



Étude du mélange de scalaires en écoulement turbulent et application à la modélisation des petites échelles

Antoine Moreau

► To cite this version:

Antoine Moreau. Étude du mélange de scalaires en écoulement turbulent et application à la modélisation des petites échelles. Mécanique [physics]. Université Claude Bernard, 2002. Français. NNT : 2002LYO10255 . tel-04144064

HAL Id: tel-04144064

<https://theses.hal.science/tel-04144064>

Submitted on 28 Jun 2023

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

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



Distributed under a Creative Commons Attribution - NonCommercial - NoDerivatives 4.0
International License

N° d'ordre : 255-2002

UNIVERSITÉ CLAUDE BERNARD - ÉCOLE CENTRALE DE LYON
Laboratoire de Mécanique des Fluides et d'Acoustique

Mémoire de thèse
pour l'obtention du grade de
Docteur de l'Université Claude Bernard
Spécialité : Mécanique
au titre de l'École doctorale de : MÉGA

présentée et soutenue publiquement le 20 décembre 2002
par Antoine MOREAU

Étude du mélange de scalaires en écoulement turbulent et application à la modélisation des petites échelles

Après avis de : M. Jean-Pierre CHOLLET
M. Luc VERVISCH

Devant la commission d'examen formée de :
M. Jean-Pierre BERTOGLIO
M. Jean-Pierre CHOLLET
M. Jean-Noël GENCE
M. Emmanuel LÉVÊQUE
M. Luc VERVISCH

Avant propos

Vous tenez entre vos mains ma thèse, fruit d'un travail de trois bonnes années.

J'ai été honoré de travailler durant cette période sous la direction de Jean-Pierre Bertoglio, qui a su m'accorder une grande confiance et une très grande liberté. Il a su douter de moi aussi quand il le fallait. Je pense qu'il l'aura fait jusqu'au bout, et jamais sans raisons !

Je voudrais le remercier pour avoir su tempérer mes excès d'enthousiasme, mon impatience et mes jugements hâtifs - travers communs sans doute à beaucoup de jeunes doctorants. J'espère qu'il m'a transmis un peu de son immense patience - elle fut mise à rude épreuve.

Je voudrais aussi le remercier pour son sens critique développé, sa rigueur, son souci de l'exactitude et du travail bien fait. Autant de choses qui font qu'aujourd'hui j'ai le sentiment de présenter un travail presque terminé - ce que je n'avais jamais espéré.

Jean-Noël Gence a accepté de présider le jury constitué pour juger mon travail. J'ai beaucoup apprécié ses cours lors de mes études. Ses synthèses bibliographiques d'une clarté légendaire sont à la base de ce travail. L'influence de ses cours sur ce manuscrit est indéniable. Je voudrais aussi le remercier pour sa disponibilité et sa gentillesse.

Les professeurs Jean-Pierre Chollet et Luc Vervisch ont accepté de rapporter mon travail et de m'écouter le présenter. Je leur suis reconnaissant d'avoir joué le rôle de regard extérieur et ceux qui me connaissent savent combien l'avis qu'ils ont rendu m'importe.

Je suis très heureux qu'Emmanuel Lévêque ait bien voulu participer à mon jury. Il représente un peu la caution de l'école qui m'a formé - un rôle d'une grande importance à mes yeux.

Compagnon d'infortune, thésard non-fumeur ou presque, hôte exceptionnellement accueillant, pronostiqueur sportif hors pair, confident, inventeur inspiré de noms pour sites de vulgarisation, amateur de plantes vertes et de nouveau flamenco, grand envoyeur de manuscrits et de transformateurs par voie postale, compagnon de virées en roller, il est un camarade sans lequel ma thèse n'aurait pas été la même. Samuel Vaux a partagé son bureau avec moi durant trois ans. Il faut souligner ici sa capacité à transformer ledit bureau en lieu de rendez-vous fréquenté, sans doute à mettre au crédit de sa grande popularité. Il était écrit que jamais Samuel ne me refuserait un service. Son aide me fut toujours précieuse, et pour cela il mérite toute ma gratitude - et ces compliments pleins d'emphase. Avec un collègue de bureau, on partage souvent plus qu'un bureau, et c'est ainsi qu'il a gagné mon amitié.

Ce travail est un prolongement de celui de Marc Elmo. Nous avons eu l'occasion de nous cotoyer un certain temps, et je lui sais gré de la confiance et du respect qu'il m'a très vite accordés. Je suis heureux d'avoir pu "donner une suite à [son] travail", qu'il sache que celui-ci a été pour

moi une référence constante.

Cette thèse doit beaucoup à Olivier Teytaud. Il m'a initié au domaine des réseaux de neurones, dans lequel il est un expert, et m'a fourni ses codes. Il a su être présent pour me guider quand il le fallait vraiment, et je garderai un excellent souvenir des discussions terriblement constructives que nous avons eues.

Shao Liang a écrit en très grande partie le code pseudo-spectral sans lequel je n'aurais rien pu faire. Je voudrais le remercier ici pour sa disponibilité et l'aide qu'il m'a accordée.

Les discussions que j'ai eues avec Wouter Bos à propos de simulations RANS ont été fructueuses, ce travail en témoigne. Qu'il en soit ici remercié.

Merci à Didier Felbacq pour nos échanges, ses avis, son soutien et la jolie propriété sur les distributions β .

Pierre Borgnat et Alice Bonhomme furent des relecteurs courageux de ce manuscrit, qui ont eu à supporter les fautes de frappe, d'orthographe et les approximations qui émaillaient les premières versions. Leur amitié m'est précieuse.

J'ai une pensée pour toutes les personnes que j'ai cotoyées pendant ces trois années au LMFA. Je pense aux permanents qui m'ont accordé depuis le début une attention et un respect qui m'ont touché, et beaucoup bien sûr aux thésards (aux discussions parfois si délicieusement éloignées de tout sujet trop sérieux). L'avantage de ne citer personne explicitement bien que l'envie m'en dérange, c'est qu'ils se reconnaîtront à la lecture de ce passage et qu'ainsi je n'en aurai oublié aucun.

Je voudrais pour finir remercier mes parents, littéraires de formation, de m'avoir laissé faire des études scientifiques - et leur rappeler ce jour où j'ai ramené ma première note du collège. Il s'agissait d'une note de sciences physiques, et je crois qu'à cet instant, ils auraient eu beaucoup de mal à croire qu'un jour ils pourraient feuilleter ce manuscrit. J'ai également une pensée pour ma petite soeur, Anne, qui s'est engagée dans un doctorat d'archéologie - et je crois en elle.

Je dédie ce manuscrit à Claire, Lucie et Quentin. Durant les trois années qui m'ont été nécessaires à l'élaboration de ce travail, je me suis aussi construit une famille. Les choses se sont ainsi déroulées que Quentin est né une semaine jour pour jour avant ma soutenance. Sans le soutien, la compréhension et la douceur de Claire, rien n'aurait été possible.

Sommaire

Avant propos	3
Introduction	7
1 Étude du scalaire passif	11
1.1 Mélange du scalaire passif	11
1.1.1 Le scalaire	11
1.1.2 Écoulement turbulent	12
1.1.3 Homogénéisation du scalaire	13
1.1.4 Le mélange turbulent	14
1.1.5 Étude du scalaire passif	14
1.1.6 Milieux réactifs et scalaires conservés	16
1.2 Simulation numérique directe	18
1.3 Description du mélange par densité de probabilité	19
1.3.1 Densité de probabilité du scalaire	20
1.3.2 Fermeture de l'équation de la densité de probabilité	21
1.4 Simulation des Grandes Échelles	22
1.4.1 Définition des grandes échelles	23
1.4.2 Équations d'évolution	25
1.4.3 Simulation de réactions chimiques et fermeture	26
1.4.4 Densité de sous-maille	30
2 Caractérisation du mélange	35
2.1 Simulations numériques directes	35
2.1.1 Méthodes d'injection	35
2.1.2 Caractérisation des simulations	37
2.1.3 Étude théorique de la moyenne temporelle de la variance	38
2.1.4 Résultats et conclusion à propos de la variance	44
2.2 Étude de l'incrément du scalaire	46
2.2.1 Fonctions de structure	46
2.2.2 Relaxation de la densité de probabilité	48
2.3 Diffusion conditionnelle	52
2.3.1 Diffusion conditionnelle affine	52
2.3.2 Résultats numériques	53
2.4 Conclusion	53

3	Analyse des modèles de sous-maille	59
3.1	Estimateur optimal	59
3.1.1	Erreur quadratique	59
3.1.2	Estimateur optimal pour $\overline{f(c)}$	60
3.1.3	Propriétés des estimateurs optimaux	61
3.1.4	Estimateur optimal de la densité de sous-maille	63
3.2	Modélisation par une loi β	65
3.2.1	Comparaisons directes	66
3.2.2	Erreurs comparées	71
3.2.3	Conclusions	75
3.3	Estimation de la variance de sous-maille	78
3.3.1	Sous-modèle de Cook et Riley	78
3.3.2	Recherche d'une meilleure approximation analytique	79
3.3.3	La dissipation du champ filtré comme paramètre du modèle	81
3.3.4	Comparaison des différents estimateurs	83
3.4	Réseaux de neurones	83
3.4.1	Présentation d'un réseau de neurones	86
3.4.2	Erreur et apprentissage	87
3.4.3	Quelques résultats	88
3.4.4	Vers une simulation idéale des grandes échelles	89
	Conclusion	95
	Annexe A : Equations d'évolution	99
	Annexe B : Méthode pseudo-spectrale	105
	Annexe C : Distributions β	109
	Annexe D : Méthode de calcul des estimateurs optimaux	113

Introduction

Un champ scalaire est une représentation mathématique permettant de décrire la distribution de la température ou la répartition d'une espèce chimique dans l'espace. La grande majorité des fluides (gaz ou liquides) qui peuplent notre quotidien sont turbulents - c'est à dire soumis à des mouvements en apparence désordonnés. C'est le cas de l'air qui nous entoure (presque toujours), et de l'eau qui s'écoule (très souvent). Il est ainsi aisé de comprendre l'importance que peut revêtir le problème du champ scalaire (qu'il décrive indifféremment un champ de température ou la concentration d'une espèce chimique) entraîné par un fluide turbulent. C'est une situation que l'on rencontre dès lors que l'on veut étudier la dispersion d'un polluant dans l'atmosphère ou dans les océans, ou encore la combustion d'un carburant dans un moteur à explosion ou dans un réacteur.

Au dix-neuvième siècle, la remarque a été émise que sans la turbulence, qui dissipe l'énergie et freine l'écoulement d'un liquide, les fleuves couleraient sans doute à des vitesses très importantes. La turbulence se caractérise en effet par la division de tourbillons en tourbillons de plus en plus petits, jusqu'à ce qu'ils disparaissent, et avec eux une bonne partie de l'énergie cinétique du fluide¹. Il serait bon d'ajouter que sans la turbulence, diluer un sucre dans un café prendrait beaucoup de temps et qu'arriver à faire couler un bain de température vraiment homogène ressemblerait à un exploit. Car la turbulence mélange très efficacement : les innombrables tourbillons de toutes tailles qu'elle engendre permettent au champ scalaire de s'homogénéiser très rapidement. C'est la turbulence qui, dans un moteur à explosion mêle intimement air et carburant - et plus les deux sont mélangés plus le moteur est efficace. Le mélange du scalaire par la turbulence est ainsi un phénomène d'une très grande importance pratique.

Dans un moteur justement, des réactions interviennent qui mettent en jeu de nombreuses espèces chimiques. Elles libèrent beaucoup d'énergie, modifient de façon importante la température, et ont une grande influence sur l'écoulement. Le scalaire **passif** est un problème théoriquement plus simple dans la mesure où le scalaire considéré n'influe pas sur l'écoulement. Cette situation correspond à celle d'un polluant dans l'atmosphère, d'un colorant dans un liquide ou d'un champ de température présentant des écarts de quelques degrés seulement. De sa compréhension dépend bien sûr la compréhension et la modélisation de phénomènes aussi complexes que la combustion turbulente. Cette thèse se concentre sur le cas du scalaire passif.

Le principe même de la démarche scientifique est de diviser les problèmes complexes en problèmes simples que l'on espère pouvoir traiter. La turbulence est un problème très difficile sans doute parce que ce processus ne peut pas vraiment lui être appliqué. Ou plutôt, il l'a déjà été : nous connaissons actuellement parfaitement les équations qui permettent de décrire le comportement des fluides. Mais, et c'est bien là le drame, ces équations ne suffisent pas pour comprendre

¹Cette vision des choses est sans doute - malgré sa simplicité apparente - la plus aboutie dont nous disposions à l'heure actuelle. Nous la devons à Kolmogorov[1], et elle date du milieu du siècle dernier.

en détail le phénomène de la turbulence ni pour prévoir quelles en seront toutes les caractéristiques. Les théories dont nous disposons actuellement pour décrire la turbulence sont d'origine essentiellement empirique - une approche déductive systématique prenant pour point de départ les équations de l'hydrodynamique manque encore en turbulence.

Le scalaire passif a de prime abord été jugé plus simple à étudier que la turbulence. Il présente cependant des comportements que l'on a en définitive plus de mal à expliquer et les méthodes théoriques classiques permettant de décrire la turbulence se révèlent parfois inadéquates pour le scalaire. La turbulence et le scalaire passif en turbulence sont ainsi des problèmes pour lesquels on connaît les équations fondamentales, mais qui restent très compliqués et d'une importance pratique énorme.

Pour décrire la turbulence et le scalaire passif, la plupart des outils théoriques développés pour la physique ont été utilisés (renormalisation et géométrie fractale par exemple). De nombreux travaux théoriques de très haut niveau ont récemment été menés sur le problème du scalaire passif plus précisément. Ils permettent de progresser dans la compréhension du mélange turbulent, mais ont en général peu, voire pas, d'implications pratiques.

Avec l'avènement des ordinateurs, de nouvelles voies d'approche des problèmes de turbulence et de mélange turbulent se sont ouvertes. Les simulations numériques directes ont permis de recréer des écoulements virtuels en utilisant les équations fondamentales de la mécanique des fluides. Ces écoulements constituent des objets d'étude extrêmement dociles, car leurs caractéristiques sont modifiables à l'envi. Il n'est de plus pas besoin d'instruments de mesure complexes pour en déterminer toutes les caractéristiques - dont beaucoup restent par ailleurs inaccessibles à l'expérience. L'ensemble de cette thèse repose sur de telles simulations, qui présentent cependant un défaut important : elles sont tellement gourmandes en capacités de calcul qu'elles ne permettent de simuler que de "petits" écoulements. Les outils informatiques offrent cependant des perspectives beaucoup plus larges que les seules simulations directes. Avant que ces outils n'existent, un modèle ne pouvait avoir de portée que tant que les équations auxquelles il aboutissait pouvaient être résolues analytiquement. Il était alors possible de conclure quant à sa qualité et à la validité des hypothèses qui le sous-tendaient. Les outils numériques permettent de s'affranchir de la contrainte de solvabilité, car ils permettent de résoudre - au moins de façon approchée - n'importe quel système d'équations, et de tester ainsi les conséquences de n'importe quelle modélisation. En pratique, on cherche à reproduire numériquement certains comportements de la turbulence ou du scalaire - en faisant appel à des modèles. Cette thèse traite par exemple d'un type de simulation qui cherche à reproduire les grandes structures de la turbulence seules. L'effet des petites échelles (des détails en somme) sur les grandes est alors modélisé.

L'étude des simulations numériques directes permet d'une part de vérifier ou d'affiner certaines prédictions théoriques - à propos du mélange du scalaire passif par exemple. Elle permet d'autre part de confirmer ou d'infirmer certaines hypothèses qui sous-tendent les modèles utilisés dans les simulations numériques à visées pratiques - concernant notamment les espèces chimiques réactives. C'est bien là l'objet de cette thèse.

Ce manuscrit est divisé en trois chapitres.

Le premier introduit les notions indispensables à la suite de l'exposé. La problématique du mélange turbulent du scalaire passif et son rapport avec les modèles de milieux réactifs sont

présentés. Cette présentation est suivie de l'évocation des principaux résultats théoriques actuels, ainsi que des résultats expérimentaux disponibles à propos du scalaire passif. Cet exposé précède une présentation de l'outil numérique à la base de tous les résultats présentés dans cette thèse : la simulation numérique directe. La description du mélange du scalaire par densités de probabilité est ensuite abordée, en évoquant son application aux simulations de type RANS ². Pour clore ce chapitre, la Simulation des Grandes Échelles est introduite, ainsi que la problématique de la modélisation de sous-maille - en particulier, un modèle est présenté, qui est basé sur l'estimation d'une "densité de sous-maille". Certaines propriétés de cette quantité sont démontrées à cette occasion, qui permettent de mieux saisir sa signification.

Le deuxième chapitre est une étude du mélange du scalaire passif, basée sur une situation stationnaire obtenue grâce à une méthode d'injection proposée dans un travail antérieur [2] par Marc Elmo - il faut d'ailleurs préciser que cette thèse est un prolongement du travail de ce dernier. La situation de mélange est caractérisée en terme de spectres et le rapport entre le degré de mélange obtenu et les caractéristiques de l'injection est étudié. L'étude des fonctions de structure permet ensuite de caractériser l'intermittence du scalaire - et de comparer les résultats obtenus avec les résultats expérimentaux disponibles. Pour finir, le comportement de la diffusion conditionnelle est abordé, car il s'agit d'une quantité centrale dans la description du mélange par densité de probabilité.

En Simulation des Grandes Échelles, le rôle des petites échelles est pris en compte à travers une modélisation dite de "sous-maille" de certaines quantités apparaissant dans les équations d'évolution. En cherchant à mesurer l'écart entre l'estimateur et la quantité à modéliser, il est possible de définir des estimateurs optimaux pour une situation précise. Le troisième chapitre permet d'introduire ces estimateurs et de proposer une série de propriétés les concernant. Ces notions sont ensuite appliquées à l'étude d'un modèle introduit au premier chapitre, le modèle par densité de sous-maille de Cook et Riley. Elles permettent de vérifier l'intérêt des distributions β . La "variance de sous-maille" apparaissant alors comme une quantité centrale, les estimateurs de cette quantité sont étudiés, et voies sont explorées pour les améliorer. La recherche de la plus petite erreur possible conduit pour finir vers l'introduction de réseaux de neurones. Ils sont utilisés ici pour évaluer les estimateurs optimaux lorsque le nombre de paramètres de base s'accroît. Leur application directe à la simulation des grandes échelles est envisagée.

Quatre annexes viennent apporter quelques compléments. Ces compléments sont d'ordre mathématique pour l'annexe A, qui concerne les équations d'évolution des quantités considérées dans ce travail, et pour l'annexe C, qui se concentre sur les propriétés des distributions β . Ils sont d'ordre plus technique pour ce qui est de l'annexe B, qui expose en détail le principe des simulations numériques directes et de l'annexe D, qui expose une technique de calcul d'estimateurs optimaux.

²Pour Reynolds Averaged Navier-Stokes

Chapitre 1

Étude du scalaire passif

Ce chapitre vise à introduire les notions et les notations qui seront abondamment utilisées par la suite. Il a pour but de donner au lecteur une idée assez générale de ce qu'est l'étude du scalaire passif, des concepts et des théories qui lui sont liés, et des méthodes auxquelles elle fait appel. L'ensemble des notions évoquées dans le titre de ma thèse est ici abordé, ce dernier faisant presque office de plan du chapitre : "Étude du mélange de scalaires en écoulement turbulent et application à la modélisation des petites échelles". J'ai supposé le lecteur au courant des notions de base de la mécanique des fluides et de la turbulence, et ignorant tout des spécificités du scalaire.

1.1 Mélange du scalaire passif

1.1.1 Le scalaire

On décrit par un champ scalaire $c(\vec{x})$ indifféremment la concentration d'une espèce chimique ou la répartition de la température. Ces quantités ont dans un fluide le même comportement : elles sont transportées par le fluide en mouvement, et elles diffusent à l'intérieur de celui-ci. L'étude du comportement de ces quantités se ramène donc à celle d'un champ scalaire c obéissant à une équation d'évolution

$$\partial_t c + \partial_i(v_i c) = D \Delta c, \quad (1.1)$$

où v_i est une composante de la vitesse.

L'équation (1.1) peut aussi s'écrire

$$\frac{dc}{dt} = D \Delta c,$$

où $\frac{d}{dt}$ désigne la dérivée lagrangienne.

Si on suit un point du fluide dans son mouvement, la variation de c associée à ce point ne dépend ainsi que du laplacien Δc . Or, quand ce point est un maximum (respectivement un minimum) local du champ scalaire, le laplacien est négatif (respectivement positif) et la valeur du champ scalaire ne peut ensuite que décroître (respectivement croître). Si le champ scalaire est majoré par c_{max} (respectivement minoré par c_{min}) à un instant donné, ce majorant (resp. minorant) reste valable à tout instant ultérieur. Or, dans presque toutes les situations physiques, le scalaire est borné ($c \in [c_{min}, c_{max}]$) dès les conditions initiales. Ces bornes restent donc valables à tout

instant ultérieur¹. De plus, l'équation (1.5) est linéaire en c . Le champ c peut donc être multiplié par une constante ou bien une constante peut lui être ajoutée, sans qu'aucune de ces opérations ne modifie l'équation d'évolution. Ainsi, si on imagine deux champs tels qu'une transformation affine permette de passer de l'un à l'autre en tout point à l'instant initial, alors cela reste vrai à tout instant ultérieur.

En conclusion, l'étude de tout scalaire borné peut se ramener à l'étude d'un champ scalaire dans un intervalle de travail arbitraire. Une transformation affine permet redéfinir cet intervalle et de retrouver celui correspondant au problème physique étudié.

1.1.2 Écoulement turbulent

Le scalaire est dit **passif** si sa présence n'influence pas l'écoulement. Dans cette étude, nous nous plaçons dans le cas d'un fluide incompressible de densité volumique ρ constante. Le mouvement fluide est alors décrit par son champ de vitesse $\vec{v}(\vec{x})$. Dans ce cadre, les équations qui gouvernent le comportement du fluide sont celles de l'incompressibilité et de Navier-Stokes

$$\partial_i v_i = 0 \quad (1.2)$$

$$\partial_t v_i + \partial_j (v_j v_i) = -\partial_i \left(\frac{p}{\rho} \right) + \nu \Delta v_i. \quad (1.3)$$

L'équation (1.2) peut être remplacée par celle qui permet de déterminer la pression à partir du champ de vitesse

$$\Delta \left(\frac{p}{\rho} \right) = -\partial_i v_j \partial_j v_i, \quad (1.4)$$

obtenue en prenant la divergence de l'équation 1.3 et qui contient donc de façon implicite l'hypothèse d'un écoulement incompressible.

À grand nombre de Reynolds, le fluide présente un comportement turbulent. C'est dans ces écoulements que l'étude du scalaire présente un grand intérêt, car la turbulence facilite le mélange.

Dans la suite de ce chapitre, nous allons nous placer dans le cas des conditions aux limites périodiques, c'est à dire dans un cube de taille L . Cette situation académique représente ce qui se passe "loin des bords" de l'écoulement, là où la turbulence est la plus homogène et isotrope possible, et où d'après Kolmogorov[3] elle a un comportement universel : les plus grosses structures sont divisées en plus petites, qui sont à leur tour divisées sans perdre d'énergie. Les plus petites structures de l'écoulement sont ensuite dissipées par la viscosité. L'énergie cinétique du fluide est donc dissipée à petite échelle.

L'échelle en dessous de laquelle on ne trouve plus de fluctuation de vitesse s'appelle **l'échelle de Kolmogorov** - elle signe l'arrêt par la viscosité de la cascade turbulente des grandes vers les petites échelles. Cette échelle s'écrit

$$\eta = \left(\frac{\nu^3}{\langle \epsilon \rangle} \right)^{\frac{1}{4}},$$

où $\epsilon = 2\nu s_{ij} s_{ij}$ (cette quantité est appelée la **dissipation**) avec $s_{ij} = \frac{1}{2}(\partial_i v_j + \partial_j v_i)$.

¹Les bornes initiales restent valables, mais la valeur maximale (respectivement minimale) prise par le champ c ne reste en général pas constante. Elle diminue (respectivement augmente) au cours du temps. Ainsi il est en général possible de trouver un encadrement plus fin du scalaire aux instants ultérieurs.

1.1.3 Homogénéisation du scalaire

Moyenne du scalaire

Nous définirons la moyenne du champ scalaire dans le domaine cubique de l'écoulement par

$$\mathcal{M} = \frac{1}{L^3} \int c(\vec{x}) \, d\vec{x}.$$

Cette quantité ne varie pas au cours du temps. Elle ne dépend donc que de la situation initiale considérée. En effet,

$$\begin{aligned} \frac{d\mathcal{M}}{dt} &= \int \partial_t c \, d\vec{x} \\ &= \int \partial_i (D \partial_i c - v_i c) \, d\vec{x} \\ &= \int (D \partial_i c - v_i c) \, d\vec{S}_i. \end{aligned}$$

Les conditions aux limites périodiques font que le flux à travers une face du cube est l'opposé de celui de la face opposée. Il s'agit en effet du même champ, mais pour des surfaces orientées à l'opposé l'une de l'autre.

Variance, énergie et dissipation du scalaire

Étudions maintenant la variance du scalaire sur notre domaine d'étude, qui s'écrit

$$\begin{aligned} \sigma^2 &= \frac{1}{L^3} \int (c - \mathcal{M})^2 \, d\vec{x} \\ &= \frac{1}{L^3} \int c^2 \, d\vec{x} - \mathcal{M}^2 \end{aligned}$$

Cette quantité donne une mesure de l'écart du champ scalaire à la moyenne, et fait intervenir la quantité c^2 . Considérons l'équation d'évolution de celle-ci :

$$\partial_t c^2 + \partial_i (v_i c^2) = D \Delta c^2 - 2D (\partial_i c)^2. \quad (1.5)$$

La quantité c^2 n'est donc pas conservée, mais dissipée, car $2D(\partial_i c)^2$, appelée la dissipation du scalaire² est toujours positive. L'évolution de la variance est alors donnée par

$$\begin{aligned} \frac{d\sigma^2}{dt} &= \int \partial_t c^2 \, d\vec{x} \\ &= - \int \epsilon_c \, d\vec{x}, \end{aligned}$$

les termes de flux disparaissant comme précédemment. On constate que la variance diminue tant que la dissipation ne s'annule pas, c'est à dire tant que le champ scalaire n'est pas uniforme et égal à $\mathcal{M}$ en tout point. C'est ce processus d'homogénéisation du champ scalaire que l'on appelle naturellement le mélange.

En résumé, un champ scalaire borné soumis à l'équation (1.1) évolue en conservant les mêmes bornes³ vers une situation d'équilibre où le champ est uniforme.

²Par analogie avec l'énergie cinétique, c^2 est parfois appelée énergie du scalaire. Certains auteurs considèrent cependant la quantité $\frac{1}{2} c^2$, toujours par analogie, et appellent dissipation du scalaire la quantité $D(\partial_i c)^2$.

³Au sens où l'encadrement initial n'est pas violé, mais n'est pas non plus le meilleur.

1.1.4 Le mélange turbulent

Le mélange est d'autant plus rapide qu'il existe des zones de fort gradient. Dans ces zones, la dissipation est forte et les flux de scalaire dûs à la diffusion sont importants. C'est toujours la diffusion qui permet le mélange. Si on prend $D = 0$, la variance du scalaire est constante.

L'influence du champ de vitesse sur le mélange est indirecte. Pour mieux la comprendre, il faut considérer l'équation d'évolution de la dissipation (établie en annexe A)

$$\partial_t \epsilon_c + \partial_i (v_i \epsilon_c) = D \Delta \epsilon_c - 4D^2 (\partial_j \partial_i c)^2 - 2D \partial_i c \partial_i v_j \partial_j c. \quad (1.6)$$

La dissipation n'est bien entendu pas une quantité conservée, et pour comprendre comment elle se crée et elle disparaît, il faut considérer les deux termes de source de l'équation précédente. Le premier, $-4D^2 (\partial_j \partial_i c)^2$, est forcément négatif. Il signifie que dans les zones de fort gradient où la dissipation est importante, comme la diffusion a tendance à faire disparaître les gradients, la dissipation elle-même disparaît lors de l'homogénéisation. Le second, $-2D \partial_i c \partial_i v_j \partial_j c$, fait explicitement intervenir la vitesse. Il représente la création de dissipation dans les zones de fort étirement (grandes valeurs de $\partial_i v_j$). La turbulence, en étirant le fluide et le champ de scalaire, crée des gradients importants et favorise ainsi l'homogénéisation du scalaire par la diffusion.

La turbulence agit donc en étirant les grandes structures du scalaire, et ce faisant, crée une dissipation importante. Elle accélère donc l'homogénéisation du scalaire : c'est ce qu'on appelle **le mélange turbulent**.

Le rapport entre la viscosité (coefficient de diffusion de la quantité de mouvement) et le coefficient de diffusion du scalaire est essentiel. Ce rapport est sans dimension et appelé le nombre de Schmidt⁴

$$Sc = \frac{\nu}{D}.$$

Il contrôle directement le rapport entre l'échelle de Kolmogorov et l'échelle de Batchelor. Cette échelle est celle en dessous de laquelle le scalaire ne présente plus de fluctuations notables, et elle est définie par

$$\eta_b = \frac{\eta}{\sqrt{Sc}}.$$

Dans ce travail, l'ordre de grandeur du nombre de Schmidt est de 1.

1.1.5 Étude du scalaire passif

Le scalaire passif a été abondamment étudié [4]. Mais si l'on a pensé initialement que son abord ne devait pas présenter trop de difficultés - il obéit en effet à une équation linéaire - il n'en est plus de même actuellement.

Approche de Corrsin et Obukhov

L'approche plus classique du scalaire passif est conceptuellement et chronologiquement très proche de l'approche par Kolmogorov de la turbulence. Elle a été proposée par Corrsin [5] et Obukhov [6], qui voient dans le mélange turbulent un phénomène de cascade. De leur point de vue, le mélange turbulent est dû à la division des structures du champ scalaire en structures plus petites, jusqu'à une échelle où ces structures peuvent être dissipées par la diffusion moléculaire.

⁴Aussi appelé Prandtl dans le cas particulier où le scalaire étudié est la température.

En suivant un raisonnement analogue à celui de Kolmogorov ils prédisent une loi de puissance pour le spectre dans la zone inertielle du champ de vitesse,

$$E_c(k) = \kappa \langle \epsilon \rangle^{-\frac{1}{3}} \langle \epsilon_c \rangle k^{-\frac{5}{3}}. \quad (1.7)$$

Cette loi de puissance a déjà été observée, y compris pour des valeurs de R_λ (le nombre de Reynolds basé sur l'échelle de Taylor λ) faibles en turbulence de grille [7]. Ceci alors que le champ de vitesse ne présentait pas encore la zone inertielle que lui prédit la théorie de Kolmogorov. Inversement, pour des écoulements qui ne sont pas aussi nettement homogènes et isotropes, le comportement ci-dessus ne semble se présenter que pour des grandes valeurs de R_λ [8].

Une loi [9, 10] analogue à la loi de Kolmogorov des quatre cinquièmes[11] est vérifiée en zone inertielle :

$$\langle [\Delta u(r)][\Delta c(r)]^2 \rangle = -\frac{4}{3} \frac{\langle \epsilon_c \rangle}{2} r, \quad (1.8)$$

où $\Delta c(r) = c(x+r) - c(x)$ et $\Delta u(r) = u(x+r) - u(x)$. Cette loi ne semble se manifester que pour de très grands nombres de Reynolds dans l'atmosphère ou en soufflerie.

Dans le cas des grands nombres de Schmidt, le spectre du scalaire s'étend au delà de l'échelle de Kolmogorov, jusqu'à l'échelle de Batchelor qui est dans ce cas nettement plus petite. À ces échelles, le scalaire n'est plus soumis au même type de mélange, et Batchelor prédit un spectre du scalaire présentant un comportement en k^{-1} [12], qui a déjà été observé[13].

Intermittence et modèles de l'advection turbulente

Le spectre ne suffit cependant pas à décrire toutes les caractéristiques du scalaire dans un écoulement turbulent. Si on considérait par exemple un champ présentant un spectre en $k^{-\frac{5}{3}}$ et des phases aléatoires indépendantes et identiquement distribuées, alors la densité de probabilité de $\Delta c(r)$ serait une gaussienne. Cela correspond à un comportement des moments de $\Delta c(r)$ prédit par la théorie d'Obukhov-Corrsin grâce à des arguments dimensionnels :

$$\langle [\Delta c(r)]^n \rangle \sim r^{\frac{n}{3}}. \quad (1.9)$$

Ce comportement n'est pas du tout observé, et en pratique l'exposant des fonctions de structure croît beaucoup moins vite que $\frac{n}{3}$. Cela signifie que pour des distances assez petites, Δc prend plus souvent que dans le cas gaussien de grandes valeurs. Cette apparition d'événements violents moins rares qu'attendu est appelée **intermittence**, et il est bien établi que le champ de vitesse turbulent et le champ de scalaire sont intermittents [7].

Le moment d'ordre deux de l'accroissement de scalaire est théoriquement lié [3] à la loi en $-\frac{5}{3}$. Ceci explique que pour $n = 2$, le comportement donné par (1.9) est bien observé - tellement bien même qu'on observe parfois un comportement en $r^{\frac{2}{5}}$ alors même que le spectre ne présente pas de loi de puissance en $-\frac{5}{3}$!

On pouvait supposer que l'intermittence du scalaire était liée à celle du champ de vitesse. Kraichnan [14] a proposé de modéliser le champ de vitesse turbulent par un champ gaussien (ne présentant donc pas d'intermittence) et décorrélié en temps. Ce modèle a ceci d'attrayant qu'il permet de donner des expressions exactes de certaines grandeurs [15]. Les simulations numériques qui ont été menées avec ce modèle [16] ont fait apparaître une intermittence importante du scalaire, alors que le champ de vitesse ne l'était pas.

Anisotropie du scalaire

L'échec relatif de la théorie d'Obukhov-Corrsin du mélange turbulent (c'est-à-dire ailleurs qu'en turbulence de grille) s'explique sans doute par l'anisotropie du scalaire à petite échelle. En effet, leur théorie est basée sur l'hypothèse qu'à petite échelle, les propriétés des champs de vitesse et de scalaire sont universelles et indépendantes des conditions aux limites. Or il semble que le scalaire connaisse une relaxation lente vers l'isotropie, c'est-à-dire que même aux temps longs et à petite échelle, il soit encore marqué par la façon dont il a été injecté. C'est en contradiction avec l'idée que les petites échelles du scalaire sont universelles.

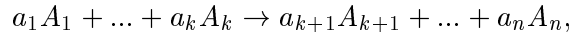
Pour conclure, il est frappant de voir combien le problème du scalaire passif, malgré une formulation simple, peut présenter des caractéristiques complexes et difficiles à aborder : une intermittence qui émerge même lorsque le champ de vitesse est gaussien et une anisotropie avérée à petite échelle.

1.1.6 Milieux réactifs et scalaires conservés

Le problème du scalaire conservé (obéissant à l'équation (1.1) sans aucun terme de source) joue paradoxalement un rôle essentiel dans les situations de milieux réactifs où les scalaires que sont les concentrations des espèces chimiques ne sont pas conservés [17]. Cette partie illustre dans quel contexte on peut se ramener à l'étude de scalaires conservés. Il n'est en général pas possible de se ramener seulement à des scalaires conservés et c'est ce qui rend les simulations difficiles.

Scalaires conservés en milieu réactif

Plaçons nous dans le cas où le milieu est réactif, c'est à dire qu'il contient des espèces chimiques A_i , qui réagissent suivant l'équation bilan



où a_i est le coefficient stœchiométrique associé à l'espèce A_i . La concentration c_{A_i} de chaque espèce obéit à une équation d'évolution reflétant la convection de l'espèce par le fluide, sa diffusion moléculaire et la réaction chimique par l'intermédiaire d'un taux de réaction σ_{A_i} :

$$\partial_t c_{A_i} + \partial_j (c_{A_i} v_j) = D_{A_i} \Delta c_{A_i} + \sigma_{A_i}. \quad (1.10)$$

L'équation bilan impose une relation entre les taux de réaction des différentes espèces, donnée par

$$\forall i, j \quad \frac{\sigma_{A_i}}{\alpha_i} = \frac{\sigma_{A_j}}{\alpha_j}, \quad (1.11)$$

où le coefficient α_i est tel que $\alpha_i = a_i$ si A_i est un réactif et $\alpha_i = -a_i$ si A_i est un produit.

En milieu gazeux par exemple, les coefficients de diffusion des différentes espèces sont tous identiques $D_{A_i} = D$. Dans ce cas, les quantités $s_{ij} = \frac{c_{A_i}}{\alpha_i} - \frac{c_{A_j}}{\alpha_j}$ obéissent alors à une équation d'advection-diffusion

$$\partial_t s_{ij} + \partial_k (s_{ij} v_k) = D \Delta s_{ij}. \quad (1.12)$$

On peut construire $\frac{n(n-1)}{2}$ de ces scalaires conservés, mais ils ne sont pas indépendants car $s_{ij} = s_{kj} - s_{ki}$. Cela signifie qu'avec tous les s_{kj} à k fixé et à $j \neq k$, on peut calculer tous les

scalaires conservés. Il y a donc seulement $n - 1$ scalaires conservés indépendants⁵, et cela signifie qu'il faut une indication sur la vitesse de réaction pour pouvoir calculer l'évolution de toutes les concentrations.

Dans le cas d'une réaction pour laquelle le taux de réaction en fonction des concentrations des différentes espèces est connu, il faut conserver l'équation d'évolution d'une des espèces chimiques avec ce taux de réaction. On peut par exemple choisir pour cela une espèce dont la concentration contrôle seule le taux de réaction (cas particulier de cinétique chimique). Elle obéit à une équation qui peut s'écrire [18]

$$\partial_t c + \partial_i (v_i c) = D \Delta c + f(c) \quad (1.13)$$

Dans bien des cas, des hypothèses sur le taux de réaction (infiniment rapide, négligeable partout sauf sur le front de flamme par exemple) permettent de ne pas avoir à résoudre cette équation lors d'une simulation.

Réaction infiniment rapide

Une hypothèse courante lors de simulations de milieux réactifs est de supposer que la réaction considérée se fait de façon instantanée [17, 19]. Cela permet de se ramener à l'étude de scalaires conservés. En revanche, les relations qui lient les scalaires conservés aux concentrations réelles ne sont pas linéaires, mais continues par morceaux.

Considérons la réaction décrite par l'équation bilan



Le scalaire $X = c_A - r c_B$ est, d'après le paragraphe précédent, un scalaire conservé. Considérons une situation dont les réactifs ne sont pas prémélangés, mais introduits dans le milieu avec une concentration c_A^0 pour A et une concentration c_B^0 pour B . Comme on sait que A et B sont forcément consommés (leur taux de réaction est négatif ou nul), et qu'une concentration est nécessairement positive, on peut écrire que

$$\begin{aligned} 0 &\leq c_A \leq c_A^0 \\ 0 &\leq c_B \leq c_B^0, \end{aligned}$$

en suivant un raisonnement analogue à celui utilisé pour montrer qu'un scalaire conservé est borné.

On fait souvent l'hypothèse que la réaction chimique considérée est infiniment rapide. Cela signifie que A et B ne peuvent coexister sans réagir. Dans ce cas, le maximum de X est obtenu pour le maximum de c_A et son minimum pour le maximum de c_B , soit

$$-r c_B^0 \leq X \leq c_A^0.$$

Dans le cas où A est en excès, $X = c_A$ et $c_A^0 - c_A = (1 + r) c_P$, et dans le cas où B est en excès, $X = -r c_B$ et $c_B^0 - c_B = \frac{1+r}{r} c_P$. Si $X = 0$, $c_P = 1$.

⁵Certains des scalaires conservés représentent en réalité la concentration d'un élément dans le milieu, et non d'une espèce. Quand les coefficients de diffusion sont les mêmes pour toutes les espèces, cela revient à dire que les éléments diffusent de la même manière, sans considération aucune de l'espèce qui les transporte. Dans une équation bilan de réaction chimique, il y a bien $n - 1$ éléments différents - ce dont on s'aperçoit lorsqu'il faut l'équilibrer et donc choisir arbitrairement un des coefficients stœchiométriques.

