



THÈSE DE DOCTORAT UNIVERSITÉ PARIS-EST

Spécialité :
SCIENCES ET TECHNIQUES DE L'ENVIRONNEMENT

Présentée et soutenue par
JEAN-MATTHIEU HAUSSAIRE

Méthodes variationnelles d'ensemble itératives pour l'assimilation de données non-linéaire. Application au transport et la chimie atmosphérique.

soutenue le 23 juin 2017

Membres du jury :

<i>Rapporteurs :</i>	Dr Emmanuel COSME	-	Université Grenoble Alpes, IGE, Grenoble
	Dr Yann MICHEL	-	Météo-France, CNRM, Toulouse
<i>Examineurs :</i>	Dr Étienne MÉMIN	-	INRIA, Fluminance, Rennes
	Pr Serge GRATTON	-	ENSEEIH, Toulouse
	Dr Emanuele EMILI	-	CERFACS, Toulouse
<i>Invité :</i>	Dr Gilles FORET	-	Université Paris-Est Créteil, LISA, Créteil
<i>Directeur de thèse :</i>	Pr Marc BOCQUET	-	École des Ponts Paristech, CERE, Marne-la-Vallée

Remerciements



Je remercie en premier lieu Emmanuel Cosme et Yann Michel pour avoir accepté d'être mes rapporteurs ainsi qu'Étienne Mémin, Serge Gratton, Emanuele Emili et Gilles Foret pour avoir accepté de prendre part au jury. Je remercie également mon directeur Marc Bocquet de m'avoir proposé ce sujet et encadré avec attention au cours de la préparation de cette thèse. Il a su m'inculquer un peu de sa rigueur qui me manquait tant. Merci également au CEREa et à l'école des Ponts de m'avoir accueilli en leur sein et financé mes travaux.

Un grand merci à Yelva Roustan qui m'a éclairé dans les profondeurs obscures de Polyphème, à Sylvain Doré et Youngseob Kim qui m'ont dépêtré des innombrables problèmes techniques inhérents au calcul scientifique, ainsi qu'à Véronique Dehlinger pour son soutien dans les méandres des tâches administratives. Je remercie enfin Vivien Mallet qui a accepté que je l'aide à encadrer ses travaux pratiques. J'ai beaucoup appris à ses côtés et j'ai pris énormément de plaisir lors de ces trop brefs événements. Finally, a great thanks to Pavel Sakov and Yun Liu with whom we coauthored two articles. The success of my thesis heavily relies on them.

Je tiens également à remercier l'ensemble de mes collègues du CEREa, ceux qui étaient là avant moi comme ceux qui resteront après. Un grand merci tout particulièrement à Laëtitia, qui a partagé avec moi l'intégralité de la durée de cette thèse. Elle a su supporter mon humour parfois un peu cinglant, mais qui n'était que la preuve de mon affection pour elle. J'espère qu'elle me le pardonnera, à défaut de m'avouer qu'elle a su apprendre à l'apprécier. Merci également à Anthony et Vincent, sur qui j'ai pu asseoir ma domination, arme à la main. Merci à tous ceux avec qui j'ai pu partager un café après le déjeuner, petit moment de détente et d'allégresse, pendant lesquels des sujets aussi variés que redondants ont été abordés. Je pense ainsi à Arnaud, qui à défaut de boire du café, nous abreuvait de sa culture éclectique, à Jérôme, qui a partagé son amour des verts, à Laëtitia, qui préférait nous faire part de ses escapades à vélo, à Carole, qui s'est parfois jointe à nous quand elle n'était pas occupée à battre les autres membres du CEREa au tennis, à Matthieu, qui a su me mater lors de la visite des mines de Saint-Étienne, et à Alban, qui a défié mes capacités logiques avec ses énigmes. J'ai apprécié ces moments avec vous, même si j'ai souvent été le premier à y couper court.

Un grand merci à mes collocataires, Pierre, Aloïs et Charles-Henri, avec qui j'ai partagé mon quotidien pendant une majorité de cette thèse. Vous avez su cantonner les discussions relatives à nos recherches respectives au minimum, ce que j'ai énormément apprécié. Même si leurs capacités étaient limitées, je les remercie pour m'avoir apporté un entraînement quotidien à un sport qui me tient à cœur. Encore toutes mes félicitations à Pierre qui a su terminer sa thèse dans les temps, et je souhaite bien du courage aux deux autres qui, je n'en doute pas, y arriveront également.

Je présente également tous mes remerciements à ma famille, qui m'a toujours aidé et soutenu. Merci Alain, pour m'avoir assisté dans cette dernière ligne droite en repérant d'un oeil d'aigle toutes les dernières coquilles qui traînaient dans ce manuscrit. Je sais qu'il en reste, mais on a fait de notre mieux ! Merci Jocelyne pour tes conseils avisés, qui m'ont permis de mettre en pratique le vieil adage *mens sana in corpore sano*. Merci à Marc, qui m'a fourni le moteur de recherche dont j'ai abusé au cours de ma thèse et qui œuvre pour rendre cette planète un petit peu meilleure. J'espère que tu réussiras dans ton entreprise et je suis fier de toi. Tous mes remerciements et félicitations à Luc, qui m'a donné une magnifique petite nièce que j'espère voir le plus souvent possible ! Enfin, mes pensées vont à ma grand-mère que je félicite pour sa forme éclatante, sa générosité et sa joie de vivre. Qu'elles vivent à tout jamais.

Enfin, je précise que rien de ceci n'aurait été possible sans celle qui m'a motivé pour commencer cette thèse en premier lieu, qui m'a soutenu au quotidien pendant les moments difficiles et avec qui je n'ai pas assez souvent partagé les moments de réussite, celle qui guide chacune de mes décisions depuis bientôt 10 ans. Cette thèse nous aura éprouvé en nous imposant une longue séparation mais son achèvement augure, par bien des aspects, de jours heureux.

De tout mon cœur, merci à tous.

Table des matières

Introduction	1
1 Modèles de chimie atmosphérique	3
1.1 Météorologie et transport	5
1.1.1 Équations de Navier-Stokes	5
1.1.2 Définition et importance de la couche limite atmosphérique	6
1.1.3 Transport : advection et diffusion	6
1.1.4 Dépôt sec et humide	6
1.2 Chimie de l’ozone troposphérique	6
1.2.1 Cycle de Leighton	6
1.2.2 Régimes NO _x et COV limité	7
1.3 Incertitudes sur les données d’entrée des modèles de chimie atmosphérique	8
1.3.1 Émissions	9
1.3.2 Condition initiale	9
1.3.3 Conditions aux limites	9
1.3.4 Météorologie dans les CTMs	10
1.4 Intégration numérique	10
1.4.1 Modélisation numérique du transport	10
1.4.2 Modélisation numérique de la chimie	12
1.4.3 Splitting	14
1.5 Résumé	14
2 Assimilation de données	17
2.1 Observations	19
2.2 BLUE et le filtre de Kalman	20
2.2.1 L’algorithme	20
2.2.2 Ses limitations et les premières solutions	21
2.3 Méthodes variationnelles	22
2.3.1 3D-Var et parallèle avec BLUE	22
2.3.2 4D-Var et parallèle avec le filtre de Kalman	23
2.4 Méthodes ensemblistes	24
2.4.1 Filtres de Kalman d’ensemble : dérivations et algorithmes	24
2.4.2 Variantes de l’EnKF	29
2.4.3 Le lisseur de Kalman d’ensemble	31
2.4.4 Localisation	34
2.4.5 Inflation	36
2.5 Méthodes variationnelles d’ensemble	39
2.5.1 Méthodes hybrides	40
2.5.2 4DEnVar	40
2.5.3 Le lisseur de Kalman d’ensemble itératif : IEnKS	42
2.5.4 Formulation en contrainte faible : l’IEnKF-Q	46
2.6 Validation et indicateurs statistiques	55
2.6.1 Erreur quadratique : RMSE	55

2.6.2	Corrélation de Pearson	56
2.6.3	Biais	56
2.6.4	Indicateurs statistiques χ^2 et RCRV	56
3	Développement de modèles jouets pour le transport et la chimie atmosphérique	59
3.1	Un premier modèle simplifié de météorologie : L95	61
3.2	Un modèle avec transport d'un traceur inerte : L95-T	61
3.3	Introduction de la chimie atmosphérique : L95-GRS	63
3.3.1	Définition du L95-GRS	64
3.3.2	Validation du modèle	67
3.3.3	Intérêt et nécessité d'un tel modèle	71
3.4	Un bref aperçu de QG-Chem	72
3.4.1	La météorologie avec un modèle quasi-géostrophique	72
3.4.2	La chimie atmosphérique avec RACM	72
3.4.3	Comparaison avec le L95-GRS	73
4	Évaluation des méthodes EnVar 4D itératives sur des modèles jouets non-linéaires	75
4.1	Expériences préliminaires sur le L95	77
4.1.1	Inter-comparaison 4D-Var, EnKS et IEnKS	77
4.1.2	Expérience avec estimation de paramètres	78
4.1.3	Expérience avec localisation	79
4.2	Expériences avec erreur modèle sur le L95	81
4.2.1	Base de comparaison	81
4.2.2	Impact de la non-linéarité	82
4.2.3	Impact de l'amplitude de l'erreur modèle	84
4.2.4	Impact de la taille de l'ensemble	85
4.2.5	Impact de la localisation	85
4.2.6	Conclusion des expériences avec erreur modèle	87
4.3	Résultats sur le L95-T	88
4.3.1	Estimation de paramètres	88
4.3.2	Régimes d'émission et de dépôt	89
4.3.3	Stratégies pour des modèles faiblement couplés	90
4.4	Assimilation de données sur le L95-GRS	93
4.4.1	Premières performances	94
4.4.2	Estimation de paramètres	96
4.4.3	Localisation	96
4.5	Conclusions	99
5	Assimilation de données avec un modèle de chimie-transport	101
5.1	Simulation de référence avec Polair3D	102
5.1.1	La configuration du modèle	102
5.1.2	Évaluation en regard des observations	102
5.2	Assimilation de données de référence avec le LETKF	105
5.2.1	Calibrage de $\mathbf{P}^f$ avec EMEP	107
5.2.2	Calibrage de $\mathbf{R}$ pour AirBase	109
5.2.3	Résumé des paramètres et des performances	109

5.3	Estimation des erreurs de représentativité des stations AirBase	111
5.4	Application de l'IEEnKS sur Polair3D	115
5.4.1	Expériences synthétiques	116
5.4.2	Expériences avec observations réelles	119
5.4.3	Assimilation des observations débiaisées	122
5.4.4	Assimilation avec l'IEEnKF-Q	122
5.4.5	Conclusions de l'application de l'IEEnKS sur un CTM	126
6	Méthodes de quantification des incertitudes liées à l'estimation d'une source de polluants	129
6.1	Introduction	131
6.1.1	Modélisation inverse et quantification d'incertitude	131
6.1.2	Modélisation inverse pour le rejet accidentel de radionucléides	131
6.1.3	Quantification d'incertitude et objectifs de cette étude	131
6.2	Les rejets de radionucléides des accidents de Tchernobyl et Fukushima	132
6.2.1	Contexte	132
6.2.2	Estimation du terme source pour les accidents de Tchernobyl et de Fukushima Daiichi	133
6.2.3	Base de données d'observations	134
6.2.4	Modélisation	134
6.3	Modélisation inverse	134
6.3.1	Statistiques d'erreur a priori log-normales	134
6.3.2	Solution bayésienne théorique	136
6.3.3	Les solutions par Bayes hiérarchique et Bayes empirique	137
6.3.4	Estimation des hyperparamètres avec l'algorithme espérance-maximisation	138
6.3.5	A priori sur les hyperparamètres	140
6.3.6	Application aux accidents de Tchernobyl et Fukushima Daiichi	141
6.4	Quantification des incertitudes	144
6.4.1	Échantillonnage préférentiel fondé sur l'approximation de Laplace	144
6.4.2	Échantillonnage par <i>randomize-then-optimize</i> naïf et sans biais	146
6.4.3	Méthode basique de Monte-Carlo par chaînes de Markov	147
6.4.4	Application aux accidents de Tchernobyl et Fukushima Daiichi	147
6.5	Hiérarchie bayésienne complète	148
6.5.1	Méthode hiérarchique de Monte-Carlo par chaînes de Markov	149
6.5.2	Analyse transdimensionnelle	149
6.5.3	Application aux accidents de Tchernobyl et Fukushima Daiichi	150
6.6	Résumé et conclusions	153
	Conclusions	155
	Bibliographie	159

Introduction

Au cours des dernières décennies, nous sommes entrés dans une ère nouvelle de l'information, où des quantités inconcevables de données sont à disposition. Il faut cependant pouvoir les utiliser intelligemment. Le domaine dit du Big Data s'est développé dans ce but. Il a permis dans certains cas de créer des modèles prédictifs en s'abstrayant d'une quelconque étude des équations de la physique. Toutefois, une branche de ce domaine, nommée assimilation de données, s'attèle à tirer profit de l'ensemble des données disponibles d'une part, mais aussi d'un modèle numérique, basé sur des équations de dynamique physique d'autre part. En effet, si pour un système donné, nous avons à notre disposition des observations de ce système ainsi que les sorties d'un modèle prévoyant son état, alors l'assimilation de données consiste à estimer l'état du système en combinant de manière optimale observations et prévisions du modèle. Les domaines d'application des méthodes d'assimilation de données sont très variés : sciences de l'atmosphère, de l'océan et du climat, hydrologie, ingénierie pétrolière, médecine ... Chacun de ces domaines a ses propres spécificités qui requièrent ses propres algorithmes. Les méthodes d'assimilation de données sont ainsi en constante évolution, depuis la création du filtre de Kalman en 1960.

Un des domaines d'application de l'assimilation de données est l'estimation de la pollution atmosphérique. Cela désigne d'une part les gaz à effet de serre, comme le dioxyde de carbone ou le méthane, qui jouent un rôle important dans le réchauffement climatique observé à l'heure actuelle. D'autre part, à l'échelle plus locale, la pollution urbaine est une préoccupation grandissante des grosses agglomérations. En effet, l'Organisation mondiale de la Santé estime qu'au niveau mondial, 1,3 million de personnes meurent chaque année en raison de la pollution de l'air des villes. Parmi les nombreux polluants présents dans l'atmosphère, l'ozone troposphérique fait l'objet de nombreuses études. En tant qu'oxydant, il est néfaste pour la santé et détériore la qualité de vie. Il est donc soumis à des réglementations européennes, qui imposent des niveaux maximaux de concentration. Pour s'assurer du respect de ces réglementations, un réseau dense d'observations a été développé en Europe pour mesurer les concentrations d'ozone troposphérique. Conjointement, des modèles de chimie atmosphérique permettent de prévoir l'évolution des concentrations au cours du temps, d'évaluer a priori les impacts de mesures préventives visant à empêcher un épisode de forte pollution à l'ozone, ou encore d'attribuer une source de polluants à l'origine de cette formation d'ozone. La disponibilité de ces modèles avancés et de ces nombreuses données, ainsi que l'intérêt sanitaire que représente une bonne connaissance des niveaux de pollution, invite à l'application de l'assimilation de données sur des modèles de chimie atmosphérique.

Cependant, ces modèles sont complexes. L'ozone n'est pas directement émis et est le résultat de réactions chimiques faisant intervenir des oxydes d'azotes, issus principalement du trafic routier, et des composés organiques volatils, provenant de sources variées comme les forêts ou l'évaporation de solvants. La relation entre les concentrations de ces espèces précurseurs de l'ozone et les concentrations d'ozone résultant de la présence de ces précurseurs n'est pas linéaire. Ainsi, une augmentation de la concentration de l'un d'entre eux peut paradoxalement conduire à une diminution des concentrations d'ozone. De plus, de très fortes incertitudes sont liées à ces modèles. Même en considérant les équations chimiques implémentées dans le modèle comme parfaites, il reste de nombreuses erreurs. Elles proviennent notamment des sources et des puits de polluants, responsables respectivement de la formation et de la destruction de ces derniers, qui sont très difficiles à estimer. Enfin, les concentrations de plusieurs dizaines ou centaines d'espèces en de très nombreux points d'un domaine tridimensionnel étant simulées par le modèle, le problème auquel l'assimilation de

données est confronté est de très grande dimension.

Les méthodes d'assimilation de données doivent donc être adaptées aux problèmes que présentent la chimie atmosphérique et le transport des polluants. Deux grandes familles de méthodes ont ainsi été utilisées avec succès dans ce contexte. La première famille de méthodes, appelées méthodes variationnelles, repose sur la minimisation numérique de l'écart entre le modèle et les observations, en recourant à la théorie du contrôle optimal. Ces méthodes permettent de tenir compte de la non-linéarité des modèles de chimie atmosphérique. L'autre famille de méthodes, dites méthodes ensemblistes, permet de représenter et propager dans le temps les incertitudes liées à ces modèles malgré leur très grande dimension. Ces deux familles de méthodes présentent donc chacune des avantages et des défauts pour ce contexte d'application. Une nouvelle catégorie de méthodes, dites EnVar 4D, qui sont à la fois ensemblistes et variationnelles, a récemment vu le jour pour tirer profit des deux approches mais n'a à ce jour pas été testée sur des modèles opérationnels de chimie atmosphérique.

En effet, face à la complexité de ces modèles, il est nécessaire de tout d'abord tester les nouvelles méthodes dans un contexte idéalisé, simple et parfaitement maîtrisé. Pour cela, des modèles d'ordre réduit, autrement appelés modèles jouets, ont été créés pour servir de banc d'essai à l'assimilation. Ils doivent être en mesure de synthétiser les phénomènes physiques à l'œuvre dans les modèles opérationnels tout en limitant certaines des difficultés liées à ces derniers. De tels modèles jouets existent pour la météorologie et sont reconnus et utilisés par la communauté des assimilateurs. En revanche, ce n'est pas encore le cas pour la chimie atmosphérique.

Ainsi, nous voyons que les besoins sont multiples pour développer l'assimilation de données sur des modèles de chimie atmosphérique. Tout d'abord, un constant effort d'amélioration des méthodes d'assimilation de données doit être fourni, afin qu'elles soient adaptées au mieux aux différents problèmes sur lesquels elles sont appliquées. Dans notre contexte d'application, ces nouvelles méthodes doivent être en mesure de tenir compte et quantifier les nombreuses incertitudes liées aux modèles de chimie atmosphérique. Pour faciliter ce développement et la validation des nouvelles méthodes, des modèles jouets adaptés doivent par conséquent exister. Après la validation de ces méthodes dans ce contexte idéalisé, les résultats doivent être étendus à des modèles plus complexes et des observations réelles de concentration de polluants.

Au cours de cette thèse, l'ensemble de ces problématiques a été abordé. En effet, le développement théorique de nouvelles méthodes d'assimilation de données, l'application de méthodes d'estimation d'incertitudes liées aux sources de polluants, la création de modèles jouets pour l'assimilation en chimie atmosphérique, ainsi que la validation des méthodes d'assimilation de données récentes sur ces modèles jouets et sur un modèle réaliste ont été au cœur de notre travail. Nous présenterons au premier chapitre les phénomènes physiques et chimiques à l'œuvre dans les modèles de chimie-transport. Dans le deuxième chapitre, nous décrirons les méthodes d'assimilation de données, aussi bien celles éprouvées et utilisées opérationnellement que des méthodes plus récentes, dont une a été développée au cours de la thèse. Le troisième chapitre décrira les modèles jouets sur lesquels les méthodes d'assimilation sont habituellement testées en météorologie, ainsi que ceux récemment introduits pour la chimie atmosphérique. Nous appliquerons au quatrième chapitre les méthodes d'assimilation de données sur ces modèles jouets. Les chapitres suivants s'attacheront à étendre ces résultats, obtenus dans des contextes simplifiés, à des cas réels, utilisant des modèles plus complexes et des observations obtenues in situ. Le cinquième chapitre présentera ainsi les expériences d'assimilation de concentrations d'ozone sur l'Europe, et le chapitre final présentera des méthodes et résultats pour le calcul d'incertitudes liées aux accidents nucléaires de Tchernobyl et Fukushima-Daiichi.

Modèles de chimie atmosphérique

On trouve dans l'atmosphère des polluants qui, selon leurs concentrations, peuvent impacter la qualité de vie des personnes qui y sont confrontées. En particulier, l'ozone troposphérique figure parmi les polluants réglementés en Europe et dans le monde. Bien que non directement émis, il se retrouve en grande quantité dans les zones urbanisées. En tant qu'oxydant, il a un impact néfaste sur la santé. Il est donc nécessaire de comprendre les phénomènes physiques et chimiques qui régissent l'évolution de la concentration des polluants dans l'atmosphère, notamment celle de l'ozone et de ses précurseurs. Leur étude permet de déterminer les équations qui composent les modèles de chimie atmosphérique. Nous ne regarderons que la chimie gazeuse et non la chimie des aérosols, plus complexe et coûteuse en temps de calcul.

L'objectif de ce chapitre est d'introduire les principaux phénomènes en jeu dans l'évolution des concentrations de polluants et leur intégration numérique dans un modèle de chimie-transport (CTM), où la météorologie n'est pas directement calculée mais est une entrée du modèle, ou dans un modèle couplé de chimie-météorologie (CCMM). Nous insisterons également tout au long de ce chapitre sur les incertitudes liées à ces modèles.

Sommaire

1.1	Météorologie et transport	5
1.1.1	Équations de Navier-Stokes	5
1.1.2	Définition et importance de la couche limite atmosphérique	6
1.1.3	Transport : advection et diffusion	6
1.1.4	Dépôt sec et humide	6
1.2	Chimie de l'ozone troposphérique	6
1.2.1	Cycle de Leighton	6
1.2.2	Régimes NO _x et COV limité	7
1.3	Incertitudes sur les données d'entrée des modèles de chimie atmosphérique	8
1.3.1	Émissions	9
1.3.2	Condition initiale	9
1.3.3	Conditions aux limites	9
1.3.4	Météorologie dans les CTMs	10
1.4	Intégration numérique	10
1.4.1	Modélisation numérique du transport	10
1.4.1.1	Schémas d'advection	10
1.4.1.2	Condition CFL	11
1.4.1.3	Diffusion numérique	12
1.4.1.4	Limiteurs de flux	12
1.4.2	Modélisation numérique de la chimie	12
1.4.2.1	Représentation par utilisation de molécules suppléantes	12
1.4.2.2	Représentation par décomposition en groupes fonctionnels	13
1.4.2.3	Approximation des états quasi-stationnaires	13
1.4.2.4	Rosenbrock d'ordre 2	13

1.4.2.5 Pas de temps adaptatif	14
1.4.3 Splitting	14
1.5 Résumé	14

Les notes des cours et livres suivants ont servi de référence pour la rédaction de ce chapitre :

- V. MALLET : Introduction aux modèles de chimie-transport pour la qualité de l'air. Notes de cours, École des Ponts ParisTech, Janvier 2008. URL <http://vivienmallet.net/static/publication/mallet08introduction.pdf>. Cours du Master SGE-AQA
- B. SPORTISSE et V. MALLET : Calcul scientifique pour l'environnement. Notes de cours, ENSTA ParisTech, 2005. URL <http://vivienmallet.net/static/publication/sportisse07calcul.pdf>. Cours de deuxième année
- C. SEIGNEUR : Environnement atmosphérique et qualité de l'air. Notes de cours, École des Ponts ParisTech, 2012. URL http://cerea.enpc.fr/fr/cours-web/POLU1_2012-2013.html. Cours de deuxième année
- B. SPORTISSE : *Pollution atmosphérique: des processus à la modélisation*. Ingénierie et développement durable. Springer-Verlag Paris, 2008.

1.1 Météorologie et transport

La météorologie est définie par l'évolution des principales variables d'état de l'atmosphère, à savoir la masse volumique, la température et la pression, ainsi que les vitesses des vents dans les trois directions. Les modèles complets de météorologie définissent également l'humidité spécifique, la hauteur des nuages ou les niveaux de pluie, entre autres, mais nous ne nous y intéresserons pas ici. Nous décrivons les équations d'évolution de ces variables et leurs interactions avec les espèces chimiques via le transport atmosphérique.

1.1.1 Équations de Navier-Stokes

Notons $\mathbf{V} = (u, v, w)$ le vecteur du vent, ρ la masse volumique de l'air, T la température et p la pression. L'équation de conservation de la masse nous donne

$$\frac{\partial \rho}{\partial t} + \text{div}(\rho \mathbf{V}) = 0. \quad (1.1)$$

L'air étant soumis à la force liée au gradient de pression, à la gravité, à la force de Coriolis et à la résistance de l'air, la deuxième loi de Newton donne

$$\rho \left(\frac{\partial \mathbf{V}}{\partial t} + \mathbf{V} \cdot \nabla \mathbf{V} \right) = -\nabla P - \rho \mathbf{g} + \rho \mathbf{F}_c + \mu \Delta \mathbf{V} \quad (1.2)$$

avec μ la constante de diffusion dynamique, $\mathbf{g}$ la gravité, et $\rho \mathbf{F}_c$ la force de Coriolis.

En ce qui concerne la température, une équation similaire peut être également déterminée

$$\rho c_p \left(\frac{\partial T}{\partial t} + \mathbf{V} \cdot \nabla T \right) = Q + \mu_\theta \Delta T \quad (1.3)$$

avec μ_θ la conductivité thermique, Q les sources de chaleur et c_p la capacité calorifique à pression constante.

Ces trois équations (1.1), (1.2), (1.3) (ainsi qu'une dernière sur l'humidité spécifique) forment le système d'équations de Navier-Stokes. Elles définissent l'évolution des principales variables d'état de l'atmosphère.

