^{(k+1)})$

où $J(\mathbf{u}_{i,j}^{(k)})$ représente le Jacobien de la fonctionnelle $F(\mathbf{u})$ évalué en la $k^{\text{ième}}$ itérée de $\mathbf{u}$.

Dans notre cadre, de discrétisations fines, la résolution du système (2.5) nécessite l'utilisation de méthodes bien appropriées. En effet, la matrice $[J(\mathbf{u}_{i,j}^{(k)})]$ est une grande matrice creuse et de plus, non symétrique. Dans le paragraphe suivant nous donnons l'algorithme qui nous a semblé le plus performant pour notre problème d'inversion.

La méthode de Newton préconditionnée par les Inconnues Incrémentales (IU-NEWT) revient à résoudre le système linéaire dans la base des I.U., l'algorithme est le suivant :

Pour chercher à résoudre $F(\mathbf{u}_{i,j}) = 0$,

- On initialise $\mathbf{u}_{i,j}$ à $\mathbf{u}_{i,j}^{(0)}$
- A l'itération k , on résout :

$$\left| \begin{array}{l} [{}^tS] \cdot [J(\mathbf{u}_{i,j}^{(k)})] \cdot [S] \cdot \overline{\delta \mathbf{u}} = [{}^tS] \cdot F(\mathbf{u}_{i,j}^{(k)}) \\ \text{et,} \quad \mathbf{u}_{i,j}^{(k+1)} = \mathbf{u}_{i,j}^{(k)} - [S] \cdot \overline{\delta \mathbf{u}} \end{array} \right. \quad (2.6)$$

- Le test d'Arrêt s'effectue sur la norme de $F(\mathbf{u}_{i,j}^{(k)})$.

où, $[S]$ est la matrice de changement de base des I.U. pour un nombre maximal de niveaux de grilles (cf. Chap. 1), et le Jacobien dans la base des I.U.,

$$[\bar{J}(\mathbf{u}_{i,j}^{(k)})] = [{}^tS] \cdot [J(\mathbf{u}_{i,j}^{(k)})] \cdot [S].$$

La matrice $[\bar{J}(\mathbf{u}_{i,j}^{(k)})]$ garde les mêmes propriétés que $[J(\mathbf{u}_{i,j}^{(k)})]$ dans la base nodale, ce qui nous a mené à inverser le système linéaire (2.6) de la même manière que le système (2.5).

L'utilisation des matrices $[S]$ et $[{}^tS]$ se fait de manière analogue à celle du chap. 1, autrement dit, ces deux matrices sont toujours simulées et les produits matrice-vecteur les impliquant se font toujours récursivement.

Remarque 4

L'inconvénient de la méthode de Newton est qu'elle requiert à chaque itération, l'inversion d'un système linéaire. En pratique, cette inversion n'est toujours faite que par une méthode itérative donc de manière approchée. Il a donc été nécessaire de quantifier la précision de l'inversion dont il fallait s'assurer pour garder une bonne convergence (peut être plus quadratique, mais au moins linéaire). *Dembo et al.* ont analysé que l'ordre de convergence de la méthode inexacte de Newton est fonction du rapport entre la précision de l'inversion et le résidu de Newton [DES82]. Nous choisissons comme d'autres auteurs [CJ84], [NAS84] une méthode appelée méthode de Newton tronquée consistant à imposer un simple critère d'arrêt de notre méthode de Krylov en fonction du résidu courant de Newton :

$$\frac{\| [J(\mathbf{u}_{i,j}^{(k)})] \cdot \delta u - F(\mathbf{u}_{i,j}^{(k)}) \|}{\| F(\mathbf{u}_{i,j}^{(k)}) \|} < 10^{-i-1}$$

Effectivement, nous nous rendons compte qu'il n'est souvent pas nécessaire d'inverser correctement le système linéaire dès les premières itérations de Newton, la solution calculée étant encore bien loin de la solution recherchée. Nous nommons les méthodes utilisant cette technique T-NEWT (comme Truncated-Newton method).

2.2.2 Résolution de systèmes linéaires non symétriques

Pour résoudre un système linéaire non symétrique qui fait intervenir une grande matrice creuse, on utilise, soit une méthode de type bi-gradient, soit une méthode de type GMRES.

Nous nous sommes penchés sur une méthode du premier type, due à Van Der Vorst [VOR92]. En fait, il s'agit d'une variante de la méthode BI-CG introduite par Fletcher [FLE76].

Dans BI-CG, l'approximation est menée de la manière suivante :

- on construit deux suites de vecteurs r_j (le résidu de BI-CG à la $j^{\text{ème}}$ itération) et $\hat{r}_j$ telles que:

r_j est orthogonal à tous les $\hat{r}_0, \dots, \hat{r}_{j-1}$ et vice versa,

$\hat{r}_j$ est orthogonal à tous les $r_0, \dots, r_{j-1}$.

- on peut écrire :

$$r_j = P_j(A) \cdot r_0 \text{ et } \hat{r}_j = \tilde{P}_j({}^t A) \cdot \hat{r}_0$$

$$\text{où, } P_j, \tilde{P}_j \in \mathcal{P}_j = \{\phi(x) = 1 + \sigma_1 x + \dots + \sigma_j x^j \mid \sigma_i \in \mathbb{R}, i \leq j\}$$

d'après construction :

$$(r_j, \hat{r}_i) = (P_j(A) \cdot r_0, \tilde{P}_i({}^t A) \cdot \hat{r}_0) = (\tilde{P}_i(A) P_j(A) \cdot r_0, \hat{r}_0) = 0, \forall i < j$$

ce qui exprime le fait que :

$$P_j(A) \cdot r_0 \perp \text{Vect}(\{\hat{r}_0, {}^t A \hat{r}_0, \dots, ({}^t A)^{j-1} \hat{r}_0\})$$

On arrive au choix du polynôme $\tilde{P}_j$ (P_j étant choisi par la méthode de gradient conjugué).

Dans un premier temps, Fletcher a choisi $\tilde{P}_j = P_j$ pour l'algorithme BI-CG, ensuite Sonneveld a eu l'idée de concentrer tout l'effort sur la convergence du résidu r_j et donc de considérer la suite de vecteurs :

$r_j = P_j^2(A) r_0$, cette version a donné naissance à l'algorithme bien connu CGS [SON89].

Ces algorithmes, bien que performants, gardaient toutefois un inconvénient majeur : une décroissance non régulière de la norme du résidu [VOR92]. Dans le but d'améliorer ceci, Van Der Vorst a eu l'idée de choisir $\tilde{P}_j$ de la façon suivante :

$$\tilde{P}_j = (1 - \omega_1 x)(1 - \omega_2 x) \dots (1 - \omega_j x)$$

où les ω_i sont des constantes bien adaptés pour la convergence de l'algorithme.

Bien qu'en pratique ces algorithmes de Bi-Gradient convergent beaucoup plus vite, théoriquement, BI-CG, CGS et BI-CGSTAB converge en au plus n itérations (où n est la taille du système à inverser), ce qui provient de la relation d'orthogonalité :

$$(P_j(A) r_0, \tilde{P}_i({}^t A) \hat{r}_0) = 0, \forall i < j$$

2.2.3 Quelques Préconditionnements pour BI-CGSTAB

Comme beaucoup de méthodes de Gradient, le coût de l'inversion peut être réduit considérablement en utilisant des algorithmes incluant un preconditionnement (pour plus de détails, on consultera [VOR92]).

En fait tout dépend du choix de la matrice qui va servir.

Dans un premier temps, nous choisissons la matrice de l'opérateur représentant la partie linéaire de F : il s'agit tout simplement de la matrice du laplacien (à la viscosité près) à laquelle on ajoute $\rho \cdot I$.

Dans un deuxième temps, en observant notre matrice Jacobienne, $J(\mathbf{u})$, nous pouvons nous rendre compte qu'elle garde une propriété intéressante : elle est pratiquement à diagonale dominante (elle est fonction de ρ , de la viscosité ν et de l'itérée de $\mathbf{u}$ et, par conséquent, dans certains cas cette propriété peut s'avérer fausse).

Le terme prépondérant, sur la diagonale est : $\delta = \frac{2\nu}{h_1^2} + \frac{2\nu}{h_2^2} + \rho$, et puisque ρ n'intervient que sur la diagonale, il s'en ressent que la propriété de dominance de la diagonale se renforcera d'autant que ρ grandira. Notons au passage que tous les autres termes élevés de la matrice seront atténués lorsque la viscosité diminuera.

De ce fait, nous décomposons $J(\mathbf{u})$ sous la forme :

$$J(\mathbf{u}) = \delta I - j(\mathbf{u})$$

et, considérons la décomposition en série de l'inverse de $J(\mathbf{u})$ comme suit :

$$J(\mathbf{u})^{-1} = \frac{1}{\delta} \left(I + \frac{1}{\delta} j(\mathbf{u}) + \frac{1}{\delta^2} j(\mathbf{u})^2 + \dots \right)$$

Comme dans la plupart des cas, δ est prépondérant par rapport aux termes de la matrice $j(\mathbf{u})$, nous pouvons nous permettre de tronquer cette série au $m^{\text{ième}}$ terme pour obtenir une matrice approchant $J(\mathbf{u})^{-1}$:

$$\text{soit, } E_m(\mathbf{u}) = \frac{1}{\delta} I + \frac{1}{\delta^2} j(\mathbf{u}) + \dots + \frac{1}{\delta^{m+1}} j(\mathbf{u})^m$$

Avec cette approximation, nous pouvons nous attendre à n'avoir aucun stockage supplémentaire et surtout ne pas perdre un point important de la réalisation informatique : le caractère vectorisable du code de calcul.

Citons Van Der Vorst [VOR82], qui a utilisé cette technique d'approximation pour développer une variante de la méthode de la factorisation Incomplète de Choleski vectorisable.

Donnons la description de l'algorithme BI-CGSTAB preconditionné, pour chercher à résoudre $Ax = b$ avec K , la matrice de preconditionnement ($K \approx A$) :

- Etape d'initialisation :

soit x_0 valeur initiale ; $r_0 = b - Ax_0$;
 $\bar{r}_0$ un vecteur arbitraire tel que $(\bar{r}_0, r_0) \neq 0$, par exemple $\bar{r}_0 = r_0$;
 $\rho_0 = \alpha = \omega_0 = 1$ (des constantes) ;
 $v_0 = p_0 = 0$ (des vecteurs) ;

- A la $i^{\text{ème}}$ itération :

$$\begin{aligned} \rho_i &= (\bar{r}_0, r_{i-1}) ; \beta = \frac{\rho_i}{\rho_{i-1}} \cdot \frac{\alpha}{\omega_{i-1}} ; \\ p_i &= r_{i-1} + \beta \cdot (p_{i-1} - \omega_{i-1} \cdot v_{i-1}) ; \\ \text{on cherche } y \text{ tel que} \\ K \cdot y &= p_i ; \end{aligned} \quad (\text{sl.1})$$

$$\begin{aligned} v_i &= Ay ; \\ \alpha &= \frac{\rho_i}{(\bar{r}_0, v_i)} ; \\ s &= r_{i-1} - \alpha \cdot v_i ; \\ \text{on cherche } z \text{ tel que} \\ K \cdot z &= s ; \end{aligned} \quad (\text{sl.2})$$

$$\begin{aligned} t &= Az ; \\ \omega_i &= \frac{(t, s)}{(t, t)} ; \\ x_i &= x_{i-1} + \alpha \cdot y + \omega_i \cdot z ; \\ r_i &= s - \omega_i \cdot t ; \end{aligned} \quad (3)$$

- le test d'arrêt s'effectue sur la norme de r_i .

Tout d'abord, précisons le choix de K dans chaque implémentation :

dans *MGNEWT*, nous avons considéré comme matrice :

$$K = \rho I - \nu \Delta_h \text{ où } \Delta_h \text{ est la matrice de discrétisation du laplacien}$$

et, nous avons effectué deux troncatures ($m = 2$ et $m = 3$) dans le développement de la série de $J(u)^{-1}$ et choisi les deux matrices de préconditionnement $E_2(u)$ et $E_3(u)$. Ces deux méthodes respectives sont nommées *P²NEWT* et *P³NEWT*.

On peut remarquer quelques points intéressants :

- Pour accroître l'efficacité de BI-CGSTAB avec préconditionnement, il est nécessaire d'utiliser des solveurs performants pour résoudre les deux systèmes

linéaires (sl.1) et (sl.2). Nous choisissons une méthode Multigrilles Correction Storage utilisant la méthode itérative de base : Gauss-Seidel noir-blanc (cf. Chap. 1). Il ne faut pas oublier qu'il ne s'agit pas de la méthode Multigrilles la plus performante à l'heure actuelle, mais sa bonne vectorisation lui permet d'être comparable aux meilleures.

- Dans *MGNEWT*, comme l'initialisation de la solution se fait à 0, à la première itération de Newton, on doit inverser (2.5), autrement dit :

$$K \cdot \delta u = R(u^{(0)}),$$

ce qui ne doit coûter qu'une seule itération de BI-CGSTAB.

A priori, l'efficacité de ce préconditionnement devrait se faire sentir surtout pour la (ou les) première(s) itération(s) de Newton.

Remarque 5

Vu que BI-CGSTAB est une méthode de projection sur un sous-espace de Krylov, on peut attirer l'attention sur un point développé dans le paragraphe suivant : par une méthode de Krylov dans un schéma itératif de type Newton, il n'est pas nécessaire de former explicitement la matrice Jacobienne, car elle n'intervient que par son produit par un vecteur, quantité qui peut être approchée par un quotient aux différences. Toutefois, il ne nous a pas paru très intéressant pour autant d'approcher ce calcul par un quotient aux différences (comme fait par d'autres [DS83a] et [SS86]) du fait qu'aucun stockage (de la Jacobienne en l'occurrence) n'est utile dans notre cas et que la complexité du produit matrice vecteur utilisant l'approximation est du même ordre que celle faisant intervenir la vraie jacobienne.

2.3 Méthode de GMRES non linéaire

Il s'agit d'une méthode due à P.N. Brown et Y. Saad, méthode non linéaire de projection sur des sous-espaces de Krylov [BS90]. Pour la décrire, il faut se reporter à un travail précédent [SS86] où Y. Saad et M. Schultz ont introduit l'algorithme GMRES pour résoudre des systèmes linéaires non symétriques. Dans un premier paragraphe, nous présentons l'algorithme de GMRES pour la résolution de systèmes linéaires. Ensuite, l'objectif du second paragraphe est, pour notre problème non linéaire, de décrire les stratégies que nous utilisons lors de nos différentes implémentations. Parmi elles, la version *Scaled-Preconditioned GMRES* nous permet d'introduire comme fait dans la section précédente nos Inconnues Incrémentales ainsi que le préconditionnement par l'inverse de la partie linéaire de la fonctionnelle.

2.3.1 Résolution de systèmes linéaires non symétriques par GMRES

Cet algorithme peut être considéré comme une généralisation de l'algorithme MINRES du à Paige et Saunder. Ils exploitent tous deux, une relation entre la méthode de Lanczos et la méthode du Gradient Conjugué (GC): que ce soit pour résoudre $Ax = b$ par GC, ou pour résoudre le problème aux valeurs propres pour la matrice A , il s'agit de méthodes de Galerkin sur le sous-espace de Krylov, $K_k = \text{Vect}(\{v_1, Av_1, \dots, A^{k-1}v_1\})$. *GMRES* utilise le procédé d'Arnoldi [ARN51], qui construit une base orthonormale par la méthode de Gram-Schmidt :

$$V_k = (\{v_1, \dots, v_k\}) \text{ du sous-espace de Krylov } K_k.$$

De plus, le procédé d'Arnoldi construit une matrice de Hessenberg $\bar{H}_k$ de dimension (k, k) , telle que :

$$\bar{H}_k = {}^tV_k \cdot A \cdot V_k.$$

Au lieu de résoudre $Ax = b$ par une méthode de Galerkin utilisant V_k , on écrit la solution x_k comme étant :

$$x_k = x_0 + z_k \text{ où } x_0 \text{ est la valeur initiale de la solution et,}$$

$$z_k \in \tilde{K}_k \equiv \text{Vect}(\{r_0, Ar_0, \dots, A^{k-1}r_0\}) \text{ avec, } r_0 = b - A \cdot x_0$$

Donnons la description du procédé :

- x_0 valeur initiale, calculons le résidu $r_0 = b - A \cdot x_0$
- calcul de $\beta = \|r_0\|$ et $v_1 = \frac{1}{\beta}r_0$
- Pour $j = 1, 2, \dots$

$$\begin{aligned} h_{i,j} &= (A \cdot v_j, v_i) \text{ pour } i = 1, \dots, j \\ \hat{v}_{j+1} &= A \cdot v_j - \sum_{i=1}^j h_{i,j} v_i \\ h_{j+1,j} &= \|\hat{v}_{j+1}\| \text{ et } v_{j+1} = \frac{\hat{v}_{j+1}}{h_{j+1,j}}. \end{aligned}$$

Ce qui revient à définir une matrice H_k de dimension $(k+1, k)$ dont les éléments non nuls sont les $h_{i,j}$ pour $1 \leq i \leq j+1$, $(1 \leq j \leq k)$.

Nous arrivons à l'étape de minimisation du résidu :

$$\mathcal{Min}_{z \in K_k} \|b - A \cdot (x_0 + z)\|_2 = \mathcal{Min}_{z \in K_k} \|r_0 - A \cdot z\|_2$$

On peut considérer z comme étant les coordonnées de y dans la base V_k ou encore :

$$z = V_k \cdot y$$

et, le problème de minimisation sur K_k revient à un problème de minimisation d'une fonctionnelle $\mathcal{J}(y)$ sur tout $\mathbb{R}^k$:

$$\underset{y \in \mathbb{R}^k}{\text{Min}} \|\beta v_1 - A \cdot V_k \cdot y\|_2 = \underset{y \in \mathbb{R}^k}{\text{Min}} \mathcal{J}(y) \text{ où } \beta = \|r_0\|_2$$

En utilisant une relation liant V_k à H_k :

$$A \cdot V_k = V_{k+1} \cdot H_k$$

$$\mathcal{J}(y) = \|V_{k+1} \cdot (\beta e_1 - H_k y)\|_2 = \|\beta e_1 - H_k y\|_2 \text{ où } e_1 = {}^t(1, 0, \dots, 0) \in \mathbb{R}^{k+1}$$

car V_{k+1} est l_2 -orthogonale.

Pour résoudre le problème de moindres carrés : $\underset{y \in \mathbb{R}^k}{\text{Min}} \mathcal{J}(y)$,

il suffit d'utiliser la méthode de Givens (de rotations planes) pour factoriser la matrice de Hessenberg H_k sous la forme $Q_k \cdot R_k$. L'implémentation de l'algorithme se fait de manière très pratique, la factorisation de Givens de H_k se fait progressivement à chaque création d'une nouvelle colonne de H_k par procédé d'Arnoldi. L'intérêt de la factorisation progressive réside dans le fait que la norme du résidu de la solution approchée peut être calculée sans même la nécessité de former la solution approchée x_k .

Ainsi, après k étapes du procédé, la décomposition de H_k est complète :

$$Q_k \cdot H_k = R_k \text{ où } Q_k \in \mathbb{R}^{(k+1) \times (k+1)} \text{ est le produit de matrices de rotations,}$$

et, $R_k \in \mathbb{R}^{(k+1) \times k}$ est constituée d'une matrice triangulaire supérieure à laquelle on rajoute une ligne de 0.

Comme Q_k est unitaire,

$$\mathcal{J}(y) = \|\beta e_1 - H_k y\|_2 = \|Q_k \cdot [\beta e_1 - H_k y]\|_2 = \|g_k - R_k y\|_2$$

$$\text{où } g_k = Q_k \beta e_1$$

Pour effectuer la minimisation de $\mathcal{J}$, il suffit d'inverser le système triangulaire supérieur provenant des k premières lignes de R_k et g_k (la dernière ligne de R_k étant constituée de 0) : ce qui nous donne directement la valeur de $\mathcal{J}(y_k)$ qui est la dernière composante de g_k . Et, en fait $\mathcal{J}(y_k)$ est la norme du résidu, obtenue sans même avoir formé la solution x_k :

$$\mathcal{J}(y_k) = \|g_k - R_k y_k\| = (g_k)_{k+1}$$

2.3.2 Les I.U. pour préconditionner GMRES non linéaire

Dans un premier temps, nous discutons du remplacement du produit de la matrice Jacobienne par un vecteur, par un quotient aux différences, de la forme :

$$J(\mathbf{u}) \cdot \mathbf{v} \approx \frac{F(\mathbf{u} + \sigma \mathbf{v}) - F(\mathbf{u})}{\sigma}$$

où $\mathbf{u}$ est l'approximation de la racine et σ un scalaire.

Ensuite, nous expliquons le procédé de recherche suivant une ligne (cf. [BS90] et [DS83b]) et, nous présentons la version de l'algorithme *GMRES* non linéaire avec préconditionnement. Dans ce cas, nous appliquons les Inconnues Incrémentales ainsi qu'un autre préconditionnement plus classique.

Approximation de la Matrice Jacobienne

Pour expliciter ce point nous citons un résultat dont la preuve pourra être vue dans [DS83b].

Théorème 2.3.1

Soient F et x^* vérifiant les hypothèses du théorème 2.2.1
il existe $\epsilon, h > 0$ tels que,
si $\{h_k\}$ est une suite de réels vérifiant $0 < |h_k| < h$ et, $x_0 \in \mathcal{N}(x^*, \epsilon)$,
la suite $x^1, x^2, \dots$ générée par :

$$\begin{aligned} (A_k)_{.j} &= \frac{F(x^k + h_k e_j) - F(x^k)}{h_k}, \quad j = 1, 2, \dots, n \\ x^{k+1} &= x^k - A_k^{-1} F(x^k) \quad k = 0, 1, \dots \end{aligned}$$

est bien définie et converge q -linéairement vers x^* , si $\lim_{k \rightarrow \infty} h_k = 0$

De plus, la convergence est q -quadratique vers x^* si, il existe une constante C_1 telle que :

$$|h_k| \leq C_1 \|F(x^k)\|$$

Ce théorème très utile ne nous enseigne pas la manière d'effectuer le choix de h_k . La situation est la suivante : en faisant l'approximation de $J(x^{(k)})_{.j}$ par $(A_k)_{.j}$ l'erreur est de l'ordre de h_k d'où la nécessité de prendre h_k petit. De plus, si l'on prend h_k trop petit, $x^{(k)} + h_k e_j$ peut être si proche de $x^{(k)}$ que $F(x^{(k)} + h_k e_j) \cong F(x^{(k)})$. En l'absence de d'autres informations, Dennis et Schnabel préconise de prendre $h_k = \sqrt{\epsilon} \cdot x^{(k)}$, et ils expliquent pourquoi l'approche par un pas uniforme peut être désastreuse dans le cas où les composantes de $x^{(k)}$ ne sont pas du même ordre de grandeur.

La technique de Recherche suivant une ligne

Cette technique a été développée dans le but "d'élargir" le bassin d'attraction de Newton. L'idée est simple : supposons que l'on ait obtenue une direction de descente p^k , on doit chercher un pas dans cette direction pour avoir $u^{(k+1)}$ la nouvelle itérée de la solution. Et, pour qu'elle soit une *bonne* itérée, il suffit donc de vérifier que cette nouvelle solution réduit la norme du résidu de Newton.

L'algorithme de backtracking line-search est le suivant :

Soit $\alpha \in [0, \frac{1}{2}]$ (par exemple $= 10^{-4}$), $0 < l < u < 1$, $\lambda_k = 1$,

Tant que $\|F(u^k + \lambda_k p^k)\| > \|F(u^k)\| + \alpha \lambda_k J(u^k) \cdot p^k$, faire :

$\begin{aligned} \lambda_k &= \eta \lambda_k \text{ pour } \eta \in [l, u] \text{ où } \eta \text{ est choisi à chaque fois comme ci-dessous} \\ u^{k+1} &= u^k + \lambda_k p^k \end{aligned}$
--

fin

Pour plus de détails, on pourra consulter [DS83b].

Le Préconditionnement

Dans le but d'améliorer leur efficacité et leur robustesse, l'usage des méthodes de projection sur des sous-espaces de Krylov nécessite l'utilisation de preconditionnements.

En général, preconditionner un algorithme revient à faire un changement linéaire de variables et à résoudre un système équivalent :

$$\begin{aligned} (Q^{-1} \cdot J \cdot P^{-1}) \cdot \tilde{\delta} &= Q^{-1} \cdot b \\ \text{où } b &= -F(u^n) \end{aligned} \tag{2.7}$$

J est toujours la Jacobienne de F évaluée en u^n , et, P, Q étant choisies de manière à ce que le système (2.7) soit plus facile à résoudre que l'original.

Dans cette version, deux systèmes linéaires doivent être résolus à chaque itération :

$$P \cdot y = c \text{ et } Q \cdot y = c$$

Il paraît évident que cette résolution supplémentaire ne doit pas pénaliser le coût de la méthode. D'où l'utilisation de solveurs très rapides pour ces deux inversions.

Dans notre cas, nous avons pris $Q^{-1} = I$ et $P^{-1} = S$ la matrice de changement de base des I.U. (cf. Chap. 1) et nommée *IUGMRES* l'implémentation. *MGGMRES* fait appel à l'algorithme qui utilise $Q = \rho I - \nu \Delta_h$, $P = I$ mais uniquement quand

le nombre d'itérations de *linesearch backtracking* précédent est nul (i.e. quand le préconditionnement vaut la peine d'être utilisé, cf. 2.2.4).

Une différence intervient, toutefois, avec l'utilisation de préconditionnements ; en effet, la norme du résidu associé au (2.7) est :

$$\|\tilde{r}\| = \|\tilde{b} - \tilde{J}\tilde{\delta}\| \quad \text{où} \quad \begin{aligned} \tilde{b} &= b, \\ \tilde{J} &= JS, \\ \text{et, } \tilde{\delta} &= S^{-1}\delta. \end{aligned}$$

L'algorithme est le suivant :

(0) Initialisation, choix de $\mathbf{u}^0$

(1) Procédé d'Arnoldi

- choix d'une tolérance ϵ
- soit $\tilde{\delta}^{(0)}$ valeur initiale ; $\mathbf{r}^{(0)} = -(F + \tilde{J}\tilde{\delta}^{(0)})$ où $F = F(\mathbf{u}^k)$ et $\tilde{J} = J(\mathbf{u}^k) \times S$
- calcul de $\beta = \|\mathbf{r}^{(0)}\|_2$ et $\tilde{\mathbf{v}}_1 = \frac{1}{\beta} S\mathbf{r}^{(0)}$
- Pour $j = 1, 2, \dots$ faire :
 - (a) calcul de $\tilde{J}\tilde{\mathbf{v}}_j$,
 - (b) $\tilde{h}_{i,j} = (\tilde{J}\tilde{\mathbf{v}}_j, \tilde{\mathbf{v}}_i)$, pour $i = 1, 2, \dots, j$
 - (c) $\hat{\mathbf{v}}_{j+1} = \tilde{J}\tilde{\mathbf{v}}_j - \sum_{i=1}^j \tilde{h}_{i,j} \tilde{\mathbf{v}}_i$
 - (d) $\tilde{h}_{j+1,j} = \|\hat{\mathbf{v}}_{j+1}\|_2$,
 - (e) $\tilde{\mathbf{v}}_{j+1} = \frac{\hat{\mathbf{v}}_{j+1}}{\tilde{h}_{j+1,j}}$
 - (f) calcul de la norme du résidu $\rho_j = \|F + \tilde{J}\tilde{\delta}^{(j)}\|_2$, de la solution $\tilde{\delta}^{(j)}$ que l'on voudrait obtenir si l'on s'arrêtait à cette étape
 - (g) si $\rho_j \leq \epsilon$ alors, on pose $m = j$ et on va en (2)

(2) Calcul de la solution approchée (par GMRES)

- on définit H_m la $(m+1) \times m$ -matrice de Hessenberg formée des éléments non nuls $\tilde{h}_{i,j}$, et $\tilde{V}_m = [\tilde{\mathbf{v}}_1, \dots, \tilde{\mathbf{v}}_m]$
- on trouve $\tilde{\mathbf{y}}_m$ solution du problème de minimisation $\text{Min} \|\beta \mathbf{e}_1 - \tilde{H}_m \tilde{\mathbf{y}}\|_2$ sur $\mathbb{R}^m$ ($\mathbf{e}_1 = (1, 0, \dots, 0)$)
- on calcule $\tilde{\delta}^{(m)} = \tilde{\delta}^{(0)} + \tilde{\mathbf{z}}^{(m)}$ où $\tilde{\mathbf{z}}^{(m)} = \tilde{V}_m \tilde{\mathbf{y}}_m$,
- on calcule (en accord avec "backtrack linesearch"), $\mathbf{u}^{k+1} = \mathbf{u}^k + S\tilde{\delta}^{(m)}$.
Si $\mathbf{u}^{k+1}$ assez précis alors on va en (3), sinon on va en (1).

(3) Fin de la résolution.

2.4 Résultats Numériques

Nous avons mis au point les méthodes en résolvant le problème discret (2.4).

En effet, comme stratégie, il nous semble intéressant de considérer des fonctions exactes discrètes, de leur appliquer l'opérateur non linéaire discret de manière à obtenir le second membre exact du problème discret. Ensuite, en prenant une solution initiale *pas trop mauvaise*, on résout le problème discret avec le second membre discret (exact). Nous proposons une autre stratégie dans [GP95].

Tout d'abord, nous présentons nos différents tests de comparaison et les motivations de notre choix. Ensuite, nous donnons les deux fonctions tests $u_{3_{ex}}$ et $u_{4_{ex}}$ (de $\mathbb{R}^2$ dans $\mathbb{R}^2$) que nous utilisons comme solution exacte de notre problème discret. Nous avons choisi 0 comme condition initiale, donc la norme de la fonction exacte nous permet de quantifier la distance initiale de la solution exacte. Il nous a semblé intéressant de choisir deux fonctions très oscillantes (de norme différente) (cf. page 121 et 122).

$$|u_{3_{ex}}|_{L^\infty} \approx 1.1 \quad |u_{4_{ex}}|_{L^\infty} \approx 2.7$$

2.4.1 Présentation des tests

Nous considérons plusieurs valeurs des deux constantes intervenant dans le problème (2.1) pour établir nos tests. Tout d'abord, pour une discrétisation de 129^2 points, nous faisons varier la constante ρ comme s'il représente l'inverse d'un pas de temps pour un schéma implicite. Nous choisissons des valeurs particulières :

- $\rho = 0$ nous rapproche des opérateurs intervenant dans le calcul d'écoulement stationnaire visqueux,
- $\rho = 128$ nous simule l'opérateur intervenant dans le calcul d'un problème évolutif non linéaire (à l'aide d'un schéma implicite) avec un pas de temps du même ordre que le pas d'espace.