On peut choisir $[0, 1]$ comme intervalle de travail en prenant pour scalaire conservé

$$\xi = \frac{c_A + r(c_B^0 - c_B)}{c_A^0 + r c_B^0}.$$

Dans ce cas, la concentration du produit de la réaction s'écrit

$$\begin{aligned} c_P &= \frac{\xi}{\xi_{st}} \quad \text{si } \xi \leq \xi_{st} \\ c_P &= \frac{1 - \xi}{1 - \xi_{st}} \quad \text{sinon,} \end{aligned}$$

avec $\xi_{st} = \frac{r c_B^0}{c_A^0 + r c_B^0}$.

Dans ce cas particulier, il est donc possible de se ramener à l'étude d'un seul scalaire conservé pour décrire une réaction chimique. De façon générale, ce n'est bien entendu pas possible. Les réactions de combustion sont en effet fortement exothermique. On ne peut plus considérer que le fluide soit incompressible et il n'est plus possible de modéliser les concentrations des espèces chimiques par un scalaire passif. L'étude du scalaire passif en écoulement compressible reste en revanche une étape obligée dans la compréhension des phénomènes associés à la combustion turbulente.

1.2 Simulation numérique directe

Les équations qui permettent de modéliser le comportement des fluides dans des conditions usuelles sont parfaitement connues – ce sont celles qui ont été présentées en 1.1. La turbulence et les phénomènes qui lui sont associés, comme le mélange turbulent, sont des comportements qui sont décrits par ces mêmes équations.

Avec l'avènement des ordinateurs, il est devenu possible d'approcher numériquement des solutions des équations fondamentales. Cette approche est appelée *Simulation Numérique Directe* d'écoulements [20]. Les écoulements simulés montrent un comportement proche de la turbulence. De telles simulations sont cependant tellement gourmandes en temps de calcul et en mémoire qu'elles sont limitées à des nombres de Reynolds modestes [21]. Par contre, lors des simulations, les champs de vitesse et de scalaire sont accessibles en tout point de l'espace. Les simulations numériques directes se sont donc avérées être un outil sans équivalent pour l'étude de la turbulence et des phénomènes qui lui sont associés. L'ensemble de mon travail est basé sur de telles simulations, et c'est pourquoi j'ai choisi de décrire en détail leur principe dans l'annexe B.

Le schéma numérique adopté pour les simulations est dit "pseudo-spectral". Il s'agit d'une situation avec des conditions aux limites périodiques. De tels écoulements permettent d'obtenir des solutions qui ressemblent à petite échelle aux écoulements dans un domaine infini, tout en étant simulables numériquement. Le fait que les conditions aux limites soient périodiques permet de représenter les champs étudiés par des séries de Fourier. Il est possible d'écrire les équations qui gouvernent les champs dans l'espace de spectral. Dans le cas du champ de vitesses les deux équations nécessaires peuvent être fusionnées en une seule, ce qui présente des avantages indéniables. Cette équation s'écrit

$$\partial_t \tilde{v}_i + \left(\delta_{il} - \frac{k_i k_l}{k^2} \right) i k_j \tilde{v}_l * \tilde{v}_j + \nu k^2 \tilde{v}_i = 0, \quad (1.15)$$

où $*$ représente une convolution. Pour le scalaire passif, l'équation (1.1) devient

$$\partial_t \tilde{c} + i k_j \tilde{v}_j * \tilde{c} + D k^2 \tilde{c} = 0. \quad (1.16)$$

On sait d'autre part qu'il existe une plus haute fréquence pour la décomposition en série de Fourier du champ de vitesse - c'est la fréquence correspondant à l'échelle de Kolmogorov. Dans le cadre d'un domaine de simulation cubique avec des conditions aux limites périodiques, cela signifie qu'il est possible de représenter le champ de vitesse sur un nombre fini de modes de Fourier (voir annexe B). Dès lors, le champ de vitesse devient représentable en machine, et on peut envisager de le faire numériquement évoluer grâce aux équations d'évolution en espace spectral, d'où le nom de la méthode.

On trouve couramment que la méthode pseudo-spectrale permet de simuler un champ de vitesse défini sur les noeuds d'une grille régulière de l'espace physique. Ce point de vue vient de l'utilisation des transformées de Fourier discrètes, qui établissent une bijection entre les valeurs du champ en ces points de la grille et les amplitudes de ses modes de Fourier. En fait, en connaissant la décomposition d'un champ en modes de Fourier, on peut accéder à sa valeur en tout point de l'espace. Cette précision est importante, car le calcul de la valeur du champ scalaire en un point quelconque de l'espace a été nécessaire à l'obtention de certains résultats du chapitre 3.

1.3 Description du mélange par densité de probabilité

L'approche traditionnelle de la turbulence consiste à décomposer les champs de vitesse ou de scalaire étudiés en une composante moyenne et une composante fluctuante. C'est évidemment le comportement moyen du fluide et du scalaire qui recueille toute l'attention du physicien. Il est aisé d'écrire les équations d'évolution des champs moyens pour la vitesse et pour un scalaire obéissant à l'équation 1.13

$$\partial_i \langle v_i \rangle = 0 \quad (1.17)$$

$$\partial_t \langle v_i \rangle + \partial_j \langle v_j v_i \rangle = -\frac{1}{\rho} \partial_i \langle p \rangle + \nu \Delta \langle v_i \rangle \quad (1.18)$$

$$\partial_t \langle c \rangle + \partial_i \langle v_i c \rangle = D \Delta \langle c \rangle + \langle f(c) \rangle \quad (1.19)$$

Un certain nombre de termes ne se déduisent pas directement des champs moyens eux-mêmes parce qu'ils sont non-linéaires. L'écriture d'équations d'évolution pour ces termes fait invariablement apparaître des termes supplémentaires d'ordre plus élevé. Ce phénomène se reproduit à chaque équation d'évolution. En définitive, on dispose d'un ensemble infini d'équations, ce qu'on appelle une hiérarchie d'équations. Ceci traduit le fait qu'il faut résoudre les équations fondamentales même pour n'avoir qu'un résultat partiel sur l'écoulement, comme les champs moyens. Il est par contre possible de modéliser certains termes inconnus, de manière à n'avoir plus à prendre en compte leur équation d'évolution. Il n'y a plus à considérer alors qu'un nombre fini d'équations, sans terme inconnu. On dit que le problème a été *fermé*, qu'on a trouvé une *fermeture*. Le système d'équations ainsi obtenu est ensuite résolu numériquement. De telles simulations sont appelées RANS ⁶

De nombreux travaux ont été menés sur ces problèmes de fermeture pour les termes non-linéaires des équations fondamentales. Ces travaux ne seront pas abordés ici, car nous allons nous

⁶Pour "Reynolds Averaged Navier-Stokes" [21].

intéresser au terme de réaction chimique dans le cas du scalaire passif. Tout d'abord, les fonctions f considérées sont en général fortement non-linéaires et une information précise sur le champ scalaire est nécessaire pour estimer $\langle f(c) \rangle$ de façon fiable - qui ne se limite pas en tous les cas à un moment d'ordre deux. De plus, une méthode utile de fermeture du terme $\langle f(c) \rangle$ ne devrait pas être propre à f - afin de pouvoir être utilisée en toute généralité. Tout ceci impose de connaître la densité de probabilité $p(\Gamma)$ du champ scalaire c [22]. Cette quantité permet en effet d'en déduire $\langle f(c) \rangle$ quelle que soit f grâce à la relation

$$\langle f(c) \rangle = \int f(\Gamma) p(\Gamma) d\Gamma \quad (1.20)$$

Il est possible d'obtenir une équation d'évolution pour la densité de probabilité du scalaire⁷ dans le cas le plus général (inhomogène et anisotrope). Valiño [23] a proposé un schéma numérique original et prometteur pour la résolution de cette équation. Cette résolution nécessite cependant quelques hypothèses de fermeture, le modèle utilisé pour la diffusion conditionnelle étant le modèle LMSE proposé par Dopazo [24].

Nous allons rappeler ici quelques résultats à propos de l'approche par densité de probabilité, et détailler quelques aspects du modèle de Dopazo. Le chapitre 2 traitera en effet de ces problèmes de fermeture.

1.3.1 Densité de probabilité du scalaire

L'approche par densité de probabilité s'est avérée fournir une description intéressante du mélange turbulent. Elle avait déjà été utilisée pour le champ de vitesse [25], mais ce sont les travaux de Dopazo [24] qui ont montré l'intérêt de cette méthode pour le champ scalaire. Nous allons nous placer ici dans le cadre d'un scalaire passif éventuellement soumis à une réaction chimique simple dans un écoulement homogène et isotrope et considérer que le scalaire obéit à l'équation d'évolution 1.13.

Dans ce cadre, la densité de probabilité pour que c prenne la valeur Γ , notée $p(\Gamma)$, obéit à une équation d'évolution qui s'écrit

$$\frac{\partial p(\Gamma)}{\partial t} = -D \frac{\partial}{\partial \Gamma} (\langle \Delta c | \Gamma \rangle p(\Gamma)) + f(\Gamma) p(\Gamma) \quad (1.21)$$

ou encore

$$\frac{\partial p(\Gamma)}{\partial t} = -\frac{\partial^2}{\partial \Gamma^2} (\langle \epsilon_c | \Gamma \rangle p(\Gamma)) + f(\Gamma) p(\Gamma), \quad (1.22)$$

où $\langle . | . \rangle$ dénote une moyenne conditionnée. Les équations ci-dessus sont démontrées en annexe A, ainsi que leur expression dans le cas inhomogène, où des termes supplémentaires liés à la vitesse apparaissent.

Si on se restreint au cas sans terme de réaction, l'équation (1.22) prend alors l'allure d'une équation de diffusion dans l'espace où évolue la variable Γ , avec un coefficient négatif. Alors qu'une équation standard de diffusion homogénéise la quantité dont elle gouverne l'évolution, une équation de diffusion inverse concentre cette quantité. L'homogénéisation du scalaire se traduit effectivement dans l'espace des densités de probabilité par une concentration. Considérons ainsi une situation initiale où le champ prend partout pour valeur l'une des deux bornes du scalaire

⁷Les détails de ces calculs sont reproduits dans l'annexe A.

($c \in [0, 1]$). La densité de probabilité correspondante pourrait s'écrire $p(\Gamma) = \alpha (\delta(\Gamma) + \delta(\Gamma - 1))$. Une évolution normale serait une relaxation de cette situation vers une densité de probabilité gaussienne, dont la variance irait en se réduisant jusqu'à ce qu'on atteigne la situation finale où $p(\Gamma) = 2\alpha \delta(\Gamma - \langle c \rangle)$, c'est-à-dire un champ c partout égal à $\langle c \rangle$ et une situation complètement mélangée.

Le fait que la densité de probabilité relaxe vers une gaussienne n'est pas absolument universel. Cette hypothèse a été appuyée par des résultats expérimentaux [26, 27], et numériques [28, 29]. Il semble cependant que la densité de probabilité présente des queues exponentielles lorsque l'échelle intégrale du scalaire est plus grande que celle de la vitesse [30].

Lorsque la densité de probabilité adopte un profil gaussien, des résultats théoriques [31, 32], expérimentaux [26] et numériques [28] laissent penser que la dissipation conditionnelle est constante. Le comportement de cette dernière dépend par contre fortement de l'état de mélange.

1.3.2 Fermeture de l'équation de la densité de probabilité

Les équations (1.21) et (1.22) ne sont hélas pas fermées, mais le problème de fermeture ne se situe plus au niveau du terme de réaction. Ce dernier apparaît naturellement fermé dans les équations. Le problème de fermeture est rejeté sur des termes comme $\langle \Delta c | \Gamma \rangle$ ou $\langle \epsilon_c | \Gamma \rangle$. Ces derniers contiennent de l'information sur la dissipation du scalaire. Bien que se calculant à partir du seul champ scalaire, la dissipation est étroitement liée aux caractéristiques de l'écoulement turbulent.

La première tentative de fermeture de ces équations est le fait de Dopazo [24]. Il propose d'approcher la diffusion conditionnelle par une fonction affine

$$\langle D\Delta c | \Gamma \rangle \simeq \frac{6D}{\lambda_c^2} (\langle c \rangle - \Gamma). \quad (1.23)$$

avec

$$\lambda_c^2 = \frac{2\sigma_c^2}{\langle (\partial_x c)^2 \rangle}.$$

Il est intéressant d'exposer rapidement le raisonnement qui conduit Dopazo à proposer cette forme pour la diffusion conditionnelle. La première étape consiste à exprimer la diffusion conditionnelle grâce à une moyenne conditionnelle en deux points

$$\langle D\Delta c | \Gamma \rangle = D\Delta_{\vec{r}} \langle c(\vec{x} + \vec{r}) | \Gamma \rangle_{\vec{r}=\vec{0}}. \quad (1.24)$$

Dans la suite de ce paragraphe la quantité $c(\vec{x} + \vec{r})$ sera notée c' , et un écoulement homogène et isotrope sera considéré. Dopazo fait alors l'hypothèse que c et c' sont des variables gaussiennes corrélées et que par conséquent la moyenne conditionnée en deux points s'écrit

$$\langle c' | \Gamma \rangle = \langle c \rangle + \frac{\langle c' c \rangle - \langle c \rangle^2}{\sigma_c^2} (\Gamma - \langle c \rangle) \quad (1.25)$$

La dépendance vis-à-vis de $\vec{r}$ est alors toute entière contenue dans la quantité $\langle c' c \rangle$, et on voit déjà apparaître la partie affine vis-à-vis de Γ . Il est alors aisé d'en déduire l'expression de la diffusion

conditionnelle :

$$\begin{aligned}
\langle D \Delta c | \Gamma \rangle &= D \Delta_{\vec{r}} \langle c(\vec{x} + \vec{r}) | \Gamma \rangle \Big|_{\vec{r}=\vec{0}} \\
&= \frac{\Gamma - \langle c \rangle}{\sigma_c^2} D \Delta_{\vec{r}} \langle c' c \rangle \Big|_{\vec{r}=\vec{0}} \\
&= -\frac{3 D \langle (\partial_x c)^2 \rangle}{\sigma_c^2} (\Gamma - \langle c \rangle)
\end{aligned}$$

L'inconvénient majeur de cette fermeture est qu'elle peut pas générer de relaxation vers une densité de probabilité gaussienne en fin de mélange car elle préserve la forme initiale de la densité de probabilité. En effet, l'équation d'évolution de p devient

$$\partial_t p + \frac{6D}{\lambda_c^2} (\Gamma - \langle c \rangle) \partial_\Gamma p = \frac{6D}{\lambda_c^2} p. \quad (1.26)$$

Les solutions de cette équation sont alors de la forme

$$p(\Gamma, t) = f \left((\Gamma - \langle c \rangle) e^{\frac{6D}{\lambda_c^2} t} + \langle c \rangle \right) e^{\frac{6D}{\lambda_c^2} t}, \quad (1.27)$$

où la fonction f est déterminée par les conditions initiales, qui s'écrivent $p(\Gamma, 0) = f(\Gamma)$. Une diffusion conditionnelle affine signifie une évolution auto-similaire pour la densité de probabilité. Cela n'empêche pas la fermeture d'être couramment utilisée [23], essentiellement pour sa simplicité.

La vitesse du mélange est ici entièrement contrôlée par le coefficient $\frac{\langle \epsilon_c \rangle}{2\sigma_c^2} = \frac{1}{2\tau_c}$ qui représente la pente de la diffusion conditionnelle. Même si elle ne fait apparaître que des quantités se rapportant au scalaire, il convient de rappeler que la dissipation est générée par la turbulence. La constante de temps du mélange $\tau_c = \frac{\sigma_c^2}{\epsilon_c}$ prend donc en compte d'autres phénomènes que la seule diffusion. L'hypothèse est couramment faite que τ_c est imposée par la turbulence et que cette quantité est égale à deux fois la constante de temps de la turbulence [21]. Dans une simulation utilisant un modèle de turbulence $k - \epsilon$, c'est en effet la seule constante de temps dont on dispose pour évaluer l'efficacité du mélange turbulent.

D'autres méthodes de fermeture existent [33], mais elles sont en général beaucoup plus lourdes que de supposer ainsi une forme *a priori* pour la diffusion conditionnelle.

1.4 Simulation des Grandes Échelles

La Simulation des Grandes Échelles (SGE), encore appelée LES (pour Large Eddy Simulation), est une approche basée sur la séparation des grandes et des petites structures des écoulements⁸. Le but est de simuler uniquement l'évolution des grandes échelles. La simulation fait nécessairement appel à la modélisation de l'effet des petites échelles sur les grandes. Les petites échelles sont en effet supposées avoir un comportement universel et se prêter ainsi particulièrement bien à la modélisation.

En ne résolvant pas les petites échelles, on gagne évidemment beaucoup de temps de calcul. On peut ainsi aborder des situations que la simulation numérique directe ne permet pas de

⁸La SGE est en général opposée dans la littérature aux méthodes RANS [21]. D'un point de vue formel, les deux sont cependant extrêmement proches : il suffit d'invertir le filtrage spatial et la moyenne statistique. De ce point de vue, les méthodes de fermeture des termes de réaction sont les mêmes en SGE et en RANS.

reproduire, et d'atteindre des valeurs plus grandes du nombre de Reynolds. Ce type de simulations a été proposé initialement par Smagorinsky [34], qui introduit une viscosité turbulente pour prendre en compte les effets des petites échelles. L'estimation de celle-ci a fait l'objet de nombreux développements [35].

La présentation des SGE - et des problèmes de fermeture qui leur sont associés - proposée ici est classique. Nous verrons que la fermeture des termes de réaction chimique requiert une approche originale, qui fait généralement intervenir la notion de densité de sous-maille. Une fois que les modèles courants basés sur cette notion auront été exposés, des démonstrations rigoureuses des propriétés de la densité de sous-maille seront proposées.

1.4.1 Définition des grandes échelles

Les grandes échelles des champs étudiés sont isolées par un filtrage passe-bas. Nous noterons par la suite $\bar{c}$ les grandes échelles du champ c . Elles peuvent ainsi s'écrire

$$\bar{c}(\vec{x}) = \int G(\vec{x}, \vec{x}', t) c(\vec{x}') d\vec{x}'.$$

Un tel filtrage est linéaire, ce qui fait que l'intervalle de travail n'a aucune influence sur l'équation finale d'évolution. Par la suite, nous utiliserons un filtrage stationnaire - G ne dépendant alors pas de t . Le filtrage commute dans ces conditions avec la dérivation temporelle. Nous considérons aussi que le filtrage est homogène. Cela signifie que G ne dépend que de $\vec{x} - \vec{x}'$. L'opération de filtrage devient alors une convolution,

$$\bar{c} = G * c = \int G(\vec{x} - \vec{x}') c(\vec{x}') d\vec{x}',$$

ce qui se traduit dans l'espace spectral par une simple multiplication

$$\tilde{\bar{c}} = \tilde{G} \cdot \tilde{c}.$$

Le filtrage commute alors avec la dérivation partielle vis-à-vis des coordonnées spatiales⁹.

Dans la littérature, tous les filtrages sont couramment considérés comme étant équivalents [36]. Il m'a paru nécessaire de présenter ici les filtrages que j'ai utilisés ou mentionnés dans mon travail. Ils ne sont pas réellement isotropes mais "cubiques". La généralisation aux filtrages isotropes de toutes les propriétés qui seront démontrées par la suite se fait cependant sans difficulté.

Filtrage porte physique

Le filtrage porte physique de taille caractéristique Δ correspond à

$$G(\vec{r}) = \frac{1}{\Delta^3} \prod_{i=1}^3 H\left(r_i - \frac{\Delta}{2}\right).$$

On peut donc l'écrire comme la moyenne du champ considéré sur un cube $D(\vec{x})$ de côté Δ centré en $\vec{x}$:

$$\bar{c}(\vec{x}) = \frac{1}{\Delta^3} \int_{D(\vec{x})} c(\vec{x}', t) d\vec{x}'.$$

⁹Dans des situations où les conditions de bord sont présentes, près d'une paroi par exemple, le maillage utilisé n'est en général plus régulier, et le filtrage considéré plus homogène. Il faut dans ce cas prendre en compte des termes dus à la non commutation du filtrage avec les dérivées spatiales. Dans le cadre de conditions aux limites périodiques, il est naturel de choisir un filtrage homogène.

La transformée de Fourier de G s'écrit alors

$$\tilde{G}(\vec{k}) = \prod_i \text{sinc} \left(\frac{1}{2} k_i \Delta \right).$$

Ce filtrage prend dans la suite de ce travail une grande importance. Dans le cas des champs générés par la méthode pseudo-spectrale, il est effectué dans l'espace spectral en multipliant chacun des modes par un sinus cardinal – comme l'indique la transformée de Fourier précédente. Un calcul plus explicite est présenté en annexe, qui permet de souligner que le filtrage ainsi effectué est un filtrage porte physique rigoureux. Bien souvent, pour des raisons pratiques évidentes, le filtrage par peigne de Dirac est utilisé à la place du filtrage porte physique.

Filtrage gaussien

Le filtrage gaussien de taille Δ correspond à

$$G(\vec{r}) = \left(\frac{8}{\pi \Delta^2} \right)^{\frac{3}{2}} e^{-\frac{8}{\Delta^2} r^2}$$

et $\tilde{G}$ est aussi une gaussienne.

Filtrage porte spectral

Le filtrage porte spectral est le plus naturel, au sens où il permet de séparer le plus simplement possible les échelles dans l'espace spectral. Il s'écrit en effet

$$\tilde{G}(\vec{k}) = \prod_i H \left(\frac{2\pi}{\Delta} - |\vec{k}_i| \right).$$

Dans les directions des trois axes choisis, seules les fréquences inférieures à la fréquence de coupure $k_c = \frac{2\pi}{\Delta}$ sont conservées.

Son expression dans l'espace physique est

$$G(\vec{r}) = \prod_i \frac{2}{r_i} \sin \frac{2\pi r_i}{\Delta}.$$

C'est souvent la version isotrope de ce filtrage qui est utilisées. Elle est simple à mettre en oeuvre : dans l'espace spectral, il suffit de ne conserver que les modes de Fourier pour lesquels $k^2 \leq k_c^2$.

Filtrage peigne de Dirac

Le filtrage “peigne de Dirac” se conçoit surtout dans le cas où l'on connaît les valeurs des champs en les noeuds d'une grille régulière. Ce filtrage n'est en définitive qu'une moyenne du champ sur les plus proches voisins, ce qui s'écrit

$$\bar{c}(\vec{x}) = \sum_{\{\vec{r}_i\}} c(\vec{x} + \vec{r}_i).$$

On peut aussi écrire l'expression précédente comme une convolution $G * c$ à condition de prendre

$$G(\vec{r}) = \sum_{\{\vec{r}_i\}} \delta(\vec{r} - \vec{r}_i).$$

1.4.2 Équations d'évolution

Pour trouver les équations d'évolution des champs filtrés, il faut appliquer l'opération de filtrage aux équations d'évolution des champs considérés. Sachant ensuite que le filtrage est linéaire et qu'il commute avec la dérivation temporelle et spatiale, il est aisé d'obtenir les équations d'évolution des quantités filtrées. Les termes qui ne sont pas fermés deviennent alors apparents.

Si l'équation d'évolution du champ c est donnée par (1.1), celle du champ scalaire filtré s'écrit ainsi

$$\partial_t \bar{c} + \partial_i (\bar{v}_i \bar{c}) = D \Delta \bar{c}. \quad (1.28)$$

Pour le champ de vitesse filtré, les équations d'évolution (1.2) et (1.3) deviennent

$$\partial_i \bar{v}_i = 0 \quad (1.29)$$

$$\partial_t \bar{v}_i + \partial_j (\bar{v}_i \bar{v}_j) = -\partial_i \left(\frac{\bar{p}}{\rho} \right) + \nu \Delta \bar{v}_i, \quad (1.30)$$

et l'équation (1.29) peut se remplacer par celle qui donne le champ de pression filtré

$$\Delta \left(\frac{\bar{p}}{\rho} \right) = -\partial_i \partial_j \bar{v}_i \bar{v}_j. \quad (1.31)$$

On note que pour le scalaire et pour la vitesse, les termes non-linéaires $\bar{v}_i \bar{c}$ et $\bar{v}_i \bar{v}_j$ ne peuvent pas se déduire simplement des seules grandes échelles. Les termes non-linéaires sont en effet ceux qui couplent les échelles entre elles - et notamment les petites avec les grandes, quel que soit le filtrage utilisé. Ces deux termes doivent être modélisés de manière à refléter l'effet des petites échelles sur les grandes. Il est bien sûr possible d'écrire des équations d'évolution pour ces quantités, mais elles ne sont elles-mêmes pas fermées. En écrivant les équations d'évolution de ces termes, on en fait apparaître de nouveaux d'ordre supérieur et qui ne sont pas fermés non plus. Ecrire récursivement des équations d'évolution pour les quantités qui ne sont pas fermées nous place donc face à une hiérarchie infinie d'équations. Le parti pris de la Simulation des Grandes Échelles est de fermer directement les premières équations d'évolution.

Les petites échelles du champ de vitesse sont caractérisées par un phénomène de cascade jusqu'à l'échelle de Kolmogorov, où l'énergie cinétique est finalement dissipée. La présence de cette cascade permet de dissiper plus rapidement l'énergie en créant des variations spatialement rapides de vitesse. C'est ce qui justifie que l'on modélise usuellement l'effet des petites échelles sur les grandes en introduisant une viscosité turbulente. Smagorinsky a proposé dans son article original [34] de prendre

$$\bar{v}_i \bar{v}_j \simeq \bar{v}_i \bar{v}_j + \frac{1}{3} (\overline{v_k^2} - \bar{v}_k^2) \delta_{ij} - 2 \nu_t(\vec{x}) \bar{s}_{ij}, \quad (1.32)$$

où $\bar{s}_{ij} = \frac{1}{2}(\partial_i \bar{v}_j + \partial_j \bar{v}_i)$. Cela revient à prendre une viscosité turbulente, qui ici n'est pas uniforme, mais varie spatialement. Pour le scalaire, les petites échelles jouent cette fois un rôle essentiel dans le mélange, puisque c'est dans les petites structures que le champ s'homogénéise le plus rapidement. On peut ainsi envisager de modéliser l'effet des petites échelles sur les grandes par une diffusivité turbulente, ce qui s'écrit

$$\bar{v}_i \bar{c} \simeq \bar{v}_i \bar{c} - D_t(\vec{x}) \partial_i \bar{c} \quad (1.33)$$

Cette idée très simple est critiquable en ce qu'il n'y a pas de raison véritable pour que la quantité $\overline{v_i v_j} - \overline{v_i} \overline{v_j}$ soit alignée avec le tenseur $\overline{\mathcal{S}_{ij}}$. C'est encore plus net dans le cas du scalaire, où l'hypothèse revient à supposer que $\overline{v_i c} - \overline{v_i} \overline{c}$ est aligné avec le gradient du scalaire filtré $\partial_i \overline{c}$. Si le modèle prend raisonnablement en compte ce qui se produit dans la direction du gradient de $\overline{c}$, ce qui se passe orthogonalement à cette direction est complètement ignoré. La plupart des améliorations proposées au modèle de Smagorinsky visent à mieux estimer la viscosité turbulente, sans prendre en compte d'autres comportements éventuels. Ces améliorations sont détaillées dans des articles de revue [35, 37], qui concluent qu'on ne dispose pas encore d'une modélisation totalement satisfaisante des phénomènes mis en jeu.

1.4.3 Simulation de réactions chimiques et fermeture

La prise en compte des réactions chimiques amène de nouveaux problèmes en simulation des grandes échelles car elle requiert généralement l'estimation de $\overline{f(c)}$, où f est une fonction fortement non-linéaire, à partir du seul champ $\overline{c}$.

Pour prendre en compte une réaction chimique, on cherche à se ramener à des scalaires conservés. Dans le cas de la réaction chimique infiniment rapide décrite en 1.1.6, on peut se ramener à un seul scalaire conservé, mais la fonction f qui permet de calculer la concentration de produit de la réaction à partir de la valeur du scalaire est non-linéaire. Le problème est alors de calculer les grandes échelles de la concentration de produit, $\overline{f(c)}$, à partir des grandes échelles du scalaire conservé $\overline{c}$.

Quand il n'est pas possible de se ramener uniquement à des scalaires conservés, on peut avoir à étudier un scalaire qui obéisse à une équation de la forme [18]

$$\partial_t \overline{c} + \partial_i \overline{v_i c} = D \Delta \overline{c} + \overline{f(c)}. \quad (1.34)$$

Dans ce cas, l'estimation de $\overline{f(c)}$ à partir du seul champ $\overline{c}$ est nécessaire pour fermer l'équation d'évolution.

Dans les deux cas décrits ici, le problème est le même : la donnée des grandes échelles d'un scalaire ne renseigne pas suffisamment sur la répartition réelle du scalaire pour qu'on ait une expression exacte de $\overline{f(c)}$. Pour modéliser $\overline{f(c)}$ sans connaître f *a priori*, et f étant en général très fortement non-linéaire, il faut nécessairement une approche fondamentalement différente de la fermeture par une viscosité ou une diffusivité turbulente. L'approche adoptée ici présente une similitude nette avec l'approche par densité de probabilité décrite en 1.3.

Densité de sous-maille

Dans un article faisant le point sur l'utilisation des méthodes statistiques pour la combustion [38], Pope introduit ce qu'il appelle la FDF, pour "Filtered Density Function", définie par

$$d_s(\Gamma) = \overline{\delta(\Gamma - c(\vec{x}', t))}(\vec{x}).$$

Nous appellerons cette quantité la **densité de sous-maille**, et la raison de cette dénomination deviendra plus claire un peu plus loin. C'est sur cette définition précise de la densité de sous-maille que repose l'essentiel des résultats présentés dans cette thèse. Cette définition ouvre en effet la voie à une formalisation rigoureuse de la modélisation de sous-maille. Il faut remarquer au passage

que la densité de sous-maille n'est absolument pas une quantité statistique, contrairement à ce que laissent couramment sous-entendre certains auteurs¹⁰.

L'importance de la densité de sous-maille vient de ce que la connaissance de cette quantité suffit à calculer $\overline{f(c)}$ grâce à la relation [38]

$$\overline{f(c)} = \int f(\Gamma) d_s(\Gamma) d\Gamma. \quad (1.35)$$

Deux approches radicalement différentes se distinguent alors à ce stade. La première est de chercher une équation d'évolution pour la densité de sous-maille et de la résoudre numériquement [39], la seconde est de chercher directement une approximation de la densité de sous-maille [19].

La première approche fait intervenir une équation d'évolution pour la densité de sous-maille qui n'est pas fermée, reflétant le fait que la densité de sous-maille contient des informations sur la sous-maille qui ne sont pas présentes dans le champ grandes échelles seul. Des tentatives [39] ont eu lieu pour fermer l'équation d'évolution de d_s et pour trouver comment la résoudre numériquement. Hélas, ces tentatives sont souvent laborieuses, demandent beaucoup de nouvelles hypothèses, et les méthodes utilisées pour estimer l'évolution de d_s sont compliquées et coûteuses en temps de calcul.

C'est ce qui justifie l'intérêt de la seconde approche, qui apparaît comme simple et relativement peu gourmande et dont on pense aujourd'hui qu'elle est la voie la plus prometteuse vers la prise en compte des réactions chimiques dans les SGE. Nous verrons dans la suite dans quelle mesure les hypothèses qui la sous-tendent sont justifiées. Cette approche est couramment utilisée en simulation [40, 41], et a fait l'objet de plusieurs études [42, 43] et développements [44].

Modèle de Cook et Riley

Cook et Riley [19] proposent d'approcher la densité de sous-maille par une fonction de même moyenne et de même variance, de la forme

$$\beta(\Gamma; \bar{c}, \bar{c}^2 - \bar{c}^2) = \frac{\Gamma^{a-1}(1-\Gamma)^{b-1}}{B(a, b)},$$

avec

$$B(a, b) = \int \Gamma^{a-1}(1-\Gamma)^{b-1} d\Gamma,$$

de manière à ce que la fonction soit normée, et

$$a = \bar{c} \left(\frac{\bar{c}(1-\bar{c})}{\bar{c}^2 - \bar{c}^2} - 1 \right) \quad (1.36)$$

$$b = \frac{a}{\bar{c}} - a \quad (1.37)$$

pour que la moyenne de β soit $\bar{c}$ et que sa variance soit $\bar{c}^2 - \bar{c}^2$. Ce type de distribution est appelé **distribution β** et la fonction $B(a, b)$ est la **fonction bêta d'Euler** [45]. La démonstration des relations (1.36) et (1.37) à partir de la moyenne et de la variance de la distribution β se trouve dans l'annexe C. La justification du choix de $\bar{c}$ et de $\bar{c}^2 - \bar{c}^2$ comme moyenne et variance se trouve un peu plus loin.

¹⁰La densité de sous-maille est parfois appelée "Sugrid Probability Density Function".

Ce modèle fournit donc naturellement une approximation de la quantité $\overline{f(c)}$ quelle que soit la fonction f . Elle est obtenue en remplaçant la densité de sous-maille par la fonction β dans l'équation (1.35) qui devient :

$$\overline{f(c)} \simeq \int f(\Gamma) \beta(\Gamma; \bar{c}, \sigma_s^2) d\Gamma.$$

Cependant, Cook et Riley soulignent que la quantité $\overline{c^2} - \bar{c}^2$, appelée **variance de sous-maille** n'est pas connue en pratique. La connaissance du champ $\bar{c}$ ne permet pas d'en obtenir une expression exacte. Aussi est-on obligé d'estimer la variance de sous-maille à partir des grandes échelles, et de faire ainsi appel à ce qu'on appelle un **sous-modèle**.

Critères de qualité du modèle

Dans les études qui ont été menées du modèle précédent, deux critères de pertinence sont utilisés. Le premier est le test de reconstruction introduit par Jimenez *et al.*[43], le second est l'erreur quadratique commise par le modèle [42].

Dans son étude, Jimenez introduit la propriété

$$\langle d_s(\Gamma) \rangle = \overline{p(\Gamma)}. \quad (1.38)$$

Supposons qu'on approche la densité de sous-maille par une fonction $p(c; \pi)$ où π est un jeu de paramètres. Le modèle de Cook et Riley entre tout à fait dans ce cadre puisqu'il propose de prendre $p(c; \pi) = \beta(c; \bar{c}, \sigma_s^2)$, avec $\pi = \{\bar{c}, \sigma_s^2\}$. Alors Jimenez en déduit qu'un bon modèle doit vérifier la **propriété de reconstruction**¹¹, directement inspirée de la propriété (1.38) :

$$\overline{p(\Gamma)} = \int p(c; \pi) p(\pi) d\pi. \quad (1.39)$$

Si un modèle vérifie la propriété de reconstruction, alors on peut écrire

$$\begin{aligned} \langle \overline{f(c)} \rangle &= \overline{\langle f(c) \rangle} \\ &= \overline{\int f(\Gamma) p(\Gamma) d\Gamma} \\ &= \int f(\Gamma) \overline{p(\Gamma)} d\Gamma \\ &= \int f(\Gamma) p(\Gamma; \pi) p(\pi) d\Gamma d\pi. \\ &= \left\langle \int f(\Gamma) p(\Gamma; \pi) d\Gamma \right\rangle \end{aligned}$$

En somme, si un modèle vérifie la propriété de reconstruction, alors la moyenne de l'estimation de $\overline{f(c)}$ qu'il fournit est exactement la moyenne de $\overline{f(c)}$ - ce qu'on est en droit d'attendre de tout bon modèle. Souvent le test de reconstruction est effectué en écoulement statistiquement homogène pour lequel $\overline{p(\Gamma)} = p(\Gamma)$.

¹¹Dans l'article de Jimenez [43] qui est le premier à ma connaissance à mentionner cette propriété, celle-ci est attribuée à Cook et Riley.

Dans l'étude menée par Moin [42] le critère de qualité du modèle est différent. Il s'agit de l'erreur quadratique commise par celui-ci sur l'estimation de $\overline{f(c)}$, qui s'écrit

$$E = \left\langle \left(\overline{f(c)} - \int f(\Gamma) p(\Gamma; \pi) d\Gamma \right)^2 \right\rangle.$$

Cette erreur fournit une mesure de la qualité d'un modèle. Elle représente l'écart moyen à la quantité à estimer, $\overline{f(c)}$. Le critère de reconstruction ne permet de se faire une idée que de l'écart entre la moyenne de $f(c)$ et la prévision du modèle pour cette même quantité. L'erreur quadratique est de plus un critère quantitatif. Notons cependant qu'il est moins intrinsèque que le précédent, puisqu'il dépend de la fonction f considérée.

Sous-modèles

Le sous-modèle initialement proposé par Cook et Riley pour l'estimation de la variance de sous-maille $\sigma_s^2 = \overline{c^2} - \bar{c}^2$ s'écrit

$$\overline{c^2} - \bar{c}^2 \simeq C_s \left(\hat{\bar{c}}^2 - \bar{c}^2 \right),$$

où $\hat{\cdot}$ est un filtre "test" de taille caractéristique deux fois supérieure à celle du filtre fondamental. Pour justifier ce choix, Cook et Riley font appel à un argument d'invariance d'échelle. En effet, la quantité $\hat{\bar{c}}^2 - \bar{c}^2$, n'est autre que la variance de sous-maille du champ $\bar{c}$ à l'échelle du filtre test.

Pierce et Moin [42, 44] ont proposé un autre sous-modèle. Pour pouvoir utiliser une procédure dynamique proche de celle de Germano [36, 46], ils ont d'abord supposé une relation de proportionnalité entre la variance de sous-maille et la dissipation du champ scalaire filtré

$$\overline{c^2} - \bar{c}^2 = C \overline{\Delta^2} (\partial_i \bar{c})^2, \quad (1.40)$$

où $\overline{\Delta}$ est la taille caractéristique du filtre considéré.

On peut poser l'égalité suivante, analogue de celle de Germano [36],

$$\widehat{\overline{c^2} - \bar{c}^2} = \hat{\bar{c}}^2 - \bar{c}^2 - \left(\hat{\bar{c}}^2 - \bar{c}^2 \right), \quad (1.41)$$

où $\hat{\cdot}$ est un filtre test de taille caractéristique $\hat{\Delta}$. Si on suppose alors que la relation (1.40) reste un bon modèle à toutes les échelles de filtrage, et si on note $\hat{\Delta}$ la taille caractéristique du filtre $\hat{\cdot}$, la relation (1.41) devient

$$C \overline{\Delta^2} (\partial_i \bar{c})^2 = C \hat{\Delta}^2 (\partial_i \hat{\bar{c}})^2 - \left(\hat{\bar{c}}^2 - \bar{c}^2 \right). \quad (1.42)$$

La seule inconnue dans cette relation étant la constante C , cette égalité fournit un moyen (voire même plusieurs) de la déterminer. On peut la prendre telle qu'elle minimise l'erreur notamment, en suivant une procédure utilisée par Lilly [47]. Cette procédure de détermination dynamique de la constante d'un modèle est caractéristique des modèles utilisant des hypothèses d'invariance d'échelle, dont le représentant le plus classique est celui de Germano pour le champ de vitesse [36]. Cette détermination pose cependant de nets problèmes de singularités. Dans le cas qui nous intéresse, la détermination dynamique de la constante peut donner des valeurs de C négatives, ce qui n'est pas compatible avec ce que l'on sait de la variance de sous-maille. Pour éliminer ces problèmes, il est couramment fait appel à des moyennes dans des directions d'homogénéité.

1.4.4 Densité de sous-maille

L'article original de Cook et Riley [19] présente quelques imprécisions à propos de la densité de sous-maille et de sa variance qui rendent sa compréhension parfois difficile. J'ai voulu dans cette partie redémontrer rigoureusement quelques propriétés de la densité de sous-maille.

Filtrages

Le choix du filtrage utilisé pour étudier la densité de sous-maille n'est pas innocent, il est même fondamental. Nous allons voir dans cette partie qu'il faut au moins prendre un filtrage positif, c'est à dire correspondant à une convolution par une fonction positive, pour que l'étude de la densité de sous-maille ait un sens. C'est un premier élément qui montre que contrairement à ce qu'on pourrait penser, les filtrages sont loin d'être tous équivalents.

Parce qu'il permet une interprétation aisée de la densité de sous-maille, j'ai choisi le filtrage porte physique comme filtrage fondamental à l'instar de Jimenez [43]. Pour mémoire, les grandes échelles d'un champ c sont alors données par

$$\bar{c}(\vec{x}) = \frac{1}{\Delta^3} \int_{D(\vec{x})} c(\vec{x}', t) \, d\vec{x}'. \quad (1.43)$$

Ce filtrage est positif, car la fonction G qui lui correspond (voir 1.4.1) est positive. Par la suite, la dépendance spatiale de $\bar{c}$ sera omise dans la notation, pour l'alléger. La notation $\overline{f(c)}$ représentera le champ $f(c)(\vec{x})$.