Cependant, lorsque l'on reste proche du sol, on peut appliquer l'approximation de Boussinesq qui consiste à supposer la densité constante (sauf dans le terme de gravité). On linéarise alors les équations de Navier-Stokes autour d'un état de référence (variables notées avec un indice r) que l'on considère

- au repos : $\mathbf{V}_r = 0$;
- hydrostatique : $\frac{\partial P_r}{\partial z} = -\rho_r g$;
- adiabatique $\frac{\partial T_r}{\partial z} = -g/c_p$;
- l'air étant un gaz parfait : $P_r = \rho_r r T_r$;

et on note par un indice 0 les valeurs moyennes de ces variables sur une hauteur donnée.

La linéarisation autour de cet état de référence donne alors, écrite dans un repère terrestre

$$\begin{cases} \frac{\partial u}{\partial x} + \frac{\partial v}{\partial y} + \frac{\partial w}{\partial z} = 0 \\ \frac{\partial u}{\partial t} + u \frac{\partial u}{\partial x} + v \frac{\partial u}{\partial y} + w \frac{\partial u}{\partial z} = \frac{\mu}{\rho_0} \Delta u - \frac{1}{\rho_0} \frac{\partial P}{\partial x} + (F_c)_x \\ \frac{\partial v}{\partial t} + u \frac{\partial v}{\partial x} + v \frac{\partial v}{\partial y} + w \frac{\partial v}{\partial z} = \frac{\mu}{\rho_0} \Delta v - \frac{1}{\rho_0} \frac{\partial P}{\partial y} + (F_c)_y \\ \frac{\partial w}{\partial t} + u \frac{\partial w}{\partial x} + v \frac{\partial w}{\partial y} + w \frac{\partial w}{\partial z} = \frac{\mu}{\rho_0} \Delta w - \frac{1}{\rho_0} \frac{\partial P}{\partial z} + \frac{T}{T_0} g + (F_c)_z \\ \rho_0 c_p \left(\frac{\partial T}{\partial t} + u \frac{\partial T}{\partial x} + v \frac{\partial T}{\partial y} + w \frac{\partial T}{\partial z} \right) = Q + \mu_\theta \Delta T \end{cases}$$

Ce sont ce type d'équations qui sont intégrées dans les modèles de météorologie (ou dans les CCMs) pour obtenir notamment les vents qui transportent les polluants dans le domaine. Ces modèles peuvent alors servir à la prévision du temps.

1.1.2 Définition et importance de la couche limite atmosphérique

Proche du sol, là où s'applique l'approximation de Boussinesq, les vents verticaux sont considérés comme quasiment nuls. Le transport vertical des masses d'air s'effectue par le flux convectif turbulent issu du réchauffement de l'atmosphère par la surface du sol en journée. On définit alors la couche limite atmosphérique (CLA) comme la hauteur d'atmosphère influencée par ces effets de surface. La hauteur de cette CLA passe de 1–2 km le jour à quelques centaines de mètres seulement la nuit, pendant laquelle la turbulence proche du sol est dû uniquement au cisaillement du vent. Cette couche est très importante en chimie atmosphérique car c'est là qu'ont lieu les mélanges de polluants. Malheureusement, sa hauteur est difficile à diagnostiquer.

1.1.3 Transport : advection et diffusion

La concentration c d'un polluant évolue avec le vent selon l'équation

$$\frac{\partial c}{\partial t} = -\operatorname{div}(\mathbf{V}c) + \operatorname{div}(\rho \mathbf{K} \nabla \frac{c}{\rho}) \quad (1.4)$$

où $\mathbf{K}$ est une matrice de diffusion turbulente. Le premier terme, dit d'advection et qui correspond au transport par le vent des polluants dans l'atmosphère, s'effectue essentiellement dans la direction horizontale. Le second terme consiste la diffusion. La diffusion horizontale n'est pas très importante, surtout étant donnée la diffusion numérique induite par l'intégration de l'advection dans les modèles (voir 1.4.1.3). En revanche, le terme diagonal de la matrice $\mathbf{K}$ correspondant au transport vertical, noté K_z , n'est pas négligeable. C'est lui qui impose le transport turbulent dans la CLA. Il fait l'objet de plusieurs modélisations différentes, notamment par [Louis \(1979\)](#) et [Troen et Mahrt \(1986\)](#).

1.1.4 Dépôt sec et humide

Les gaz peuvent sortir de l'atmosphère en se déposant au sol par diffusion ou sous l'effet de la pluie. On parle alors respectivement de dépôt sec et dépôt humide (ou lessivage). Ces phénomènes dépendent donc des données météorologiques comme la turbulence verticale ou la pluviométrie. Ils font l'objet de plusieurs modélisations différentes. En ce qui concerne le dépôt sec, on distingue notamment les modélisations de [Zhang *et al.* \(2003\)](#) ou [Wesely \(1989\)](#).

1.2 Chimie de l'ozone troposphérique

L'atmosphère est composée en grande majorité de diazote (N_2 , environ 80%) et de dioxygène (O_2 , environ 20%). Cependant, de très nombreuses autres espèces chimiques se retrouvent dans l'atmosphère avec des concentrations allant du picogramme à quelques milligrammes par mètre cube. Parmi elles, l'ozone est celle qui fait l'objet de notre étude et dont nous allons décrire brièvement le cycle de vie.

1.2.1 Cycle de Leighton

Sous l'action du rayonnement solaire, le dioxyde d'azote (NO_2) se photolyse pour former du monoxyde d'azote et de l'oxygène :



Cet oxygène va ensuite s’associer au dioxygène de l’atmosphère pour former de l’ozone :



où M est une molécule d’air. La forte réactivité de (R 1.2) fait que les deux réactions précédentes sont parfois fusionnées en



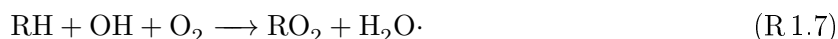
La réaction inverse existe également et consomme l’ozone en retour :



Ces deux équations (R 1.3) et (R 1.4) constituent le cycle de Leighton et conduisent à un équilibre entre O_3 et oxydes d’azote (NO_x). Cependant, l’ozone n’est pas le seul oxydant capable de transformer le NO en NO_2 . Les radicaux peroxyels tels que l’hydroperoxye (HO_2) ou les peroxyels organiques (notés RO_2) peuvent produire une réaction similaire à (R 1.4)



Ces radicaux sont issus de l’oxydation ou de la photolyse de composés organiques volatiles (COV). On note parmi les COV les alcanes, alcènes, aldéhydes, ou composés aromatiques entre autres. Un exemple d’oxydation d’un alcane RH est donné par



Pareillement, un exemple de photolyse d’un aldéhyde comme le formaldéhyde HCHO est donné par



Les COV sont principalement oxydés par les radicaux hydroxyle OH, comme dans l’équation (R 1.7). Ces radicaux peuvent être issus par exemple de la photolyse du peroxyde d’hydrogène (H_2O_2) ou de l’acide nitreux (HNO_2) :



En résumé, l’ozone est majoritairement produit par photolyse du NO_2 . Ce dernier est produit par oxydation du NO par des radicaux, eux même produits par oxydation et photolyse des COV. Les concentrations d’ozone sont donc intriquées avec celles des NO_x et COV. La multitude d’équations rentrant en considération dans une définition complète du système chimique de l’ozone et la grande amplitude de vitesses de réaction liées à chacune de ces équations rend le système raide et donc difficile à intégrer numériquement.

1.2.2 Régimes NO_x et COV limité

Si on schématise le bilan d’ozone du paragraphe précédent, le NO issu de la réaction (R 1.1) peut être reconverti en NO_2 soit par l’ozone, via l’équation (R 1.4), conduisant à un bilan d’ozone nul, soit par les équations (R 1.5) et (R 1.6), conduisant à un bilan d’ozone positif. On rappelle que ces dernières équations sont le résultat de l’oxydation des COV.

Ces deux processus sont en compétition et on peut alors distinguer deux régimes :

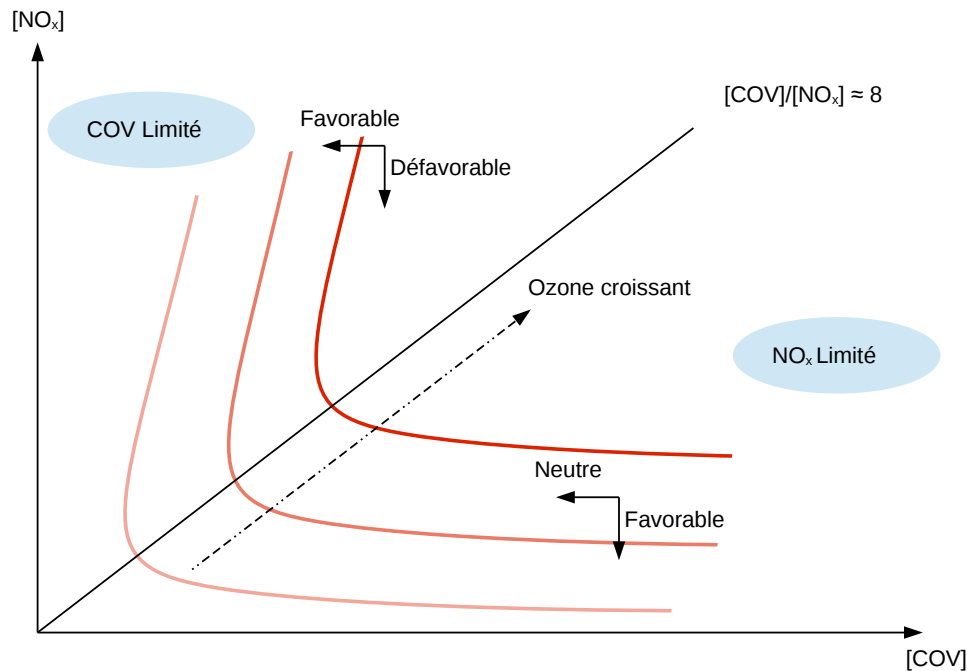


FIGURE 1.1 – Isopleths d’ozone en fonction des concentrations de NO_x et de COV. La valeur $[\text{COV}]/[\text{NO}_x] \simeq 8$ est typique du continent nord-américain.

- un régime riche en NO_x et pauvre en COV, favorisant le premier processus ;
- un régime riche en COV et pauvre en NO_x , favorisant le deuxième processus.

Dans le premier régime, dit “COV limité”, le système est très sensible aux concentrations de COV, pour lesquels une augmentation de leurs concentrations implique une forte augmentation de l’ozone. Une augmentation des concentrations de NO_x dans ce régime déplace l’équilibre du cycle de Leighton conduisant à une faible diminution des concentrations d’ozone. Réciproquement, dans le deuxième régime, dit “ NO_x limité”, l’ozone est très sensible à l’augmentation des concentrations de NO_x , mais ne l’est pas à celles des COV. La distinction des deux régimes peut se schématiser par le tracé des isopleths d’ozone (isovaleurs) en fonction des concentrations de NO_x et de COV (Fig. 1.1). Cette particularité de la dépendance de l’ozone aux concentrations de NO_x et de COV implique une forte non-linéarité vis-à-vis de ses précurseurs.

1.3 Incertitudes sur les données d’entrée des modèles de chimie atmosphérique

Outre les approximations et les différents choix de modélisation présentés jusqu’ici, les CTMs et les CCMMs nécessitent de nombreuses données d’entrée, présentant de fortes incertitudes. Le choix de ces données impacte grandement la qualité des simulations, d’autant plus que le modèle n’a aucun contrôle sur elles. Nous listons donc ici certaines des principales sources d’incertitude liées aux données nécessaires aux modèles de chimie atmosphérique.

1.3.1 Émissions

On distingue les polluants primaires, c'est-à-dire ceux qui sont directement émis dans l'atmosphère, comme par exemple certains COV et le monoxyde d'azote, des polluants secondaires, c'est-à-dire ceux qui sont issus de réactions chimiques dans l'atmosphère, comme l'ozone entre autres. L'origine des polluants primaires est principalement liée à l'activité humaine (émissions anthropiques), comme le trafic, l'activité industrielle ou le chauffage urbain. Mais on trouve également des émissions liées à la biomasse (émissions biogéniques) ainsi que les volcans et feux de forêts qui produisent majoritairement des COV, ou encore la foudre qui produit des NO_x .

Étant donnée l'importance des NO_x et des COV dans le cycle de l'ozone, il est nécessaire d'estimer les quantités de polluants émis par les différentes sources. En ce qui concerne la pollution anthropique, des inventaires existent. A l'échelle européenne, le European Monitoring and Evaluation Programme (EMEP) (Vestreng *et al.*, 2004) récupère des données d'émissions qu'il couple à des résultats de simulations et des observations in situ, permettant de produire un inventaire par classes Standard Nomenclature for Air Pollution (SNAP). Un autre inventaire d'émissions, GEMS-TNO, est décrit dans Visschedijk *et al.* (2007).

En ce qui concerne la pollution biogénique, différents modèles ont par ailleurs été développés. Le Model of Emissions of Gases and Aerosols from Nature (MEGAN) (Guenther *et al.*, 2006) ou encore le modèle de Simpson *et al.* (1999) fournissent ces émissions en fonction de l'occupation des sols (*land use coverage* (LUC)).

Enfin, il est possible d'utiliser des méthodes de modélisation inverse pour réestimer certaines de ces émissions. Koohkan *et al.* (2013) ont par exemple diagnostiqué les flux de COV à l'échelle européenne en utilisant des méthodes de modélisation inverse.

Ces émissions, aussi bien anthropiques que biogéniques, sont difficiles à estimer et sont donc une forte source d'incertitude dans les CTMs et les CCMMs.

1.3.2 Condition initiale

Avant de lancer une simulation, toutes les variables du système doivent être initialisées. Étant donné que les équations de la chimie sont stables, la trajectoire du modèle converge vers une trajectoire unique quelque soit la condition initiale. Elle n'est donc pas très importante et peut être erronée vu qu'elle aura été oubliée au bout de quelques semaines de modélisation. Il est par contre nécessaire d'appliquer une période de spin-up pour se retrouver sur cette trajectoire.

Cependant, cela ne se passe pas ainsi pour la météorologie, qui est chaotique. Ainsi, les erreurs sur la condition initiale se retrouvent amplifiées avec le temps et les trajectoires issues de deux conditions initiales différentes divergent.

La condition initiale n'est donc pas une forte source d'incertitude dans un CTM, mais elle est cruciale dans le cas d'un CCMM.

1.3.3 Conditions aux limites

À part dans le cas des modèles globaux, il est nécessaire de définir les conditions au bord du domaine. Ceci est souvent effectué en interpolant aux bords les résultats d'une simulation globale. Il s'agit alors de conditions aux limites de type Dirichlet. Dans notre cas, les conditions aux limites sont issues du modèle Model for OZone And Related chemical Tracers (MOZART) dans la version 2 ou 4 de son implémentation (Horowitz *et al.*, 2003; Emmons *et al.*, 2010).

1.3.4 Météorologie dans les CTMs

Dans un modèle de chimie atmosphérique, la météorologie peut soit être modélisée *online*, c'est-à-dire en même temps que la chimie, soit elle peut être modélisée séparément (*offline*) et les résultats de simulation sont alors une donnée d'entrée du modèle de chimie. Dans le premier cas, on parle d'un modèle couplé de chimie-météorologie (CCMM), où la météorologie et la chimie sont fortement couplées, tandis que le second s'appelle un modèle de chimie-transport (CTM), où la météorologie et la chimie ne sont que faiblement couplées.

Dans le cas d'un CTM, la météorologie est donc une donnée d'entrée sur laquelle le modèle n'a aucun contrôle. Les variables météorologiques correspondent entre autres aux vents, température, hauteur de la CLA, intensité du rayonnement solaire, ou encore les précipitations. Ces variables sont fournies sur un domaine qui ne coïncide souvent pas avec celui du CTM, ce qui nécessite une interpolation. Par ailleurs, les concentrations simulées des polluants ne peuvent pas avoir de rétroaction sur les variables météorologiques (particulièrement la température), contrairement à un CCMM. Ajouté aux erreurs sur les champs météorologiques fournis, l'ensemble de ces limites dues au couplage faibles des deux modèles fait de la météorologie une forte source d'incertitude des CTMs.

Les données météorologiques en Europe peuvent être issues des réanalyses du Centre européen pour les prévisions météorologiques à moyen terme (CEPMMT), du modèle MM5 (Dudhia *et al.*, 2005), ou encore du Weather Research and Forecasting model (WRF) (Skamarock *et al.*, 2008).

1.4 Intégration numérique

L'intégration numérique des équations des modèles de chimie atmosphérique nécessite d'effectuer quelques approximations, entre autres une discrétisation en temps et en espace des équations. Cette section présente quelques éléments couramment utilisés qui nous seront utiles par la suite.

1.4.1 Modélisation numérique du transport

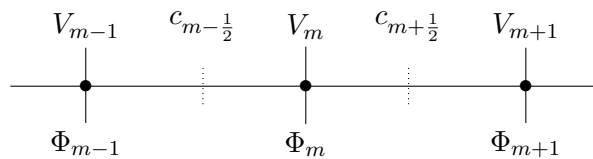
1.4.1.1 Schémas d'advection

Il s'agit ici de discrétiser en temps et en espace l'équation d'advection

$$\frac{\partial c}{\partial t} = -\text{div}(\mathbf{V}c), \quad (1.5)$$

avec c la concentration et $\mathbf{V}$ le vent. Nous présentons deux schémas d'advection couramment utilisés dans les modèles de chimie atmosphérique, comme Polair3D de la plateforme Polyphemus (Mallet *et al.*, 2007) ou CHIMERE (Menut *et al.*, 2013).

Schéma *upwind* Dans le schéma *upwind*, le vent est défini à l'interface entre les mailles des concentrations, selon une grille-C d'Arakawa, que nous illustrons ici dans un cas unidimensionnel.



où c_m désigne la concentration au point m , V_m le vent, et ϕ_m les flux, qui sont de la sorte :

$$\Phi_m = V_m c_{m-\frac{1}{2}} \quad \text{si } V_m \geq 0, \quad (1.6)$$

$$\Phi_m = V_m c_{m+\frac{1}{2}} \quad \text{si } V_m < 0. \quad (1.7)$$

L'équation d'advection *upwind* pour le transport des concentrations est :

$$\frac{dc_{m+\frac{1}{2}}}{dt} = \frac{\Phi_m - \Phi_{m+1}}{\Delta x}, \quad (1.8)$$

avec Δx la distance entre deux points de grille. L'avantage de ce schéma est d'une part sa simplicité, et d'autre part sa nature conservatrice et positive. En effet, il n'y a pas de perte de masse et les concentrations après advection restent positives (sous une certaine condition, voir 1.4.1.2). En revanche, il est fortement diffusif.

Schéma DST3 Le schéma précédent correspond à une approche où le problème est d'abord discrétisé en espace, et est ensuite résolu en temps selon un schéma choisi (Euler explicite ou Runge-Kutta 4 par exemple). Cependant, on peut faire le choix de discrétiser conjointement en temps et en espace. On parle alors de schéma *direct space time* (DST).

Si on note $a = V \frac{\Delta t}{\Delta x}$ le nombre de Courant-Friderichs-Lewy (CFL), avec Δt et Δx les pas de temps et d'espace, et c_i^n la concentration au point i et au temps t_n , on peut alors définir une méthode d'ordre 3 dite DST3 avec

$$c_i^{n+1} = \gamma_{-2} c_{i-2}^n + \gamma_{-1} c_{i-1}^n + \gamma_0 c_i^n + \gamma_1 c_{i+1}^n, \quad (1.9)$$

où

$$\begin{aligned} \gamma_{-2} &= -a(1 - a^2)/6, \\ \gamma_{-1} &= a(2 - a)(1 + a)/2, \\ \gamma_0 &= (2 - a)(1 - a^2)/2, \\ \gamma_1 &= -a(2 - a)(1 - a)/6. \end{aligned}$$

Ce schéma a donc un ordre élevé, mais ne conserve pas la positivité. Ainsi, il ne peut pas être utilisé tel quel comme schéma d'advection pour un CTM. Il nécessite d'être associé à un limiteur de flux (1.4.1.4).

1.4.1.2 Condition CFL

Le coefficient $a = V \frac{\Delta t}{\Delta x}$ défini précédemment intervient dans la condition dite CFL. Cette condition impose que pour un vent V donné, le domaine spatio-temporel du modèle vérifie

$$V \Delta t < \Delta x. \quad (1.10)$$

Intuitivement, si on a un vent de vitesse $V > 0$ advectant une masse d'air pendant Δt , il déplacera la masse de $V \Delta t$. Il s'agit alors de s'assurer que l'intégralité de la masse ne sorte pas de la maille, c'est-à-dire étant donné l'équation (1.8), que $\frac{\Phi}{\Delta x} = \frac{Vc}{\Delta x} < \frac{c}{\Delta t}$, ce qui correspond à (1.10). C'est cette condition qui assure par exemple la positivité du schéma d'intégration *upwind*, ainsi que sa stabilité et celle du DST3. Une démonstration plus rigoureuse de l'origine de cette condition pourra être trouvée dans Sportisse et Mallet (2005).

La vérification de cette condition est non seulement primordiale, mais aussi contraignante quand on souhaite que la résolution spatiale du domaine devienne fine, car elle impose également un pas de temps d'intégration plus faible.

1.4.1.3 Diffusion numérique

Comme énoncé dans la section 1.1.3, on néglige souvent la diffusion horizontale car elle est amplement supplantée par ce qu'on appelle la diffusion numérique. Cette diffusion apparaît au cours de l'intégration numérique de l'advection, qui en interpolant sur les points de grille, diffuse les concentrations sur plusieurs mailles. Ceci s'explique notamment par le fait que pour un schéma d'advection et une discrétisation donnés, la solution de (1.5) qui est effectivement obtenue se rapproche plus de la solution d'une équation aux dérivées partielles (EDP) dite équivalente à (1.5), cette EDP différant de l'équation originelle.

1.4.1.4 Limiteurs de flux

On a vu à la section 1.4.1.1 deux schémas d'intégration numérique de l'advection, à savoir l'*upwind* et le DST3. Le premier garantit la positivité de la solution, mais au prix d'une forte diffusion numérique, tandis que le deuxième, d'ordre plus élevé, est moins diffusif mais n'est plus positif. Il s'avère que la diffusion numérique est un mal nécessaire qui permet de garantir la positivité systématique du schéma *upwind*, qui est le schéma de diffusion minimale ayant cette propriété. Pour pouvoir limiter la diffusion numérique tout en garantissant la positivité des schémas dans la majorité des cas, on fait appel à des limiteurs de flux. Le principe est de choisir un schéma d'ordre faible (ϕ^1) quand on est proche du gradient, et un schéma d'ordre supérieur (ϕ^2) moins diffusif quand la solution est lisse.

On écrit le flux total choisi de la forme

$$\phi_{i+1/2} = \phi_{i+1/2}^1 - \Psi(\theta_i)(\phi_{i+1/2}^1 - \phi_{i+1/2}^2), \quad \text{avec} \quad \theta_i = \frac{c_i - c_{i-1}}{c_{i+1} - c_i}. \quad (1.11)$$

où Ψ est le limiteur de flux, qui peut être choisi entre autres selon la formulation de Koren-Sweby, ou selon la formulation de Van Leer.

$$\text{Koren-Sweby : } \Psi(\theta) = \max(0, \min(\beta\theta), \min(\theta, \beta)) \quad \text{avec} \quad 1 \leq \beta \leq 2 \quad (1.12)$$

$$\text{Van Leer : } \Psi(\theta) = \frac{\theta + |\theta|}{1 + |\theta|} \quad (1.13)$$

Le modèle Polair3D de la plateforme Polyphemus utilise par exemple un schéma d'advection DST3 avec un limiteur de flux de Koren-Sweby.

1.4.2 Modélisation numérique de la chimie

Le très grand nombre de molécules et de réactions chimiques différentes entrant en jeu dans la pollution atmosphérique rend difficile une implémentation exhaustive de l'intégralité de cette chimie, surtout en ce qui concerne les COV. Par conséquent, des méthodes adaptées ont été développées pour réduire la taille du mécanisme chimique modélisé.

1.4.2.1 Représentation par utilisation de molécules suppléantes

Cette approche consiste à faire l'hypothèse que les molécules organiques similaires vont se comporter chimiquement de manière similaire. Il s'agit alors de définir une molécule suppléante qui jouera le rôle d'une combinaison de plusieurs espèces chimiques, avec une cinétique correspondant à une moyenne pondérée des cinétiques des espèces suppléées. La pondération se fonde sur les émissions relatives des différentes espèces et cette cinétique dépend donc des données d'émissions.

L'avantage de cette approche est qu'elle permet de bien représenter l'évolution des molécules chimiques. Son désavantage est qu'elle ne permet pas la conservation de la masse de carbone puisque des molécules de nombre de carbone différent sont regroupées en une seule molécule suppléante. Cette approche est notamment développée dans les deux versions du Regional Atmospheric Chemistry Mechanism (RACM) (Goliff *et al.*, 2013).

1.4.2.2 Représentation par décomposition en groupes fonctionnels

Cette deuxième approche consiste à décomposer une molécule en groupes fonctionnels et à supposer qu'elle se comporte comme la somme des comportements de ces groupes fonctionnels. L'avantage de cette méthode est qu'elle conserve la masse de carbone, contrairement à la représentation précédente. Cependant, on perd la structure des molécules avec cette approche et on ne peut pas suivre les COV sur plusieurs étapes d'oxydation. Elle est notamment développée dans le système Carbon Bond 2005 (CB05) (Yarwood *et al.*, 2005).