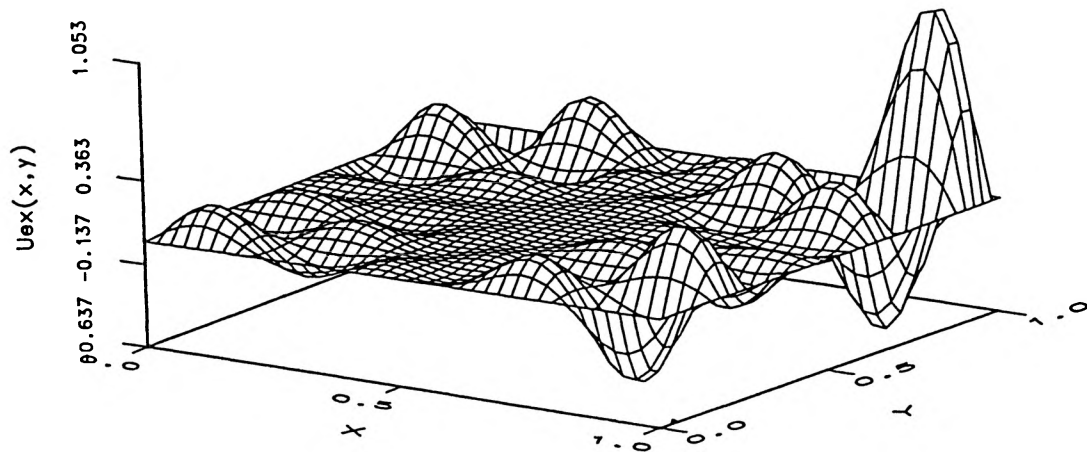
Après quelques expériences $\rho = 64$ s'est révélée une valeur suffisante.

Quatre premiers tests nous permettent de tester nos algorithmes pour chacun des deux problèmes discrets et pour chacune des fonctions exactes testées. Au test T_5 , nous nous mettons dans le cas d'une faible viscosité et d'un coefficient $\rho = 128$, des constantes pouvant approcher un calcul d'écoulement instationnaire peu visqueux. Le dernier test reprend le test T_5 pour un problème *stationnaire*.

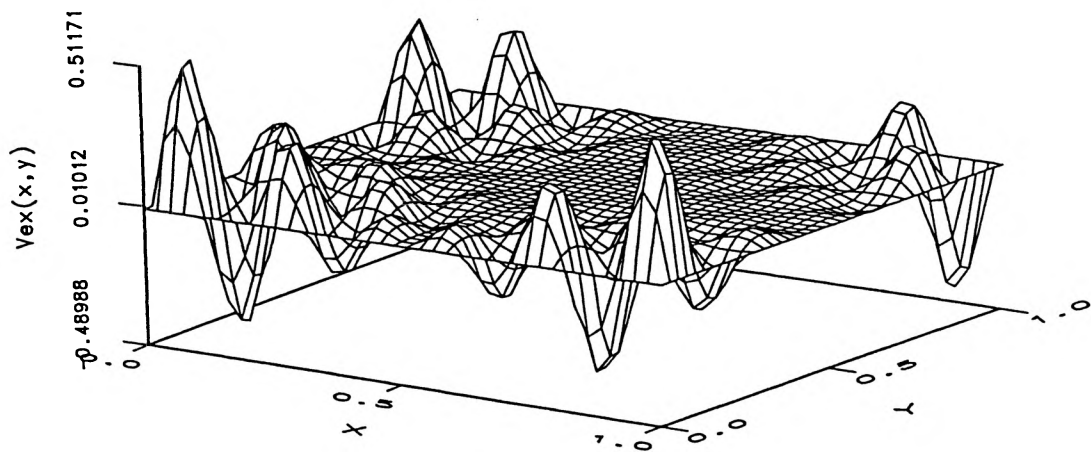
$$u_{ex}(x,y) = \sin(5\pi x)\sin(7\pi y)\exp(x+y)(x-0.5)(y-0.5)$$

$$v_{ex}(x,y) = \sin(10\pi x)\sin(7\pi y)\exp(\cos(2\pi xy))(x-0.5)(y-0.5)$$

3rd test function (1st component)



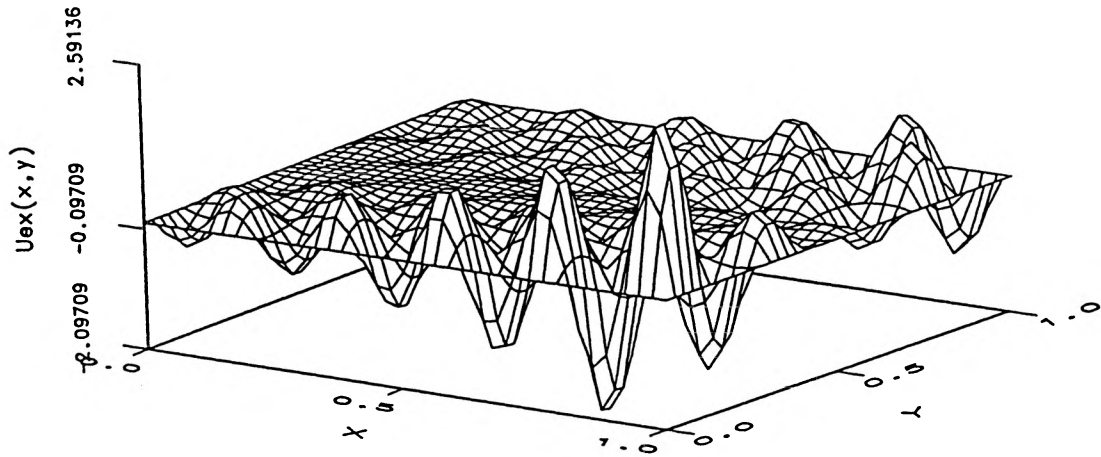
3rd test function (2nd component)



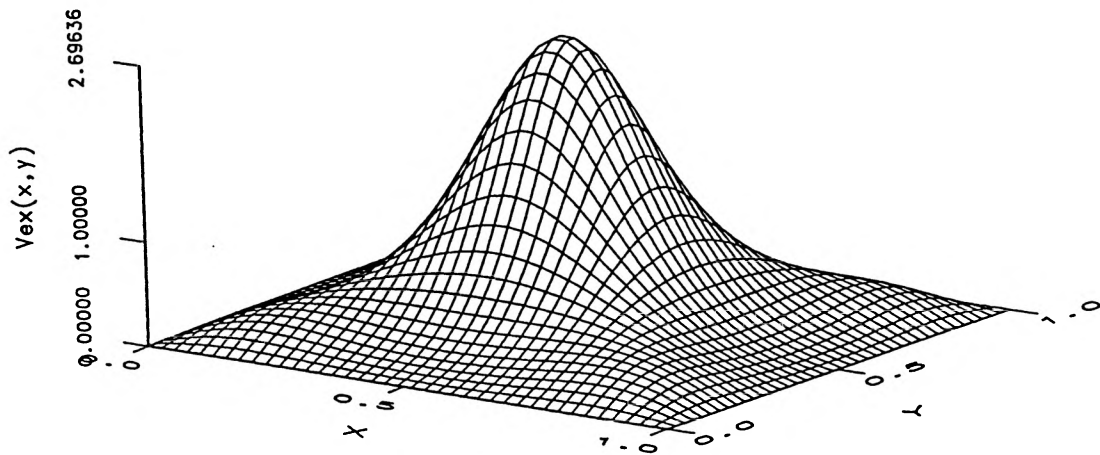
$$U_{ex}(x,y) = \sin(10\pi x)\exp(2x-y)\sin(7\pi y)(y-0.5)$$

$$V_{ex}(x,y) = \sin(\pi x)\sin(\pi y)\exp(\cos(2\pi x)\cos(2\pi y))$$

4th test function (1st component)



4th test function (2nd component)



2.4.2 Résultats pour les méthodes de Newton

Nous rappelons dans un premier temps les noms de méthodes utilisées et les paragraphes où elles sont explicitées, ensuite la série de tests qui nous ont semblé très significatifs, suivi de deux tableaux de résultats à ces tests.

- *NEWT* et *IUNEW*T : Paragraphe 2.2.1 p.31
- *MGNEWT* et *P²NEWT* : Paragraphe 2.2.5 p.34
- *T – NEWT* : Paragraphe 2.2.2 p.32

Le tableau ci-dessous rappelle la série de tests présentés précédemment servant de base à notre étude :

Tests	Fonction	Discrétisation	ρ	ν
T_1	3	127^2	0	0.1
T_2	3	127^2	64	0.1
T_3	4	127^2	0	0.1
T_4	4	127^2	64	0.1
T_5	3	127^2	128	10^{-3}
T_6	3	127^2	0	0.01

Nous récapitulons les résultats aux tests proposés ci-dessus, pour chacune des méthodes de Newton : le chiffre entre parenthèses représente le nombre d'itérations de Newton, le nombre suivant, le temps CPU pour obtenir la convergence et le dernier le nombre total d'itérations de BI-CGSTAB.

Tests	<i>NEWT</i>	<i>MGNEWT</i>	<i>IUNEW</i> T	<i>P²NEWT</i>
T_1	(6)73/1049	(4)34/20	(4)17/127	(6)124/655 ^{s 1}
T_2	(4)20/290	(4)23/16	(4)34/262	(4)41/215
T_3	(7)139/1775	(8)660/373	(6)48/334	(6)226 ^s
T_4	(4)34/429	(4)58/58	(4)46/321	(7)61/327
T_5	(4)5/63	(4)29/68	(5)113/936	(4)16/84
T_6	(9)146/2198	(9)700/400 ^s	(9)72/560	(9)303/1689

Le test de sortie sur la norme L^2 du Résidu de l'inversion par BI-CGSTAB est 5.10^{-10} , pour les méthodes : *NEWT*, *MGNEWT* et *IUNEW*T. Pour *P²NEWT* et *P³NEWT*, ce test est de 10^{-13} , ce qui se justifie par le changement d'échelle de ces algorithmes atténuant l'ordre de grandeur du résidu.

A certaines itérations de Newton, nous nous sommes rendus compte que l'inversion

1. Cas de stagnation du résidu de Newton a 5.10^{-9}

s'avérait difficile sans pour autant empêcher la convergence de la méthode. Pour pallier à cette difficulté, nous nous sommes fixés un nombre maximal d'itérations (2000) de Bi-Gradient, et un test relatif de stagnation du résidu. Ce dernier teste la valeur du rapport entre la différence de deux résidus consécutifs de Gradient et le résidu courant. Il faut bien se rendre compte que le test de convergence, le nombre maximal d'itérations et le test de stagnation doivent quand même être fixé pour permettre l'inversion la plus précise de la matrice jacobienne. En effet, une résolution imprécise du système linéaire nous ramènerait au concept d'une méthode de Newton tronquée (truncated-Newton method) ou quasi-Newton.

Remarque sur le test de stagnation

Si on observe le comportement du résidu par la méthode BI-CGSTAB, on constate visiblement que dans certains cas la décroissance peut être perturbée. Cette décroissance reste quand même beaucoup plus régulière que celle des anciens algorithmes de gradient pour les problèmes non symétriques [VOR92]. D'autre part, comme toutes les méthodes de gradient, la courbe de décroissance du résidu de BI-CGSTAB présente les caractéristiques habituelles : une première zone assez plate, le temps de trouver une coordonnée substantielle du vecteur de base de l'espace de Krylov (direction de descente), une zone un peu plus accentuée où effectivement cette bonne direction a été trouvée puis, une troisième zone de stagnation présentant parfois des oscillations, difficilement interprétables à cause des erreurs d'arrondi qui découlent de la précision des moyens informatiques. Nous avons choisi pour le test de stagnation, $\epsilon_{\text{stag}} = 10^{-4}$. Avec r_l représentant le résidu de BI-CGSTAB à l'itération l , le test de stagnation est :

$$\frac{|\|r_l\|_2 - \|r_{l+1}\|_2|}{\|r_l\|_2} \leq \epsilon_{\text{stag}}.$$

De plus, nous nous sommes fixés un nombre maximal de 10 itérations pour l'inversion par Multigrilles de chaque système linéaire. Enfin, le test de sortie porte encore sur la norme L^2 du résidu : il est de 10^{-14} .

Analyse des tests relatifs à la fonction 3

Les résultats sont très clairs : si l'on regarde les temps CPU mis pour obtenir la convergence, pour le test T_1 , les méthodes *IUNEW*T et *MGNEWT* sont les deux plus rapides avec des facteurs d'accélération de la convergence (f.a.c.) respectifs de 4,3 et 2,1 que la méthode non préconditionnée *NEWT*, par contre, quand on augmente ρ (au test T_2 , $\rho = 64$), *MGNEWT* reste inchangée quand on constate que le préconditionnement par les I.U. retarde la convergence (f.a.c. = 0,6).

On explique ce phénomène par la variation du nombre de condition de la matrice Jacobienne. En augmentant ρ , on renforce la diagonale ce qui a pour effet de diminuer le conditionnement de $J(u)$ et par conséquent facilite l'inversion dans la base nodale. Cette propriété est perdue dans la base des Inconnues Incrémentales car après multiplication de $J(u)$ à gauche par $'S$ et à droite par S , on ne peut plus rien dire du conditionnement de la matrice obtenue $\tilde{J}(u)$, la matrice S n'étant pas

orthogonale. Mais, utiliser *MGNEWT* reste toujours intéressant, du fait que les deux résolutions des systèmes linéaires (sl.1) et (sl.2) (cf. page 113) font intervenir une matrice $(\rho I - \nu \Delta_h)$ prise dans la base nodale, donc, facilement inversible.

Si l'on regarde les résultats obtenus pour *P²NEWT*, on s'aperçoit que les temps CPU sont même plus grands que ceux mis par la méthode *NEWT*. Toutefois, l'analyse de la page 112 reste vraie, on peut le constater avec la variation du nombre d'itérations de BI-CGSTAB lors des tests T_1 , T_2 et T_6 . (Le test T_5 , bien qu'étayant notre propos ci-dessus sur la variation du conditionnement de la matrice Jacobienne sera examiné par la suite). Nous avons vérifié le gain sur le nombre d'itérations pour *P³NEWT* par rapport à *P²NEWT*, ce qui s'explique par l'approximation plus précise de l'inverse de $J(u)$ par la matrice de préconditionnement $E_3(u)$ que $E_2(u)$ (cf. 2.2.3).

Il faut tout de même s'interroger sur les résultats des méthodes faisant intervenir un préconditionnement polynômial. Nous obtenons toujours un gain sur le nombre d'itérations mais une perte de temps CPU. En fait, il suffit de remarquer que Van Der Vorst, lors de son implémentation de ce type de méthodes, à effectuer une partie de son code de calcul en langage assembleur, en l'occurrence, le produit matrice-vecteur. Cette pratique a eu pour effet de diminuer de façon significative le coût du produit matrice-vecteur. Cette opportunité, ne faisant pas partie de nos compétences, nous laisserons tomber l'étude de *PⁿNEWT*.

Au test T_5 , la Jacobienne dans la base nodale a sa diagonale tellement dominante, que l'inversion est quasiment immédiate. Aucun préconditionnement ne peut améliorer cette convergence, tous ceux que nous avons étudié, la retarde. A cause de la grande valeur de ρ , *IUNEW*T est forcément la méthode la plus lente.

En diminuant la viscosité (du test T_1 au test T_6) nous obtenons bien que l'inverse de la partie linéaire de la fonctionnelle non linéaire ne peut pas approcher de façon aussi précise l'inverse de la jacobienne, si bien que le résidu de *MGNEWT* stagne.

Analyse des tests relatifs à la fonction 4

Nous observons le même comportement des méthodes *NEWT* et *IUNEW*T que pour la fonction 3. Par contre, *MGNEWT* ne converge plus aussi rapidement. Nous remarquons que cette dégradation provient d'une stagnation du résidu de BI-CGSTAB dès la troisième inversion. Nous nous attendions bien à la diminution des performances de ce préconditionnement à mesure de l'augmentation des itérations de Newton. Le test T_3 est le plus difficile, d'une part à cause du second membre discret qui y intervient, et d'autre part par la faible valeur de ρ . Pour adapter notre préconditionnement, il serait intéressant de l'utiliser avec une méthode tronquée ou uniquement lors des premières inversions.

Nous présentons les résultats, aux mêmes tests, des méthodes de Newton tronquées obtenues en faisant varier le test de sortie de l'inversion en fonction de la valeur du résidu de Newton :

Tests	$T-NEWT$	$T-MGNEWT$	$T-IUNEWT$
T_1	(4)45/551	(4)29/16	(4)8/55
T_2	(4)8/117	(4)21/18	(4)14/97
T_3	(8)84/1060	(6)250 ^{s 2}	(6)17/118
T_4	(5)16/189	(4)40/21	(6)23/162
T_5	(4)4/44	(4)22/44	(7)148/1078
T_6	(10)132/1641	(10)393/164	(8)33/234

On peut dire, d'après le tableau des résultats ci-dessus que tous les temps CPU sont améliorés sauf lors du test T_5 , où la convergence est pratiquement immédiate par $NEWT$. Le nombre d'itérations de Newton est quelquefois plus grand que celui obtenu pour les vraies méthodes de Newton, ce qui est un résultat classique [DS83a]. Nous rajouterons juste que les résultats varient de la même manière face aux différents tests que les anciens. Comme prévu, $MGNEWT$ est améliorée mais n'est toujours pas adaptée pour les dernières itérations de Newton.

2.4.3 Résultats pour les méthodes de GMRES

On propose la même série de tests que ceux effectués pour les méthodes de Newton, rappelée dans le tableau ci-dessous :

Tests	Fonction	Discrétisation	ρ	ν
T_1	3	127^2	0	0.1
T_2	3	127^2	64	0.1
T_3	4	127^2	0	0.1
T_4	4	127^2	64	0.1
T_5	3	127^2	128	10^{-3}
T_6	3	127^2	0	0.01

2. Cas de stagnation du résidu de Newton à 5.10^{-9}

Ensuite, le tableau ci-après regroupe les résultats aux tests obtenus :

Tests	<i>GMRES</i>	<i>MGGMRES</i>	<i>IUGMRES</i>
T_1	(1073) 231	(515) 117	(73) 14
T_2	(182) 31	(52) 26	(144) 29
T_3	(1217) 244	(421) 144	(129) 25
T_4	(235) 42	(105) 48	(174) 34
T_5	(46) 7	(45) 14	(1023) 248
T_6	(1255) 272	(653) 209	(156) 30

Les chiffres entre parenthèses représentent le nombre d'évaluations de la fonctionnelle non linéaire F , et le nombre qui suit, le temps CPU sur CRAY-YMP mis pour obtenir la norme L^2 du Résidu de Newton inférieure à 1.10^{-9} . Comme dans le paragraphe 2.4.2, nous avons fixé un nombre maximal de 10 itérations pour l'inversion par Multigrilles de chaque système linéaire et le test de sortie porte sur la norme L^2 du résidu 10^{-14} .

Tout d'abord, explicitons les divers paramètres de l'implémentation :

la taille maximale de la dimension de l'espace de Krylov est $tmk = 50$, et l'étape d'initialisation se fait avec une projection sur deux vecteurs de base. Le critère d'augmentation/diminution de la dimension de la base orthonormale V de l'espace de Krylov sur lequel on fait la projection :

- si $F(\mathbf{u}^k) \leq 1,3 \cdot F(\mathbf{u}^{k-1})$ alors, on augmente la taille de deux
- si $F(\mathbf{u}^k) \geq 5 \cdot F(\mathbf{u}^{k-1})$ alors, on diminue la taille de deux.

De plus, pour l'implémentation de la procédure de "linesearch backtracking" (cf. page 118), nous avons pris les paramètres suivants :

$$l = 0, 1, u = 2, 0 \text{ et } \alpha = 10^{-9}$$

Nous avons limité le nombre d'itérations de backtracking à 4, et demandé au coefficient η d'être toujours supérieur à une valeur 5.10^{-3} . Le paramètre h_k du théorème 2.3.1 a été fixé comme l'a suggéré Dennis et Schnabel dans [DS83b], i.e. de la forme :

$$h_k = \sqrt{\epsilon} \cdot \max\{|\mathbf{u}_k|, typ\mathbf{u}_k\} \text{sign}(\mathbf{u}_k)$$

où, ϵ représente l'erreur relative en calculant $F(\mathbf{u})$, $typ\mathbf{u}_k$ la taille approximative de l'inconnue cherchée, et $\text{sign}(\mathbf{u}_k)$ le signe de l'inconnue $\mathbf{u}_k$. En l'absence de d'autres informations, ϵ sera choisi égal à la précision des moyens informatiques en notre possession, sinon, nous avons utilisé la proposition de Saad, à savoir :

$$h_k^{(i)} = \frac{\frac{\text{resid}^{(i-1)} \cdot \sqrt{\epsilon}}{\mathcal{M}ax_j \|\vec{H}_j\|^2} + h_k^{(i-1)}}{2}$$

Analyse des tests utilisant la Fonction 3

D'emblée, nous pouvons remarquer que le préconditionnement par Multi-grilles uniquement pour les premières itérations de Newton gagne en efficacité. En effet, il améliore la convergence aux tests T_1, T_2 , et T_6 . De plus, on s'aperçoit toujours de l'étroit lien de la rapidité de convergence et du conditionnement de la matrice Jacobienne simulée dans les méthodes de GMRES (cf. 2.4.2).

Pour le test T_1 , le facteur d'accélération de la convergence (f.a.c) est de 16,5 pour IUGMRES et de 2 pour MGGMRES; pour le test T_6 , le f.a.c est de 9,1 pour IUGMRES et de 1,3 pour MGGMRES. Ce qui nous permet de redire que le préconditionnement par les I.U. est mieux adapté aux tests utilisant la fonction 3 avec $\rho = 0$. Même au test T_2 , nous obtenons un f.a.c pour IUGMRES supérieur à 1 donc meilleur que celui obtenu pour les méthodes de Newton.

Analyse des tests utilisant la Fonction 4

Du fait des versions implémentées de l'algorithme GMRES (versions utilisant la procédure de line-search-backtracking), on pouvait s'attendre à de meilleurs résultats (pour les tests utilisant la fonction 4) que pour ceux provenant des méthodes de Newton. Ce qui a été le cas uniquement pour les algorithmes préconditionnés.

D'abord, les méthodes GMRES et MGGMRES ont le comportement attendu, à savoir une vitesse de convergence accrue quand le paramètre ρ , simulant l'évolution, augmente. Ensuite, la méthode MGGMRES converge pour les deux tests, par contre, elle n'améliore en rien la vitesse de convergence de GMRES, bien au contraire, elle la ralentit (f.a.c = 0,4 pour T_3 et f.a.c = 0,3 pour T_4). L'autre méthode IUGMRES nous donne des résultats plus intéressants, elle augmente la rapidité de convergence pour les tests T_3 et T_4 (f.a.c = 2,8 et 1.2).

2.5 Conclusions

Tout d'abord, il faut reconnaître l'intérêt de la méthode de GMRES pour ce genre de résolution pour les attraits suivants: sa modulabilité, sa robustesse et sa rapidité. En effet, il est aisé de changer de problème du fait qu'il suffit seulement de changer le sous-programme effectuant l'évaluation de la fonctionnelle non linéaire et d'adapter quelques paramètres. De plus, les résultats prouvent bien que les Inconnues Incrémentales sont un bon candidat pour accélérer la convergence lors de la résolution du problème (2.1). Elle demeure une méthode un peu plus lente que les méthodes de Newton tronquées, et plus gourmande en espace mémoire. En effet, l'implémentation utilisant une taille maximale du sous-espace de Krylov de 50 utilise deux fois plus d'espace mémoire que les implémentations de Newton-Bicgstab.

Nos résultats prouvent aussi qu'il est plus important de rechercher de bons préconditionnements pour résoudre le système linéaire intervenant dans des schémas de

type Newton que de s'intéresser aux différents types d'algorithmes de résolution.

Au sujet des différents préconditionnements que nous avons expérimenté, nous avons montré qu'il était préférable d'utiliser les Inconnues Incrémentales au lieu du préconditionnement par l'inverse de la partie linéaire quand le paramètre qui simulait l'évolution était petit (en valeur absolue) et que la préférence s'inversait quand ce paramètre devenait de l'ordre du nombre de points. Nous avons aussi montré que le préconditionnement par l'inverse de la partie linéaire s'avérait être performant pendant seulement le début de la résolution, et qu'il vallait mieux s'en passer par la suite.

Annexe A

Discrétisation du problème de Neumann

Tout d'abord précisons quelques notations en dimension 2 :

Soit $\Omega =]a, b[\times]c, d[$, on rappelle (1.2) :

$$\begin{cases} -\Delta u + u = f & \text{dans } \Omega \\ \frac{\partial u}{\partial n} = 0 & \text{sur } \partial\Omega \end{cases}$$

Soient $M \begin{pmatrix} x_k \\ y_l \end{pmatrix}$, $R_h = (x_k, y_l)_{k,l \in \mathbb{Z}}$

$$\sigma_{h_1 h_2}(M) = [x_k - \frac{h_1}{2}, x_k + \frac{h_1}{2}[\times [y_l - \frac{h_2}{2}, y_l + \frac{h_2}{2}[$$

$$\sigma_{h_1 h_2}(M, r) = \bigcup_{\substack{(i,j) \in \mathbb{Z}^2 \\ |i| + |j| \leq r}} \sigma_{h_1 h_2}(x_k + i\frac{h_1}{2}, y_l + j\frac{h_2}{2})$$

d'où l'espace d'approximation constitué des fonctions étagées suivantes :

$$w_{hM}(x, y) = \mathcal{X}_{\sigma_{h_1 h_2}(M)}(x, y)$$

soient

$$(\partial_1 \phi)(x, y) = \frac{1}{h_1} \left[\phi(x + \frac{h_1}{2}, y) - \phi(x - \frac{h_1}{2}, y) \right],$$

$$(\partial_2 \phi)(x, y) = \frac{1}{h_2} \left[\phi(x, y + \frac{h_2}{2}) - \phi(x, y - \frac{h_2}{2}) \right],$$

le problème approché s'écrit :

$$\begin{cases} \text{trouver } u_h = \sum_{M \in \Omega_h^1} u_h(M) w_{hM} & \text{tel que,} \\ a_h(u_h, v_h) = \langle l_h, v_h \rangle \quad \forall v_h \in \mathcal{Vect}(w_{hM}) \end{cases}$$

où :

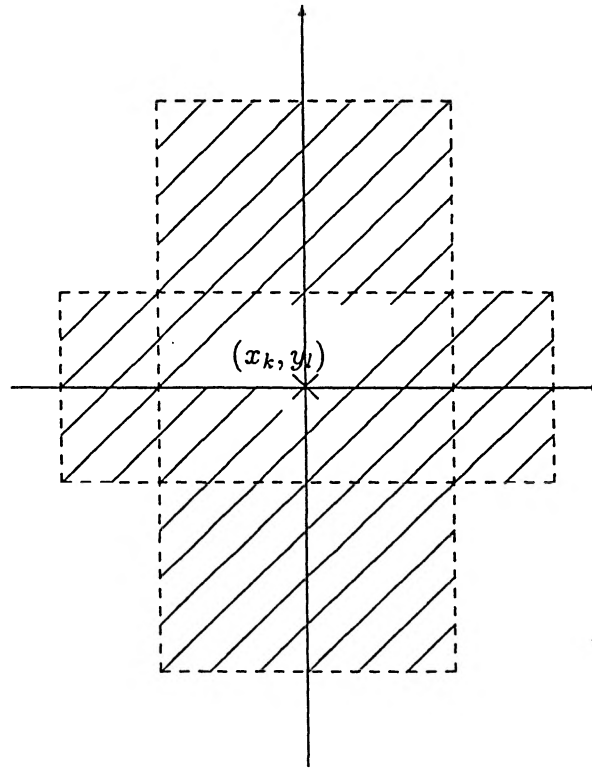
$$a_h(u_h, v_h) = \sum_{i=1}^2 \int_{\Omega} (\partial_i u_h(x, y)) (x, y) (\partial_i v_h(x, y)) (x, y) dx dy + \int_{\Omega} u_h v_h dx dy,$$

$$\langle l_h, v_h \rangle = \int_{\Omega} f(x, y) v_h(x, y) dx dy,$$

$$\Omega_h^1 = \{M/M \in R_h \text{ et } \sigma_{h_1 h_2}(M, 1) \cap \Omega \neq \emptyset\}$$

et,

$$\begin{aligned} \sigma_{h_1 h_2}(M, 1) = & \sigma_{h_1 h_2}(M) \cup \sigma_{h_1 h_2}(x_k - \frac{h_1}{2}, y_l) \cup \sigma_{h_1 h_2}(x_k + \frac{h_1}{2}, y_l) \\ & \cup \sigma_{h_1 h_2}(x_k, y_l + \frac{h_2}{2}) \cup \sigma_{h_1 h_2}(x_k, y_l - \frac{h_2}{2}) \end{aligned}$$



$$\sigma_{h_1 h_2}(M, 1)$$

$$\text{en posant : } P \begin{pmatrix} x_i \\ y_j \end{pmatrix} \quad A(x_i) = \left[x_i - \frac{h_1}{2}, x_i + \frac{h_1}{2} \right[,$$

$$B(y_j) = \left[y_j - \frac{h_2}{2}, y_j + \frac{h_2}{2} \right[, \quad C^+(x_k) = [x_k, x_{k+1}[, \quad C^-(y_j) = [y_{j-1}, y_j[$$

$$a_h(w_h M, w_h P) =$$

$$\begin{aligned} & \frac{1}{h_1 h_2} \int_a^b \int_a^b \left\{ \mathcal{X}_{C^-(x_k)} \times B(y_l)(x, y) - \mathcal{X}_{C^+(x_k)} \times B(y_l)(x, y) \right\} \\ & \quad \times \left\{ \mathcal{X}_{C^-(x_i)} \times B(y_j)(x, y) - \mathcal{X}_{C^+(x_i)} \times B(y_j)(x, y) \right\} dx dy \\ & \frac{1}{h_1 h_2} \int_a^b \int_a^b \left\{ \mathcal{X}_{A(x_i)} \times C^-(y_j)(x, y) - \mathcal{X}_{A(x_i)} \times C^+(y_j)(x, y) \right\} \\ & \quad \times \left\{ \mathcal{X}_{A(x_k)} \times C^-(y_l)(x, y) - \mathcal{X}_{A(x_k)} \times C^+(y_l)(x, y) \right\} dx dy \\ & + \int_a^b \int_a^b \mathcal{X}_{A(x_i) \times A(y_j) \cup A(x_k) \times A(y_l)}(x, y) dx dy \end{aligned}$$

Bibliographie

- [ARN51] W.E. ARNOLDI. The principle of minimized iteration in the solution of the matrix eigenvalue problem. *Quart. Applied Mathematics*, 9:17–29, 1951.
- [BPX90] J.H. BRAMBLE, J.E. PASCIAK, and J. XU. Parallel multilevel preconditioners. *Mathematics of Computation*, 55(191):1–22, 1990.
- [BS90] P.N. BROWN and Y. SAAD. Hybrid Krylov methods for nonlinear systems of equations. *SIAM J. Sci. Stat. Comput.*, 11(3):450–481, 1990.
- [CEA64] J. CEA. Approximation variationnelle des problèmes aux limites. *Ann. Inst. Fourier*, 14:345–444, 1964.
- [CJ84] T.F. CHAN and K.R. JACKSON. Nonlinearly preconditioned Krylov subspace methods for discrete Newton algorithms. *SIAM J. Sci. Stat. Comput.*, 5(3):533–542, 1984.
- [CHE93] J.P. CHEHAB. *Méthode des Inconnues Incrémentales. Applications au calcul des bifurcations*. PhD thesis, Université Paris-XI, 01/1993.
- [CMT94] M. CHEN, A. MIRANVILLE, and R. TEMAM. Incremental Unknowns in finite difference in three space dimensions. *Computational and Applied Mathematics*, 14(3):217, 1995.
- [CT91] M. CHEN and R. TEMAM. Incremental Unknowns for solving partial differential equations. *Numer. Math.*, 59:255–271, 1991.
- [CT93] M. CHEN and R. TEMAM. Incremental Unknowns in finite differences: Condition number of the matrix. *SIAM J. Matrix Anal. Appl.*, 14(2):432–455, 1993.
- [DES82] R.S. DEMBO, S.C. EISENSTAT, and T. STEIHAUG. Inexact Newton methods. *SIAM J. Numer. Anal.*, 19(2):400–408, 1982.
- [DS83a] R.S. DEMBO and T. STEIHAUG. Truncated-Newton algorithms for large-scale unconstrained optimization. *Mathematical Programming*, 26:190–212, 1983.