Une première propriété permet de se faire une idée de ce qui sépare les filtres positifs des autres.

Proposition 1.1 *Si le filtrage τ est un filtre positif, alors*

$$\forall \vec{x} \, a(\vec{x}) \leq b(\vec{x}) \implies \forall \vec{x} \, \overline{a}(\vec{x}) \leq \overline{b}(\vec{x})$$

◇ *Preuve.*

$$\begin{aligned} a(\vec{x}') \leq b(\vec{x}') &\implies G(\vec{x} - \vec{x}') a(\vec{x}') \leq G(\vec{x} - \vec{x}') b(\vec{x}') \\ &\implies \int G(\vec{x} - \vec{x}') a(\vec{x}') \, d\vec{x}' \leq \int G(\vec{x} - \vec{x}') b(\vec{x}') \, d\vec{x}', \end{aligned}$$

soit

$$\overline{a}(\vec{x}) \leq \overline{b}(\vec{x}).$$

□

Comme $0 \leq c(\vec{x}) \leq 1$, et au vu de la propriété précédente, $0 \leq \bar{c}(\vec{x}) \leq 1$. **Les filtrages positifs respectent donc les bornes du scalaire**, ce qui n'est pas le cas des filtrages quelconques. On peut donc déjà remarquer que comme tous les modèles supposent que les grandes échelles ont les mêmes bornes que le champ scalaire, ils supposent implicitement que le filtrage avec lequel ils sont utilisés est positif.

Densité de sous-maille

La fraction de volume du cube $D(\vec{x})$ occupée par un champ scalaire de valeur $c \leq \Gamma$, notée $V_\Delta(\Gamma, \vec{x})$ est l'intégrale sur le cube d'une fonction qui vaut 1 si $c(\vec{x}, t) \leq \Gamma$ et 0 sinon, divisée par le volume du cube. Cette fonction s'exprime simplement à l'aide d'une fonction de Heavyside, et l'on peut écrire

$$V_\Delta(\Gamma, \vec{x}) = \frac{1}{\Delta^3} \int_{D(\vec{x})} H(\Gamma - c(\vec{x}', t)) \, d\vec{x}'.$$

On appelle **densité de c dans le cube $D(\vec{x})$** ou encore **densité de sous-maille de c en $\vec{x}$** la dérivée de $V_\Delta(\Gamma, \vec{x})$ par rapport à Γ , et on la note $d_s(\Gamma, \vec{x})$.

Elle peut aussi s'écrire

$$\begin{aligned} d_s(\Gamma, \vec{x}) &= \frac{\partial V_\Delta(\Gamma, \vec{x})}{\partial \Gamma} \\ &= \frac{1}{\Delta^3} \int_{D(\vec{x})} \frac{\partial}{\partial \Gamma} H(\Gamma - c(\vec{x}', t)) \, d\vec{x}' \\ &= \frac{1}{\Delta^3} \int_{D(\vec{x})} \delta(\Gamma - c(\vec{x}', t)) \, d\vec{x}' \\ &= \overline{\delta(\Gamma - c(\vec{x}'))}(\vec{x}). \end{aligned}$$

La densité de sous-maille donne une idée de la répartition en valeur du scalaire à l'intérieur du domaine de filtrage $D(\vec{x})$. Elle ne contient en réalité aucune information sur la répartition spatiale du scalaire. Par la suite, la dépendance spatiale des quantités étudiées ne sera plus indiquée pour alléger la notation. La densité de sous-maille sera par exemple notée $\overline{\delta(\Gamma - c)}$.

Propriétés

On attend de la densité de sous-maille qu'elle possède un certain nombre de propriétés, notamment qu'elle soit positive et normée.

Proposition 1.2 *La densité de sous-maille est positive : $\forall \Gamma \, d_s(\Gamma) \geq 0$*

◊ *Preuve.*

$$\begin{aligned} V_\Delta(\Gamma + \gamma) &= \overline{H(\Gamma + \gamma - c)} \\ &= \overline{H(\Gamma - c) + H(\Gamma + \gamma - c) - H(\Gamma - c)} \end{aligned}$$

Or, tant que $\gamma \geq 0$, $H(\Gamma + \gamma - c) - H(\Gamma - c) \geq 0$, et comme le filtre est positif,

$$\overline{H(\Gamma + \gamma - c) - H(\Gamma - c)} \geq 0.$$

Donc V_Δ est une fonction croissante. Comme $\frac{\partial V_\Delta}{\partial \Gamma} = d_s$, la densité de sous-maille est une fonction positive. $\square$

Proposition 1.3 *La densité de sous-maille est normée*

$$\int_0^1 d_s(\Gamma) \, d\Gamma = 1$$

◇ *Preuve.*

Comme $\frac{\partial V_\Delta}{\partial \Gamma} = d_s$, $\int_0^\gamma d_s(\Gamma) \, d\Gamma = V_\Delta(\gamma)$. Donc $\int_0^1 d_s(\Gamma) \, d\Gamma = V_\Delta(1)$. Or $V_\Delta(1) = 1$, car $c \leq 1$. $\square$

Il faut bien remarquer que ces deux dernières propriétés, généralement considérées comme vérifiées, ne le sont que parce que le filtrage utilisé est positif. Les propriétés de la densité de sous-maille énoncées par Pope ou Jimenez ne sont pas difficiles à établir. La première est celle qui fait l'intérêt de l'étude de la densité de sous-maille :

Proposition 1.4

$$\forall f, \overline{f(c)}(\vec{x}) = \int f(\Gamma) \, d_s(\Gamma, \vec{x}) \, d\Gamma.$$

◇ *Preuve.*

En effet,

$$\begin{aligned} \int f(\Gamma) \, d_s(\Gamma, \vec{x}) \, d\Gamma &= \int f(\Gamma) \, \overline{\delta(\Gamma - c)} \, d\Gamma \\ &= \overline{\int f(\Gamma) \, \delta(\Gamma - c) \, d\Gamma} \\ &= \overline{f(c)} \end{aligned}$$

$\square$

Elle signifie que connaître la répartition en valeur du scalaire dans la sous-maille est suffisant pour calculer $\overline{f(c)}$ quelle que soit f . Avoir une approximation de la densité de sous-maille fournit ainsi automatiquement une approximation des grandes échelles de n'importe quelle fonction du champ scalaire. Cela donne une méthode de fermeture des membres inconnus liés à des taux de réaction, par exemple.

La moyenne de la densité de sous-maille se calcule aisément, grâce à la propriété 1.4, qui permet d'écrire en prenant $f(c) = c$ que

$$\int \Gamma \, d_s(\Gamma) \, d\Gamma = \overline{c}. \quad (1.44)$$

La variance de la densité de sous-maille est appelée **variance de sous-maille** et notée σ_s^2 .

$$\begin{aligned} \sigma_s^2 &= \int (\Gamma - \overline{c})^2 \, d_s(\Gamma) \, d\Gamma \\ &= \int \Gamma^2 d_s(\Gamma) \, d\Gamma - 2\overline{c} \int \Gamma d_s(\Gamma) \, d\Gamma + \overline{c}^2 \int d_s(\Gamma) \, d\Gamma \\ &= \overline{c^2} - \overline{c}^2 \end{aligned}$$

Proposition 1.5 *La variance de sous-maille peut être encadrée par*

$$0 \leq \sigma_s^2 \leq \bar{c} \cdot (1 - \bar{c})$$

◇ *Preuve.*

Comme $\forall \Gamma, d_s \geq 0$, et que $\forall \Gamma, (\Gamma - \bar{c})^2 \geq 0$, alors $\int (\Gamma - \bar{c})^2 d_s(\Gamma) d\Gamma \geq 0$.

D'autre part, on sait que le scalaire est borné, donc

$$\begin{aligned} 0 \leq c \leq 1 &\Rightarrow c^2 \leq c \\ c^2 \leq c &\Rightarrow \bar{c}^2 \leq \bar{c}. \end{aligned}$$

D'où $\bar{c}^2 - \bar{c}^2 \leq \bar{c}(1 - \bar{c})$.

□

La seconde inégalité reflète le fait que le champ scalaire est borné. Ceci interdit une variance trop grande, notamment lorsqu'on s'approche des bornes.

Proposition 1.6

$$\langle d_s(\Gamma) \rangle = \overline{p(\Gamma)}(\vec{x})$$

◇ *Preuve.*

La moyenne statistique et le filtrage commutent, donc

$$\langle d_s(\Gamma) \rangle = \left\langle \overline{\delta(\Gamma - c)} \right\rangle = \overline{\langle \delta(\Gamma - c) \rangle} = \overline{p(\Gamma)}(\vec{x})$$

□

La propriété suivante fournit pour finir une méthode effective pour accéder à la densité de sous-maille.

Proposition 1.7 *Considérons la variable aléatoire $c(M)$ où M est un point tiré aléatoirement dans le cube $D(\vec{x})$ de façon à ce que tous les points du cube soient équiprobables. Alors la densité de probabilité de cette variable est égale à la densité de sous-maille en $\vec{x}$.*

◇ *Preuve.*

M étant tiré aléatoirement dans le cube et de façon identiquement distribuée, la probabilité $P(c(M) \leq \Gamma)$ que la variable $c(M)$ soit inférieure à la valeur Γ est égale à la fraction du volume de $D(\vec{x})$ occupé par un champ scalaire de valeur inférieure à Γ . Cela s'écrit

$$P(c(M) \leq \Gamma) = V_\Delta(\Gamma). \quad (1.45)$$

Et par suite, la densité de probabilité de $c(M)$ s'écrit

$$p(c(M) = \Gamma) = \frac{dP(c(M) \leq \Gamma)}{d\Gamma} = \frac{dV_\Delta(\Gamma)}{d\Gamma} = d_s(\Gamma).$$

□

On peut donc tirer des points aléatoirement dans un cube de sous-maille et relever la densité de probabilité de la valeur du champ scalaire en ces points. Elle donne accès à la densité de sous-maille. Bien entendu, cela ne peut se faire qu'à la condition d'avoir accès au champ scalaire en tout point de l'espace - et donc en dehors des noeuds d'une éventuelle grille. Comme dans le cadre d'une simulation pseudo-spectrale le champ scalaire est défini en tout point de l'espace physique, il devient possible d'accéder à la variance de sous-maille. C'est cette particularité qui m'a permis de mener l'étude présentée dans le chapitre trois.

Chapitre 2

Caractérisation du mélange

Ce chapitre vise à caractériser l'état de mélange obtenu pour un cas d'injection particulier - en termes de spectres, puis de fonctions de structure pour mesurer l'intermittence du scalaire. Au passage, le rapport entre le degré de mélange (caractérisé par la variance du scalaire) et les caractéristiques de l'injection est étudié. Une attention particulière est pour finir portée à la diffusion conditionnelle, en raison de l'importance qu'elle revêt dans l'équation d'évolution de la densité de probabilité.

2.1 Simulations numériques directes

La partie de mon travail que décrit ce chapitre est essentiellement un prolongement de la thèse de Marc Elmo. De la même manière que ce dernier, j'ai mené des simulations numériques directes d'écoulements turbulents avec la méthode d'injection de scalaire qu'il a développée. J'ai cependant pu atteindre des résolutions (et donc des nombres de Reynolds) plus élevées grâce au développement des moyens de calcul. L'étude menée ici est en définitive complémentaire de celle de Marc Elmo.

2.1.1 Méthodes d'injection

Injection d'énergie cinétique

L'énergie cinétique est naturellement dissipée dans un écoulement turbulent. Dans une simulation sans apport extérieur d'énergie cinétique, dite "en déclin", il est plus difficile de voir émerger des comportements de cascade auto-similaire prévus par la théorie de Kolmogorov. Un tel phénomène de cascade suppose en effet une certaine forme d'équilibre entre les échelles, équilibre qui peut ne pas avoir le temps de s'instaurer dans une simulation en déclin. Pour ces raisons, il est intéressant d'obtenir une situation statistiquement stationnaire. On espère donner à une éventuelle cascade le temps de se manifester en forçant le champ de vitesse. Cette injection d'énergie se fait à grande échelle, plus particulièrement pour des fréquences comprises entre 0,05 et 0,1. Ceci assure que la présence d'énergie aux petites échelles provient d'une cascade (quand bien même elle ne serait pas auto-similaire).

L'injection d'énergie cinétique est ici basée sur l'utilisation d'équations de Langevin[28, 48, 49], ce qui lui confère une certaine mémoire, et un caractère aléatoire correspondant à ce qu'on peut imaginer de plus efficace pour générer de la turbulence à valeurs de R_λ élevées.

Injection de scalaire

Pour obtenir une situation statistiquement stationnaire pour le champ scalaire, ce qui facilite nettement l'étude du mélange, il est aussi nécessaire d'injecter de la variance. Jusqu'au travail de Marc Elmo [2, 50, 51], inspiré par Sabel'nikov, le moyen constamment utilisé pour obtenir une situation stationnaire était l'injection par gradient moyen imposé [52, 53, 29]. Une telle injection ne permet pas de choisir une éventuelle échelle d'injection du scalaire. De plus, elle ne correspond pas à l'idée que l'on peut se faire de l'injection d'un scalaire dans un écoulement : on imagine plutôt une injection localisée, même s'il s'agit d'une grosse structure, où le scalaire prend des valeurs extrêmes (qui en général déterminent ses bornes). C'est ce type d'injection que la méthode décrite ci-dessous vise à reproduire.

Paramètres de l'injection

Considérons un champ scalaire prenant ses valeurs entre 0 et 1. La méthode d'injection mise au point par Elmo consiste à injecter de façon aléatoire spatialement et périodique en temps des cubes de taille fixée qui soient emplis de scalaire prenant la valeur 0 ou 1. Pour conserver une moyenne à 0,5 et des statistiques symétriques par rapport à cette moyenne, il faut injecter autant de cubes à $c = 0$ que de cubes à $c = 1$. Il convient cependant de remarquer que pour éviter tout phénomène de Gibbs (qui ne manquerait pas d'apparaître si les cubes étaient injectés tels quels) on lisse le bord des cubes injectés de manière à ce qu'ils ne produisent pas de discontinuités du champ scalaire.

Les trois paramètres qui permettent de contrôler l'injection sont

- Δt , la période temporelle de l'injection,
- n , le nombre total de cubes injectés, en général pris pair,
- ℓ , le côté des cubes injectés.

Caractérisation de l'injection

A partir de ces paramètres, on peut construire un temps caractéristique de l'injection : celui que mettrait théoriquement l'injection pour renouveler l'ensemble du domaine. C'est donc la période Δt divisée par la fraction du cube renouvelée à chaque injection :

$$T_{res} = \Delta t \frac{L^3}{n \ell^3}.$$

On peut prendre pour taille caractéristique de l'injection l'échelle intégrale du scalaire

$$L_c = \frac{\pi}{2 \sigma^2} \int_0^{+\infty} \frac{1}{k} |\tilde{c}(k)|^2 dk.$$

Deux nombres sans dimensions peuvent ainsi être construits qui caractérisent l'injection

$$\begin{aligned} R_t &= \frac{T_{res}}{T_{turb}} \\ R_\ell &= \frac{L_c}{L_u}, \end{aligned}$$

où L_u est l'échelle intégrale de vitesse¹, et $T_{turb} = \frac{L_u}{u'}$ le temps caractéristique de la turbulence².

¹Définie par $L_u = \frac{3 \pi}{4 \langle u^2 \rangle} \int_0^{+\infty} \frac{1}{k} |\tilde{u}(k)|^2 dk$

²Avec $u' = \sqrt{\langle u^2 \rangle}$.

2.1.2 Caractérisation des simulations

Mon travail est basé sur des simulations numériques directes qui ont été menées à l'Idris³ sur une machine vectorielle. Les principes fondamentaux de ces simulations (représentation des champs et méthodes d'injection) ont été décrits dans le chapitre précédent. Les caractéristiques de ces simulations sont résumées dans le tableau 2.1.

L'étude menée par Marc Elmo a permis de montrer que le taux d'injection de scalaire pouvait plus ou moins marquer les caractéristiques du champ scalaire. Les taux utilisés ici sont parmi les plus faibles utilisés par Marc Elmo, ceux qui marquent le moins le champ de scalaire et, entre autres choses, laissent se développer des densités de probabilités proche de gaussiennes au centre. L'échelle d'injection a été choisie pour que l'échelle intégrale du scalaire soit la plus grande possible, afin que puisse apparaître une zone inertielle la plus importante possible. Pour finir l'énumération des principales différences entre les simulations de Marc et les miennes il faut préciser que les capacités de calcul ayant évolué, j'ai pu mener des simulations d'une plus grande résolution (256^3 contre 128^3).

Pour chaque simulation, une fois vérifié qu'un état stationnaire avait été atteint, des champs ont été sélectionnés à des instants différents (de manière à ce que les champs prélevés soient décorrélés). C'est de cette manière qu'on a obtenu les données de références pour ce travail. Un post-traitement de ces données a ensuite été effectué pour obtenir les résultats présentés ici. Cependant, comme les grandeurs qui décrivent le degré de mélange (comme la variance du scalaire) varient légèrement au cours du temps, il n'est pas possible de s'assurer que les champs prélevés soient parfaitement caractéristiques de la situation étudiée - et par exemple que la variance du scalaire est la moyenne temporelle de cette variance lors de la simulation. Par opposition, Marc Elmo effectuait des statistiques à chaque pas de temps au cours de la simulation - ce qui imposait de mener une nouvelle simulation pour étudier toute nouvelle quantité, et c'est pourquoi cette solution n'a pas été retenue. Cette méthode fournissait par contre les moyennes temporelles des quantités étudiées.

Pour caractériser mes simulations, les densités de probabilité du scalaire associées aux séries 256-I-1 et 256-II-1 ont été présentées sur la figure 2.3. On y constate qu'une injection plus faible rend la densité de probabilité plus piquée, c'est-à-dire conduit à un mélange plus important. L'évolution temporelle de l'écart-type σ est présentée figure 2.4, et nous montre qu'après un court régime transitoire, l'écart-type semble osciller faiblement autour d'une valeur moyenne⁴. La figure 2.5 nous montre plus précisément l'évolution aux temps courts de la variance et servira de base à la réflexion menée dans le paragraphe suivant.

Les spectres permettent également de caractériser les écoulements et les situations de mélange. La figure 2.6 présente des spectres pour deux simulations (128-II-1 et 256-II-1) correspondant à la même injection mais à des résolutions différentes, et l'on peut constater qu'une zone inertielle plus longue semble effectivement se développer quand la résolution est maximale, c'est-à-dire quand le nombre de Reynolds est plus grand. On peut cependant constater que l'exposant de la loi de puissance pour le spectre en 256^3 est légèrement inférieur à $\frac{5}{3}$, ce que confirme la figure 2.7. Elle montre qu'un spectre multiplié par $k^{1,45}$ présente un plateau très net. La figure 2.8 permet de comparer le spectre du scalaire à celui du champ de vitesse et on peut constater que le champ de

³La machine utilisée est appelée *uqbar*, il s'agit d'une SX-5 de NEC. Elle est classée 162ème dans le classement de juin 2002 des 500 machines les plus puissantes du monde. Cette liste est disponible sur le site <http://www.top500.org>.

⁴L'évolution présentée ici représente quelques 50 heures de calcul.

<i>Désignation</i>	<i>Résolution</i>	ν	D	Sc	Δt	$\frac{\ell}{L}$	n	T_{res}	Re_λ
128-I-1	128^3	0,75	0,75	1	$2, 2 \cdot 10^{-2}$	$\frac{1}{4}$	2	0,704	40,3
128-II-1	128^3	0,75	0,75	1	$3, 0 \cdot 10^{-2}$	$\frac{1}{4}$	2	0,960	40,8
128-II-0,5	128^3	0,75	1,5	0,5	$3, 0 \cdot 10^{-2}$	$\frac{1}{4}$	2	0,960	40,8
256-I-1	256^3	0,3	0,3	1	$2, 2 \cdot 10^{-2}$	$\frac{1}{4}$	2	0,704	67,3
256-II-1	256^3	0,3	0,3	1	$3, 0 \cdot 10^{-2}$	$\frac{1}{4}$	2	0,960	66,9
256-II-2,5	256^3	0,75	0,3	2,5	$3, 0 \cdot 10^{-2}$	$\frac{1}{4}$	2	0,960	39,9

TAB. 2.1 – Simulations Numériques Directes – Caractéristiques essentielles.

vitesse ne semble pas présenter de comportement auto-similaire. Il semble que cela corresponde à des observations expérimentales[54, 55, 7], qui font état d'une pente en $-\frac{5}{3}$ (ou plutôt légèrement inférieure, comme c'est le cas ici) pour le spectre du scalaire sans que le champ de vitesse puisse déjà présenter ces caractéristiques - et ce, pour des nombres de Reynolds R_λ de l'ordre de ceux obtenus ici[4]. Pour finir, la figure 2.9 présente les spectres normalisés pour trois nombres de Schmidt différents, mais avec la même injection de scalaire et le même champ de vitesse. On constate comme on pouvait s'y attendre que les spectres s'allongent quand le nombre de Schmidt (qui contrôle l'échelle la plus petite que les structures du scalaire puissent atteindre) augmente.

2.1.3 Étude théorique de la moyenne temporelle de la variance

Nous allons considérer le cas d'une injection de n cubes de côté ℓ dans le domaine cubique de côté L avec une période de Δt . Le scalaire est borné ($c \in [0, 1]$) et on injecte autant de cubes de 0 que de cubes de 1. Ainsi, la fraction du domaine qui se trouve renouvelée tous les Δt sera notée ϕ et s'écrit

$$\phi = \frac{n \ell^3}{L^3}.$$

L'injection est considérée comme fréquente, c'est à dire qu'on a $\Delta t \ll \tau_c$. Une telle condition implique que les caractéristiques statistiques de l'écoulement varient relativement peu entre deux injections. Cela explique l'évolution de la variance, figure 2.5, linéaire en temps entre deux injections. Tout se passe comme si la dissipation était assez peu affectée par l'injection. La fraction du domaine renouvelée ϕ est petite ($\frac{1}{64}$ du volume).

A l'instant t_0 , on suppose qu'une injection vient d'avoir lieu et on note σ_0^2 la valeur prise par la variance. Jusqu'à l'instant $t_0 + \Delta t$ auquel a lieu une autre injection, la variance évolue selon l'équation

$$\frac{d\sigma^2}{dt} = \langle \epsilon_c \rangle. \quad (2.1)$$

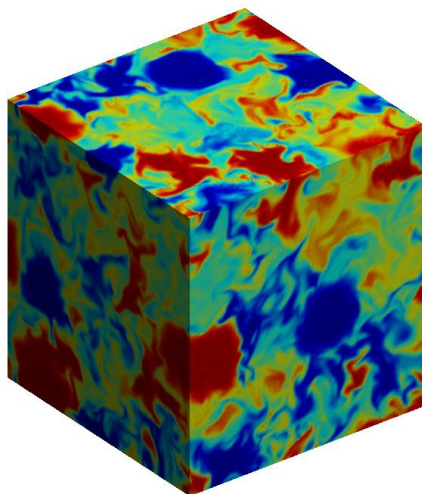


FIG. 2.1 – Vue du domaine cubique de simulation - on remarquera la périodicité des conditions aux limites et la présence de blocs de scalaires, trace d’une récente injection. Simulation 128-I-1.

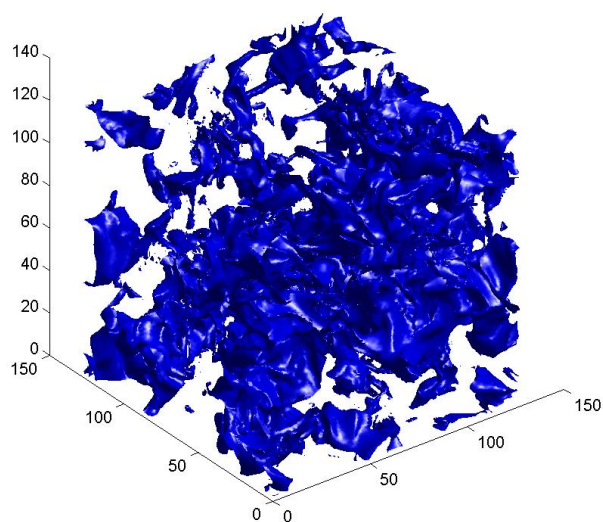


FIG. 2.2 – Isosurface $c = 0,75$ du scalaire - vue en 3 dimensions pour la simulation 128-I-1

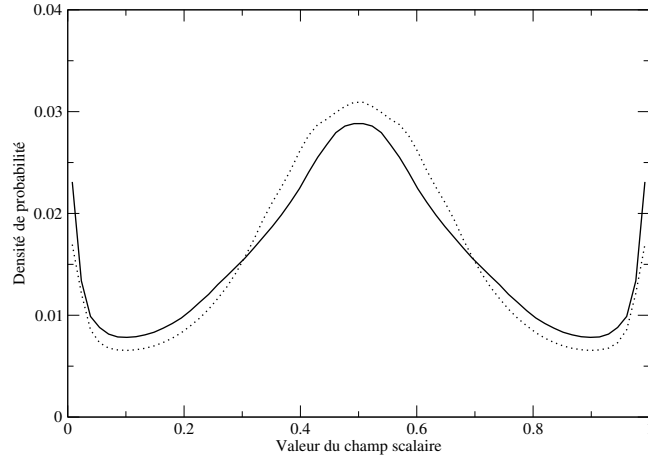


FIG. 2.3 – Densités de probabilité pour les simulations 256-I-1 (trait continu) et 256-II-1 (trait pointillé).

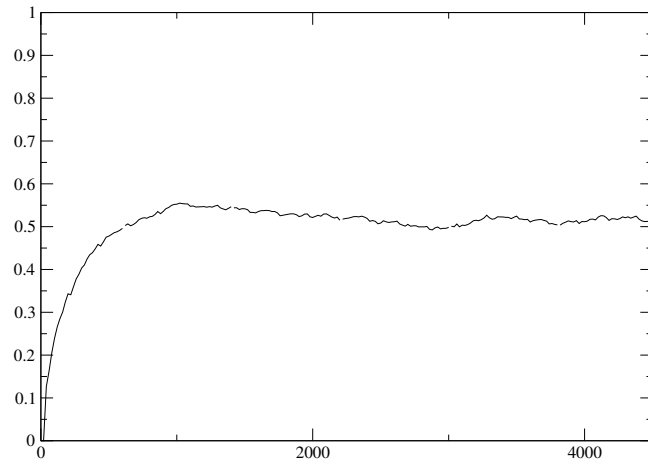


FIG. 2.4 – Évolution temporelle de l'écart-type σ aux temps longs. L'unité de temps est le pas de discrétisation. Un temps τ_c représente environ 700 pas.

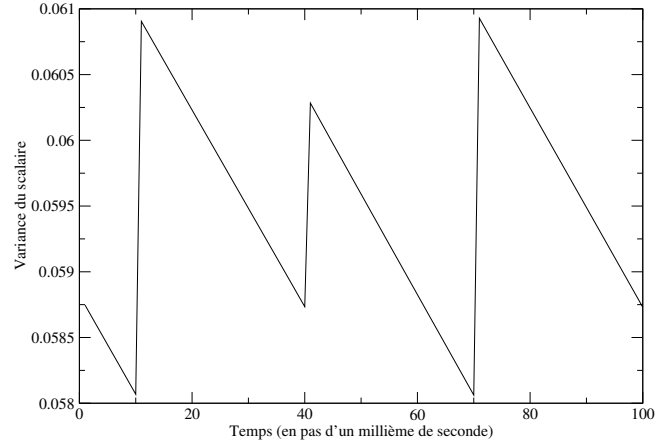


FIG. 2.5 – Évolution temporelle de la variance en situation stationnaire - agrandissement. L'injection se produit tous les 30 pas de temps.

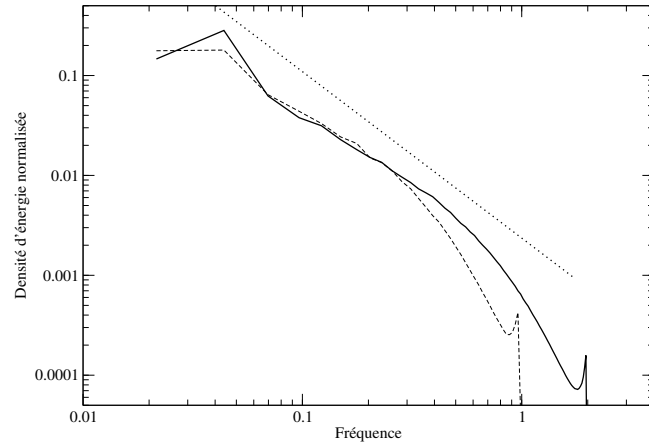


FIG. 2.6 – Spectres normalisés par la variance pour les simulations 128-II-1 (tirets) et 256-II-1 (ligne continue), ainsi que loi en $-\frac{5}{3}$ (pointillés). Les deux simulations connaissent la même injection et le même Schmidt, mais des valeurs de R_λ de 41 et 67 respectivement.

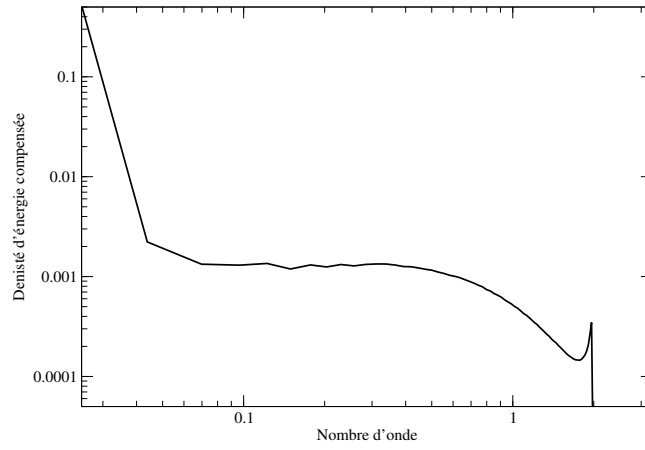


FIG. 2.7 – Spectre normalisé par la variance et compensé (multiplié par $k^{1,45}$) pour la simulation 256-I-1.

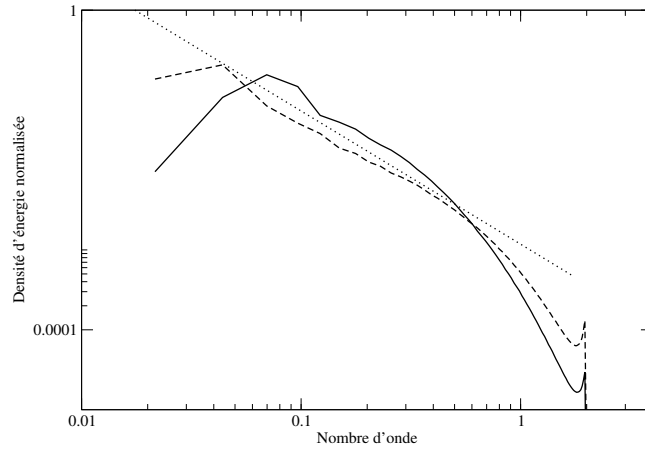


FIG. 2.8 – Spectres normalisés de la vitesse (ligne continue) ; du scalaire (tirets) ; spectre en $-\frac{5}{3}$ (pointillés) - simulation 256-I-1

On peut donc écrire

$$\sigma^2(t_0 + t) = \sigma_0^2 - \langle \epsilon_c \rangle t. \quad (2.2)$$

Lorsque se produit l'injection suivante, une fraction ϕ du domaine est renouvelée par du scalaire frais. En moyenne, la variance du scalaire dans la partie du domaine qui n'est pas renouvelée est de $\sigma^2(t_0 + \Delta t)$. Nous noterons γ la variance dans la partie renouvelée du domaine. Si on considère une injection parfaite où le champ vaut 0 ou 1 en tout point du domaine d'injection, la variance dans ce domaine est de $\frac{1}{4}$. En pratique le lissage de l'injection, évoqué lors de la présentation de la méthode, impose $\gamma < \frac{1}{4}$ de façon très nette.

La variance du champ scalaire juste après l'injection peut ainsi s'écrire

$$\sigma^2(t_0 + \Delta t) (1 - \phi) + \gamma \phi.$$

Si on considère maintenant qu'on a obtenu un régime périodique (plus que réellement stationnaire) alors la variance après injection est égale à la variance à t_0 , ce qui s'écrit

$$(\sigma_0^2 - \langle \epsilon_c \rangle) (1 - \phi) + \gamma \phi = \sigma_0^2. \quad (2.3)$$

La moyenne temporelle de la variance notée $\langle \sigma^2 \rangle_t$ peut alors être déterminée par

$$\begin{aligned} \langle \sigma^2 \rangle_t &= \frac{1}{\Delta t} \int_{t_0}^{t_0 + \Delta t} \sigma^2(t) dt \\ &= \sigma_0^2 - \frac{1}{2} \langle \epsilon_c \rangle \Delta t, \end{aligned}$$

et l'équation (2.3) devient

$$\frac{\Delta t \langle \sigma^2 \rangle_t}{\tau_c} = \phi \left(\gamma - \langle \sigma^2 \rangle_t + \frac{1}{2} \langle \epsilon_c \rangle \Delta t \right), \quad (2.4)$$

avec $\tau_c = \frac{\langle \sigma^2 \rangle_t}{\langle \epsilon_c \rangle}$. En pratique, les fluctuations de variance autour de sa moyenne sont relativement petites devant la variance elle-même. La quantité $\frac{1}{2} \langle \epsilon_c \rangle \Delta t$ représente ces fluctuations et elle peut être négligée devant la différence $\gamma - \langle \sigma^2 \rangle$. L'équation précédente permet alors d'écrire une relation entre la variance moyenne et la quantité $T_{res} = \frac{\Delta t}{\phi}$

$$\langle \sigma^2 \rangle_t = \frac{\gamma}{1 + \frac{T_{res}}{\tau_c}}. \quad (2.5)$$

Il faut signaler qu'un résultat absolument identique peut être obtenu en supposant qu'entre deux injections c'est τ_c qui reste constante et que la variance connaît donc une décroissance exponentielle.

La moyenne temporelle de la variance est donc directement déterminée par le rapport de temps $\frac{T_{res}}{\tau_c}$ au travers de la relation (2.5). Dans la limite d'une injection très importante (quand $T_{res} \ll \tau_c$) la variance vaut γ , c'est à dire que l'injection ne laisse pas le temps au scalaire d'être mélangé et que le domaine est renouvelé rapidement. La variance obtenue est alors celle des cubes d'injection.

<i>Simulation</i>	$\langle \sigma^2 \rangle_t$	$\langle \epsilon_c \rangle = \langle 2D (\partial_i c)^2 \rangle$	τ_c	γ^{-1}
128-I-1	0,0805	0,1032	0,780	6,53
128-II-1	0,06523	0,1015	0,643	6,16
128-II-0,5	0,05897	0,0953	0,619	6,65
256-I-1	0,06603	0,0860	0,768	7,90
256-II-1	0,05710	0,0761	0,750	7,68
256-II-2,5	0,06119	0,0696	0,879	7,81

TAB. 2.2 – Caractérisation de l'état de mélange

2.1.4 Résultats et conclusion à propos de la variance

Les résultats pour toutes les simulations sont présentés tableau 2.2 : la variance, la dissipation et le temps caractéristique du mélange ont été calculés, ainsi que l'estimation de γ^{-1} fournie pour chaque simulation par la formule

$$\gamma^{-1} = \frac{1}{\langle \sigma^2 \rangle_t \left(1 + \frac{T_{res}}{\tau_c}\right)}.$$

Les valeurs obtenues sont toutes relativement proches, car γ n'est a priori déterminé que par l'injection (taille des cubes et forme du lissage) - et toutes les simulations sont soumises au même type d'injection. La valeur moyenne de γ^{-1} obtenue vaut environ 7,12. Comme la relation (2.5) peut aussi s'écrire

$$\frac{1}{\langle \sigma^2 \rangle_t} = \gamma^{-1} \left(1 + \frac{T_{res}}{\tau_c}\right), \quad (2.6)$$

la figure 2.10 représente une droite de pente et d'ordonnée à l'origine de 7,12, ainsi que pour chaque simulation un point $\left(\frac{T_{res}}{\tau_c}, \frac{1}{\langle \sigma^2 \rangle_t}\right)$. Cela permet de vérifier que la relation (2.6) et par conséquent (2.5) sont grossièrement vérifiées par les simulations.

Si la concordance avec la relation (2.6) reste grossière, c'est sans doute parce que les variances présentées dans le tableau 2.2 ne sont pas exactement les variances moyennes caractéristiques de la situation de mélange imposée par l'injection. La figure 2.5 montre en effet que les fluctuations temporelles de variance ne sont pas complètement négligeables.

Cette étude montre que l'on peut comprendre assez simplement la relation qui lie les caractéristiques de l'injection à la situation de mélange obtenue. Notons au passage que le paramètre $\frac{T_{res}}{\tau_c}$ semble plus directement lié au degré de mélange que $\frac{T_{res}}{T_{turb}}$.

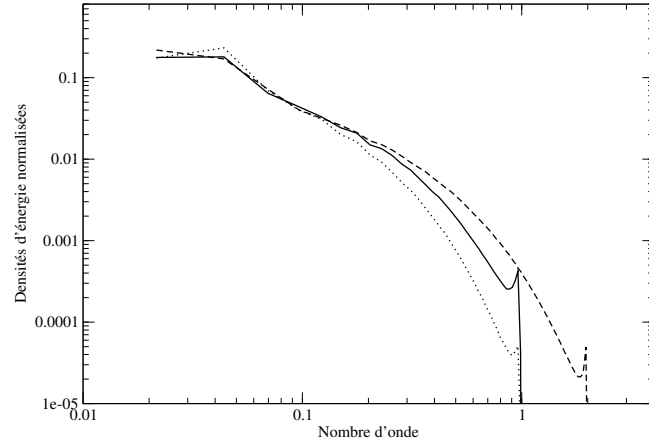


FIG. 2.9 – Spectres normalisés du champ scalaire pour les simulations 128-II-1 (ligne continue ; $Sc = 1$) ; 256-II-2,5 (tirets ; $Sc = 2,5$) ; 128-II-0,5 (pointillés ; $Sc = 0,5$). Les champs de vitesses sont tous identiques ($R_\lambda \simeq 40$).

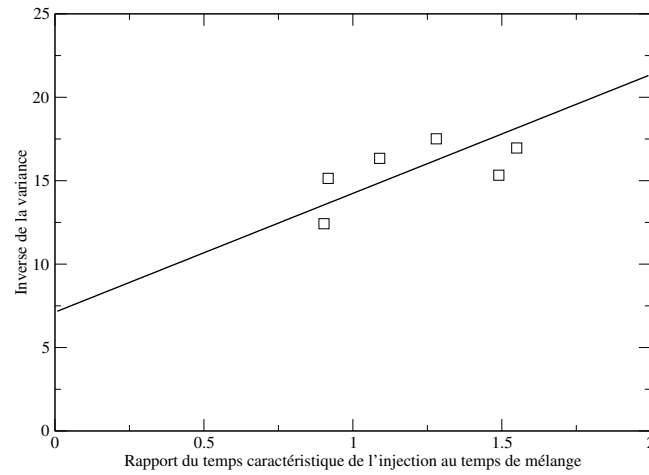


FIG. 2.10 – Tracé de $\frac{1}{\sigma^2}$ en fonction de $\frac{T_{res}}{\tau_c}$.

2.2 Étude de l'incrément du scalaire

L'incrément de scalaire, noté θ est défini par

$$\theta = c(\vec{x} + \vec{r}) - c(\vec{x}).$$

La densité de probabilité de cette quantité apporte des informations importantes sur le comportement du scalaire. Les fonctions de structure sont les moments $\langle \theta^n \rangle$ de cette densité de probabilité. D'après une analyse dimensionnelle à la Obukhov-Corrsin, ces quantités devraient se comporter pour les moments d'ordre pair comme

$$\langle \theta^n \rangle \sim r^{\frac{n}{3}},$$

ce qui revient à supposer une densité de probabilité gaussienne. Il est bien établi que ce n'est pas le cas pour de courtes distances (c'est un signe de ce qu'on appelle l'intermittence), mais on s'attend à ce que pour r suffisamment grand la densité de probabilité relaxe effectivement vers une gaussienne [4].