1.4.2.3 Approximation des états quasi-stationnaires

Certaines espèces chimiques sont extrêmement réactives, de telle sorte que dès qu'elles sont produites, elles sont quasiment instantanément détruites par d'autres réactions chimiques. Leur durée de vie est alors si faible qu'il est possible de faire une approximation, dite approximation des états quasi-stationnaires (QSSA), qui consiste à considérer l'état stationnaire pour ces espèces. Il suffit alors d'écrire l'équation d'évolution de la concentration de l'espèce X_i

$$\frac{\partial[X_i]}{\partial t} = P_i - L_i[X_i], \quad (1.14)$$

où P_i et L_i sont les termes de production et de consommation par les réaction chimiques, et de supposer $\frac{\partial[X_i]}{\partial t} = 0$, c'est-à-dire

$$[X_i]_{\text{QSSA}} = \frac{P_i}{L_i}. \quad (1.15)$$

1.4.2.4 Rosenbrock d'ordre 2

L'approximation précédente aide à diminuer la raideur du système d'équations du modèle de chimie. Cependant, cela ne suffit pas à pouvoir envisager une méthode d'intégration explicite, qui ne serait pas stable. Il faut alors recourir à une méthode semi-implicite, comme la méthode de Rosenbrock d'ordre 2, communément utilisée en chimie atmosphérique (Hundsdoerfer et Verwer, 2003), notamment par le modèle Polair3D de Polyphemus (Mallet *et al.*, 2007).

Si on a l'équation différentielle autonome pour la concentration $\mathbf{c}$

$$\frac{d\mathbf{c}}{dt} = f(\mathbf{c}), \quad (1.16)$$

que l'on discrétise en temps, notons $\mathbf{c}_n$ la concentration au temps t_n . Alors la méthode d'intégration de Rosenbrock autonome d'ordre 2 est

$$\mathbf{c}_{n+1} = \mathbf{c}_n + \frac{3}{2}\Delta t \mathbf{k}_1 + \frac{1}{2}\Delta t \mathbf{k}_2, \quad (1.17)$$

$$(\mathbf{I} - \gamma \Delta t \mathbf{A}) \mathbf{k}_1 = f(\mathbf{c}_n), \quad (1.18)$$

$$(\mathbf{I} - \gamma \Delta t \mathbf{A}) \mathbf{k}_2 = f(\mathbf{c}_n + \Delta t \mathbf{k}_1) - 2\mathbf{k}_1, \quad (1.19)$$

où $\gamma = 1 \pm \frac{1}{\sqrt{2}}$, $\Delta t = t_{n+1} - t_n$ et $\mathbf{A}$ est la matrice jacobienne $\mathbf{A} = f'(\mathbf{c}_n)$. On remarque que ce schéma est très coûteux car nécessitant deux résolutions de systèmes linéaires pour obtenir $\mathbf{k}_1$ et $\mathbf{k}_2$.

1.4.2.5 Pas de temps adaptatif

Une dernière méthode pour se prévenir contre la raideur des systèmes d'équations chimiques est de faire appel à un pas de temps adaptatif. L'idée est de fixer comme critère une variation relative ou absolue maximale des concentrations des espèces lors de l'intégration d'un pas de temps. Pour satisfaire ce critère, on doit alors parfois diminuer la taille du pas de temps en fonction de la variation des concentrations.

Concrètement, on vérifie si la différence entre l'approximation du premier ordre et celle du second ordre est plus petite qu'une tolérance ε fixée. Notons $\bar{\mathbf{c}}_{n+1}$, la solution d'un schéma d'ordre 1, $\mathbf{c}_{n+1}$ la solution d'un schéma d'ordre 2 et $E_{n+1} = \|\mathbf{c}_{n+1} - \bar{\mathbf{c}}_{n+1}\|_\infty$, avec $\|\mathbf{x}\|_\infty = \max(|x_1|, \dots, |x_n|)$ la norme infinie. Soit $E_{n+1} \leq \varepsilon$, auquel cas le pas de temps utilisé était suffisamment petit et on continue l'intégration, soit $E_{n+1} > \varepsilon$, auquel cas on recommence l'intégration avec un Δt plus petit. Dans les deux cas, pour l'intégration suivante, on sélectionne $\Delta t_{\text{new}} = \nu \Delta t \sqrt{\frac{\varepsilon}{E_n}}$ où $\nu < 1$ sert à éviter une succession de rejets de pas de temps. On remarque également que $\Delta t_{\text{new}} < \Delta t$ si $E_n \leq \varepsilon$. (Savcenko *et al.*, 2007).

Cette méthode permet d'adapter le pas de temps d'intégration dans les cas où les concentrations évoluent très rapidement, comme par exemple le matin quand la photolyse se met en place.

1.4.3 Splitting

L'équation d'évolution de la concentration d'une espèce chimique c_i , si on prend en considération l'intégralité des processus d'advection, de diffusion, de réaction chimique, de dépôts et d'émissions, est de la forme

$$\frac{\partial c_i}{\partial t} = \underbrace{-\text{div}(\mathbf{V}c_i)}_{\text{advection}} + \underbrace{\text{div}\left(\rho \mathbf{K} \nabla \frac{c_i}{\rho}\right)}_{\text{diffusion}} + \underbrace{\chi_i(c)}_{\text{chimie}} + \underbrace{E_i}_{\text{émissions}} - \underbrace{\lambda_i}_{\text{dépôts}}. \quad (1.20)$$

On la nomme équation du transport réactif. Chacun de ces processus ne nécessite pas forcément les mêmes méthodes d'intégration. Là où une méthode explicite peut suffire pour intégrer l'advection, on a vu qu'une méthode semi-implicite d'ordre 2 était nécessaire pour intégrer la chimie. Par ailleurs, il peut être numériquement coûteux d'intégrer l'ensemble des processus en une fois quand la méthode requiert des inversions de matrices, qui seront alors beaucoup plus grosses.

Ainsi, on applique ce que l'on appelle un *splitting* de l'équation. Au premier ordre, cela revient à d'abord intégrer la partie advection ainsi que les phénomènes de dépôt et d'émission, puis la partie diffusion avant de finir avec la chimie. Il est également possible d'appliquer un *splitting* d'ordre supérieur, dit *splitting* de Strang (Strang, 1968), où l'advection et la diffusion sont appliquées en deux fois, une partie avant la chimie et une autre après.

1.5 Résumé

Nous avons introduit au cours de ce chapitre les phénomènes physico-chimiques modélisés dans les CTMs et les CCMMs. Pour nombre d'entre eux, des approximations et des choix de modélisation doivent être faits. Ceci représente une première source d'incertitude. À cela s'ajoute le choix des méthodes d'intégration numérique des équations définies par la modélisation. De part la discrétisation spatio-temporelle de ces équations, de nouvelles imprécisions s'accumulent. Enfin, de nombreuses données d'entrée doivent être fournies aux CTMs et CCMMs, qui sont elles aussi entachées d'erreurs et amplifient encore l'incertitude autour des résultats de simulation.

Il peut alors être jugé nécessaire de faire appel à des mesures de concentrations pour recentrer ces simulations et se rapprocher de la vérité de l'état de l'atmosphère. Cette démarche d'utiliser des mesures pour améliorer des résultats de simulations est appelée l'assimilation de données et fait l'objet de notre prochain chapitre.

Assimilation de données

Les modèles introduits au chapitre 1 donnent des résultats de simulation qui diffèrent de la vérité, en raison de toutes les incertitudes et approximations mentionnées. Il y a par ailleurs à disposition de multiples mesures des concentrations de polluants dans l'atmosphère. Cependant, ces mesures sont également partiellement erronées étant donnée l'erreur d'instrumentation. De plus, elles sont fournies avec une fréquence temporelle et une résolution spatiale inférieures à celles de nos modèles. On recombine alors ces deux informations grâce à l'assimilation de données, qui est définie dans [Le Dimet et Blum \(2002\)](#) (dans le cadre de l'assimilation de données pour l'étude d'un fluide géophysique) de la manière suivante : « L'assimilation de données est l'ensemble des techniques qui permettent de combiner, de façon optimale (en un sens à définir), l'information mathématique contenue dans les équations et l'information physique provenant des observations en vue de reconstituer l'état de l'écoulement. ».

Ce chapitre présente les principales méthodes d'assimilation de données qui ont été appliquées à des modèles de météorologie et de chimie atmosphérique, en donnant leur justification théorique, leurs hypothèses et leurs limitations. Parmi ces méthodes, nous insistons sur le fait que l'IEKF-Q, présenté à la section 2.5.4, est une nouvelle méthode qui a fait l'objet d'une soumission ([Sakov et al., 2017, soumis](#)).

Sommaire

2.1	Observations	19
2.2	BLUE et le filtre de Kalman	20
2.2.1	L'algorithme	20
2.2.2	Ses limitations et les premières solutions	21
2.3	Méthodes variationnelles	22
2.3.1	3D-Var et parallèle avec BLUE	22
2.3.2	4D-Var et parallèle avec le filtre de Kalman	23
2.4	Méthodes ensemblistes	24
2.4.1	Filtres de Kalman d'ensemble : dérivations et algorithmes	24
2.4.1.1	EnKF stochastique	24
2.4.1.2	EnKF déterministe	26
2.4.1.3	Formulation variationnelle de l'ETKF	28
2.4.2	Variantes de l'EnKF	29
2.4.2.1	Un EnKF sans inversion de matrice : le DEnKF	30
2.4.2.2	Un ETKF avec opérateur d'observation non-linéaire : le MLEF	30
2.4.3	Le lisseur de Kalman d'ensemble	31
2.4.4	Localisation	34
2.4.4.1	Localisation des covariances	35
2.4.4.2	Localisation par domaine	35
2.4.5	Inflation	36
2.4.5.1	Inflation multiplicative	36
2.4.5.2	L'EnKF-N, un exemple d'inflation adaptative	37

2.4.5.3	Inflation additive et erreur modèle	38
2.5	Méthodes variationnelles d'ensemble	39
2.5.1	Méthodes hybrides	40
2.5.2	4DEnVar	40
2.5.3	Le lisseur de Kalman d'ensemble itératif : IEnKS	42
2.5.3.1	L'algorithme	42
2.5.3.2	Le MDA : assimilation multiple des observations	45
2.5.3.3	Localisation de l'IEnKS	45
2.5.4	Formulation en contrainte faible : l'IEnKF-Q	46
2.5.4.1	Dérivation	48
2.5.4.2	Découplage de $\mathbf{u}$ et $\mathbf{v}$ quand l'opérateur d'observation est linéaire	50
2.5.4.3	Algorithmes	51
2.5.4.4	Discussion	53
2.6	Validation et indicateurs statistiques	55
2.6.1	Erreur quadratique : RMSE	55
2.6.2	Corrélation de Pearson	56
2.6.3	Biais	56
2.6.4	Indicateurs statistiques χ^2 et RCRV	56

Nous nous sommes appuyés pour la rédaction de ce chapitre sur les notes de cours, livres et articles suivants :

- M. BOCQUET : Introduction to the principles and methods of data assimilation in the geosciences. Notes de cours, École des Ponts ParisTech, 2014. URL <http://cerea.enpc.fr/HomePages/bocquet/teaching/assim-mb-en.pdf>. Cours du Master2 OACOS & WAPE
- P. SAKOV : EnKF lectures. Notes de cours, NERSC, 2011. URL <http://enkf.nersc.no/Presentations/EnKFlecturesbyPavelSakov/>. School on Data assimilation
- M. ASCH, M. BOCQUET et M. NODET : Chapter 6: The ensemble Kalman filter. *In Data assimilation: methods, algorithms, and applications* Asch et al. (2016c), chap. 6, p. 153–193
- M. ASCH, M. BOCQUET et M. NODET : Chapter 7: Ensemble variational methods. *In Data assimilation: methods, algorithms, and applications* Asch et al. (2016c), chap. 7, p. 195–216
- P. SAKOV, J.-M. HAUSSAIRE et M. BOCQUET : An iterative ensemble Kalman filter in presence of additive model error. *Q. J. R. Meteorol. Soc.*, 2017, soumis.

2.1 Observations

Avant de rentrer dans les détails des différentes méthodes d'assimilation de données, nous jugeons utile de présenter brièvement certains des différents types d'observations d'ozone et de polluants disponibles. Les sources d'observations sont en effet très variées. Pour ce qui est des données météorologiques aussi bien que de la pollution atmosphérique, en plus des satellites qui apportent une masse immense d'informations, on peut noter les nombreuses mesures in situ, issues des vols commerciaux, ou encore de bouées en mer. Chacune de ces observations apporte des informations complémentaires mais nécessite des traitements différents et soulève leur lot de problèmes quand on essaie de les incorporer dans des simulations issues de modèles. En effet, une mesure in situ de concentration d'ozone par exemple pourra être très précise mais fortement influencée par une source très proche de la station de mesure et ainsi ne pas être représentative de la concentration moyenne à l'intérieur d'une grande maille d'un modèle. Par ailleurs, ces données sont souvent fournies avec une fréquence horaire et sont inégalement réparties spatialement. Les données satellitaires quant à elles, bien que précises, en grande quantité et fréquemment mises à jour, apportent de nouvelles questions avec chaque nouveau satellite mis en orbite. C'est le cas par exemple de la mission *Surface Water & Ocean Topography* (SWOT), prochainement lancée, qui soulève la question de la corrélation spatiale des erreurs d'observation (Ruggiero *et al.*, 2016). Ainsi, pour une observation donnée, on notera deux types d'erreur (Cohn, 1997) :

- l'erreur de mesure instrumentale - relativement faible, elle dépend de l'appareil utilisé pour effectuer la mesure ;
- l'erreur de représentativité - comme son nom l'indique, elle exprime dans quelle mesure une observation est représentative du phénomène modélisé. Elle dépend donc du modèle utilisé (particulièrement sa résolution) et de la qualité de la mesure (distance d'une source influente par exemple).

Nous nous sommes uniquement intéressés aux mesures in situ de concentrations d'ozone au sol. Il existe plusieurs réseaux de mesure dans le monde, notamment la Base de Données de la Qualité de l'Air (BDQA) en France ou European Monitoring and Evaluation Programme (EMEP) (Hjellbrekke *et al.*, 2011) et AirBase (Simeoens, 2013) à l'échelle européenne. Les stations de mesure d'AirBase d'ozone sont catégorisées en fonction de leur localisation (urbaine, périurbaine, rurale) et du type de pollution qu'elles observent (industrielle, de trafic, de fond). Ainsi, pour une simulation à l'échelle européenne, seules les stations de mesure de la pollution de fond sont adéquates, mais peuvent néanmoins avoir une erreur de représentativité importante en fonction de la résolution du domaine de simulation. On montre la position des stations de mesure d'AirBase qualifiées de rurales et de fond sur la Fig. 2.1.

Gaubert *et al.* (2014) ont requalifié les stations d'observation d'AirBase en leur attribuant une nouvelle localisation. Ils ont pour cela regroupé entre elles des stations qui avaient une valeur moyenne et une variabilité journalière similaires. Ainsi, des stations qu'AirBase considère comme rurales et de fond peuvent être qualifiées d'urbaines par la nouvelle classification utilisée dans Gaubert *et al.* (2014) et détaillée dans Gaubert (2013). Ils distinguent ainsi 5 classes d'observations, à savoir les stations de fond "MOU", séparées en fonction de leur altitude en "MOU<300m" et "MOU>300m", les stations rurales "RUR", les stations périurbaines "PUR" et les stations urbaines "URB". Quand cette classification d'observations sera utilisée, on les nommera les observations de Gaubert.

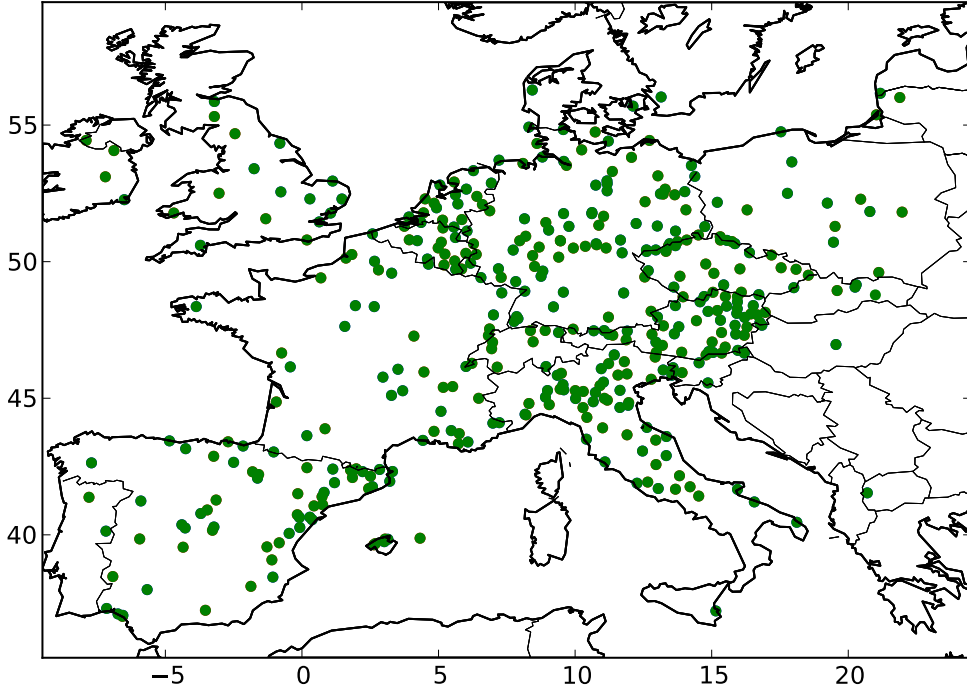


FIGURE 2.1 – Sites de mesure d’ozone du réseau AirBase notés ruraux et de fond.

2.2 BLUE et le filtre de Kalman

2.2.1 L’algorithme

Le filtre de Kalman tel qu’il a été développé initialement il y a de ça plus de 50 ans (Kalman, 1960) est désormais devenu suffisamment connu pour que nous nous contentions de présenter ici, succinctement, les hypothèses et résultats de celui-ci.

Étape d’analyse :

Notons $\mathbf{x}_f \in \mathbb{R}^n$ le vecteur d’état issu d’une simulation (appelé ébauche, f pour *forecast* en anglais), $\mathbf{y} \in \mathbb{R}^p$ le vecteur d’observation et $\mathbf{x}_t \in \mathbb{R}^n$ la vérité de l’état physique que l’on essaie d’approcher. On suppose que les erreurs d’ébauche et d’observation sont gaussiennes et sans biais, c’est-à-dire

$$\mathbf{x}_f = \mathbf{x}_t + \boldsymbol{\varepsilon}_f, \quad \text{avec} \quad \boldsymbol{\varepsilon}_f \sim \mathcal{N}(\mathbf{0}, \mathbf{P}^f), \quad (2.1)$$

$$\mathbf{y} = \mathbf{H}\mathbf{x}_t + \boldsymbol{\varepsilon}_o, \quad \text{avec} \quad \boldsymbol{\varepsilon}_o \sim \mathcal{N}(\mathbf{0}, \mathbf{R}), \quad (2.2)$$

$$(2.3)$$

où $\mathbf{H} \in \mathbb{R}^{p \times n}$ est un opérateur linéaire d’observation qui projette l’état sur la position des observations, et où $\mathbf{P}^f \in \mathbb{R}^{n \times n}$ et $\mathbf{R} \in \mathbb{R}^{p \times p}$ sont respectivement les matrices de covariance des erreurs d’ébauche et d’observation. Par ailleurs, on suppose les erreurs d’ébauche et d’observation indépendantes, c’est-à-dire $\mathbb{E}[\boldsymbol{\varepsilon}_o \boldsymbol{\varepsilon}_f] = \mathbf{0}$, où $\mathbb{E}$ désigne l’espérance.

L’état dit analysé $\mathbf{x}_a \in \mathbb{R}^n$, qui est de biais nul et tel que la trace de la matrice de covariance des erreurs d’analyse est minimale, est donné par

$$\mathbf{x}_a = \mathbf{x}_f + \mathbf{K}(\mathbf{y} - \mathbf{H}\mathbf{x}_f), \quad (2.4)$$

où $\mathbf{K} \in \mathbb{R}^{n \times p}$ est le gain de Kalman, qui s'écrit

$$\mathbf{K} = \mathbf{P}^f \mathbf{H}^T (\mathbf{R} + \mathbf{H} \mathbf{P}^f \mathbf{H}^T)^{-1}. \quad (2.5)$$

La matrice de covariance des erreurs d'analyse est donnée par

$$\mathbf{P}^a = (\mathbf{I}_n - \mathbf{K} \mathbf{H}) \mathbf{P}^f, \quad (2.6)$$

où $\mathbf{I}_n \in \mathbb{R}^{n \times n}$ est la matrice identité. Cette analyse est appelée *best linear unbiased estimator* (BLUE).

La formule de Sherman-Morrison-Woodbury (SMW), qui établit que pour des matrices $\mathbf{A}$, $\mathbf{U}$, $\mathbf{C}$, $\mathbf{V}$ de bonnes tailles

$$(\mathbf{A} + \mathbf{U} \mathbf{C} \mathbf{V})^{-1} = \mathbf{A}^{-1} - \mathbf{A}^{-1} \mathbf{U} (\mathbf{C}^{-1} + \mathbf{V} \mathbf{A}^{-1} \mathbf{U})^{-1} \mathbf{V} \mathbf{A}^{-1}, \quad (2.7)$$

nous permet d'avoir une formulation équivalente du gain de Kalman

$$\mathbf{K} = (\mathbf{P}^{f-1} + \mathbf{H}^T \mathbf{R}^{-1} \mathbf{H})^{-1} \mathbf{H}^T \mathbf{R}^{-1}. \quad (2.8)$$

Étape de propagation :

Pour cycloper le filtre de Kalman, on propage alors l'état analysé $\mathbf{x}_a$ en utilisant le modèle $\mathbf{M}_k$, que l'on suppose linéaire, du temps k jusqu'au temps $k + 1$ où de nouvelles observations seront disponible. Ainsi, on obtient,

$$\mathbf{x}_f^{k+1} = \mathbf{M}_k \mathbf{x}_a^k \quad (2.9)$$

et la matrice de covariance des erreurs d'ébauche est donnée par

$$\mathbf{P}_{k+1}^f = \mathbf{M}_k \mathbf{P}_k^a \mathbf{M}_k^T + \mathbf{Q}_k, \quad (2.10)$$

où $\mathbf{Q}_k \in \mathbb{R}^{n \times n}$ est la matrice de covariance de l'erreur modèle induite par la propagation, définie comme

$$\mathbf{x}_t^{k+1} = \mathbf{M}_k \mathbf{x}_t^k + \boldsymbol{\varepsilon}_q^k, \quad \text{avec} \quad \boldsymbol{\varepsilon}_q^k \sim \mathcal{N}(\mathbf{0}, \mathbf{Q}_k). \quad (2.11)$$

On se retrouve alors avec une nouvelle ébauche, un nouveau jeu d'observations et de nouvelles statistiques d'erreur permettant de recommencer l'étape d'analyse. C'est le cycle alliant analyse puis propagation que l'on nomme filtre de Kalman.

2.2.2 Ses limitations et les premières solutions

On rappelle que les hypothèses principales de ce filtre sont la linéarité des opérateurs d'observation $\mathbf{H}$ et de propagation $\mathbf{M}$, ainsi que la gaussianité et l'absence de biais des statistiques des erreurs d'ébauche, d'observation et de modèle. Toutes ces hypothèses sont simplificatrices et ne tiennent pas dans beaucoup de cas pratiques. Le but des différentes méthodes d'assimilation de données est alors de s'abstraire de certaines de ces hypothèses. Nous supposons dans la majeure partie des cas que $\mathbf{Q} = \mathbf{0}$, c'est-à-dire que le modèle est parfait, qu'il n'y a pas d'erreur modèle.

Le filtre de Kalman étendu propose par exemple, dans le cas d'un modèle $\mathcal{M}$ ou d'un opérateur d'observation $\mathcal{H}$ non-linéaire, d'utiliser leur tangent linéaire $\mathbf{H}$ et $\mathbf{M}$ dans les formules (2.5), (2.6) et (2.10). Le pseudo-code correspondant à un cycle du filtre de Kalman étendu est explicité dans l'algorithme 1. Nous noterons par la suite les opérateurs en cursif s'ils sont non-linéaires ou en gras s'ils sont linéaires ou si l'on utilise leur tangent linéaire.

Par ailleurs, pour des modèles, notamment de météorologie ou de chimie atmosphérique, la taille des vecteurs d'état (et dans une moindre mesure d'observation) peut être de l'ordre du milliard de variables. Ceci prohibe absolument la propagation complète par le modèle des matrices de covariance des erreurs d'analyse comme dans l'équation (2.10), équation qui en outre nécessite l'adjoint du modèle (ou de son tangent linéaire). C'est pourquoi un autre champ de développement des méthodes d'assimilation de données s'attèle à alléger le coût numérique du calcul des matrices de covariance des erreurs d'analyse et d'ébauche. La première possibilité d'algorithme allant dans ce sens est l'interpolation optimale, qui consiste à fixer tout au long des cycles d'analyse la matrice $\mathbf{P}^f$. Les méthodes dites ensemblistes que nous présenterons dans la section 2.4 permettent une représentation et une propagation allégée des erreurs d'ébauche et d'analyse.

Algorithme 1 Algorithme pour un cycle du filtre de Kalman étendu.