- [DS83b] J.E. DENNIS and R.B. SCHNABEL. *Numerical methods for unconstrained optimization and nonlinear equations*. Prentice-Hall, Englewood Cliffs NJ, 1983.
- [FLE76] R. FLETCHER. Conjugate Gradient methods for indefinite systems. *Lecture Notes in Mathematics*, 506:73–89, 1976.
- [FLE88] C.A.J. FLETCHER. *Computational techniques for fluid dynamics*, volume 1 of *Springer Series in Computational Physics*. Springer-Verlag, 1988.
- [GL90] G.H. GOLUB and VAN LOAN. *Matrix computations*. Johns Hopkins Series in the Mathematical Sciences. The Johns Hopkins University Press, 1990.
- [GOY94] O. GOYON. *Résolution numérique de problèmes stationnaires et évolutifs non linéaires par la méthode des Inconnues Incrémentales*. PhD thesis, Université Paris-XI, 12/1994.
- [GP95] O. GOYON and P. POULLET. Numerical resolutions using Incremental Unknowns methods of steady-state Navier-Stokes like equations. Prépublication de l'Université PARIS-SUD, 1995. 95-87.
- [GRMM92] M.A. GOMES-RUGGIERO, J.M. MARTINEZ, and A.C. MORETTI. Comparing algorithms for solving sparse nonlinear systems of equations. *SIAM J. Sci. Stat. Comput.*, 13(2):459–483, 1992.
- [HAC85] W. HACKBUSCH. *Multigrid methods and applications*. Springer-Verlag, 1985.
- [HIR92] C. HIRSCH. *Fundamentals of Numerical Discretization*, volume 1 of *Numerical Computation of Internal and External Flows*. John Wiley & Sons, 1992.
- [MED95] T. TACHIM MEDJO. *Sur les formulations vitesse-vorticité pour les équations de Navier-Stokes en dimension deux - Implémentation des Inconnues Incrémentales Oscillantes dans les équations de Navier-Stokes et de réaction-diffusion*. PhD thesis, Université Paris-XI, 06/ 1995.
- [NAS84] S.G. NASH. Newton-type minimization via Lanczos method. *SIAM J. Numer. Anal.*, 21(4):770–788, 1984.
- [PAS92] F. PASCAL. *Méthodes de Galerkin non linéaires en discrétisation par éléments finis et pseudo-spectrale. Application à la Mécanique des Fluides*. PhD thesis, Université Paris-XI, 01/ 1992.

- [SON89] P. SONNEVELD. Cgs: A fast Lanczos-type solver for nonsymmetric linear systems. *SIAM J. Sci. Stat. Comput.*, 10:36–52, 1989.
- [SS86] Y. SAAD and M.H. SCHULTZ. A generalized minimal residual algorithm for solving nonsymmetric linear systems. *SIAM J. Sci. Stat. Comput.*, 7(3):856–869, 1986.
- [TEM66] R. TEMAM. Sur l’approximation des équations de Navier-Stokes. *C. R. Acad. Sci. Paris, Sér. A*, 262:219–221, 1966.
- [TEM84] R. TEMAM. *Navier-Stokes Equations*. North-Holland, 1984.
- [TEM90] R. TEMAM. Inertial manifolds and multigrid methods. *SIAM J. Math. Anal.*, 21(1):154–178, 1990.
- [VOR82] H.A. VAN DER VORST. A vectorizable variant of some ICCG methods. *SIAM J. Sci. Stat. Comput.*, 3(3):350–356, 1982.
- [VOR92] H.A. VAN DER VORST. Bi-cgstab: A fast and smoothly converging variant of Bi-CG for the solution of nonsymmetric linear systems. *SIAM J. Sci. Stat. Comput.*, 13(2):631–644, 1992.
- [WES82] P. WESSELING. Theoretical and practical aspects of a multigrid method. *SIAM J. Sci. Stat. Comput.*, 3:387–407, 1982.
- [YSE86] H. YSERANTANT. On the multi-level splitting of finite element spaces. *Numerische Mathematik*, 49:379–412, 1986.

Partie III

RESOLUTION DES EQUATIONS DE NAVIER-STOKES BI-DIMENSIONNELLES

Introduction

This part is divided into three chapters.

The first chapter presents the equations constituting the evolutive Navier-Stokes problem which models the two-dimensional incompressible flow of a Newtonian viscous fluid, and the physical problem which is treated. We precise our way to visualize the numerical results.

In the second chapter, we develop a numerical resolution method based on a staggered finite differences discretization, like the Marker And Cell method (MAC) one, that we valid for the squared lid driven cavity problem for Reynolds numbers until 5000. At each time step, we solve a generalized Stokes problem by an orthogonal projection algorithm which consists in apply a conjugate gradient method onto the kernel of the discrete divergence operator. The resolution of the pressure problem requires a preconditioner, then we use the classical multigrid algorithm used in the first part to perform the inner iterations. Several results prove that our resolution method is efficient and robust.

The last chapter deals with the development of a multilevel method based on a concept emerging from new results from the dynamical systems theory. We set in a case of a solution computed for a quasi periodical state, the Reynolds number is 10^4 and the discretization is fine and still uniform (257^2 points). Although, basis ideas are valid for asymptotic state, we try developing the method in the transient state of the flow.

The first step is dedicated to introduce new I.U. which have the special feature giving the staggered property of the M.A.C. mesh, onto the all grids - we called them Staggered Incremental Unknowns (S.I.U.), which are of second order and permit to hierarchize the velocity field and the pressure. Moreover, an analysis of the norm of many quantities coming from the discrete system enable us to obtain nonlinear dynamic indicators.

We think that it is possible to freeze in time the small-scale quantities in order to compute quickly the large-scale one. By frozen the nonlinear interaction term, we project the discrete system on the coarse grid. The restriction choice stands before solving the coarse system.

Finally, we propose a scheme class of Non Linear Galerkin type.

Chapitre 1

Présentation du problème

1.1 Rappels théoriques et formulation

Les équations de Navier-Stokes sont des équations provenant de la mécanique des milieux continus qui modélisent l'écoulement d'un fluide newtonien. Nous limiterons notre étude au cas d'écoulements bi-dimensionnels visqueux incompressibles. L'incompressibilité signifiant une densité constante du fluide, deux variables sont nécessaires à la description de son état :

$\mathbf{u}(\mathbf{x}, t)$ est le vecteur vitesse de la particule au point $\mathbf{x}$ à l'instant t ,

$p(\mathbf{x}, t)$ est la pression du fluide au point $\mathbf{x}$ à l'instant t ,

$\mathbf{x} \in \Omega$ ouvert borné de $\mathbb{R}^2$ et $t \in \mathbb{R}^+$.

Par le *principe fondamental de la dynamique* et la *conservation de la masse*, nous obtenons les équations dites de Navier-Stokes, régissant l'écoulement bi-dimensionnel, dans Ω , d'un fluide newtonien visqueux incompressible soumis à la densité des forces extérieures $\mathbf{f}$:

$$\begin{cases} \rho \frac{d\mathbf{u}}{dt} - \mu \Delta \mathbf{u} + \nabla p = \mathbf{f} & \text{dans } \Omega \times \mathbb{R}^+ \\ \nabla \cdot \mathbf{u} = 0 & \text{dans } \Omega \times \mathbb{R}^+ \end{cases}$$

où ρ représente la densité constante du fluide,

μ le coefficient de viscosité dynamique,

et $\frac{d}{dt}$ représente la dérivée particulière.

En fixant une longueur, un temps et une vitesse caractéristiques de l'écoulement, on présente la forme adimensionnée des équations en reprenant les mêmes notations pour les quantités réduites :

$$\begin{cases} \frac{\partial \mathbf{u}}{\partial t} - \frac{1}{Re} \Delta \mathbf{u} + (\mathbf{u} \cdot \nabla) \mathbf{u} + \nabla p = \mathbf{f} & \text{dans } \Omega \times \mathbb{R}^+ \\ \nabla \cdot \mathbf{u} = 0 & \text{dans } \Omega \times \mathbb{R}^+ \end{cases} \quad (1.1)$$

La constante Re est le nombre de Reynolds (sans dimension) qui est fonction de la viscosité dynamique et de la densité du fluide donc caractérisera son état.

Pour obtenir la formulation vitesse-pression, il reste à munir le système (1.1) de conditions initiales et de conditions aux limites. Nous considérerons une forme de conditions aux limites de type Dirichlet.

$$\mathbf{u}(\mathbf{x}, 0) = \mathbf{u}_0(\mathbf{x}) \quad \forall \mathbf{x} \in \Omega, \quad (1.2)$$

où $\mathbf{u}_0$ doit nécessairement vérifier $\nabla \cdot \mathbf{u}_0 = 0$, et

$$\mathbf{u}(\mathbf{x}, t) = \mathbf{g}(\mathbf{x}, t) \quad \forall \mathbf{x} \in \partial\Omega \quad (1.3)$$

$\mathbf{g}$ doit aussi vérifier la condition de flux nul: $\int_{\partial\Omega} \mathbf{g} \cdot \mathbf{n} d\sigma = 0$, où $\mathbf{n}$ est la normale extérieure au domaine Ω .

Il existe bien d'autres formulations pour les équations de Navier-Stokes, comme la formulation fonction de courant-vorticité utilisé par [GG82],[GOY94], la formulation vitesse-tourbillon utilisé par [DAU92],[MED95] mais l'intérêt premier de la formulation vitesse-pression est de permettre l'extension la plus facile possible aux écoulements tri-dimensionnels.

Il nous est nécessaire aussi de rappeler aussi le problème de Stokes, considéré comme problème approché de Navier-Stokes pour des écoulements stationnaires s'effectuant à très faible nombre de Reynolds, s'écrivant sous la forme suivante dans le cas de conditions aux limites non homogènes :

$$\left\{ \begin{array}{ll} -\frac{1}{Re} \Delta \mathbf{u} + \nabla \mathbf{u} = \mathbf{f} & \text{dans } \Omega \\ \nabla \cdot \mathbf{u} = 0 & \text{dans } \Omega \\ \mathbf{u} = \mathbf{g} & \text{sur } \partial\Omega \end{array} \right. \quad (1.4)$$

Le résultat d'existence et d'unicité de ce problème (1.4) est classique, soit

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

on obtient une unique solution de (1.4) $(\mathbf{u}, p) \in \mathcal{H} \times L^2(\Omega)/\mathbb{R}$.

Pour présenter le résultat d'existence et d'unicité de la solution des équations de Navier-Stokes incompressibles homogènes, rappelons au préalable quelques notations :

soit Ω un ouvert borné de $\mathbb{R}^2$, de frontière localement lipschitzienne, posons :

$$\begin{aligned} \mathcal{V} &= \{\mathbf{u} \in \mathcal{D}(\Omega), \operatorname{div} \mathbf{u} = 0\}, \\ V &= \overline{\mathcal{V}}^{H_0^1(\Omega)}, \\ H &= \overline{\mathcal{V}}^{L^2(\Omega)}, \end{aligned}$$

si l'on munit V (respectivement H) du produit scalaire induit par $H_0^1(\Omega)$ noté $((.,.))$ (respectivement $L^2(\Omega)$ et noté $(.,.)$) on obtient deux espaces de Hilbert.

On définit aussi la forme trilinéaire, continue et antisymétrique sur V par,

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

Nous énonçons une première formulation variationnelle obtenue comme au chapitre 2 après intégration par parties, et une deuxième équivalente à la première faisant intervenir la pression :

Pour $f \in L^2(0, T; V')$ et $\mathbf{u}_0 \in H$,
chercher $\mathbf{u} \in L^2(0, T; V)$ telle que

$$\begin{cases} \frac{d}{dt}(\mathbf{u}, \mathbf{v}) + \frac{1}{Re}((\mathbf{u}, \mathbf{v})) + b(\mathbf{u}, \mathbf{u}, \mathbf{v}) = \langle f, \mathbf{v} \rangle, \quad \forall \mathbf{v} \in V \\ \mathbf{u}(0) = \mathbf{u}_0, \end{cases} \quad (1.5)$$

Pour $f \in L^2(0, T; V')$ et $\mathbf{u}_0 \in H$,
chercher $(\mathbf{u}, p) \in L^2(0, T; H_0^1(\Omega)) \times L^2(0, T; L^2(\Omega))$ tel que

$$\begin{cases} \frac{d}{dt}(\mathbf{u}, \mathbf{v}) + \frac{1}{Re}((\mathbf{u}, \mathbf{v})) + b(\mathbf{u}, \mathbf{u}, \mathbf{v}) + (\nabla p, \mathbf{v}) = \langle f, \mathbf{v} \rangle, \quad \forall \mathbf{v} \in H_0^1(\Omega) \\ (q, \operatorname{div} \mathbf{u}) = 0, \quad \forall q \in L^2(\Omega) \\ \mathbf{u}(0) = \mathbf{u}_0. \end{cases} \quad (1.6)$$

Nous sommes en mesure d'énoncer le résultat suivant qui est prouvé dans [TEM84].

Théorème 1.1.1

Pour $f \in L^2(0, T; V')$ et $\mathbf{u}_0 \in H$, il existe une unique solution $\mathbf{u} \in L^\infty(0, T; H) \cap L^2(0, T; V)$ satisfaisant (1.5).
De plus, $\mathbf{u}$ est faiblement continue de $[0, T]$ dans H ($\mathbf{u} \in C_w([0, T]; H)$).

1.2 Problèmes modèles et Visualisation

1.2.1 La cavité carrée entraînée

Nous décrivons dans ce paragraphe le premier problème modèle qui va nous permettre de tester nos implémentations. Nous choisirons pour débiter un domaine carré, la configuration est celle d'un fluide visqueux incompressible s'écoulant dans un canal de section carrée infiniment long pour que l'on puisse supposer que la vitesse et la pression s'annule dans la direction longitudinale, et que l'écoulement soit bi-dimensionnel. Avec des conditions aux bord de type entraînement, l'écoulement s'effectue sans apport direct de forces extérieures ($f = 0$).

Nous considérerons deux types de conditions aux bords, l'un (USLC acronyme anglais de *Unregularized Squared Lid-driven Cavity*) permettant d'avoir un écoulement présentant des singularités aux deux coins supérieurs et l'autre (RSLC) permettant d'avoir un écoulement plus régulier.

USLC

$$\mathbf{u} = \begin{pmatrix} u(\mathbf{x}, t) \\ v(\mathbf{x}, t) \end{pmatrix} = \begin{pmatrix} 1 \\ 0 \end{pmatrix}$$

RSLC

$$\mathbf{u} = \begin{pmatrix} u(\mathbf{x}, t) \\ v(\mathbf{x}, t) \end{pmatrix} = \begin{pmatrix} (1 - (2x - 1)^2)^2 \\ 0 \end{pmatrix}$$

Nous comparerons nos résultats avec ceux obtenus par d'autres auteurs et regroupés dans [GG82], [BJ90], [PAS92] et [GOY94] pour (USLC) et [SHE87], [SHE91], [PAS92] pour (RSLC).

La figure (3.1) présente la configuration de référence ainsi que la nomenclature habituelle pour le problème de la cavité entraînée.

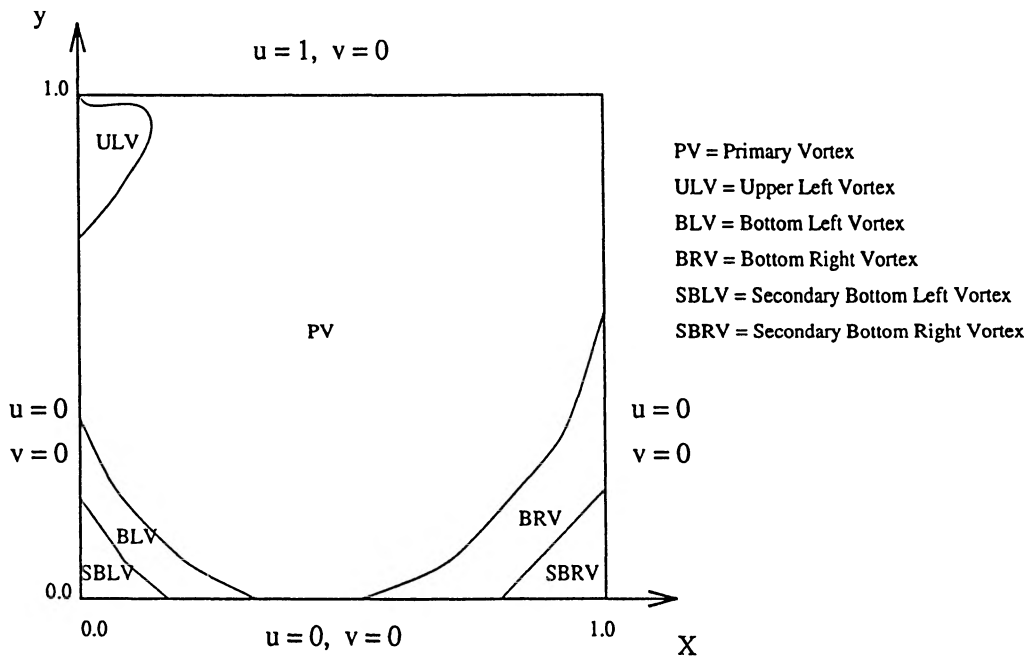


FIG. 1.1 - Configuration de la cavité carrée entraînée

Nous utiliserons cette terminologie pour présenter nos résultats numériques.

1.2.2 Visualisation des résultats

Pour permettre d'apprécier qualitativement les solutions de chacun des problèmes résolus, il est nécessaire de visualiser différentes quantités [GCLU84]. Nous énumérons ci-dessous nos indicateurs et les raisons qui ont motivé leur choix.

Isovorticité

Il s'agit des lignes de niveaux de la fonction scalaire la vorticité ($\omega(\mathbf{x}, t)$) dont nous rappelons la définition.

Comme nous avons un mouvement plan, le rotationnel de la vitesse est uniquement porté par l'axe Oz , et :

$$\omega(\mathbf{x}, t) = (\nabla \wedge \mathbf{u}(\mathbf{x}, t)).z$$

Ce tracé est fort utile dans notre contexte. D'une part, il renseigne sur la dynamique des tourbillons, et d'autre part, il met surtout en évidence, par des oscillations, les éventuelles faiblesses de la méthode de résolution.

Lignes de courant

Ce sont les lignes de niveaux de la fonction de courant $\psi(\mathbf{x}, t)$. Cette fonction scalaire, dont le tracé renseigne sur les différentes structures présentes dans l'écoulement, est définie par :

$$\begin{aligned} u(\mathbf{x}, t) &= \frac{\partial \psi}{\partial y}(\mathbf{x}, t) \\ v(\mathbf{x}, t) &= -\frac{\partial \psi}{\partial x}(\mathbf{x}, t), \end{aligned}$$

et calculée, en substituant $\mathbf{u}$ dans la définition de la vorticité, par l'équation de Poisson (1.7) avec des conditions aux limites de type Dirichlet homogène.

$$-\Delta \psi = \omega \quad \text{dans } \Omega_h \quad (1.7)$$

Champ de vecteurs de la vitesse

Il s'agit du graphique couramment représenté, nous utiliserons AVS¹ pour l'effectuer. Nous avons tenté de le compléter en ajoutant des zooms dans les zones faisant apparaître les structures de plus basse échelle. Pour des raisons de clarté nous oterons de nos tracés, pour les problèmes *USLC*, la ligne supérieure correspondant aux points où est imposée la condition limite.

Isobares

La représentation des lignes de niveaux de la pression est aussi utile. Elle permet de rendre compte de la bonne approximation de la condition aux limites du problème de la pression.

1. Logiciel de visualisation commercialisé par TETHYS

Profils médians de la vitesse

Les profils que nous dessinerons sont les coupes médianes de chacune des composantes de la vitesse : à $x = 0.5$ pour u , et $y = 0.5$ pour v . Par ces coupes, on peut s'apercevoir nettement de la vitesse de transition du caractère de l'écoulement. Proche du bord on s'attend à noter l'apparition de couches limites, tandis qu'en s'en éloignant, une zone plus régulière symboliserait la partie principale de l'écoulement.

Chapitre 2

Méthode de différences finies et méthode M.A.C.

Nous présentons, dans ce chapitre une méthode de différences finies pour résoudre les équations de Navier-Stokes bi-dimensionnelles incompressibles. Le paragraphe 2.1 donne notre choix de discrétisations spatiale et temporelle. Pour l'espace, il s'agit d'une discrétisation mixte dite de grilles décalées identiques à celle utilisée par la méthode MAC (nous dirons de type MAC).

L'algorithme de résolution est un algorithme de projection qui est décrit au paragraphe 2.3. L'implémentation que nous en faisons nécessite l'utilisation d'un préconditionnement.

Nous appliquons cette méthode au problème modèle énoncé dans le chapitre précédent et discutons des résultats au paragraphe 2.4.

L'implémentation de cette méthode s'est faite par le développement d'un code de calcul de référence qui nous permettra de mener une analyse "a posteriori" fine pour une hiérarchisation détaillée au chapitre suivant. Cette étude doit être considérée comme la base du développement d'une méthode de type Galerkin non linéaire.

2.1 Problème discret

Après avoir présenté les différents opérateurs discrets, nous détaillons la méthode en utilisant une discrétisation de l'équation du moment par un schéma semi-implicite pour le terme dissipatif et explicite pour les termes convectifs. Ensuite nous mentionnons notre choix de la condition aux limites, pour le problème faisant intervenir la pression. Une étude de stabilité linéaire terminera cette section, pour notre schéma numérique.

Une des méthodes les plus communément utilisées est la discrétisation (M.A.C.) introduite par Harlow et Welch [HW65] et étudiée par de nombreux auteurs et ingénieurs [PT83], [DAV83], [FLE88], [HIR92]. D'un point de vue théorique, nous mentionnons l'estimation d'erreur provenant de V. Girault dans [GIR76], qui trouve que

la méthode des éléments finis adaptée à la discrétisation de type MAC fournie une méthode d'ordre 1, pour approcher les équations de Navier-Stokes incompressibles stationnaires.

2.1.1 Schéma numérique

Cette méthode utilise une discrétisation qui est caractérisée par l'utilisation de grilles décalées et la résolution d'une équation de Poisson pour la pression provenant de l'équation discrète du moment. La figure 2.2 donne la configuration (MAC) pour les points intérieurs.

Soient Δx , Δy les pas d'espace dans chacune des directions, nous approchons par différences finies centrées les opérateurs de dérivation et notons :

$$(div_h \mathbf{u}_h)_{i,j} = \frac{1}{\Delta x}(u_{i+1/2,j} - u_{i-1/2,j}) + \frac{1}{\Delta y}(v_{i,j+1/2} - v_{i,j-1/2})$$

$$\nabla_{hx} p_{i+1/2,j}^n = \frac{1}{\Delta x}(p_{i+1,j}^n - p_{i,j}^n), \quad \nabla_{hy} p_{i,j+1/2}^n = \frac{1}{\Delta y}(p_{i,j+1}^n - p_{i,j}^n)$$

$$\begin{aligned} \Delta_h u_{i+1/2,j}^n &= \frac{1}{\Delta x^2}(u_{i+3/2,j}^n - 2u_{i+1/2,j}^n + u_{i-1/2,j}^n) \\ &\quad + \frac{1}{\Delta y^2}(u_{i+1/2,j+1}^n - 2u_{i+1/2,j}^n + u_{i+1/2,j-1}^n), \end{aligned}$$

où ∇_h et div_h sont évalués aux points intérieurs des maillages (cf. FIG. 2.3) $(i + 1/2, j) \in T_h^1$ et $(i, j + 1/2) \in T_h^2$.

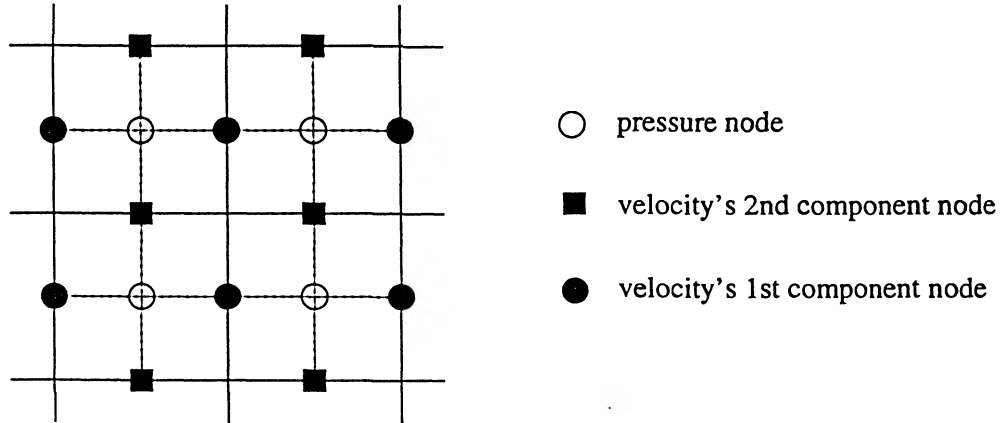


FIG. 2.2 - Configuration MAC pour les points intérieurs

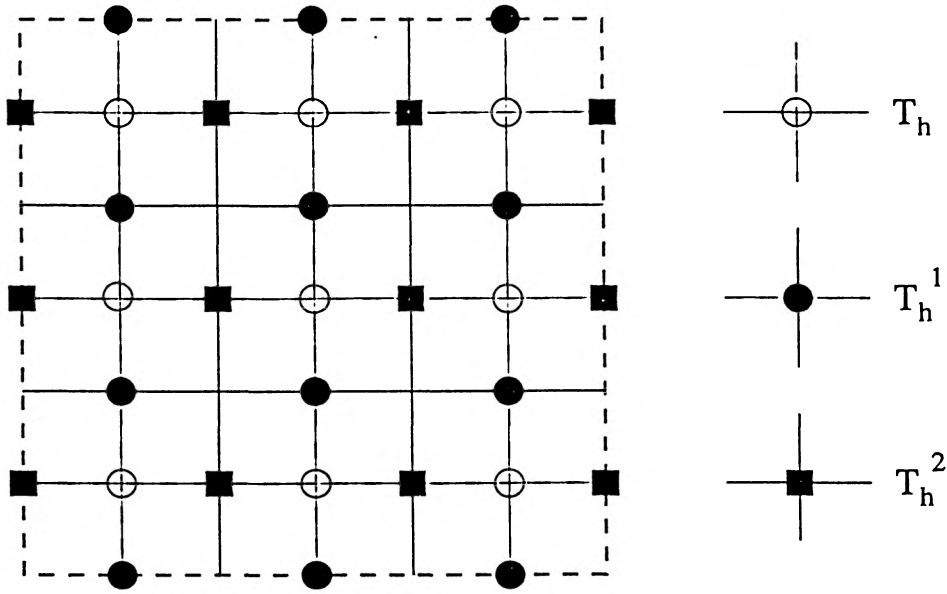


FIG. 2.3 - Les trois maillages différents

Nous avons choisi d'approcher l'opérateur laplacien par le schéma centré habituel à 5 points, d'ordre 2. Concernant la discrétisation temporelle, nous avons choisi de traiter le terme de diffusion de façon semi-implicite par le schéma de Cranck-Nicholson et le terme non-linéaire de manière explicite (par Euler d'abord et Adams-Bashforth ensuite). Nous reviendrons sur ce dernier point quand nous vérifierons la stabilité de notre schéma. Pour le schéma de Cranck-Nicholson-Euler, après projection sur chacune des directions spatiales, nous avons le système discret suivant, pour les points intérieurs au domaine Ω .

$$\left\{ \begin{array}{l} \frac{1}{\Delta t}(u_{i+1/2,j}^{n+1} - u_{i+1/2,j}^n) + \nabla_{h_x} p_{i+1/2,j}^{n+1} + \frac{1}{2\Delta x} u_{i+1/2,j}^n (u_{i+3/2,j}^n - u_{i-1/2,j}^n) \\ - \frac{1}{Re} (\Delta_h u_{i+1/2,j}^{n+1} + \Delta_h u_{i+1/2,j}^n) + \frac{1}{2\Delta y} \hat{v}_{i+1/2,j}^n (u_{i+1/2,j+1}^n - u_{i+1/2,j-1}^n) = 0 \\ \frac{1}{\Delta t}(v_{i,j+1/2}^{n+1} - v_{i,j+1/2}^n) + \nabla_{h_y} p_{i,j+1/2}^{n+1} + \frac{1}{2\Delta y} v_{i,j+1/2}^n (v_{i,j+3/2}^n - v_{i,j-1/2}^n) \\ - \frac{1}{Re} (\Delta_h v_{i,j+1/2}^{n+1} + \Delta_h v_{i,j+1/2}^n) + \frac{1}{2\Delta x} \hat{u}_{i,j+1/2}^n (v_{i+1,j+1/2}^n - v_{i-1,j+1/2}^n) = 0 \\ \nabla_{h_x} u_{i,j}^{n+1} + \nabla_{h_y} v_{i,j}^{n+1} = 0, \end{array} \right. \quad (2.1)$$

où Δt est le pas de temps et les variables $\hat{u}_{i,j+1/2}^n$ et $\hat{v}_{i+1/2,j}^n$ sont obtenus par interpolation,

$$\hat{u}_{i,j+1/2}^n = \frac{1}{4}(u_{i+1/2,j}^n + u_{i+1/2,j+1}^n + u_{i-1/2,j}^n + u_{i-1/2,j+1}^n)$$

$$\hat{v}_{i+1/2,j}^n = \frac{1}{4}(v_{i+1,j+1/2}^n + v_{i,j+1/2}^n + v_{i,j-1/2}^n + v_{i+1,j-1/2}^n)$$

Remarque 6

A cause de la spécificité du maillage décalé MAC, nous introduisons une norme discrète équivalente à la norme L^2 que nous nommons norme de W_h (ou W -norm), définie comme ci-dessous :

$$(u_h, \theta_h)_h = \Delta x \Delta y \left(\sum_{(i,j) \in T_h^1} u_{i,j} \theta_{i,j}^1 + \sum_{(i,j) \in T_h^2} v_{i,j} \theta_{i,j}^2 \right)$$

2.1.2 Conditions aux limites

En substituant $u_{i+1/2,j}^{n+1}$ et $v_{i,j+1/2}^{n+1}$ dans la relation d'incompressibilité discrète, on retrouve l'équation de Poisson de la pression (en anglais Pressure Poisson Equation). Une difficulté réside dans la condition aux limites à imposer à cette équation. Gresho et Sani assurent qu'il faut prendre une condition de type Neumann non homogène faisant intervenir la composante normale de l'équation du moment, pour obtenir une bonne solution à tout $t \geq 0$ [GS87]. Ce calcul étant tout de même assez coûteux, nous avons préféré approcher comme le suggèrent Peyret et Taylor dans [PT83] par des conditions aux limites homogènes pour le gradient de la pression (ce qui se fait de façon implicite avec notre algorithme de résolution) pour notre problème modèle. D'autres auteurs, comme Fletcher, affirment que ce choix est cohérent pour des simulations d'écoulements à grand nombre de Reynolds [FLE88], et qu'il permet aussi de bien poser mathématiquement le problème pour des temps $t > 0$ [GS87].