Différentes densités de probabilité de l'incrément sont présentées figure 2.12 et 2.11, pour différentes valeurs de r . Plus une densité est piquée plus elle correspond à une valeur de r petite. Il est évident que ces densités de probabilité sont assez éloignées des fonctions gaussiennes pour ces petites valeurs. Ce qui se passe précisément à des distances plus grandes sera étudié plus bas.

2.2.1 Fonctions de structure

Pour caractériser l'écart de la densité de probabilité des incréments à la gaussienne, il est d'usage de considérer les moments de cette densité, et d'y rechercher un comportement en loi de puissance. Ce comportement n'est pas nettement décelable pour une résolution de 128^3 , mais il le devient avec une résolution plus élevée. On peut constater sur la figure 2.13 que le moment d'ordre deux présente un comportement en loi de puissance relativement net : une fois tracé en coordonnées logarithmiques, une portion linéaire apparaît assez nettement. Une régression linéaire permet d'estimer l'exposant de la loi de puissance pour l'ordre deux par régression linéaire. Le résultat est présenté dans le tableau 2.3, et il est relativement proche de l'exposant théorique, $\frac{2}{3}$. Cet exposant est lié à la loi de puissance proche de $-\frac{5}{3}$ présentée par le spectre[3].

Les fonctions de structure d'ordre plus élevé présentent pour la même gamme d'échelle (jusqu'à la taille d'injection du scalaire, soit $\frac{L}{4}$) un comportement en loi de puissance (voir figure 2.14).

Pour évaluer les exposants de ces lois de puissance, on peut représenter les fonctions de structure d'ordre $2n$ avec $n \geq 2$ en fonction de la fonction de structure d'ordre 2. Les courbes obtenues, qui sont tracées figure 2.16 en coordonnées logarithmiques, présentent des portions linéaires beaucoup plus développées, dont il est aisé de déterminer la pente. Une fois l'exposant déterminé, il suffit de le multiplier par celui obtenu pour le moment d'ordre deux pour avoir une détermination de l'exposant de la fonction de structure d'ordre $2n$. Le tableau 2.3 présente l'ensemble des résultats obtenus par cette méthode, appelée *Extended Self-Similarity*.

Les exposants ainsi calculés semblent être tout à fait compatibles avec les résultats tant expérimentaux que numériques présentés dans l'article de revue de Warhaft[4] et qui laissaient présager une intermittence très importante du scalaire. La figure 2.16 représente les exposants obtenus pour les fonctions de structure d'ordre pair et permet de comparer à la prévision de la théorie d'Obukhov-Corrsin pour le scalaire.

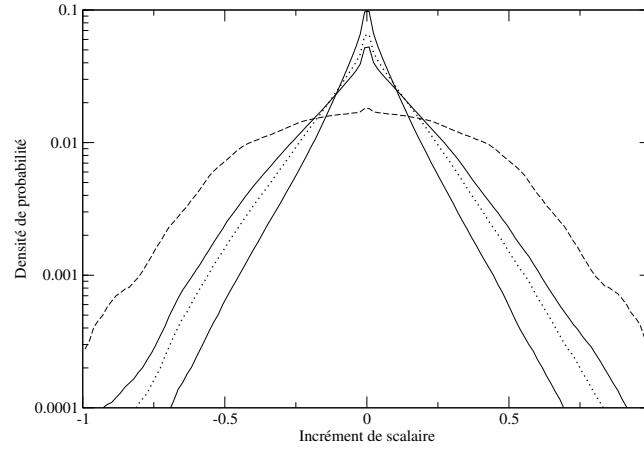


FIG. 2.11 – Densité de probabilité de l'incrément de scalaire en coordonnées semi-logarithmiques – elle s'étale quand r augmente. Les valeurs de $\frac{r}{L}$ choisies (L est la taille du domaine) sont $\frac{1}{64}$, $\frac{1}{32}$, $\frac{3}{32}$ et $\frac{1}{2}$. Simulation 256-I-1.

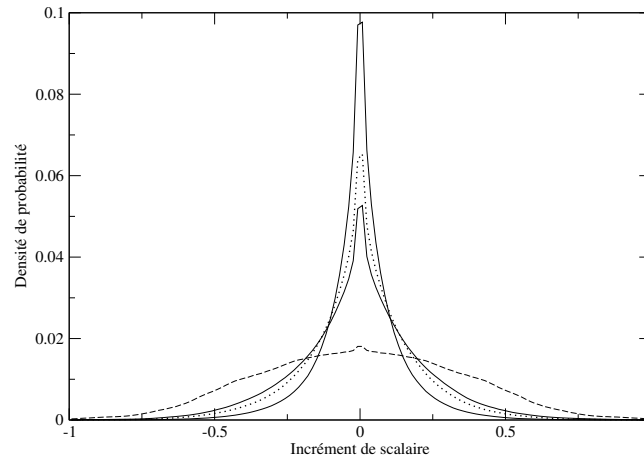


FIG. 2.12 – Densité de probabilité de l'incrément de scalaire en coordonnées linéaires – elle s'étale quand r augmente. Les valeurs de $\frac{r}{L}$ choisies (L est la taille du domaine) sont $\frac{1}{64}$, $\frac{1}{32}$, $\frac{3}{32}$ et $\frac{1}{2}$. Simulation 256-I-1.

<i>Fonction</i>	<i>Coeff. ESS⁵</i>	<i>Exposant</i>
$\langle \theta^2 \rangle$	—	0,689
$\langle \theta^4 \rangle$	1,55	1,07
$\langle \theta^6 \rangle$	1,91	1,32
$\langle \theta^8 \rangle$	2,14	1,47

TAB. 2.3 – Exposants des fonctions de structures du scalaire

Ces résultats confirment également qu’une zone inertielle se développe nettement pour le champ scalaire (pour une résolution de 256^3) alors qu’aucun comportement de ce type ne se dégage pour le champ de vitesse. On a déjà vu[4] des spectres qui ne présentaient pas de zone inertielle alors que les fonctions de structure suivaient des lois de puissance. La figure 2.17 montre que ce n’est pas le cas ici : pas plus que le spectre la fonction de structure d’ordre deux ne présente de loi de puissance. C’est également vrai pour les autres fonctions de structure, qui ne sont pas représentées ici.

2.2.2 Relaxation de la densité de probabilité

Lorsque r devient suffisamment grand, $c' = c(\vec{x} + \vec{r})$ et $c = c(\vec{x})$ sont considérés comme indépendants. Leur densité de probabilité conjointe $p_{c,c'}$ est telle que $p_{c,c'}(\alpha, \beta) = p(\alpha)p(\beta)$, où p est la densité de probabilité du scalaire (nous sommes dans une situation statistiquement homogène). Dans ce cas, la densité de probabilité de l’incrément de scalaire, p_θ s’écrit

$$\begin{aligned}
p_\theta(\Gamma) &= \langle \delta(\Gamma - \theta) \rangle \\
&= \int \delta(\Gamma - (\alpha - \beta)) p_{c,c'}(\alpha, \beta) \, d\alpha \, d\beta \\
&= \int \delta(\Gamma - (\alpha - \beta)) p(\alpha) p(\beta) \, d\alpha \, d\beta \\
&= \int p(\alpha) p(\Gamma + \alpha) \, d\alpha.
\end{aligned}$$

Cette expression est une convolution[56]. Lorsque la densité de probabilité du scalaire est une gaussienne, la densité de probabilité des incréments de scalaire relaxe logiquement vers une gaussienne. Connaissant la densité de probabilité du scalaire dans notre cas (voir figure 2.3), il ne faut pas s’attendre à voir celle des incréments relaxer vers une gaussienne. Le kurtosis est présenté figure 2.18 et il semble tendre vers une valeur inférieure à 3 (valeur obtenue avec une densité de probabilité gaussienne). Ceci confirme immédiatement l’analyse ci-dessus. La figure 2.19 permet de comparer la densité de probabilité de l’incrément obtenue pour la plus grande distance possible et la quantité $\int p(\alpha) p(\Gamma + \alpha) \, d\alpha$. On constate tout d’abord que les deux courbes sont très proches l’une de l’autre, et ensuite que $\int p(\alpha) p(\Gamma + \alpha) \, d\alpha$ n’est effectivement pas une gaussienne.

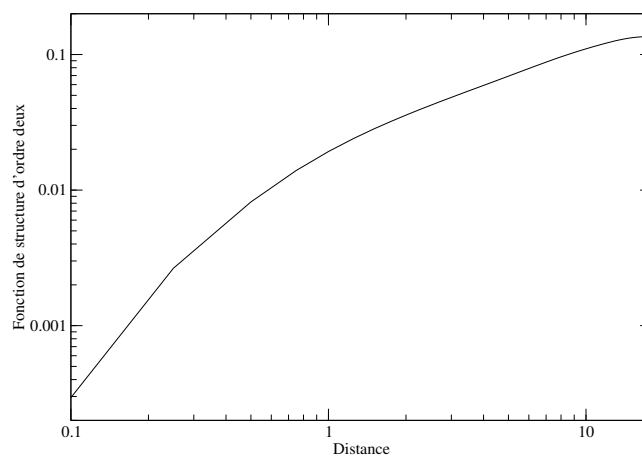


FIG. 2.13 – Tracé de $\langle \theta^2 \rangle$ en fonction de r – Une nette portion linéaire apparaît.

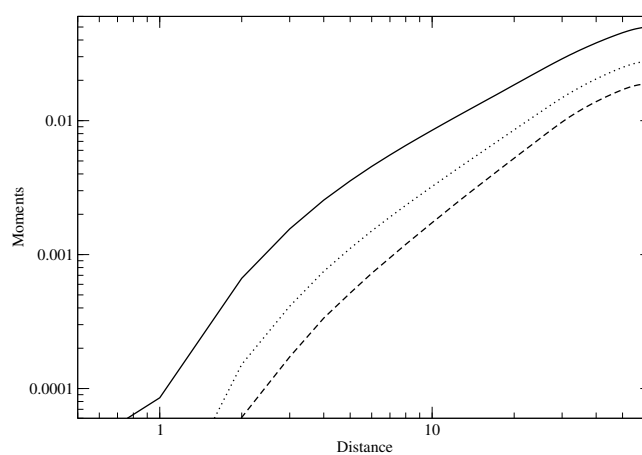


FIG. 2.14 – Moments d'ordre 4 (ligne continue), 6 (pointillés) et 8 (tirets) en fonction de r

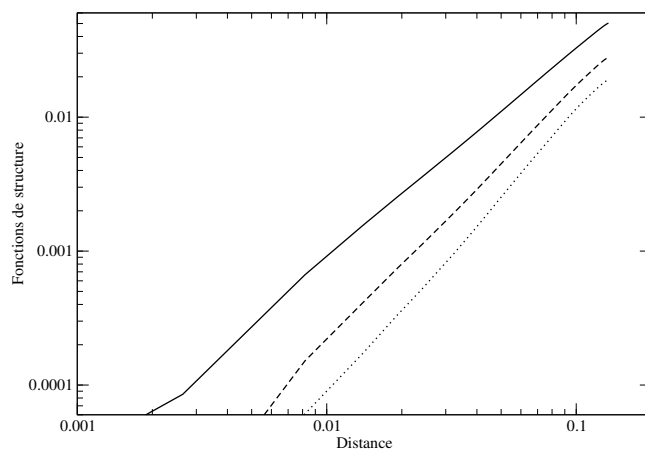


FIG. 2.15 – Moments d'ordre 4 (ligne continue) , 6 (tirets) et 8 (pointillés) en fonction du moment d'ordre deux

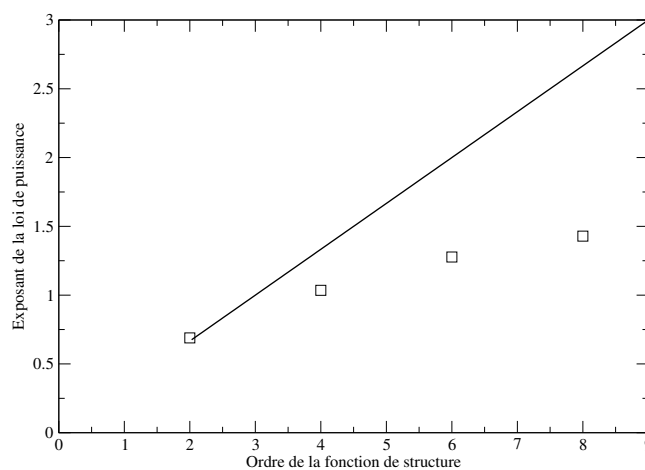


FIG. 2.16 – Moments de l'incrément du scalaire (carrés); Prévisions de la théorie KOC (ligne continue)

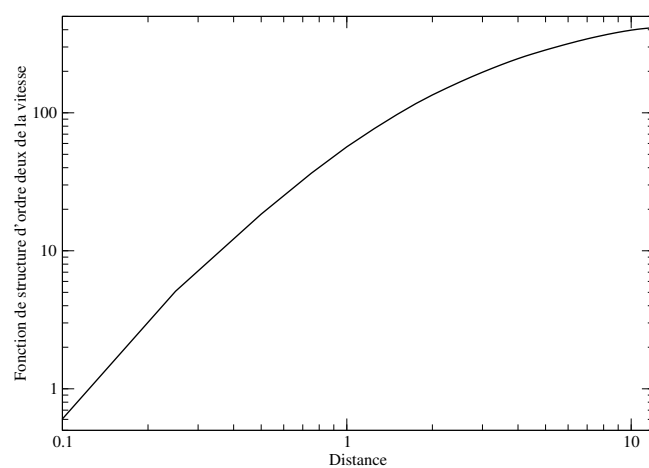


FIG. 2.17 – Tracé de $\langle(\delta u)^2\rangle$ en fonction de r – Aucune loi de puissance ne se dégage.

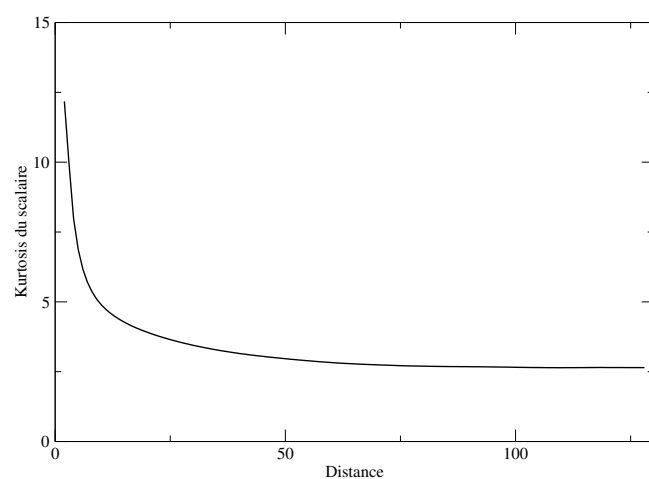


FIG. 2.18 – Kurtosis de l'incrément en fonction de la distance

2.3 Diffusion conditionnelle

La diffusion conditionnelle, comme nous l'avons vu au paragraphe 1.3, est une quantité centrale pour comprendre l'évolution de la densité de probabilité du scalaire. Estimer cette quantité permet de fermer l'équation de la densité de probabilité. Le modèle LMSE étant la première tentative de fermeture de cette équation - et sans doute la plus simple -, il sert ici de référence.

2.3.1 Diffusion conditionnelle affine

L'expression du modèle de Dopazo est lié à la constante de temps caractéristique du mélange car elle s'écrit

$$\langle D \Delta c | \Gamma \rangle = \frac{1}{2 \tau_c} (\Gamma - \langle c \rangle)$$

avec⁶

$$\tau_c = \frac{\sigma_c^2}{\langle \epsilon_c \rangle}.$$

Notons ici que dès qu'on suppose une forme affine pour la diffusion conditionnelle, alors il est aisé de retrouver le coefficient $\frac{1}{2 \tau_c}$. Pour cela, nous allons établir une première relation entre la dissipation moyenne et la diffusion conditionnelle :

$$\begin{aligned} \langle \epsilon_c \rangle &= 2 D (\partial_i c)^2 \\ &= -2 \langle D c \Delta c \rangle \\ &= - \int 2 \langle D c \Delta c | \Gamma \rangle p(\Gamma) d\Gamma \\ &= - \int 2 \Gamma \langle D \Delta c | \Gamma \rangle p(\Gamma) d\Gamma \end{aligned}$$

La relation s'écrit ainsi

$$\langle \epsilon_c \rangle = -2 \int \Gamma \langle D \Delta c | \Gamma \rangle p(\Gamma) d\Gamma. \quad (2.7)$$

Cette relation exacte permet de saisir le lien direct entre la diffusion conditionnelle et la dissipation totale. Une des conséquences essentielles est qu'il peut être pertinent de normaliser la diffusion conditionnelle par la dissipation totale.

Nous allons maintenant faire l'hypothèse que la diffusion conditionnelle est une fonction affine de Γ . Nous noterons $-\alpha$ la pente de cette fonction. L'hypothèse s'écrit donc

$$\langle D \Delta c | \Gamma \rangle \simeq -\alpha (\Gamma - \langle c \rangle). \quad (2.8)$$

En utilisant la relation (2.7), il vient

$$\begin{aligned} \langle \epsilon_c \rangle &= \int 2 \alpha \Gamma (\Gamma - \langle c \rangle) p(\Gamma) d\Gamma \\ &= 2 \alpha \left[\int \Gamma^2 p(\Gamma) d\Gamma - \int \Gamma \langle c \rangle p(\Gamma) d\Gamma \right] \\ &= 2 \alpha [\langle c^2 \rangle - \langle c \rangle^2] \\ &= 2 \alpha \sigma_c^2. \end{aligned}$$

⁶Pour mémoire, j'ai choisi de prendre $\epsilon_c = 2 D (\partial_i c)^2$.

Ce qui s'écrit en définitive

$$\alpha = \frac{\langle \epsilon_c \rangle}{2 \sigma_c^2}$$

La méthode pour retrouver le modèle de Dopazo exposée ici est simple, et montre que l'hypothèse principale du modèle est la forme affine de la diffusion. Le travail de Dopazo donne en plus des éléments d'explication à propos du développement de la portion linéaire.

2.3.2 Résultats numériques

Dans ce paragraphe, nous ne considérerons que des diffusions conditionnelles normalisées, c'est à dire multipliées par $2 \tau_c$, la pente théorique du modèle LMSE. Ceci afin de pouvoir comparer les diffusions conditionnelles pour différents états de mélange et différentes valeurs de D ou de ν .

Portion linéaire et injection

En regardant en détail les résultats de Marc Elmo, on peut être amené à penser que la diffusion conditionnelle laisse souvent apparaître une portion linéaire importante. Les résultats présentés figure 2.20 confirment cette impression.

Plus précisément, une portion linéaire semble apparaître pour les faibles taux d'injection, ce qui laisse présager d'une certaine universalité car ce sont les situations qui dépendent *a priori* le moins de la technique d'injection. La figure 2.21 permet de voir comment pour une injection plus forte (simulation 256-I-1) la portion linéaire semble légèrement déstabilisée.

Dans des conditions de forçage très important, le travail de Marc Elmo montre que la portion linéaire est même indécélable.

Forme de la diffusion conditionnelle

La forme de la diffusion conditionnelle semble être communes à toutes les situations étudiées ici. Elle ne dépend donc pas du taux d'injection, dans la mesure où celui-ci reste relativement faible. C'est ce que montre la figure 2.22 où sont présentées ensemble toutes les diffusions normalisées. Le tableau 2.4 confirme cette analyse de façon quantitative : il présente le rapport entre la pente de la portion linéaire et la pente du cas où la diffusion conditionnelle est complètement affine. La constance de ce rapport pour des états de mélanges divers ou des nombres de Schmidt variés est frappante.

Cette forme caractéristique est peut-être à rapprocher de celle obtenue par Dopazo[57]. En considérant un champ scalaire unidimensionnel initialement égal à $H(x)$ convecté par un champ de vitesse aléatoire de la forme $\alpha(t)x$, il obtient pour la diffusion conditionnelle normalisée la fonction

$$\langle \Delta c | \Gamma \rangle \sim \text{erf}^{-1}(2\Gamma - 1) e^{-[\text{erf}^{-1}(2\Gamma - 1)]^2}$$

à un facteur près. Celle-ci est présentée figure 2.23, et la ressemblance avec la forme de la diffusion conditionnelle obtenue précédemment est nette.

2.4 Conclusion

La résolution la plus grande (256^3 , et un nombre de Reynolds R_λ de l'ordre de 66) des simulations a donc permis de mettre en évidence de façon indiscutable une zone auto-similaire pour

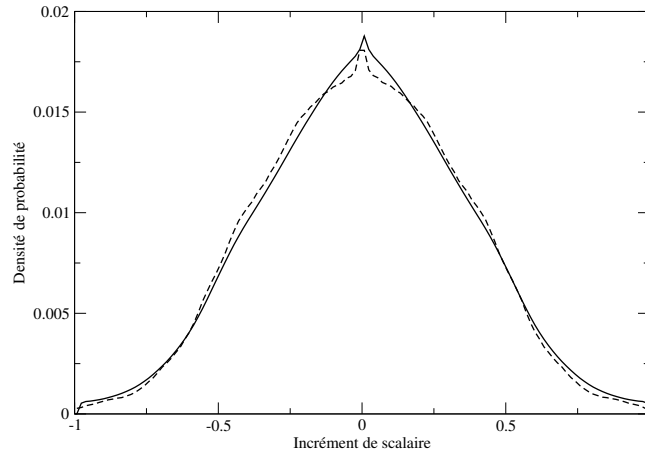


FIG. 2.19 – Comparaison entre la densité de probabilité des incréments (tirets) et la convolution de la densité de probabilité du scalaire (ligne continue).

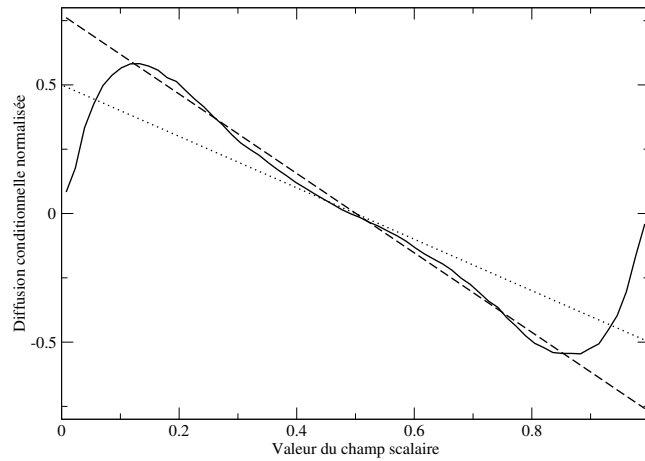


FIG. 2.20 – Diffusion conditionnelle pour la simulation 256-II-1 (ligne continue) - LMSE (pointillés) - Pente $1,5 \frac{1}{2\tau}$ (tirets)

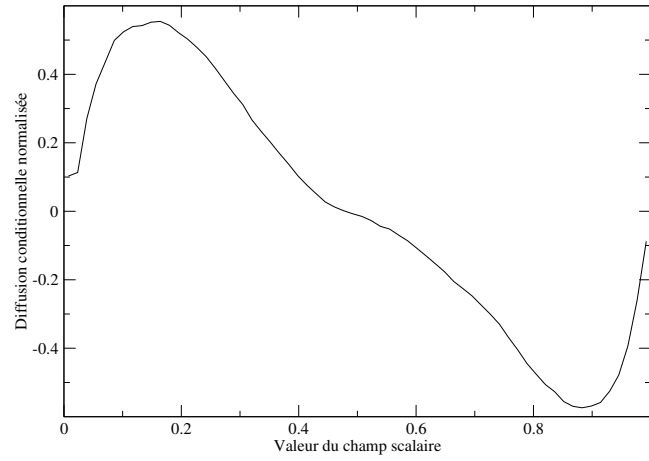


FIG. 2.21 – Diffusion conditionnelle pour l'injection la plus importante (simulation 256-I-1)

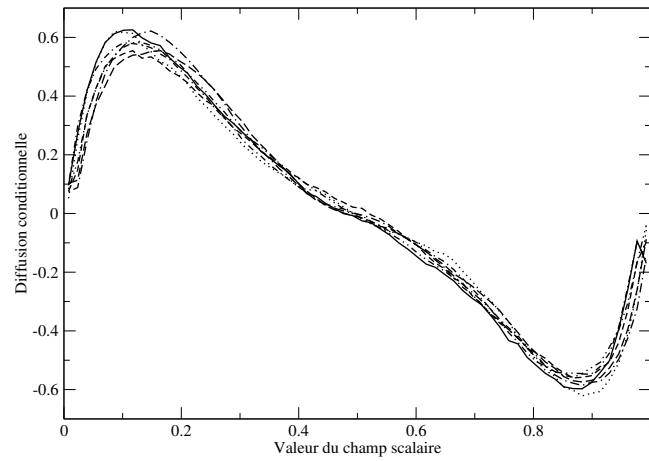


FIG. 2.22 – Diffusions conditionnelles normalisées par $\frac{1}{2\tau_c}$ pour toutes les simulations

<i>Simulation</i>	$\frac{1}{2\tau_c}$	Pente : p	$2\tau_c p$
128-I-1	0,641	0,974	1,52
128-II-1	0,778	1,159	1,49
128-II-0,5	0,808	1,284	1,59
256-I-1	0,651	1,035	1,59
256-II-1	0,667	0,987	1,48
256-II-2,5	0,569	0,842	1,48

TAB. 2.4 – Pente de la diffusion conditionnelle et rapport à $\frac{1}{2\tau_c}$.

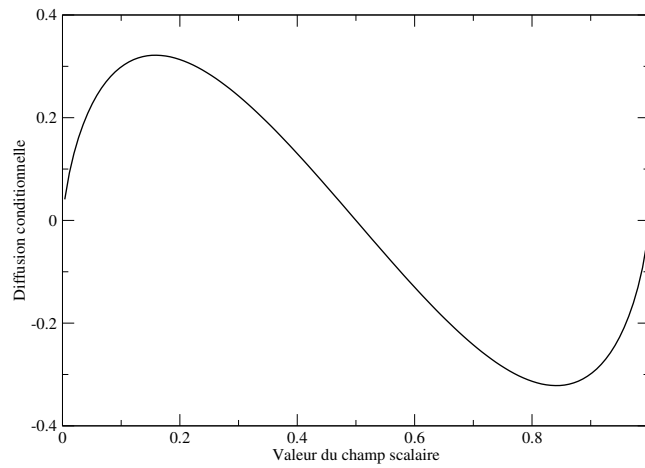


FIG. 2.23 – Diffusion conditionnelle obtenue par Dopazo[57]

le scalaire. Ce comportement se caractérise par des lois de puissance pour le spectre et pour les fonctions de structure du scalaire. Il est par contre frappant que le champ de vitesse, pour la même gamme d'échelle, ne présente aucun comportement auto-similaire - ni pour le spectre, qui ne présente pas le comportement en $-\frac{5}{3}$ attendu en zone inertielle, ni pour les fonctions de structure. Cet état des choses a déjà été rapporté en turbulence de grille, il n'est donc pas étonnant en soi. Mais ceci confirme que les simulations numériques directes reproduisent correctement certains des comportements des écoulements réels. Et également que la technique d'injection n'altère pas lesdits comportements.

La densité de probabilité des incréments relaxe effectivement à grande distance vers $\int p(\alpha) p(\Gamma + \alpha) d\alpha$, comme on pouvait s'y attendre. Cette dernière quantité n'est cependant pas une gaussienne. Ceci est dû à la méthode d'injection, et fournit un exemple de comportement altéré par l'injection.

Concernant la diffusion conditionnelle, on peut conclure qu'il semble courant de voir se développer une portion linéaire - ce qui est conforté par le cas de l'injection par gradient moyen. On pourrait imaginer qu'on arrête subitement l'injection. Dans ce cas, l'existence d'une portion linéaire signifie que la partie centrale de la densité de probabilité évoluerait (au moins aux temps courts) de façon auto-similaire. L'existence de portion linéaire confirme qu'un régime auto-similaire peut être attendu pour le scalaire, et permet de penser que le modèle LMSE a des fondements raisonnables.

La diffusion conditionnelle présente de plus un comportement caractéristique, et il faut sans doute y voir là aussi un effet de l'injection. Le cas théorique de Dopazo et l'injection de scalaire telle qu'elle est pratiquée dans les simulations ont un point commun qui explique sans doute pourquoi les diffusions conditionnelles se ressemblent : l'injection de scalaire dans les deux cas se fait par des blocs de valeurs extrêmes. A l'intérieur de tels blocs la diffusion est faible même si elle est nécessairement non nulle ⁷. Les blocs sont en effet homogènes. Il n'est ainsi pas étonnant de trouver que c'est pour les valeurs extrêmes de Γ que la diffusion prend des valeurs plus faibles. Une telle injection correspond tout à fait à l'idée que l'on peut se faire d'une injection réelle de scalaire (colorant, réactif). Cela pourrait éventuellement justifier le remplacement d'une diffusion conditionnelle purement affine par une diffusion conditionnelle plus "réaliste".

⁷Si on il n'y aurait pas de mélange.

<i>Simulation</i>	e_c	ϵ	τ	$\frac{\tau_c}{\tau}$	Sc
128-I-1	283,4	396,5	0,715	1,09	1
128-II-1	294,6	416,5	0,707	0,91	1
128-II-0,5	289,5	410,7	0,705	0,88	0,5
256-I-1	299,8	397,0	0,755	1,02	1
256-II-1	293,6	385,2	0,762	0,98	1
256-II-2,5	280,3	394,0	0,711	1,24	2,5

TAB. 2.5 – Caractérisation du champ de vitesse

Chapitre 3

Analyse des modèles de sous-maille

Une modélisation est avant tout une approximation, un estimateur d'une quantité, qui commet donc implicitement une erreur. Nous allons voir qu'en choisissant l'erreur quadratique pour mesurer la qualité d'un modèle, il est possible de définir un estimateur optimal - qui approche au mieux une quantité à partir de paramètres fixés. Ces notions sont utilisées pour étudier le modèle de Cook et Riley pour $\overline{f(c)}$, couramment utilisé en simulation des grandes échelles. Nous verrons que les estimateurs optimaux permettent de sonder *séparément* les différentes hypothèses qui sous-tendent ce modèle, et ainsi de mieux comprendre son fonctionnement.

3.1 Estimateur optimal

Nous considérons dans ce paragraphe le problème de l'estimation de la quantité $\overline{f(c)}$ lorsque f n'est pas connue *a priori*. Nous avons vu que la démarche adoptée par Cook et Riley[19] revient à essayer d'approcher une quantité donnée, $\overline{f(c)}$, avec des fonctions de quelques paramètres fondamentaux. Dans le cas du modèle de Cook et Riley sans sous-modèle, les paramètres fondamentaux sont de façon évidente $\overline{c}$ et σ_s^2 , tandis qu'il s'agit de $\overline{c}$ et α lorsqu'on utilise le sous-modèle qu'ils proposent. De façon plus générale, nous allons considérer qu'il s'agit d'estimer une quantité $\overline{f(c)}$ grâce à une fonction d'un jeu de paramètres fondamentaux π . Les considérations que nous allons mener ici se généralisent bien entendu à l'estimation de toute autre quantité que $\overline{f(c)}$ et sont valables dans un autre cadre que celui de la simulation des grandes échelles.

3.1.1 Erreur quadratique

Pour dire si un estimateur est optimal, il faut pouvoir comparer deux estimateurs, et donc disposer d'une mesure de leur qualité. Celle-ci est fournie par l'erreur quadratique commise par un estimateur $g(\pi)$ de $\overline{f(c)}$, qui s'écrit

$$E_g = \left\langle \left(\overline{f(c)} - g(\pi) \right)^2 \right\rangle. \quad (3.1)$$

Cette erreur est nécessairement positive ou nulle. Ainsi, plus elle est petite, meilleur est l'estimateur. Comme on l'a vu au paragraphe 1.4.3, l'erreur quadratique est couramment utilisée en modélisation de sous-maille comme mesure de la qualité d'un modèle.

Estimation par une constante

Considérons tout d'abord le cas où l'on cherche à estimer la quantité $\overline{f(c)}$ par une constante. Dans ce cas, l'erreur commise par l'estimateur constant C peut s'écrire

$$\begin{aligned} E_C &= \left\langle \left(\overline{f(c)} - C \right)^2 \right\rangle \\ &= \left\langle \left(\overline{f(c)} - \langle \overline{f} \rangle \right)^2 \right\rangle + \left\langle \left(\langle \overline{f} \rangle - C \right)^2 \right\rangle - 2 \left\langle \overline{f(c)} - \langle \overline{f} \rangle \right\rangle \left(\langle \overline{f} \rangle - C \right) \\ &= \left\langle \left(\overline{f(c)} - \langle \overline{f} \rangle \right)^2 \right\rangle + \left\langle \left(\langle \overline{f} \rangle - C \right)^2 \right\rangle. \end{aligned}$$

Cette égalité signifie que la valeur de la constante qui minimise l'erreur quadratique est la moyenne de $\overline{f(c)}$. La variance de $\overline{f(c)}$ apparaît ici comme l'erreur que l'on commet en prenant la moyenne de $\overline{f(c)}$ comme estimateur.

Les erreurs commises par les différents estimateurs ne seront par la suite jamais présentées telles quelles. Elles seront toujours normalisées par la variance de la quantité que les estimateurs tentent d'approcher. De cette façon, on peut se faire une idée de la qualité absolue d'un estimateur. Il va en effet de soi qu'en estimant une quantité dont la variance est très grande, on risque de commettre une erreur importante. Alors qu'inversement, l'erreur commise en estimant une quantité dont la variance est très faible a toutes les chances d'être petite. Normaliser l'erreur par la variance a pour but de corriger cette distorsion. Les erreurs commises deviennent ainsi comparables.

Les erreurs normalisées ont pour ordre de grandeur l'unité voire moins. En effet, la variance représente une erreur importante, puisque c'est celle commise par un estimateur très fruste (la moyenne de la quantité). Un estimateur commettant une erreur normalisée supérieure à l'unité est donc considéré comme très mauvais.

3.1.2 Estimateur optimal pour $\overline{f(c)}$

Proposition 3.1 *L'erreur commise par un estimateur $g(\pi)$ de $\overline{f(c)}$ vérifie l'égalité suivante :*

$$\left\langle \left(\overline{f(c)} - g(\pi) \right)^2 \right\rangle = \left\langle \left(\overline{f(c)} - \left\langle \overline{f(c)} \middle| \pi \right\rangle \right)^2 \right\rangle + \left\langle \left(\left\langle \overline{f(c)} \middle| \pi \right\rangle - g(\pi) \right)^2 \right\rangle.$$

◇ *Preuve.*

$$\begin{aligned} E_g &= \left\langle \left(\overline{f(c)} - \left\langle \overline{f(c)} \middle| \pi \right\rangle \right)^2 \right\rangle + \left\langle \left(\left\langle \overline{f(c)} \middle| \pi \right\rangle - g(\pi) \right)^2 \right\rangle \\ &\quad - 2 \cdot \left\langle \left(\left\langle \overline{f(c)} \middle| \pi \right\rangle - \overline{f(c)} \right) \cdot \left(\left\langle \overline{f(c)} \middle| \pi \right\rangle - g(\pi) \right) \right\rangle \end{aligned}$$

Or on peut écrire

$$\begin{aligned}
& \langle (\langle \bar{f} | \pi \rangle - \bar{f}) \cdot (\langle \bar{f} | \pi \rangle - g(\pi)) \rangle \\
&= \int (\langle \bar{f} | \pi \rangle - \bar{f}) \cdot (\langle \bar{f} | \pi \rangle - g(\pi)) p(\pi, \bar{f}) d\pi d\bar{f} \\
&= \int d\pi p(\pi) (\langle \bar{f} | \pi \rangle - g(\pi)) \int (\langle \bar{f} | \pi \rangle - \bar{f}) p(\bar{f} | \pi) d\bar{f} \\
&= \int d\pi p(\pi) (\langle \bar{f} | \pi \rangle - g(\pi)) \left(\langle \bar{f} | \pi \rangle - \int \bar{f} p(\bar{f} | \pi) d\bar{f} \right)
\end{aligned}$$

et comme par définition $\langle \bar{f} | \pi \rangle = \int \bar{f} p(\bar{f} | \pi) d\bar{f}$, le terme précédent est nul. $\square$

La propriété 3.1 montre que l'erreur commise par un estimateur $g(\pi)$ de $\overline{f(c)}$ peut se décomposer en la somme de deux erreurs. La première erreur est l'**erreur irréductible pour le jeu de paramètres** π , et elle s'écrit

$$\left\langle \left(\overline{f(c)} - \langle \overline{f(c)} | \pi \rangle \right)^2 \right\rangle. \quad (3.2)$$

On remarquera que cette quantité ne dépend pas de g mais seulement du jeu de paramètre π . La seconde est l'**erreur supplémentaire commise par g** :

$$\left\langle \left(\langle \overline{f(c)} | \pi \rangle - g(\pi) \right)^2 \right\rangle. \quad (3.3)$$

Cette quantité est positive ou nulle. Tout estimateur $g(\pi)$ commet donc une erreur au moins égale à l'erreur irréductible pour π , ce qui justifie le qualificatif d'irréductible.

Cette seconde erreur s'annule si et seulement si $g(\pi) = \langle \overline{f(c)} | \pi \rangle$ quelle que soit la valeur prise par les paramètres. Par conséquent, il existe un seul et unique estimateur qui minimise l'erreur quadratique, et il s'agit de la moyenne de $\overline{f(c)}$ conditionnée par π . On peut ainsi définir un estimateur optimal pour le problème qui nous intéresse :

Définition 1 L'estimateur optimal de $\overline{f(c)}$ pour le jeu de paramètres π est la fonction

$$\langle \overline{f(c)} | \pi \rangle$$

3.1.3 Propriétés des estimateurs optimaux

Les propriétés détaillées ici ne sont pas des plus fondamentales. Elles permettent cependant de se forger une intuition du comportement des estimateurs optimaux.

Estimateur optimal et nuage de points

Il est courant d'illustrer la qualité d'un estimateur en reproduisant un nuage de points sur un diagramme. Pour chaque estimation, on place un point dont l'ordonnée est la valeur de la quantité f à estimer, et l'abscisse la valeur de l'estimateur $g(\pi)$. La propriété suivante permet de dire que si l'estimateur est optimal (indépendamment du jeu de paramètres π) alors le barycentre des points situés sur une même verticale (à $g(\pi) = G$ fixé) est porté par la droite $x = y$.

Proposition 3.2 *Soit f une quantité à approcher, et $g(\pi)$ une fonction à valeur dans R . Si $g(\pi) = \langle f | \pi \rangle$, alors*

$$\langle f | g(\pi) = G \rangle = G.$$

◇ *Preuve.*

La densité de probabilité de g est reliée à celle de π par la relation

$$p(G) = \int p(G, \pi) \, d\pi = \int p(\pi) \, \delta(G - g(\pi)) \, d\pi.$$

Dans le cas où g est inversible (*i.e.* où chaque valeur des paramètres donne une valeur différente pour $g(\pi)$, ce qui n'est en général pas du tout vrai), alors on a l'égalité $p(g(\pi)) = p(\pi)$.