Requis: : $\mathbf{x}_f, \mathbf{P}^f, \mathbf{y}, \mathbf{R}, \mathcal{M}, \mathbf{M}, \mathcal{H}, \mathbf{H}$

- 1: $\mathbf{K} = \mathbf{P}^f \mathbf{H}^T (\mathbf{H} \mathbf{P}^f \mathbf{H}^T + \mathbf{R})^{-1}$
 - 2: $\mathbf{x}_a = \mathbf{x}_f + \mathbf{K}(\mathbf{y} - \mathcal{H}(\mathbf{x}_f))$
 - 3: $\mathbf{P}^a = (\mathbf{I}_n - \mathbf{K} \mathbf{H}) \mathbf{P}^f$
 - 4: $\mathbf{x}_f = \mathcal{M}(\mathbf{x}_a)$
 - 5: $\mathbf{P}^f = \mathbf{M} \mathbf{P}^a \mathbf{M}^T$
-

2.3 Méthodes variationnelles

Pour s'abstraire de l'hypothèse des opérateurs d'observation et de modèle linéaires, on peut recourir aux méthodes dites variationnelles. Il s'agit de définir une fonction de coût à minimiser, avec ces opérateurs non-linéaires. Une fois cette fonctionnelle définie, nous avons à notre disposition l'ensemble des méthodes itératives d'optimisation (Nocedal et Wright, 2006) pour en déterminer le minimum, en linéarisant à chaque étape les opérateurs non-linéaires autour de l'optimum en cours. Nous cherchons donc à définir un formalisme où les opérateurs $\mathcal{H}$ et $\mathcal{M}$ sont potentiellement non-linéaires et qui soit équivalent au formalisme du filtre de Kalman quand ils sont linéaires.

2.3.1 3D-Var et parallèle avec BLUE

Définissons la fonction de coût suivante

$$\mathcal{J}(\mathbf{x}) = \frac{1}{2} \|\mathbf{x} - \mathbf{x}_f\|_{\mathbf{P}^f}^2 + \frac{1}{2} \|\mathbf{y} - \mathcal{H}(\mathbf{x})\|_{\mathbf{R}}^2, \quad (2.12)$$

où $\|\mathbf{x}\|_{\mathbf{A}}^2 = \mathbf{x}^T \mathbf{A}^{-1} \mathbf{x}$. La formule de son gradient est

$$\nabla \mathcal{J}(\mathbf{x}) = \mathbf{P}^{f-1}(\mathbf{x} - \mathbf{x}_f) - \mathbf{H}^T \mathbf{R}^{-1}(\mathbf{y} - \mathcal{H}(\mathbf{x})). \quad (2.13)$$

où $\mathbf{H}$ est le linéarisé de $\mathcal{H}$ autour de $\mathbf{x}$. Si $\mathcal{H}$ est linéaire, cette fonctionnelle est convexe étant donné que $\mathbf{P}^f$ et $\mathbf{R}$ sont définies positives. Elle a donc un minimum, qui est atteint en $\mathbf{x}^*$ tel que $\nabla \mathcal{J}(\mathbf{x}^*) = \mathbf{0}$, c'est-à-dire

$$\mathbf{x}^* = \mathbf{x}_f + (\mathbf{P}^{f-1} + \mathbf{H}^T \mathbf{R}^{-1} \mathbf{H})^{-1} \mathbf{H}^T \mathbf{R}^{-1}(\mathbf{y} - \mathcal{H}(\mathbf{x}_f)). \quad (2.14)$$

On retrouve alors la formule du gain de Kalman (2.8), et $\mathbf{x}^*$ est égal à l'analyse BLUE (2.4). Sinon, on peut utiliser un algorithme de minimisation itératif de type Newton, ou quasi-Newton

par exemple. On remarquera par ailleurs que l'inverse de la Hessienne de la fonction de coût dans le cas linéaire, noté $\mathbf{D}$, est égal à

$$\mathbf{D} = (\mathbf{P}^{\mathbf{f}-1} + \mathbf{H}^T \mathbf{R}^{-1} \mathbf{H})^{-1} = \mathbf{P}^{\mathbf{a}}. \quad (2.15)$$

L'algorithme consistant à minimiser cette fonction de coût est appelé 3D-Var, et il fournit un résultat équivalent au BLUE (dans le cas d'un opérateur d'observation linéaire). Si on cycle en propageant l'état obtenu après la minimisation et en conservant la matrice $\mathbf{P}^{\mathbf{f}}$ intacte, on retrouve l'algorithme de l'interpolation optimale.

2.3.2 4D-Var et parallèle avec le filtre de Kalman

On souhaiterait également avoir une équivalence avec le filtre de Kalman complet, au moins dans les cas $\mathbf{H}$ et $\mathbf{M}$ linéaires. On suppose également que l'on est ici en contrainte forte, c'est-à-dire $\mathbf{Q} = \mathbf{0}$. On considère une fonctionnelle définie sur une fenêtre de temps de t_0 à t_L , dans laquelle on a accès aux observations $\mathbf{y}_0, \dots, \mathbf{y}_L$:

$$\mathcal{J}(\mathbf{x}_0) = \frac{1}{2} \|\mathbf{x}_0 - \mathbf{x}_{\mathbf{f}}\|_{\mathbf{P}^{\mathbf{f}}}^2 + \frac{1}{2} \sum_{l=0}^L \|\mathbf{y}_l - \mathcal{H}_l(\mathbf{x}_l)\|_{\mathbf{R}_l}^2, \quad (2.16)$$

où $\mathbf{x}_l$ est défini par récurrence, $\mathbf{x}_l = \mathcal{M}_{l \leftarrow (l-1)}(\mathbf{x}_{l-1})$, c'est-à-dire $\mathbf{x}_l = \mathcal{M}_{l \leftarrow 0}(\mathbf{x}_0)$ avec $\mathcal{M}_{l \leftarrow 0} = \mathcal{M}_{l \leftarrow (l-1)} \circ \mathcal{M}_{(l-1) \leftarrow (l-2)} \circ \dots \circ \mathcal{M}_{1 \leftarrow 0}$ et $\circ$ désigne la composition d'opérateurs. Par convention, on notera $\mathcal{M}_{0 \leftarrow 0} = \mathbf{I}_n$.

Si on note encore une fois $\mathbf{M}$ et $\mathbf{H}$ les linéarisés du modèle et de l'opérateur d'observation, le gradient de $\mathcal{J}$ s'écrit

$$\nabla_{\mathbf{x}_0} \mathcal{J}(\mathbf{x}_0) = \mathbf{P}^{\mathbf{f}-1}(\mathbf{x}_0 - \mathbf{x}_{\mathbf{f}}) - \sum_{l=0}^L \mathbf{M}_{l \leftarrow 0}^T \mathbf{H}_l^T \mathbf{R}_l^{-1} (\mathbf{y}_l - \mathcal{H}_l \circ \mathcal{M}_{l \leftarrow 0}(\mathbf{x}_0)). \quad (2.17)$$

Une expression approchée de l'inverse de la Hessienne est alors donnée par

$$\mathbf{D} = \left(\mathbf{P}^{\mathbf{f}-1} + \sum_{l=0}^L \mathbf{M}_{l \leftarrow 0}^T \mathbf{H}_l^T \mathbf{R}_l^{-1} \mathbf{H}_l \mathbf{M}_{l \leftarrow 0} \right)^{-1}. \quad (2.18)$$

À nouveau, dans le cas $\mathcal{H}$ et $\mathcal{M}$ linéaires, la solution du problème quadratique est donné directement par

$$\mathbf{x}_{\mathbf{a}} = \mathbf{x}_{\mathbf{f}} - \mathbf{D} \nabla_{\mathbf{x}_{\mathbf{f}}} \mathcal{J}(\mathbf{x}_{\mathbf{f}}). \quad (2.19)$$

Sinon, on doit recourir à des méthodes itératives de minimisation. Enfin, la matrice de covariance des erreurs d'analyse est, comme dans (2.15), l'inverse de la Hessienne.

Une propriété importante de la minimisation de cette fonctionnelle est le principe de transfert d'optimalité. Celui-ci exprime que le résultat $\mathbf{x}_L$ au temps t_L , à la fin de la fenêtre, est le même si (1) on effectue la minimisation de $\mathcal{J}$ en une seule fois à partir de $\mathbf{x}_0$, ou (2) on divise l'intervalle $[t_0, t_L]$ en deux intervalles $[t_0, t_k]$ et $[t_k, t_L]$ et on minimise une première fois en partant de $\mathbf{x}_0$ et en allant jusqu'à t_k , pour obtenir $\mathbf{x}_k^{\mathbf{a}}$ et $\mathbf{P}_k^{\mathbf{a}}$, puis on minimise à nouveau de t_k jusqu'à t_L en utilisant ces $\mathbf{x}_k^{\mathbf{a}}$ et $\mathbf{P}_k^{\mathbf{a}}$ comme ébauche. Cela ne signifie pas pour autant que les deux trajectoires sont équivalentes à tous les pas de temps, mais seulement à la fin de la fenêtre.

C'est cette propriété, couplée avec le fait que l'analyse du 3D-Var est équivalente au BLUE, qui fait que le 4D-Var est équivalent au filtre de Kalman dans le cas linéaire. En effet, si on subdivise la

fenêtre $[t_0, t_L]$ en L intervalles de taille 1, chaque analyse de type 3D-Var est équivalente au BLUE, et le résultat du 4D-Var au temps final t_L sera ainsi équivalent au résultat du filtre de Kalman grâce au transfert d'optimalité.

Le principal désavantage du 4D-Var est la nécessité d'avoir à sa disposition l'adjoint du tangent linéaire du modèle pour calculer le gradient de $\mathcal{J}$. Il est malheureusement très coûteux en terme de main d'œuvre et de rigueur de développer en parallèle un modèle et son adjoint. Les modèles sont en effet en constante évolution et progrès tandis que l'implémentation de l'adjoint du modèle est souvent en retard par rapport à ces évolutions. Une contrainte supplémentaire est la nécessité de bien paramétrer la matrice $\mathbf{P}^f$ dans l'algorithme. En revanche, cette méthode effectue une analyse variationnelle qui permet de bien gérer la non-linéarité du modèle.

2.4 Méthodes ensemblistes

2.4.1 Filtres de Kalman d'ensemble : dérivations et algorithmes

La première dérivation du filtre de Kalman d'ensemble (ensemble Kalman filter (EnKF)), date de Evensen (1994) qui le corrigera par la suite dans Burgers *et al.* (1998). Cet algorithme ne remet pas en question l'hypothèse de gaussianité des statistiques d'erreur, mais se concentre sur une représentation simplifiée de ces statistiques. En effet, il repose sur l'idée d'avoir un échantillon de vecteurs d'état dont la dispersion représente les incertitudes associées à l'ébauche. Cet ensemble de vecteurs peut alors être propagé par le modèle, qu'il soit linéaire ou non, ce qui permet une propagation de l'erreur d'analyse sans avoir à recourir à l'adjoint du tangent linéaire du modèle, ni à définir de grosses matrices d'erreur.

Il existe deux variantes principales de l'EnKF. La première est dite stochastique, car elle a recours à une perturbation stochastique des observations dans l'analyse. La deuxième est dite déterministe, par opposition, car il n'y a pas de telle perturbation stochastique dans sa dérivation.

2.4.1.1 EnKF stochastique

Étape d'analyse :

Nous avons à notre disposition un ensemble de m vecteurs d'ébauche $\mathbf{x}_f^i$ où $i \in [1, m]$ représente l'indice du membre dans l'ensemble. En règle générale, $m \ll n$. La dispersion de l'ensemble est censée représenter l'incertitude associée à notre ébauche, de telle sorte que l'on estime $\mathbf{P}^f$ à partir de l'ensemble :

$$\mathbf{P}^f = \frac{1}{m-1} \sum_{i=1}^m (\mathbf{x}_f^i - \bar{\mathbf{x}}_f)(\mathbf{x}_f^i - \bar{\mathbf{x}}_f)^T, \quad \text{avec} \quad \bar{\mathbf{x}}_f = \frac{1}{m} \sum_{i=1}^m \mathbf{x}_f^i. \quad (2.20)$$

On peut également factoriser $\mathbf{P}^f$ sous la forme

$$\mathbf{P}^f = \mathbf{X}_f \mathbf{X}_f^T, \quad (2.21)$$

où $\mathbf{X}_f \in \mathbb{R}^{n \times m}$ désigne les anomalies normalisées, telles que, pour la colonne $i \in [1, m]$,

$$[\mathbf{X}_f]_i = \frac{\mathbf{x}_f^i - \bar{\mathbf{x}}_f}{\sqrt{m-1}}. \quad (2.22)$$

Si on effectue une analyse en appliquant la formule du BLUE (2.4) pour chacun des éléments de l'ensemble, on obtient alors un ensemble de vecteurs d'état analysés $\{\mathbf{x}_a^i\}_{i \in [1, m]}$, dont les anomalies

sont de la forme

$$[\mathbf{X}_a]_i = \frac{\mathbf{x}_a^i - \bar{\mathbf{x}}_a}{\sqrt{m-1}} = [\mathbf{X}_f]_i + \mathbf{K}(\mathbf{0} - \mathbf{H}[\mathbf{X}_f]_i) = (\mathbf{I}_n - \mathbf{KH})[\mathbf{X}_f]_i. \quad (2.23)$$

Ceci donne alors une matrice de covariance des erreurs d'analyse

$$\mathbf{P}^a = \mathbf{X}_a \mathbf{X}_a^T = (\mathbf{I}_n - \mathbf{KH}) \mathbf{X}_f \mathbf{X}_f^T (\mathbf{I}_n - \mathbf{KH})^T = (\mathbf{I}_n - \mathbf{KH}) \mathbf{P}^f (\mathbf{I}_n - \mathbf{KH})^T. \quad (2.24)$$

Cependant, la formule théorique du filtre de Kalman donne une erreur $\mathbf{P}^a$ selon la formule (2.6)

$$\mathbf{P}^a = (\mathbf{I}_n - \mathbf{KH}) \mathbf{P}^f = (\mathbf{I}_n - \mathbf{KH}) \mathbf{X}_f \mathbf{X}_f^T (\mathbf{I}_n - \mathbf{KH})^T + \mathbf{K} \mathbf{R} \mathbf{K}^T. \quad (2.25)$$

Ainsi, la matrice fournie par (2.24) est sous-estimée car il manque le deuxième terme de (2.25), ce qui risque de conduire à la sous-dispersion de l'ensemble et à la divergence du filtre.

Pour pallier ce problème, dans le cas d'une erreur d'observation non-biaisée et symétrique, une solution est de perturber le vecteur d'observation $\mathbf{y}$ pour chacun des membres de l'ensemble :

$$\mathbf{y} \rightarrow \mathbf{y}_i \equiv \mathbf{y} + \mathbf{u}_i, \quad (2.26)$$

où $\mathbf{u}_i$ est tiré selon la distribution gaussienne $\mathbf{u}_i \sim \mathcal{N}(\mathbf{0}, \mathbf{R})$. Notons $\bar{\mathbf{u}}$ la moyenne de l'ensemble des $\mathbf{u}_i$. La formule (2.23) devient

$$[\mathbf{X}_a]_i = (\mathbf{I}_n - \mathbf{KH})[\mathbf{X}_f]_i + \mathbf{K} \mathbf{E}_i, \quad \text{avec} \quad \mathbf{E}_i = \frac{(\mathbf{u}_i - \bar{\mathbf{u}})}{\sqrt{m-1}}. \quad (2.27)$$

Si on note alors $\mathbf{E} = [\mathbf{E}_1, \dots, \mathbf{E}_m]$, on obtient $\mathbf{X}_a = (\mathbf{I}_n - \mathbf{KH}) \mathbf{X}_f + \mathbf{K} \mathbf{E}$, soit

$$\mathbf{P}^a = (\mathbf{I}_n - \mathbf{KH}) \mathbf{X}_f \mathbf{X}_f^T (\mathbf{I}_n - \mathbf{KH})^T + \mathbf{K} \mathbf{E} \mathbf{E}^T \mathbf{K}^T + (\mathbf{I}_n - \mathbf{KH}) \mathbf{X}_f \mathbf{E}^T \mathbf{K}^T + \mathbf{K} \mathbf{E} \mathbf{X}_f^T (\mathbf{I}_n - \mathbf{KH})^T, \quad (2.28)$$

dont l'espérance statistique donne la valeur escomptée

$$\mathbb{E}[\mathbf{P}^a] = (\mathbf{I}_n - \mathbf{KH}) \mathbf{P}^f (\mathbf{I}_n - \mathbf{KH})^T + \mathbf{K} \mathbb{E}[\mathbf{E} \mathbf{E}^T] \mathbf{K}^T \quad (2.29)$$

$$= (\mathbf{I}_n - \mathbf{KH}) \mathbf{X}_f \mathbf{X}_f^T (\mathbf{I}_n - \mathbf{KH})^T + \mathbf{K} \mathbf{R} \mathbf{K}^T = (\mathbf{I}_n - \mathbf{KH}) \mathbf{P}^f. \quad (2.30)$$

Le qualificatif de stochastique décrivant cet algorithme provient de cette étape de perturbation des observations (2.26). Bien qu'étant une solution élégante au problème qu'elle résout, elle introduit un bruit numérique d'échantillonnage.

Si $\mathcal{H}$ est non-linéaire, la formule de $\mathbf{K}$ a recours au tangent linéaire de $\mathcal{H}$. Mais l'ensemble peut nous servir à faire des différences finies pour approcher le tangent linéaire et son adjoint à partir de l'opérateur non-linéaire. Les facteurs $\mathbf{P}^f \mathbf{H}^T$ ou $\mathbf{H} \mathbf{P}^f \mathbf{H}^T$ qui apparaissent dans la formule de $\mathbf{K}$ peuvent s'approcher en remarquant que

$$\mathbf{P}^f \mathbf{H}^T = \mathbf{X}_f (\mathbf{H} \mathbf{X}_f)^T, \quad (2.31)$$

$$\mathbf{H} \mathbf{P}^f \mathbf{H}^T = (\mathbf{H} \mathbf{X}_f) (\mathbf{H} \mathbf{X}_f)^T, \quad (2.32)$$

et que

$$[\mathbf{H} \mathbf{X}_f]_i = \mathbf{H}(\mathbf{x}_f^i - \bar{\mathbf{x}}_f) \simeq \mathcal{H}(\mathbf{x}_f^i) - \bar{\mathbf{y}}_f \quad \text{pour } i \in [1, m], \quad \text{avec } \bar{\mathbf{y}}_f = \frac{1}{m} \sum_{i=1}^m \mathcal{H}(\mathbf{x}_f^i). \quad (2.33)$$

$$(2.34)$$

On peut alors noter

$$[\mathbf{Y}_f]_i = \frac{\mathcal{H}(\mathbf{x}_f^i) - \bar{\mathbf{y}}_f - \mathbf{u}_i + \bar{\mathbf{u}}}{\sqrt{m-1}}, \quad (2.35)$$

et remarquer que le gain de Kalman peut s'écrire sous la forme

$$\mathbf{K} = \mathbf{X}_f \mathbf{Y}_f^T (\mathbf{Y}_f \mathbf{Y}_f^T)^{-1}, \quad (2.36)$$

puisque $\mathbf{X}_f \mathbf{Y}_f^T$ échantillonne $\mathbf{P}^f \mathbf{H}^T$ et $\mathbf{Y}_f \mathbf{Y}_f^T$ échantillonne $\mathbf{R} + \mathbf{H} \mathbf{P}^f \mathbf{H}^T$.

Étape de propagation :

Contrairement au filtre de Kalman où il est nécessaire de propager les matrices de covariance des erreurs d'analyse, il suffit ici de propager l'intégralité des $m \ll n$ vecteurs de l'ensemble avec le modèle, potentiellement non-linéaire, pour obtenir l'ensemble des vecteurs d'ébauche pour l'analyse suivante

$$\mathbf{x}_f^{i,k+1} = \mathcal{M}_k(\mathbf{x}_a^{i,k}), \quad \text{pour } i \in [1, m]. \quad (2.37)$$

On évite ainsi le recours au tangent linéaire du modèle requis par le filtre de Kalman étendu.

Le pseudo-code correspondant à un cycle de l'EnKF stochastique est explicité dans l'algorithme 2.

Algorithme 2 Algorithme pour un cycle de l'EnKF stochastique.

Requis: : $m, \{\mathbf{x}_f^i\}_{i \in [1, m]}, \mathbf{y}, \mathbf{R}, \mathcal{M}, \mathcal{H}$

```

1: for  $i = 1 \dots m$  do
2:    $\mathbf{y}_i = \mathbf{y} + \mathbf{u}_i$  avec  $\mathbf{u}_i \sim \mathcal{N}(\mathbf{0}, \mathbf{R})$ 
3: end for
4:  $\bar{\mathbf{x}}_f = \frac{1}{m} \sum_{i=1}^m \mathbf{x}_f^i$ 
5:  $\bar{\mathbf{u}} = \frac{1}{m} \sum_{i=1}^m \mathbf{u}_i$ 
6:  $\bar{\mathbf{y}}_f = \frac{1}{m} \sum_{i=1}^m \mathcal{H}(\mathbf{x}_f^i)$ 
7: for  $i = 1 \dots m$  do
8:    $[\mathbf{X}_f]_i = \frac{\mathbf{x}_f^i - \bar{\mathbf{x}}_f}{\sqrt{m-1}}$ 
9:    $[\mathbf{Y}_f]_i = \frac{\mathcal{H}(\mathbf{x}_f^i) - \bar{\mathbf{y}}_f - \mathbf{u}_i + \bar{\mathbf{u}}}{\sqrt{m-1}}$ 
10: end for
11:  $\mathbf{K} = \mathbf{X}_f \mathbf{Y}_f^T (\mathbf{Y}_f \mathbf{Y}_f^T)^{-1}$ 
12: for  $i = 1 \dots m$  do
13:    $\mathbf{x}_a^i = \mathbf{x}_f^i + \mathbf{K}(\mathbf{y}_i - \mathcal{H}(\mathbf{x}_f^i))$ 
14:    $\mathbf{x}_f^i = \mathcal{M}(\mathbf{x}_a^i)$ 
15: end for
```

2.4.1.2 EnKF déterministe

Pour s'abstraire de la perturbation stochastique des observations, une variété de filtres de Kalman d'ensemble dits déterministes a été développée. Ces filtres sont réunis sous le titre d'ensemble square root Kalman filter (EnSRKF) (Tippett *et al.*, 2003). Nous nous concentrerons plus particulièrement sur une implémentation nommée ensemble transform Kalman filter (ETKF) (Bishop *et al.*, 2001; Ott *et al.*, 2004; Hunt *et al.*, 2007) qui a été appliquée dans un large éventail de domaines. La majorité des filtres EnSRKF sont algébriquement équivalents et leurs implémentations

ne donnent des résultats différents que par l'usage de la localisation (voir 2.4.4). Une bonne illustration de ce phénomène, pour deux lisseurs, est mis en avant par Raanes (2016), qui démontre l'équivalence entre deux variétés de lisseurs de Kalman d'ensemble.

Nous définissons $\mathbf{P}^f$ et $\mathbf{X}_f$ selon les équations (2.21) et (2.22), et nous notons $\mathbf{Y}_f = \mathbf{H}\mathbf{X}_f$. Si l'opérateur d'observation est non-linéaire, cette expression est soit définie avec le tangent linéaire de ce dernier, soit en utilisant la formule (2.33). Nous supposons par ailleurs que l'analyse se situe dans le sous-espace de l'ensemble, c'est-à-dire le sous-espace affine engendré par les anomalies, dont les vecteurs $\mathbf{x}$ sont de la forme

$$\mathbf{x} = \bar{\mathbf{x}}_f + \mathbf{X}_f \mathbf{w}, \quad (2.38)$$

où $\mathbf{w} \in \mathbb{R}^m$ est le vecteur des coordonnées dans le sous-espace de l'ensemble. Ces coordonnées ne sont pas uniques car le vecteur $\mathbf{w} + \lambda \mathbf{1}$ avec $\mathbf{1} = (1, \dots, 1)^T \in \mathbb{R}^m$ et $\lambda \in \mathbb{R}$ donne le même vecteur $\mathbf{x}$ puisque les anomalies sont centrées ($\mathbf{X}_f \mathbf{1} = \mathbf{0}$).

Plutôt que d'effectuer une analyse pour chaque membre de l'ensemble, nous allons effectuer une unique analyse et engendrer de nouvelles anomalies centrées autour de cette analyse. On cherche $\mathbf{w}_a \in \mathbb{R}^m$ les coordonnées dans l'espace de l'ensemble de $\mathbf{x}_a$

$$\mathbf{x}_a = \bar{\mathbf{x}}_f + \mathbf{X}_f \mathbf{w}_a. \quad (2.39)$$

Si on réécrit l'équation (2.4), on obtient

$$\bar{\mathbf{x}}_f + \mathbf{X}_f \mathbf{w}_a = \bar{\mathbf{x}}_f + \mathbf{X}_f \mathbf{X}_f^T \mathbf{H}^T (\mathbf{H} \mathbf{X}_f \mathbf{X}_f^T \mathbf{H}^T + \mathbf{R})^{-1} \boldsymbol{\delta}, \quad (2.40)$$

avec $\boldsymbol{\delta} = \mathbf{y} - \mathcal{H}(\bar{\mathbf{x}}_f)$ qui donne

$$\mathbf{w}_a = \mathbf{X}_f^T \mathbf{H}^T (\mathbf{H} \mathbf{X}_f \mathbf{X}_f^T \mathbf{H}^T + \mathbf{R})^{-1} \boldsymbol{\delta} = \mathbf{Y}_f^T (\mathbf{Y}_f \mathbf{Y}_f^T + \mathbf{R})^{-1} \boldsymbol{\delta}. \quad (2.41)$$

Cette première solution est exprimée dans l'espace des observations, avec le recours à l'inversion d'une matrice de taille $\mathbb{R}^{p \times p}$. Cependant, on peut appliquer la formule de SMW (2.7) et ainsi obtenir

$$\mathbf{w}_a = (\mathbf{I}_m + \mathbf{Y}_f^T \mathbf{R}^{-1} \mathbf{Y}_f)^{-1} \mathbf{Y}_f^T \mathbf{R}^{-1} \boldsymbol{\delta}. \quad (2.42)$$

Cette formule est cette fois exprimée dans l'espace de l'ensemble, réduisant ainsi le coût de calcul de l'inversion de matrice dans le cas où $m \ll p$.