Il faut bien se rendre compte que la discrétisation MAC est bien adaptée au traitement de la pression car elle n'est définie qu'à l'intérieur du domaine, mais présente quelques inconvénients concernant celui de la vitesse. En effet, si l'on regarde la figure 2.4, le calcul de v_1 ou u_1 nécessite l'évaluation des termes convectifs et dissipatifs faisant intervenir des points extérieurs au domaine $\bar{\Omega}$.

Nous avons utilisé une technique d'extrapolation d'ordre 0 qui ne semble pas détériorer complètement la précision de la solution numérique [PT83].

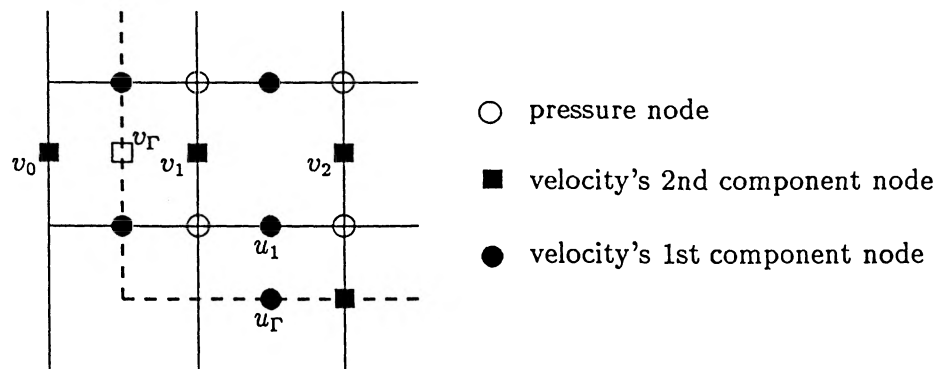


FIG. 2.4 - Traitement des conditions aux limites pour un maillage MAC

En posant, $v_0 = 2v_\Gamma - v_1$, (avec les notations de la Fig. 2.4) et en utilisant un schéma centré, on obtient comme approximation des dérivées les expressions suivantes :

$$\begin{aligned} \left. \frac{\partial v}{\partial x} \right|_1 &= \frac{1}{2\Delta x} (v_2 + v_1 - 2v_\Gamma) \\ \text{et } \left. \frac{\partial^2 v}{\partial x^2} \right|_1 &= \frac{4}{3\Delta x^2} (v_2 - 3v_1 + 2v_\Gamma). \end{aligned}$$

Par un schéma traitant explicitement le terme non linéaire comme pour le schéma donnant le système discret (2.1), nous obtenons à chaque itération, un problème dit de Stokes généralisé à résoudre. Soit,

$$\begin{cases} \alpha \mathbf{u} - \frac{1}{Re} \Delta \mathbf{u} + \nabla p = \mathbf{j} & \text{dans } \Omega \times [0, T] \\ \nabla \cdot \mathbf{u} = 0 & \text{dans } \Omega \times [0, T] \\ \mathbf{u} = \mathbf{g} & \text{sur } \Gamma \times [0, T], \end{cases} \quad (2.2)$$

avec une constante $\alpha > 0$, représentant l'inverse d'un pas de temps. Comme nous le verrons par la suite, la stabilité de notre schéma explicite nous limite à de petits pas de temps donc fait intervenir de grands α .

2.1.3 Stabilité linéaire

Un point essentiel qui mérite notre attention est la stabilité de notre schéma numérique. Plusieurs méthodes existent pour l'analyser mais se limitent souvent aux problèmes linéaires [HIR92]. De plus, il est souvent très difficile de tenir compte des conditions aux limites. Toutefois, ce type d'analyse permet, pour un problème non-linéaire, d'obtenir des conditions nécessaires de stabilité, ce qui nous semble intéressant même si nous ne pouvons traiter qu'un problème linéaire ayant des conditions aux limites périodiques. Dans cette optique, nous avons choisit une analyse de Von Neumann unidimensionnelle pour l'équation d'advection-diffusion suivante munie de conditions aux limites périodiques dans un domaine $\Omega = [0, 2\pi]^2$:

$$\begin{cases} \frac{\partial u}{\partial t} + a \cdot \frac{\partial u}{\partial x} + b \cdot \frac{\partial u}{\partial y} - \nu \Delta u = f & \text{dans } \Omega \\ u(x + 2\pi e_i, t) = u(x, t) & \text{sur } \partial\Omega \end{cases} \quad (2.3)$$

Nous discrétisons l'équation (2.3) par le schéma numérique que nous avons utilisé pour la résolution des équations de Navier-Stokes, soit Cranck-Nicholson sur le terme diffusif et d'Adams-Bashforth pour le terme d'advection. Chacun des opérateurs de dérivation sera approché à l'aide de différences finies centrées, d'où le système discret suivant :

$$\begin{aligned}
u_{i,j}^{n+1} - u_{i,j}^n &= -a \cdot \frac{\Delta t}{4\Delta x} (3u_{i+1,j}^n - 3u_{i-1,j}^n - u_{i+1,j}^{n-1} + u_{i-1,j}^{n-1}) \\
&\quad - b \cdot \frac{\Delta t}{4\Delta y} (3u_{i,j+1}^n - 3u_{i,j-1}^n - u_{i,j+1}^{n-1} + u_{i,j-1}^{n-1}) \\
&\quad - \frac{\Delta t}{2Re\Delta x^2} (2u_{i,j}^{n+1} - u_{i-1,j}^{n+1} - u_{i+1,j}^{n+1} + 2u_{i,j}^n - u_{i-1,j}^n - u_{i+1,j}^n) \\
&\quad - \frac{\Delta t}{2Re\Delta y^2} (2u_{i,j}^{n+1} - u_{i,j-1}^{n+1} - u_{i,j+1}^{n+1} + 2u_{i,j}^n - u_{i,j-1}^n - u_{i,j+1}^n)
\end{aligned} \tag{2.4}$$

où nous reprenons les mêmes notations pour les pas de discrétisation, a et b sont des réels fixés.

Par une décomposition de Fourier discrète pour une seule harmonique, nous obtenons :

$$u_{i,j}^n = \sum_{k_x, k_y} u^n e^{Ik_x i \Delta x} e^{Ik_y j \Delta y} = \sum_{k_x, k_y} u^n e^{Ii\phi_x} e^{Ij\phi_y}$$

avec $I^2 = -1$, u^n l'amplitude de l'harmonique et ${}^i(k_x, k_y)$ le vecteur nombre d'onde; ϕ_x, ϕ_y sont deux paramètres qui varieront de 0 à 2π .

Comme le schéma (2.4) est un schéma à deux pas, on introduit une nouvelle variable :

$$v_{i,j}^n = u_{i,j}^{n-1},$$

de manière à réécrire le système,

$$\left\{ \begin{array}{l}
u_{i,j}^{n+1} - u_{i,j}^n = -\alpha_x (3u_{i+1,j}^n - 3u_{i-1,j}^n - v_{i+1,j}^n + v_{i-1,j}^n) \\
\quad - \alpha_y (3u_{i,j+1}^n - 3u_{i,j-1}^n - v_{i,j+1}^n + v_{i,j-1}^n) \\
\quad - \beta_x (2u_{i,j}^{n+1} - u_{i-1,j}^{n+1} - u_{i+1,j}^{n+1} + 2u_{i,j}^n - u_{i-1,j}^n - u_{i+1,j}^n) \\
\quad - \beta_y (2u_{i,j}^{n+1} - u_{i,j-1}^{n+1} - u_{i,j+1}^{n+1} + 2u_{i,j}^n - u_{i,j-1}^n - u_{i,j+1}^n) \\
v_{i,j}^{n+1} = u_{i,j}^n
\end{array} \right. \tag{2.5}$$

avec les nouvelles notations,

$$\alpha_x = \frac{a \cdot \Delta t}{4\Delta x}, \quad \alpha_y = \frac{b \cdot \Delta t}{4\Delta y}, \quad \beta_x = \frac{\Delta t}{2Re\Delta x^2} \quad \text{et} \quad \beta_y = \frac{\Delta t}{2Re\Delta y^2}.$$

$$(2.5) \iff \begin{pmatrix} u_{i,j}^{n+1} \\ v_{i,j}^{n+1} \end{pmatrix} = M^{-1} \cdot N \begin{pmatrix} u_{i,j}^n \\ v_{i,j}^n \end{pmatrix}$$

où M et N peuvent s'écrire en fonction de ϕ_x et ϕ_y ,

$$M = \left[\begin{array}{c|c} 1 + 4(\beta_x \sin^2(\frac{\phi_x}{2}) + \beta_y \sin^2(\frac{\phi_y}{2})) & 0 \\ \hline 0 & 1 \end{array} \right],$$

$$N = \left[\begin{array}{c|c} 1 - 6I(\alpha_x \sin(\phi_x) + \alpha_y \sin(\phi_y)) & 6I(\alpha_x \sin(\phi_x) + \alpha_y \sin(\phi_y)) \\ \hline -4(\beta_x \sin^2(\frac{\phi_x}{2}) + \beta_y \sin^2(\frac{\phi_y}{2})) & 0 \end{array} \right],$$

ce qui nous donne la matrice d'amplification $G = M^{-1}.N$. Avec la complexité de cette matrice nous ne pouvons chercher à mettre en évidence un critère de stabilité indépendamment des valeurs de a et de b . Nous fixerons donc $a = b = Cte$ et chercherons à savoir si les paramètres que nous considérons pour la résolution de notre problème de Navier-Stokes vérifie bien la condition nécessaire de stabilité.

Par exemple, le test de la cavité carrée (RSLC), pour un nombre de Reynolds de 10000, une discrétisation de 257 points et $a = b = 1$, le fait de prendre un nombre de Courant ($\frac{a\Delta t}{\Delta x}$) de 0.5 nous restreint le pas de temps à 2.10^{-3} . En utilisant un logiciel de calcul formel, on obtient la vérification suivante, que chacune des deux valeurs propres de la matrice d'amplification reste bien dans la boule unité (cf. Fig. 3.5).

Il est clair que ce résultat n'est donné qu'à titre indicatif, d'une part notre problème (RSLC) n'est pas muni des mêmes conditions aux limites, et d'autre part, ni unidimensionnel ni linéaire. De plus l'analyse de Von-Neuman a aussi été remise en question par certains auteurs. Toutefois, nous pensons qu'il n'est pas inutile de vérifier l'adéquation du schéma numérique même sur un problème plus simple.

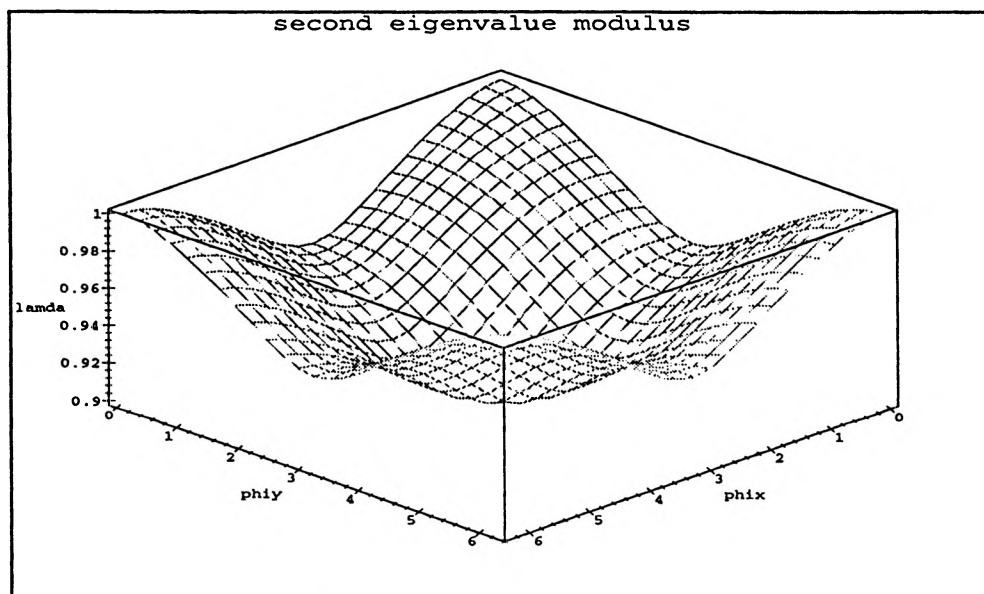
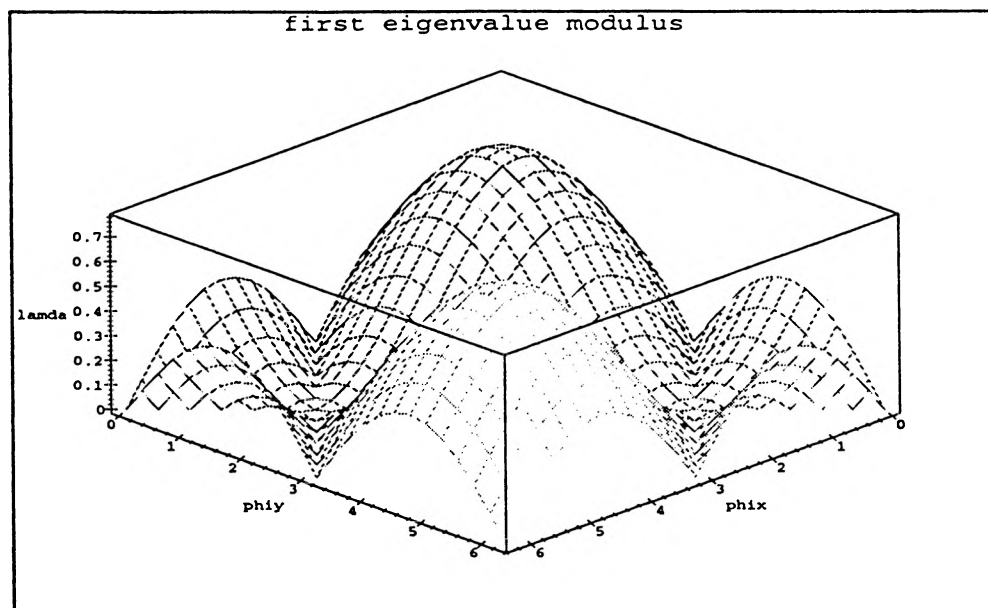


FIG. 2.5 - Vérification de la stabilité linéaire

2.2 Algorithme de projection et préconditionnements

Pour résoudre notre système discret nous utilisons une méthode de projection orthogonale due à Bramley [BRA92]. Elle a l'avantage par rapport aux méthodes classiques de projection de ne pas nécessiter l'imposition explicite de conditions aux limites pour l'équation de Poisson de la pression.

Pour décrire cet algorithme, nous donnons une réécriture plus appropriée du système discret de l'équation de Navier-Stokes (2.1),

$$\left[\begin{array}{c|c} A & B \\ \hline {}^tB & 0 \end{array} \right] \begin{pmatrix} \mathbf{u} \\ p \end{pmatrix} = \begin{pmatrix} \mathbf{f} \\ 0 \end{pmatrix} \quad (2.6)$$

où $\mathbf{u}$ et p sont les inconnues au temps $n + 1$ et $\mathbf{f}$ regroupe les termes explicites ainsi que le second membre. A est la matrice symétrique d'un opérateur défini positif sur le noyau de tB , car $\text{Ker}({}^tB)$ est sous-espace fermé pour la norme H_0^1 induite. B est la matrice rectangulaire de l'opérateur *gradient*, donc elle n'est pas de rang total.

En fait, l'algorithme est une méthode de Gradient Conjugué appliquée à la restriction de A à $\text{Ker}({}^tB)$, définie à l'aide d'un projecteur orthogonal P sur $\text{Ker}({}^tB)$.

Il suffit de résoudre la première équation de (2.6) pour p au sens des moindres carrés, ce qui donne : $p = B^\dagger(\mathbf{f} - A\mathbf{u})$, où $B^\dagger = ({}^tBB)^{-1}{}^tB$ est le pseudoinverse de Moore-Penrose de B [BIG74]. Après substitution dans la même équation, on obtient :

$$PA\mathbf{u} = P\mathbf{f} \quad (2.7)$$

où $P = I - BB^\dagger$ est le projecteur orthogonal sur $\text{Ker}({}^tB)$. De plus,

$${}^tB\mathbf{u} = 0 \Leftrightarrow \mathbf{u} = P\mathbf{u}, \quad (2.8)$$

ce qui peut être pris en compte dans (2.7) pour donner le système suivant :

$$\begin{cases} PAP\mathbf{u} = P\mathbf{f} \\ p = B^\dagger(\mathbf{f} - AP\mathbf{u}) + v \quad \forall v \in \text{Ker}({}^tB) \end{cases} \quad (2.9)$$

La réciproque est donnée dans [BRA92] et fait de $(P\mathbf{u}, p)$ la solution de (2.6) si et seulement si $(\mathbf{u}, p)$ vérifie (2.9).

L'utilisation d'une méthode de gradient conjugué (CG) pour résoudre le système projeté va consister à construire un espace de Krylov K_i qui va contenir la solution. De manière analogue au chapitre 2, (pour une seule suite de vecteurs dans CG),

$$K_i = \text{Vect}(\mathbf{d}_0, (PAP)\mathbf{d}_0, \dots, (PAP)^{i-1}\mathbf{d}_0) \subset \text{Ker}({}^tB),$$

or, d'après (2.8) il suffit de choisir $\mathbf{d}_0 \in \text{Ker}({}^tB)$, pour avoir

$$\tilde{K}_i = \text{Vect}(\mathbf{d}_0, (PA)\mathbf{d}_0, \dots, (PA)^{i-1}\mathbf{d}_0) = K_i \subset \text{Ker}({}^tB),$$

ce qui nous permettra de réduire la complexité de l'algorithme que nous utiliserons.

Une astuce permet aussi de retrouver la pression sans avoir à la calculer par la deuxième équation de (2.9). En l'initialisant à $p_0 = B^\dagger(f - Au_0)$, il suffit d'itérer comme suit :

$$p_{i+1} = B^\dagger(f - A(u_i + \alpha_i d_i)) = p_i - \alpha_i B^\dagger A d_i,$$

avec p_i la $i^{\text{ème}}$ itérée de la pression, d_i toujours un vecteur de descente et α_i la $i^{\text{ème}}$ coordonnée de la solution dans la base de K_{i+1} , et remarquer que la correction s'obtient par une étape intermédiaire du calcul de u_{i+1} .

Le point coûteux de cet algorithme est l'étape de projection, plus précisément, l'inversion de la matrice $({}^t B.B)$.

Bramley préconise d'en faire une factorisation complète une fois pour toute. Mais, si l'on envisage de faire des résolutions avec des discrétisations très fines, le stockage des deux matrices triangulaires requises risque d'encombrer, voire de dépasser la taille de la mémoire des super-calculateurs à notre disposition. Il suffit de considérer une discrétisation fine classique de 257 points dans chaque direction, et nous nous rendons compte qu'il faudrait beaucoup plus que la mémoire qui est à notre disposition. Toutefois, il nous est encore possible d'améliorer le coût de cette inversion en utilisant un préconditionnement adapté. Labadie propose d'utiliser comme matrice de préconditionnement un laplacien discret [CC87]. Ce qui a été implémenté avec succès lors du développement d'un code de calcul par F. Pascal résolvant les mêmes équations par un algorithme d'Uzawa [PAS92]. Comme le maillage de la pression est de pas $\Delta x, \Delta y$, nous choisirons l'inverse de la matrice discrète du laplacien de pas $\Delta x, \Delta y$, comme préconditionneur de l'inversion de ${}^t B.B$.

Tout comme le préconditionnement réalisé au chapitre précédent, il est évident que l'inversion de l'opérateur laplacien discret doit être rapide. D'après nos conclusions du premier chapitre, nous pouvons penser utiliser une inversion de type multi-grilles classique, ou encore une inversion par les Inconnues Incrémentales. Nous choisissons la première solution pour sa plus grande efficacité. Nous verrons dans la section 3 qu'il aurait aussi été possible de mener une résolution par l'algorithme de projection dans la base des Inconnues Incrémentales.

2.3 Résultats numériques

Comme l'initialisation de notre méthode de résolution nous demande un champ de vitesse solénoïdal, nous proposons dans le premier paragraphe, les résultats du problème de Stokes (1.4) avec les mêmes conditions aux limites que celles des problèmes modèles *RSLC* et *USLC*. Ensuite, afin d'apprécier l'efficacité de notre préconditionnement nous présentons les courbes de décroissance du résidu de l'étape de projection pour le problème de Stokes généralisé. Le deuxième paragraphe donne les résultats du problème de Navier-Stokes pour différents cas tests. La validation de notre code de calcul se fera par vérification des caractéristiques de la solution avec les résultats connus pour chacun des tests pour les problèmes modèles.

2.3.1 Problèmes de Stokes et de Stokes généralisé

Pour résoudre le problème de Stokes, nous utilisons le même algorithme de projection, initialisé par une vitesse nulle. Pour un nombre de Reynolds de 5000 et une discrétisation de 129^2 points, nous représentons les quantités physiques caractéristiques habituelles solutions du problème *USLC* (FIG. 2.6) (*RSLC* étant du même type) : les isobares, les lignes de courant et d'énergie, l'isovorticité ainsi que le champ de vecteurs de la vitesse. Nous pouvons nous rendre compte de la parfaite symétrie par rapport à l'axe $x_1 = 0.5$ de chacune des quantités.

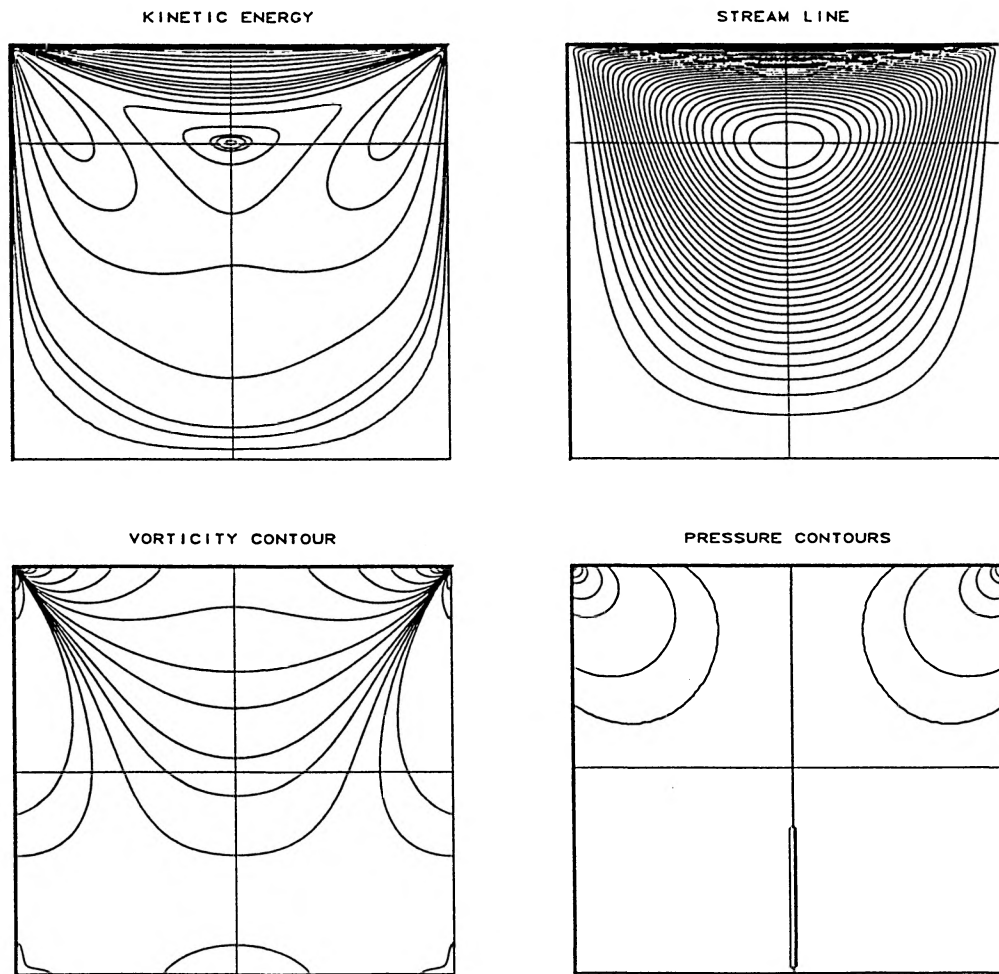


FIG. 2.6 - *Energie cinétique, lignes de courant, isovorticité et isobares pour USLC à $Re = 5000$ et 129^2 points*

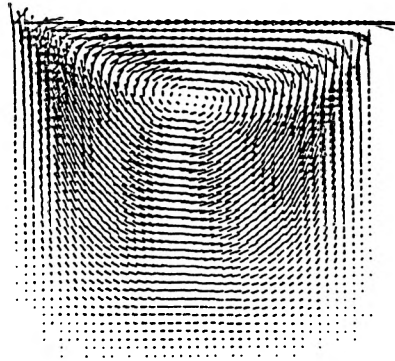


FIG. 2.7 - Champ de vitesse pour USLC à $Re = 5000$ et 129^2 points

Comme le problème de Stokes généralisé intervient à chaque itération, il est très intéressant de choisir une méthode efficace pour sa résolution. Notre implémentation de l'algorithme de projection pour la formulation vitesse-pression ne prétend pas permettre la résolution la plus rapide possible, mais il nous sert à faire des simulations à grand nombre de Reynolds en un temps raisonnable.

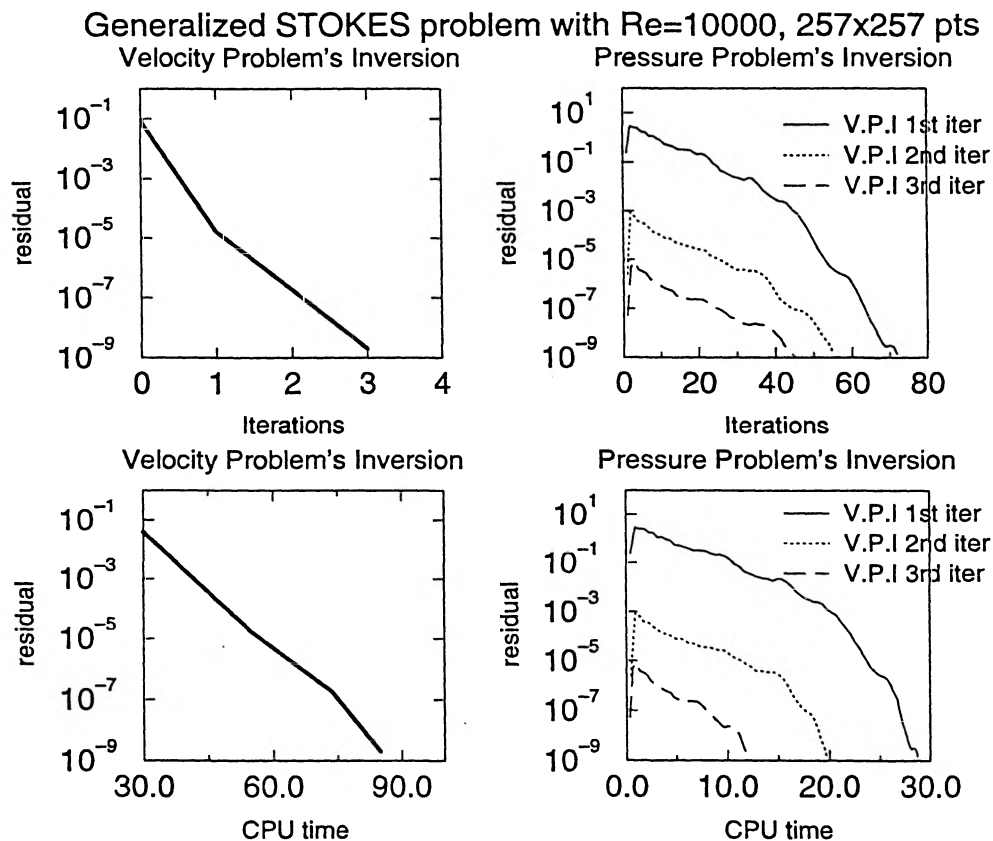


FIG. 2.8 - Résolution avec préconditionnement

Nous justifions numériquement que nous avons choisi un bon préconditionnement pour l'inversion du système faisant intervenir la pression en fournissant les courbes de décroissance du résidu. Nous montrons que le système projeté est un système bien conditionné et qu'il s'inverse en peu d'itérations.