On peut aussi relier la moyenne d'une quantité g à estimer conditionnée par la valeur g prise par la fonction $g(\pi)$, à la moyenne de g conditionnée par la valeur prise par π

$$\begin{aligned} \langle f | g(\pi) = G \rangle p(G) &= \int f \, p(f, G) \, df \\ &= \int f \, p(f, \pi) \, \delta(g(\pi) - G) \, df \, d\pi \\ &= \int \delta(g(\pi) - G) \, d\pi \int f \, p(f, \pi) \, df \\ &= \int \delta(g(\pi) - G) \langle f | \pi \rangle \, p(\pi) \, d\pi \end{aligned}$$

Dans le cas où $g(\pi) = \langle f | \pi \rangle$,

$$\begin{aligned} \langle f | G \rangle p(G) &= \int \delta(\langle f | \pi \rangle - G) \langle f | \pi \rangle \, p(\pi) \, d\pi \\ &= G \int \delta(\langle f | \pi \rangle - G) \, p(\pi) \, d\pi \\ &= G p(G) \end{aligned}$$

D'où $g(\pi) = \langle f | \pi \rangle \Rightarrow \langle f | G \rangle = G$. □

Hierarchie des jeux de paramètres

Si l'on dispose d'un jeu de paramètres π_1 , tout jeu de paramètre $\pi_2 = f(\pi_1)$ construit à partir de π_1 risque d'être moins pertinent. C'est ce que traduit la propriété suivante :

Proposition 3.3 *Soient π_1 et π_2 deux jeux de paramètres, $g(\pi_2)$ une fonction à valeur dans R et f la quantité à estimer. Alors*

$$\langle (f - \langle f | \pi_1, \pi_2 \rangle)^2 \rangle \leq \langle (f - \langle f | \pi_1, g(\pi_2) \rangle)^2 \rangle.$$

◇ *Preuve.*

L'erreur commise par le premier estimateur, $\langle f | \pi_1, \pi_2 \rangle$, s'écrit aussi

$$\langle (f - \langle f | \pi_1, \pi_2 \rangle)^2 \rangle = \int (\langle f | \pi_1, \pi_2 \rangle - f)^2 \, p(\pi_1, \pi_2, f) \, d\pi_1 \, d\pi_2 \, df$$

Et celle commise par le second peut se développer ainsi

$$\begin{aligned}
& \langle (f - \langle f | \pi_1, g(\pi_2) \rangle)^2 \rangle \\
&= \int (\langle f | \pi_1, g \rangle - f)^2 p(\pi_1, g) d\pi_1 dg df \\
&= \int (\langle f | \pi_1, g \rangle - f)^2 p(\pi_1, \pi_2, f) \delta(g - g(\pi_2)) d\pi_1 dg d\pi_2 df \\
&= \int (\langle f | \pi_1, g(\pi_2) \rangle - f)^2 p(\pi_1, \pi_2, f) d\pi_1 d\pi_2 df \\
&= \int d\pi_2 d\pi_1 \int (\langle f | \pi_1, g(\pi_2) \rangle - f)^2 p(\pi_1, \pi_2, f) df \\
&= \int d\pi_2 d\pi_1 p(\pi_1, \pi_2) \int \{ (\langle f | \pi_1, \pi_2 \rangle - f)^2 + (\langle f | \pi_1, \pi_2 \rangle - \langle f | \pi_1, g(\pi_2) \rangle)^2 \\
&\quad - 2(\langle f | \pi_1, \pi_2 \rangle - f)(\langle f | \pi_1, \pi_2 \rangle - \langle f | \pi_1, g(\pi_2) \rangle) \} p(f | \pi_1, \pi_2) df \\
&= \langle (f - \langle f | \pi_1, \pi_2 \rangle)^2 \rangle + \int (\langle f | \pi_1, \pi_2 \rangle - \langle f | \pi_1, g(\pi_2) \rangle)^2 p(\pi_1, \pi_2) d\pi_2 d\pi_1
\end{aligned}$$

La différence entre les deux erreurs est donc le terme

$$\int (\langle f | \pi_1, \pi_2 \rangle - \langle f | \pi_1, g(\pi_2) \rangle)^2 p(\pi_1, \pi_2) d\pi_2 d\pi_1,$$

qui est positif de façon évidente. La propriété est donc bien vérifiée. $\square$

3.1.4 Estimateur optimal de la densité de sous-maille

Pour approcher la densité de sous-maille en utilisant un jeu de paramètre π , on utilise une fonction $h(\Gamma, \pi)$. Cependant, on ne peut pas définir aussi simplement que précédemment une erreur commise par h sur l'estimation de d_s . Pour comparer quantitativement la qualité de deux estimateurs de la densité de sous-maille, il faut se donner une fonction f et comparer les erreurs des estimateurs de $\overline{f(c)}$ générés par les deux approximations différentes de la densité de sous-maille. L'estimateur de $\overline{f(c)}$ généré en prenant h comme approximation de la densité de sous-maille s'écrit

$$e_{f,h}(\pi) = \int f(\Gamma) h(\Gamma, \pi) d\Gamma. \quad (3.4)$$

Si on considère deux estimateurs quelconques de la densité de sous-maille utilisant un même jeu de paramètres π , il n'est pas possible de façon générale de dire que l'un d'entre eux générera de meilleurs estimateurs de $\overline{f(c)}$ quelle que soit la fonction f considérée. Cependant, si on prend $h(\Gamma, \pi) = \langle d_s(\Gamma) | \pi \rangle$, et qu'on regarde l'estimateur généré au moyen de la relation 3.4, on obtient l'égalité suivante, valable quelle que soit f :

$$\langle \overline{f(c)} | \pi \rangle = \left\langle \int f(\Gamma) d_s(\Gamma) d\Gamma \middle| \pi \right\rangle = \int f(\Gamma) \langle d_s(\Gamma) | \pi \rangle d\Gamma. \quad (3.5)$$

On sait que les estimateurs générés par $\langle d_s(\Gamma) | \pi \rangle$ sont tous optimaux au sens de l'erreur quadratique, quelle que soit f . Il est donc possible de dire que quelle que soit la fonction f considérée, cet estimateur génère de meilleurs estimateurs de $\overline{f(c)}$ que les autres. On peut considérer la définition suivante en guise de conclusion de ce raisonnement :

Définition 2 L'estimateur optimal de la densité de sous-maille $d_s(\Gamma)$ pour un jeu de paramètres π s'écrit

$$\langle d_s(\Gamma) | \pi \rangle$$

Unicité

La propriété suivante garantit qu'il n'existe pas d'autre estimateur de la densité de sous-maille qui présente la propriété de générer des estimateurs optimaux de $\overline{f(c)}$.

Proposition 3.4 *L'estimateur optimal de la densité de sous-maille pour un jeu de paramètres π donné est unique.*

◇ *Preuve.*

Supposons qu'il existe une fonction $h(\Gamma, \pi)$ telle que $\forall f, \int f(\Gamma) h(\Gamma, \pi) d\Gamma = \langle \overline{f(c)} | \pi \rangle$, alors pour toute valeur de π , et pour tout n entier,

$$\int \Gamma^n h(\Gamma, \pi) d\Gamma = \int \Gamma^n \langle d_s(\Gamma) | \pi \rangle d\Gamma.$$

Cela signifie que pour cette valeur des paramètres, tous les moments des deux fonctions sont égaux. Les deux fonctions sont donc égales pour toute valeur de Γ , et comme cette propriété est vraie pour toute valeur de π , elle est vraie pour toute valeur de π et de Γ . □

Moyenne et variance de l'estimateur optimal

La moyenne de l'estimateur optimal à π fixé est obtenue en prenant $f(c) = c$ dans la relation (3.5)

$$\begin{aligned} \int \Gamma \langle d_s(\Gamma) | \pi \rangle d\Gamma &= \left\langle \int \Gamma d_s(\Gamma) d\Gamma \middle| \pi \right\rangle \\ &= \langle \bar{c} | \pi \rangle. \end{aligned}$$

Par la suite, nous considérerons que le paramètre $\bar{c}$ fait partie intégrante de π quelque soit le jeu de paramètres considéré. La quantité $\bar{c}$ est en effet le paramètre le plus naturel. Dans ce cas,

$$\langle \bar{c} | \pi \rangle = \bar{c}.$$

Cela revient à dire que l'on considère que la moyenne de l'estimateur optimal est simplement $\bar{c}$.

Par contre, σ_s^2 ne fait en général pas partie des paramètres, car il s'agit d'une quantité qui ne se déduit pas directement du champ grande échelle du scalaire. La variance de l'estimateur optimal de la densité à π fixé est donnée par

$$\begin{aligned} \int (\Gamma - \bar{c})^2 \langle d_s(\Gamma) | \pi \rangle d\Gamma &= \left\langle \int (\Gamma - \bar{c})^2 d_s(\Gamma) d\Gamma \middle| \pi \right\rangle \\ &= \langle \sigma_s^2 | \pi \rangle. \end{aligned}$$

Dans le cas où la variance de sous-maille fait partie du jeu de paramètres, c'est à dire qu'elle est supposée connue, on a bien sûr : $\langle \sigma_s^2 | \pi \rangle = \sigma_s^2$.

Nouvelle expression de l'estimateur optimal

Il existe une autre écriture pour l'estimateur optimal. En effet,

$$\begin{aligned}
 \langle d_s(\Gamma) | \pi \rangle &= \left\langle \overline{\delta(\Gamma - c(\vec{x} + \vec{r}))} \middle| \pi \right\rangle \\
 &= \frac{1}{\Delta^3} \int d\vec{r} \langle \delta(\Gamma - c(\vec{x} + \vec{r})) | \pi \rangle \\
 &= \frac{1}{\Delta^3} \int d\vec{r} \frac{1}{p(\pi)} \int \delta(\Gamma - c(\vec{x} + \vec{r})) p(c(\vec{x} + \vec{r}) = C, \pi) dC \\
 &= \frac{1}{\Delta^3} \int d\vec{r} p(c(\vec{x} + \vec{r}) = \Gamma | \pi) \\
 &= \overline{p(\Gamma, \vec{r} | \pi)}
 \end{aligned}$$

où $p(\Gamma, \vec{r} | \pi)$ est la densité de probabilité pour que le scalaire prenne en $\vec{x} + \vec{r}$ la valeur Γ , sachant la valeur des paramètres π pris en $\vec{x}$. Le filtrage porte bien évidemment ici sur la variable $\vec{r}$ seule.

Propriété de reconstruction

De façon évidente, l'estimateur optimal pour la densité de sous-maille vérifie la relation

$$\int \langle d_s(\Gamma) | \pi \rangle p(\pi) d\pi = \langle d_s(\Gamma) \rangle = \overline{p(\Gamma)}, \quad (3.6)$$

qui n'est autre que la propriété de reconstruction. Ainsi lorsqu'on cherche à minimiser l'erreur commise par un estimateur, on fait également en sorte qu'il vérifie la propriété de reconstruction, si bien que les deux critères usuels de qualité d'un modèle sont intimement liés. Tout estimateur optimal vérifie cependant cette propriété. En somme, le fait qu'un estimateur donné vérifie la propriété de reconstruction ne nous renseigne que sur la qualité de cet estimateur à π donné, mais pas sur la pertinence du jeu de paramètres lui-même. Le test de reconstruction ne permet pas de comparer des estimateurs optimaux, et de conclure ainsi quant à la pertinence des jeux de paramètres qu'ils font intervenir. Pour finir, il faut également remarquer que les estimateurs optimaux ne sont probablement pas les seuls à pouvoir réussir le test de reconstruction. Ainsi dans le cas homogène, il est aisé de voir que la quantité $p(c | \pi)$, qui est bien différente de l'estimateur optimal, vérifie également la propriété de reconstruction.

3.2 Modélisation par une loi β

Nous allons voir comment les estimateurs optimaux peuvent être utilisés pour jauger finement les hypothèses qui sous-tendent un modèle. Nous allons plus particulièrement étudier à quel point la modélisation de la densité de sous-maille par une loi β (voir 1.4.3) est fondée (**hypothèse β**).

L'ensemble des résultats présentés ici le sont pour un filtrage porte physique, comme dans le premier chapitre, et pour une taille de filtrage correspondant à 8 mailles élémentaires, soit une taille caractéristique $\overline{\Delta} \simeq 16 \eta$. Les résultats semblent cependant peu dépendre de la taille de filtre choisie. Afin d'illustrer l'effet des différents filtrages utilisés dans cette étude, les spectres des champs scalaires c et $\bar{c}$ sont représentés figure 3.1 pour des données issues de la simulation 128-I-1.

Ce sont d'ailleurs les données issues de cette simulation qui ont été utilisées dans ce chapitre. La figure montre également le spectre de $\widehat{\bar{\epsilon}}$, qui fait intervenir le filtre test utilisé dans le sous-modèle de Cook et Riley pour la variance de sous-maille.

3.2.1 Comparaisons directes

Les distributions β étant utilisées pour approcher directement les densités de sous-maille, il est naturel de vouloir directement comparer les deux. Des densités de sous-maille présentant la même moyenne et la même variance ont été calculées grâce à la propriété 1.7 et représentées sur la figure 3.2. Elles sont comparées à une distribution β de même moyenne et de même variance. On constate que les densités de sous-maille n'ont pas toutes la même forme, et qu'il est par conséquent difficile de jauger ainsi de la pertinence de la distribution β . Pour savoir quel est le meilleur choix de fonction pour approcher la densité de sous-maille connaissant seulement $\bar{\epsilon}$ et σ_s^2 , il faut considérer la moyenne des densités présentées précédemment, c'est à dire l'estimateur optimal de la densité de sous-maille pour le jeu de paramètres $\{\bar{\epsilon}, \sigma_s^2\}$. C'est cet estimateur qu'il faut comparer à une distribution β de même moyenne $\bar{\epsilon}$, et de même variance σ_s^2 .

Dans le cas d'un jeu de paramètres π quelconque (contenant cependant $\bar{\epsilon}$), la meilleure fonction susceptible d'approcher la densité de sous-maille est l'estimateur optimal pour π , soit $\langle d_s(\Gamma) | \pi \rangle$. Pour savoir alors si le choix des distributions β est fondé, il faut comparer cet estimateur optimal à une distribution β ayant strictement la même moyenne et la même variance que l'estimateur optimal. Les résultats du paragraphe 3.1.4 permettent de dire qu'il faut prendre $\bar{\epsilon}$ pour moyenne, et $\langle \sigma_s^2 | \pi \rangle$ pour variance. La distribution β à laquelle il faut comparer l'estimateur optimal s'écrit ainsi

$$\beta(\Gamma; \bar{\epsilon}, \langle \sigma_s^2 | \pi \rangle).$$

Il faut remarquer que cela revient à prendre pour estimateur de la variance de sous-maille le meilleur qui soit pour le jeu de paramètres π considéré. Par la suite, nous verrons ce qu'implique de prendre un estimateur de σ_s^2 qui ne soit pas optimal.

Méthode de calcul des densités de sous-maille et des estimateurs

Pour calculer l'estimateur optimal de la densité de sous-maille, il faut tout d'abord isoler des noeuds de la grille physique pour lesquels les paramètres π prennent des valeurs déterminées (à une certaine incertitude près). Par exemple, si $\bar{\epsilon}$ et σ_s^2 sont les paramètres, on cherche tous les points de la grille pour lesquels $\bar{\epsilon} = \bar{\epsilon}_0 \pm \delta\bar{\epsilon}$ et $\sigma_s^2 = \sigma_{s0}^2 \pm \delta\sigma_s^2$. Chacun de ces points est appelé **événement**.

Une fois un événement repéré en $\vec{x}$, on tire dans le cube $D(\vec{x})$ aléatoirement quelques points, on relève la valeur du champ en ces points. On calcule la densité de probabilité de cette variable aléatoire. Cela revient à calculer la densité de sous-maille pour chaque événement (grâce à la propriété 1.7), puis à moyenner sur tous les événements ¹.

¹Dans l'exemple précédent, pour s'assurer qu'on a choisi $\delta\bar{\epsilon}$ et $\delta\sigma_s^2$ suffisamment petits, on les diminue pour voir si le résultat est altéré. S'il ne l'est pas, on peut considérer le résultat comme fiable. Il faut également remarquer qu'on a pas besoin que les statistiques aient convergé pour chaque événement (c'est à dire qu'on ait suffisamment approché la densité de sous-maille en ce point) pour que avoir un estimateur optimal convergé. En pratique une dizaine de points par événement peut suffire si on en dispose de plusieurs milliers.

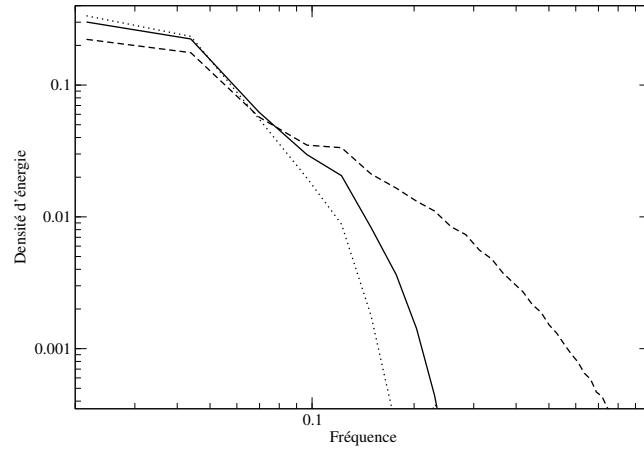


FIG. 3.1 – Spectres des champs c (tirets), $\bar{c}$ (ligne continue) et $\hat{c}$ (pointillés).

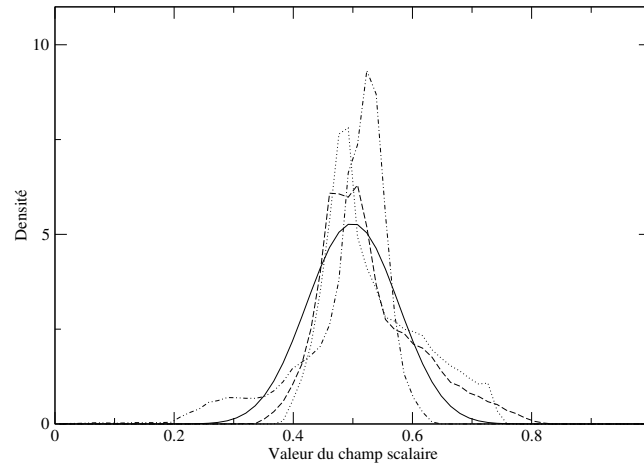


FIG. 3.2 – Comparaison entre une distribution β et des densités de sous-maille de même moyenne ($\bar{c} = 0,5$) et variance ($\sigma_s^2 = 0,0055$)

Lorsque les paramètres fondamentaux sont par exemple $\bar{c}$ et α , il faut aussi calculer $\langle \sigma_s^2 | \bar{c}, \alpha \rangle$. Ceci est fait en calculant $\sigma_s^2 = \overline{c^2} - \bar{c}^2$ pour chaque événement, et en moyennant sur tous les événements.

Il faut souligner que les courbes obtenues par cette méthode sont très sensibles à la technique d'interpolation utilisée. Il a été souligné précédemment que l'on disposait dans le cadre du schéma pseudo-spectral de simulation d'une interpolation "naturelle". On peut en effet considérer que le schéma donne en fait par les modes de Fourier accès à la valeur du champ scalaire en tous les points de l'espace physique. Le temps de calcul nécessaire à l'obtention de ces valeurs est cependant important². Il est tentant d'utiliser une méthode d'interpolation plus légère entre les points de la grille donnée par la TFD. Les estimateurs ainsi calculés ne sont cependant pas aussi proches des distributions β . Ceci pour souligner la sensibilité des résultats obtenus.

Variance de sous-maille connue

On considère dans ce paragraphe que les paramètres fondamentaux sont $\bar{c}$ et σ_s^2 , que l'on suppose connus. Il s'agit donc de comparer l'estimateur optimal de la densité de sous-maille à une distribution β de moyenne $\bar{c}$ et de variance σ_s^2 :

$$\beta(c; \bar{c}, \sigma_s^2).$$

L'étude des fonctions β montre que lorsque leur moyenne est proche de 0,5 et leur variance réduite, elles ressemblent beaucoup à des gaussiennes de même moyenne et de même variance (voir l'annexe C). Étant donné le rôle central que les gaussiennes sont souvent amenées à jouer en statistique et dans l'étude du scalaire passif, on pouvait s'attendre à une ressemblance des fonctions β et des estimateurs optimaux dans ce cas. La figure 3.3 montre que c'est tout à fait le cas : la correspondance entre les deux courbes est frappante. En prenant une moyenne plus proche des bornes, on ne fait pas sensiblement varier ni l'allure générale des fonctions β encore relativement proches des gaussiennes, ni d'après la figure 3.4 le bon accord entre la fonction β et l'estimateur optimal.

Par contre, les fonctions β présentent parfois des allures plus exotiques. Elles peuvent en effet diverger en 0 et/ou en 1. On pouvait alors s'attendre à ce que quand $\bar{c}$ se rapproche des bornes, la correspondance avec les estimateurs optimaux ne soit pas aussi bonne que précédemment. La comparaison directe présentée figure 3.5 montre finalement que si.

La comparaison directe des fonctions β et des estimateurs optimaux de la densité de sous-maille a donc montré une ressemblance très nette des deux estimateurs, au delà de toute attente. Et ceci alors que les densités réelles ne présentent pas de ressemblance particulière avec les distributions β comme le montre la figure 3.2. Il me semble possible que de tels résultats reflètent une propriété mathématique des lois β - que pour l'instant j'ignore. Quoi qu'il en soit, **les fonctions β semblent être un excellent choix quand la variance de sous-maille est supposée connue.**

Un exemple de reconstruction utilisant les distributions β lorsque la variance de sous-maille est supposée connue est présenté sur la figure 3.6. Il montre que la densité de probabilité du scalaire est convenablement reconstruite, notamment près des bornes. Ce type de résultat a été particulièrement discuté par Elmo *et al.*[50].

²Sur un AMD à 1,3 GHz il faut une dizaine d'heures de calcul pour obtenir un estimateur optimal par la méthode décrite ici.

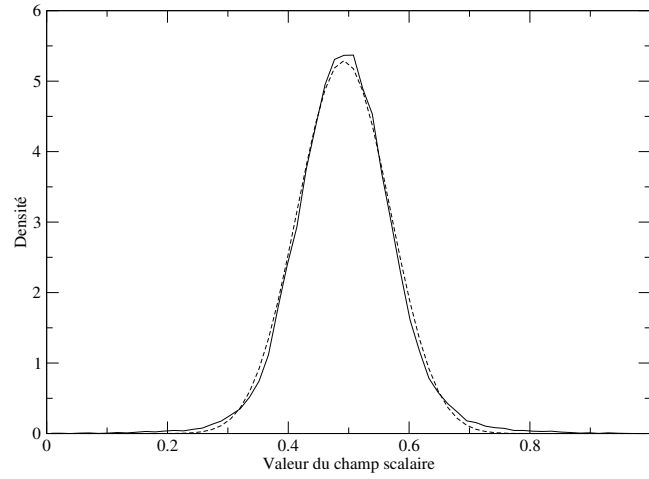


FIG. 3.3 – Comparaison entre $\beta(\Gamma; \bar{c}, \sigma_s^2)$ (trait continu) et $\langle d_s(\Gamma) | \bar{c}, \sigma_s^2 \rangle$ (pointillé) pour $\bar{c} = 0,5 \pm 0,01$ et $\sigma_s^2 = 0,0055 \pm 0,0001$.

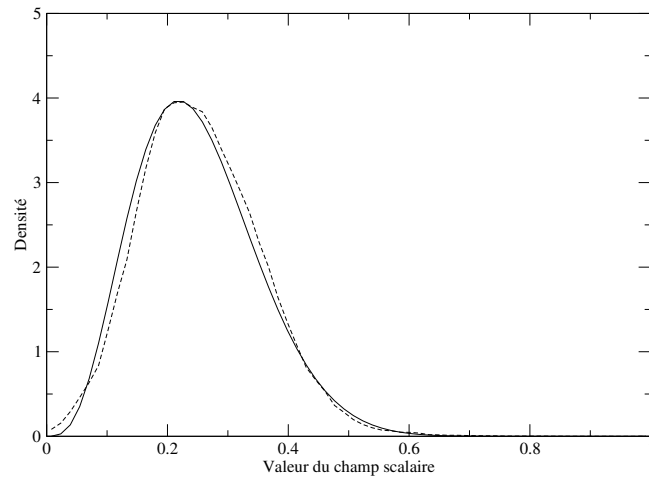


FIG. 3.4 – Comparaison entre $\beta(\Gamma; \bar{c}, \sigma_s^2)$ (trait continu) et $\langle d_s(\Gamma) | \bar{c}, \sigma_s^2 \rangle$ (pointillé) pour $\bar{c} = 0,25 \pm 0,01$ et $\sigma_s^2 = 0,01 \pm 0,001$.

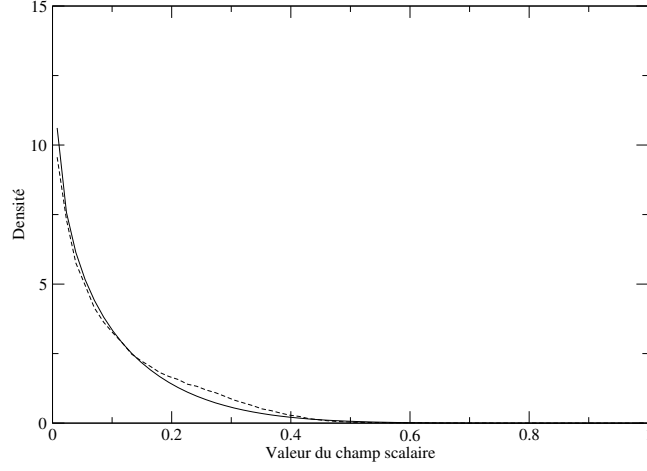


FIG. 3.5 – Comparaison entre $\beta(\Gamma; \bar{c}, \sigma_s^2)$ (trait continu) et $\langle d_s(\Gamma) | \bar{c}, \sigma_s^2 \rangle$ (pointillé) pour $\bar{c} = 0, 1 \pm 0, 01$ et $\sigma_s^2 = 0, 01 \pm 0, 001$.

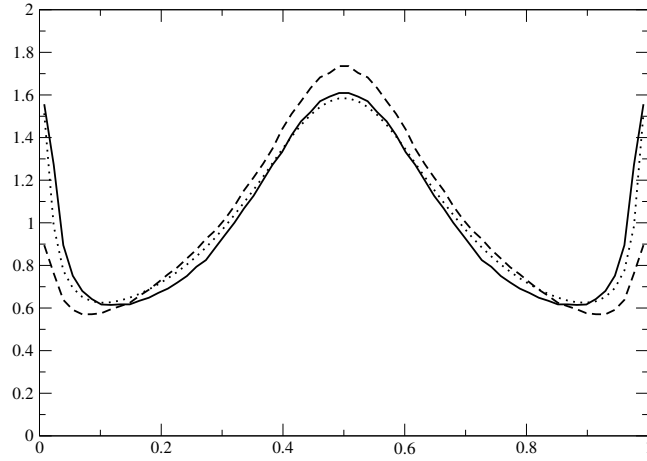


FIG. 3.6 – Reconstruction (pointillés) de la densité de probabilité du scalaire (ligne continue) et comparaison à la densité de probabilité de $\bar{c}$, qui permet de voir que le filtrage diminue l'importance des pics près des bornes, mais que ceux-ci sont bien reconstruits.

Variance de sous-maille estimée

Lorsque Cook et Riley utilisent un sous-modèle pour estimer σ_s^2 grâce à α , cela revient du point de vu des estimateurs optimaux à faire un changement dans les paramètres fondamentaux - qui deviennent alors $\bar{c}$ et α . Le meilleur estimateur de la densité de sous-maille devient donc $\langle d_s(\Gamma) | \bar{c}, \alpha \rangle$, qu'il va falloir comparer à une distribution β de même moyenne ($\bar{c}$) et même variance ($\langle \sigma_s^2 | \bar{c}, \alpha \rangle$)

$$\beta(\Gamma; \bar{c}, \langle \sigma_s^2 | \bar{c}, \alpha \rangle) \quad (3.7)$$

En choisissant pour variance $\langle \sigma_s^2 | \bar{c}, \alpha \rangle$, on suppose qu'on connaît l'estimateur optimal de la variance de sous-maille pour le jeu de paramètres considéré, et qu'on dispose ainsi d'un sous-modèle idéal. C'est ce qui permet de séparer les conséquences de l'hypothèse β et celles du sous-modèle. La validité du sous-modèle de Cook et Riley sera étudiée plus loin.

Quand la variance de sous-maille est supposée connue, l'estimateur optimal est alors la moyenne de fonctions (différentes densités de sous-maille) qui ont toutes la même moyenne et la même variance. Ici, comme la variance de sous-maille est inconnue, l'estimateur optimal est la moyenne de fonctions de même moyenne, mais de variances différentes. Cela explique sans doute que les fonctions β soient des candidates moins parfaites, qu'elles prennent des allures de gaussienne (figure 3.9), qu'elles soient un peu plus excentrées (figure 3.8) ou qu'elles divergent (figure 3.7).

Les figures 3.10 et 3.11 montrent quelques densités réelles comparées à des lois bêta de la forme (3.7). Dans le cas où la moyenne est proche de 0, 5, on constate bien que comme c'est le paramètre α qui contrôle statistiquement la situation, les densités réelles présentent des variances différentes. Cette dispersion de la variance de sous-maille est ce qui explique que les densités réelles soient plus éloignées de la fonction bêta que dans le cas où σ_s^2 est connue (voir figure 3.2).

Cependant, il faut noter que dans le cas où la moyenne s'approche des bornes, et bien que les densités n'aient pas vraiment l'allure de distributions bêta (voir figure 3.11), elles en sont cependant relativement proches. Cela explique que l'estimateur optimal ne soit en définitive pas très éloigné des lois bêta dans ce cas particulier (voir figure 3.7).

3.2.2 Erreurs comparées

Les comparaisons directes ne sont cependant qu'un test partiel du bien-fondé de l'hypothèse de Cook et Riley. Pour avoir une vision plus précise et surtout quantitative de ce qu'entraîne cette hypothèse, il convient d'examiner l'erreur supplémentaire dont elle est la source dans l'estimation de $\overline{f(c)}$ - pour différentes fonctions f . Il va donc falloir comparer l'estimateur optimal de $\overline{f(c)}$ à l'estimateur de $\overline{f(c)}$ généré par l'hypothèse d'une fonction β comme estimateur optimal de la variance de sous-maille (3.7). Cet estimateur, noté $A_f(\pi)$ pour signifier qu'il permet d'estimer $\overline{f(c)}$ à partir d'un jeu de paramètre π s'écrit

$$A_f(\pi) = \int f(\Gamma) \beta(\Gamma; \bar{c}, \langle \sigma_s^2 | \pi \rangle) d\Gamma \quad (3.8)$$

Cet estimateur est difficile à calculer numériquement, surtout dans le cas relativement courant où la fonction β diverge en 0 ou en 1. Cependant, pour des fonctions f particulières (polynômes surtout), les estimateurs A_f prennent une forme analytique beaucoup plus simple.

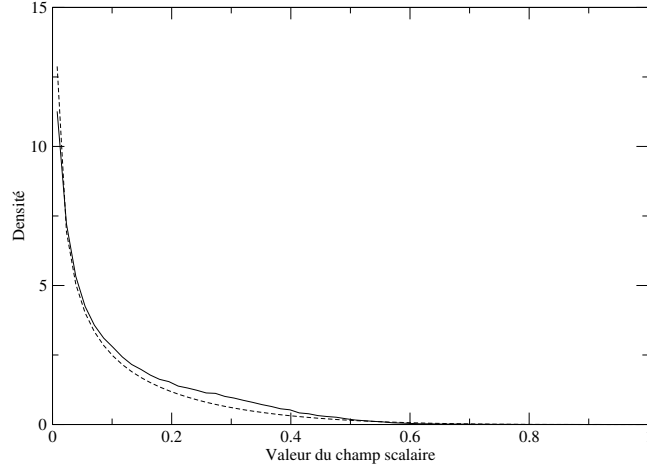


FIG. 3.7 – Comparaison entre $\beta(\Gamma; \bar{c}, \langle \sigma_s^2 | \bar{c}, \alpha \rangle)$ (pointillé) et $\langle d_s(\Gamma) | \bar{c}, \alpha \rangle$ (trait continu) pour $\bar{c} = 0,1 \pm 0,01$ et $\alpha = 0,01 \pm 0,001$.

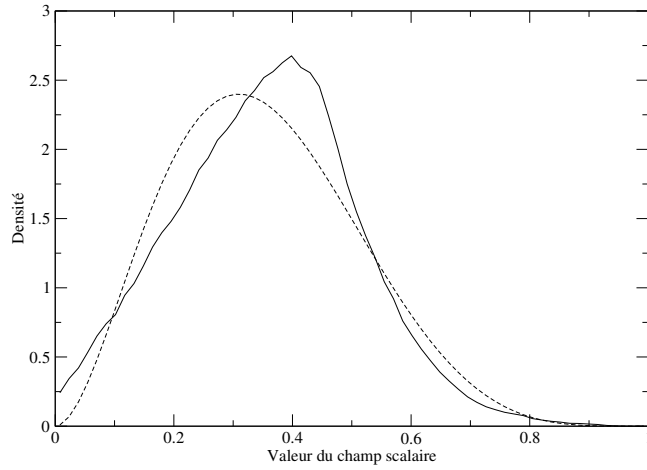


FIG. 3.8 – Comparaison entre $\beta(\Gamma; \bar{c}, \langle \sigma_s^2 | \bar{c}, \alpha \rangle)$ (trait pointillé) et $\langle d_s(\Gamma) | \bar{c}, \alpha \rangle$ (trait continu) pour $\bar{c} = 0,355 \pm 0,005$ et $\alpha = 0,01 \pm 0,001$.

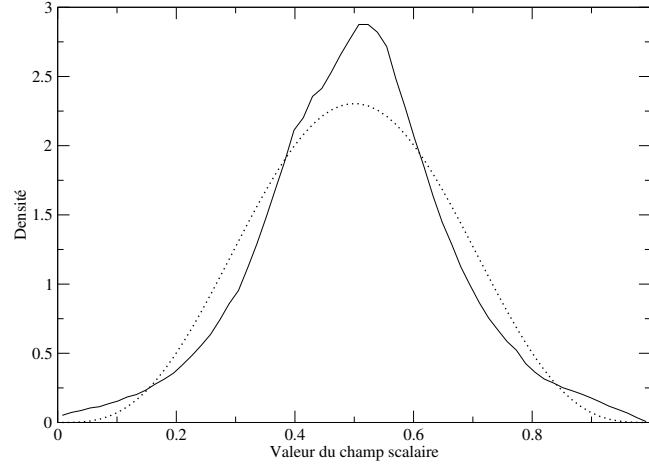


FIG. 3.9 – Comparaison entre $\beta(\Gamma; \bar{c}, \langle \sigma_s^2 | \bar{c}, \alpha \rangle)$ (trait pointillé) et $\langle d_s(\Gamma) | \bar{c}, \alpha \rangle$ (trait continu) pour $\bar{c} = 0,5 \pm 0.01$ et $\alpha = 0,0125 \pm 0.0005$.

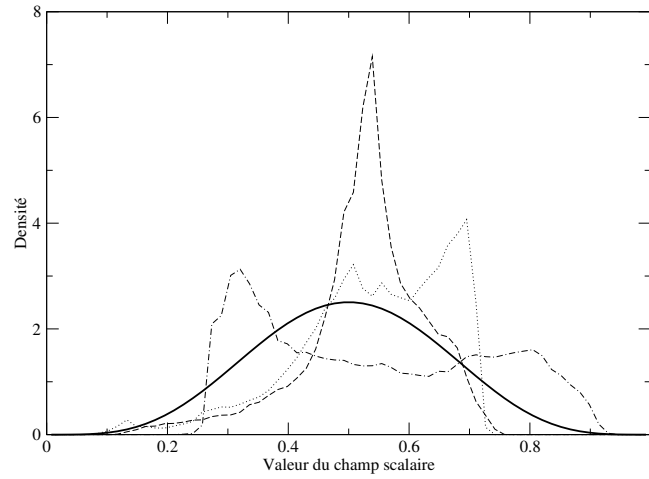


FIG. 3.10 – Comparaison entre $\beta(\Gamma; \bar{c}, \langle \sigma_s^2 | \bar{c}, \alpha \rangle)$ (ligne continue) et des densités de sous-maille pour des points présentant les mêmes valeurs de $\bar{c}$ et α . Les variances des densités sont donc différentes.

Estimateurs admettant une expression simplifiée

Les moments des fonctions β ont une expression analytique relativement simple. Cela signifie que pour $f(c) = c^n$, l'estimateur généré s'écrit (voir annexe C)

$$A_{c^n}(\pi) = \int \Gamma^n \frac{\Gamma^{a-1}(1-\Gamma)^{b-1}}{B(a,b)} d\Gamma \quad (3.9)$$

$$= \prod_{k=0}^{n-1} \frac{a+k}{a+b+k} \quad (3.10)$$

Les estimateurs sont linéaires vis à vis de la fonction estimée f , c'est à dire qu'on a

$$A_{\lambda_1 f_1 + \lambda_2 f_2} = \lambda_1 \cdot A_{f_1} + \lambda_2 \cdot A_{f_2}.$$

Avec ces deux propriétés, on comprend qu'il est possible d'obtenir une expression simple pour les polynômes de la forme $P = \sum_{j=0}^d \alpha_j \Gamma^j$. Cette expression s'écrit

$$A_P = \alpha_0 + \sum_{j=1}^d \alpha_j \prod_{k=0}^{j-1} \frac{a+k}{a+b+k}. \quad (3.11)$$

Il faut remarquer que cette expression se factorise très facilement. On aboutit ainsi à une forme factorisée de l'estimateur ressemblant à la factorisation de Horner du polynôme, qui devrait nettement accélérer le calcul numérique d'un tel estimateur :

$$A_P = \alpha_0 + \frac{a}{a+b} \left(\alpha_1 + \frac{a+1}{a+b+1} \left(\dots + \alpha_d \frac{a+d-1}{a+b+d-1} \right) \right) \quad (3.12)$$

Parmi les polynômes, il en est de particuliers de la forme $h_\kappa = (x(1-x))^\kappa$. Ces polynômes sont aussi des fonctions β non normalisées avec pour exposants $a = b = \kappa + 1$, et forment des approximations acceptables de taux de réaction [44]. L'expression des estimateurs générés pour ces fonctions est

$$\begin{aligned} A_{h_\kappa} &= \int \frac{\Gamma^{\kappa+a-1} (1-\Gamma)^{\kappa+b-1}}{B(a,b)} d\Gamma \\ &= \frac{B(a+\kappa, b+\kappa)}{B(a,b)} \\ &= \prod_{k=0}^{\kappa-1} \frac{a+k}{a+b+\kappa+k} \cdot \frac{B(a, b+\kappa)}{B(a,b)} \\ &= \prod_{k=0}^{\kappa-1} \frac{a+k}{a+b+\kappa+k} \cdot \prod_{k=0}^{\kappa-1} \frac{b+k}{a+b+k} \\ &= \left(\prod_{k=0}^{\kappa-1} (a+k)(b+k) \right) \left(\prod_{k=0}^{2\kappa-1} (a+b+k) \right)^{-1} \end{aligned}$$

C'est donc bien entendu parmi ces fonctions dont l'estimateur généré par la relation (3.8) connaît une expression relativement simple qu'ont été choisies les fonctions pour lesquelles les calculs d'erreur ont été menés.

$f(c) =$	$h_2 = (4c(1-c))^2$	$h_3 = (4c(1-c))^3$
$\langle \overline{f(c)} \overline{c}, \sigma_s^2 \rangle$	$4,969.10^{-3}$	$1,139.10^{-2}$
$A_f(\overline{c}, \sigma_s^2)$	$5,201.10^{-3}$	$1,312.10^{-2}$
$\langle \overline{f(c)} \overline{c}, \alpha \rangle$	$6,317.10^{-2}$	$7,874.10^{-2}$
$A_f(\overline{c}, \alpha)$	$6,358.10^{-2}$	$8,087.10^{-2}$
$\langle \overline{f(c)} \overline{c}, \epsilon_{\overline{c}} \rangle$	$5,140.10^{-2}$	$6,487.10^{-2}$
$A_f(\overline{c}, \epsilon_{\overline{c}})$	$5,170.10^{-2}$	$6,620.10^{-2}$

TAB. 3.1 – Erreurs normalisées commises par différents estimateurs de $\overline{f(c)}$ pour différentes fonctions f

$f(c) =$	$(4c(1-c))^2$	$(4c(1-c))^3$
$\langle \overline{f(c)} \overline{c}, \alpha \rangle$	$6,063.10^{-2}$	$7,416.10^{-2}$
$A_f(\overline{c}, \alpha)$	$6,341.10^{-2}$	$7,775.10^{-2}$
$\langle \overline{f(c)} \overline{c}, \epsilon_{\overline{c}} \rangle$	$5,028.10^{-2}$	$6,175.10^{-2}$
$A_f(\overline{c}, \epsilon_{\overline{c}})$	$5,248.10^{-2}$	$6,495.10^{-2}$

TAB. 3.2 – Erreurs normalisées commises par différents estimateurs de $\overline{f(c)}$ pour différentes fonctions f , avec un filtrage peigne de Dirac.

Résultats

Les résultats sont présentés table 3.1 en ce qui concernent les erreurs commises pour l'estimation de $\overline{f(c)}$ pour deux fonctions f différentes. Les deux fonctions choisies sont h_2 et h_3 , qui sont représentées figure 3.12. La fonction h_1 n'a pas été utilisée, car $\langle h_1 | \pi \rangle$ se calcule simplement à partir de $\overline{c}$ et de $\langle \sigma_s^2 | \pi \rangle$. Le tableau donne pour chaque fonction l'erreur fondamentale et l'erreur commise par un estimateur issu de l'hypothèse β , $A_f(\pi)$ - et ce, pour trois jeux de paramètres : $\{\overline{c}, \sigma_s^2\}$, $\{\overline{c}, \alpha\}$ et $\{\overline{c}, \epsilon_{\overline{c}}\}$. L'estimateur optimal est calculé grâce à une méthode décrite en détail dans l'annexe D, qui utilise des histogrammes.