L'analyse étant obtenue, il reste encore à engendrer les nouvelles anomalies. On va pour cela factoriser $\mathbf{P}^a = \mathbf{X}_a \mathbf{X}_a^T$ à partir de la formulation escomptée (2.25)

$$\mathbf{P}^a = (\mathbf{I}_n - \mathbf{K}\mathbf{H}) \mathbf{P}^f \quad (2.43)$$

$$= (\mathbf{I}_n - \mathbf{X}_f \mathbf{Y}_f^T (\mathbf{Y}_f \mathbf{Y}_f^T + \mathbf{R})^{-1} \mathbf{H}) \mathbf{X}_f \mathbf{X}_f^T \quad (2.44)$$

$$= \mathbf{X}_f (\mathbf{I}_m - \mathbf{Y}_f^T (\mathbf{Y}_f \mathbf{Y}_f^T + \mathbf{R})^{-1} \mathbf{Y}_f) \mathbf{X}_f^T. \quad (2.45)$$

On choisit donc alors

$$\mathbf{X}_a = \mathbf{X}_f (\mathbf{I}_m - \mathbf{Y}_f^T (\mathbf{Y}_f \mathbf{Y}_f^T + \mathbf{R})^{-1} \mathbf{Y}_f)^{1/2} \mathbf{U}, \quad (2.46)$$

où $\mathbf{U}$ est une matrice orthogonale, vérifiant $\mathbf{U}\mathbf{1} = \mathbf{1}$ pour que les anomalies restent centrées. La racine carrée est définie comme l'unique racine carrée (semi-)définie positive d'une matrice symétrique (semi-)définie positive.

On peut simplifier cette expression en appliquant de nouveau la même transformation que pour passer de (2.41) à (2.42)

$$\mathbf{X}_a = \mathbf{X}_f(\mathbf{I}_m - \mathbf{Y}_f^T(\mathbf{Y}_f\mathbf{Y}_f^T + \mathbf{R})^{-1}\mathbf{Y}_f)^{\frac{1}{2}}\mathbf{U} \quad (2.47)$$

$$= \mathbf{X}_f(\mathbf{I}_m - (\mathbf{I}_m + \mathbf{Y}_f^T\mathbf{R}^{-1}\mathbf{Y}_f)^{-1}\mathbf{Y}_f^T\mathbf{R}^{-1}\mathbf{Y}_f)^{\frac{1}{2}}\mathbf{U} \quad (2.48)$$

$$= \mathbf{X}_f[(\mathbf{I}_m + \mathbf{Y}_f^T\mathbf{R}^{-1}\mathbf{Y}_f)^{-1}((\mathbf{I}_m + \mathbf{Y}_f^T\mathbf{R}^{-1}\mathbf{Y}_f) - \mathbf{Y}_f^T\mathbf{R}^{-1}\mathbf{Y}_f)]^{\frac{1}{2}}\mathbf{U} \quad (2.49)$$

$$= \mathbf{X}_f(\mathbf{I}_m + \mathbf{Y}_f^T\mathbf{R}^{-1}\mathbf{Y}_f)^{\frac{1}{2}}\mathbf{U}. \quad (2.50)$$

On remarque que l'on transforme les anomalies d'ébauche $\mathbf{X}_f$ en les multipliant par la droite par une matrice qui est donc dans le sous-espace de l'ensemble. Les anomalies d'analyse sont donc une combinaison linéaire des anomalies d'ébauche. Ainsi, une fois en possession de l'analyse et des anomalies analysées, nous avons à notre disposition l'ensemble analysé complet :

$$\mathbf{x}_a^i = \mathbf{x}_a + \sqrt{m-1}\mathbf{X}_f[(\mathbf{I}_m + \mathbf{Y}_f^T\mathbf{R}^{-1}\mathbf{Y}_f)^{\frac{1}{2}}\mathbf{U}]_i \quad (2.51)$$

$$= \bar{\mathbf{x}}_f + \mathbf{X}_f(\mathbf{w}_a + \sqrt{m-1}[(\mathbf{I}_m + \mathbf{Y}_f^T\mathbf{R}^{-1}\mathbf{Y}_f)^{\frac{1}{2}}\mathbf{U}]_i), \quad (2.52)$$

pour $i \in [1, m]$.

Étape de propagation :

L'étape de propagation est totalement identique à celle de l'EnKF stochastique.

Le pseudo-code correspondant à un cycle de l'EnKF déterministe, dans sa version ETKF, est explicité dans l'algorithme 3.

Algorithme 3 Algorithme pour un cycle de l'EnKF déterministe (ETKF).

Requis: : $m, \mathbf{E}_f = [\mathbf{x}_f^1, \dots, \mathbf{x}_f^m], \mathbf{y}, \mathbf{R}, \mathcal{M}, \mathcal{H}, \mathbf{U}$

- 1: $\bar{\mathbf{x}} = \mathbf{E}_f \mathbf{1}/m$
 - 2: $\mathbf{X} = (\mathbf{E}_f - \bar{\mathbf{x}}\mathbf{1}^T)/\sqrt{m-1}$
 - 3: $\mathbf{Z} = \mathcal{H}(\mathbf{E}_f)$
 - 4: $\bar{\mathbf{y}} = \mathbf{Z}\mathbf{1}/m$
 - 5: $\mathbf{S} = \mathbf{R}^{-\frac{1}{2}}(\mathbf{Z} - \bar{\mathbf{y}}\mathbf{1}^T)/\sqrt{(m-1)}$
 - 6: $\boldsymbol{\delta} = \mathbf{R}^{-\frac{1}{2}}(\mathbf{y} - \bar{\mathbf{y}})$
 - 7: $\mathbf{T} = (\mathbf{I}_m + \mathbf{S}\mathbf{S}^T)^{-1}$
 - 8: $\mathbf{w} = \mathbf{T}\mathbf{S}^T\boldsymbol{\delta}$
 - 9: $\mathbf{E}_a = \bar{\mathbf{x}}\mathbf{1}^T + \mathbf{X}(\mathbf{w}\mathbf{1}^T + \sqrt{m-1}\mathbf{T}^{\frac{1}{2}}\mathbf{U})$
 - 10: $\mathbf{E}_f = \mathcal{M}(\mathbf{E}_a)$
-

2.4.1.3 Formulation variationnelle de l'ETKF

On a montré dans la section 2.3.1 l'équivalence entre le 3D-Var et le BLUE. Il est nécessaire d'en faire de même avec l'ETKF et de définir la fonctionnelle qui correspond à cet algorithme.

On réécrit ici la fonctionnelle (2.12)

$$\mathcal{J}(\mathbf{x}) = \frac{1}{2}\|\mathbf{x} - \bar{\mathbf{x}}_f\|_{\mathbf{P}_f}^2 + \frac{1}{2}\|\mathbf{y} - \mathcal{H}(\mathbf{x})\|_{\mathbf{R}}^2. \quad (2.53)$$

Dans le contexte de l'ETKF, on écrit le vecteur d'état dans les coordonnées de l'ensemble selon (2.38). On note alors la fonction de la variable $\mathbf{w}$

$$\mathcal{L}(\mathbf{w}) = \mathcal{J}(\bar{\mathbf{x}}_f + \mathbf{X}_f\mathbf{w}). \quad (2.54)$$

Il nous faut pour cela définir le sens de $\mathbf{P}^f$ qui intervient dans ces fonctionnelles alors que pourtant $\mathbf{P}^f$ n'est pas de rang plein. Si on fait la décomposition en valeurs singulières de $\mathbf{X}_f = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^T$ avec $\mathbf{\Sigma} \in \mathbb{R}^{m' \times m'}$ la matrice diagonale des $m' \leq m$ valeurs singulières positives, $\mathbf{U} \in \mathbb{R}^{n \times m'}$ telle que $\mathbf{U}^T\mathbf{U} = \mathbf{I}_{m'}$ et $\mathbf{V} \in \mathbb{R}^{m \times m'}$ telle que $\mathbf{V}^T\mathbf{V} = \mathbf{I}_{m'}$, alors on peut noter l'inverse de Moore-Penrose $\mathbf{P}^{f\dagger} = \mathbf{U}\mathbf{\Sigma}^{-2}\mathbf{U}^T$. C'est cet inverse qui est utilisé dans la définition de la fonction de coût. La fonctionnelle s'écrit alors

$$\mathcal{L}(\mathbf{w}) = \frac{1}{2} \|\mathbf{X}_f \mathbf{w}\|_{\mathbf{P}^f}^2 + \frac{1}{2} \|\mathbf{y} - \mathcal{H}(\bar{\mathbf{x}}_f + \mathbf{X}_f \mathbf{w})\|_{\mathbf{R}}^2 \quad (2.55)$$

$$= \frac{1}{2} \mathbf{w}^T \mathbf{V} \mathbf{V}^T \mathbf{w} + \frac{1}{2} \|\mathbf{y} - \mathcal{H}(\bar{\mathbf{x}}_f + \mathbf{X}_f \mathbf{w})\|_{\mathbf{R}}^2 \quad (2.56)$$

Comme on l'a déjà remarqué précédemment, $\mathcal{L}(\mathbf{w}) = \mathcal{L}(\mathbf{w} + \lambda \mathbf{1})$, ce qui fait qu'il n'y a pas unicité du minimum de la fonction. Pour y remédier on ajoute alors un terme de jauge qui force la solution $\mathbf{w}$ sur le noyau de $\mathbf{X}_f$ et $\mathbf{V}$

$$\mathcal{L}^g(\mathbf{w}) = \frac{1}{2} \mathbf{w}^T (\mathbf{I}_m - \mathbf{V} \mathbf{V}^T) \mathbf{w}, \quad (2.57)$$

et la forme régularisée de la fonction de coût devient

$$\mathcal{L}(\mathbf{w}) = \frac{1}{2} \|\mathbf{y} - \mathcal{H}(\bar{\mathbf{x}}_f + \mathbf{X}_f \mathbf{w})\|_{\mathbf{R}}^2 + \frac{1}{2} \|\mathbf{w}\|^2. \quad (2.58)$$

Cette fonction de coût a bien le même minimum mais celui-ci est atteint en un unique $\mathbf{w}^*$ tel que $(\mathbf{I}_m - \mathbf{V} \mathbf{V}^T) \mathbf{w}^* = \mathbf{0}$.

Son gradient est alors

$$\nabla \mathcal{L}(\mathbf{w}) = -\mathbf{Y}^T \mathbf{R}^{-1} (\mathbf{y} - \mathcal{H}(\bar{\mathbf{x}}_f + \mathbf{X}_f \mathbf{w})) + \mathbf{w}, \quad (2.59)$$

et on peut approcher l'inverse de sa Hessienne par

$$\mathbf{D} = (\mathbf{I}_m + \mathbf{Y}^T \mathbf{R}^{-1} \mathbf{Y})^{-1}, \quad (2.60)$$

où $\mathbf{Y} = \mathbf{H} \mathbf{X}_f$ avec $\mathbf{H}$ le tangent linéaire de $\mathcal{H}$ en $\bar{\mathbf{x}}_f + \mathbf{X}_f \mathbf{w}$. Dans le cas d'une fonctionnelle quadratique, c'est-à-dire avec un opérateur d'observation linéaire, la solution est directement fournie par

$$\mathbf{w}^* = -\mathbf{D} \nabla \mathcal{L}(\mathbf{0}). \quad (2.61)$$

et on retrouve l'optimum de l'ETKF (2.42). Dans le cas contraire, on a recours à des méthodes itératives de minimisation. On peut enfin engendrer les anomalies a posteriori comme dans l'expression (2.50), en utilisant la racine carrée de la Hessienne

$$\mathbf{X}_a = \mathbf{X}_f \mathbf{D}^{\frac{1}{2}} \mathbf{U}, \quad (2.62)$$

avec $\mathbf{U}$ à nouveau une matrice orthogonale vérifiant $\mathbf{U} \mathbf{1} = \mathbf{1}$.

2.4.2 Variantes de l'EnKF

De nombreuses variantes des filtres de Kalman d'ensemble existent et nous ne tenons pas à les mentionner toutes, mais il est utile pour la suite que nous nous attardions sur certains algorithmes.

2.4.2.1 Un EnKF sans inversion de matrice : le DEnKF

Nous avons jusqu'ici qualifié de déterministe tout algorithme ne recourant pas à une perturbation stochastique des observations. Cependant, il existe un algorithme, décrit dans [Sakov et Oke \(2008\)](#), qui est appelé deterministic EnKF (DEnKF) par les auteurs, et qui diffère algébriquement de l'ETKF que nous avons décrit précédemment.

L'analyse dans le DEnKF est identique à celle de l'ETKF, mais la construction des anomalies et de l'ensemble autour de l'analyse est différente. Tout d'abord, il nous faut remarquer qu'il est possible d'engendrer les anomalies d'analyse dans l'espace des états, en multipliant $\mathbf{X}_f$ par la gauche. En effet, on souhaite avoir

$$\mathbf{P}^a = (\mathbf{I}_n - \mathbf{K}\mathbf{H})\mathbf{P}^f. \quad (2.63)$$

Or

$$\begin{aligned} (\mathbf{I}_n - \mathbf{K}\mathbf{H})\mathbf{P}^f &= \left[\mathbf{I}_n - \mathbf{P}^f \mathbf{H}^T (\mathbf{R} + \mathbf{H} \mathbf{P}^f \mathbf{H}^T)^{-1} \mathbf{H} \right] \mathbf{P}^f \\ &= \left[\mathbf{I}_n - \mathbf{P}^f \mathbf{H}^T (\mathbf{R} + \mathbf{H} \mathbf{P}^f \mathbf{H}^T)^{-1} \mathbf{H} \right]^{\frac{1}{2}} \left[\mathbf{I}_n - \mathbf{P}^f \mathbf{H}^T (\mathbf{R} + \mathbf{H} \mathbf{P}^f \mathbf{H}^T)^{-1} \mathbf{H} \right]^{\frac{1}{2}} \mathbf{P}^f \\ &= \left[\mathbf{I}_n - \mathbf{P}^f \mathbf{H}^T (\mathbf{R} + \mathbf{H} \mathbf{P}^f \mathbf{H}^T)^{-1} \mathbf{H} \right]^{\frac{1}{2}} \mathbf{P}^f \left[\mathbf{I}_n - \mathbf{H}^T (\mathbf{R} + \mathbf{H} \mathbf{P}^f \mathbf{H}^T)^{-1} \mathbf{H} \mathbf{P}^f \right]^{\frac{1}{2}} \\ &= \left[\mathbf{I}_n - \mathbf{P}^f \mathbf{H}^T (\mathbf{R} + \mathbf{H} \mathbf{P}^f \mathbf{H}^T)^{-1} \mathbf{H} \right]^{\frac{1}{2}} \mathbf{P}^f \left[\mathbf{I}_n - \mathbf{P}^f \mathbf{H}^T (\mathbf{R} + \mathbf{H} \mathbf{P}^f \mathbf{H}^T)^{-1} \mathbf{H} \right]^{\frac{T}{2}} \\ &= (\mathbf{I}_n - \mathbf{K}\mathbf{H})^{\frac{1}{2}} \mathbf{P}^f (\mathbf{I}_n - \mathbf{K}\mathbf{H})^{\frac{T}{2}}. \end{aligned} \quad (2.64)$$

Ainsi, on a l'expression suivante pour les anomalies analysées

$$\mathbf{X}_a = (\mathbf{I}_n - \mathbf{K}\mathbf{H})^{\frac{1}{2}} \mathbf{X}_f. \quad (2.65)$$

Intuitivement et sans rigueur, si on suppose $\mathbf{K}\mathbf{H}$ petit devant $\mathbf{I}_n$ (sans donner de définition exacte à petit), alors le développement limité au première ordre de l'équation (2.65) donne

$$\mathbf{X}_a = \left(\mathbf{I}_n - \frac{1}{2} \mathbf{K}\mathbf{H} \right) \mathbf{X}_f. \quad (2.66)$$

On évite ainsi le coup prohibitif du calcul de la racine carrée, tout en faisant une approximation qui peut s'avérer mineure. Par ailleurs, si on pousse le développement limité pour intégrer le terme de deuxième ordre, on obtient alors une matrice de covariance des erreurs d'analyse légèrement plus grande, comparé à l'utilisation de la racine carrée exacte. Ceci permet dans certaines situations également de s'abstenir de l'inflation (voir 2.4.5).

Nous nous reporterons à [Sakov et Oke \(2008\)](#) ou [Asch et al. \(2016a\)](#) pour une justification plus rigoureuse de ce schéma.

2.4.2.2 Un ETKF avec opérateur d'observation non-linéaire : le MLEF

Une première façon de traiter un opérateur d'observation non-linéaire est d'utiliser son tangent linéaire. Malheureusement, cela peut avoir un coup de calcul prohibitif. Il est possible d'utiliser l'ensemble pour calculer des différences finies afin d'approximer le tangent linéaire de l'opérateur d'observation. C'est ce qui est fait dans le maximum likelihood ensemble filter (MLEF) ([Zupanski, 2005](#); [Carrassi et al., 2009](#)).

En effet, si on utilise le formalisme variationnel de l'ETKF décrit dans la section 2.4.1.3, le problème consiste à calculer $\mathbf{Y} = \mathbf{H}_{|\mathbf{x}} \mathbf{X}_f$, qui apparaît dans (2.59), sans faire appel à $\mathbf{H}_{|\mathbf{x}}$ le tangent linéaire de $\mathcal{H}$ au point $\mathbf{x} = \bar{\mathbf{x}} + \mathbf{X}_f \mathbf{w}$. [Sakov et al. \(2012\)](#) présentent deux façons d'y parvenir.

Nous avons déjà vu que l'on peut à cette fin utiliser des différences finies avec l'équation (2.33). La première façon s'en inspire et réduit les perturbations par un facteur $0 < \varepsilon \ll 1$ avant d'appliquer l'opérateur non-linéaire $\mathcal{H}$ puis de retransformer par ε^{-1} . On a alors

$$\mathbf{Y} = \frac{1}{\varepsilon} \mathcal{H}(\mathbf{x}\mathbf{1}^T + \varepsilon \mathbf{X}_f) \left(\mathbf{I}_m - \frac{\mathbf{1}\mathbf{1}^T}{m} \right). \quad (2.67)$$

Cette méthode est nommée *bundle* par les auteurs.

L'autre méthode consiste à transformer l'ensemble en lui appliquant la matrice $\mathbf{T} = [\mathbf{I}_m + \mathbf{Y}^T \mathbf{R}^{-1} \mathbf{Y}]^{-\frac{1}{2}}$ issue de la précédente itération de minimisation, lui appliquer $\mathcal{H}$ et le retransformer en lui appliquant $\mathbf{T}^{-1}$. Cela donne

$$\mathbf{Y} = \mathcal{H}(\mathbf{x}\mathbf{1}^T + \mathbf{X}_f \mathbf{T} \sqrt{m-1}) \left(\mathbf{I}_m - \frac{\mathbf{1}\mathbf{1}^T}{m} \right) \mathbf{T}^{-1} / \sqrt{m-1}. \quad (2.68)$$

Par l'application de $\mathbf{T}$ sur les anomalies, on se place dans un domaine où $\mathcal{H}$ est proche du régime linéaire. Cette méthode est nommée *transform*.

On peut alors utiliser cette expression de $\mathbf{Y}$ dans le calcul du gradient et du Hessien, qui peuvent alors être utilisés par un algorithme de minimisation de (2.58) comme un Gauss-Newton, une méthode de type quasi-Newton, ou un Levenberg-Marquardt. On obtiendra alors un $\mathbf{w}^*$ optimal et on pourra terminer l'algorithme de l'ETKF en engendrant les anomalies d'analyse comme dans (2.62).

Le pseudo-code correspondant à un cycle du MLEF *bundle* ou *transform*, minimisé avec un algorithme de Gauss-Newton, est explicité dans l'algorithme 4.

2.4.3 Le lisseur de Kalman d'ensemble

Au fur et à mesure que les observations deviennent disponibles, elles sont assimilées par le filtre pour fournir le meilleur état possible à l'instant présent. Si on note t_1 le début de l'expérience d'assimilation de données, notons $\mathbf{y}_{L:1} = \mathbf{y}_L, \dots, \mathbf{y}_1$ l'intégralité des observations depuis le temps t_1 jusqu'au temps présent t_L . Le filtrage consiste ainsi à obtenir la fonction de densité de probabilité (pdf) de l'état $\mathbf{x}_L$ conditionné par les observations, à savoir $p(\mathbf{x}_L | \mathbf{y}_{L:1})$. On n'utilise donc que des observations passées ou présentes. Cette démarche est utilisée lorsque l'on souhaite faire de la prévision météorologique opérationnelle par exemple.

Le lissage en revanche considère la pdf $p(\mathbf{x}_{L:1} | \mathbf{y}_{L:1})$, où $\mathbf{x}_{L:1} = \mathbf{x}_L, \dots, \mathbf{x}_1$ est la trajectoire parcourue par l'état. On utilise donc des observations passées, présentes et potentiellement futures pour estimer l'état à un temps donné. Cette démarche est utilisée pour des réanalyses qui visent à obtenir la meilleure estimation de l'état étant donnée l'intégralité des informations disponibles.

On peut étendre les algorithmes de filtre de Kalman et de filtre de Kalman d'ensemble au cas du lissage. Parmi la multitude d'algorithmes disponibles, nous nous concentrerons sur l'ensemble Kalman smoother (EnKS) (Evensen et van Leeuwen, 2000), dont nous donnerons une formulation en termes déterministes plutôt que stochastiques.

L'algorithme tel qu'il est décrit par exemple dans Evensen (2003), dans un contexte stochastique, ou Cosme *et al.* (2010) dans un contexte déterministe, est simple et algorithmiquement peu coûteux, même s'il nécessite beaucoup de mémoire. Nous allons travailler par récurrence. On imagine que l'on a conservé, pour $l \in [t_1, t_{L-1}]$ le résultat de l'analyse $\mathbf{x}_l$ et ses anomalies $\mathbf{X}_l$ qui représentent la moyenne et la dispersion des pdf gaussiennes $p(\mathbf{x}_l | \mathbf{y}_{L-1:1})$. Nous sommes désormais au temps t_L

Algorithme 4 Algorithme pour un cycle du MLEF, en version *bundle* ou *transform*, avec minimisation de Gauss-Newton.

Requis : $m, \mathbf{E}_f = [\mathbf{x}_f^1, \dots, \mathbf{x}_f^m], \mathbf{y}, \mathbf{R}, \mathcal{M}, \mathcal{H}, e, j_{\max}, \varepsilon, \mathbf{U}$

- 1: $\bar{\mathbf{x}} = \mathbf{E}_f \mathbf{1} / m$
 - 2: $\mathbf{X} = (\mathbf{E}_f - \bar{\mathbf{x}} \mathbf{1}^T) / \sqrt{m-1}$
 - 3: Transform : $\mathbf{T} = \mathbf{I}_m$
 - 4: $j = 0, \mathbf{w} = \mathbf{0}$
 - 5: **repeat**
 - 6: $\mathbf{x} = \bar{\mathbf{x}} + \mathbf{X} \mathbf{w}$
 - 7: Bundle : $\mathbf{E} = \mathbf{x} \mathbf{1}^T + \varepsilon \mathbf{X}$
 - 8: Transform : $\mathbf{E} = \mathbf{x} \mathbf{1}^T + \mathbf{X} \mathbf{T} \sqrt{m-1}$
 - 9: $\mathbf{Z} = \mathcal{H}(\mathbf{E})$
 - 10: $\bar{\mathbf{y}} = \mathbf{Z} \mathbf{1} / m$
 - 11: Bundle : $\mathbf{Y} = (\mathbf{Z} - \bar{\mathbf{y}} \mathbf{1}^T) / \varepsilon$
 - 12: Transform : $\mathbf{Y} = (\mathbf{Z} - \bar{\mathbf{y}} \mathbf{1}^T) \mathbf{T}^{-1} / \sqrt{m-1}$
 - 13: $\nabla \mathcal{J} = \mathbf{w} - \mathbf{Y}^T \mathbf{R}^{-1} (\mathbf{y} - \bar{\mathbf{y}})$
 - 14: $\mathbf{D} = [\mathbf{I}_m + \mathbf{Y}^T \mathbf{R}^{-1} \mathbf{Y}]^{-1}$
 - 15: $\Delta \mathbf{w} = \mathbf{D} \nabla \mathcal{J}$
 - 16: $\mathbf{w} := \mathbf{w} - \Delta \mathbf{w}$
 - 17: Transform : $\mathbf{T} = \mathbf{D}^{\frac{1}{2}}$
 - 18: $j := j + 1$
 - 19: **until** $j \geq j_{\max}$ **or** $\|\Delta \mathbf{w}\| < e$
 - 20: Bundle : $\mathbf{T} = \mathbf{D}^{\frac{1}{2}}$
 - 21: $\mathbf{E}_a = \mathbf{x} \mathbf{1}^T + \sqrt{m-1} \mathbf{X} \mathbf{T} \mathbf{U}$
 - 22: $\mathbf{E}_f = \mathcal{M}(\mathbf{E}_a)$
-

et on applique l'ETKF. L'analyse est alors

$$\mathbf{x}_L = \bar{\mathbf{x}}_{f,L} + \mathbf{X}_{f,L}(\mathbf{I}_m + \mathbf{Y}_f^T \mathbf{R}^{-1} \mathbf{Y}_f)^{-1} \mathbf{Y}_f^T \mathbf{R}^{-1} \boldsymbol{\delta} \quad (2.69)$$

$$\mathbf{X}_L = \mathbf{X}_{f,L}(\mathbf{I}_m + \mathbf{Y}_f^T \mathbf{R}^{-1} \mathbf{Y}_f)^{\frac{1}{2}} \mathbf{U}. \quad (2.70)$$

On propage cette analyse dans le passé de sorte que pour tout $l \in [t_1, t_{L-1}]$, on remplace $\mathbf{x}_l$ et $\mathbf{X}_l$ par

$$\mathbf{x}_l \leftarrow \mathbf{x}_l + \mathbf{X}_l(\mathbf{I}_m + \mathbf{Y}_f^T \mathbf{R}^{-1} \mathbf{Y}_f)^{-1} \mathbf{Y}_f^T \mathbf{R}^{-1} \boldsymbol{\delta} \quad (2.71)$$

$$\mathbf{X}_l \leftarrow \mathbf{X}_l(\mathbf{I}_m + \mathbf{Y}_f^T \mathbf{R}^{-1} \mathbf{Y}_f)^{\frac{1}{2}} \mathbf{U}. \quad (2.72)$$

On applique donc la même transformation à l'analyse et aux perturbations dans le passée que celle appliquée au temps présent. Ces nouvelles analyses et anomalies représentent désormais la pdf $p(\mathbf{x}_l | \mathbf{y}_{L:1})$. L'initialisation de la récurrence correspond à la passe avant d'un ETKF qui fournit une estimation de $p(\mathbf{x}_L | \mathbf{y}_{L:1})$.