Bien que la plupart des calculs aient été effectués sur CRAY C98 d'IDRIS¹, nous présentons à titre indicatif, aux figures 2.8 et 2.9 la décroissance du résidu de chacune des méthodes de gradient conjugué. Les temps C.P.U sont ceux mis pour faire ces calculs sur le CRAY YMP/EL du C.R.I.² d'Orsay. Nous obtenons les mêmes gains sur le CRAY C98, soit un facteur d'accélération de convergence d'un peu plus de 2 pour les premières itérations et d'un peu plus de 3 pour les dernières itérations.

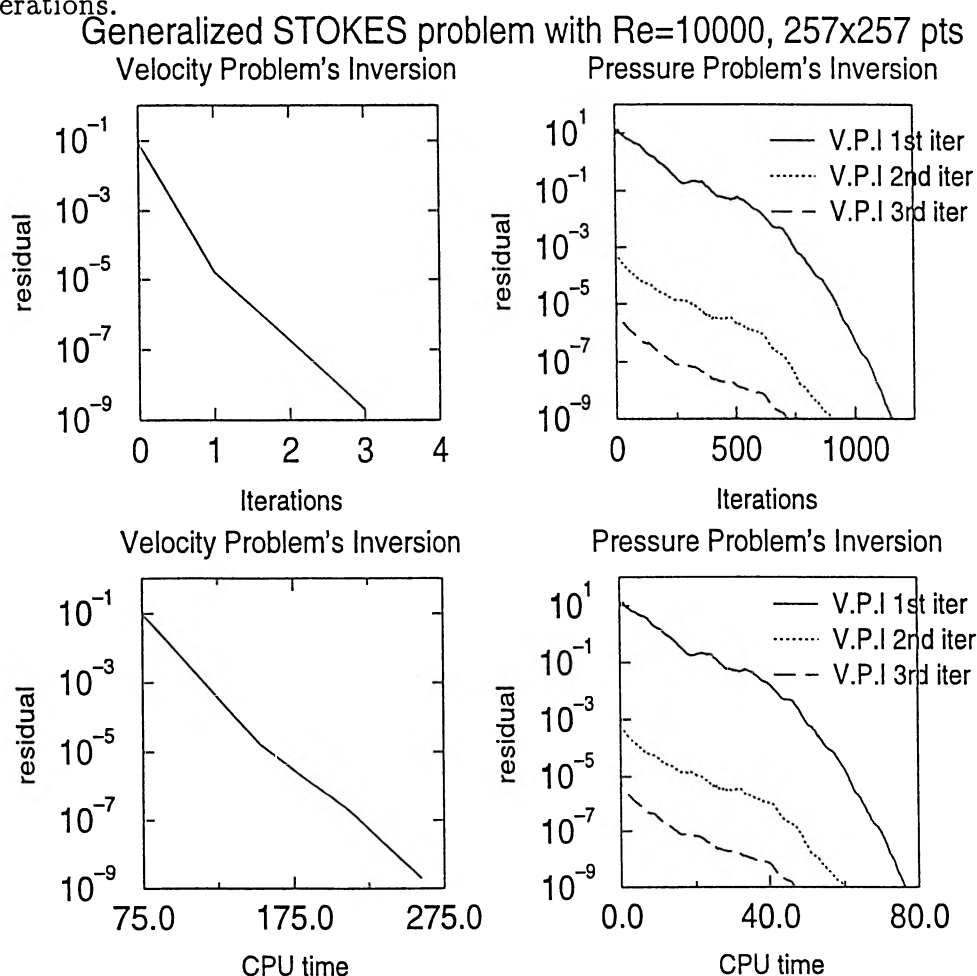


FIG. 2.9 - Résolution sans préconditionnement

1. Institut de Développement et des Ressources en Informatique Scientifique

2. Centre de Ressources en Informatique

2.3.2 Problème de Navier-Stokes

Tout d'abord il faut rappeler les caractéristiques physiques des simulations d'écoulements visqueux incompressibles. Le cas test de la cavité entraînée fait partie des plus classiques et nous savons que selon les valeurs du nombre de Reynolds, nous pouvons obtenir des comportements asymptotiques différents des solutions.

Les auteurs ne sont pas toujours d'accord sur les nombres limites permettant ces changements de comportement, mais tous sont unanimes sur le fait que la solution est stationnaire pour un nombre de Reynolds inférieur à 7500. De ce fait, certains auteurs choisissent de rechercher cette solution en résolvant le problème stationnaire (cf. [GGS82], [SHE87], [BJ90]). Cette stratégie, d'un avantage économique certain, a l'inconvénient d'entraîner des convergences d'algorithmes difficiles pour des nombres de Reynolds de l'ordre de 10^4 . Il n'y a malheureusement pas d'alternatives pour justifier ces calculs : il faut résoudre le problème évolutif !

Une simulation de l'écoulement va permettre à la solution itérative de transiter par plusieurs phases. Quand nous atteindrons la convergence nous dirons que nous avons un régime établi, qui s'identifie à mesure de l'augmentation du Reynolds par une solution stationnaire, puis périodique, puis quasi-périodique pour finalement atteindre la turbulence. Auparavant, après une phase très courte que nous pouvons qualifier d'amorçage, nous aurons le très long régime transitoire qui ne cessera de s'étendre plus on se rapprochera des grandes valeurs du Reynolds. Il faut bien se rendre compte que ces calculs sont très coûteux et que l'on ne peut que rarement dépasser le premier nombre de Reynolds critique. Toutefois ce challenge est surmonté par certains, (cf. [SHE91] pour la cavité carrée et [GGH90] pour la cavité double) qui pensent avoir franchi la dite bifurcation de Hopf.

Afin de valider notre code de calcul, nous comparons avec d'autres auteurs les caractéristiques des solutions pour chacun des problèmes modèles à différents nombre de Reynolds : 100, 1000, 5000.

Nous représentons 20 lignes de courant significatives, aux intensités fixées, proposées dans [GOY94], dont nous rappelons les valeurs, 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 les courbes d'isovorticité, les lignes de niveau d'énergie, et les isobares, nous nous chargerons de tracer les plus significatives.

Remarques sur l'implémentation

Le code de calcul a été optimisé avec le souci d'avoir une vectorisation efficace. L'analyseur de performances du CRAY C98, HPM³ nous donne une vitesse d'exécution de notre code de calcul de 360Mflops⁴. La puissance théorique du calculateur étant de 1000Mfops, il est considéré comme faisant partie des programmes performants. Nous pensons qu'il peut encore être amélioré en affinant l'écriture de la procédure de relaxation (faite par Gauss-Seidel damier et cadencée à 450Mfops) et en utilisant un prolongement de Multi-grilles C.S. (260Mfops) plus performant que le "prolongement par neuf points".

Remarque sur le schéma numérique

Nous avons utilisé dans un premier temps, le schéma de Cranck-Nicholson sur la partie linéaire couplé avec celui d'Euler explicite sur le terme non linéaire. Bien qu'il soit d'ordre 1, il nous permet de bien capter la solution stationnaire, mais le temps auquel nous obtenons la convergence est beaucoup plus long qu'un schéma plus précis : Cranck-Nicholson couplé avec Adams-Bashforth. Ce qui nous amène à préciser que dans ce chapitre, le temps adimensionné de simulation est pris pour ce dernier schéma d'ordre 2.

Solutions pour Reynolds = 100

Les tableaux suivants (2.1 et 2.2) récapitulent les caractéristiques de la solution stationnaire trouvée pour les problèmes *RSLC* et *USLC* pour un nombre de Reynolds de 100. Nous pouvons remarquer que les résultats varient très peu d'un auteur à un autre, hormis la localisation des tourbillons que Shen a proposé dans [SHE87]. Toutefois nous précisons que les tests de convergence du schéma numérique peuvent faire varier légèrement l'intensité ainsi que la position des extrema de la fonction de courant (et par conséquent de la vorticité). A titre d'exemple, si l'on augmente la précision sur le test de convergence, on constate une augmentation sensible de la valeur de l'intensité de la fonction de courant du centre du tourbillon principal. Ce point, déjà remarqué par d'autres auteurs, explique les valeurs que nous obtenons aux tableaux (2.1 et 2.2), le test d'arrêt étant :

$$\frac{\|u^n - u^{n+1}\|_W}{\|u^n\|_W} \leq 8.10^{-4}$$

Il est utile de présenter la solution des équations de Navier-Stokes sous plusieurs aspects [GGH90]. Aussi, nous montrons à la figure 2.10 ses lignes de courant, son isovorticité, son énergie cinétique, ses isobares ainsi que le champ de vitesse.

3. Hardware Performance Monitor

4. Million floating operation per second

Regularized Squared Lid-driven Cavity problem for $Re = 100$				
Auteurs	Poulet	Pascal [PAS92]	Shen [SHE87]	Shen [SHE91]
Formulation	(u, p)	(u, p)	(u, p)	(u, p)
Problème	évolutif	évolutif	stationnaire	évolutif
Discrétisation	Diff. finies	Elts finis	Spectral-Cheb.	Spectral-Cheb.
Maillage	65×65	65×65	64×64	17×17
Δt	0.1	0.1		4.0
P.V.				
(x, y)	(0.606, 0.756)	(0.609, 0.750)	(0.63, 0.77)	(0.609, 0.750)
$\psi(x, y)$	-0.0844	-0.08374	-0.08366	-0.0836
$\omega(x, y)$	-2.958	-2.933		
B.L.V.				
(x, y)	(0.031, 0.039)	(0.031, 0.031)	(0.05, 0.05)	(0.031, 0.031)
$\psi(x, y)$	$1.46 \cdot 10^{-6}$	$4.29 \cdot 10^{-8}$	$1.27 \cdot 10^{-6}$	$1.4 \cdot 10^{-6}$
$\omega(x, y)$	$1.231 \cdot 10^{-2}$	$7.49 \cdot 10^{-3}$		
B.R.V.				
(x, y)	(0.953, 0.047)	(0.953, 0.0469)	(0.97, 0.06)	(0.953, 0.047)
$\psi(x, y)$	$5.10 \cdot 10^{-6}$	$4.56 \cdot 10^{-6}$	$4.9 \cdot 10^{-6}$	$4.67 \cdot 10^{-6}$
$\omega(x, y)$	$1.90 \cdot 10^{-2}$	$1.79 \cdot 10^{-2}$		

TAB. 2.1 - Résultats du problème RSLC pour un nombre de Reynolds de 100, P.V.: Primary Vortex, B.L.V.: Bottom Left Vortex, B.R.V.: Bottom Right Vortex

Unregularized Squared Lid-driven Cavity problem for $Re = 100$					
Auteurs	Poullet	Pascal [PAS92]	B.-J. [BJ90]	Goyon [GOY94]	Ghia [GGS82]
Formulation	$(\mathbf{u}, p)$	$(\mathbf{u}, p)$	$(\mathbf{u}, p)$	(ψ, ω)	(ψ, ω)
Problème	évolutif	évolutif	stationnaire	évolutif	stationnaire
Discrétisation	Diff. finies	Elt. finis	Diff. finies	Diff. finies	Diff. finies
Maillage	65×65	65×65	129×129	129×129	129×129
Δt	$8 \cdot 10^{-3}$	0.1		$5 \cdot 10^{-3}$	
P.V.					
(x, y)	(0.619, 0.746)	(0.609, 0.734)	(0.617, 0.734)	(0.617, 0.734)	(0.617, 0.734)
$\psi(x, y)$	-0.105	-0.104	-0.103	-0.103	-0.103
$\omega(x, y)$	-3.297	-3.113			
B.L.V.					
(x, y)	(0.032, 0.032)	(0.031, 0.031)	(0.031, 0.039)	(0.031, 0.039)	(0.031, 0.039)
$\psi(x, y)$	$2.4 \cdot 10^{-6}$	$1.2 \cdot 10^{-6}$	$1.6 \cdot 10^{-6}$	$1.2 \cdot 10^{-6}$	$1.8 \cdot 10^{-6}$
$\omega(x, y)$	$1.17 \cdot 10^{-2}$	$9.8 \cdot 10^{-3}$			
B.R.V.					
(x, y)	(0.937, 0.063)	(0.953, 0.062)	(0.945, 0.062)	(0.945, 0.062)	(0.945, 0.062)
$\psi(x, y)$	$1.29 \cdot 10^{-5}$	$1.09 \cdot 10^{-5}$	$1.25 \cdot 10^{-5}$	$1.25 \cdot 10^{-5}$	$1.25 \cdot 10^{-5}$
$\omega(x, y)$	$4.29 \cdot 10^{-2}$	$2.45 \cdot 10^{-2}$			

TAB. 2.2 - Résultats du problème USLC pour un nombre de Reynolds de 100, P.V.: Primary Vortex, B.L.V.: Bottom Left Vortex, B.R.V.: Bottom Right Vortex

Solutions pour Reynolds = 1000

Comme l'on peut le remarquer au tableau 2.3, plus le nombre de Reynolds augmente et plus la dynamique recherchée devient complexe. Un contre-tourbillon apparaît pour deux auteurs [GGS82] et [GOY94] ayant pris une discrétisation deux fois plus fine que la notre, tandis qu'il n'existe ni pour nous ni pour d'autres (ayant le même nombre de points que nous).

Tout de même, nous obtenons les résultats de localisation très proches de la moyenne des auteurs, et surtout les plus proches de ceux qui utilisent la discrétisation la plus fine.

Cela montre qu'avec l'augmentation du nombre de Reynolds, il est nécessaire de faire croître le nombre de degrés de liberté pour capter correctement la solution. On l'explique par la finesse de l'épaisseur de la couche limite, qui est de l'ordre de $O(Re^{\frac{1}{2}})$ [SHE91].

A la figure 2.11, nous traçons les différents graphiques habituels qui présentent la solution.

Unregularized Squared Lid-driven Cavity problem for $Re = 1000$					
Auteurs	Poulet	Pascal [PAS92]	B.-J. [BJ90]	Goyon [GOY94]	Ghia [GGS82]
Formulation	(u, p)	(u, p)	(u, p)	(ψ, ω)	(ψ, ω)
Problème	évolutif	évolutif	stationnaire	évolutif	stationnaire
Discrétisation	Diff. finies	Elts finis	Diff. finies	Diff. finies	Diff. finies
Maillage	65×65	65×65	129×129	129×129	129×129
Δt	$8 \cdot 10^{-3}$	0.1		$2 \cdot 10^{-2}$	
P.V.					
(x, y)	(0.539, 0.571)	(0.531, 0.562)	(0.531, 0.559)	(0.531, 0.562)	(0.531, 0.562)
$\psi(x, y)$	-0.112	-0.123	-0.116	-0.116	-0.118
$\omega(x, y)$	-1.99	-0.022			
B.L.V.					
(x, y)	(0.079, 0.079)	(0.078, 0.078)	(0.086, 0.082)	(0.086, 0.078)	(0.086, 0.078)
$\psi(x, y)$	$1.83 \cdot 10^{-4}$	$2.28 \cdot 10^{-4}$	$3.25 \cdot 10^{-4}$	$2.11 \cdot 10^{-4}$	$2.31 \cdot 10^{-4}$
$\omega(x, y)$	0.295	0.324			
B.R.V.					
(x, y)	(0.857, 0.111)	(0.859, 0.109)	(0.871, 0.109)	(0.867, 0.117)	(0.851, 0.109)
$\psi(x, y)$	$1.71 \cdot 10^{-3}$	$1.69 \cdot 10^{-3}$	$1.91 \cdot 10^{-3}$	$1.63 \cdot 10^{-3}$	$1.75 \cdot 10^{-3}$
$\omega(x, y)$		1.22			
S.B.R.V.					
(x, y)	-	-	-	(0.99, 0.008)	(0.99, 0.008)
$\psi(x, y)$	-	-	-	$-8.8 \cdot 10^{-8}$	$-9.3 \cdot 10^{-8}$

TAB. 2.3 - Résultats du problème USLC pour un nombre de Reynolds de 1000, P.V.: Primary Vortex, B.L.V.: Bottom Left Vortex, B.R.V.: Bottom Right Vortex, S.B.R.V.: Second Bottom Right Vortex

Solutions à Reynolds = 5000

Pour ce nombre de Reynolds, nous avons choisi une discrétisation très fine soit 257^2 points. Nous avons choisi un pas de temps de $2 \cdot 10^{-3}$ permettant de vérifier la condition de type CFL :

$$\|u\|_{L^\infty} \frac{\Delta t}{\Delta x} \leq \frac{1}{2}.$$

Nous considérons que la solution stationnaire est atteinte à $T = 59$, ce qui est en total accord avec d'autres auteurs (F. Pascal, qui initialise aussi sa méthode d'éléments finis avec la solution du problème de Stokes, la considère stationnaire à $T = 60$).

Nous obtenons une localisation de nos tourbillons conforme à celle trouvée dans la littérature (*cf.* Tab 2.4, ci-après), avec l'intensité de contre-tourbillons variant sensiblement. Toutefois, nous pouvons noter l'apparition d'un contre-contre-tourbillon négatif en bas à gauche, ce qui explique la différence sur l'intensité des structures cohérentes. Notre discrétisation étant plus fine que celles utilisées par les autres auteurs ayant fait ce test, tout porte à croire que nos résultats sont les plus proches de la "réalité".

Les figures 2.12 et 2.13 donnent les différentes caractéristiques de la solution.

En regardant le tracé des isobares, on peut s'apercevoir de la bonne approximation de la condition aux limites de type Neuman homogène : les tangentes des isobares, au bord du domaine, sont quasiment toutes parallèles à la normale.

Le tracé de la vorticit   ne pr  sente aucune oscillation, donc notre discr  tisation semble suffisamment fine pour permettre une r  solution correcte.

On s'aper  oit de la taille des tourbillons pr  sent dans l'  coulement en observant le trac   des lignes de courant et du champ de vecteurs de la vitesse.

Les couches limites sont bien situ  es autour du bord du domaine.

Temps calcul et Impl  mentation

Nous donnons    titre d'exemple, le temps calcul pour le probl  me RSLC    $Re = 5000$ (maillage de 257^2 points). La solution stationnaire a n  cessit   30000 pas de temps d'int  gration, avec un co  t de 5 s par pas de temps, ce qui nous donne environ 40h de calcul sur CRAY C90 d'IDRIS.

Ce code de calcul a   t   compl  tement vectoris   et optimis  . Toutefois, nous avons utilis   un adressage index   car nous avons utilis   les r  sultats num  riques que donne cet algorithme    chaque it  ration, pour faire un diagnostic utile pour la suite de cette   tude (*cf.* 3.2 Analyse "a posteriori").

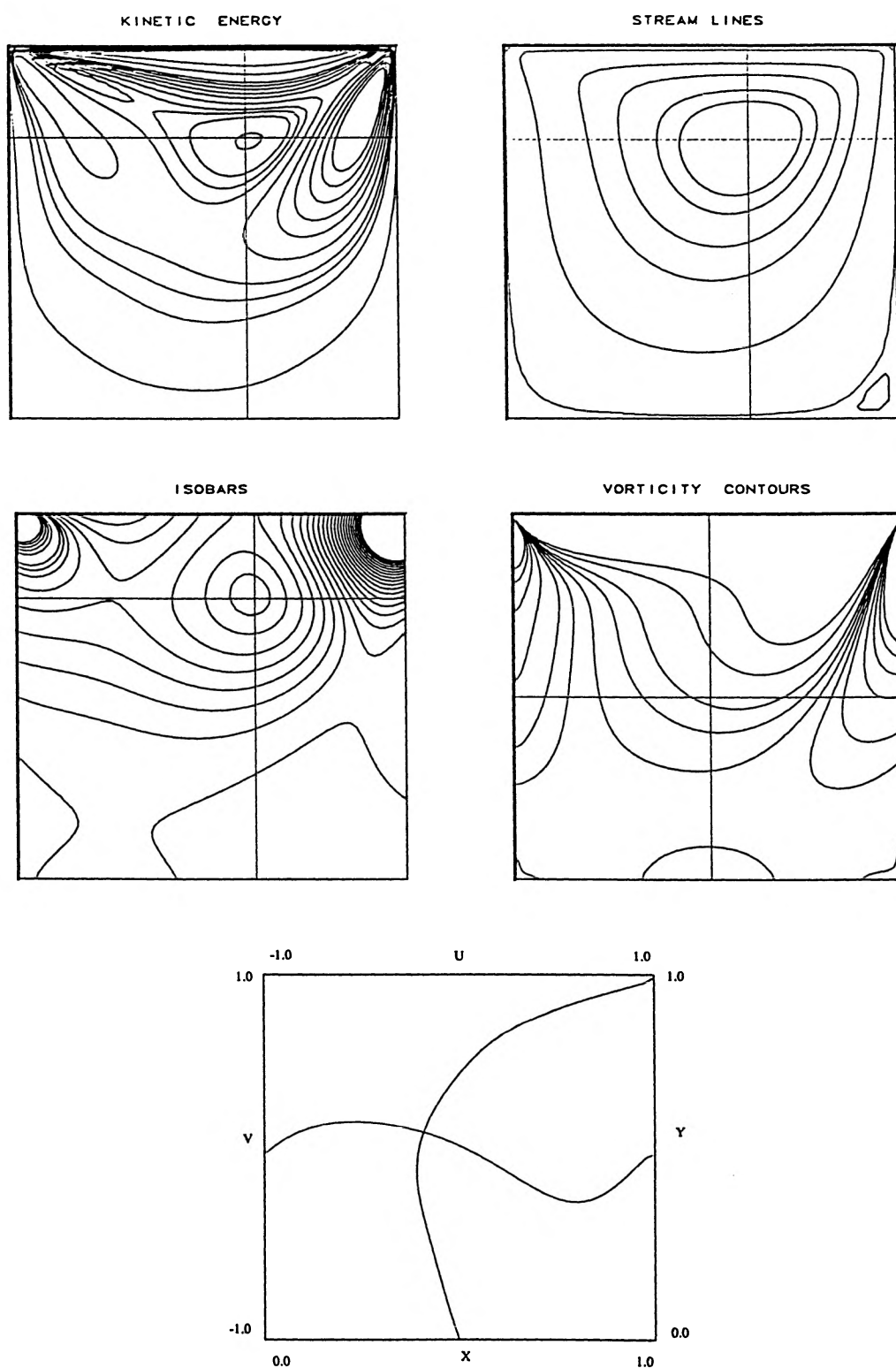


FIG. 2.10 - *Energie cinétique, lignes de courant, isovorticité, isobares et profils médians de la vitesse pour USLC à $Re = 100$ et 65^2 points*

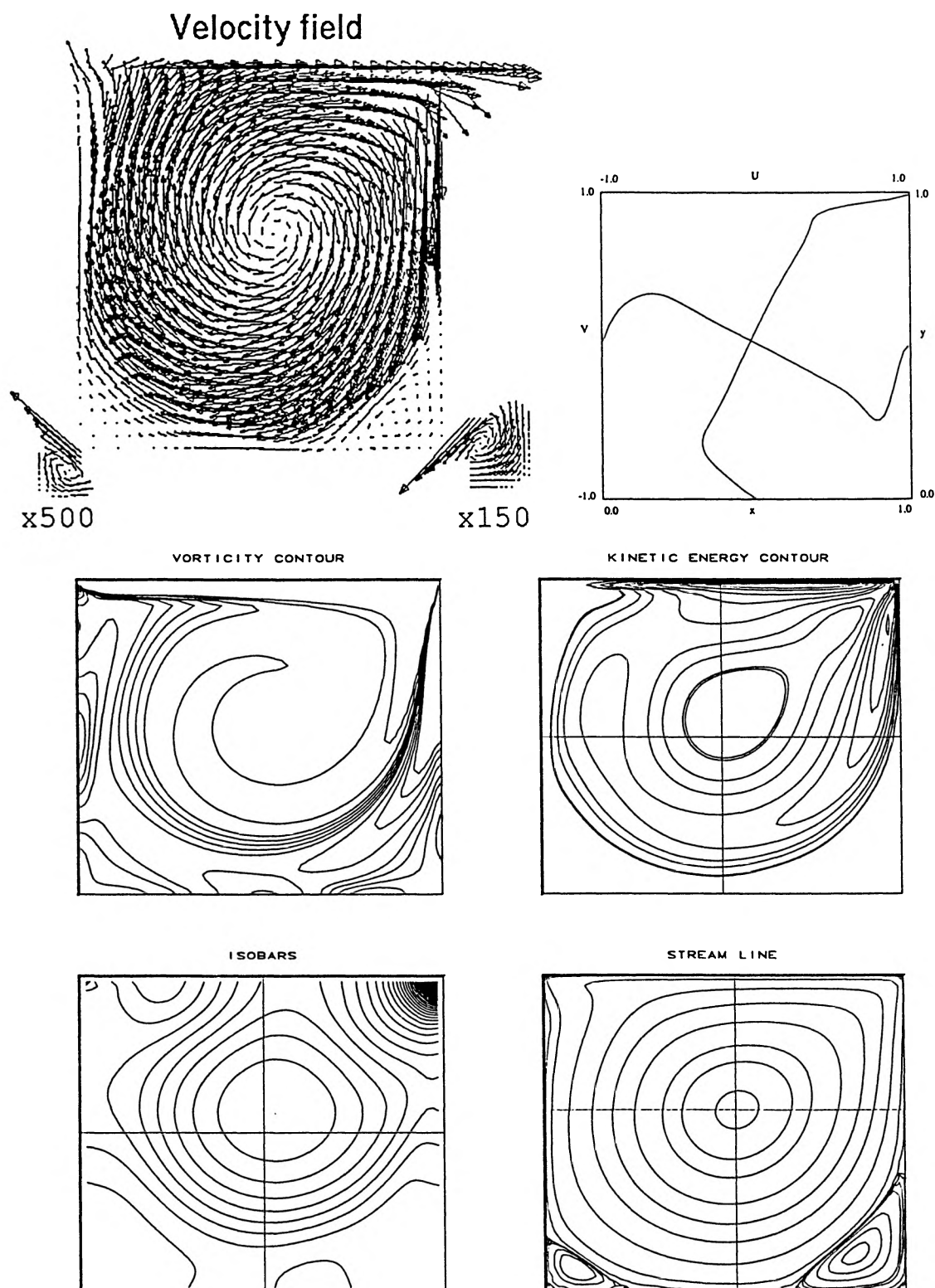


FIG. 2.11 - Energie cinétique, lignes de courant, isovorticité, isobares, champ de la vitesse et profils médians à partir du centre pour USLC à $Re = 1000$ et 257^2 points

Regularized Squared Lid-driven Cavity problem for $Re = 5000$			
Auteurs	Pouillet	Pascal [PAS92]	Shen [SHE91]
Formulation	(u, p)	(u, p)	(u, p)
Problème	évolutif	évolutif	évolutif
Discretisation	Différences finies	Eltéments finis	Spectral-Chebyshev
Maillage	257×257	129×129	33×33
Δt	$2 \cdot 10^{-3}$	0.05	0.03
P.V.			
(x, y)	(0.522, 0.569)	(0.539, 0.531)	(0.516, 0.531)
$\psi(x, y)$	$-9.5 \cdot 10^{-2}$	$-9.7 \cdot 10^{-2}$	$-8.8 \cdot 10^{-2}$
$\omega(x, y)$	-2.198	-2.169	
B.L.V.			
(x, y)	(0.071, 0.153)	(0.086, 0.117)	(0.094, 0.094)
$\psi(x, y)$	$1.37 \cdot 10^{-3}$	$6.72 \cdot 10^{-4}$	$7.53 \cdot 10^{-4}$
$\omega(x, y)$	1.461	0.731	
B.R.V.			
(x, y)	(0.816, 0.082)	(0.805, 0.078)	(0.922, 0.94)
$\psi(x, y)$	$2.04 \cdot 10^{-3}$	$2.42 \cdot 10^{-3}$	$7.75 \cdot 10^{-4}$
$\omega(x, y)$	1.648	2.009	
U.L.V.			
(x, y)	(0.082, 0.910)	(0.078, 0.906)	(0.078, 0.92)
$\psi(x, y)$	$5.74 \cdot 10^{-4}$	$7.86 \cdot 10^{-4}$	$6.78 \cdot 10^{-4}$
$\omega(x, y)$	0.965	1.159	
S.B.L.V.			
(x, y)	(0.016, 0.016)	-	-
$\psi(x, y)$	$-1.02 \cdot 10^{-6}$	-	-
$\omega(x, y)$	$-5.01 \cdot 10^{-2}$	-	-
S.B.R.V.			
(x, y)	(0.992, 0.012)	(0.984, 0.016)	-
$\psi(x, y)$	$-1.23 \cdot 10^{-7}$	$-2.39 \cdot 10^{-7}$	-
$\omega(x, y)$	$-8.26 \cdot 10^{-3}$	$-2.34 \cdot 10^{-2}$	

TAB. 2.4 - Résultats du problème RSLC pour un nombre de Reynolds de 5000, P.V.: Primary Vortex, B.L.V.: Bottom Left Vortex, B.R.V.: Bottom Right Vortex, U.L.V.: Upper Left Vortex, S.B.R.V.: Second Bottom Right Vortex, S.B.L.V.: Second Bottom Left Vortex

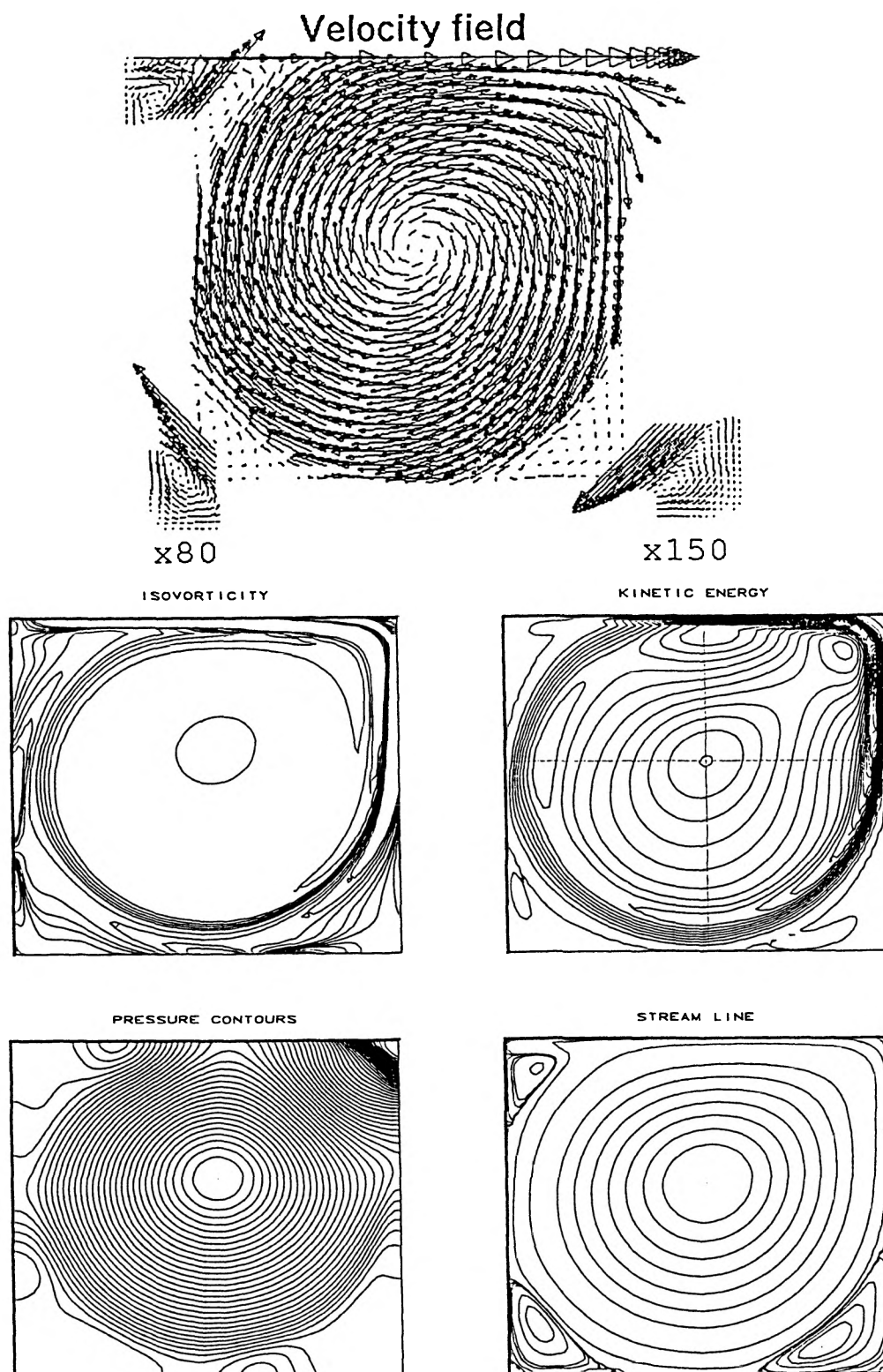


FIG. 2.12 - *Energie cinétique, lignes de courant, isovorticité, isobares et champ de vecteurs de la vitesse pour RSLC à $Re = 5000$ et 257^2 points*