Le tableau 3.2 présente le même type de résultats que précédemment, mais en prenant comme filtre fondamental un filtrage peigne de Dirac. Il est manifeste que les erreurs trouvées sont proches des valeurs exposées dans le tableau 3.1. Cela conduit à penser que l'hypothèse β reste excellente pour des filtrages différents du filtrage porte physique, et notamment un filtrage pour lequel la densité de sous-maille n'est plus une fonction continue mais une distribution. L'erreur supplémentaire commise par l'hypothèse β semble donc négligeable et peu dépendre du filtrage considéré.

3.2.3 Conclusions

Variance de sous-maille connue

Les comparaisons directes nous ont donc permis de comparer les estimateurs optimaux pour la densité de sous-maille à des fonctions bêta. Elles ont montré que les deux étaient extrêmement

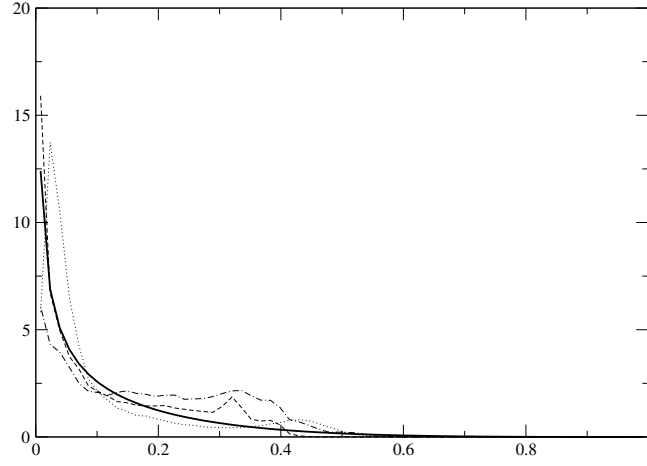


FIG. 3.11 – Comparaison entre $\beta(\Gamma; \bar{c}, \langle \sigma_s^2 | \bar{c}, \alpha \rangle)$ (ligne continue) et des densités de sous-maille pour des points présentant les mêmes valeurs de $\bar{c}$ et α . Les variances des densités sont donc différentes.

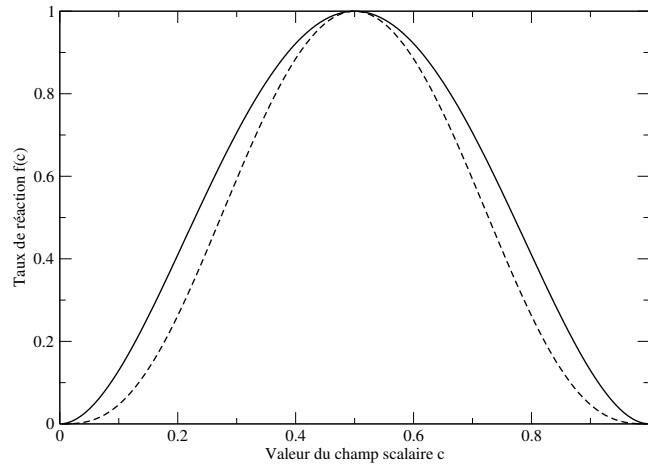


FIG. 3.12 – Fonctions h_2 (ligne continue) et h_3 (pointillés) utilisées comme fonctions test pour l'estimation de $\overline{f(c)}$

proches. Le calcul des erreurs commises par les différents estimateurs a permis de savoir précisément l'erreur supplémentaire due à l'hypothèse bêta. Les résultats obtenus laissent supposer qu'elle est parfaitement négligeable. Ces résultats laissent ainsi à penser que les fonctions bêta sont un excellent choix pour reproduire la forme des estimateurs optimaux.

Des résultats présentés en annexe C laissent supposer que le choix de distributions β est pertinent utilisé dans le même cadre (étude des estimateurs optimaux de la densité de sous-maille) mais pour des champs différents comme c^2 ou $\sqrt{c}$. Peut-être est-ce là le signe d'une propriété mathématique des distributions β , qui ferait que la moyenne de distributions bornées de même moyenne et de même variance tendrait vers une distribution β . Disons que si les estimateurs optimaux ne sont pas exactement des fonctions β , cette distinction est au delà de la précision dont nous disposons [58].

Variance de sous-maille estimée

On constate tout d'abord que l'erreur des estimateurs optimaux associés à des jeux de paramètres exclusivement issus des grandes échelles (et dont la variance de sous-maille ne fait donc pas partie) est plus importante d'un ordre de grandeur. Cela signifie que les paramètres grandes échelles considérés contraignent statistiquement beaucoup moins la situation locale et sont donc moins pertinents. A cette erreur irréductible s'ajoute une erreur due à l'utilisation d'une distribution β , qui reste cependant modeste.

On constate en effet que les distributions β reproduisent moins bien la forme de l'estimateur optimal de la densité de sous-maille que lorsque la variance de sous-maille est connue. On sait que l'estimateur optimal est la moyenne des densités de sous-maille pour des événements associés à des valeurs identiques des paramètres considérés. Par rapport au cas précédent, l'estimateur optimal est donc la moyenne de densités ayant toutes la même moyenne $\bar{c}$, mais des variances σ_s^2 qui sont cette fois différentes. Il se produit ainsi une dispersion statistique en variance des densités, inhérente au jeu de paramètre considéré.

Plus un jeu de paramètres permet d'estimer pertinemment σ_s^2 , plus la dispersion en variance des densités est réduite. Plus on réduit cette dispersion, plus les estimateurs optimaux prennent une forme proche de celle qu'ils ont lorsque la variance de sous-maille est connue, forme qui est alors bien reproduite par les distributions β .

En conclusion, plus un jeu de paramètres permet d'estimer finement σ_s^2 en la contraignant statistiquement, meilleure se révèle l'hypothèse β . L'amélioration du modèle ne passe donc pas tant par un abandon de l'hypothèse β (et la recherche de distributions plus adaptées) que par une meilleure estimation de la variance de sous-maille.

Au vu de ces résultats, on pourrait avoir envie de résoudre numériquement l'équation d'évolution de la variance de sous-maille, qui s'écrit

$$\partial_t \sigma_s^2 + \partial_i (\overline{v_i c^2}) - 2 \bar{c} \partial_i (\overline{v_i c}) = D \Delta \sigma_s^2 - 2D \overline{(\partial_i c)^2} + 2D \partial_i \bar{c}^2.$$

Résoudre cette équation donnerait une approximation directe de la variance de sous-maille. Dans le cas où l'équation d'évolution du scalaire ferait intervenir un terme $f(c)$, il faudrait ajouter des termes à l'équation de σ_s^2 :

$$2c \overline{f(c)} - 2\bar{c} \overline{f(c)}.$$

Ceux-ci peuvent se fermer immédiatement et de façon quasiment optimale à l'aide d'une hypothèse β . Les autres termes non-fermés peuvent être estimés en faisant intervenir une diffusivité

turbulente, mis à part $\overline{\epsilon_c}$ qui demande une modélisation particulière. Ce problème a été abordé par Girimaji *et al.* [59]. Le fait que l'hypothèse β soit si efficace incite donc à résoudre l'équation d'évolution de la densité de sous-maille, plutôt qu'à l'estimer directement.

3.3 Estimation de la variance de sous-maille

Dans la section précédente, l'étude a été effectuée en supposant qu'on connaissait l'estimateur optimal de la variance de sous-maille, quantité centrale du modèle. Dans la pratique, il n'en est rien et des sous-modèles doivent être utilisés, qui sont parfois simplistes. Nous allons voir comment la démarche que nous avons retenue permet d'étudier la validité de ces modèles et de proposer d'éventuelles améliorations.

3.3.1 Sous-modèle de Cook et Riley

Le premier sous-modèle proposé l'a été par Cook et Riley [19] dans leur article initial. En invoquant un argument d'auto-similarité, ils posent une simple relation de proportionnalité entre σ_s^2 et α

$$\sigma_s^2 \simeq C_\alpha \alpha. \quad (3.13)$$

Le filtre test qui intervient dans la définition de $\alpha = \widehat{c}^2 - \widehat{c}^2$ est un filtre peigne de Dirac effectué sur une grille correspondant à une grille de simulation des grandes échelles fictive, de pas 8 fois celui de la grille réelle.

Cook et Riley proposent plusieurs façons de déterminer la constante C_α dans les conditions de la simulation des grandes échelles. Dans notre cas, nous allons prendre la constante qui minimise l'erreur quadratique commise par l'estimateur (3.13) à l'instar de Moin [42]. Cette erreur s'écrit

$$\begin{aligned} E_v &= \langle (\sigma_s^2 - C_\alpha \alpha)^2 \rangle \\ &= \langle \sigma_s^4 \rangle + C_\alpha^2 \langle \alpha^2 \rangle - 2C_\alpha \langle \sigma_s^2 \cdot \alpha \rangle. \end{aligned}$$

On cherche le minimum de cette erreur, fonction de C_α , en considérant sa dérivée

$$\frac{\partial E_v}{\partial C_\alpha} = 2 C_\alpha \langle \alpha^2 \rangle - 2 \langle \sigma_s^2 \cdot \alpha \rangle$$

qui s'annule pour

$$C_\alpha = \frac{\langle \sigma_s^2 \cdot \alpha \rangle}{\langle \alpha^2 \rangle}. \quad (3.14)$$

C'est pour cette valeur de la constante que l'erreur est minimale. Ce type de raisonnement est très courant [47], et il se généralise immédiatement à d'autres quantités que σ_s^2 et α - ce qui sera couramment utilisé dans la suite de ce travail.

Comme l'a déjà fait Jimenez [43], on peut se demander dans quelle mesure l'hypothèse d'une simple relation de proportionnalité est justifiée. Pour tenter de répondre à cette question, on peut considérer l'estimateur optimal de σ_s^2 pour α , c'est à dire $\langle \sigma_s^2 | \alpha \rangle$. Cette fonction de α est présentée figure 3.13. On en conclut que *si on se limite à l'utilisation du paramètre α* , la relation de proportionnalité entre σ_s^2 et α est un très bon choix.

3.3.2 Recherche d'une meilleure approximation analytique

On a vu dans les paragraphes précédents que le modèle de Cook et Riley avec sous-modèle revenait à utiliser un jeu de *deux* paramètres constitué par $\bar{\tau}$ et α . Dans ce cadre, l'objectif d'un estimateur analytique est de s'approcher le plus possible de l'estimateur optimal $\langle \sigma_s^2 | \bar{\tau}, \alpha \rangle$. Pour cela, on peut faire intervenir le paramètre $\bar{\tau}$ dans l'estimation de σ_s^2 (alors qu'on a vu ci-dessus qu'il était ignoré dans l'estimation de σ_s^2 de Cook et Riley). Il faut alors rechercher un estimateur qui ne soit plus une simple relation de proportionnalité entre σ_s^2 et α . Comme la relation de proportionnalité est cependant loin d'être infondée, on peut prendre pour estimateur

$$\sigma_s^2 \simeq h(\bar{\tau}) \alpha, \quad (3.15)$$

en adoptant une démarche souvent qualifiée de "dynamique". Cela revient à dire qu'à $\bar{\tau}$ fixé, la relation de proportionnalité reste valable. L'erreur commise par cet estimateur s'écrit alors

$$\begin{aligned} E_h &= \langle (\sigma_s^2 - h(\bar{\tau}) \alpha)^2 \rangle \\ &= \int (\sigma_s^2 - \bar{\tau} \alpha)^2 p(\sigma_s^2, \bar{\tau}, \alpha) d\bar{\tau} d\sigma_s^2 d\alpha \\ &= \int p(\bar{\tau}) d\bar{\tau} \int (\sigma_s^2 - \bar{\tau} \alpha)^2 p(\sigma_s^2, \alpha | \bar{\tau}) d\sigma_s^2 d\alpha \\ &= \int p(\bar{\tau}) d\bar{\tau} \{ \langle \sigma_s^4 | \bar{\tau} \rangle + h(\bar{\tau})^2 \langle \alpha^2 | \bar{\tau} \rangle - 2 h(\bar{\tau}) \langle \alpha \sigma_s^2 | \bar{\tau} \rangle \} \end{aligned}$$

Cela signifie qu'il existe une unique fonction de $\bar{\tau}$ qui minimise l'erreur : en effet, comme $p(\bar{\tau}) \geq 0$, pour que l'erreur soit minimale, il faut que le terme $\langle \sigma_s^4 | \bar{\tau} \rangle + h(\bar{\tau})^2 \langle \alpha^2 | \bar{\tau} \rangle - 2 h(\bar{\tau}) \langle \alpha \sigma_s^2 | \bar{\tau} \rangle$ le soit pour toute valeur de $\bar{\tau}$. Quand $\bar{\tau}$ est fixée, $h(\bar{\tau})$ peut être considérée comme une constante, et le raisonnement précédent peut être répété. En définitive, pour minimiser l'erreur commise par l'estimateur, il faut prendre la fonction

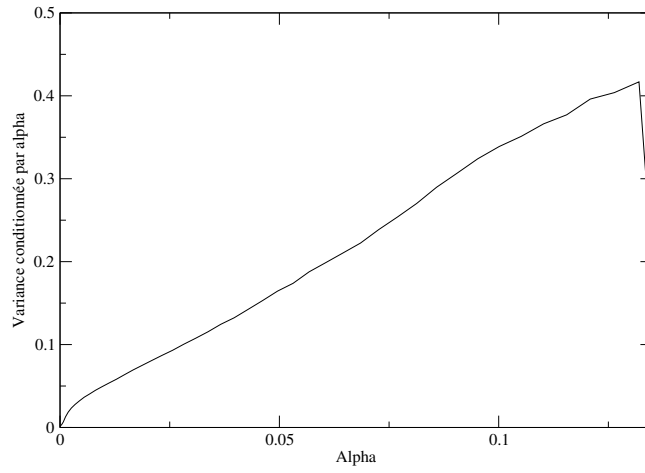
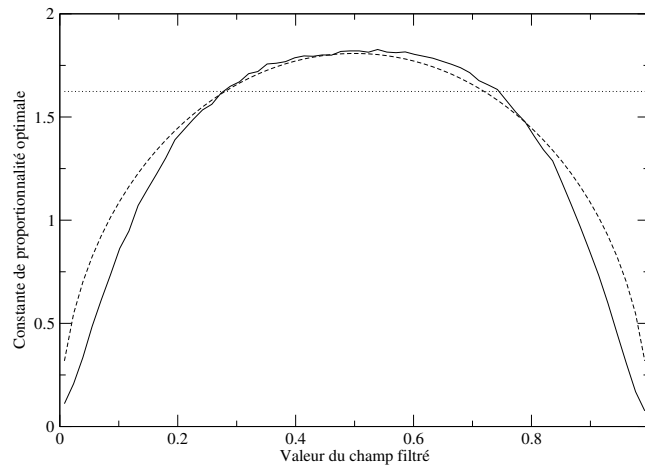
$$h(\bar{\tau}) = \frac{\langle \sigma_s^2 \cdot \alpha | \bar{\tau} \rangle}{\langle \alpha^2 | \bar{\tau} \rangle}. \quad (3.16)$$

Cette fonction est présentée figure 3.14, et on constate que sa dépendance en $\bar{\tau}$ est très loin d'être négligeable. Supposer la relation (3.13) revient à prendre $h(\bar{\tau}) = C_\alpha$, et si cela peut se justifier loin des bords où $h(\bar{\tau})$ semble connaître un plateau, on comprend bien combien cette hypothèse est erronée dès lors que $\bar{\tau}$ s'approche des bornes. L'allure de la fonction $h(\bar{\tau})$ ne semble pas dépendre de la taille de filtrage utilisée (voir figure 3.15). L'importance de la proximité des bornes dans les valeurs prises par $h(\bar{\tau})$ s'explique par le fait que près des bornes la variance de sous-maille est limitée en excursion. C'est ce que reflète la propriété 1.5, qui indique que

$$\sigma_s^2 \leq \bar{\tau}(1 - \bar{\tau}). \quad (3.17)$$

Si on veut prendre en compte ces effets de bord, on peut essayer de reproduire empiriquement le comportement de h en écrivant que $h(\bar{\tau}) \sim \sqrt{\bar{\tau}(1 - \bar{\tau})}$. On constate en effet que cette forme est plus efficace pour approcher h que $\bar{\tau}(1 - \bar{\tau})$. Cela revient à proposer comme nouvel estimateur analytique de la variance de sous-maille

$$\sigma_s^2 \simeq C_0 \sqrt{\bar{\tau}(1 - \bar{\tau})} \alpha \quad (3.18)$$

FIG. 3.13 – $\langle \sigma_s^2 | \alpha \rangle$ FIG. 3.14 – Fonction $h(\bar{c})$ optimale (trait continu); $h(\bar{c}) = C_\alpha$ (trait pointillé); $h(\bar{c}) = C_0 \sqrt{\bar{c}(1 - \bar{c})}$ (trait discontinu)

en choisissant C_0 de façon à minimiser l'erreur commise par l'estimateur, soit

$$C_0 = \frac{\langle \sigma_s^2 \sqrt{\bar{c}(1-\bar{c})} \alpha \rangle}{\langle \bar{c}(1-\bar{c}) \alpha^2 \rangle}. \quad (3.19)$$

Pour savoir s'il est toujours fondé de supposer que la quantité α intervienne comme un simple facteur dans l'expression de ce nouvel estimateur il faut considérer l'estimateur optimal

$$\left\langle \frac{\sigma_s^2}{\sqrt{\bar{c}(1-\bar{c})}} \middle| \alpha \right\rangle.$$

Cette fonction est représentée figure 3.16. Elle laisse penser qu'il est tout à fait fondé de prendre pour estimateur (3.18).

L'estimateur construit ci-dessus, s'il s'avère être bien meilleur que l'hypothèse d'une simple proportionnalité à α n'est pas pour autant parfait. Cette étude a cependant le mérite de souligner que l'hypothèse d'un comportement auto-similaire de la variance de sous-maille est sans doute fautive dans le cas du scalaire borné. Ceci ne dépend pas de la taille du filtrage utilisé, comme le montre la figure 3.15 qui présente les fonctions $h(\bar{c})$ optimales pour différentes tailles de filtrages. On peut penser que même dans la zone inertielle la mieux développée, le comportement de $h(\bar{c})$ perdurera et que l'hypothèse auto-similaire n'est ainsi pas tenable. Ce sont sans doute les bornes du scalaire qui jouent ici un rôle crucial. Il faut noter d'ailleurs que le fait que le scalaire soit borné ne peut pas se déduire du spectre.

3.3.3 La dissipation du champ filtré comme paramètre du modèle

Il existe un autre sous-modèle pour σ_s^2 , proposé par Moin *et al.*[42]. Celui-ci fait intervenir non plus le paramètre α mais la dissipation des grandes échelles du scalaire,

$$\epsilon_{\bar{c}} = 2 D (\partial_i \bar{c})^2. \quad (3.20)$$

De la même manière que pour le sous-modèle précédent, on suppose une simple relation de proportionnalité entre la variance de sous-maille et cette quantité

$$\sigma_s^2 \simeq \lambda \epsilon_{\bar{c}}, \quad (3.21)$$

avec

$$\lambda = \frac{\langle \sigma_s^2 \epsilon_{\bar{c}} \rangle}{\langle \epsilon_{\bar{c}}^2 \rangle}. \quad (3.22)$$

La moyenne de la variance de sous-maille conditionnée par la dissipation du champ filtré, représentée figure 3.17, semble indiquer que la relation de proportionnalité est un bon choix.

On peut mener une étude comparable à celle qui a été menée pour le paramètre α . En effet, ce changement de sous-modèle correspond du point de vue des estimateurs optimaux à un changement des paramètres fondamentaux, qui sont maintenant $\bar{c}$ et $\epsilon_{\bar{c}}$. On peut donc rechercher un nouvel estimateur de la forme

$$\sigma_s^2 \simeq k(\bar{c}) \epsilon_{\bar{c}}.$$

Et en prenant

$$k(\bar{c}) = \frac{\langle \sigma_s^2 \epsilon_{\bar{c}} | \bar{c} \rangle}{\langle \epsilon_{\bar{c}}^2 | \bar{c} \rangle},$$

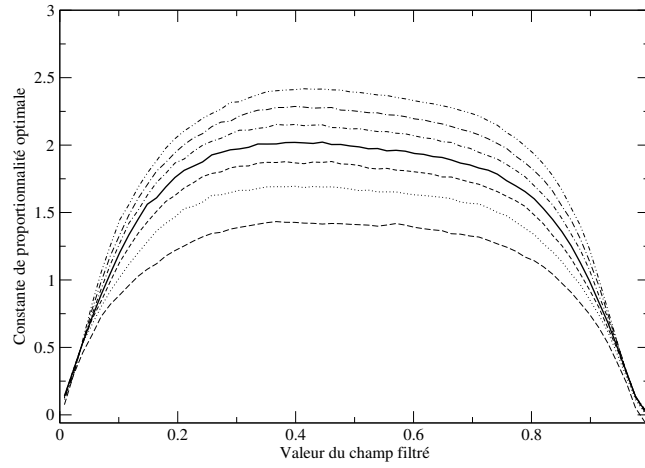


FIG. 3.15 – Fonction $h(\bar{c})$ optimale pour différentes tailles de filtrage (de la plus petite en dessous à la plus grande au dessus). La courbe en gras correspond au filtrage utilisé.

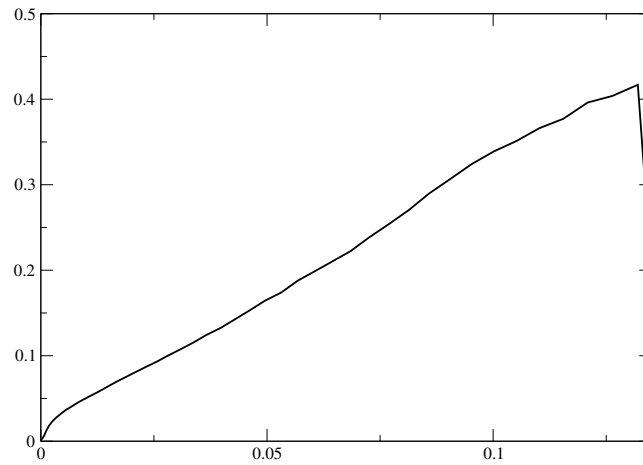


FIG. 3.16 – $\left\langle \frac{\sigma_s^2}{\sqrt{\bar{c}(1-\bar{c})}} \middle| \alpha \right\rangle$

on sait qu'on va minimiser l'erreur quadratique commise par l'estimateur. Cette fonction est représentée figure 3.18. Au vu du résultat, on peut prendre pour nouvel estimateur

$$\sigma_s^2 \simeq \Lambda \sqrt{\bar{c}(1-\bar{c})} \epsilon_{\bar{c}} \quad (3.23)$$

avec

$$\Lambda = \frac{\left\langle \sigma_s^2 \sqrt{\bar{c}(1-\bar{c})} \epsilon_{\bar{c}} \right\rangle}{\langle \bar{c}(1-\bar{c}) \epsilon_{\bar{c}}^2 \rangle}.$$

Pour finir, comme pour le paramètre α , la quantité $\left\langle \frac{\sigma_s^2}{\sqrt{\bar{c}(1-\bar{c})}} \middle| \epsilon_{\bar{c}} \right\rangle$ est calculée et présentée figure 3.19. La forme de l'estimateur analytique semble d'après ce graphique un peu moins optimale.

3.3.4 Comparaison des différents estimateurs

Pour comparer les nouveaux estimateurs qui viennent d'être introduits, il convient de considérer l'erreur qu'ils commettent sur l'estimation de la variance de sous-maille. Ces résultats sont présentés dans le tableau 3.3 par ordre d'erreur croissante. Ils incluent les estimateurs optimaux pour pouvoir juger de la proximité des estimateurs analytiques et des estimateurs optimaux.

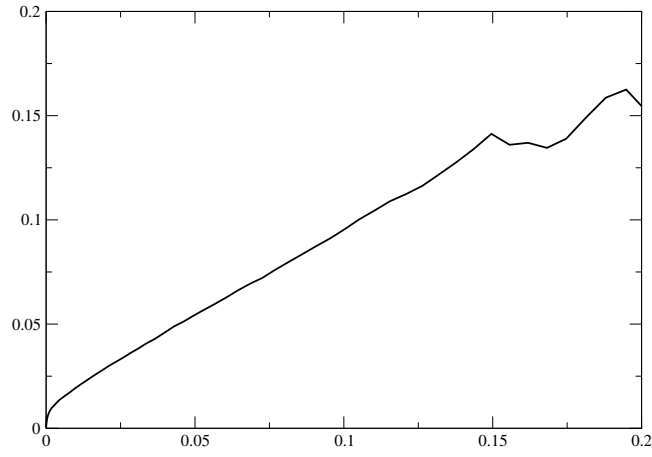
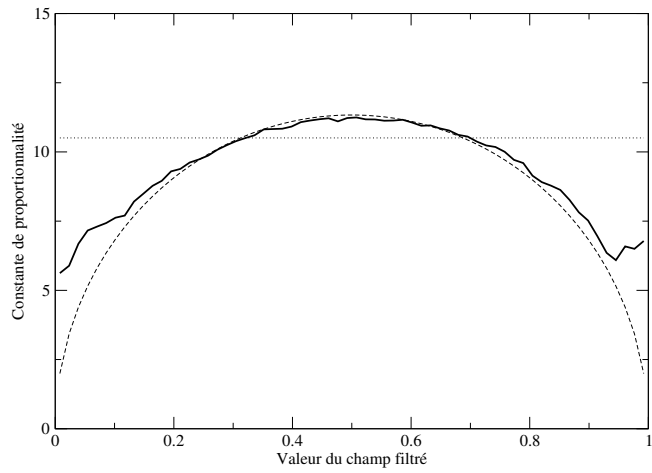
Il apparaît immédiatement que les estimateurs analytiques sont encore assez loin des estimateurs optimaux. Dans le cas où $\bar{c}$ et α sont les paramètres fondamentaux, le nouvel estimateur³ représente une diminution importante de l'erreur par rapport à $C_\alpha \alpha$, et il parvient ainsi à rivaliser avec les estimateurs utilisant $\bar{c}$ et $\epsilon_{\bar{c}}$. Les erreurs commises par les estimateurs utilisant $\epsilon_{\bar{c}}$ et le nouvel estimateur en α sont en effet tout à fait comparables. Des calculs d'erreur partiels qui ne sont pas présentés ici laissent à penser qu'une injection plus faible, parce qu'elle diminue la fréquence des cas où $\bar{c}$ est proche des bornes, pourrait désavantager le nouvel estimateur en α qui commettrait alors une erreur plus grande que les estimateurs en $\epsilon_{\bar{c}}$. Cette erreur serait dans ce cas proche de celle commise par la proportionnalité à α . C'est dire s'il est difficile de conclure qu'il existe un estimateur analytique vraiment meilleur que les autres. Il convient de remarquer également que l'introduction d'un nouvel estimateur utilisant $\bar{c}$ et $\epsilon_{\bar{c}}$ n'apporte pas d'amélioration très significative, même si elle est avérée.

Dans le tableau 3.4 sont présentées les erreurs commises en utilisant les différents estimateurs de la variance de sous-maille et l'hypothèse β pour estimer $\overline{f(c)}$. Une chose est réellement frappante : quelle que soit la fonction f considérée, meilleure est l'approximation de σ_s^2 par un estimateur, meilleure est l'approximation de $\overline{f(c)}$ générée par cet estimateur et une hypothèse β . Et ce même lorsque la différence entre deux erreurs pour l'estimation de σ_s^2 est très faible (cas des nouveaux estimateurs). On peut donc envisager que tout progrès dans l'estimation de σ_s^2 se traduise de façon générale par une amélioration de l'estimation de $\overline{f(c)}$.

3.4 Réseaux de neurones

Les réseaux de neurones sont un vaste sujet en informatique (on consultera [60] pour un ouvrage de référence pas trop théorique sur le sujet). Ils sont encore relativement peu utilisés en turbulence, sauf peut-être pour des usages assez spécifiques comme le contrôle [61] de certaines

³ $\sqrt{\bar{c}(1-\bar{c})} \alpha$

FIG. 3.17 – $\langle \sigma_s^2 | \epsilon_{\bar{c}} \rangle$ FIG. 3.18 – Fonction $k(\bar{c})$ optimale (trait continu) ; $k(\bar{c}) = \lambda$ (trait pointillé) ; $h(\bar{c}) = \Lambda \sqrt{\bar{c}(1 - \bar{c})}$ (trait discontinu)

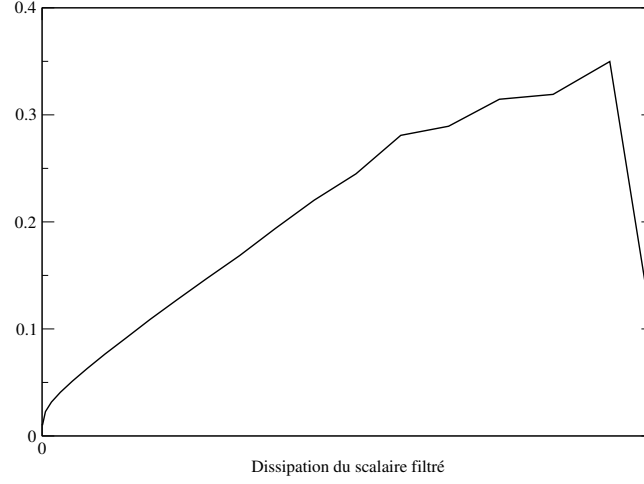


FIG. 3.19 – $\left\langle \frac{\sigma_s^2}{\sqrt{\bar{c}(1-\bar{c})}} \middle| \epsilon_{\bar{c}} \right\rangle$

<i>Estimateur</i>	<i>Erreur normalisée</i>
$\langle \sigma_s^2 \bar{c}, \epsilon_{\bar{c}} \rangle$	0,194
$\langle \sigma_s^2 \bar{c}, \alpha \rangle$	0,240
$C_0 \sqrt{\bar{c}(1-\bar{c})} \alpha$	0,299
$\Lambda \sqrt{\bar{c}(1-\bar{c})} \epsilon_{\bar{c}}$	0,304
$\lambda \epsilon_{\bar{c}}$	0,316
$C_\alpha \alpha$	0,363

TAB. 3.3 – Erreurs commises par différents estimateurs de f , classées par erreur décroissante

caractéristiques d'écoulements turbulents, ou la reconnaissance de schémas d'évolution de tourbillons [62]. L'utilisation qui est faite ici des réseaux de neurones est beaucoup plus simple que dans les deux articles cités ci-dessus, mais elle est par contre très différente dans l'esprit.

Si l'utilisation des réseaux de neurones est apparue ici comme naturelle, c'est parce que la méthode décrite en annexe D à base d'histogramme devient trop coûteuse dès lors que l'on veut l'utiliser avec plus de 3 paramètres. Dans le cadre de ce travail et pour l'utilisation qui va en être faite, les réseaux de neurones ne sont qu'une méthode efficace pour approcher une moyenne conditionnelle. Étant donné cependant l'importance que peuvent avoir les moyennes conditionnelles en modélisation, les réseaux de neurones peuvent être un outil très puissant. Nous envisagerons également la possibilité d'utiliser les réseaux "tels quels" dans des simulations, pour estimer les termes non fermés des équations d'évolution.

3.4.1 Présentation d'un réseau de neurones

Le réseau de neurones décrit ici est d'un type très commun et relativement ancien : le perceptron multi-couche. Pour l'usage qui en est fait ici, il n'est pas nécessaire d'utiliser des types de réseaux très compliqués comme ceux utilisés par Lee[61].

Un réseau de neurones n'est rien d'autre qu'une fonction d'un vecteur $\vec{x}$ à valeurs dans un espace vectoriel et qui s'écrit (avec la convention de sommation sur les indices répétés)

$$F_i(\vec{x}) = A_{ij} \tanh(B_{jk} x_k + b_j) + a_i. \quad (3.24)$$

Un réseau de neurones donné est donc déterminé par les matrices A_{ij} et B_{jk} et les vecteurs a_i et b_j , qu'on désigne usuellement comme étant les **poids** associés au réseau. La fonction décrite ci-dessus peut prendre des allures complètement différentes selon les poids. On peut ainsi s'en servir pour approcher à peu près n'importe quelle autre fonction, et c'est cette souplesse qui fait la puissance et l'intérêt des réseaux de neurones.

Ce type de fonction a historiquement été construit par analogie avec un réseau de neurones réels. Dans notre cas, un neurone artificiel est une fonction à valeurs réelles d'un vecteur x_i qui s'écrit

$$f(\vec{x}) = \tanh(w_i x_i)$$

et qui se représente usuellement comme sur la figure 3.20.

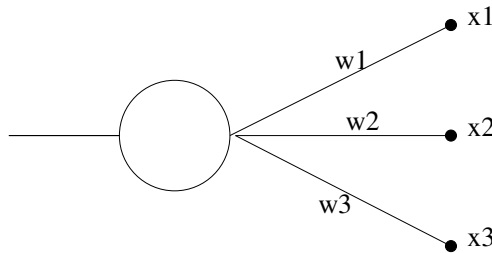


FIG. 3.20 – Représentation d'un neurone artificiel

Un tel neurone accorde un poids w_i à chacune de ses entrées et utilise une fonction sigmoïde (ici une tangente hyperbolique) pour obtenir une sortie. La fonction $\vec{F}$ peut être représentée par un réseau de ces neurones artificiels. La figure 3.21 montre ce que serait ce réseau si l'entrée $\vec{x}$

était de dimension deux, la sortie de dimension un, et si la couche intermédiaire (appelée souvent couche cachée) comptait trois neurones.

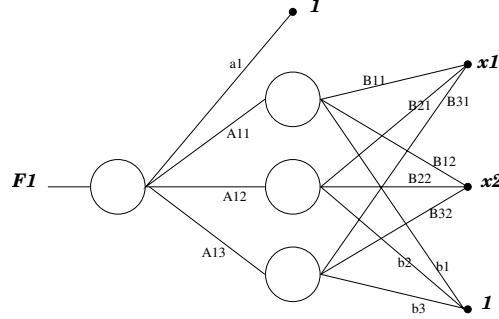


FIG. 3.21 – Schéma d'un réseau de neurones

Les neurones de sortie sont ici particuliers en ce qu'ils n'utilisent pas de sigmoïdes, parce qu'elles n'ont pour eux pas d'utilité particulière - au contraire, elles limiteraient les valeurs de sortie. On remarquera la présence d'entrées prenant la valeur constante 1. Ce sont elles qui permettent de faire intervenir sur le schéma les vecteurs $\vec{a}$ et $\vec{b}$ introduits dans la fonction (3.24). Il faut noter pour finir que le nombre de neurones dans la couche cachée est choisi arbitrairement et qu'il n'est pas lié à la dimension de l'entrée ou de la sortie. Dans les essais qui ont été menés plus loin, le nombre de neurones de la couche intermédiaire varie de 30 à 50.

Il est possible d'imaginer des réseaux plus complexes notamment en augmentant le nombre de couches (ce qui est relativement courant pour les problèmes compliqués), ou en abandonnant la structure en couche pour une structure plus désordonnée (ce qui ne se fait pas en général, car on ne dispose pas d'arguments solides en faveur d'une structure plutôt que d'une autre).

3.4.2 Erreur et apprentissage

Supposons que nous disposions de données constituées de N couples $\{\vec{y}, \vec{x}\}$. Il s'agit de trouver la meilleure façon d'approcher $\vec{y}$ par une fonction de la forme de (3.24) de $\vec{x}$, c'est à dire de trouver les poids les plus adaptés. C'est cette phase qui est appelée **apprentissage**. Pour mesurer l'acuité d'un réseau de neurones, on calcule l'erreur quadratique qu'il commet en estimant $\vec{y}$, et qui s'écrit

$$E = \frac{1}{N} \sum_{\{\vec{y}, \vec{x}\}} \sum_i (y_i - F_i(\vec{x}))^2.$$

L'apprentissage consiste en pratique en la recherche d'un minimum de cette fonction. En minimisant l'erreur, on s'approche ainsi le plus possible de l'estimateur optimal $\langle \vec{y} | \vec{x} \rangle$.

L'erreur est fonction des différents poids (matrices A_{ij} et B_{jk} , vecteurs a_i et b_j), ce qui fait qu'elle détermine une hypersurface dans un espace où les coordonnées sont les poids et l'erreur. Trouver le minimum absolu de cette surface en termes d'erreur est une tâche très difficile en général, car c'est un problème NP-complet. De nombreux algorithmes existent qui peuvent être utilisés pour le résoudre (par exemple des algorithmes génétiques [63]). Dans le cas qui nous intéresse, des solutions plus élaborées n'apportent pas d'avantage tangible par rapport à une descente de gradient. C'est d'autant plus vrai que l'erreur est une fonction infiniment dérivable des différents

poids, et que ses dérivées partielles s'expriment relativement simplement :

$$\begin{aligned}
\frac{\partial E}{\partial A_{ij}} &= \frac{2}{N} \sum_{\{\vec{y}, \vec{x}\}} (F_i(\vec{x}) - y_i) \cdot \tanh \left(\sum_k B_{jk} x_k \right) \\
\frac{\partial E}{\partial B_{jk}} &= \frac{2}{N} \sum_{\{\vec{y}, \vec{x}\}} \sum_i (F_i(\vec{x}) - y_i) \cdot A_{ij} x_k \left(1 - \tanh^2 \left(\sum_k B_{jk} x_k \right) \right) \\
\frac{\partial E}{\partial a_i} &= \frac{2}{N} \sum_{\{\vec{y}, \vec{x}\}} (F_i(\vec{x}) - y_i) \\
\frac{\partial E}{\partial b_j} &= \frac{2}{N} \sum_{\{\vec{y}, \vec{x}\}} \sum_i (F_i(\vec{x}) - y_i) \cdot A_{ij} \left(1 - \tanh^2 \left(\sum_k B_{jk} x_k \right) \right)
\end{aligned}$$

Disposer d'une expression explicite du gradient de cette surface permet de mettre en oeuvre des méthodes du type "descente de gradient". Il s'agit de choisir un point de départ sur la surface, et de suivre la direction indiquée par le gradient (la plus grande pente). Une telle méthode garantit que l'on puisse trouver un minimum local relativement rapidement. Dans le cas précis des réseaux de neurones, on prend les couples $\{\vec{y}, \vec{x}\}$ un à un, et on modifie la position du point sur la surface en fonction de la contribution de chaque couple au gradient. Il s'agit d'une descente de gradient stochastique.

L'erreur commise par les réseaux de neurones présentée ci-après est une erreur en généralisation. Par exemple, sur un champ de scalaire en 128^3 , un huitième des points seulement seront utilisés pour l'apprentissage. Par contre, l'erreur calculée l'est en utilisant tous les points disponibles. Comme dans le cas des histogrammes, et pour les mêmes raisons, il s'agit d'une erreur en généralisation. Les paramètres du réseau de neurones (qui ne sont pas tous détaillés ici, notamment ceux qui contrôlent la descente de gradient) sont ajustés pour minimiser cette erreur.

Les réseaux de neurones constituent une méthode très intéressante comparée à la méthode par histogramme. Ils n'ont besoin que de relativement peu de statistiques pour produire un estimateur fiable. Ceci parce qu'ils supposent que la fonction à approcher est continue. L'influence d'un couple $\{\vec{x}, \vec{y}\}$ donné ne se limite ainsi pas à une case de discrétisation qui lui correspondrait, mais il participe à l'élaboration de la fonction un peu moins "localement". De plus, l'information nécessaire pour approcher l'estimateur optimal est réduite : elle se limite aux poids du réseau de neurones.

3.4.3 Quelques résultats

Comme l'ont montré les paragraphes précédents, l'hypothèse β induit une connexion entre l'estimateur de la variance de sous-maille et les estimateurs de $\overline{f(c)}$ quelle que soit f . Meilleur est l'estimateur de σ_s^2 , meilleurs sont ceux générés par l'hypothèse β . De plus, l'erreur commise par les estimateurs optimaux avec deux paramètres sur l'estimation de σ_s^2 est encore assez importante. Tout ceci laisse à penser qu'il est possible d'améliorer la qualité de l'estimation de la variance en utilisant des estimateurs optimaux de jeux de paramètres plus pertinents - et que cette amélioration se traduira par une amélioration de l'estimation de $f(c)$.