Pour justifier cet algorithme, commençons par déterminer $p(\mathbf{x}_{L-1} | \mathbf{y}_{L:1})$. On sait que l'ETKF nous fournit une estimation de $p(\mathbf{x}_L | \mathbf{y}_{L:1})$. La formule de Bayes nous dit que

$$p(\mathbf{x}_{L-1} | \mathbf{y}_{L:1}) \propto p(\mathbf{y}_L | \mathbf{x}_{L-1}, \mathbf{y}_{L-1:1}) p(\mathbf{x}_{L-1} | \mathbf{y}_{L-1:1}). \quad (2.73)$$

La pdf $p(\mathbf{x}_{L-1} | \mathbf{y}_{L-1:1})$ est estimée par le résultat de l'ETKF, et est supposée gaussienne, centrée en $\mathbf{x}_{a,L-1}$ et de covariance $\mathbf{P}_{L-1}^a = \mathbf{X}_{a,L-1} \mathbf{X}_{a,L-1}^T$. Si on écrit $\mathbf{x}_{L-1}$ avec les coordonnées dans l'espace de l'ensemble $\mathbf{w}$, on peut dire que cette pdf est proportionnelle à $\exp(-\frac{1}{2} \|\mathbf{w}\|^2)$.

En ce qui concerne $p(\mathbf{y}_L | \mathbf{x}_{L-1}, \mathbf{y}_{L-1:1})$, on a

$$p(\mathbf{y}_L | \mathbf{x}_{L-1}, \mathbf{y}_{L-1:1}) \propto \exp\left\{-\frac{1}{2} \|\mathbf{y}_L - \mathcal{H}_L \circ \mathcal{M}_{L \leftarrow L-1}(\mathbf{x}_{a,L-1} + \mathbf{X}_{a,L-1} \mathbf{w})\|_{\mathbf{R}_L}^2\right\}. \quad (2.74)$$

La fonction de coût associée à la pdf $p(\mathbf{x}_{L-1} | \mathbf{y}_{L:1})$ est donc

$$\mathcal{J}(\mathbf{w}) = \frac{1}{2} (\|\mathbf{y}_L - \mathcal{H}_L \circ \mathcal{M}_{L \leftarrow L-1}(\mathbf{x}_{a,L-1} + \mathbf{X}_{a,L-1} \mathbf{w})\|_{\mathbf{R}_L}^2 + \|\mathbf{w}\|^2). \quad (2.75)$$

Si on la linéarise autour de $\mathbf{x}_{a,L-1}$, on obtient

$$\mathcal{J}(\mathbf{w}) = \frac{1}{2} (\|\mathbf{y}_L - \mathcal{H}_L \circ \mathcal{M}_{L \leftarrow L-1}(\mathbf{x}_{a,L-1} + \mathbf{X}_{a,L-1} \mathbf{w})\|_{\mathbf{R}_L}^2 + \|\mathbf{w}\|^2) \quad (2.76)$$

$$\simeq \frac{1}{2} (\|\mathbf{y}_L - \mathcal{H}_L \circ \mathcal{M}_{L \leftarrow L-1}(\mathbf{x}_{a,L-1}) + \mathbf{H}_L \mathbf{M}_{L \leftarrow L-1} \mathbf{X}_{a,L-1} \mathbf{w}\|_{\mathbf{R}_L}^2 + \|\mathbf{w}\|^2) \quad (2.77)$$

$$\simeq \frac{1}{2} (\|\mathbf{y}_L - \mathcal{H}_L(\mathbf{x}_{f,L}) + \mathbf{H}_L \mathbf{X}_{f,L} \mathbf{w}\|_{\mathbf{R}_L}^2 + \|\mathbf{w}\|^2) \quad (2.78)$$

$$\simeq \frac{1}{2} (\|\boldsymbol{\delta}_L + \mathbf{Y}_L \mathbf{w}\|_{\mathbf{R}_L}^2 + \|\mathbf{w}\|^2), \quad (2.79)$$

dont on dérive naturellement l'analyse présentée en (2.71). La dérivation pour les états $\mathbf{x}_l$ avec $1 \leq l < L-1$ est identique à celle-ci et récursive.

On note que dans le cas d'un modèle parfait et linéaire, la repropagation jusqu'à t_L par le modèle de l'état analysé à un temps t_l redonne la valeur de l'ETKF, c'est-à-dire

$$\mathbf{M}_{L \leftarrow l} \mathbf{x}_{a,l} = \mathbf{x}_{a,L}. \quad (2.80)$$

Cette propriété fait que l'EnKS n'est qu'un lisseur et ne fournit pas de résultat de filtrage différent de celui de l'ETKF.

Nous avons dans ce paragraphe présenté la version dite *fixed-interval* du lisseur de Kalman, à savoir que l'ensemble des états de t_1 à t_L est analysé à chaque nouvelle observation. On peut se contenter de ne mettre à jour que les K derniers états, de t_{L-K} à t_L . L'algorithme est alors qualifié de *fixed-lag*. Il est également possible de concatener l'ensemble des états et des observations de t_1 à t_L et d'effectuer une unique analyse avec ces états augmentés, avec par conséquent des matrices d'erreur d'ébauche et d'observations également augmentées. Cette démarche est présentée dans [van Leeuwen et Evensen \(1996\)](#) en temps qu'*ensemble smoother* et reprise dans les algorithmes de 4D-ETKF ([Hunt et al., 2004, 2007](#)). Enfin, on peut aussi se concentrer sur l'estimation de $p(\mathbf{x}_k | \mathbf{y}_{L:1})$, avec $k < L$, au lieu de $p(\mathbf{x}_{L:1} | \mathbf{y}_{L:1})$. Cette démarche est qualifiée de *marginal smoothing*, en opposition au *joint smoothing* présenté jusqu'ici. Ces variantes du lisseur de Kalman sont résumées dans [Cosme et al. \(2012\)](#).

Le pseudo-code correspondant au cycle du *fixed-interval* EnKS au temps t_L est explicité dans l'algorithme 5.

Algorithme 5 Algorithme pour un cycle de l'EnKS au temps t_L .

Requis: : $m, L, \mathbf{E}_f^L, \{\mathbf{E}_s^l\}_{l \in [1, L-1]}, \mathbf{y}, \mathbf{R}, \mathcal{M}, \mathcal{H}, \mathbf{U}$

- 1: $\mathbf{E} = \mathbf{E}_f^L$
 - 2: $\bar{\mathbf{x}} = \mathbf{E}\mathbf{1}/m$
 - 3: $\mathbf{X} = (\mathbf{E} - \bar{\mathbf{x}}\mathbf{1}^T)/\sqrt{m-1}$
 - 4: $\mathbf{Z} = \mathcal{H}(\mathbf{E})$
 - 5: $\bar{\mathbf{y}} = \mathbf{Z}\mathbf{1}/m$
 - 6: $\mathbf{S} = \mathbf{R}^{-\frac{1}{2}}(\mathbf{Z} - \bar{\mathbf{y}}\mathbf{1}^T)/\sqrt{m-1}$
 - 7: $\boldsymbol{\delta} = \mathbf{R}^{-\frac{1}{2}}(\mathbf{y} - \bar{\mathbf{y}})$
 - 8: $\mathbf{T} = (\mathbf{I}_m + \mathbf{S}\mathbf{S}^T)^{-1}$
 - 9: $\mathbf{w} = \mathbf{T}\mathbf{S}^T\boldsymbol{\delta}$
 - 10: $\mathbf{E}_s^L = \bar{\mathbf{x}}\mathbf{1}^T + \mathbf{X}(\mathbf{w}\mathbf{1}^T + \sqrt{m-1}\mathbf{T}^{\frac{1}{2}}\mathbf{U})$
 - 11: **for** $l = 1 \dots L-1$ **do**
 - 12: $\mathbf{E} = \mathbf{E}_s^l$
 - 13: $\bar{\mathbf{x}} = \mathbf{E}\mathbf{1}/m$
 - 14: $\mathbf{X} = (\mathbf{E} - \bar{\mathbf{x}}\mathbf{1}^T)/\sqrt{m-1}$
 - 15: $\mathbf{E}_s^l := \bar{\mathbf{x}}\mathbf{1}^T + \mathbf{X}(\mathbf{w}\mathbf{1}^T + \sqrt{m-1}\mathbf{T}^{\frac{1}{2}}\mathbf{U})$
 - 16: **end for**
 - 17: $\mathbf{E}_f^{L+1} = \mathcal{M}(\mathbf{E}_s^L)$
-

2.4.4 Localisation

Le nombre de membres de l'ensemble étant limité ($m \ll n$), les matrices estimées par les équations de la forme de (2.21) ne sont que des approximations et ne sont pas de rang plein. Il peut alors apparaître dans la matrice $\mathbf{P}^f$ des covariances entre variables qui n'ont pas lieu d'être. La localisation consiste à faire l'hypothèse que les covariances tendent vers 0 quand deux variables sont distantes en espace ou en temps et ainsi transformer les matrices de covariance des erreurs d'ébauche ou d'observation pour faire disparaître les covariances erronées. Pour ce faire, on peut soit réduire les valeurs des coefficients loin de la diagonale de la matrice $\mathbf{P}^f$, soit faire une analyse autour de chaque point de grille, en ne gardant que les observations et variables de l'état les plus proches de ce point de grille. Ces deux méthodes s'appellent respectivement localisation des covariances et localisation par domaine.

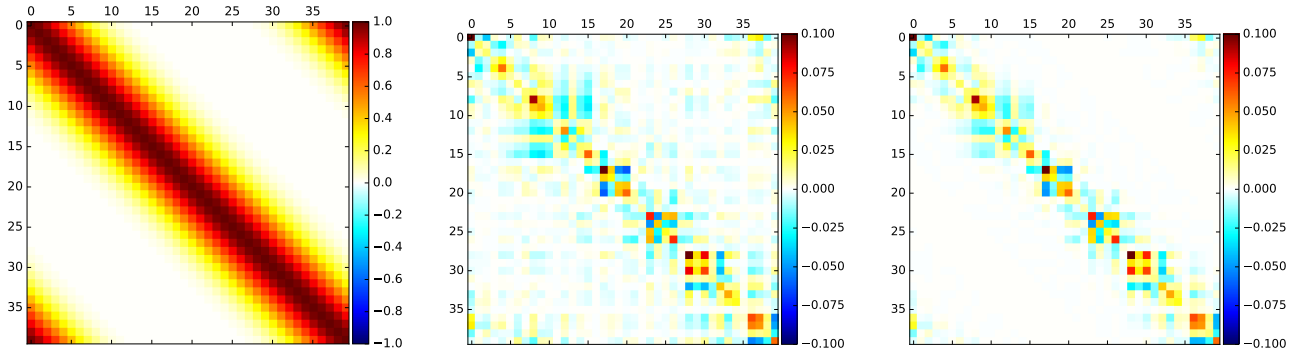


FIGURE 2.2 – À gauche, la matrice ρ . Au milieu, la matrice $\mathbf{P}^f$ avant régularisation, dont on voit les corrélations bruitées loin de la diagonale. À droite, la nouvelle matrice après application de la localisation.

2.4.4.1 Localisation des covariances

La localisation des covariances (ou localisation par produit de Schur) modifie la matrice $\mathbf{P}^f$ pour lisser les corrélations erronées loin de la diagonale. Pour ce faire, on fait un produit de Schur (ou de Hadamard) de $\mathbf{P}^f$ avec une matrice de corrélation $\rho \in \mathbb{R}^{n \times n}$

$$[\rho \circ \mathbf{P}^f]_{i,j} = [\mathbf{P}^f]_{i,j} [\rho]_{i,j}. \quad (2.81)$$

Ainsi, si ρ est telle que les corrélations décroissent avec la distance, on aura filtré les corrélations de longue distance erronées dans $\mathbf{P}^f$. On rappelle que le théorème de Schur garantit que si ρ et $\mathbf{P}^f$ sont définies positives, alors $[\rho \circ \mathbf{P}^f]$ est également définie positive. L'action du produit de Schur sur les covariances peut-être visualisée sur la Fig. 2.2. On se place dans le cas d'un domaine unidimensionnel de taille 40.

2.4.4.2 Localisation par domaine

Avec la localisation par domaine (ou analyse locale), on effectue une analyse pour chaque point de grille du domaine au lieu d'une analyse globale. On ne garde pour ce faire que les observations situées dans un rayon, dit rayon de localisation, autour du point de grille où a lieu l'analyse. Si ce rayon est trop grand, on revient sur une situation proche de l'analyse globale, et si ce rayon est trop petit, on risque de perdre des corrélations de courtes ou moyennes distances, ce qui peut détériorer les performances.

Le coût de calcul de cette méthode peut paraître prohibitif vu le grand nombre d'analyses requises pour de grands domaines, mais la taille réduite de chacun des sous-systèmes ainsi que la possibilité de paralléliser cette tâche rend cette méthode viable. On pourra se reporter à [Hunt et al. \(2007\)](#) pour l'algorithme précis du localized ensemble transform Kalman filter (LETKF) et sa complexité algorithmique.

Une première façon simple de sélectionner les observations est de définir un cercle autour du point d'analyse et de tronquer la matrice $\mathbf{R}$ pour ne garder que les termes correspondant aux observations dans ce cercle. Cette troncature est qualifiée de *boxcar*.

Cependant, deux analyses à des points proches vont partager certaines observations et en inclure ou exclure d'autres, provoquant des discontinuités dans la transition. La deuxième façon de procéder remédie à ce problème en appliquant une fonction de transition à support compact continue qui

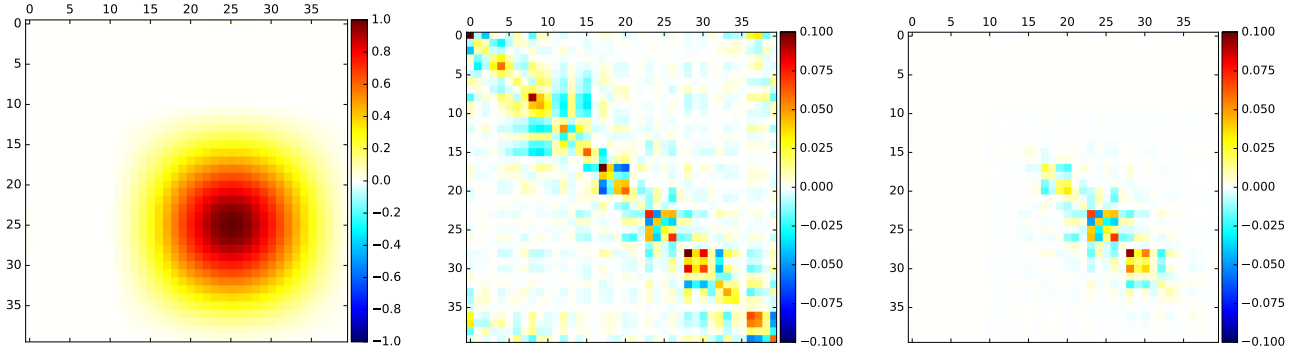


FIGURE 2.3 – À gauche, la fonction de Gaspari-Cohn centrée autour du point de grille 25. Au milieu, la matrice $\mathbf{P}^f$ avant régularisation, dont on voit les corrélations bruitées loin de la diagonale. À droite, la nouvelle matrice après application de la localisation. Cette matrice ne sera utilisée que pour l’analyse locale autour du point de grille 25.

tend vers 0 quand la distance croît. La fonction d’ordre 5 de Gaspari-Cohn (GC) (Gaspari et Cohn, 1999) est fréquemment utilisée

$$G(r) = \begin{cases} 1 - \frac{5}{3}r^2 + \frac{5}{8}r^3 + \frac{1}{2}r^4 - \frac{1}{4}r^5 & \text{si } 0 \leq r \leq 1 \\ 4 - 5r + \frac{5}{3}r^2 + \frac{5}{8}r^3 - \frac{1}{2}r^4 + \frac{1}{12}r^5 - \frac{2}{3r} & \text{si } 1 \leq r \leq 2 \\ 0 & \text{si } r \geq 2. \end{cases} \quad (2.82)$$

Si on note d la distance entre une observation et le point de grille où a lieu l’analyse, on appliquera la fonction de GC avec $\frac{d}{c}$ où c est le demi rayon de localisation. On notera par ailleurs qu’utiliser G pour localiser la matrice d’erreur d’observation $\mathbf{R}$ revient à transformer les anomalies observées et les innovations par $G^{-\frac{1}{2}}$ (Sakov et Bertino, 2011). L’action de cette transformation sur les covariances peut-être visualisée sur la Fig. 2.3. Le cas illustré est le même que celui présenté dans la localisation de Schur. On note que si c’est deux méthodes de localisation sont équivalentes d’un point de vue théorique, il s’avère que la localisation par domaine est inadaptée à l’assimilation d’observations intégrales, comme les observations satellites dépendant de toute la colonne atmosphérique, les observations des données GPS sol et radio-occultation et des données radar.

2.4.5 Inflation

Le nombre limité de membres de l’ensemble pour engendrer les matrices de covariances amène toujours des erreurs d’échantillonnage, qui peuvent s’accumuler au fur et à mesure des cycles d’assimilation. Ces erreurs peuvent aller jusqu’à déclencher la divergence du filtre. Pour y remédier, on accroît empiriquement les matrices de covariances.

2.4.5.1 Inflation multiplicative

On peut accroître les matrices de covariance, avant ou après l’analyse, en les multipliant par un facteur $\lambda^2 \geq 1$. On appelle cela l’inflation multiplicative.

$$\mathbf{P}^f \leftarrow \lambda^2 \mathbf{P}^f. \quad (2.83)$$

On peut également de manière équivalente disperser l’ensemble en augmentant ses anomalies

$$\mathbf{X}_f \leftarrow \lambda \mathbf{X}_f. \quad (2.84)$$

Dans le cas d'un modèle parfait mais non-linéaire, les erreurs d'échantillonnage sont une conséquence indirecte de la non-linéarité du modèle lors de la propagation de l'ensemble (Bocquet, 2011; Whitaker et Hamill, 2012; Bocquet *et al.*, 2015). L'inflation multiplicative pallie alors une erreur intrinsèque à l'EnKF.

La valeur de l'inflation dépend du type d'application utilisé, mais elle varie également en fonction de l'espace et du temps. C'est un paramètre qui doit être optimisé, ce qui est très coûteux. Ainsi, de multiples algorithmes ont été développés pour avoir une inflation adaptative. Parmi eux, nous nous attarderons sur l'algorithme du *finite-size* (Bocquet, 2011; Bocquet et Sakov, 2012) que nous avons appliqué durant la thèse.

2.4.5.2 L'EnKF-N, un exemple d'inflation adaptative

L'inflation, comme la localisation, est une solution pour pallier le fait que la taille de l'ensemble étant réduite, des erreurs d'échantillonnage apparaissent, notamment en raison de la propagation au cours des cycles d'assimilation par le modèle non-linéaire. Ainsi, les estimations de $\mathbf{x}_f$ et $\mathbf{P}^f$ à partir de l'ensemble $\mathbf{E}$, via $\bar{\mathbf{x}}$ et $\mathbf{P} = \mathbf{X}_f \mathbf{X}_f^T$, sont erronées. Lors de l'assimilation d'une observation $\mathbf{y}$, nous cherchons habituellement à estimer la pdf $p(\mathbf{x}|\mathbf{y}, \boldsymbol{\theta})$, où $\boldsymbol{\theta}$ représente des paramètres tels que $\mathbf{x}_f$ ou $\mathbf{P}^f$. Mais si on suppose que ces paramètres sont méconnus et sont eux-mêmes des variables aléatoires, on doit alors chercher à estimer $p(\mathbf{x}, \boldsymbol{\theta}|\mathbf{y})$.

La formule de Bayes nous dit que

$$p(\mathbf{x}, \boldsymbol{\theta}|\mathbf{y}) \propto p(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta})p(\mathbf{x}, \boldsymbol{\theta}). \quad (2.85)$$

Si on suppose que $\mathbf{y}$ ne dépend de $\boldsymbol{\theta}$ qu'à travers $\mathbf{x}$, et qu'on a une distribution a priori $p(\boldsymbol{\theta})$, on peut faire un choix hiérarchique et écrire

$$p(\mathbf{x}, \boldsymbol{\theta}|\mathbf{y}) \propto p(\mathbf{y}|\mathbf{x})p(\mathbf{x}|\boldsymbol{\theta})p(\boldsymbol{\theta}). \quad (2.86)$$

La distribution $p(\mathbf{x}|\boldsymbol{\theta})$ est toujours notre connaissance sur l'ébauche, et on a désormais ce qui est nommé un *hyperprior* $p(\boldsymbol{\theta})$. L'EnKF classique présuppose que la connaissance que nous donne l'ensemble $\mathbf{E}$ sur les moments d'ordre 1 et 2 de l'erreur d'ébauche est parfaite, c'est-à-dire $p(\boldsymbol{\theta}) = p(\mathbf{x}_f, \mathbf{P}^f) = \delta(\mathbf{P}^f - \mathbf{P})\delta(\mathbf{x}_f - \bar{\mathbf{x}})$, où δ est la distribution de Dirac. On va ici ne pas faire de telle hypothèse et marginaliser la pdf $p(\mathbf{x}|\boldsymbol{\theta})$ sur tous les $\boldsymbol{\theta}$ possibles, tels que $\mathbf{P}^f$ soit semi-définie positive. La distribution de l'ébauche donne alors

$$p(\mathbf{x}|\mathbf{E}) = \int d\boldsymbol{\theta} p(\mathbf{x}|\mathbf{E}, \boldsymbol{\theta})p(\boldsymbol{\theta}|\mathbf{E}). \quad (2.87)$$

$p(\mathbf{x}|\mathbf{E}, \boldsymbol{\theta})$ ne dépend pas de $\mathbf{E}$ puisqu'il dépend déjà des vraies statistiques, et donc

$$p(\mathbf{x}|\mathbf{E}) = \int d\boldsymbol{\theta} p(\mathbf{x}|\boldsymbol{\theta})p(\boldsymbol{\theta}|\mathbf{E}). \quad (2.88)$$

En appliquant la formule de Bayes sur $p(\boldsymbol{\theta}|\mathbf{E})$, on a

$$p(\mathbf{x}|\mathbf{E}) = \frac{1}{p(\mathbf{E})} \int d\boldsymbol{\theta} p(\mathbf{x}|\boldsymbol{\theta})p(\mathbf{E}|\boldsymbol{\theta})p(\boldsymbol{\theta}). \quad (2.89)$$

On a donc fait apparaître l'*hyperprior* dans la formule. Il faut que l'on fasse un choix pour cette distribution, sachant qu'elle correspond à notre connaissance a priori sur les erreurs d'ébauche avant de commencer toute assimilation. Dans Bocquet (2011), une distribution de Jeffrey est choisie :

$$p_J(\boldsymbol{\theta}) = p_J(\mathbf{x}_f, \mathbf{P}^f) \propto |\mathbf{P}^f|^{-\frac{n+1}{2}}, \quad (2.90)$$

avec $|\mathbf{P}^f|$ le déterminant de la matrice $\mathbf{P}^f$, de dimension $n \times n$. Avec cet *hyperprior*, il trouve

$$p(\mathbf{x}|\mathbf{E}) \propto \exp \left\{ 1 + \frac{(\mathbf{x} - \bar{\mathbf{x}})^T (\varepsilon_m \mathbf{P})^\dagger (\mathbf{x} - \bar{\mathbf{x}})}{m-1} \right\}^{-\frac{m}{2}}, \quad (2.91)$$

pour $\mathbf{x}$ dans le sous-espace de l'ensemble, $p(\mathbf{x}|\mathbf{E}) = 0$ sinon, avec $\varepsilon_m = 1 + 1/m$ et $\mathbf{P}^\dagger$ l'inverse de Moore-Penrose de $\mathbf{P}$. On remarque que ε_m a ici un rôle d'inflation multiplicative, et provient de l'incertitude sur $\mathbf{x}_f$. Cette expression est à comparer à

$$p(\mathbf{x}|\mathbf{E}) \propto \exp \left\{ -\frac{1}{2} (\mathbf{x} - \bar{\mathbf{x}})^T (\mathbf{P})^\dagger (\mathbf{x} - \bar{\mathbf{x}}) \right\} \quad (2.92)$$

pour $\mathbf{x}$ dans le sous-espace de l'ensemble, $p(\mathbf{x}|\mathbf{E}) = 0$ sinon, obtenue dans le cas habituel de l'EnKF.