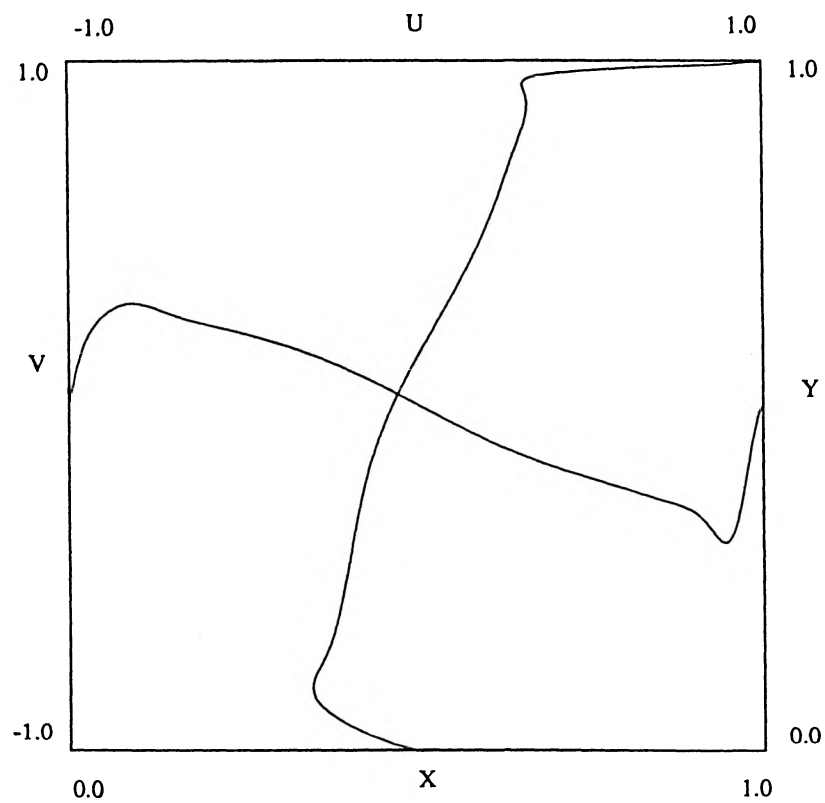


FIG. 2.13 - *Profils médians à partir du centre pour RSLC, $Re = 5000$*

Chapitre 3

Développement d'une méthode multi-niveaux

Dans le cas spectral, il est naturel de hiérarchiser ces structures par une ou plusieurs troncatures dans le spectre de l'inconnue cherchée. Nous voyons dans le paragraphe 3.1 qu'il n'est pas aussi aisé de l'envisager par les Inconnues Incrémentales classiques pour un maillage décalé, et pour cette raison nous définissons de nouvelles I.U. : les Inconnues Incrémentales décalées ou *Staggered Incremental Unknowns* (S.I.U.).

En utilisant cette nouvelle hiérarchisation, et les résultats provenant de la méthode de différences finies classique développée au chapitre précédent, nous présentons au paragraphe 3.2 une estimation "a posteriori" de la solution numérique.

Au paragraphe 3.3, nous utilisons cette hiérarchisation avec une grille de S.I.U. et analysons les opérateurs discrets considérés dans chacune des bases. Ensuite, nous donnons deux procédés pour restreindre le problème de Stokes généralisé survenant à chaque itération de Navier-Stokes à un problème grossier. L'un provient d'une restriction classique, tandis que l'autre, de la restriction sur la forme variationnelle. Une discussion nous permet d'effectuer le bon choix du problème grossier afin qu'il soit possible de le résoudre par l'algorithme de projection orthogonale présenté au chapitre précédent, via un changement de variables des inconnues.

A la section 3.4, nous proposons une gamme de schémas multi-niveaux non classiques, nous semblant adaptés à la dynamique non linéaire présente dans ce type de problème.

Finalement, nous concluons sur ces points et proposons le sens dans lequel notre méthode mérite d'être poursuivie.

3.1 Hiérarchisation des inconnues pour un maillage MAC à l'aide des *Staggered Incremental Unknowns* (S.I.U.)

Comme nous l'avons vu précédemment, la particularité des maillages décalés MAC est de définir la première composante de la vitesse au milieu des faces inférieure et supérieure d'une cellule; la deuxième composante au milieu des faces droite et gauche et la pression aux sommets de cette cellule (cf. figure 3.3). Et la définition classique des I.U. (cf. paragraphe 1.2.2) nécessite pour deux niveaux de grilles, l'utilisation de G_h (grille fine de pas h) et G_{2h} (grille grossière): les I.U. sont localisées aux points appartenant à $G_h \setminus (G_{2h} \cup \Gamma_h)$.

Il suffit de remarquer que $G_h \cap G_{2h} = \emptyset$, pour comprendre que cette construction ne peut être envisagée si l'on veut garder la même structure MAC sur la grille grossière et une structure multi-niveaux.

Pour y remédier, nous considérons G_g comme étant une grille grossière de pas $3h$ et G_h la grille fine de pas h . Cette décomposition triadique a déjà été proposée lors de la conception d'une méthode performante, de type multigrilles, par Liu *et al.* (cf. [LLM90]) en vue de résoudre le problème d'un écoulement dans un canal plan (*plane channel flow*) avec les conditions en amont et en aval de type Poiseuille.

Nous adopterons les notations, G_h (respectivement G_{3h}) sera la grille constituée de points intérieurs au domaine de pas h (resp. de pas $3h$). Les nouvelles I.U. que nous appellerons S.I.U. seront donc situées aux points du maillage appartenant à $G_h \setminus (G_{3h} \cup \partial\Omega_h)$.

Avant de donner leur description pour un niveau de grille de S.I.U., nous mentionnerons que nous les avons construites en utilisant la formule de Taylor, pour qu'elles soient d'ordre 2 comme les I.U. classiques. Par rapport aux remarques qui ont été faites dans un travail précédent (cf. [GAR92]) sur les I.U. dans un autre concept, nous pensons que la matrice de changement de base qui garantirait aux S.I.U. d'être d'un ordre plus élevé sera beaucoup trop pleine et le fait de l'appliquer pénaliserait beaucoup le coût de nombreuses résolutions.

Soient $\Omega = (a, b) \times (c, d)$ un domaine rectangle, $h_1 = 1/(3.N_1)$ et $h_2 = 1/(3.N_2)$ les pas de discrétisation dans chacune des directions ($N_1, N_2 \in \mathbb{N}$). Pour (u, p) solution discrète de (2.1), on redéfinit les vecteurs des S.I.U. $\bar{u}_h, \bar{p}_h$, les matrices de changement de bases respectives S et T de la vitesse et de la pression, de la base des S.I.U. à la base classique (nodale):

$$u = S\bar{u}, \quad p = T\bar{p}, \quad \text{où } \bar{u} = {}^t(y_{i,j}^u, z_{i,j}^u, y_{i,j}^v, z_{i,j}^v), \quad \bar{p} = {}^t(p_{i,j}^y, p_{i,j}^z), \quad \text{avec}$$

- $y_{i,j}^u$ et $y_{i,j}^v$ les valeurs associées aux points de la grille grossière de la vitesse,
- $z_{i,j}^u$ et $z_{i,j}^v$ les S.I.U. de chacune des composantes de la vitesse, et
- $p_{i,j}^y$ et $p_{i,j}^z$ les composantes grossière et incrémentale de la pression.

3.1.1 S.I.U. pour la vitesse

Commençons par donner les formules de changement de base pour chacune des composantes de la vitesse,

G_g se constitue des points auxquels sont affectés les valeurs suivantes :

$$y_{3i,3j+3/2}^u = u_{3i,3j+3/2} \quad \text{pour } i = 0, \dots, N_1; j = 0, \dots, N_2 - 1$$

$$y_{3i+3/2,3j}^v = v_{3i+3/2,3j} \quad \text{pour } i = 0, \dots, N_1 - 1; j = 0, \dots, N_2$$

complétée des points de $\partial\Omega$:

$$y_{3i,0}^u = u_{3i,0} \text{ et } y_{3i,3N_2}^u = u_{3i,3N_2} \quad \forall i = 0, \dots, N_1$$

$$y_{0,3j}^v = v_{0,3j} \text{ et } y_{3N_1,3j}^v = v_{3N_1,3j} \quad \forall j = 0, \dots, N_2$$

Nous pouvons toujours regrouper nos S.I.U. en trois types comme au chapitre 1 (cf. FIG. 3.2).

Type (FV) deux points de G_h situés verticalement entre deux pts de G_g ainsi que celui qui se trouve verticalement entre un point interne de G_g et celui du bord qui l'a complétée :

$$z_{3i,3j+3/2+1}^u = u_{3i,3j+3/2+1} - \frac{1}{3}(2u_{3i,3j+3/2} + u_{3i,3(j+1)+3/2})$$

$$z_{3i,3j+3/2+2}^u = u_{3i,3j+3/2+2} - \frac{1}{3}(u_{3i,3j+3/2} + 2u_{3i,3(j+1)+3/2})$$

$$\forall i = 1, \dots, N_1 - 1; j = 0, \dots, N_2 - 2$$

$$z_{3i,1/2}^u = u_{3i,1/2} - \frac{1}{3}(u_{3i,3/2} + 2u_{3i,0})$$

$$z_{3i,3N_2-1/2}^u = u_{3i,3N_2-1/2} - \frac{1}{3}(u_{3i,3N_2-3/2} + 2u_{3i,3N_2})$$

$$\forall i = 1, \dots, N_1 - 1$$

$$z_{3i+3/2,3j+1}^v = v_{3i+3/2,3j+1} - \frac{1}{3}(2v_{3i+3/2,3j} + v_{3i+3/2,3j+3})$$

$$z_{3i+3/2,3j+2}^v = v_{3i+3/2,3j+2} - \frac{1}{3}(v_{3i+3/2,3j} + 2v_{3i+3/2,3j+3})$$

$$\forall i = 0, \dots, N_1 - 1; j = 0, \dots, N_2 - 1$$

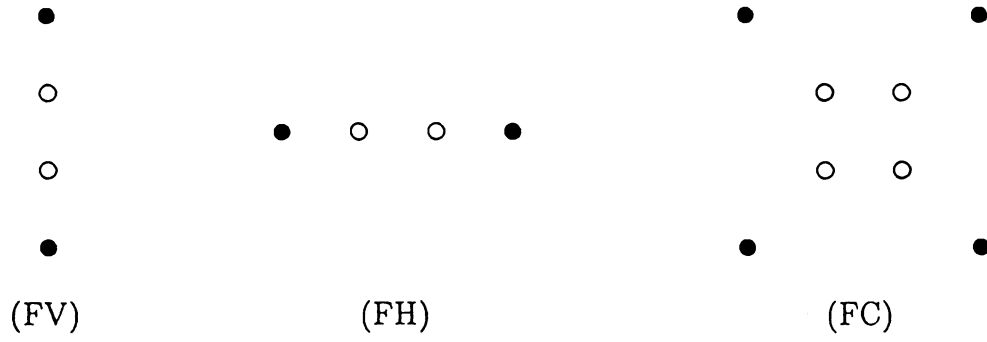


FIG. 3.2 - Familles de S.I.U. pour la vitesse en dim 2

Type (FH) deux points de G_h situés horizontalement entre deux pts de G_g ainsi que celui qui se trouve horizontalement entre un point interne de G_g et celui du bord qui l'a complétée :

$$z_{3i+1,3j+3/2}^u = u_{3i+1,3j+3/2} - \frac{1}{3}(2u_{3i,3j+3/2} + u_{3i+3,3j+3/2})$$

$$z_{3i+2,3j+3/2}^u = u_{3i+2,3j+3/2} - \frac{1}{3}(u_{3i,3j+3/2} + 2u_{3i+3,3j+3/2})$$

$$\forall i = 0, \dots, N_1 - 1; j = 0, \dots, N_2 - 1$$

$$z_{3i+3/2+1,3j}^v = v_{3i+3/2+1,3j} - \frac{1}{3}(2v_{3i+3/2,3j} + v_{3(i+1)+3/2,3j})$$

$$z_{3i+3/2+2,3j}^v = v_{3i+3/2+2,3j} - \frac{1}{3}(v_{3i+3/2,3j} + 2v_{3(i+1)+3/2,3j})$$

$$\forall i = 0, \dots, N_1 - 2; j = 1, \dots, N_2 - 1$$

$$z_{1/2,3j}^v = v_{1/2,3j} - \frac{1}{3}(v_{3/2,3j} + 2v_{0,3j})$$

$$z_{3N_1-1/2,3j}^v = v_{3N_1-1/2,3j} - \frac{1}{3}(v_{3N_1-3/2,3j} + 2v_{3N_1,3j})$$

$$\forall j = 1, \dots, N_2 - 1$$

Type (FC) quatre points de G_h situés au milieu de quatre pts de G_g ainsi que les deux autres points proches du bord :

$$z_{3i+1,3j+3/2+1}^u = u_{3i+1,3j+3/2+1} - \frac{1}{12} [5u_{3i,3j+3/2} + 3u_{3i+3,3j+3/2} + 3u_{3i,3(j+1)+3/2} + u_{3i+3,3(j+1)+3/2}]$$

$$z_{3i+2,3j+3/2+1}^u = u_{3i+2,3j+3/2+1} - \frac{1}{12} [5u_{3i+3,3j+3/2} + 3u_{3i,3j+3/2} + 3u_{3i+3,3(j+1)+3/2} + u_{3i,3(j+1)+3/2}]$$

$$z_{3i+1,3j+3/2+2}^u = u_{3i+1,3j+3/2+2} - \frac{1}{12} [5u_{3i,3(j+1)+3/2} + 3u_{3i,3j+3/2} + 3u_{3i+3,3(j+1)+3/2} + u_{3i+3,3j+3/2}]$$

$$z_{3i+2,3j+3/2+2}^u = u_{3i+2,3j+3/2+2} - \frac{1}{12} [5u_{3i+3,3(j+1)+3/2} + 3u_{3i+3,3j+3/2} + 3u_{3i,3(j+1)+3/2} + u_{3i,3j+3/2}]$$

$$\forall i = 0, \dots, N_1 - 1; \quad j = 1, \dots, N_2 - 1$$

$$z_{3i+1,1/2}^u = u_{3i+1,1/2} - \frac{1}{6} (6u_{3i+1,0} - u_{3i,3/2} + u_{3i+3,3/2})$$

$$z_{3i+2,1/2}^u = u_{3i+2,1/2} - \frac{1}{6} (6u_{3i+2,0} - u_{3i+3,3/2} + u_{3i,3/2})$$

$$z_{3i+1,3N_2-1/2}^u = u_{3i+1,3N_2-1/2} - \frac{1}{6} (6u_{3i+1,3N_2} - u_{3i,3N_2-3/2} + u_{3i+3,3N_2-3/2})$$

$$z_{3i+2,3N_2-1/2}^u = u_{3i+2,3N_2-1/2} - \frac{1}{6} (6u_{3i+2,3N_2} - u_{3i+3,3N_2-3/2} + u_{3i,3N_2-3/2})$$

$$\forall i = 0, \dots, N_1 - 1$$

Nous avons dû compléter chaque type de S.I.U. par un certain nombre de valeurs aux nœuds situés entre les points qui ont complété la grille grossière et ceux immédiatement les plus proches. Ces S.I.U. particulières n'ont pas d'équivalents dans le maillage classique considéré au chapitre 1, mais ont été obtenues encore par interpolation d'ordre 2, en faisant intervenir les valeurs aux bords du domaine - ce qui ne présente pas de difficultés grâce aux conditions aux limites de type Dirichlet -.

$$z_{3i+3/2+1,3j+1}^v = v_{3i+3/2+1,3j+1} - \frac{1}{12} [5v_{3i+3/2,3j} + 3v_{3(i+1)+3/2,3j} + 3v_{3i+3/2,3j+3} + v_{3(i+1)+3/2,3j+3}]$$

$$z_{3i+3/2+2,3j+1}^v = v_{3i+3/2+2,3j+1} - \frac{1}{12} [5v_{3(i+1)+3/2,3j} + 3v_{3i+3/2,3j} + 3v_{3(i+1)+3/2,3j+3} + v_{3i+3/2,3j+3}]$$

$$z_{3i+3/2+1,3j+2}^v = v_{3i+3/2+1,3j+2} - \frac{1}{12} [5v_{3i+3/2,3j+3} + 3v_{3(i+1)+3/2,3j+3} + 3v_{3i+3/2,3j} + v_{3(i+1)+3/2,3j}]$$

$$z_{3i+3/2+2,3j+2}^v = v_{3i+3/2+2,3j+2} - \frac{1}{12} [5v_{3(i+1)+3/2,3j+3} + 3v_{3(i+1)+3/2,3j} + 3v_{3i+3/2,3j+3} + v_{3i+3/2,3j}]$$

$$\forall i = 1, \dots, N_1 - 1; \quad j = 0, \dots, N_2 - 1$$

$$z_{1/2,3j+1}^v = v_{1/2,3j+1} - \frac{1}{6} (6u_{0,3j+1} - v_{3/2,3j} + v_{3/2,3j+3})$$

$$z_{1/2,3j+2}^v = v_{1/2,3j+2} - \frac{1}{6} (6v_{0,3j+2} - v_{3/2,3j+3} + v_{3/2,3j})$$

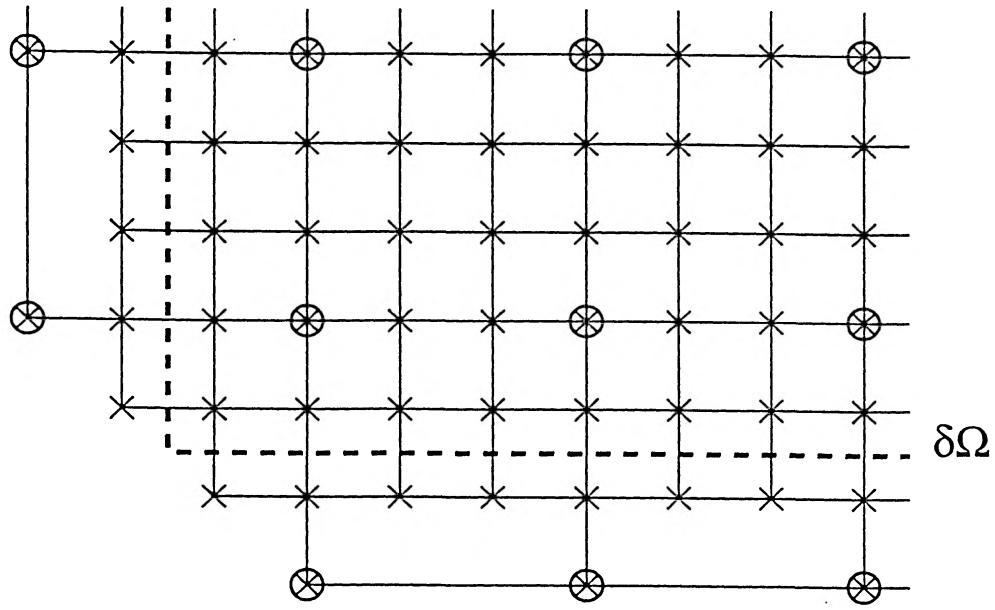
$$z_{3N_1-1/2,3j+1}^v = v_{3N_1-1/2,3j+1} - \frac{1}{6} (6u_{3N_1,3j+1} - v_{3N_1-3/2,3j} + v_{3N_1-3/2,3j+3})$$

$$z_{3N_1-1/2,3j+2}^v = v_{3N_1-1/2,3j+2} - \frac{1}{6} (6u_{3N_1,3j+2} - v_{3N_1-3/2,3j+3} + v_{3N_1-3/2,3j})$$

$$\forall j = 0, \dots, N_2 - 1$$

3.1.2 S.I.U. pour la pression

La grille grossière de la pression étant aussi une grille d'inconnues de pas $3h$ et la grille fine constituée de valeurs aux nœuds de $G_h \setminus G_{3h}$, la configuration des S.I.U. d'ordre 2 pour la pression paraît être la même que celle de la vitesse. Toutefois, une différence notable intervient sur les valeurs aux points proches du bord du domaine, car nous n'avons plus de condition aux limites de type Dirichlet pour la pression. Nous sommes donc obligés d'extrapoler les valeurs intérieures de la grille grossière au bord du domaine. Ces S.I.U. pour la pression peuvent être regroupées en trois types analogues à ceux de la vitesse, et un quatrième type traitant de manière spécifique les quatre valeurs aux nœuds des coins internes du maillage. Nous récapitulons les formules définissant les S.I.U. classiques ainsi que celles obtenues par extrapolation linéaire. La figure 3.3 donne la configuration d'un maillage proche du bord pour la pression.



⊗ pressure coarse grid component

× pressure fine grid component

FIG. 3.3 - S.I.U. pour la pression

Type (FV) : deux points de G_h situés verticalement entre deux pts de G_g ainsi que celui qui se trouve verticalement entre un point externe de G_g et le point de $\partial\Omega_h$:

$$p_{3i+3/2,3j+5/2}^z = p_{3i+3/2,3j+5/2} - \frac{1}{3}[2p_{3i+3/2,3j+3/2} + p_{3i+3/2,3(j+1)+3/2}]$$

$$p_{3i+3/2,3j+7/2}^z = p_{3i+3/2,3j+7/2} - \frac{1}{3}[2p_{3i+3/2,3(j+1)+3/2} + p_{3i+3/2,3j+3/2}]$$

$$\forall i = 0, \dots, N_1 - 1; j = 0, \dots, N_2 - 2$$

$$p_{3i+3/2,1/2}^z = p_{3i+3/2,1/2} - \frac{1}{3}[4p_{3i+3/2,3/2} - p_{3i+3/2,3+3/2}]$$

$$p_{3i+3/2,3N_2-1/2}^z = p_{3i+3/2,3N_2-1/2} - \frac{1}{3}[4p_{3i+3/2,3N_2-3/2} - p_{3i+3/2,3(N_2-1)-3/2}]$$

$$\forall i = 0, \dots, N_1 - 1$$

Type (FH) : deux points de G_h situés horizontalement entre deux pts de G_g ainsi que celui qui se trouve horizontalement entre un point externe de G_g et le point de $\partial\Omega_h$:

$$p_{3i+5/2,3j+3/2}^z = p_{3i+5/2,3j+3/2} - \frac{1}{3}[2p_{3i+3/2,3j+3/2} + p_{3(i+1)+3/2,3j+3/2}]$$

$$p_{3i+7/2,3j+3/2}^z = p_{3i+7/2,3j+3/2} - \frac{1}{3}[2p_{3(i+1)+3/2,3j+3/2} + p_{3i+3/2,3j+3/2}]$$

$$\forall i = 0, \dots, N_1 - 2; j = 0, \dots, N_2 - 1$$

$$p_{1/2,3j+3/2}^z = p_{1/2,3j+3/2} - \frac{1}{3}[4p_{3/2,3j+3/2} - p_{3+3/2,3j+3/2}]$$

$$p_{3N_1-1/2,3j+3/2}^z = p_{3N_1-1/2,3j+3/2} - \frac{1}{3}[4p_{3N_1-3/2,3j+3/2} - p_{3(N_1-1)-3/2,3j+3/2}]$$

$$\forall i = 0, \dots, N_2 - 1$$

Type (FB) : quatre coins de G_h :

$$p_{1/2,1/2}^z = p_{1/2,1/2} - \frac{1}{3}[4p_{3/2,3/2} - p_{3+3/2,3+3/2}]$$

$$p_{1/2,3N_2-1/2}^z = p_{1/2,3N_2-1/2} - \frac{1}{3}[4p_{3/2,3N_2-3/2} - p_{3+3/2,3(N_2-1)-3/2}]$$

$$p_{3N_1-1/2,1/2}^z = p_{3N_1-1/2,1/2} - \frac{1}{3}[4p_{3N_1-3/2,3/2} - p_{3(N_1-1)-3/2,3+3/2}]$$

$$p_{3N_1-1/2,3N_2-1/2}^z = p_{3N_1-1/2,3N_2-1/2} - \frac{1}{3}[4p_{3N_1-3/2,3N_2-3/2} - p_{3(N_1-1)-3/2,3(N_2-1)-3/2}]$$

Type (FC) : quatre points de G_h situés au milieu de quatre pts de G_g ainsi que les deux autres points proches du bord :

$$p_{3i+5/2,3j+5/2}^z = p_{3i+5/2,3j+5/2} - \frac{1}{12} [5p_{3i+3/2,3j+3/2} + 3p_{3(i+1)+3/2,3j+3/2} + 3p_{3i+3/2,3(j+1)+3/2} + p_{3(i+1)+3/2,3(j+1)+3/2}]$$

$$p_{3i+7/2,3j+5/2}^z = p_{3i+7/2,3j+5/2} - \frac{1}{12} [5p_{3(i+1)+3/2,3j+3/2} + 3p_{3(i+1)+3/2,3(j+1)+3/2} + 3p_{3i+3/2,3j+3/2} + p_{3i+3/2,3(j+1)+3/2}]$$

$$p_{3i+7/2,3j+7/2}^z = p_{3i+7/2,3j+7/2} - \frac{1}{12} [5p_{3(i+1)+3/2,3(j+1)+3/2} + 3p_{3(i+1)+3/2,3j+3/2} + 3p_{3i+3/2,3(j+1)+3/2} + p_{3i+3/2,3j+3/2}]$$

$$p_{3i+5/2,3j+7/2}^z = p_{3i+5/2,3j+7/2} - \frac{1}{12} [5p_{3i+3/2,3(j+1)+3/2} + 3p_{3i+3/2,3j+3/2} + 3p_{3(i+1)+3/2,3(j+1)+3/2} + p_{3(i+1)+3/2,3j+3/2}]$$

$\forall i = 0, \dots, N_1 - 2; j = 0, \dots, N_2 - 2$

$$p_{3i+5/2,1/2}^z = p_{3i+5/2,1/2} - \frac{1}{9} [8p_{3i+3/2,3/2} + 4p_{3i+3/2,3+3/2} - 2p_{3(i+1)+3/2,3/2} - p_{3(i+1)+3/2,3+3/2}]$$

$$p_{3i+7/2,1/2}^z = p_{3i+7/2,1/2} - \frac{1}{9} [8p_{3(i+1)+3/2,3/2} + 4p_{3(i+1)+3/2,3+3/2} - 2p_{3i+3/2,3/2} - p_{3i+3/2,3+3/2}]$$

$$p_{3i+5/2,3N_2-1/2}^z = p_{3i+5/2,3N_2-1/2} - \frac{1}{9} [8p_{3i+3/2,3N_2-3/2} + 4p_{3i+3/2,3(N_2-1)-3/2} - 2p_{3(i+1)+3/2,3N_2-3/2} - p_{3(i+1)+3/2,3(N_2-1)-3/2}]$$

$$p_{3i+7/2,3N_2-1/2}^z = p_{3i+7/2,3N_2-1/2} - \frac{1}{9} [8p_{3(i+1)+3/2,3N_2-3/2} + 4p_{3(i+1)+3/2,3(N_2-1)-3/2} - 2p_{3i+3/2,3N_2-3/2} - p_{3i+3/2,3(N_2-1)-3/2}]$$

$\forall i = 0, \dots, N_1 - 2$

$$p_{1/2,3j+5/2}^z = p_{1/2,3j+5/2} - \frac{1}{9} [8p_{3/2,3j+3/2} + 4p_{3+3/2,3j+3/2} - 2p_{3/2,3(j+1)+3/2} - p_{3+3/2,3(j+1)+3/2}]$$

$$p_{1/2,3j+7/2}^z = p_{1/2,3j+7/2} - \frac{1}{9} [8p_{3/2,3(j+1)+3/2} + 4p_{3+3/2,3(j+1)+3/2} - 2p_{3/2,3j+3/2} - p_{3+3/2,3j+3/2}]$$

$$p_{3N_1-1/2,3j+5/2}^z = p_{3N_1-1/2,3j+5/2} - \frac{1}{9} [8p_{3N_1-3/2,3j+3/2} + 4p_{3(N_1-1)-3/2,3j+3/2} - 2p_{3N_1-3/2,3(j+1)+3/2} - p_{3(N_1-1)-3/2,3(j+1)+3/2}]$$

$$p_{3N_1-1/2,3j+7/2}^z = p_{3N_1-1/2,3j+7/2} - \frac{1}{9} [8p_{3N_1-3/2,3(j+1)+3/2} + 4p_{3(N_1-1)-3/2,3(j+1)+3/2} - 2p_{3N_1-3/2,3j+3/2} - p_{3(N_1-1)-3/2,3j+3/2}]$$

$\forall j = 0, \dots, N_2 - 2$

3.2 Analyse “a posteriori” de la solution

Dans cette section, nous présentons une analyse de la taille des différents termes provenant du système discret (2.1). Il ne s'agit pas exactement d'une estimation a posteriori de la solution pour la simple raison que nous n'avons pas atteint le régime établi, mais nous verrons que cette vision n'en est pas très éloignée.