Pour un nombre de paramètres réduit, il a été vérifié que les histogrammes produisaient une erreur comparable à celle commise par les réseaux de neurones. En pratique, cette erreur est

identique à la cinquième décimale près, ce qui conduit à conclure que les histogrammes sont efficaces pour l'estimation de la variance de sous-maille.

Dans le tableau 3.4.3 sont présentés quelques uns des résultats obtenus avec les réseaux de neurones pour plusieurs jeux de paramètres. L'apprentissage se fait sur un huitième d'une image de 128^3 points, et le calcul de l'erreur sur la totalité de ce champ. Le code⁴ utilisé tourne sous Octave, un équivalent de Matlab. Si les erreurs calculées ne sont pas exactement comparables à celles présentées pour les histogrammes, c'est parce qu'Octave est limité quant à la taille des données qu'il peut manipuler.

Les deux premières lignes permettent de confirmer que $\epsilon_{\bar{\tau}}$ est bien un meilleur paramètre que α . La troisième permet d'affirmer que ces deux paramètres sont complémentaires : l'erreur commise avec α et $\epsilon_{\bar{\tau}}$ est plus petite que celle commise avec $\epsilon_{\bar{\tau}}$ seule.

La quatrième ligne vise à étudier l'influence de paramètres concernant la vitesse. En pratique, le gain en erreur semble négligeable (bien que non nul) si on compare à la première ligne. Il ne semble pas pertinent d'essayer d'incorporer des informations sur le champ de vitesse dans le jeu de paramètres.

La cinquième ligne du tableau utilise comme paramètres ceux décrits au paragraphe 1.4.3 pour le modèle dynamique de Moin - qui vise à estimer la densité de sous-maille. Elle donne donc l'erreur la plus petite qu'on puisse commettre en utilisant ces paramètres. Si le modèle pose des problèmes de singularité, du moins peut-on constater ici que les paramètres utilisés sont pertinents. Le gain en erreur est en effet conséquent. On peut remarquer cependant que la quantité $\bar{\tau}$ n'intervient pas dans ce modèle, alors qu'on a vu précédemment qu'elle avait une grande importance dans l'estimation analytique de σ_s^2 . En ne faisant que rajouter $\bar{\tau}$ aux paramètres précédents, on s'aperçoit que l'erreur commise est de nouveau très nettement réduite, comme le montre la huitième ligne. L'avant dernière ligne et l'anté-pénultième ligne du tableau montrent que les deux paramètres $\epsilon_{\bar{\tau}}$ et $\hat{\epsilon}_{\bar{\tau}}$ introduits par Moin conduisent à des réductions de l'erreur sensiblement identiques. Ceci dit, il est plus efficace d'utiliser $\bar{\tau}$ et un de ces deux paramètres que ces deux paramètres sans $\bar{\tau}$.

De façon générale, il apparaît qu'en ajoutant des paramètres qui caractérisent l'allure locale du champ scalaire, on obtient une diminution de l'erreur commise par les estimateurs optimaux.

3.4.4 Vers une simulation idéale des grandes échelles

Ce paragraphe repose pour partie sur les travaux d'Adrian[64] puis de Langford et Moser[65], qui proposent une analyse très formelle de la simulation des grandes échelles, dans un but explicitement prospectif : il s'agit de savoir si la démarche de minimisation systématique de l'erreur quadratique nous oriente dans la bonne direction. Peut-on vraiment dire que diminuer l'erreur quadratique des estimateurs rend la simulation des grandes échelles plus fiable ? Nous verrons que la réponse est sans doute positive, et que l'utilisation directe de réseaux de neurones est peut-être un moyen de progresser dans la voie d'une simulation des grandes échelles idéale.

Inversibilité et positivité du filtrage

Nous avons vu dans les chapitres précédents qu'on pouvait distinguer deux types de filtrages : les filtrages positifs qui conservent les inégalités et par conséquent les bornes du scalaire, et ceux

⁴Fourni par Olivier Teytaud.

$f(c) =$	$h_2 = (4c(1-c))^2$	$h_3 = (4c(1-c))^3$
$C_0 \sqrt{\bar{c}(1-\bar{c})} \alpha$	$7,148.10^{-2}$	$8,873.10^{-2}$
$\Lambda \sqrt{\bar{c}(1-\bar{c})} \epsilon_{\bar{c}}$	$7,211.10^{-2}$	$9,074.10^{-2}$
$\lambda \epsilon_{\bar{c}}$	$7,554.10^{-2}$	$9,424.10^{-2}$
$C_\alpha \alpha$	$7,829.10^{-2}$	$9.473.10^{-2}$

TAB. 3.4 – Erreurs normalisées commises par de $\overline{f(c)}$ pour différentes fonctions f

<i>Paramètres fondamentaux</i>	<i>Erreur commise</i>
$\bar{c}, \alpha$	0,259
$\bar{c}, \epsilon_{\bar{c}}$	0,215
$\bar{c}, \alpha, \epsilon_{\bar{c}}$	0,192
$\bar{c}, \alpha, \epsilon$	0,258
$\alpha, \epsilon_{\bar{c}}, \epsilon_{\widehat{\bar{c}}}, \widehat{\epsilon_{\bar{c}}}$	0,173
$\bar{c}, \alpha, \epsilon_{\bar{c}}, \epsilon_{\widehat{\bar{c}}}$	0,158
$\bar{c}, \alpha, \epsilon_{\bar{c}}, \widehat{\epsilon_{\bar{c}}}$	0,152
$\bar{c}, \alpha, \epsilon_{\bar{c}}, \epsilon_{\widehat{\bar{c}}}, \widehat{\epsilon_{\bar{c}}}$	0,137

TAB. 3.5 – Erreur commise sur l'estimation de σ_s^2 par des réseaux de neurones pour différents jeux de paramètres

qui ne les conservent pas. Nous allons voir qu'il existe une autre distinction essentielle entre les filtrages.

Considérons un filtrage noté $\bar{c} = G * c$. Dans l'espace spectral, cette égalité devient simplement

$$\tilde{\bar{c}}(\vec{k}) = \tilde{G}(\vec{k}) \cdot \tilde{c}(\vec{k}).$$

Si $\forall \vec{k} \tilde{G}(\vec{k}) \neq 0$, alors on peut écrire

$$\tilde{c} = \frac{\tilde{\bar{c}}}{\tilde{G}}.$$

Il est théoriquement possible dans ces conditions de retrouver le champ c à partir du champ filtré $\bar{c}$. Les filtrages pour lesquels $\tilde{G}$ est non-nulle presque partout sont donc théoriquement inversibles. C'est évidemment le cas du filtrage gaussien, et c'est aussi le cas du filtre porte physique. En effet, même si $\tilde{G}$ s'annule, elle le fait en des points isolés. Et comme en général le spectre de $\tilde{c}$ est continu, on peut retrouver $\tilde{c}$ intégralement en le prolongeant par continuité quand $\tilde{G}$ s'annule⁵. C'est aussi le cas du peigne de Dirac, pour lequel on peut voir les choses un peu différemment. Comme ce filtrage est intimement lié à l'existence d'une grille, plaçons nous sur une grille cubique de N^3 points. La valeur de $\bar{c}$ en un point est la somme (éventuellement pondérée) sur P points de la grille des valeurs prises par le champ c . Si on écrit cette somme pour chacun des N^3 points de la grille, alors on voit apparaître un système de N^3 équations linéaires à N^3 inconnues (les valeurs de c en tous les points de la grille). Ce système est a priori inversible, donc le filtrage l'est.

Le seul filtrage qui ait été évoqué jusqu'à présent et qui ne soit pas inversible est évidemment le filtre porte spectral. Il se trouve que c'est également le seul filtre considéré qui ne soit pas positif. Mais il s'agit d'un hasard, puisqu'on peut parfaitement construire un filtrage qui soit à la fois positif et non inversible. Prenons

$$\tilde{G}(\vec{k}) = \prod_i H\left(\frac{2\pi}{\Delta} - |\vec{k}_i|\right),$$

et considérons le filtrage dont la fonction de transfert T s'écrit

$$T = G^2 = \prod_i H\left(\frac{2}{r_i} \sin \frac{2\pi r_i}{\Delta}\right)^2,$$

et qui par conséquent détermine un filtrage positif. Dans l'espace spectral, $\tilde{T} = \tilde{G} * \tilde{G}$ et comme $\tilde{G}$ est à support compact, $\tilde{T}$ l'est aussi⁶. Un filtrage qui n'est pas inversible est donc un filtrage qui "oublie" des fréquences, et dont la fonction $\tilde{G}$ s'annule sur un ensemble de points de mesure non-nulle. Comme dans le cas de la SGE on considère des filtrages passe-bas, l'ensemble des filtrages à support compact représente bien les filtrages non-inversibles susceptibles d'être utilisés en SGE.

Comme le soulignent Langford et Moser [65], lorsqu'on mène une SGE, on ne prend pas en compte toutes les fréquences jusqu'à celle correspondant à l'échelle de Kolmogorov. C'est d'ailleurs ce qui fait l'intérêt de telles simulations. Cela signifie implicitement que le filtrage fondamental utilisé "oublie" les plus petites fréquences, et qu'il s'agit donc d'un filtrage qui n'est pas inversible. Une véritable simulation des grandes échelles repose sur l'utilisation d'un filtrage non inversible.

⁵Il est vrai que dans notre cas de conditions aux limites périodiques le spectre n'est pas continu. Il existe qu'un ensemble de mesure nulle de valeurs de Δ pour lesquelles les zéros des sinus cardinaux correspondent exactement à un mode. Dans ces cas très précis, on ne peut pas inverser le filtrage.

⁶De la même manière qu'on appelle en anglais le filtrage porte physique le "top hat filter", on pourrait appeler celui-ci le filtrage du chapeau pointu étant donné sa forme.

Formulation optimale de la SGE

La vision de la SGE exposée par Langford et Moser [65] est résumée ici de façon moins formelle. Le filtrage considéré ici est à support compact dans l'espace spectral, il n'est donc pas inversible et est caractérisé par une fréquence de coupure.

Admettons que nous disposions d'un champ grande échelle⁷ (C, V_i) . Comme le filtre n'est pas inversible, il existe *a priori* une infinité de champ complets (c, v_i) dont les grandes échelles sont (C, V_i) . Cet ensemble s'écrit

$$\mathcal{S}(C, V_i) = \{(c, v_i) / \forall \vec{x}, \bar{c}(\vec{x}) = C(\vec{x}) \text{ et } \bar{v}_i(\vec{x}) = V_i(\vec{x})\}.$$

Nous supposons par la suite que nous disposons d'une mesure de probabilité sur cet ensemble. Cela permet de définir une moyenne sur cet ensemble - cette moyenne correspondant en réalité à une moyenne conditionnée par la connaissance de (C, V_i) :

$$\langle \cdot \rangle_{\mathcal{S}(C, V_i)} = \langle \cdot | \forall \vec{x} C(\vec{x}), V_i(\vec{x}) \rangle.$$

On a donc nécessairement les égalités

$$\begin{aligned} C &= \langle \bar{c} \rangle_{\mathcal{S}(C, V_i)} \\ V_i &= \langle \bar{v}_i \rangle_{\mathcal{S}(C, V_i)} \end{aligned}$$

Langford et Moser[65] ont montré qu'il n'existe qu'une seule évolution de (C, V_i) qui reproduise les statistiques en plusieurs point que l'on pourrait obtenir avec des simulations directes. Cette condition s'écrit pour le champ scalaire

$$\langle G(c(\vec{x})) \rangle_{SGE} = \langle G(c(\vec{x})) \rangle_{SND},$$

où $G(c(\vec{x}))$ est une fonction quelconque de valeurs prises par le champ c en différents points, $\langle \cdot \rangle_{SGE}$ est la moyenne effectuée sur les simulations grandes échelles, et $\langle \cdot \rangle_{SND}$ la moyenne obtenue avec des simulations numériques directes.

L'évolution qui fait que la simulation des grandes échelles vérifie la propriété précédente est telle que

$$\begin{cases} \partial_t V_i = \langle \partial_t \bar{v}_i \rangle_{\mathcal{S}(C, V_i)} \\ \partial_t C = \langle \partial_t \bar{c} \rangle_{\mathcal{S}(C, V_i)} \end{cases}$$

Dans le cas d'un filtrage homogène, l'équation d'évolution de C devient :

$$\begin{aligned} \partial_t C &= \langle \partial_t \bar{c} \rangle_{\mathcal{S}(C, V_i)} \\ &= \langle \overline{\partial_t c} \rangle_{\mathcal{S}(C, V_i)} \\ &= \left\langle -\partial_i \bar{v}_i \bar{c} + D \Delta \bar{c} + \overline{f(c)} \right\rangle_{\mathcal{S}(C, V_i)} \\ &= -\partial_i \left\langle \bar{v}_i \bar{c} + \overline{f(c)} \right\rangle_{\mathcal{S}(C, V_i)} + D \Delta C. \end{aligned}$$

⁷C'est-à-dire qu'il pourrait être le filtrage d'un champ réel. Cela lui interdit notamment de présenter dans sa décomposition de Fourier des fréquences supérieures à la fréquence de coupure du filtre.

Il est possible de mener le même type de calcul pour le champ de vitesse. Autrement dit, la meilleure façon d'estimer les termes "non-fermés" au sens de la présentation classique de la SGE et de prendre leur moyenne sur l'ensemble $\mathcal{S}(C, V_i)$ - et donc leur moyenne conditionnée. Les équations d'évolution de (C, V_i) deviennent donc

$$\partial_i V_i = 0 \quad (3.25)$$

$$\partial_t V_i + \partial_j \langle \overline{v_i v_j} | \forall \vec{x} V_i(\vec{x}) \rangle = \nu \Delta V_i - \frac{1}{\rho} \partial_i \langle \overline{p} | \forall \vec{x} V_i(\vec{x}) \rangle \quad (3.26)$$

$$\partial_t C + \partial_i \langle \overline{v_i c} | \forall \vec{x} C(\vec{x}), V_i(\vec{x}) \rangle = D \Delta C + \left\langle \overline{f(c)} \middle| \forall \vec{x} C(\vec{x}), V_i(\vec{x}) \right\rangle, \quad (3.27)$$

avec

$$\partial_i \partial_j \langle \overline{v_i v_j} | \forall \vec{x} V_i(\vec{x}) \rangle = -\frac{1}{\rho} \Delta \langle \overline{p} | \forall \vec{x} V_i(\vec{x}) \rangle. \quad (3.28)$$

S'il était possible de mener une telle simulation, elle serait "en tous sens idéale" [65].

Estimateurs locaux

On suppose en général qu'il existe une distance R telle que les quantités statistiques calculées en deux points distants de R ou plus soient indépendantes. Considérons deux quantités statistiques a et b calculées respectivement en deux points $\vec{x}_0$ et $\vec{x}_1$ séparés d'une distance supérieure à R et π un ensemble de paramètres quelconque. Il semble intuitivement⁸ évident que si les situations statistiques en $\vec{x}_0$ et en $\vec{x}_1$ sont indépendantes on puisse écrire

$$\langle a | \pi, b \rangle = \langle a | \pi \rangle. \quad (3.29)$$

Nous allons supposer que la relation (3.29) est vérifiée pour a et b calculés en des points suffisamment distants simplement parce que cela correspond à l'intuition. Cela revient à écrire par exemple que

$$\langle \overline{v_i v_j}(\vec{x}) | \forall \vec{r} V_i(\vec{r}) \rangle = \langle \overline{v_i v_j}(\vec{x}) | V_i(\vec{x} + \vec{r}), \forall \vec{r} \text{ tq } |\vec{r}| \leq R \rangle.$$

Les autres moyennes conditionnées peuvent être transformées de la même manière. Dans ces conditions, une SGE idéale ne demande plus que la connaissance d'estimateurs locaux - c'est ainsi que nous les désignerons par la suite. Cette simplification doit cependant être relativisée : les estimateurs locaux sont toujours des fonctions d'une infinité de paramètres.

On peut penser cependant qu'il est possible d'approcher ces estimateurs locaux par des estimateurs optimaux basés sur un jeu de paramètres suffisamment pertinent. Si le jeu de paramètres caractérise suffisamment la situation locale, on peut penser qu'il va s'approcher nettement de l'estimateur local⁹.

⁸En réalité, le fait que a et b soient indépendants (ce qui par définition donne que $p(a, b) = p(a) p(b)$) n'est pas suffisant pour en déduire que la propriété (3.29) est vraie. Supposer cette propriété vraie est donc plus fort que supposer l'indépendance de a et b : c'est supposer que a et b sont indépendants conditionnellement à tout jeu de paramètres π (selon les termes consacrés).

⁹On sait grâce à la propriété 3.3 que l'erreur commise par l'estimateur local sera toujours inférieure à celle de l'estimateur optimal, car les paramètres de celui-ci sont forcément des fonctions des paramètres de l'estimateur local.

Utilisation directe des réseaux de neurones

Pour obtenir les erreurs commises par les estimateurs optimaux qui ont été présentées ci-dessus des réseaux de neurones ont été “entraînés”. Des matrices donnant les poids respectifs de chaque réseau ont été obtenues. Il est envisageable, dans le cadre d’une simulation des grandes échelles, d’utiliser les poids calculés et donc les réseaux de neurones obtenus pour estimer la variance de sous-maille en tout point puis, par l’intermédiaire d’une hypothèse β , les grandes échelles de n’importe que taux de réaction $\overline{f(c)}$. La simple donnée des poids du réseau suffit à pouvoir mener cette opération. Les réseaux de neurones - avec cependant certaines restrictions - sont donc utilisables directement en simulation des grandes échelles. De la même manière, on peut envisager d’utiliser des réseaux de neurones pour estimer des quantités comme $\overline{v_i c}$ ou même $\overline{v_i v_j}$. En cherchant des jeux de paramètres toujours plus efficaces pour l’estimation d’une de ces quantités, on diminue l’erreur quadratique commise sur l’estimation des termes inconnus et on se rapproche des estimateurs locaux, donc d’une simulation des grandes échelles idéales. Le point de vue de Langford et Moser justifie ainsi la démarche de minimisation systématique de l’erreur quadratique.

L’utilisation directe de réseaux de neurones est cependant contraignante. Tout d’abord, il faut avoir accès à des données issues de simulations numériques directes pour pouvoir entraîner les réseaux de neurones. Cette condition limite le plus haut nombre de Reynolds que l’on puisse espérer atteindre avec une simulation “guidée” par des réseaux de neurones. D’autre part, un réseau entraîné l’est pour une situation très particulière, caractérisée par un nombre de Reynolds, un nombre de Schmidt et une taille de filtrage Δ fixés. Une façon de contourner cet obstacle est d’intégrer ces trois grandeurs dans les paramètres. Pour la taille de filtrage par exemple, on peut compter sur le fait que le comportement des estimateurs semble varier continuellement lorsque Δ varie (voir figure 3.15). Le réseau ferait alors lui-même la “transition” entre les différentes tailles de filtrage. Une autre solution pour prendre en compte la variation de Δ serait sans doute de n’utiliser comme paramètres que des grandeurs adimensionnalisées par Δ .

Quoi qu’il en soit, les réseaux de neurones représentent une voie possible vers une simulation des grandes échelles idéale...

Conclusion

Le premier aspect de ce travail est une illustration supplémentaire de ce que la simulation numérique directe d'écoulement peut apporter à l'étude de la turbulence et des phénomènes qui lui sont associés, comme le mélange turbulent. On peut en effet considérer que ces simulations permettent de reproduire des comportements qui sont observés expérimentalement. Les spectres obtenus pour le scalaire et les fonctions de structure présentent indubitablement des lois de puissance. Les exposants de ces lois, qui caractérisent l'intermittence, sont tout à fait en accord avec les données expérimentales disponibles. Il a de plus été observé expérimentalement que le champ de vitesse pour les nombres de Reynolds atteints en simulation ne présentait pas de loi de puissance pour le spectre ou les fonctions de structure - ce que laissent également à penser les simulations. Ceci renforce la confiance que l'on peut placer dans les données issues de simulations numériques directes. Disposer de ces données permet donc de sonder finement la validité de certaines hypothèses, de vraiment mesurer la qualité d'une modélisation - quand aucun dispositif expérimental ne le permet.

Dans une première partie, ce travail a permis de caractériser la situation de mélange produite par la méthode d'injection de blocs de scalaire de position aléatoire, développée par Marc Elmo [2]. On peut affirmer que cette dernière est désormais bien caractérisée - en compilant les nombreux résultats obtenus par Marc Elmo et ceux issus des simulations réalisées pour cette thèse, d'une plus grande résolution. On sait notamment lier (grossièrement) la variance du scalaire (et donc le degré de mélange obtenu) aux caractéristiques mêmes de l'injection. Il est également possible de distinguer plus clairement ce qui est propre à l'injection de ce qui semble relever de comportements plus généraux.

Ainsi la technique d'injection ne semble pas altérer l'apparition de comportements auto-similaires, ni l'intermittence du scalaire. Elle ne semble pas non plus (à la condition qu'elle reste modérée) affecter le développement d'une portion linéaire pour la diffusion conditionnelle. Ces comportements, que l'on observe également avec une injection par gradient moyen imposé, peuvent donc être considérés comme relativement universels.

Au contraire, il apparaît qu'avec la méthode d'injection utilisée ici, la densité de probabilité de l'accroissement de scalaire ne relaxe pas vers une gaussienne quand la distance croît - ce qu'on peut lier à la densité de probabilité du scalaire lui-même. Le comportement de la diffusion conditionnelle près des bornes ne semble pas dépendre des caractéristiques détaillées de l'injection mais être inhérent au fait que l'on injecte le scalaire par blocs de valeurs extrêmes - ce que confirme une analyse théorique de Dopazo. De façon générale, et comme le montre également le dernier chapitre (importance des distributions β notamment), la technique d'injection donne un rôle fondamental aux bornes du scalaire, qu'elle impose. On peut envisager que certaines de ces caractéristiques puissent être celles d'une espèce chimique dans un écoulement turbulent. Les espèces chimiques sont en général introduites dans un milieu en structures de valeurs extrêmes. Leurs bornes sont donc fixées *a priori*. Autant d'éléments que ne reproduit pas l'injection par

gradient moyen et qui distinguent les deux méthodes.

Pour finir, cette technique d'injection devrait pouvoir constituer un outil très efficace pour l'étude de l'anisotropie du scalaire. L'injection de blocs parallélépipédiques offrirait en effet un contrôle sur l'anisotropie du forçage. En faisant varier le temps entre chaque injection, il devrait être possible d'observer l'isotropisation du scalaire, et d'en déterminer la vitesse. L'anisotropie du scalaire étant actuellement considérée comme une étape incontournable dans la compréhension fine du mélange turbulent, une telle étude serait sans doute riche d'enseignements.

Dans la seconde partie de ce travail, la modélisation de sous-maille en Simulation des Grandes Échelles a été explorée. L'erreur quadratique a été utilisée pour mesurer la qualité des estimateurs de quantités inconnues. Pour un ensemble donné de paramètres, nous avons montré qu'il existe un unique estimateur qui minimise cette erreur quadratique. Cet type d'estimateur, appelé estimateur optimal a été étudié de façon théorique. Un estimateur optimal pour la densité de sous-maille a également été défini. L'utilisation de ces notions a permis de mieux comprendre le fonctionnement du modèle de Cook et Riley en évaluant séparément les conséquences des diverses hypothèses qui le sous-tendent.

La première d'entre elles est que l'on peut approcher la densité de sous-maille grâce à des distributions β . Dans le cas où la variance de la densité de sous-maille σ_s^2 est connue, il s'avère que les distributions β soient un excellent choix¹⁰. La variance de sous-maille σ_s^2 n'étant en pratique pas connue, on dispose de plusieurs estimateurs pour l'approcher. Comme nous l'avons montré, cette connaissance imparfaite de σ_s^2 produit une dispersion en variance. Elle a pour conséquence de rendre le choix des distributions β en apparence moins pertinent. Cependant, plus l'erreur commise sur l'estimation de la variance est faible, plus la dispersion en variance est faible et meilleur s'avère le choix des distributions β . En définitive, l'amélioration du modèle ne passe sans doute que par l'amélioration de l'estimation de la variance de sous-maille. Le choix des distributions β ne semble pas devoir être remis en cause.

La variance de sous-maille reste cependant une quantité difficile à estimer. L'hypothèse d'auto-similarité invoquée pour justifier l'estimateur de Cook et Riley de σ_s^2 semble mise en défaut. Une stratégie de minimisation systématique de l'erreur quadratique apporte des améliorations parfois sensibles aux estimateurs. Ces derniers restent cependant loin d'atteindre les performances des estimateurs optimaux approchés numériquement.

La démarche de minimisation de l'erreur quadratique, qui devrait mener à une SGE "en tous sens idéale" au sens de Langford et Moser, pousse à enrichir la famille de paramètres utilisés pour la modélisation - notamment pour l'estimation de σ_s^2 . Cette démarche conduit à l'utilisation de réseaux de neurones, seuls à même d'approcher des estimateurs optimaux faisant intervenir un nombre élevé de paramètres. Les réseaux permettent de calculer l'erreur commise par un estimateur optimal. Cette erreur renseigne alors sur la pertinence intrinsèque du jeu de paramètres choisi. Les calculs effectués montrent que l'ajout de paramètres diminue en général l'erreur commise de façon appréciable. À la condition de pouvoir utiliser des données issues de simulations numériques directes pour l'apprentissage, on peut imaginer des SGE dans lesquelles les termes inconnus seraient estimés par des réseaux de neurones. Devoir mener une simulation directe pour entraîner les réseaux est cependant une condition très contraignante, qui limite le nombre de

¹⁰Les résultats présentés dans ce travail laissent supposer l'existence d'une propriété mathématique des distributions β . Si tel était bien le cas, cela justifierait l'utilisation de distributions β dans toutes les situations où le champ considéré est borné, et notamment pour un scalaire réactif.

Reynolds que l'on est susceptible d'atteindre. Plus généralement, les réseaux peuvent être utilisés pour fermer des termes inconnus, mais seulement si l'on dispose de suffisamment de données pour l'apprentissage¹¹.

Il convient de remarquer cependant que les résultats présentés dans ce travail ont été obtenus pour un type d'injection particulier et un nombre de Schmidt très proche de l'unité (en conservant un champ de vitesse identique, il faudrait augmenter de beaucoup la résolution pour accéder à de grands nombres de Schmidt). Il serait intéressant de mener une étude comparable pour d'autres types de situation (un scalaire soumis à un taux de réaction ou un autre type d'injection par exemple).

Les simulations numériques, telle la SGE, ont d'ores et déjà une grande importance pratique. Des modélisations fines sont en général nécessaires pour assurer la fiabilité des simulations et leur capacité à prendre en compte des phénomènes aussi complexes que la combustion. Dans ce cadre, l'approche par densité présumée est utilisée avec succès (qu'il s'agisse d'une densité de probabilité, de sous-maille ou autre) pour la modélisation de phénomènes liés aux réactions chimiques. Elle devrait à l'avenir permettre de simuler des situations encore plus complexes comme en combustion diphasique, où les propriétés statistiques des gouttelettes devront nécessairement être finement prises en compte. Ce travail propose justement une analyse rigoureuse d'une approche par densité présumée et des méthodes systématiques pour l'étude des problèmes de modélisation - permettant l'évaluation de la pertinence des paramètres utilisés, ou l'indication de directions possibles d'amélioration. Autant d'éléments qui laissent espérer que les résultats présentés ici trouveront une grande utilité à l'avenir.

¹¹C'est peut-être vers la météorologie qu'il faudrait se tourner pour trouver d'autres situations où des données abondantes pourraient permettre l'emploi des réseaux de neurones.

Annexe A : Equations d'évolution

Evolution de c^2

$$\begin{aligned}
 \partial_t c^2 &= 2c \partial_t c \\
 &= 2Dc \partial_i \partial_i c - 2c v_i \partial_i c \\
 &= D\partial_i (2c \partial_i c) - 2D(\partial_i c)^2 - \partial_i (v_i c^2)
 \end{aligned}$$

Soit

$$\partial_t c^2 + \partial_i (v_i c^2) = D\Delta c^2 - \epsilon_c$$

Évolution de la dissipation du scalaire

$$\begin{aligned}
 \partial_t (\partial_i c)^2 &= 2 \partial_i c \partial_t \partial_i c \\
 &= 2 \partial_i c \partial_i [D\Delta c - \partial_j (v_j c)] \\
 &= 2 \partial_i c [D\partial_j \partial_j \partial_i c - \partial_j (v_j \partial_i c) - \partial_i v_j \partial_j c] \\
 &= D\Delta (\partial_i c)^2 - 2D(\partial_j \partial_j \partial_i c)^2 - \partial_j (v_j (\partial_i c)^2) - \partial_i c \partial_i v_j \partial_j c
 \end{aligned}$$

Soit

$$\partial_t (\partial_i c)^2 + \partial_j (v_j (\partial_i c)^2) = D\Delta (\partial_i c)^2 - 2D(\partial_j \partial_i c)^2 - \partial_i c \partial_i v_j \partial_j c$$

ou

$$\partial_t \epsilon_c + \partial_i (v_i \epsilon_c) = D\Delta \epsilon_c - 4D^2 (\partial_j \partial_i c)^2 - 2D \partial_i c \partial_i v_j \partial_j c$$

Évolution de la transformée de Fourier

Considérons l'équation de conservation de la quantité de mouvement, en l'absence de forces extérieures *a priori*, et en écoulement incompressible. Elle peut s'écrire

$$\partial_t v_i + \partial_j (v_j v_i) = \nu \Delta v_i - \frac{1}{\rho} \partial_i p.$$

D'où en prenant la divergence de cette équation,

$$\partial_t (\partial_i v_i) + \partial_i \partial_j (v_i v_j) = \nu \Delta (\partial_i v_i) - \frac{1}{\rho} \Delta p.$$

Soit, comme $\partial_i v_i = 0$,

$$\Delta p = \rho \partial_i v_j \partial_j v_i.$$

On prend alors la transformée de cette équation :

$$\tilde{p} = -\rho \frac{k_l k_j}{k^2} \tilde{v}_l * \tilde{v}_j.$$

Cette première équation nous donne la pression en fonction du champ de vitesse, ce qui va nous permettre d'éliminer la pression de l'équation de conservation de la quantité de mouvement. Si on prend la transformée de cette dernière, on obtient¹² :

$$\begin{aligned} \partial_t \tilde{v}_i + i k_j \tilde{v}_j * \tilde{v}_i &= -i k_i \frac{\tilde{p}}{\rho} - k^2 \tilde{v}_i \\ \partial_t \tilde{v}_i + \delta_{il} i k_j \tilde{v}_j * \tilde{v}_l &= -i k_i \frac{k_l k_j}{k^2} \tilde{v}_j * \tilde{v}_l - k^2 \tilde{v}_i \end{aligned}$$

Soit, finalement, une équation d'évolution dans l'espace de Fourier ne faisant intervenir que le champ de vitesse

$$\partial_t \tilde{v}_i + \left(\delta_{il} - \frac{k_i k_l}{k^2} \right) (i k_j \tilde{v}_l * \tilde{v}_j) + \nu k^2 \tilde{v}_i = 0$$

Pour ce qui est du scalaire, la simple transformée de Fourier de son équation d'évolution donne immédiatement l'équation d'évolution de sa transformée de Fourier :

$$\partial_t \tilde{c} + i k_j \tilde{v}_j * \tilde{c} + D k^2 \tilde{c} = 0.$$

Équation d'évolution de σ_s^2 , la variance de sous-maille

Le champ scalaire filtré $\bar{c}$ obéit ici à l'équation d'évolution

$$\partial_t \bar{c} + \partial_i (\bar{v}_i \bar{c}) = D \Delta \bar{c} + \overline{f(c)}. \quad (\text{A-1})$$

Dans ce cas, l'équation d'évolution de $\bar{c}^2$ s'obtient par

$$\begin{aligned} \partial_t c^2 + \partial_i (v_i c^2) &= D \Delta c^2 - \epsilon_c \\ \partial_t \bar{c}^2 + \partial_i (\bar{v}_i \bar{c}^2) &= D \Delta \bar{c}^2 - \bar{\epsilon}_c \end{aligned}$$

L'équation d'évolution de $\bar{c}^2$ s'obtient en multipliant l'équation (A-1) par $\bar{c}$

$$\begin{aligned} \partial_t c + \partial_i (v_i c) &= D \Delta c + \overline{f(c)} \\ \partial_t \bar{c} + \partial_i (\bar{v}_i \bar{c}) &= D \Delta \bar{c} + \overline{f(c)} \\ 2 \bar{c} \partial_t \bar{c} + 2 \bar{c} \partial_i (\bar{v}_i \bar{c}) &= 2 D \bar{c} \Delta \bar{c} + 2 \bar{c} \overline{f(c)} \\ \partial_t \bar{c}^2 + 2 \bar{c} \partial_i (\bar{v}_i \bar{c}) &= D \Delta \bar{c}^2 - 2 D \partial_i \bar{c}^2 + 2 \bar{c} \overline{f(c)}. \end{aligned}$$

Pour finir, l'équation de $\sigma_s^2 = \bar{c}^2 - \bar{c}^2$ est obtenue en faisant la différence des deux équations précédentes :

$$\begin{aligned} \partial_t (\bar{c}^2 - \bar{c}^2) + \partial_i (\bar{v}_i \bar{c}^2) - 2 \bar{c} \partial_i (\bar{v}_i \bar{c}) &= D \Delta \bar{c}^2 - D \Delta \bar{c}^2 - 2 D \overline{(\partial_i c)^2} + 2 D \partial_i \bar{c}^2 + 2 \bar{c} \overline{f(c)} - 2 \bar{c} \overline{f(c)} \\ \partial_t \sigma_s^2 + \partial_i (\bar{v}_i \bar{c}^2) - 2 \bar{c} \partial_i (\bar{v}_i \bar{c}) &= D \Delta \sigma_s^2 - 2 D \overline{(\partial_i c)^2} + 2 D \partial_i \bar{c}^2 + 2 \bar{c} \overline{f(c)} - 2 \bar{c} \overline{f(c)}. \end{aligned}$$

¹²On désigne la convolution par le signe $*$.

Évolution de la densité de probabilité du scalaire

Ce paragraphe permet d'introduire le formalisme probabiliste, et en dérive les équations d'évolution pour la fonction de densité de probabilité du champ de concentration.

Pour une géométrie donnée, on considère Ω , l'ensemble des écoulements possibles, caractérisés par leurs champs à l'instant initial, par exemple. On dispose d'une mesure de probabilité μ , qui associe à toute partie de Ω sa probabilité. Par définition, la mesure de probabilité est telle que

$$\mu(\emptyset) = 0,$$

$$\mu(\Omega) = 1,$$

et

$$\mu(\cup_i A_i) = \sum_i \mu(A_i).$$

La donnée de Ω et de μ définit ce qu'on appellera un *ensemble statistique*. La mesure de probabilité donne en fait le poids statistique de chaque évènement considéré.

Une variable aléatoire est une fonction qui associe à chaque élément de Ω une valeur réelle. Dans le cadre de la turbulence, cela peut être la vitesse en un point donné, ou toute autre quantité (énergie cinétique, variance du scalaire). On notera $\langle f \rangle$ ou $\int f d\mu$ la moyenne statistique de la variable aléatoire f .

Soit a , une variable aléatoire. On définit sa fonction densité de probabilité par

$$p(A) = \frac{\partial P(a \leq A)}{\partial A},$$

où $P(a \leq A)$ est la probabilité que a prenne une valeur inférieure à A .

Or $P(a \leq A) = \mu(\mathcal{E})$, avec $\mathcal{E} = \{e \in \Omega | a(e) \leq A\}$. Et $\mu(\mathcal{E}) = \langle 1_{\mathcal{E}} \rangle$, où $1_{\mathcal{E}}$ vaut 1 si $e \in \mathcal{E}$ et 0 sinon. On peut aussi l'écrire

$$P(a \leq A) = \langle H(A - a(e)) \rangle,$$

et donc

$$p(A) = \left\langle \frac{\partial H(A - a)}{\partial A} \right\rangle = \langle \delta(A - a) \rangle.$$

La grandeur $c(\vec{x}, t)$ est une variable aléatoire, dont on notera $p(\Gamma, \vec{x}, t)$ la fonction densité de probabilité que $c(\vec{x}, t)$ prenne la valeur Γ . L'évolution de chaque écoulement obéit à des équations d'évolution déterminées. La densité de probabilité évolue naturellement avec les champs. Elle obéit ainsi à une équation d'évolution qui se déduit des équations d'évolution des champs.

On sait en effet que $p(\Gamma, \vec{x}, t) = \langle \delta(\Gamma - c(\vec{x}, t)) \rangle$. On posera pour simplifier $\delta(\Gamma - c(\vec{x}, t)) = \Pi$. On peut donc écrire que $\partial_t p(\Gamma, \vec{x}, t) = \langle \partial_t \Pi \rangle$. Or

$$\begin{aligned} \partial_t \Pi &= -\frac{\partial \Pi}{\partial \Gamma} \partial_t c \\ &= \frac{\partial}{\partial \Gamma} (\Pi v_j \partial_j c - \Pi D \Delta c - \Pi f(c)) \\ &= -v_j \left(-\frac{\partial \Pi}{\partial \Gamma} \partial_j c \right) - \frac{\partial}{\partial \Gamma} (\Pi D \Delta c + \Pi f(c)) \\ &= -v_j \partial_j \Pi - D \frac{\partial}{\partial \Gamma} (\Pi \Delta c) - \frac{\partial}{\partial \Gamma} (\Pi f(c)) \end{aligned}$$

Et

$$\begin{aligned}
\Delta \Pi &= \partial_i \partial_i \Pi \\
&= -\frac{\partial}{\partial \Gamma} (\partial_i (\Pi \partial_i c)) \\
&= -\frac{\partial}{\partial \Gamma} (\partial_i \Pi \partial_i c + \Pi \Delta c) \\
&= \frac{\partial^2}{\partial \Gamma^2} (\Pi \partial_i c \partial_i c) - \frac{\partial}{\partial \Gamma} (\Pi \Delta c)
\end{aligned}$$

Soit

$$-D \frac{\partial}{\partial \Gamma} (\Pi \Delta c) = D \Delta \Pi - \frac{\partial^2}{\partial \Gamma^2} (\Pi \epsilon),$$

avec $\epsilon = D \partial_i c \partial_i c$, la dissipation. Il vient donc

$$\partial_t \Pi + v_j \partial_j \Pi = D \Delta \Pi - \frac{\partial^2}{\partial \Gamma^2} (\Pi \epsilon) - \frac{\partial}{\partial \Gamma} (\Pi f(c)).$$

Et on prend alors la moyenne de cette équation pour trouver l'équation d'évolution de p .

$$\partial_t p + \partial_j \langle v_j \Pi \rangle = D \Delta p - \frac{\partial^2}{\partial \Gamma^2} (\langle \Pi \epsilon \rangle) - \frac{\partial}{\partial \Gamma} (\langle \Pi f(c) \rangle).$$

On peut réécrire l'équation sachant que pour deux variables aléatoires a et b ,

$$\begin{aligned}
\langle a \delta(B - b) \rangle &= \left\langle \int A \delta(A - a) \delta(B - b) \, dA \right\rangle \\
&= \int A \langle \delta(A - a) \delta(B - b) \rangle \, dA \\
&= \int A p(A, B) \, dA \\
&= \int A p(A|B) p(B) \, dA \\
&= \langle A|B \rangle p(B),
\end{aligned}$$

où $\langle A|B \rangle$ est la moyenne de a sachant que $b = B$. Il est à noter cependant que cette moyenne n'est définie que tant que $p(B)$ n'est pas nulle. L'écriture la plus sûre est donc

$$\langle a \delta(B - b) \rangle = \int A p(A, B) \, dA,$$

avec $p(A, B)$ densité de probabilité conjointe que $a = A$ et $b = B$. Cette écriture est cependant peu utilisée dans la littérature ; on lui préfère bien souvent les expressions avec les moyennes conditionnées, plus parlantes.