Cette expression (2.91) nous permet d'écrire une nouvelle fonction de coût pour notre ETKF, similaire à (2.58), mais où le terme d'ébauche a été modifié en fonction (Bocquet *et al.*, 2015) :

$$\mathcal{L}(\mathbf{w}) = \frac{1}{2} \|\mathbf{y} - \mathcal{H}(\bar{\mathbf{x}}_f + \mathbf{X}_f \mathbf{w})\|_{\mathbf{R}}^2 + \frac{1}{2} (m+1) \ln \left(\varepsilon_m + \frac{\|\mathbf{w}\|^2}{m-1} \right). \quad (2.93)$$

Le schéma qui est issu de la résolution de cette fonction de coût a été nommé *finite-size* EnKF, ou EnKF-N.

Il a été montré, sur des modèles jouets et considérés comme parfaits, que ce schéma ne requiert pas d'inflation multiplicative et donne des résultats équivalents à ceux obtenus avec un ETKF dont l'inflation multiplicative est optimisée empiriquement.

2.4.5.3 Inflation additive et erreur modèle

On peut également utiliser l'inflation pour corriger une erreur modèle, indépendante du schéma de l'EnKF. On choisit alors souvent de l'inflation additive, sur les matrices de covariances

$$\mathbf{P}^f \leftarrow \mathbf{P}^f + \mathbf{Q}. \quad (2.94)$$

La façon ensembliste et stochastique pour appliquer cette inflation additive est d'ajouter, à chaque membre de l'ensemble, une perturbation corrélée selon $\mathbf{Q}$:

$$\mathbf{x}_f^i \leftarrow \mathbf{x}_f^i + \mathbf{Q}^{\frac{1}{2}} \boldsymbol{\varepsilon}_i, \quad \text{avec } \boldsymbol{\varepsilon}_i \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_n), \quad \text{et pour } i \in [1, m]. \quad (2.95)$$

On peut imposer que l'ensemble des perturbations $\mathbf{D} = [\mathbf{Q}^{\frac{1}{2}} \boldsymbol{\varepsilon}_1, \dots, \mathbf{Q}^{\frac{1}{2}} \boldsymbol{\varepsilon}_m]$ soit de moyenne nulle. Ceci réduit la variance qui peut être alors ajustée en multipliant par $\sqrt{m/(m-1)}$. Les nouvelles erreurs d'ébauche sont alors, en espérance,

$$\mathbf{P}_f = \mathbf{X}_f \mathbf{X}_f^T + \mathbf{Q} \quad (2.96)$$

$$+ (\bar{\mathbf{Q}} - \mathbf{Q}) + \frac{1}{m-1} (\mathbf{X}_f \mathbf{D}^T + \mathbf{D} \mathbf{X}_f^T), \quad (2.97)$$

avec $\bar{\mathbf{Q}} = \frac{1}{m-1} \mathbf{D} \mathbf{D}^T$. Ainsi, cette formulation ne donne pas parfaitement la matrice de covariance des erreurs d'ébauche escomptée car elle contient un bruit stochastique sous la forme de la ligne (2.97).

Cependant, un tel écart est inévitable car le problème n'est pas bien posé depuis le départ. En effet, il s'agit de constater que $\mathbf{Q}$ pouvant être de rang plein, il est impossible de la représenter complètement avec les seuls m membres de l'ensemble. En revanche, on peut se concentrer sur

la projection de $\mathbf{Q}$ sur l'espace de l'ensemble. Suivant cette idée, une méthode déterministe de construction des perturbations est présentée dans [Raanes et al. \(2015\)](#) et nommée SQRT-CORE.

Notons $\mathbf{X}_f^\dagger$ l'inverse de Moore-Penrose de $\mathbf{X}_f$, et $\mathbf{\Pi} = \mathbf{X}_f \mathbf{X}_f^\dagger$ le projecteur orthogonal sur l'espace des colonnes de $\mathbf{X}_f$. Si on applique cette projection à $\mathbf{Q}$ pour former $\hat{\mathbf{Q}} = \mathbf{\Pi} \mathbf{Q} \mathbf{\Pi}$, on peut alors essayer d'atteindre une nouvelle erreur d'ébauche suivant

$$\mathbf{P}^f = \mathbf{X}_f \mathbf{X}_f^T + (m-1) \hat{\mathbf{Q}} \quad (2.98)$$

$$= \mathbf{X}_f \mathbf{X}_f^T + (m-1) \mathbf{\Pi} \mathbf{Q} \mathbf{\Pi} \quad (2.99)$$

$$= \mathbf{X}_f \mathbf{X}_f^T + (m-1) \mathbf{X}_f \mathbf{X}_f^\dagger \mathbf{Q} \mathbf{X}_f^{\dagger T} \mathbf{X}_f^T \quad (2.100)$$

$$= \mathbf{X}_f (\mathbf{I}_m + (m-1) \mathbf{X}_f^\dagger \mathbf{Q} \mathbf{X}_f^{\dagger T}) \mathbf{X}_f^T. \quad (2.101)$$

Ainsi, on peut définir une nouvelle anomalie

$$\mathbf{X}_f \leftarrow \mathbf{X}_f (\mathbf{I}_m + (m-1) \mathbf{X}_f^\dagger \mathbf{Q} \mathbf{X}_f^{\dagger T})^{\frac{1}{2}}, \quad (2.102)$$

qui permet d'obtenir l'erreur d'ébauche (2.98).

Comparée à l'erreur d'ébauche de (2.94), il manque un terme résiduel dans l'espace orthogonal à celui des anomalies, qui n'est pas atteint par cette méthode. [Raanes et al. \(2015\)](#) considère des algorithmes plus avancés pour mieux incorporer ce terme résiduel.

2.5 Méthodes variationnelles d'ensemble

Nous avons vu dans les sections précédentes les algorithmes variationnels, de type 4D-Var, et d'ensemble, de type EnKF. L'avantage des algorithmes de type EnKF est que grâce à un ensemble de taille limité, on peut définir des matrices d'erreur qui évoluent avec le temps pour un coût de calcul raisonnable. On n'a par ailleurs pas besoin du tangent linéaire et de l'adjoint de l'opérateur d'observation car l'ensemble peut permettre, par différences finies, de l'estimer. Cependant, la taille limitée de l'ensemble, qui fait la force de ces algorithmes, vient avec son lot de contreparties. Notamment, il faut faire appel à de la localisation et de l'inflation pour pallier les erreurs d'échantillonnage. Parallèlement, les méthodes variationnelles, qui dans certains contextes idéalisés sont équivalents au filtre de Kalman et à l'EnKF, permettent de prendre en compte les non-linéarités des opérateurs d'observation et des modèles grâce aux méthodes itératives de minimisation. Un de leurs problèmes est de définir correctement la matrice des erreurs d'ébauche, qui impacte le résultat de l'analyse. Mais le défaut majeur de ces méthodes reste la nécessité de recourir au calcul du tangent linéaire et de l'adjoint de ces modèles et opérateurs d'observation.

Il s'agit alors de développer des méthodes dites variationnelles d'ensemble, ou EnVar 4D ([Lorenc, 2013](#); [Asch et al., 2016b](#)), qui tirent profit des deux formalismes pour essayer de garder leurs avantages en s'abstrayant de leur défauts. Nous présenterons dans ce contexte les méthodes hybrides dans un premier temps, puis le 4DEnVar ([Desroziers et al., 2014](#)), qui utilise un ensemble pour ne pas avoir à recourir à l'adjoint du modèle, et qui est développé pour la prévision opérationnelle à Météo-France entre autres. Enfin nous présenterons l'itérative ensemble Kalman smoother (IEnKS) ([Bocquet et Sakov, 2014](#); [Bocquet, 2016](#)), qui utilise un ensemble de trajectoires 4D pour faire une analyse variationnelle non-linéaire tout en propageant les erreurs. Toutes ces méthodes ont été développées dans le contexte de la contrainte forte, c'est-à-dire sans erreur modèle. Aussi, terminerons-nous par la présentation d'une nouvelle méthode variationnelle d'ensemble pensée pour la contrainte faible, nommée IEnKF-Q.

2.5.1 Méthodes hybrides

On appelle hybrides des méthodes où les erreurs d'ébauche ont une partie statique, comme dans l'interpolation optimale, le 3D-Var et le 4D-Var, et une partie ensembliste produite par un EnKF. L'intérêt de ces méthodes est de faire bénéficier un algorithme de type 3D-Var par exemple d'une représentation dynamique des erreurs et réciproquement, d'éviter la déficience de rang des matrices produites par un algorithme de type EnKF en lui adjoignant une matrice de rang plein définissant, a priori, les erreurs climatologiques.

On notera donc $\mathbf{P}^f = \mathbf{X}_f \mathbf{X}_f^T$ la partie ensembliste des erreurs et $\mathbf{C}$ la partie climatologique. On choisit alors $0 \leq \gamma \leq 1$ et on définit la nouvelle matrice des erreurs d'ébauche comme

$$\mathbf{B} = \gamma \mathbf{C} + (1 - \gamma) \mathbf{P}^f. \quad (2.103)$$

Si on prend $\gamma = 1$, on retombe sur le 3D-Var, tandis que $\gamma = 0$ correspond à l'EnKF.

Dans le cas stochastique, on peut simplement minimiser pour chaque membre de l'ensemble avec $i \in [1, m]$

$$\mathcal{J}_i(\mathbf{x}) = \frac{1}{2} \|\mathbf{y} + \mathbf{u}_i - \mathcal{H}(\mathbf{x})\|_{\mathbf{R}}^2 + \frac{1}{2} \|\mathbf{x} - \mathbf{x}_f^i\|_{\mathbf{B}}^2, \quad (2.104)$$

en reprenant les notations de la section 2.4.1.1. On obtient ainsi un ensemble d'analyses qui correspond parfaitement aux erreurs voulues. Cet algorithme a été proposé par Hamill et Snyder (2000).

Dans le cas déterministe, la situation est plus compliquée, du moins pour construire l'ensemble a posteriori. On peut faire une analyse pour l'état moyen en minimisant

$$\mathcal{J}(\mathbf{x}) = \frac{1}{2} \|\mathbf{y} - \mathcal{H}(\mathbf{x})\|_{\mathbf{R}}^2 + \frac{1}{2} \|\mathbf{x} - \bar{\mathbf{x}}_f\|_{\mathbf{B}}^2. \quad (2.105)$$

La méthode simpliste pour réengendrer l'ensemble est de suivre littéralement la formule de l'ETKF pour obtenir les analyses (2.50), sans tenir compte de la matrice $\mathbf{C}$ dans la formulation. Cela revient à fixer $\gamma = 0$ pour la génération des anomalies, ce qui risque de sous-estimer la variance des erreurs d'analyse, et imposer de recourir à de l'inflation.

Une autre solution est de partir de la formule de la transformée à gauche des anomalies analysées (2.65), à laquelle on applique la formule de SMW, pour avoir

$$\mathbf{X}_a = (\mathbf{I}_n + \mathbf{B} \mathbf{H}^T \mathbf{R}^{-1} \mathbf{H})^{-\frac{1}{2}} \mathbf{X}_f \mathbf{U}. \quad (2.106)$$

Cette formule fait apparaître la matrice $\mathbf{B}$ complète cette fois-ci, mais malheureusement pour de gros systèmes, il est quasiment impossible de calculer cette racine carrée explicitement.

2.5.2 4D-EnVar

De nombreux centres de prévisions météorologiques se sont tournés vers l'utilisation du 4D-Var dans les années 90. Étant donnés les efforts investis dans cette méthode par ces centres et devant le problème que représente le maintien d'un adjoint, ils ont récemment fait émerger le 4D-EnVar qui se fonde sur l'algorithme du 4D-Var tout en évitant l'utilisation de l'adjoint et également en ayant des erreurs évoluant avec le temps (Liu *et al.*, 2008; Buehner *et al.*, 2010; Desroziers *et al.*, 2014; Lorenc *et al.*, 2015; Buehner *et al.*, 2015).

Rappelons la fonction de coût du 4D-Var

$$\mathcal{J}(\mathbf{x}_0) = \frac{1}{2} \|\mathbf{x}_0 - \bar{\mathbf{x}}_f\|_{\mathbf{P}^f}^2 + \frac{1}{2} \sum_{l=0}^L \|\mathbf{y}_l - \mathcal{H}_l \circ \mathcal{M}_{l \leftarrow 0}(\mathbf{x}_0)\|_{\mathbf{R}_l}^2. \quad (2.107)$$

Comme dans un ETKF, on suppose que l'on a $\mathbf{P}^f = \mathbf{X}_f \mathbf{X}_f^T$, et que l'on réécrit la fonction de coût en fonction des coordonnées $\mathbf{w}$ dans l'espace de l'ensemble engendré par $\bar{\mathbf{x}}_f$ et les anomalies $\mathbf{X}_f$, en ajoutant le terme de jauge comme pour (2.58),

$$\mathcal{J}(\mathbf{w}) = \frac{1}{2} \|\mathbf{w}\|^2 + \frac{1}{2} \sum_{l=0}^L \|\mathbf{y}_l - \mathcal{H}_l \circ \mathcal{M}_{l \leftarrow 0}(\bar{\mathbf{x}}_f + \mathbf{X}_f \mathbf{w})\|_{\mathbf{R}_l}^2. \quad (2.108)$$

Si on linéarise autour de $\bar{\mathbf{x}}_f$, on a alors

$$\mathcal{J}_{\text{LIN}}(\mathbf{w}) = \frac{1}{2} \|\mathbf{w}\|^2 + \frac{1}{2} \sum_{l=0}^L \|\mathbf{y}_l - \mathcal{H}_l \circ \mathcal{M}_{l \leftarrow 0}(\bar{\mathbf{x}}_f) + \mathbf{H}_l \mathbf{M}_{l \leftarrow 0} \mathbf{X}_f \mathbf{w}\|_{\mathbf{R}_l}^2, \quad (2.109)$$

Le minimum de cette fonction est atteint en

$$\mathbf{w}_a = \left[\mathbf{I}_m + \sum_{l=0}^L (\mathbf{H}_l \mathbf{M}_{l \leftarrow 0} \mathbf{X}_f)^T \mathbf{R}_l^{-1} (\mathbf{H}_l \mathbf{M}_{l \leftarrow 0} \mathbf{X}_f) \right]^{-1} \sum_{l=0}^L (\mathbf{H}_l \mathbf{M}_{l \leftarrow 0} \mathbf{X}_f)^T \mathbf{R}_l^{-1} (\mathbf{y}_l - \mathcal{H}_l \circ \mathcal{M}_{l \leftarrow 0}(\bar{\mathbf{x}}_f)). \quad (2.110)$$

Le point clé est le fait que le terme $\mathbf{Y}_l = (\mathbf{H}_l \mathbf{M}_{l \leftarrow 0} \mathbf{X}_f)$ peut être estimé grâce à une simple propagation de l'ensemble, avec une formule similaire à (2.67)

$$(\mathbf{H}_l \mathbf{M}_{l \leftarrow 0} \mathbf{X}_f) \simeq \frac{1}{\varepsilon} \mathcal{H}_l \circ \mathcal{M}_{l \leftarrow 0} (\bar{\mathbf{x}}_f \mathbf{1}^T + \varepsilon \mathbf{X}_f) \left(\mathbf{I}_m - \frac{\mathbf{1} \mathbf{1}^T}{m} \right), \quad (2.111)$$

où on a à nouveau $0 < \varepsilon \ll 1$ pour réduire les anomalies et calculer les différences finies. Si la dispersion de l'ensemble est déjà suffisamment faible, on peut aller jusqu'à prendre $\varepsilon = 1$, ce qui est le cas dans les implémentations du 4D-EnVar à Météo-France.

Ce type d'analyse est communément appelé 4D-EnVar. Comme pour le 4D-Var, la définition de la matrice des erreurs d'ébauche $\mathbf{P}_f$, ou ici plus précisément des anomalies d'ébauche $\mathbf{X}_f$, reste un problème. Elles peuvent par exemple être engendrées grâce à l'ensemble issu d'un EnKF indépendant. En outre, il est possible de produire des anomalies d'analyse grâce à une formule déterministe de type ETKF, comme le fait le 4D-ETKF (Hunt *et al.*, 2004, 2007). Par ailleurs, le fait que l'on ne minimise pas la fonction de coût (2.108) mais sa version linéarisée (2.109) empêche de prendre en compte la non-linéarité du modèle.

Un dernier point problématique de l'algorithme est l'implémentation de la localisation. Dans l'application du 4D-EnVar, comme dans le 4D-Var, des covariances asynchrones, par exemple du type $\mathbf{X}_f \mathbf{Y}_l$, apparaissent et doivent être localisées. Malheureusement, même si on peut localiser les erreurs d'ébauche par un produit de Schur par exemple (voir 2.4.4.1), la propagation par le modèle et la localisation ne commutent pas (Fairbairn *et al.*, 2014). Ainsi, conserver un opérateur de localisation unique et fixe tout au long de la fenêtre n'est pas une solution satisfaisante, bien qu'applicable. Bocquet (2016) explique que si la localisation est effectuée par un produit de Schur avec une matrice $\boldsymbol{\rho}$ définie positive, et que la dynamique du modèle est hyperbolique, une équation d'évolution de la longueur de localisation de $\boldsymbol{\rho}$ peut être dérivée. Mais les cas pratiques où ces conditions sont réunies sont limités. Notamment, dès que le modèle est diffusif, la dynamique devient parabolique et non plus hyperbolique. Il propose alors, comme Desroziers *et al.* (2016), d'advecter les matrices de localisation tout au long de la fenêtre avec un modèle de substitution nécessitant un temps de calcul moins élevé que le modèle complet.

2.5.3 Le lisseur de Kalman d'ensemble itératif : IEnKS

Contrairement aux méthodes EnVar 4D précédentes, dont le développement était heuristique, l'itérative ensemble Kalman smoother (IEnKS) a une dérivation bayésienne de son algorithme, présentée dans [Bocquet et Sakov \(2014\)](#). Ceci permet d'avoir une connaissance précise des conditions optimales d'utilisation, limitations et approximations associées à ce schéma. Il étend à une analyse variationnelle 4D l'itérative ensemble Kalman filter (IEnKF) de [Sakov et al. \(2012\)](#). Bien qu'il soit appelé lisseur (*smoother* en anglais), il peut également fournir une estimation en filtrage et en prévision. Nous présentons tout d'abord la version *single data assimilation* (SDA) de l'algorithme, où les observations sont assimilées une unique fois.

2.5.3.1 L'algorithme

L'algorithme n'a pas besoin de recourir au tangent linéaire et à l'adjoint des opérateurs d'observation et de propagation, car il les estime comme le 4D-EnVar. Nous supposons alors qu'ils sont non-linéaires.

Étape d'analyse :

Nous reprenons le formalisme de l'ETKF et réécrivons la fonction de coût, écrite dans l'espace de l'ensemble, qui a déjà été utilisée pour le 4D-EnVar

$$\mathcal{J}(\mathbf{w}) = \frac{1}{2} \|\mathbf{w}\|^2 + \frac{1}{2} \sum_{l=L-S+1}^L \|\mathbf{y}_l - \mathcal{H}_l \circ \mathcal{M}_{l \leftarrow 0}(\bar{\mathbf{x}}_0 + \mathbf{X}_0 \mathbf{w})\|_{\mathbf{R}_l}^2. \quad (2.112)$$

Dans cette fonction de coût, $\bar{\mathbf{x}}_0$ et $\mathbf{X}_0$ désignent respectivement la moyenne et les anomalies de l'ensemble au temps t_0 et on assimile les $S \leq L + 1$ observations $\mathbf{y}_{L-S+1}, \dots, \mathbf{y}_L$ comprises entre les temps t_{L-S+1} et t_L . $\mathbf{w}$ désigne toujours les coordonnées du vecteur d'état dans l'espace de l'ensemble. Le cas où $L = 0$ et $S = 1$ correspond au MLEF (voir 2.4.2.2). Le cas $L = 1$ et $S = 1$ correspond à l'IEnKF de [Sakov et al. \(2012\)](#). Enfin, si on se limite à une itération de minimisation avec $S = L + 1$, on retombe sur le 4D-ETKF ([Hunt et al., 2007](#)).

Plutôt que de linéariser cette fonction de coût autour de l'ébauche comme pour le 4D-EnVar et le 4D-ETKF, nous allons minimiser complètement la fonction de coût. Ceci permet de mieux prendre en compte les non-linéarités des opérateurs d'observation et de propagation. Dans la suite, nous choisissons d'utiliser un algorithme de Gauss-Newton pour la minimisation mais d'autre choix sont possibles (notamment l'algorithme de Levenberg-Marquardt comme dans [Bocquet et Sakov \(2012\)](#)).

Nous utilisons la notation (j) pour désigner la $j^{\text{ème}}$ itération de Gauss-Newton. Pour la première itération, on sélectionne $\mathbf{w}^{(0)} = \mathbf{0}$. Nous avons alors

$$\mathbf{x}_0^{(j)} = \bar{\mathbf{x}}_0 + \mathbf{X}_0 \mathbf{w}^{(j)} \quad (2.113)$$

$$\nabla \mathcal{J}^{(j)} = \mathbf{w}^{(j)} - \sum_{l=L-S+1}^L \mathbf{Y}_l^{(j)\top} \mathbf{R}_l^{-1} (\mathbf{y}_l - \mathcal{H}_l \circ \mathcal{M}_{l \leftarrow 0}(\mathbf{x}_0^{(j)})) \quad (2.114)$$

$$\mathbf{D}^{(j)} = \left[\mathbf{I}_m + \sum_{l=L-S+1}^L \mathbf{Y}_l^{(j)\top} \mathbf{R}_l^{-1} \mathbf{Y}_l^{(j)} \right]^{-1}, \quad (2.115)$$

où $\mathbf{Y}_l^{(j)} = \nabla[\mathcal{H}_l \circ \mathcal{M}_{l \leftarrow 0}]_{|\mathbf{x}_0^{(j)}} \mathbf{X}_0$. Si le modèle adjoint n'est pas disponible, on peut à nouveau utiliser la formulation *bundle* pour estimer ce tangent linéaire, comme pour le MLEF et le 4D-EnVar,

$$\mathbf{Y}_l^{(j)} \simeq \frac{1}{\varepsilon} \mathcal{H}_l \circ \mathcal{M}_{l \leftarrow 0}(\mathbf{x}_0^{(j)} \mathbf{1}^\top + \varepsilon \mathbf{X}_0) (\mathbf{I}_m - \frac{\mathbf{1} \mathbf{1}^\top}{m}). \quad (2.116)$$

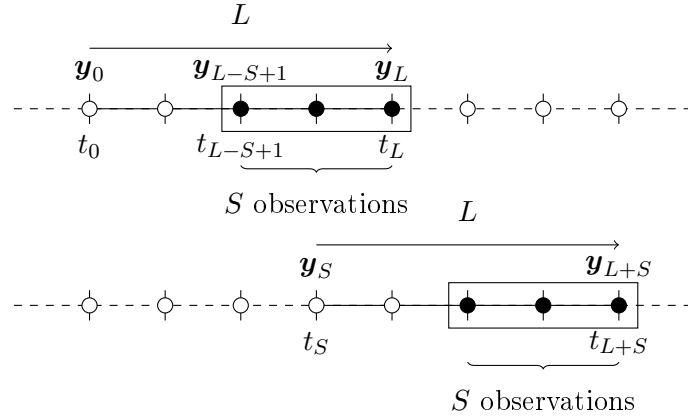


FIGURE 2.4 – Cyclage de l’IEnKS SDA avec $L = 4$ et $S = 3$, en termes de Δt , l’intervalle de temps entre deux observations. Seules les S observations marquées en noir sont assimilées, mais le lissage s’effectue sur toute la fenêtre de longueur $L\Delta t$.

On incrémente $\mathbf{w}^{(j)}$ de la sorte

$$\mathbf{w}^{(j+1)} = \mathbf{w}^{(j)} - \mathbf{D}^{(j)} \nabla \mathcal{J}^{(j)} \quad (2.117)$$

et on boucle jusqu’à ce que l’incrément $\|\mathbf{w}^{(j+1)} - \mathbf{w}^{(j)}\|$ soit plus petit qu’une tolérance fixée. On a alors atteint le minimum $\mathbf{w}^*$ et l’analyse suit celle de l’ETKF

$$\mathbf{x}_a = \bar{\mathbf{x}}_0 + \mathbf{X}_0 \mathbf{w}^* \quad (2.118)$$

$$\mathbf{X}_a = \mathbf{X}_f (\mathbf{D}^*)^{\frac{1}{2}} = \mathbf{X}_f \left(\mathbf{I}_m + \sum_{l=L-S+1}^L \mathbf{Y}_l^{*T} \mathbf{R}_l^{-1} \mathbf{Y}_l^* \right)^{-\frac{1}{2}}. \quad (2.119)$$

Cette analyse nous donne une estimation de l’état au temps t_0 . Elle peut être propagée jusqu’au bout de la fenêtre pour obtenir l’estimation de l’état au temps t_L et le score en filtrage.

Étape de propagation :

Pour cycler l’algorithme, on propage l’ensemble de S pas de temps, pour que S nouvelles observations rentrent dans la fenêtre et que l’on puisse faire une nouvelle analyse. L’ensemble $\mathbf{E}_S$ au temps t_S de la nouvelle analyse est donc simplement

$$\mathbf{E}_S = \mathcal{M}_{S \leftarrow 0} (\mathbf{x}_a \mathbf{1}^T + \sqrt{m-1} \mathbf{X}_a). \quad (2.120)$$

Le cyclage de l’algorithme est représenté schématiquement dans la Fig. 2.4.