Nous adopterons une écriture formelle de ce système directement hiérarchisé dans la base des S.I.U., qui sera détaillée à la section suivante. Notre outil d'analyse est le plus simple: la norme $L^2(\Omega)$ de chacune des images des opérateurs discrets. Nous l'appliquerons dans différentes simulations obtenues par la méthode de résolution développée dans la section précédente. Aux figures 3.4 et 3.5, nous proposons l'évolution de la norme de W_h de la vitesse sur chacune des grilles (grossière et fines). On vérifie sans peine que les S.I.U. sont d'ordre 2, en retrouvant le rapport entre la norme de la grille des S.I.U. de la vitesse et la norme de la grille grossière de la vitesse de l'ordre 9.

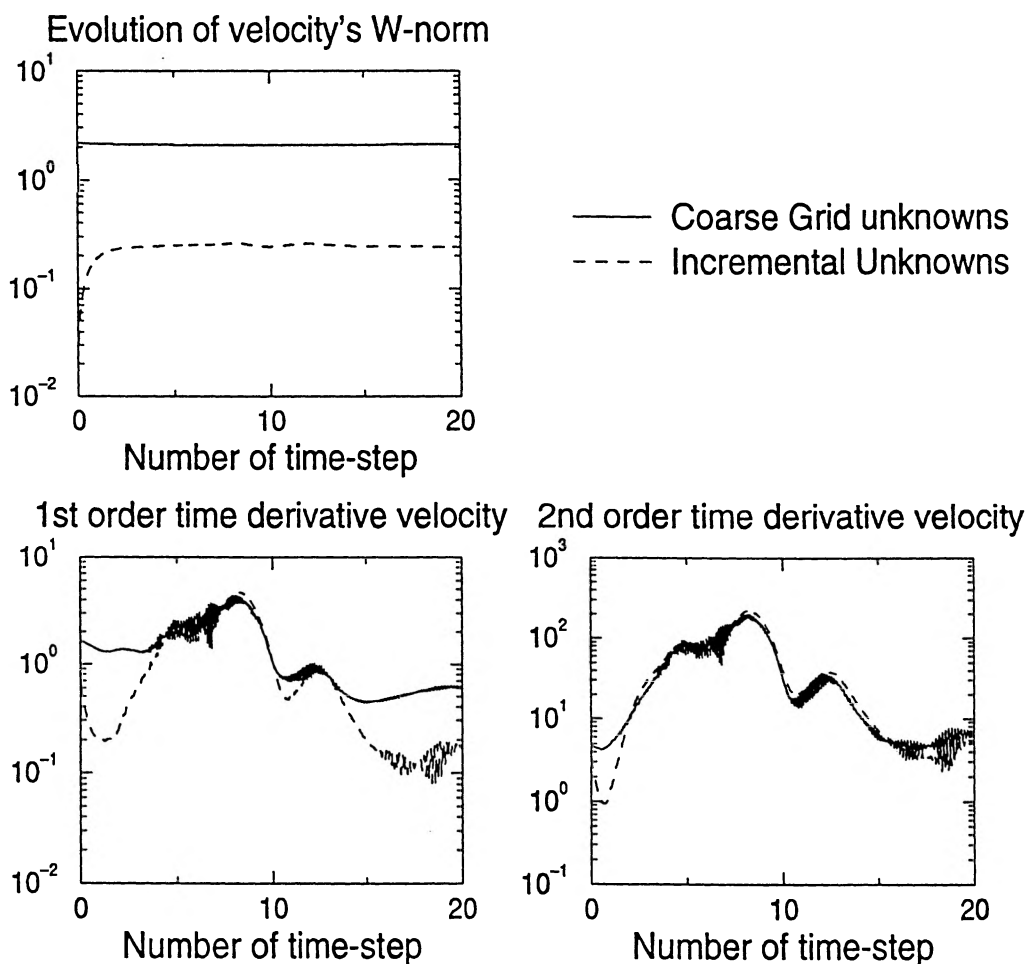


FIG. 3.4 - Cas test de la cavité entraînée régularisée à $Re=10^4$, 257×257 pts

Nous pensons que l'évolution même globale des quantités détaillée ci-dessous donne une information de la dynamique du phénomène que nous simulons. Par exemple, on peut déjà s'apercevoir d'une nette différence entre l'évolution de la norme des taux d'accroissement de la vitesse. A faible nombre de Reynolds, on observe que pour un écoulement devenant très vite stationnaire, les taux d'accroissement de la vitesse décroissent très régulièrement. Alors que pour un assez grand nombre de Reynolds, (figure précédente), on observe une longue période transitoire pendant laquelle la dynamique s'installe, entraînant alors, des variations de ces taux.

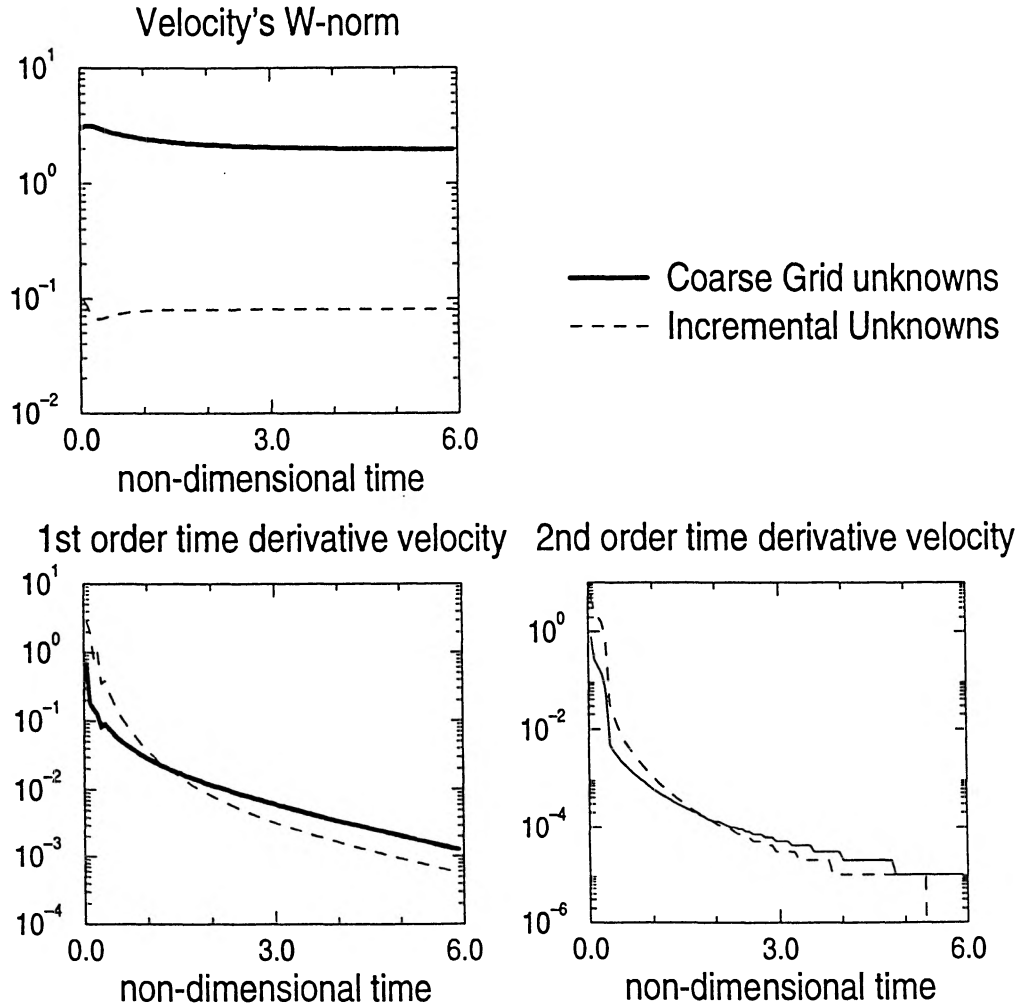


FIG. 3.5 - Cas test de la cavité entraînée régularisée à $Re=100$, 65×65 pts

De façon formelle, le problème de Navier-Stokes directement hiérarchisé peut s'écrire sous la forme suivante :

$$\begin{cases} \frac{\partial \mathbf{u}_h}{\partial t} + \frac{1}{Re} A.S\bar{\mathbf{u}}_h + B.T.\bar{p}_h + b_h(\mathbf{u}_h, \mathbf{u}_h) = \mathbf{f}_h & \text{dans } \Omega_h \\ \mathbf{u}_h = \mathbf{g}_h & \text{sur } \Gamma_h, \end{cases} \quad (3.1)$$

A partir de ce système (3.1), nous pouvons comptabiliser plusieurs actions de chacun des opérateurs sur la grille grossière. Nous regarderons l'effet du terme régularisant sur les inconnues grossières, décomposé en celui provenant effectivement de la grille grossière ($Ay(t)$) et de celui obtenu par le changement de base des S.I.U. ($A.R.y(t)$). L'évolution de l'effet du laplacien sur les S.I.U., restreint à la grille grossière ($A.z(t)$) est aussi prise en compte (cf. 3.6).

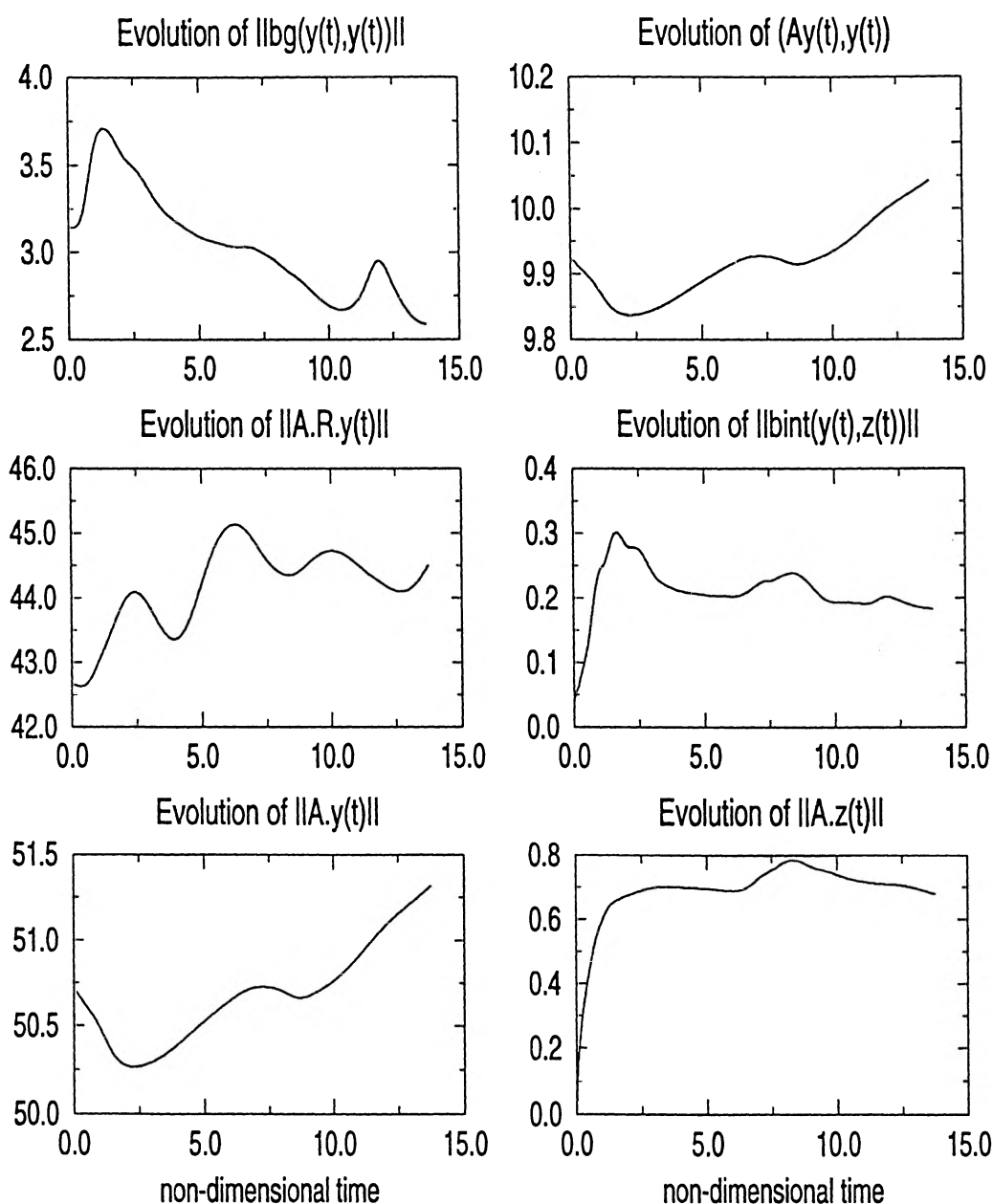


FIG. 3.6 - Cas test de la cavité entraînée régularisée à $Re=10^4$, 257×257 pts

Nous faisons intervenir un nouvel opérateur discret, le terme non linéaire d'interaction entre la vitesse sur la grille grossière et les S.I.U. de la vitesse, défini comme suit :

$$b_{int}(y_{3h}(t), z_h(t)) = \mathcal{R}_g b_h(u_h, u_h) - b_{3h}(y_{3h}, y_{3h}),$$

$\mathcal{R}_g$ étant la restriction sur la grille grossière. Nous considérons l'évolution de ce terme non-linéaire d'interaction, ainsi que du terme inertiel classique approché sur la grille grossière (noté $bg(y_{3h}(t), y_{3h}(t))$).

Nous présentons aussi la variation relative de chacune de ces quantités au cours du temps à la figure suivante.

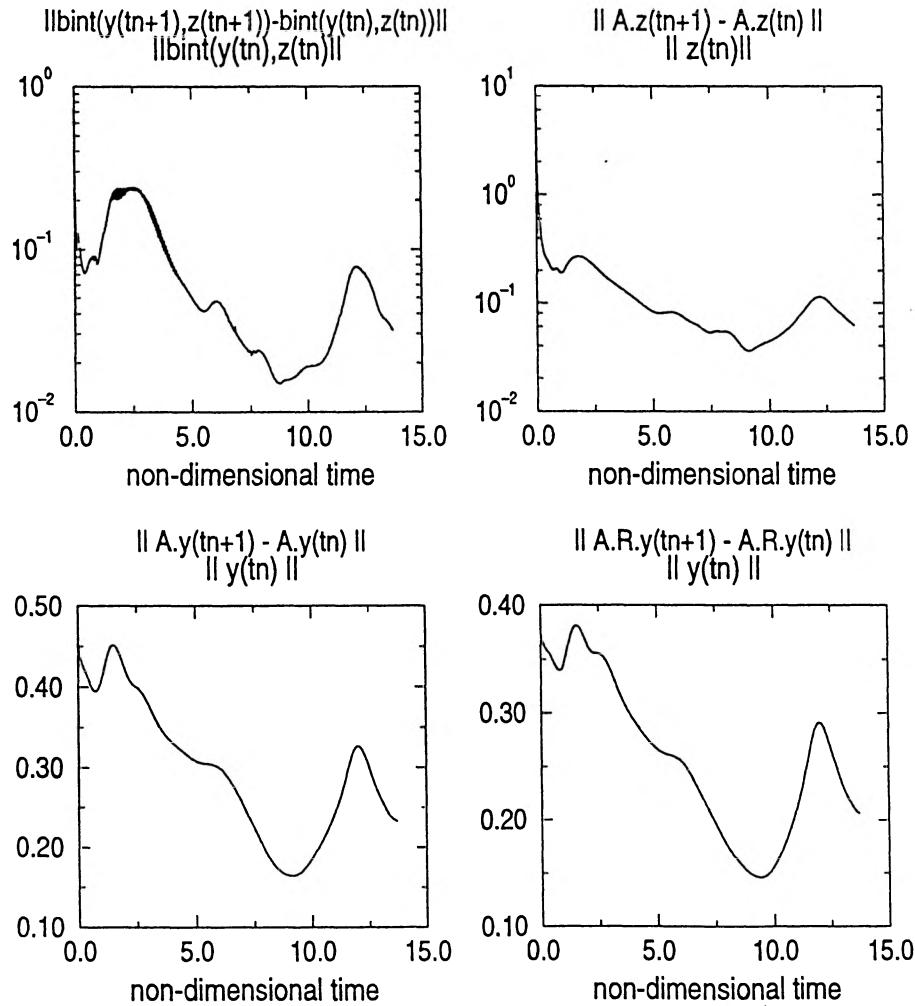


FIG. 3.7 - Cas test de la cavité entraînée régularisée à $Re=10^4$, 257×257 pts

Dans le cas de la simulation pour un nombre de Reynolds de 10^4 , pour l'intervalle de temps considéré $[0, 15]$, on se rend compte que toutes les courbes de variations des quantités (cf. 3.7) considérées présentent un minimum à peu près au même temps (environ entre $T = 8$ et $T = 9$), et un maximum local à $T = 12$.

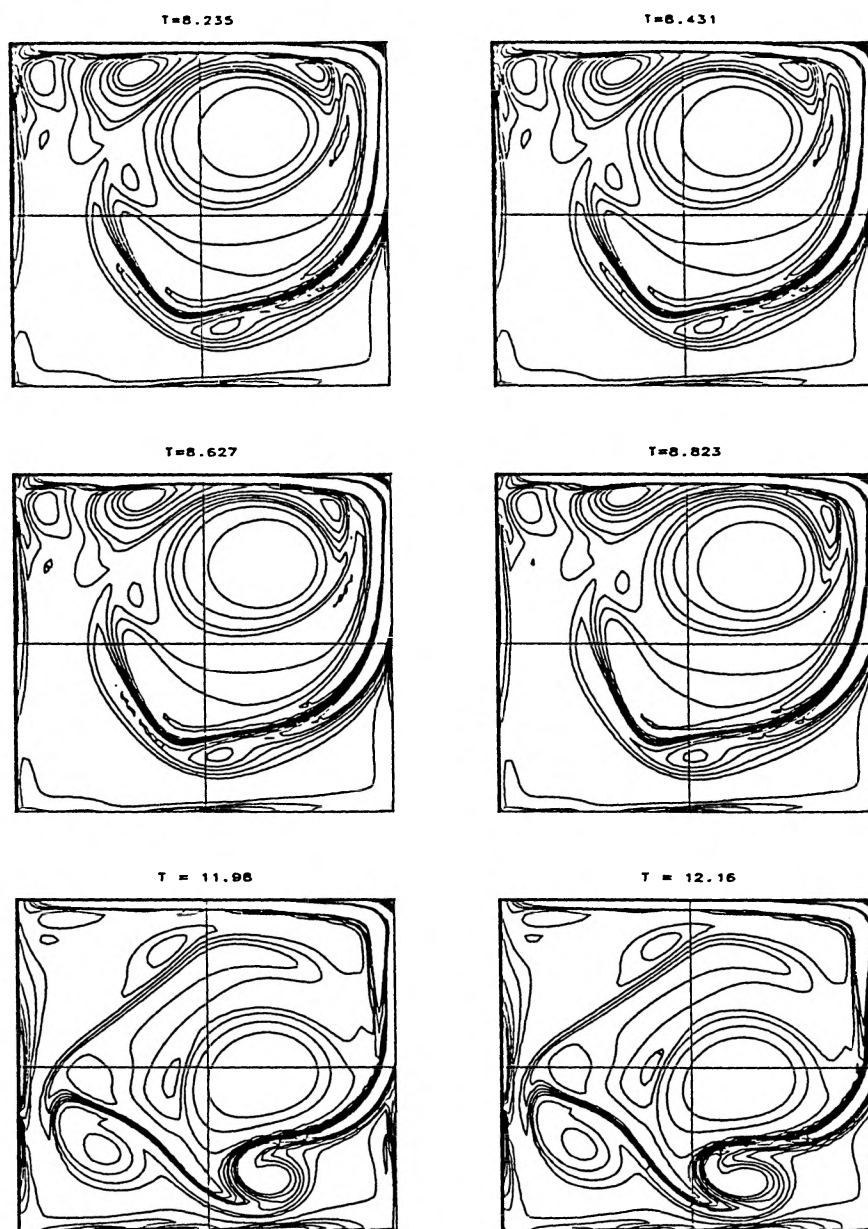


FIG. 3.8 - Lignes d'Isovorticité à $T \approx 8.5$ et $T \approx 12.0$

On se rend bien compte que la dynamique intérieure de la cavité qui, une fois les structures cohérentes créées, est encore très faible à $T \approx 8.5$ et, qui s'installe complètement à T proche de 12.

3.3 Problème de Stokes généralisé sur la grille grossière

La considération que nous cherchons à développer se rapprochant du concept des méthodes de type Galerkin non linéaire consiste à ne calculer à certains instants, que les inconnues de la grille grossière du système discret provenant des équations de Navier-Stokes. Les inconnues de la grille complémentaire plus fine sont les S.I.U. qui seront figées à la valeur de leur itérée précédente. La grande difficulté numérique à surmonter vient du couplage des inconnues des différentes grilles lié, d'une part à la non orthogonalité des S.I.U. pour le produit scalaire de $H_0^1(\Omega)$ induit et d'autre part à la présence du terme non linéaire. Aussi, il nous apparaît fondamental de s'interroger sur les opérateurs discrets intervenant sur chacune des grilles du système afin d'obtenir un problème sur la grille grossière proche du vrai problème de Stokes généralisé (itération de Navier-Stokes) sur la grille en question, bien posé et surtout, facile à résoudre numériquement.

Après avoir rappeler et analyser la structure des opérateurs classiques intervenant dans la base nodale, nous présentons deux procédés permettant d'obtenir un système sur la grille grossière, et discutons de leur résolution. Le premier, inhérent aux méthodes de différences finies, s'obtient par restriction sur la grille grossière du système total hiérarchisé après une approximation des opérateurs discrets comme dans la base nodale. Le deuxième procédé s'obtient par la hiérarchisation faite sur la formulation variationnelle du problème, de manière analogue à la technique utilisée pour la discrétisation par éléments finis hiérarchiques implémenté par [PAS92] et [CLT95].

3.3.1 Analyse des opérateurs intervenant dans la base nodale

Considérons l'écriture habituelle formalisée d'une itération d'un schéma numérique de Cranck-Nicholson sur les termes diffusifs et d'Euler-explicite sur le terme convectif, ce qui nous donne un système discret obtenu comme à la section 2.1 :

$$\left\{ \begin{array}{ll} (I + \frac{\Delta t}{2.Re} A).u_h^{n+1} + B.p_h^{n+1} &= (I - \frac{\Delta t}{2.Re} A).u_h^n + \Delta t b(u_h^n, u_h^n) \quad \text{dans } \Omega_h \\ {}^t B.u_h^{n+1} &= 0 \quad \text{dans } \Omega_h \\ u_h^{n+1} &= g_h \quad \text{sur } \Gamma_h \end{array} \right. \quad (3.2)$$

On reconnaît le problème de Stokes généralisé pour lequel nous détaillerons les différents opérateurs de (3.2) et ceci pour un niveau de grille de S.I.U. Nous reprenons le même nombre de points de discrétisation que celui considéré au paragraphe 3.1.1, soit $\Omega =]a, b[\times]c, d[$, $h_k = \frac{1}{3N_k}$, $k = 1, 2$ le pas de discrétisation dans chacune des directions ($N_k \in \mathbb{N}$). Chacun des opérateurs de changement de base, de la base des S.I.U. à la base nodale, pour la vitesse et la pression garde la même structure que

celui plus classique de la base des I.U. à la base nodale. Nous notons S celui agissant sur la vitesse et T celui de la pression, qui s'écrivent sous forme de matrice blocs :

$$S = \left[\begin{array}{c|c} \frac{I}{R_1} & 0 \\ \hline 0 & \frac{I}{R_2} \end{array} \right], \quad T = \left[\frac{I}{\tilde{R}} \right],$$

où,

R_1 et R_2 sont les deux interpolateurs agissant indépendamment sur chacune des composantes grossières de la vitesse de tailles respectives $((N_1 - 1) \times N_2, 8N_1N_2 - 2N_2)$ et $((N_1 \times (N_2 - 1), 8N_1N_2 - 2N_1)$,

$\tilde{R}$ est l'interpolateur agissant sur la grille grossière de la pression de dimension $(N_1 \times N_2, 4N_1 \times N_2)$,

I dénote les blocs identité indépendamment de leur dimension.

Nous conviendrons de faire l'analogie de la matrice S avec son premier grand bloc (car l'analyse que nous faisons des opérateurs du système hiérarchisée à une composante n'a aucun mal à se généraliser pour un problème vectoriel) et de noter :

$$S = \left[\frac{I}{R} \right] \text{ quand il le sera possible.}$$

La discrétisation centrée usuelle du laplacien (à 5 points) avec le rangement grille par grille fait intervenir la matrice blocs suivante :

$$A = \left[\begin{array}{c|c} A_{11} & A_{12} \\ \hline A_{21} & A_{22} \end{array} \right],$$

et on a les informations suivantes sur les différents blocs :

$$A_{11} = \frac{4}{h_1^2} I, \text{ de dimension } ((N_1 - 1) \cdot N_2)^2$$

$A_{12} = {}^t A_{21}$, est une matrice rectangulaire creuse, formée de quatre coefficients non nuls (égaux à $-1/h_1^2$) par ligne dont deux situés dans les $2(N_1 - 1)(N_2 - 1)$ premières colonnes (S.I.U. de type (FV)), et les deux autres dans les $2N_1N_2$ colonnes suivantes (S.I.U. de type (FH)).

A_{22} est une matrice creuse carrée de dimension $(2N_2 \cdot (4N_1 - 1))^2$ que nous n'explicitons pas plus pour l'instant.

Il est utile aussi de rappeler la structure de l'opérateur approchant le gradient de la pression (la même convention de l'équation scalaire sera adoptée), après avoir réordonné les différentes grilles d'inconnues :

$$B = \left[\begin{array}{c|c} 0 & B_{12} \\ \hline B_{21} & B_{22} \end{array} \right]$$

de dimension $(3N_2(3N_1 - 1), 9N_1N_2)$,

le bloc nul à pour dimension $(N_2(N_1 - 1), N_1.N_2)$ et exprime que le gradient approché au second ordre par différences centrées n'est calculé qu'à partir de contributions de la grille fine complémentaire de la pression.

B_{12} est un bloc creux de dimension $(N_2(N_1 - 1), 8N_1N_2)$ formé de lignes contenant deux coefficients non nuls $(\pm 1/(2h_1))$

B_{21} est un bloc creux de dimension $(2N_2(4N_1 - 1), N_1N_2)$ ayant seulement une valeur non nulle $(\pm 1/(2h_1))$ par ligne (pour les lignes des S.I.U. de type (FH)).

B_{22} bloc creux de dimension $(2N_2(4N_1 - 1), N_1N_2)$ ayant deux valeurs non nulles pour les lignes correspondant aux S.I.U. de type (FV) et (FC) et d'une valeur pour les lignes des S.I.U. de type (FH).

Pour l'autre dimension qu'il a été convenu d'omettre, les blocs correspondant à B_{21} et B_{22} auront une structure équivalente avec les lignes correspondant aux S.I.U. de type (FV) et (FH) interverties.

3.3.2 Restriction directe

On considère le problème de Stokes généralisé (3.2) directement hiérarchisé, soit :

$$\left\{ \begin{array}{ll} (I + \alpha A)S\bar{u}_h^{n+1} + BT\bar{p}_h^{n+1} &= f_h^n \quad \text{dans } \Omega_h \\ {}^tBS\bar{u}_h^{n+1} &= 0 \quad \text{dans } \Omega_h \\ \bar{u}_h^{n+1} &= g_h \quad \text{sur } \Gamma_h, \end{array} \right. \quad (3.3)$$

avec $\alpha = \frac{\Delta t}{2Re}$ et $f_h^n = (I - \alpha A)u_h^n + \Delta tb(u_h^n, u_h^n)$. Il est possible d'analyser les différents blocs intervenant sur la restriction de la première équation de (3.3) à la grille grossière. Le premier bloc de AS est $A_{11} + A_{12}R$,

$$AS = \left[\begin{array}{c|c} A_{11} + A_{12}R & A_{12} \\ \hline A_{21} + A_{22}R & A_{22} \end{array} \right],$$

et on peut se rendre compte qu'il s'agit d'un bloc symétrique. En effet, si l'on décrit l'action de l'opérateur discret du laplacien sur la grille grossière de la vitesse on obtient, en utilisant les formules du paragraphe 3.1.1 pour les S.I.U. de type (FV) et (FH) situés au moins à une distance d'un pas de Γ :

$$\begin{aligned}
 (-\Delta_h u_h)_{3i,3j+3/2} &= \frac{1}{h_1^2} \left(4u_{3i,3j+3/2} - u_{3i-1,3j+3/2} - u_{3i+1,3j+3/2} \right. \\
 &\quad \left. - u_{3i+1,3j+5/2} - u_{3i,3j+1/2} \right) \\
 &= \frac{1}{h_1^2} \left(4u_{3i,3j+3/2} \right. \\
 &\quad - z_{3(i-1)+2,3j+3/2}^u - \frac{1}{3}(u_{3(i-1),3j+3/2} + 2u_{3i,3j+3/2}) \\
 &\quad - z_{3i+1,3j+3/2}^u - \frac{1}{3}(2u_{3i,3j+3/2} + u_{3(i+1),3j+3/2}) \\
 &\quad - z_{3i,3j+5/2}^u - \frac{1}{3}(2u_{3i,3j+3/2} + u_{3(i+1),3(j+1)+3/2}) \\
 &\quad \left. - z_{3i+1,3(j-1)+2+3/2}^u - \frac{1}{3}(u_{3i,3(j-1)+3/2} + 2u_{3i,3j+3/2}) \right) \\
 &= \frac{1}{3h_1^2} \left(4y_{3i,3j+3/2}^u - y_{3(i-1),3j+3/2}^u - y_{3(i+1),3j+3/2}^u \right. \\
 &\quad \left. - y_{3i,3(j+1)+3/2}^u - y_{3i,3(j-1)+3/2}^u \right) \\
 &\quad - \frac{1}{h_1^2} (z_{3i-1,3j+3/2}^u + z_{3i+1,3j+3/2}^u + z_{3i,3j+5/2}^u + z_{3i,3j+3/2}^u),
 \end{aligned}$$

donc, pour ces points de la grille grossière, le bloc $A_{11} + A_{12}R = -3\Delta_{3h}$. Pour les autres points de la grille grossière, nous obtenons une approximation de l'opérateur grossier du laplacien identique à celle faite sur la grille fine d'inconnues par extrapolation linéaire. Par exemple, regardons le traitement pour la première composante de la vitesse aux points $(3ih_1, 3h_2/2)$:

$$\begin{aligned}
 (-\Delta_h u_h)_{3i,3/2} &= \frac{1}{h_1^2} (4u_{3i,3/2} - u_{3i,1/2} - u_{3i,5/2} - u_{3i-1,3/2} - u_{3i+1,3/2}) \\
 &= \frac{1}{h_1^2} (4u_{3i,3/2} - z_{3i,1/2}^u - \frac{1}{3}(u_{3i,3/2} + 2u_{3i,0}) \\
 &\quad - z_{3i,5/2}^u - \frac{1}{3}(u_{3i,3+3/2} + 2u_{3i,3/2}) \\
 &\quad - z_{3i-1,3/2}^u - \frac{1}{3}(u_{3(i-1),3/2} + 2u_{3i,3/2}) \\
 &\quad - z_{3i+1,3/2}^u - \frac{1}{3}(u_{3(i+1),3/2} + 2u_{3i,3/2})) \\
 &= \frac{1}{3h_1^2} (2u_{3i,3/2} - u_{3(i-1),3/2} - u_{3(i+1),3/2}) \\
 &\quad + \frac{1}{3h_1^2} (3u_{3i,3/2} - 2u_{3i,0} - u_{3i,3+3/2}) \\
 &\quad - \frac{1}{h_1^2} (z_{3i,1/2}^u + z_{3i,5/2}^u + z_{3i-1,3/2}^u + z_{3i+1,3/2}^u),
 \end{aligned}$$

on reconnaît les coefficients obtenus par l'extrapolation linéaire pour approcher le terme de dérivation du second ordre suivant l'axe des y . Nous pouvons conclure que $A_{11} + A_{12}R = -3\Delta_{3h}$.