L'équation d'évolution de p s'écrit donc finalement

$$\partial_t p(\Gamma) + \partial_j (\langle v_j | \Gamma \rangle p(\Gamma)) = D \Delta p(\Gamma) - \frac{\partial^2}{\partial \Gamma^2} (\langle \epsilon | \Gamma \rangle p(\Gamma)) - \frac{\partial}{\partial \Gamma} (f(\Gamma) p(\Gamma)).$$

Elle fait apparaître un terme de convection $\partial_j (\langle v_j | \Gamma \rangle p)$, un terme de diffusion $D \Delta p$, qui tend à homogénéiser la densité de probabilité, un terme de diffusion inverse traduisant le mélange par processus diffusif, $-\frac{\partial^2}{\partial \Gamma^2} (\langle \epsilon | \Gamma \rangle p)$, ainsi qu'un terme lié à une éventuelle réaction chimique mais qui est fermé. Le terme de diffusion inverse traduit un courant de densité de probabilité vers c_0 , la moyenne du scalaire. Le signe négatif du terme produit l'effet inverse d'un terme de diffusion usuel. Celui-ci concentre donc la fonction de probabilité au lieu de l'étaler. Il traduit la diminution de l'écart-type, et fait très logiquement intervenir la dissipation ϵ .

Dans le cas où l'ensemble statistique est homogène – c'est à dire que les quantités statistiques construites à partir des champs ne dépendent pas de $\vec{x}$ – l'équation se réduit à

$$\partial_t p(\Gamma) = -\frac{\partial^2}{\partial \Gamma^2} (\langle \epsilon | \Gamma \rangle p(\Gamma)) - \frac{\partial}{\partial \Gamma} (f(\Gamma) p(\Gamma)).$$

Il faut également noter qu'en moyennant directement l'équation

$$\partial_t \Pi = -v_j \partial_j \Pi - D \frac{\partial}{\partial \Gamma} (\Pi \Delta c) - \frac{\partial}{\partial \Gamma} (\Pi f(c)),$$

on obtient

$$\partial_t p(\Gamma) + \partial_j (\langle v_j | \Gamma \rangle p(\Gamma)) = -\frac{\partial}{\partial \Gamma} (\langle D \Delta c | \Gamma \rangle p(\Gamma)) - \frac{\partial}{\partial \Gamma} (f(\Gamma) p(\Gamma)),$$

équation qui fait intervenir la diffusion conditionnelle. En écoulement homogène, cette équation devient

$$\partial_t p(\Gamma) = -\frac{\partial}{\partial \Gamma} (\langle D \Delta c | \Gamma \rangle p(\Gamma)) - \frac{\partial}{\partial \Gamma} (f(\Gamma) p(\Gamma)).$$

Annexe B : Méthode pseudo-spectrale

Comme nous l'avons vu dans l'annexe A, l'équation d'évolution du champ de vitesse dans l'espace spectral s'écrit

$$\partial_t \tilde{v}_i + \left(\delta_{il} - \frac{k_i k_l}{k^2} \right) i k_j \tilde{v}_l * \tilde{v}_j + \nu k^2 \tilde{v}_i = 0,$$

où $*$ représente une convolution. Pour le scalaire passif, elle s'écrit

$$\partial_t \tilde{c} + i k_j \tilde{v}_j * \tilde{c} + D k^2 \tilde{c} = 0.$$

Plaçons nous maintenant dans le cadre d'un écoulement avec conditions aux limites périodiques, un cube de côté L , comme au paragraphe 1.1. Des tels écoulements permettent d'obtenir des solutions qui ressemblent aux écoulements dans un domaine infini, tout en étant simulables numériquement. Dans un tel cadre, les champs étudiés (vitesse ou scalaire) peuvent alors se décomposer en série de Fourier, ce qui s'écrit pour $c(\vec{x})$

$$c(x, y, z) = \sum_{i=-\infty}^{+\infty} \sum_{j=-\infty}^{+\infty} \sum_{k=-\infty}^{+\infty} a_{i,j,k} e^{i\omega_0(i x + j y + k z)},$$

avec $\omega_0 = \frac{2\pi}{L}$. Les coefficients $a_{i,j,k}$ sont alors donnés par la relation

$$a_{i,j,k} = \frac{1}{L^3} \int c(x, y, z) e^{-i\omega_0(i x + j y + k z)}.$$

Considérons pour simplifier le cas où le nombre de Schmidt est pris égal à l'unité. La théorie de Kolmogorov stipule qu'en dessous d'une certaine échelle η , les champs ne présentent plus de variation significative. Cela se traduit ici par le fait que les champs se décomposent sur un nombre fini de modes de Fourier. En pratique, l'échelle la plus petite pour laquelle les champs présentent encore des variations significatives est de l'ordre de η . Une telle échelle correspond à une fréquence spatiale de $\frac{1}{\eta}$. Si on définit

$$\kappa \simeq \frac{L}{2\eta},$$

la décomposition du champ $c(\vec{x})$ s'écrit alors

$$c(x, y, z) = \sum_{i=-\kappa}^{+\kappa} \sum_{j=-\kappa}^{+\kappa} \sum_{k=-\kappa}^{+\kappa} a_{i,j,k} e^{i\omega_0(i x + j y + k z)}. \quad (\text{B-1})$$

Nous allons maintenant définir une grille dans l'espace physique et voir qu'il existe une bijection entre l'ensemble des valeurs prises par le champ sur les noeuds de cette grille et les coefficients $a_{i,j,k}$. Nous allons définir $N = 2\kappa + 1$. Nous noterons $c(m, p, q)$ la valeur du champ $c(\vec{x})$ au noeud de la grille repéré par m, p, q . Un tel noeud a pour coordonnées

$$\begin{aligned} x &= m \frac{L}{N} \\ y &= p \frac{L}{N} \\ z &= q \frac{L}{N}. \end{aligned}$$

Il est alors possible d'écrire que

$$c(m, p, q) = \sum_{i=-\kappa}^{+\kappa} \sum_{j=-\kappa}^{+\kappa} \sum_{k=-\kappa}^{+\kappa} a_{i,j,k} e^{\frac{2i\pi}{N}(im+jp+kq)},$$

et on peut vérifier que

$$a_{i,j,k} = \frac{1}{N^3} \sum_{m=0}^{N-1} \sum_{p=0}^{N-1} \sum_{q=0}^{N-1} c(m, p, q) e^{-\frac{2i\pi}{N}(im+jp+kq)}. \quad (\text{B-2})$$

La connaissance des $c(m, p, q)$ permet ainsi de retrouver les coefficients de Fourier de la décomposition de $c(\vec{x})$ en un nombre fini de modes. Cette transformation est appelée *Transformée de Fourier Discrète* (abrégé TFD).

La définition des coefficients $a_{i,j,k}$ peut s'étendre à n'importe quel entier pour i, j ou k . Les coefficients ne sont dans ce cas pas tous différents, car ils vérifient les égalités suivantes :

$$\begin{aligned} a_{i,j,k}^* &= a_{-i,-j,-k} \\ a_{i,j,k} &= a_{i+N,j,k} \\ a_{i,j,k} &= a_{i,j+N,k} \\ a_{i,j,k} &= a_{i,j,k+N} \end{aligned}$$

Dans la pratique d'ailleurs les coefficients sont indexés par des entiers compris entre 0 et $N - 1$. Grâce aux relations précédentes, il est relativement simple de relier ces coefficients aux coefficients de Fourier. Les valeurs du champ sur les noeuds de la grille peuvent finalement s'écrire

$$c(m, p, q) = \sum_{i=0}^{N-1} \sum_{j=0}^{N-1} \sum_{k=0}^{N-1} a_{i,j,k} e^{\frac{2i\pi}{N}(im+jp+kq)}, \quad (\text{B-3})$$

et cette transformation est appelée *Transformée de Fourier Discrète Inverse*.

La TFD est utilisée dans la méthode pseudo-spectrale pour pouvoir calculer plus efficacement les convolutions. Une convolution dans l'espace spectral correspond en effet à un simple produit dans l'espace physique. Le principe est donc de passer en espace physique grâce à une TFD inverse, d'y effectuer la multiplication sur les noeuds de la grille et d'utiliser une TFD directe pour retourner

en espace spectral. C'est à ce processus que la méthode doit son nom, car elle utilise tout de même l'espace physique. Il convient de souligner cependant que si deux champs sont représentés par un nombre fixé de modes, leur produit nécessite *a priori* un plus grand nombre de modes pour être représenté. Il est tout à fait normal que les termes non-linéaires posent problème dans ce cadre, car ce sont eux qui sont responsables du transfert de l'énergie du signal vers les plus grandes fréquences (processus de cascade). Ces termes génèrent donc des hautes fréquences, et dans le cas où l'avancement en temps est discret comme dans toute simulation, il n'est pas étonnant que cela pose problème. Dans notre cas, plus de modes qu'il n'est strictement nécessaire sont utilisés pour représenter les champs. Ces modes supplémentaires sont utiles lors du calcul de la convolution pour représenter les termes non-linéaires de façon plus fiable. En pratique, le tiers des modes de plus haute fréquence ne sont utilisés que pour le calcul des termes non-linéaires.

Si les modes de plus haute fréquence du produit (et qui ne sont pas représentés) sont importants les coefficients $a_{i,j,k}$ n'ont plus vraiment de sens. Ils n'en conservent que si les modes qui ne sont pas représentables apportent une contribution négligeable au champ. Il existe des algorithmes très efficaces pour calculer une TFD (directe ou inverse). Ceux-ci sont appelés algorithmes de Transformée de Fourier Rapide, ce qu'on voit souvent abrégé FFT pour "Fast Fourier Transform". Leur efficacité explique que le passage par l'espace physique pour le calcul des convolutions soit en définitive moins coûteux. Ces algorithmes imposent le plus souvent de prendre pour N une puissance de 2. La plus haute fréquence n'est alors représentée que par un coefficient, réel. Une TFD ne permet pas de calculer les valeurs du champ considéré en tout point de l'espace, mais seulement sur les noeuds de la grille. Il ne faut pas oublier pour autant que le champ est bien défini par ses modes de Fourier en tout point de l'espace physique. Pour pouvoir calculer le champ en un point $\vec{x}$ donné, il faut relier chaque coefficient au mode de Fourier qui lui correspond et sommer les contributions de tous les modes selon l'équation (B-1).

Pour effectuer un filtrage porte physique, il n'est pas équivalent de se placer dans l'espace spectral (où l'on dispose de toutes les informations sur le champ) et dans l'espace physique en utilisant les valeurs sur les noeuds d'une grille. Plaçons-nous dans le cas d'un champ unidimensionnel $c(x)$ et considérons le filtrage porte physique de taille Δ . Le champ filtré en x_0 s'écrit

$$\bar{c}(x_0) = \frac{1}{\Delta} \int_{-\frac{\Delta}{2}}^{\frac{\Delta}{2}} c(x_0 + x) \, dx$$

Dans le cas où l'on peut écrire $c(x)$ comme une série de Fourier, avec un nombre fini de modes on peut écrire

$$\begin{aligned} \bar{c}(x_0) &= \frac{1}{\Delta} \int_{-\frac{\Delta}{2}}^{\frac{\Delta}{2}} \sum_{-K}^K a_k e^{i\omega_0 k(x_0+x)} \, dx \\ &= \sum_{-K}^K a_k e^{i\omega_0 k x_0} \frac{1}{\Delta} \int_{-\frac{\Delta}{2}}^{\frac{\Delta}{2}} e^{i\omega_0 k x} \, dx \\ &= \sum_{-K}^K a_k e^{i\omega_0 k x_0} \frac{1}{i\omega_0 k \Delta} \left[e^{i\omega_0 k x} \right]_{-\frac{\Delta}{2}}^{\frac{\Delta}{2}} \\ &= \sum_{-K}^K a_k e^{i\omega_0 k x_0} \operatorname{sinc} \left(\frac{1}{2} k \omega_0 \Delta \right) \end{aligned}$$

Annexe C : Distributions β

Fonction beta d'Euler

La fonction beta d'Euler peut-être définie comme suit[45] :

$$B(a, b) = \int x^{a-1} (1-x)^{b-1} dx = \frac{\Gamma(a)\Gamma(b)}{\Gamma(a+b)},$$

où $\Gamma(x)$ est la fonction gamma d'Euler, qui vérifie $\Gamma(x+1) = x\Gamma(x)$. Pour plus de détails sur les fonctions β , les distributions β et la fonction Γ , on consultera avec avantage le livre de Laurent Schwarz[45]. De l'expression des fonctions beta d'Euler à l'aide de la fonction $\Gamma(x)$, on déduit immédiatement que $B(a, b) = B(b, a)$. Les identités que l'on va établir ci-dessous se généralisent donc sans problème.

On peut donc exprimer

$$\begin{aligned} B(a+1, b) &= \frac{\Gamma(a+1)\Gamma(b)}{\Gamma(a+b+1)} \\ &= \frac{a\Gamma(a)\Gamma(b)}{(a+b)\Gamma(a+b)} \\ &= \frac{a}{a+b} B(a, b). \end{aligned}$$

Donc on peut de façon générale écrire que

$$\begin{aligned} B(a+n, b) &= \frac{a+n-1}{a+b+n-1} B(a+n-1, b) \\ &= \prod_{k=0}^{n-1} \frac{a+k}{a+b+k} B(a, b). \end{aligned}$$

Distribution β normalisée

On introduit la fonction

$$g(\Gamma, a, b) = \frac{\Gamma^{a-1} (1-\Gamma)^{b-1}}{B(a, b)}$$

dont les moments sont donnés par

$$\begin{aligned} \int x^n g(x, a, b) dx &= \frac{B(a+n, b)}{B(a, b)} \\ &= \prod_{k=0}^{n-1} \frac{a+k}{a+b+k} \end{aligned}$$

Si on souhaite que g ait comme moyenne $\bar{c}$, alors on écrit $\bar{c} = \frac{a}{a+b}$, ce qui fournit immédiatement

$$b = \frac{a}{\bar{c}} - a \quad (\text{C-1})$$

Si on souhaite que sa variance soit σ^2 , alors il faut écrire que

$$\sigma^2 + \bar{c}^2 = \frac{B(a+2, b)}{B(a, b)} = \frac{a+1}{a+b+1} \cdot \bar{c}$$

En remplaçant b grâce à l'équation C-1, on obtient

$$\begin{aligned} \sigma^2 + \bar{c}^2 &= \bar{c} \frac{a+1}{a+b+1} \\ \text{Or } a+b &= \frac{a}{\bar{c}} \\ \text{Donc } \sigma^2 + \bar{c}^2 &= \bar{c}^2 \frac{a+1}{a+\bar{c}} \\ &= \bar{c}^2 \left(1 + \frac{1-\bar{c}}{a+\bar{c}} \right) \\ \text{D'où } a+\bar{c} &= \frac{\bar{c}^2(1-\bar{c})}{\sigma^2} \\ a &= \bar{c} \left(\frac{\bar{c}(1-\bar{c})}{\sigma^2} - 1 \right) \end{aligned}$$

Distribution beta sur l'intervalle $[-1, 1]$

Une distribution beta non normalisée sur l'intervalle $[-1, 1]$ s'écrit (grâce à un simple changement de variable)

$$\left(\frac{1+x}{2} \right)^{a-1} \left(\frac{1-x}{2} \right)^{b-1}.$$

L'intégrale de cette dernière sur son intervalle de définition vaut

$$\begin{aligned} \int_{-1}^1 \left(\frac{1+x}{2} \right)^{a-1} \left(\frac{1-x}{2} \right)^{b-1} dx &= 2 \int_0^1 c^{a-1} (1-c)^{b-1} dc \\ &= 2 B(a, b), \end{aligned}$$

en effectuant le changement de variable défini par $x = 2c - 1$. La distribution β normalisée sur $[-1, 1]$ s'écrit donc [50]

$$\frac{1}{2 B(a, a)} \left(\frac{1+x}{2} \right)^{a-1} \left(\frac{1-x}{2} \right)^{b-1}.$$

Parenté des distributions β et des gaussiennes

La propriété suivante¹³ permet d'affirmer que lorsque la variance devient petite, et que l'influence des bornes se fait par conséquent moins sentir, une distribution β se comporte comme une gaussienne.

¹³Due à Didier Felbacq.

Proposition C-5

$$e^{-x^2} = \lim_{n \rightarrow \infty} \left[\left(1 + \frac{x}{\sqrt{n}}\right) \left(1 - \frac{x}{\sqrt{n}}\right) \right]^n$$

◇ *Preuve.* L'exponentielle vérifie la propriété

$$e^x = \lim_{n \rightarrow \infty} \left[1 + \frac{x}{n}\right]^n.$$

Si on remplace x par $-x^2$, celle-ci devient

$$e^{-x^2} = \lim_{n \rightarrow \infty} \left[1 - \frac{x^2}{n}\right]^n = \lim_{n \rightarrow \infty} \left[\left(1 + \frac{x}{\sqrt{n}}\right) \left(1 - \frac{x}{\sqrt{n}}\right) \right]^n.$$

□

Plus précisément, cela signifie que si on prend une distribution β centrée dont les bornes sont en $-\sqrt{n}$ et en $\sqrt{n}$, alors quand n tend vers $+\infty$, la distribution β tend vers une gaussienne. Autrement dit, si on considère des bornes fixes, lorsque la variance de la distribution β tend vers 0 et que l'influence des bornes se fait moins sentir, elle prend l'allure d'une gaussienne.

Universalité des distributions β

Ce paragraphe présente une façon très simple de vérifier si les distributions β peuvent effectivement prétendre à une certaine universalité ou non. L'idée est simplement de répéter sur un autre champ que c les opérations ayant mené aux courbes 3.3, 3.4 et 3.5.

Pour espérer trouver des résultats identiques, il faut que le champ choisi présente une certaine régularité et soit représentable sur un nombre fini de modes de Fourier - en gros le même nombre que celui nécessaire à la représentation de c . Ceci pour pouvoir utiliser l'interpolation grâce aux modes de Fourier. Le champ étudié doit surtout avoir pour bornes 0 et 1, puisqu'on suppose que les bornes sont l'élément déterminant dans l'apparition des distributions *beta*.

Le champ choisi ici est c^2 , car il correspond aux critères donnés ci-dessus, mais ses caractéristiques statistiques sont cependant radicalement différentes de celles de c - densité de probabilité, spectre, variance, moyenne, ... etc.

Les résultats sont présentés sur les figures C-1 et C-2 et ils ne laissent aucun doute sur le fait que les lois beta sont dans ce cas également très proches des estimateurs optimaux de la densité de sous-maille. Les résultats obtenus sont très comparables à ceux exposés pour le champ c . Ces résultats laissent à penser qu'ils reflètent une propriété fondamentale des distributions β . Des résultats partiels qui ne sont pas présentés ici laissent penser que les distributions β sont tout aussi pertinentes pour le champ $\sqrt{c}$.

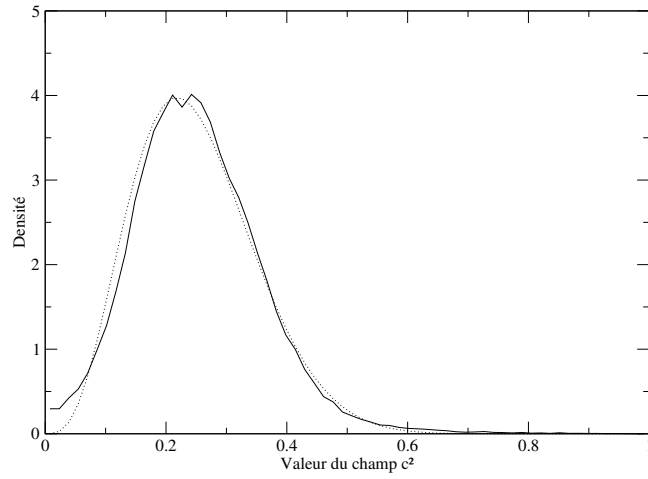


FIG. C-1 – Comparaison entre $\langle d_s | \bar{c}, \sigma_s^2 \rangle$ et une distribution β

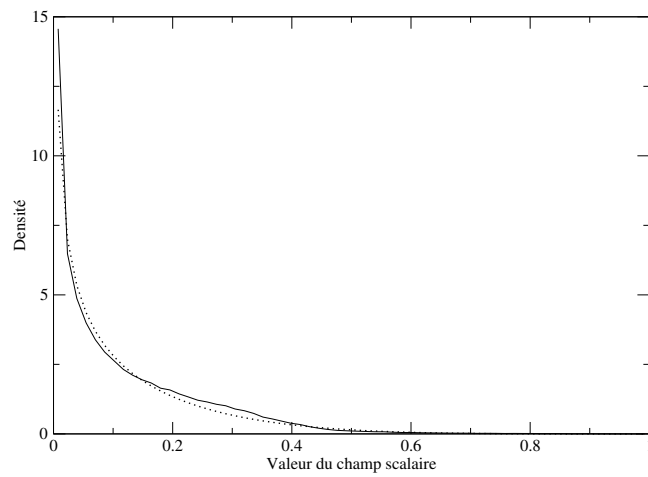


FIG. C-2 – Comparaison entre $\langle d_s | \bar{c}, \sigma_s^2 \rangle$ et une distribution β

Annexe D : Méthode de calcul des estimateurs optimaux

Pour mener des calculs d'erreur, il faut accéder à des quantités de la forme $\langle a|b, c \rangle$ à partir d'un ensemble de triplets (a, b, c) extraits des données de la simulation numérique directe. Pour approcher $\langle a|b, c \rangle$ il est possible d'utiliser des fonctions constantes par morceaux, autrement dit des histogrammes. Pour cela, il faut discrétiser l'espace à deux dimensions dans lequel évoluent b et c en cellules. L'ensemble des triplets peut alors être partitionné en sous-ensembles, chaque sous-ensemble correspondant à une cellule. A chaque cellule est attribuée comme valeur la moyenne de a calculée sur le sous-ensemble correspondant. C'est ainsi qu'on dispose d'une fonction constante par morceaux : lorsqu'on veut une estimation de $\langle a|b, c \rangle$, il faut calculer la cellule désignée par b et c et prendre pour valeur de l'estimation celle associée à la cellule. Le principe général de la méthode exposé, il convient d'examiner en détail la technique utilisée pour s'assurer qu'elle permet d'approcher du mieux possible $\langle a|b, c \rangle$.

La discrétisation de l'espace ne se fait tout d'abord pas à pas constant. La discrétisation a au contraire été choisie pour essayer de faire que toutes les cellules aient un sous-ensemble correspondant comportant grossièrement le même nombre de triplets. On évite ainsi d'avoir des cellules surpeuplées dont les statistiques convergent parfaitement et des cellules sous-peuplées où les statistiques ne sont pas fiables.

Plaçons nous dans l'espace B où évolue le paramètre b . Les paramètres étudiés sont tels que $b \in [0, b_{max}]$. Cet espace est discrétisé en N cellules. Si on définit b_i comme étant la frontière entre la cellule i et la cellule $i + 1$, la cellule i correspond à des valeurs de b telles que $b_{i-1} \leq b \leq b_i$. On prend pour les frontières des cellules

$$b_i = b_{max} \left(\frac{i}{N} \right)^p.$$

L'entier p est choisi de telle façon que chaque cellule de l'espace B reçoive un nombre de triplet à peu près comparable. C'est parfois très grossier, mais cela permet de discrétiser plus finement là où sont présents beaucoup de triplets. Par exemple, pour le paramètre α on peut prendre $p = 2$ alors que pour ϵ_{τ} qui présente une densité de probabilité très piquée près de 0, $p = 4$ est plus indiqué.

Une fois ce processus effectué pour le paramètre c dans l'espace C ¹⁴, la discrétisation de l'espace $B \times C$ est alors naturelle. Chaque cellule est repérée par deux entiers i et j et correspond

¹⁴En prenant le même entier N mais avec un entier p différent.

à des valeurs b et c telles que

$$\begin{aligned} b_i &\leq b \leq b_{i+1} \\ c_j &\leq c \leq c_{j+1}. \end{aligned}$$

C'est ainsi qu'est obtenue une discrétisation dont la taille des cellules, bien que n'étant pas constante, est contrôlée par l'entier N . Il convient maintenant de choisir convenablement cet entier. Imaginons que nous disposons d'un ensemble de triplets (a_k, b_k, c_k) , et qu'une fois notre histogramme constitué, nous évaluons sa qualité en calculant l'erreur quadratique qu'il commet sur notre échantillon de triplets. Si on note $H(i, j)$ la valeur prise par notre histogramme pour la cellule repérée par i et j , cette erreur s'écrit

$$E = \frac{1}{N} \sum_{k=1}^N [H(i(b_k), j(c_k)) - a_k]^2.$$

Si N est trop petit, cette erreur est grande parce que la fonction constante par morceaux n'est pas une approximation assez souple. Il est bien entendu qu'il faut prendre N relativement grand. Si N est choisi très grand (N^2 par exemple de l'ordre du nombre de triplets disponibles pour les statistiques) il finira par ne rester que très peu de triplets par cellule. Dans le cas où il resterait un seul triplet par cellule l'erreur obtenue serait ainsi nulle !

Pour déterminer le meilleur N , il faut diviser l'ensemble des triplets en deux parties. L'une sert pour calculer l'histogramme, la seconde pour calculer l'erreur. Cette erreur est dite *erreur en généralisation*. Si N est pris trop grand, l'histogramme s'éloigne de la "véritable" fonction $\langle a | b, c \rangle$ et par conséquent l'erreur en généralisation augmente. Cela signifie que l'histogramme est trop dépendant de l'échantillon sur lequel on le calcule. L'erreur en généralisation est représentée figure D-1, et l'on constate effectivement qu'elle connaît un minimum – dont la position dépend du type de triplet considéré. **La valeur choisie pour N est celle qui minimise l'erreur en généralisation.**

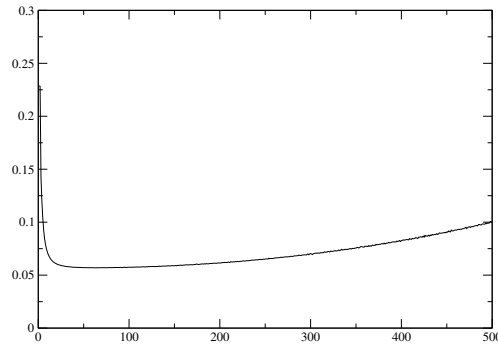


FIG. D-1 – Erreur en généralisation commise par l'histogramme en fonction de N pour l'estimation de $\langle \sigma_s^2 | \bar{c}, \epsilon_{\bar{c}} \rangle$

Cette méthode relativement complexe nous permet de garantir que les histogrammes ainsi construits sont les meilleurs possibles avec les données dont nous disposons. Le seul problème de cette méthode est qu'elle nécessite beaucoup de statistiques pour fournir un résultat convergé – mais qu'il reste difficile de garantir que pour chacune des "cases" de l'histogramme il y ait suffisamment d'évènements.

Bibliographie

- [1] A.N. Kolmogorov. Local turbulence structure in incompressible viscous fluid at very high reynolds numbers. *Dokl. AN SSSR*, 30 :299–303, 1941.
- [2] Marc Elmo. *Étude du mélange en turbulence isotrope par fonctions densité de probabilité et interprétation géométrique*. PhD thesis, École Centrale de Lyon, 2000.
- [3] U. Frisch. *Turbulence : the legacy of A.N. Kolmogorov*. Cambridge University Press, 1995.
- [4] Z. Warhaft. Passive scalars in turbulent flows. *Annu. Rev. Fluid. Mech.*, 32 :203–240, 2000.
- [5] S. Corrsin. On the spectrum of isotropic temperature fluctuations in isotropic turbulence. *J. Appl. Phys.*, 22 :469–473, 1951.
- [6] A.M. Obukhov. Structure of the temperature field in turbulent flow. *Izv. Akad. Nauk. SSSR, Geogr. Geophys.*, 13 :58–69, 1949.
- [7] L. Mydlarski and Z. Warhaft. Passive scalar statistics in high-péclet-number grid turbulence. *J. Fluid Mech.*, 358 :135–175, 1998.
- [8] K.R. Sreenivasan. The passive scalar spectrum and the obukhov-corrsin spectrum. *Phys. Fluids*, 8 :189–196, 1996.
- [9] Monin and Yaglom. *Statistical Fluid Mechanics*, volume 2. The MIT Press, 1975.
- [10] A.M. Yaglom. On the local structure of a temperature field in a turbulent flow. *Dokl. AN SSSR*, 69 :743–746, 1949.
- [11] H.O. Rasmussen. A new proof of kolmogorov four-fifth law. *Phys. Fluids*, 11 (11) :3495–3498, 1999.
- [12] G.K. Batchelor. Small-scale variation of convected quantities like temperature in a turbulent field. part 1. general discussion and the case of small conductivity. *J. Fluid Mech.*, 5 :113, 1959.
- [13] M.C. Jullien, P. Castiglione, and P. Tabeling. Observation of the batchelor regime. *Phys. Rev. Lett.*, 85 :3636, 2000.
- [14] R.H. Kraichnan. Anomalous scaling of a randomly advected passive scalar. *Phys. Rev. Lett.*, 72 :1016, 1994.
- [15] K. Gawedzki and A. Kupiainen. Anomalous scaling of the passive scalar. *Phys. Rev. Lett.*, 75 :3834–3837, 1995.
- [16] S. Chen and R.H. Kraichnan. Simulations of a randomly advected passive scalar field. *Phys. Fluids*, 10 :2867–2884, 1998.
- [17] R.W. Bilger. *Turbulent Reacting Flows*, chapter Turbulent flows with nonpremixed reactants, pages 65–113. Springer Verlag, 1980.
- [18] F. Gao and E.E. O’Brien. A large-eddy simulation scheme for turbulent reacting flows. *Phys. Fluids A*, 5 (6) :1282–1284, 1993.

-
- [19] A.W. Cook and J.J. Riley. A subgrid model for equilibrium chemistry in turbulent flows. *Phys. Fluids*, 6 (8) :2868–2870, 1994.
- [20] P. Moin and K. Mahesh. Direct numerical simulations : a tool in turbulence research. *Annu. Rev. Fluid Mech.*, 30 :539–578, 1998.
- [21] S.B. Pope. *Turbulent Flows*. Cambridge University Press, 2000.
- [22] E.E. O’Brien. *Turbulent Reacting Flows*, chapter The pdf approach for turbulent reactive flows, pages 185–218. Springer Verlag, 1985.
- [23] L. Valiño. A field monte carlo formulation for calculating the probability density function of en single scalar in a turbulent flow. *Flow, Turb. Comb.*, 1998 :157–172, 1998.
- [24] C. Dopazo. Relaxation of initial probability density functions in the turbulent convection of scalar fields. *Phys. Fluids*, 22 (1) :20, 1979.
- [25] T.S. Lundgren. Distribution functions in the statistical theory of turbulence. *Phys. Fluids*, 10 (5) :969–975, 1967.
- [26] Jayesh and Z. Warhaft. Probability distribution, conditional dissipation and transport of passive temperature fluctuations in grid-generated turbulence. *Phys. Fluids A*, 4 :2292–2307, 1992.
- [27] B.R. Lane, O.N. Mesquita, S.R. Meyers, and J.P. Gollub. Probability distributions and thermal transport in a turbulent grid flow. *Phys. Fluids A*, 5 (9) :2255–2263, 1993.
- [28] V. Eswaran and S.B. Pope. Direct numerical simulations of the turbulent mixing of a passive scalar. *Phys. Fluids*, 31 (3) :506–520, 1987.
- [29] M.R. Overholt and S.B. Pope. Direct numerical simulation of a passive scalar with imposed mean gradient in isotropic turbulence. *Phys. Fluids*, 8 (11) :3128–3148, 1996.
- [30] F.A. Jaber, R.S. Miller, and P. Givi. Conditional statistics in turbulent scalar mixing and reaction. *AIChE Journal*, 42 (4) :1149–1152, 1996.
- [31] F. Gao. Mapping closure and non-gaussianity of the scalar probability density functions in isotropic turbulence. *Phys. Fluids A*, 3 (10) :2438–2444, 1991.
- [32] E.E. O’Brien and T.L. Jiang. The conditional dissipation rate of an initially binary scalar in homogeneous turbulence. *Phys. Fluids A*, 3 (12) :3121–3123, 1991.
- [33] W.P. Jones. *Closure Strategies for Turbulent and Transitionnal Flows*, chapter The Joint Scalar Probability Density Function. Cambridge University Press, 2001.
- [34] J. Smagorinsky. General circulation experiments with the primitive equations. *Mon. Weather Rev.*, 91 (3) :99–164, 1963.
- [35] M. Lesieur and O. Métais. New trends in large-eddy simulations of turbulence. *Ann. Rev. Fluid. Mech.*, 28 :45–82, 1996.
- [36] M. Germano. Turbulence : the filtering approach. *J. Fluid Mech.*, 238 :325–336, 1992.
- [37] C. Meneveau and J. Katz. Scale invariance and turbulence models for large eddy simulation. *Annu. Rev. Fluid Mech.*, 32 :1–32, 2000.
- [38] S.B. Pope. Pdf methods for turbulent reactive flows. *Prog. Energy Combust. Sci.*, 11 :119–192, 1985.
- [39] P.J. Colucci, F.A. Jaber, P. Givi, and S.B. Pope. Filtered density function for large eddy simulation of turbulent reacting flows. *Phys. Fluids*, 10 (2) :499–515, 1998.

- [40] S.M. De Bruyn Kops, J.J. Riley, G. Kosály, and A.W. Cook. Investigation of modeling for non-premixed turbulent combustion. *Flow, Turb. and Comb.*, 60 :105–122, 1998.
- [41] J. Réveillon and L. Vervisch. Response of the dynamic les model to heat release induced effects. *Phys. Fluids*, 8 (8) :2248–2250, 1996.
- [42] C.D. Pierce and P. Moin. A dynamic model for subgrid-scale variance and dissipation rate of a conserved scalar. *Phys. Fluids*, 10 (12) :3041–3044, 1998.
- [43] J. Jimenez, A. Liñán, M.M. Rogers, and F.J. Higuera. *A priori* testing of subgrid models for chemically reacting non-premixed turbulent shear flows. *J. Fluid Mech.*, 1997 :149–171, 1997.
- [44] C. Wall, B.J. Boersma, and P. Moin. An evaluation of the assumed beta probability density function subgrid-scale model for large eddy simulation of nonpremixed turbulent combustion with heat release. *Phys. Fluids*, 12 (10) :2522–2529, 2000.
- [45] L. Schwartz. *Méthodes mathématiques pour la physique*. Editions Hermann, 1965.
- [46] M. Germano, U. Piomelli, P. Moin, and W.H. cabot. A dynamic subgrid-scale eddy viscosity model. *Phys. Fluids A*, 3 (7) :1760–1765, 1991.
- [47] D.K. Lilly. A proposed modification of the germano subgrid-scale closure method. *Phys. Fluids A*, 4 (3) :633–634, 1992.
- [48] L. Shao, J.P. Perlat, and J.P. Bertoglio. Large-eddy simulation of the longtime behaviour of the weakly compressible turbulence. In *Proc. 9th ISTP Conf.*, 1996.
- [49] F. Gauthier, L. Shao, R. Wunenburger, and J.P. Bertoglio. An improved two-point closure for weakly compressible turbulence and comparisons with large-eddy simulations. In *Proc. 11th Turbulent Shear Flow*, 1997.
- [50] M. Elmo, A. Moreau, J.P. Bertoglio, and V.A. Sabel'nikov. Mixing in isotropic turbulence with scalar injection and applications to subgrid modeling. *Flow, Turbulence and Combustion*, 65 :113–131, 2000.
- [51] A. Moreau, M. Elmo, and J.P. Bertoglio. A priori tests of subgrid models for the scalar fluctuations in statistically stationary isotropic turbulence. In *IUTAM Symposium on Turbulent Mixing and Combustion*, 3-6 juin 2000.
- [52] A. Pumir. A numerical study of the mixing of a passive scalar in three dimension in presence of a mean gradient. *Phys. Fluids*, 6 (6) :2118–2132, 1994.
- [53] M. Holzer and E. Siggia. Turbulent mixing of a passive scalar. *Phys. Fluids*, 6 (5) :1820–1837, 1994.
- [54] Z. Warhaft and J.L. Lumley. An experimental study of temperature fluctuations in grid-generated turbulence. *J. Fluid Mech.*, 88 :659–684, 1978.
- [55] Jayesh, C. Tong, and Z. Warhaft. On temperature spectra in grid turbulence. *Phys. Fluids*, 6 :306–312, 1994.
- [56] W. Appel. *Mathématiques pour la physique et les physiciens*. H&K Éditions, 2002.
- [57] J.N. Gence. *Sur la description des mélanges turbulents par densités de probabilité*. Communication personnelle, 1994.
- [58] A. Moreau, M. Elmo, and J.P. Bertoglio. Analyse des modèles de sous-maille par densité présumée en simulation des grandes échelles. *Soumis aux C. R. Acad. Sci. Paris*, 2002.

-
- [59] S.S. Girimaji and Y. Zhou. Analysis and modeling of subgrid scalar mixing using numerical data. *Phys. Fluids*, 8 (5) :1224–1236, 1996.
 - [60] C.M. Bishop. *Neural Networks for Pattern Recognition*. Oxford University Press, 1995.
 - [61] C. Lee, J. Kim, and D. Babcock. Application of neural networks for turbulence control for drag reduction. *Phys. Fluids*, 9 (6) :1740–1747, 1997.
 - [62] F. Giralt, A. Arenas, J. Ferre-Giné, R. Rallo, and G.A. Kopp. The simulation and interpretation of free turbulence with a cognitive neural system. *Phys. Fluids*, 12 (7) :1826–1835, 2000.
 - [63] V. Périquet, A. Moreau, S. Carles, J.P. Schermann, and C. Desfrancois. Cluster size effects upon anion solvation of n-heterocyclic molecules and nucleic acid bases. *J. of Electron Spect. and Rel. Phen.*, 106 :141–151, 2000.
 - [64] R.J. Adrian. Stochastic estimation of sub-grid scale motions. *Appl. Mech. Rev.*, 43 (5) :S214–S218, 1990.
 - [65] J.A. Langford and R.D. Moser. Optimal les formulations for isotropic turbulence. *J. Fluid Mech.*, 398 :321–346, 1999.

Résumé

Le problème du mélange de quantités scalaires en écoulement turbulent est abordé dans un but d'application à la modélisation des petites échelles pour la prédiction d'écoulements réactifs. Des Simulations Numériques Directes d'un scalaire passif convecté par la turbulence ont tout d'abord été menées. Les données recueillies permettent d'étudier une situation de mélange obtenue grâce à une injection aléatoire de scalaire frais. Un comportement auto-similaire du scalaire a été mis en évidence, et la caractérisation de l'intermittence qui en découle est en bon accord avec les résultats expérimentaux disponibles. Par ailleurs, ces données permettent de sonder finement la pertinence d'hypothèses en modélisation pour la Simulation des Grandes Échelles. A cet effet, la notion d'estimateur optimal est introduite et généralisée. Ces estimateurs permettent alors d'étudier séparément les conséquences des hypothèses qui peuvent sous-tendre un modèle. Le modèle de Cook et Riley est ainsi exploré, qui fait appel à des distributions bêta. L'utilisation de ces fonctions se révèle en général fondée - dans certaines conditions elle s'avère même être un choix remarquable. Il est souligné que la qualité du modèle repose sur l'estimation de la variance de sous-maille. Les estimateurs usuels de cette quantité sont comparés, et des améliorations sont proposées grâce à une démarche de minimisation systématique de l'erreur quadratique. Cette démarche conduit naturellement à introduire des réseaux de neurones simples. Il est mis en évidence que ces réseaux permettent d'approcher certains estimateurs optimaux ce qui ouvre la voie à leur utilisation comme outil pour la mise au point de Simulations des Grandes Échelles.

Mots-clés : Scalaire passif, turbulence, simulation des grandes échelles, réseaux de neurones

Abstract

The problem of the mixing of passive scalars in a turbulent flow has been investigated, in order to model the small scales. Direct Numerical Simulations of this situation have been performed, with a periodical injection of fresh scalar. The data generated has been used to investigate the statistically stationary situation obtained. The spectra and the structure functions of the scalar actually present a clear power-law range. The intermittent properties of the scalar have thus been computed and are in good agreement with the available experimental results. In order to take chemical reactions into account in Large Eddy Simulations, a fine modelling of the subgrid scales is required. Many assumptions are then made, which deserve to be evaluated using the DNS data. The notion of optimal estimate in the sense of the quadratic error is introduced and extended. This allows to evaluate separately the assumptions which underlie the model of Cook and Riley. This model uses beta distributions, which is found to be well-chosen. Under some conditions, the beta distributions are even an excellent choice. The estimation of the subgrid variance then appears to be the flaw of the model. The usual estimates for this quantity are compared and improvements are proposed, using a strategy of systematic reduction of the quadratic error. This finally leads to the use of neural networks, which are able to approach optimal estimates and thus could be used as tools in the search for a reliable Large Eddy Simulation.

Keywords : Passive scalar, turbulence, large eddy simulation, neural networks