On trouve aussi une expression simple de l'opérateur intervenant sur la grille grossière de la pression. Après avoir trouvé la structure bloc de la matrice rectan-

gulaire BT , on utilise les formules de définition de la hiérarchisation de la pression (pour les S.I.U. de type (FH)) au paragraphe 3.1.2, et on obtient :

$$BT = \left[\frac{B_{12}\tilde{R}}{B_{21} + B_{22}\tilde{R}} \middle| \frac{B_{12}}{B_{22}} \right],$$

$$\begin{aligned} \left(\frac{\partial p}{\partial x} \right)_{3i,3j+3/2}^h &= \frac{1}{h_1} \left(p_{3(i-1)+7/2,3j+3/2} - p_{3(i-1)+5/2,3j+3/2} \right) \\ &= \frac{1}{h_1} \left(p_{3(i-1)+7/2,3j+3/2}^z + \frac{1}{3} (2p_{3i+3/2,3j+3/2} + p_{3(i-1)+3/2,3j+3/2}) \right. \\ &\quad \left. - p_{3(i-1)+5/2,3j+3/2}^z - \frac{1}{3} (2p_{3(i-1)+3/2,3j+3/2} + p_{3i+3/2,3j+3/2}) \right) \\ &= \frac{1}{3h_1} \left(p_{3i+3/2,3j+3/2}^y - p_{3i-3/2,3j+3/2}^y \right) \\ &\quad + \frac{1}{h_1} \left(p_{3i+1/2,3j+3/2}^z - p_{3i-1/2,3j+3/2}^z \right) \\ &= \left(\frac{\partial p^y}{\partial x} \right)_{3i,3j+3/2}^{3h} + \left(\frac{\partial p^z}{\partial x} \right)_{3i,3j+3/2}^h, \end{aligned}$$

pour tous les points de la grille grossière intérieurs au domaine.

On obtient l'égalité analogue pour la deuxième composante du gradient en utilisant la définition des S.I.U. de type (FV) soit,

$$\left(\frac{\partial p}{\partial y} \right)_{3i+3/2,3j}^h = \left(\frac{\partial p^y}{\partial y} \right)_{3i+3/2,3j}^{3h} + \left(\frac{\partial p^z}{\partial y} \right)_{3i+3/2,3j}^h,$$

toujours pour les points de la grille grossière intérieurs au domaine.

Il reste à analyser l'opérateur intervenant dans l'équation de (3.3) traduisant l'incompressibilité :

$${}^tBS = \left[\frac{{}^tB_{21}R}{{}^tB_{12} + {}^tB_{22}R} \middle| \frac{{}^tB_{21}}{{}^tB_{22}} \right],$$

La contribution de ${}^tB_{21}Ry$ se calcule au point de discrétisation de la grille grossière de la pression et, en utilisant la définition de R (pour les S.I.U. de type (FV) et (FH)) comme précédemment, on a :

$$\begin{aligned} \left({}^tB_{21}Ry \right)_{3i+3/2,3j+3/2} &= \frac{1}{h_1} \left[\left(R_1 y^u \right)_{3i+2,3j+3/2} - \left(R_1 y^u \right)_{3i+1,3j+3/2} \right] \\ &\quad + \frac{1}{h_2} \left[\left(R_2 y^v \right)_{3i+3/2,3j+2} - \left(R_2 y^v \right)_{3i+3/2,3j+1} \right] \\ &= \left(\nabla_{3h} \cdot u_h \right)_{3i+3/2,3j+3/2}. \end{aligned}$$

Après avoir explicité ces opérateurs, on peut restreindre le système (3.3) sur la grille grossière et obtenir le problème suivant :

$$\begin{cases} (I - 3\alpha\Delta_{3h})\mathbf{y}_{3h}^{n+1} + \nabla_{3h} \left(p_{3h}^y \right)^{n+1} = \mathbf{f}_{3h}^n + \alpha C_h \mathbf{z}_h^{n+1} - \mathcal{R}_g \nabla_h \left(p_h^z \right)^{n+1} & \text{dans } \Omega_{3h} \\ \nabla_{3h} \cdot \mathbf{y}_{3h}^{n+1} = -{}^tB_{21} \mathbf{z}_h^{n+1} & \text{dans } \Omega_{3h} \\ \mathbf{y}_{3h}^{n+1} = \mathbf{g}_{3h} & \text{sur } \Gamma_{3h} \end{cases} \quad (3.4)$$

où, nous notons f_{3h}^n la restriction à la grille grossière du second membre de la première équation de (3.3). $\mathcal{R}_g$ constitue la restriction à la grille grossière, et C_h est l'opérateur qui, aux quatre S.I.U. chacune des composantes de la vitesse situés à la distance d'un pas d'un point de la grille grossière, associe en ce point leur somme réduite d'un facteur h_1^2 .

3.3.3 Restriction sur la forme variationnelle

Comme nous l'avons mentionné au premier chapitre, les schémas numériques en discrétisation par différences finies peuvent découler d'une formulation variationnelle. Mais, une analyse détaillée de la structure de l'opérateur laplacien dans la base des S.I.U. nous montre qu'il n'est pas toujours possible de faire l'analogie entre les éléments finis hiérarchiques et les S.I.U. (ou I.U.) en différences finies. Ensuite, nous discuterons de la restriction du problème discret sur la forme variationnelle.

Analyse

Pour se convaincre de ce point, il suffit de faire le calcul du bloc intervenant sur la grille grossière du laplacien dans la base des S.I.U. En reprenant les notations du paragraphe précédent, avec A le laplacien à 5 points, on obtient :

$$\hat{A} = {}^t S.A.S = \left(\begin{array}{c|c} \hat{A}_{11} & \\ \hline \frac{A_{11} + A_{12}R + {}^t R(A_{21} + A_{22}R)}{A_{21} + A_{22}R} & \frac{A_{12} + {}^t R A_{22}}{A_{22}} \end{array} \right)$$

On a vu que $A_{11} + A_{12}R = 3(-\Delta_{3h})$ où l'approximation se fait par le même schéma à 5 points, mais avec le pas grossier. Pour calculer $\hat{A}_{11}$, il suffit de détailler le produit scalaire de W_h suivant, où les contributions interviennent uniquement sur la grille des S.I.U. :

$$((A_{21} + A_{22}R)y^u, Ry^v) \quad \text{où} \quad \bar{u} = \begin{pmatrix} y^u \\ z^u \end{pmatrix} \quad \text{et,} \quad \bar{v} = \begin{pmatrix} y^v \\ z^v \end{pmatrix}.$$

Selon la figure 3.9, il faut calculer les contributions localisées aux points de la grille des S.I.U. de type (FV), (FH) et (FC), soient v_{ij} , h_{ij} et c_{ij} si l'on considère une seule

composante de la vitesse.

$$\begin{aligned}
 h^2(A_{21}y^u + A_{22}Ry^u)_{v_{02}} &= -y_5^u + 4(R_1y^u)_{v_{02}} - (R_1y^u)_{c_{14}} - (R_1y^u)_{v_{01}} - (R_1y^u)_{c_{23}} \\
 &= \frac{1}{6}(3y_5^u + y_2^u) - \frac{1}{12}(3y_4^u + 3y_6^u + y_1^u + y_3^u), \\
 h^2(A_{21}y^u + A_{22}Ry^u)_{v_{01}} &= \frac{1}{6}(3y_2^u + y_5^u) - \frac{1}{12}(3y_1^u + 3y_3^u + y_4^u + y_6^u), \\
 h^2(A_{21}y^u + A_{22}Ry^u)_{h_{02}} &= \frac{1}{6}(3y_5^u + y_4^u) - \frac{1}{12}(3y_2^u + 3y_8^u + y_1^u + y_7^u), \\
 h^2(A_{21}y^u + A_{22}Ry^u)_{c_{14}} &= 4(R_1y^u)_{c_{14}} - (R_1y^u)_{h_{02}} - (R_1y^u)_{v_{02}} - (R_1y^u)_{c_{12}} \\
 &\quad - (R_1y^u)_{c_{13}} \\
 &= \frac{1}{6}(y_2^u + y_4^u - y_5^u - y_1^u).
 \end{aligned}$$

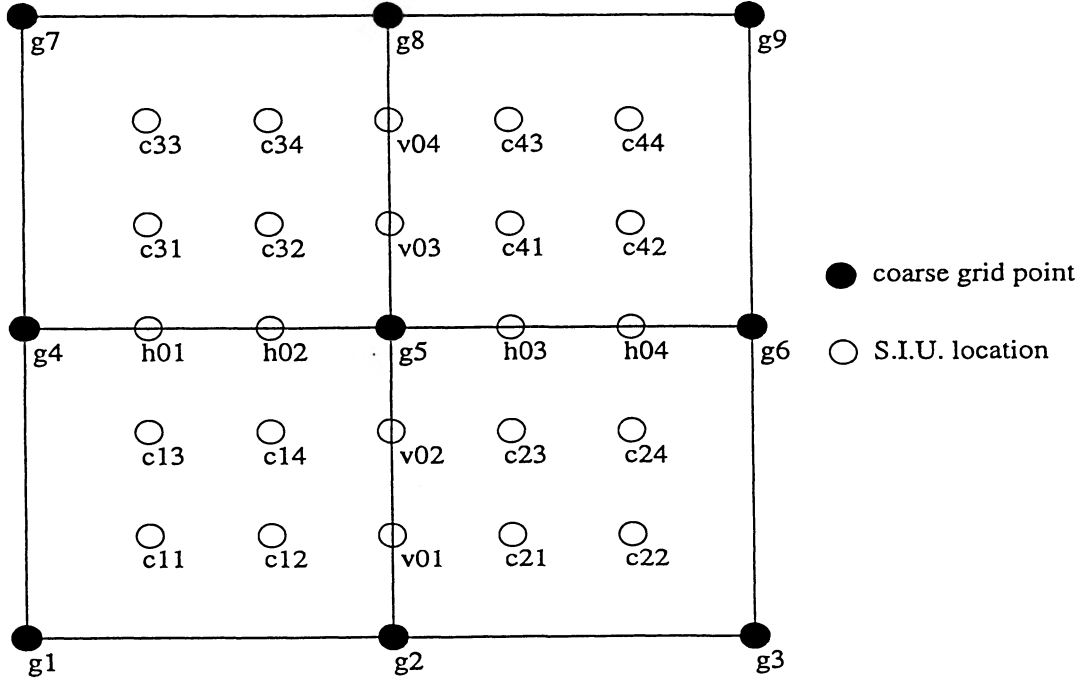


FIG. 3.9 - Localisation des S.I.U. liées à un point grossier

Pour retrouver la structure de $\hat{A}_{11}$, il suffit de repérer le coefficient de y_5^u dans le produit scalaire, soit, sommer les différentes expressions de $h^2(A_{21}y^u + A_{22}R_1y^u)$ avec les poids donnés par l'interpolateur R_1 . Pour cela, il suffit de regrouper les S.I.U. en trois groupes auxquels sont affectés les mêmes coefficients. Le premier est constitué de $\{v_{02}, v_{03}, h_{02}, h_{03}\}$ pour lequel le coefficient est $\frac{2}{3}$, le deuxième, $\{v_{01}, v_{04}, h_{01}, h_{04}\}$ de coefficient $\frac{1}{3}$, et le dernier constitué des c_{ij} qui par symétries ne donnent aucune contribution.

$$\sum_{1 \leq i, j \leq 4} h^2(A_{21}y^u + A_{22}R_1y^u)_{c_{ij}} \cdot (R_1y^u)_{c_{ij}} = 0,$$

car,

$$\begin{aligned} (A_{21}y^u + A_{22}R_1y^u)_{c_{11}} &= (A_{21}y^u + A_{22}R_1y^u)_{c_{14}} \\ &= -(A_{21}y^u + A_{22}R_1y^u)_{c_{12}} = -(A_{21}y^u + A_{22}R_1y^u)_{c_{13}} \end{aligned}$$

et $(R_1y^u)_{c_{12}} + (R_1y^u)_{c_{13}} = (R_1y^u)_{c_{11}} + (R_1y^u)_{c_{14}} \quad \forall 1 \leq i \leq 4$ et ceci quelles que soient les quatre cellules du maillage grossier.

On finit par trouver l'expression de l'opérateur grossier :

$$\hat{A}_{11}y_5^u = \frac{1}{18}(52y_5^u - 8y_2^u - 8y_4^u - 8y_6^u - 8y_8^u - 5y_1^u - 5y_3^u - 5y_7^u - 5y_9^u).$$

Il s'agit bien d'une discrétisation d'ordre 2 de l'opérateur laplacien sur la grille grossière (la somme nulle des coefficients et leur symétrie permet la déduction évidente en utilisant la formule de Taylor) mais l'on ne retrouve pas celle que nous avons utilisé sur la grille fine. Ceci vient du fait que l'approximation du laplacien par un schéma à 5 points communément utilisée pour des raisons d'économie (le coût d'un produit matrice-vecteur étant quasiment réduit de moitié) ne provient pas d'une formulation variationnelle avec une discrétisation de type MAC.

Par contre, si nous avons choisi la discrétisation standard à neuf points donnée par le calcul des produits scalaires des fonctions de base de W_h par une méthode classique de Galerkin (comme en discrétisation par éléments finis) nous aurions obtenu le même opérateur discret.

Restriction

Considérons le problème de Navier-Stokes sous forme faible (1.6) avec une discrétisation temporelle obtenue de façon explicite, dans la nouvelle base des S.I.U. Ce qui nous mène au problème de Stokes généralisé discret suivant :

$$\begin{cases} {}^tS.(M + \alpha A).S\bar{u}_h^{n+1} + {}^tS.B.T\bar{p}_h^{n+1} = {}^tSf_h^n & \text{dans } \Omega_{3h} \\ {}^tT.B.S\bar{u}_h^{n+1} = 0 & \text{dans } \Omega_{3h} \\ \bar{u}_h = g_h & \text{sur } \Gamma_h, \end{cases} \quad (3.5)$$

où M est la matrice de masse obtenue par produit scalaire dans $L^2(\Omega)$ des fonctions de base de la vitesse, et jouant le même rôle que l'identité. D'après l'étude

menée au chapitre précédent, nous savons que le conditionnement de l'opérateur $I + \alpha A$ est bien plus petit que celui de ${}^tS.(M + \alpha A).S$, ce qui explique en partie, que nous n'ayons pas fait le choix de cette restriction. L'autre raison concerne la lourdeur d'implémentation, de la résolution par l'algorithme de projection du problème de Stokes généralisé une fois restreint. Il suffit de remarquer que le projecteur sur l'espace à divergence nulle est défini comme

$$\hat{P} = I - {}^tSBT[{}^tT{}^tBS{}^tSBT]^{-1}{}^tT{}^tBS,$$

et il nécessite pour son application un trop grand nombre d'opérations.

3.3.4 Résolution du problème grossier

On suppose que l'on se situe à un instant de la simulation, où les structures incrémentales des inconnues évoluent très peu. Par conséquent nous allons les laisser figées à leur valeur précédente et résoudre le problème grossier suivant :

$$\left\{ \begin{array}{ll} (I - 3\alpha\Delta_{3h})\tilde{y}_{3h}^{n+1} + \nabla_{3h}\left(\tilde{p}_{3h}^y\right)^{n+1} &= k_{3h}^n \quad \text{dans } \Omega_{3h} \\ \nabla_{3h}.\tilde{y}_{3h}^{n+1} &= -{}^tB_{21}z_h^n \quad \text{dans } \Omega_{3h} \\ \tilde{y}_{3h}^{n+1} &= g_{3h} \quad \text{sur } \Gamma_{3h} \end{array} \right. \quad (3.6)$$

où, $k_{3h}^n = f_{3h}^n + C_h z_h^n - \mathcal{R}_g \nabla_h(p_h^z)^n$, et f^n , C_h et $\mathcal{R}_g$ ont été définis lors de l'analyse précédente.

Nous reconnaissons les mêmes opérateurs intervenant dans le problème classique de Stokes généralisé, mais la spécificité intervient sur un second membre non nul de la deuxième équation de (3.6). Or la résolution par l'algorithme de projection orthogonale du chapitre précédent nécessite une réelle relation d'incompressibilité.

Pour y remédier, nous allons considérer le problème grossier résolu à l'itération précédente (3.7) et effectuer un relèvement à l'aide de la vraie solution (sur la grille grossière) de Navier-Stokes trouvée à l'itération précédente.

$$\left\{ \begin{array}{ll} (I - 3\alpha\Delta_{3h})y_{3h}^n + \nabla_{3h}\left(p_{3h}^y\right)^n &= f_{3h}^{n-1} + \alpha C_h z_h^n - \mathcal{R}_g \nabla_h(p_h^z)^n \quad \text{dans } \Omega_{3h} \\ \nabla_{3h}.y_{3h}^n &= -{}^tB_{21}z_h^n \quad \text{dans } \Omega_{3h} \\ y_{3h}^n &= g_{3h} \quad \text{sur } \Gamma_{3h} \end{array} \right. \quad (3.7)$$

Ce qui revient à faire la différence entre chacune des équations du problème (3.6) et celles de (3.7), pour obtenir (3.8) avec les notations suivantes :

$$\begin{aligned} Y &= \tilde{y}_{3h}^{n+1} - y_{3h}^n \\ P^Y &= \left(\tilde{p}_{3h}^y\right)^{n+1} - \left(p_{3h}^y\right)^n \end{aligned}$$

$$\left\{ \begin{array}{ll} (I - 3\alpha\Delta_{3h})Y + \nabla_{3h}P^Y &= f_{3h}^n - f_{3h}^{n-1} \quad \text{dans } \Omega_{3h} \\ \nabla_{3h}.Y &= 0 \quad \text{dans } \Omega_{3h} \\ Y &= 0 \quad \text{sur } \Gamma_{3h}. \end{array} \right. \quad (3.8)$$

On peut remarquer que la dynamique habituellement alimentée aux itérations de Navier-Stokes par la condition limite et se propageant par le second membre, semble se transmettre uniquement par le second membre. On s'attend à ce que la résolution du système (3.8) soit plus difficile que celle du système classique (pour le même nombre de points de discrétisation), car le coefficient devant l'opérateur Δ_{3h} est trois fois plus important ce qui ne peut qu'augmenter le conditionnement de la matrice à inverser.

Il faut préciser qu'il est aussi nécessaire de préconditionner, avec le laplacien discret, l'algorithme de gradient conjugué. L'algorithme de multigrilles classique développé à la première partie effectue des cycles V de longueur maximale, par conséquent il nécessite un nombre de points intérieurs de discrétisation sur la grille la plus fine de $2^d - 1$ (et un seul point intérieur sur la grille la plus grossière). Donc, cela nous restreint dans le choix du nombre de points sur la grille fine et, d'autre part, nous demande de changer de stratégie sur la grille grossière. Il est possible de considérer un cycle de longueur 1 et d'inverser le laplacien grossier (discrétisé avec 42 points) par une méthode directe (factorisation de Gauss). La factorisation n'étant faite qu'une fois, à l'étape d'initialisation, nous devons noter un gain dans l'inversion par notre nouvelle implémentation.

3.4 Schéma de type Galerkin Non Linéaire

Dans cette section, nous explicitons le premier schéma que nous pouvons qualifier de type Galerkin Non Linéaire, bien que problème discret ne provienne pas directement d'une formulation faible discrète.

Pour l'instant, nous pensons qu'il faille déterminer de nombreuses constantes de manière empirique et nous avons utilisé comme schéma de référence pour la résolution classique, le schéma de Cranck-Nicholson et Euler explicite (ordre 1 en temps).

La figure 3.10 présente schématiquement la stratégie dont nous envisageons la mise au point.

Pour que la méthode garde son intérêt (i.e. qu'elle permette d'obtenir un gain de complexité et/ou de temps), on cherche une configuration qui permette de choisir k le plus grand, l le plus petit avec un nombre de cycles ($nbcycle$) le plus grand.

De plus, nous ne voulons pas que l'on puisse noter une perte significative de précision sur le calcul de la solution.

Comme nous l'avons fait en détail à la section précédente, pour la restriction du problème sur la grille fine au problème grossier, nous décrivons l'étape de prolongement. Il s'agit du prolongement le plus classique que l'on puisse envisager avec des S.I.U. figées, soit :

$$\tilde{y}_{3h}^m \xrightarrow{\text{prolongement}} \begin{pmatrix} \tilde{y}_{3h}^m \\ \tilde{z}_h^{m-k} \end{pmatrix} \xrightarrow{\text{changement de base}} \begin{pmatrix} \tilde{y}_{3h}^m \\ R\tilde{y}_{3h}^m + \tilde{z}_h^{m-k} \end{pmatrix} \equiv \tilde{u}_h^m,$$

Remarque 7

Bien qu'il ne s'agisse que d'une étude faite sur une configuration d'une grille grossière et une grille de S.I.U., nous pouvons sans mal, définir une stratégie plus générale avec un nombre de niveaux de grilles de S.I.U. d_s quelconque ($d_s \geq 1$), et un nombre de points quelconque sur la grille grossière. Mais, cela ne nous paraît pas fondamental d'augmenter le nombre de grilles, dans la mesure où si l'on cherche à simuler assez longtemps sur la grille grossière, il est nécessaire de pouvoir capter avec cette discrétisation, suffisamment d'informations sur la solution.

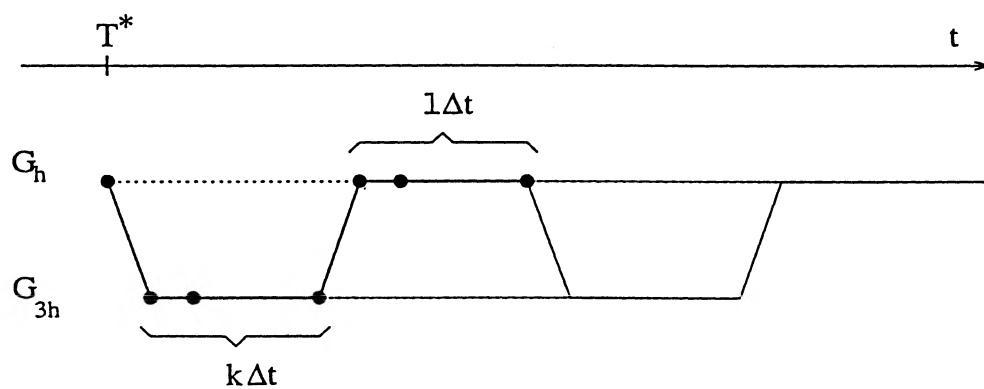


FIG. 3.10 - Deux cycles de résolution

Conclusions et développements

Dans cette partie, nous pouvons noter la création puis la mise en œuvre de nouveaux algorithmes efficaces pour la résolution des problèmes de Stokes et Navier-Stokes instationnaire.

La développement du premier algorithme repose sur plusieurs techniques performantes : il utilise la discrétisation inhérente à la méthode MAC, un algorithme de projection récent préconditionné par une méthode multigrilles efficace. De plus, cette implémentation a été validée dans le cas de la cavité carrée pour un nombre de Reynolds de 5000.

Ensuite, nous avons créé une nouvelle hiérarchisation d'un maillage MAC, et nous nous sommes attachés à proposer une méthode de type Galerkin non linéaire en discrétisation par différences finies. Cela nous a permis de concevoir des estimateurs de dynamique à partir du terme inertiel non linéaire, et de définir une stratégie multi-niveaux. La descente sur la grille grossière est obtenue en figeant les S.I.U. et le terme d'interaction entre les différentes structures. Ensuite, la résolution du problème grossier approché peut se faire en utilisant la même méthode de projection orthogonale que celle utilisée pour le problème sur la grille fine.

Nous pensons affiner notre étape de prolongement en effectuant un lissage des S.I.U. avant de prolonger les composantes grossières sur la grille fine. On pourrait utiliser une **variété inertielle approchée** comme opérateur de lissage, un peu différente de celle introduite par Foias *et al.* (cf. [FMT88]), car le coût d'implémentation de cette dernière, en différences finies, est beaucoup trop grand.

Il serait aussi très intéressant de changer de problème modèle et d'envisager des indicateurs de dynamique non linéaire plus localisés dans l'espace. Cela nous permettrait d'envisager le raffinement local dans les zones de forts gradients, de diminuer dans ce cas la taille de chacune de nos S.I.U., et surtout de résoudre efficacement des écoulements de couche limite.

Finalement, nous pouvons dire que le concept des Inconnues Incrémentales permet de nombreuses avancées et qu'il mérite d'être utilisé dans des domaines variés.

Bibliographie

- [BIG74] A. BEN-ISRAEL and T.N.E. GREVILLE. *Generalized Inverses theory and applications*. Pure & Applied Mathematics. John Wiley & Sons, 1974.
- [BJ90] C.H. BRUNEAU and C. JOURON. An efficient scheme for solving partial differential equations. *Journal of Computational Physics*, 89(2):389–413, 1990.
- [BRA92] R. BRAMLEY. An orthogonal projection algorithm for linear Stokes problems. Technical Report 1190, Center for Supercomputing Research and Development, University of Illinois, Urbana, 1992.
- [CC87] K. CAHOUE and J. CHABARD. Some fast 3d finite element solvers for generalized Stokes problem. Technical Report HE-41/87.03, Electricité De France, D.E.R., 6 quai Watier, 78400 Chatou, 1987.
- [CLT95] C. CALGARO, J. LAMINIE, and R. TEMAM. Dynamical multilevel schemes for the solution of evolution equations by hierarchical finite element methods. Prépublication de l'Université PARIS-SUD, 1995. 95-77.
- [DAU92] O. DAUBE. Resolution of the 2d Navier-Stokes equation in velocity-vorticity form by means of an influence matrix technique. *Journal of Computational Physics*, 103:402–414, 1992.
- [DAV83] R.W. DAVIS. Finite difference methods for fluid flow. In J. Noye and C. Fletcher, editors, *Computational Techniques and Applications: CTAC-83*. North-Holland, 1983.
- [FLE88] C.A.J. FLETCHER. *Computational Techniques for Fluid Dynamics*, volume 2 of *Springer Series in Computational Physics*. Springer-Verlag, 1988.
- [FMT88] C. FOIAS, O. MANLEY, and R. TEMAM. Modelling of the interaction of small and large eddies in two-dimensional turbulent flow. *Math. Mod. and Num. Anal.*, 22:93–114, 1988.

- [GAR92] S. GARCIA. *Etude algébrique et numérique de la méthode des Inconnues Incrémentales*. PhD thesis, Université Paris-XI, 06/1992.
- [GCLU84] P.M. GRESHO, S.T. CHAN, R.L. LEE, and C.D. UPSON. A modified finite element method for solving the time-dependant, incompressible Navier-Stokes equations. part 2: applications. *Int. J. Numerical Methods in Fluids*, 4:619–640, 1984.
- [GGH90] J.W. GOODRICH, K. GUSTAFSON, and K. HALASI. Hopf bifurcation in the driven cavity. *Journal of Computational Physics*, 90:219–261, 1990.
- [GGS82] U. GHIA, K.N. GHIA, and C.T. SHIN. High-re solutions for incompressible flow using the Navier-Stokes equations and a multigrid method. *Journal of Computational Physics*, 48:387–411, 1982.
- [GIR76] V. GIRAULT. A combined finite element and Marker And Cell method for solving Navier-Stokes equations. *Numer. Math.*, 26:39–59, 1976.
- [GOY94] O. GOYON. *Résolution Numérique de problèmes stationnaires et évolutifs non linéaires par la méthode des Inconnues Incrémentales*. PhD thesis, Université Paris-XI, 12/1994.
- [GS87] P.M. GRESHO and R.L. SANI. On pressure boundary conditions for the Navier-Stokes equations. *Int. J. Numerical Methods in Fluids*, 7:1111–1145, 1987.
- [HIR92] C. HIRSCH. *Fundamentals of Numerical Discretization*, volume 1 of *Numerical Computation of Internal and External Flows*. John Wiley & Sons, 1992.
- [HW65] F. HARLOW and J. WELCH. Numerical calculation of time-dependent viscous incompressible flow of fluid with free surfaces. *Phys. Fluids*, 8(12):2181–2189, 1965.
- [LLM90] C. LIU, Z. LIU, and S. McCORMICK. Multigrid methods for flow transition in a planar channel. In *Computational Physics, Plenary Lectures and Expanded Selected Papers from the IMACS 1st International Conference*, University of Colorado at Boulder, USA, 1990. North-Holland.
- [MED95] T. TACHIM MEDJO. *Sur les formulations vitesse-vorticité pour les équations de Navier-Stokes en dimension deux - Implémentation des Inconnues Incrémentales Oscillantes dans les équations de Navier-Stokes et de réaction-diffusion*. PhD thesis, Université Paris-XI, 06/1995.

- [PAS92] F. PASCAL. *Méthodes de Galerkin non linéaires en discrétisation par éléments finis et pseudo-spectrale. Application à la Mécanique des Fluides*. PhD thesis, Université Paris-XI, 01/1992.
- [PT83] R. PEYRET and T.D. TAYLOR. *Computational Methods for Fluid Flow*. Springer-Verlag, 1983.
- [SHE87] J. SHEN. *Résolution numérique des équations de Stokes et Navier-Stokes par les méthodes spectrales*. PhD thesis, Université Paris-XI, 1987.
- [SHE91] J. SHEN. Hopf bifurcation of the unsteady regularized driven cavity flow. *Journal of Computational Physics*, 95:228–245, 1991.
- [TEM84] R. TEMAM. *Navier-Stokes Equations*. North-Holland, 1984.