La combinaison de la propagation des erreurs via l’ensemble et d’une analyse variationnelle quadri-dimensionnelle rend cette méthode robuste et apte à battre le 4D-Var et l’EnKF/EnKS sur de multiples modèles d’ordre réduit (voir chapitre 4; Bocquet et Sakov (2013, 2014); Haussaire et Bocquet (2016)). Ce résultat a été validé sur tous les régimes de non-linéarité testés, en lissage, filtrage et prévision. L’IEnKS s’est également avéré utile pour l’estimation de paramètres des modèles. La seule approximation de ce schéma est l’hypothèse gaussienne lors de la reconstruction des anomalies d’analyse et concernant l’ébauche, ainsi que la contrainte forte (absence d’erreur modèle). Enfin, son formalisme permet d’inclure une variété de méthodes EnVar 4D en tant que sous-cas, ce qui fait de cet algorithme un parfait archétype de ces méthodes. Il ne manque que son extension à une matrice des erreurs d’ébauche hybride, comme à la section 2.5.1.

Le pseudo-code correspondant à un cycle de l’IEnKS est explicité dans l’algorithme 6.

Algorithme 6 Algorithme pour un cycle de l'IEEnKS SDA *bundle*, de lag L , assimilant S observations par cycle, avec minimisation de Gauss-Newton.

Requis: $m, \mathbf{E}_{0,f}, \{\mathbf{y}_l\}_{l \in [L-S+1, L]}, \{\mathbf{R}_l\}_{l \in [L-S+1, L]}, \{\mathcal{M}_{l \leftarrow l-1}\}_{l \in [1, L]}, \{\mathcal{H}_l\}_{l \in [L-S+1, L]}, e, j_{\max}, \varepsilon, \mathbf{U}, S, L$

- 1: $\bar{\mathbf{x}}_0 = \mathbf{E}_{0,f} \mathbf{1} / m$
- 2: $\mathbf{X}_0 = (\mathbf{E}_{0,f} - \bar{\mathbf{x}}_0 \mathbf{1}^T) / \sqrt{m-1}$
- 3: $j = 0, \mathbf{w} = \mathbf{0}, P = L - S + 1$
- 4: **repeat**
- 5: $\mathbf{x}_0 = \bar{\mathbf{x}}_0 + \mathbf{X}_0 \mathbf{w}$
- 6: $\mathbf{E}_0 = \mathbf{x}_0 \mathbf{1}^T + \varepsilon \mathbf{X}_0$
- 7: $\mathbf{E}_P = \mathcal{M}_{P \leftarrow 0}(\mathbf{E}_0)$
- 8: $\bar{\mathbf{y}}_P = \mathcal{H}_P(\mathbf{E}_P) \mathbf{1} / m$
- 9: $\mathbf{Y}_P = (\mathcal{H}_P(\mathbf{E}_P) - \bar{\mathbf{y}}_P \mathbf{1}^T) / \varepsilon$
- 10: **for** $l = P + 1 \dots L$ **do**
- 11: $\mathbf{E}_l = \mathcal{M}_{l \leftarrow l-1}(\mathbf{E}_{l-1})$
- 12: $\bar{\mathbf{y}}_l = \mathcal{H}_l(\mathbf{E}_l) \mathbf{1} / m$
- 13: $\mathbf{Y}_l = (\mathcal{H}_l(\mathbf{E}_l) - \bar{\mathbf{y}}_l \mathbf{1}^T) / \varepsilon$
- 14: **end for**
- 15: $\nabla \mathcal{J} = \mathbf{w} - \sum_{l=P}^L \mathbf{Y}_l^T \mathbf{R}_l^{-1} (\mathbf{y}_l - \bar{\mathbf{y}}_l)$
- 16: $\mathbf{D} = [\mathbf{I}_m + \sum_{l=P}^L \mathbf{Y}_l^T \mathbf{R}_l^{-1} \mathbf{Y}_l]^{-1}$
- 17: $\Delta \mathbf{w} = \mathbf{D} \nabla \mathcal{J}$
- 18: $\mathbf{w} := \mathbf{w} - \Delta \mathbf{w}$
- 19: $j := j + 1$
- 20: **until** $j \geq j_{\max}$ **or** $\|\Delta \mathbf{w}\| < e$
- 21: $\mathbf{E}_{0,s} = \mathbf{x}_0 \mathbf{1}^T + \sqrt{m-1} \mathbf{X}_0 \mathbf{D}^{\frac{1}{2}} \mathbf{U}$
- 22: $\mathbf{E}_{S,f} = \mathcal{M}_{S \leftarrow 0}(\mathbf{E}_{0,s})$

2.5.3.2 Le MDA : assimilation multiple des observations

L'intérêt que présente l'IEnKS est d'analyser des observations éloignées dans le temps, de telle sorte que l'on focalise l'analyse sur la partie instable du modèle chaotique. Cependant, pour de trop grandes longueurs de fenêtre d'assimilation de données, ceci échoue. Pour stabiliser le schéma, une idée est d'assimiler plus d'observations, quitte à les observer plusieurs fois. La difficulté est alors de rester statistiquement cohérent, ce que [Bocquet et Sakov \(2014\)](#) s'attellent à démontrer. On dénommera ainsi le schéma *multiple data assimilation* (MDA) en opposition au précédent que l'on a qualifié de *single data assimilation* (SDA).

On peut généraliser la formule (2.112) en autorisant l'assimilation de n'importe quelles observations dans la fenêtre avec un poids $0 \leq \beta_l \leq 1$

$$\mathcal{J}(\mathbf{w}) = \frac{1}{2} \|\mathbf{w}\|^2 + \frac{1}{2} \sum_{l=0}^L \beta_l \|\mathbf{y}_l - \mathcal{H}_l \circ \mathcal{M}_{l \leftarrow 0}(\bar{\mathbf{x}}_0 + \mathbf{X}_0 \mathbf{w})\|_{\mathbf{R}_l}^2. \quad (2.121)$$

Si on choisit $\beta_l = 0$ pour $0 \leq l < L - S + 1$ et $\beta_l = 1$ pour $L - S + 1 \leq l \leq L$, on retombe sur l'équation (2.112) du SDA, et les observations sont assimilées une unique fois.

Assimiler une observation $\mathbf{y}_l$ avec un poids β_l revient à considérer lors de l'analyse une probabilité gaussienne de l'observation de la forme

$$p(\mathbf{y}_l^{\beta_l} | \mathbf{x}) = \frac{\exp\{-\frac{\beta_l}{2} \|\mathbf{y}_l - \mathcal{H}(\mathbf{x})\|_{\mathbf{R}}^2\}}{\sqrt{(2\pi/\beta_l)^p |\mathbf{R}|}}, \quad (2.122)$$

où $|\mathbf{R}|$ est le déterminant de $\mathbf{R}$. La notation $\mathbf{y}_l^{\beta_l}$ fait référence à l'analyse de $\mathbf{y}_l$ partielle de poids β_l . Pour faire une assimilation statistiquement cohérente d'une observation, il faut que la somme de ses poids soit égale à 1, c'est-à-dire $\sum_{l=0}^L \beta_l = 1$.

Il est bon de préciser que pour obtenir une estimation d'un état à un temps $t_l > t_0$, une simple propagation de l'état optimal à t_0 ne suffit plus. Il faut appliquer ce que [Bocquet et Sakov \(2014\)](#) appellent une *balancing step*, qui consiste en une analyse supplémentaire, avec des poids différents, pour obtenir un état à t_0 dont la propagation donne les bonnes statistiques.

Un dernier point à préciser est que l'on peut choisir les β_l tels que $\sum_{l=0}^L \beta_l \neq 1$ mais de telle sorte que l'on obtienne une puissance de la pdf précédente, ce qui peut présenter des avantages, notamment lors de l'application du *finite-size*. Mon collègue Anthony Fillion poursuit ces travaux en étudiant une version quasi-statique de l'IEnKS.

2.5.3.3 Localisation de l'IEnKS

En ce qui concerne les notations, dans cette section, la taille de l'ensemble sera notée m_e . La notation m sera réservée pour indexer les analyses locales.

La localisation avec l'IEnKS est aussi compliquée que pour le 4DEnVar. Dans le contexte de la localisation par domaine cependant, [Bocquet \(2016\)](#) propose un algorithme. Dans un premier temps, il fait fi du problème de la covariance des domaines de localisation. Ce schéma est nommé la localisation par domaines statiques. Pour chaque itération de la minimisation, l'algorithme consiste donc, en chaque point de grille η_m , à effectuer une analyse locale pour obtenir un $\mathbf{w}^m$, comme pour le LETKF. Il utilise les observations situées dans un domaine D_m centré autour de η_m , en remplaçant la matrice des erreurs d'observation $\mathbf{R}$ par une matrice locale $\mathbf{R}^m$, qui ne se rapporte qu'aux observations proches et dont les variances augmentent avec la distance. L'itération se finit en engendrant un nouvel état global au temps t_0 qui sert de point de départ pour l'itération suivante, et qui est issu de la reconstruction à partir de tous les $\mathbf{w}^m$ locaux. Les sensibilités sont calculées à

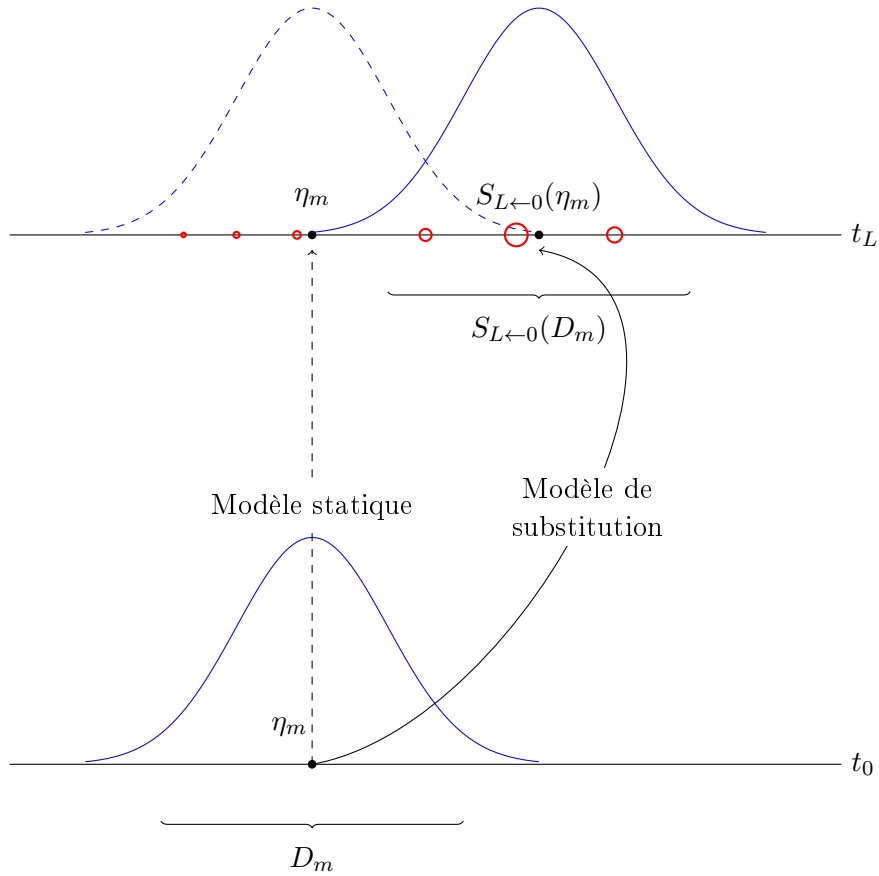


FIGURE 2.5 – Schéma de l'analyse avec une localisation par domaines covariants et statiques. Les variables d'état au point η_m à t_0 sont analysées avec les observations au temps t_L . Leur précision est diminuée en fonction de la distance entre le centre du domaine initial D_m et leur position équivalente après advection, par l'application de la fonction de Gaspari-Cohn (2.82). La localisation par domaines statiques ne bouge pas le domaine. Les observations sont marquées par des cercles rouges dont la taille indique la précision conformément à la fonction de Gaspari-Cohn.

partir d'une propagation des anomalies globales, aussi cette étape du calcul n'est pas plus coûteuse qu'avec l'algorithme global. Le détail de l'algorithme est décrit par l'algorithme 7.

Cependant, avec cette méthode, pour analyser le point η_m au temps t_0 , on utilise les observations au temps t_L situées dans le domaine D_m centré autour de η_m . Comme il a été suggéré à la fin de la section 2.5.2, il faudrait plutôt advecter ce domaine D_m avec un modèle de substitution $S_{L←0}$ jusqu'au temps t_L et prendre les observations situées dans ce domaine $S_{L←0}(D_m)$. Plus simplement, on peut advecter en sens arrière (de t_L jusqu'à t_0) la position des observations, pour obtenir leur position équivalente à t_0 , et assimiler les observations dont la position équivalente est située dans le domaine D_m . Ce schéma est nommé la localisation par domaines covariants. Ces deux procédés de localisation sont schématisés sur la Fig. 2.5.

2.5.4 Formulation en contrainte faible : l'IEnKF-Q

Jusqu'ici, tous les algorithmes présentés faisaient l'hypothèse de modèle parfait, autrement appelé contrainte forte, à savoir que la propagation de la vérité par le modèle selon l'équation (2.11) utilisait une matrice $\mathbf{Q} = \mathbf{0}$. Cependant, cette hypothèse simplificatrice est bien souvent injustifiée

Algorithme 7 Algorithme pour un cycle de l'EnKS local SDA *bundle*, de lag L , assimilant S observations par cycle, avec minimisation de Gauss-Newton. La taille de l'ensemble est notée m_e dans cet algorithme.

Requis: : $m_e, S, L, M, \mathbf{E}_{0,f}, \{\mathbf{y}_l\}_{l \in [L-S+1, L]}, \{\mathbf{R}_l\}_{l \in [L-S+1, L]}, \{\mathcal{M}_{l \leftarrow l-1}\}_{l \in [1, L]}, \{\mathcal{H}_l\}_{l \in [L-S+1, L]},$
 $e, j_{\max}, \varepsilon, \{\mathbf{U}_m\}_{m \in [1, M]}$

```

1:  $\bar{\mathbf{x}}_0 = \mathbf{E}_{0,f} \mathbf{1} / m_e$ 
2:  $\mathbf{X}_0 = (\mathbf{E}_{0,f} - \bar{\mathbf{x}}_0 \mathbf{1}^T) / \sqrt{m_e - 1}$ 
3:  $j = 0, P = L - S + 1$ 
4: for  $m = 1 \dots M$  do
5:    $\mathbf{w}^m = \mathbf{0}$ 
6: end for
7: repeat
8:   for  $m = 1 \dots M$  do
9:      $[\mathbf{x}_0]_m = [\bar{\mathbf{x}}_0]_m + [\mathbf{X}_0 \mathbf{w}^m]_m$ 
10:   end for
11:    $\mathbf{E}_0 = \mathbf{x}_0 \mathbf{1}^T + \varepsilon \mathbf{X}_0$ 
12:    $\mathbf{E}_P = \mathcal{M}_{P \leftarrow 0}(\mathbf{E}_0)$ 
13:    $\bar{\mathbf{y}}_P = \mathcal{H}_P(\mathbf{E}_P) \mathbf{1} / m_e$ 
14:    $\mathbf{Y}_P = (\mathcal{H}_P(\mathbf{E}_P) - \bar{\mathbf{y}}_P \mathbf{1}^T) / \varepsilon$ 
15:   for  $l = P + 1 \dots L$  do
16:      $\mathbf{E}_l = \mathcal{M}_{l \leftarrow l-1}(\mathbf{E}_{l-1})$ 
17:      $\bar{\mathbf{y}}_l = \mathcal{H}_l(\mathbf{E}_l) \mathbf{1} / m_e$ 
18:      $\mathbf{Y}_l = (\mathcal{H}_l(\mathbf{E}_l) - \bar{\mathbf{y}}_l \mathbf{1}^T) / \varepsilon$ 
19:   end for
20:   for  $m = 1 \dots M$  do
21:      $\nabla \mathcal{J}_m = \mathbf{w}^m - \sum_{l=P}^L \mathbf{Y}_l^T (\mathbf{R}_l^m)^{-1} (\mathbf{y}_l - \bar{\mathbf{y}}_l)$ 
22:      $\mathbf{D}_m = [\mathbf{I}_{m_e} + \sum_{l=P}^L \mathbf{Y}_l^T (\mathbf{R}_l^m)^{-1} \mathbf{Y}_l]^{-1}$ 
23:      $\Delta \mathbf{w}^m = \mathbf{D}_m \nabla \mathcal{J}_m$ 
24:      $\mathbf{w}^m := \mathbf{w}^m - \Delta \mathbf{w}^m$ 
25:   end for
26:    $j := j + 1$ 
27: until  $j \geq j_{\max}$  or  $\sum_{m=1}^M \|\Delta \mathbf{w}^m\| < e$ 
28: for  $m = 1 \dots M$  do
29:    $[\mathbf{E}_{0,s}]_m = [\mathbf{x}_0 \mathbf{1}^T]_m + \sqrt{m_e - 1} \mathbf{X}_0 \mathbf{D}_m^{\frac{1}{2}} \mathbf{U}_m$ 
30: end for
31:  $\mathbf{E}_{S,f} = \mathcal{M}_{S \leftarrow 0}(\mathbf{E}_{0,s})$ 

```

et une branche de l'assimilation de données se penche sur le développement d'algorithmes avec erreur modèle, autrement qualifiée de contrainte faible. Cette erreur modèle peut être simplement incorporée dans les algorithmes ensemblistes stochastiques. C'est le cas pour l'équivalent stochastique de l'IEnKF, nommé ensemble randomized maximum likelihood filter (EnRML) et introduit par Gu et Oliver (2007), qui a vu son extension à la contrainte faible dans Mandel *et al.* (2016). Les filtres itératifs déterministes comme l'IEnKF n'ont cependant jusqu'à maintenant pas inclu l'erreur modèle. En effet, l'assimilation asynchrone des observations est plus aisée dans le cas du modèle parfait (Evensen et van Leeuwen, 2000; Hunt *et al.*, 2004; Sakov *et al.*, 2010). De plus, se placer dans un contexte avec erreur modèle a pour conséquence de considérablement augmenter l'espace de contrôle des algorithmes de type 4D-Var ou IEnKS. En effet, quand dans le cas de la contrainte forte, estimer l'état $\mathbf{x}_0$ au temps t_0 suffisait à reconstruire toute la trajectoire jusqu'à t_L par simple propagation, il faut désormais estimer l'ensemble des états $\mathbf{x}_0, \dots, \mathbf{x}_L$ (Trémolet, 2006).

L'article Sakov *et al.* (2017, soumis), que nous avons soumis au cours de la thèse, étend l'IEnKF au cas de la contrainte faible. C'est cet algorithme que nous présentons dans cette partie.

2.5.4.1 Dérivation

La fonction de coût (2.112) de l'IEnKS devient, en contrainte faible,

$$\mathcal{J}(\mathbf{x}_1, \mathbf{x}_2) = \frac{1}{2} \|\mathbf{x}_1 - \bar{\mathbf{x}}_f\|_{\mathbf{P}_f}^2 + \frac{1}{2} \|\mathbf{y}_2 - \mathcal{H}(\mathbf{x}_2)\|_{\mathbf{R}}^2 + \frac{1}{2} \|\mathbf{x}_2 - \mathcal{M}(\mathbf{x}_1)\|_{\mathbf{Q}}^2, \quad (2.123)$$

où $\mathcal{M} = \mathcal{M}_{2 \leftarrow 1}$. Puisque nous sommes dans le contexte de l'IEnKF, nous avons $L = S = 1$, ce qui explique qu'il ne reste que les variables $\mathbf{x}_1$ et $\mathbf{x}_2$. Les opérateurs d'observation $\mathcal{H}$ et de modèle $\mathcal{M}$ peuvent ici être non-linéaires. On remarque qu'un terme supplémentaire est apparu pour tenir compte de l'incertitude sur le modèle, selon la matrice de covariance d'erreur $\mathbf{Q}$. La minimisation de cette fonction de coût vise à obtenir les états optimaux $\{\mathbf{x}_1^*, \mathbf{x}_2^*\}$ tels que

$$\mathbf{x}_1^*, \mathbf{x}_2^* = \arg \min_{\{\mathbf{x}_1, \mathbf{x}_2\}} \mathcal{J}(\mathbf{x}_1, \mathbf{x}_2). \quad (2.124)$$

Nous pouvons à nouveau nous placer dans l'espace de l'ensemble et définir les coordonnées $\mathbf{u} \in \mathbb{R}^m$ et $\mathbf{v} \in \mathbb{R}^{m_q}$ de $\mathbf{x}_1$ et $\mathbf{x}_2$ respectivement selon

$$\mathbf{x}_1 = \bar{\mathbf{x}}_f + \mathbf{X}_1^f \mathbf{u}, \quad (2.125)$$

$$\mathbf{x}_2 = \mathcal{M}(\mathbf{x}_1) + \mathbf{X}^q \mathbf{v}, \quad (2.126)$$

avec

$$\mathbf{X}_1^f (\mathbf{X}_1^f)^T = \mathbf{P}_1^f, \quad (2.127)$$

$$\mathbf{X}^q (\mathbf{X}^q)^T = \mathbf{Q}. \quad (2.128)$$

$\mathbf{X}^q \in \mathbb{R}^{n \times m_q}$ est une anomalie d'erreur modèle. La fonction de coût, selon ces nouvelles coordonnées devient donc

$$\mathcal{J}(\mathbf{u}, \mathbf{v}) = \frac{1}{2} \mathbf{u}^T \mathbf{u} + \frac{1}{2} \|\mathbf{y}_2 - \mathcal{H}(\mathbf{x}_2)\|_{\mathbf{R}}^2 + \frac{1}{2} \mathbf{v}^T \mathbf{v}. \quad (2.129)$$

On peut concaténer $\mathbf{u}$ et $\mathbf{v}$ en un unique vecteur $\mathbf{w} = \begin{pmatrix} \mathbf{u} \\ \mathbf{v} \end{pmatrix}$, et ainsi

$$\mathcal{J}(\mathbf{w}) = \frac{1}{2} \mathbf{w}^T \mathbf{w} + \frac{1}{2} \|\mathbf{y}_2 - \mathcal{H}(\mathbf{x}_2)\|_{\mathbf{R}}^2, \quad (2.130)$$

avec $\mathbf{x}_2(\mathbf{w}) = \mathcal{M}(\bar{\mathbf{x}}_f + \mathbf{X}_1^f \mathbf{w}_{[1:m]}) + \mathbf{X}^q \mathbf{w}_{[m+1:m+m_q]}$ et $\mathbf{w}_{[a:b]}$ désigne un sous-vecteur de $\mathbf{w}$ formé des éléments compris entre les indices a et b .

Au minimum, le gradient de cette fonction de coût s'annule, conduisant à l'équation

$$\mathbf{w} - (\mathbf{H}\mathbf{X})^T \mathbf{R}^{-1} [\mathbf{y}_2 - \mathcal{H}(\mathbf{x}_2)] = \mathbf{0}, \quad (2.131)$$

avec

$$\mathbf{X} = [\mathbf{M}\mathbf{X}_1^f, \mathbf{X}^q], \quad (2.132)$$

$$\mathbf{H} = \nabla \mathcal{H}(\mathbf{x}_2), \quad (2.133)$$

$$\mathbf{M} = \nabla \mathcal{M}(\mathbf{x}_1). \quad (2.134)$$

On peut résoudre l'équation (2.131) itérativement par une méthode de Gauss-Newton, c'est-à-dire incrémenter $\mathbf{w}^{(j)}$ suivant l'équation (2.117)

$$\mathbf{w}^{(j+1)} = \mathbf{w}^{(j)} - \mathbf{D}^{(j)} \nabla \mathcal{J}(\mathbf{w}^{(j)}), \quad (2.135)$$

avec $\mathbf{D}$ l'inverse de l'Hessien de (2.130), que l'on approche de la sorte :

$$\mathbf{D}^{(j)} \approx [\mathbf{I}_{m+m_q} + (\mathbf{H}^{(j)} \mathbf{X}^{(j)})^T \mathbf{R}^{-1} \mathbf{H}^{(j)} \mathbf{X}^{(j)}]^{-1}. \quad (2.136)$$

Cela donne les incréments

$$\begin{aligned} \mathbf{w}^{(j+1)} - \mathbf{w}^{(j)} &= [\mathbf{I}_{m+m_q} + (\mathbf{H}^{(j)} \mathbf{X}^{(j)})^T \mathbf{R}^{-1} \mathbf{H}^{(j)} \mathbf{X}^{(j)}]^{-1} \\ &\times \left\{ (\mathbf{H}^{(j)} \mathbf{X}^{(j)})^T \mathbf{R}^{-1} [\mathbf{y}_2 - \mathcal{H}(\mathbf{x}_2^{(j)})] - \mathbf{w}^{(j)} \right\}. \end{aligned} \quad (2.137)$$

Cette équation, qui permet d'obtenir l'analyse, est le cœur de l'IEnKF-Q.

Il reste cependant à obtenir les anomalies lissées $\mathbf{X}_1^s$ et filtrées $\mathbf{X}_2^a$, la première pour réduire l'incertitude à t_1 après l'assimilation d'une nouvelle observation, et la deuxième pour pouvoir commencer un nouveau cycle d'assimilation.

On définit pour cela les perturbations des états à t_1 et t_2 en fonction des perturbations $\delta \mathbf{u}$ et $\delta \mathbf{v}$

$$\delta \mathbf{x}_1 = \mathbf{X}_1^f \delta \mathbf{u}, \quad \delta \mathbf{x}_2 = \mathbf{M}\mathbf{X}_1^f \delta \mathbf{u} + \mathbf{X}^q \delta \mathbf{v}. \quad (2.138)$$

Dans le cas linéaire, l'inverse de l'Hessien après convergence $\mathbf{D}^*$ représente la covariance des erreurs dans l'espace de l'ensemble de telle sorte que $\mathbb{E}[\mathbf{w}^*(\mathbf{w}^*)^T] = \mathbf{D}^*$. En utilisant l'expression (2.138), on a alors

$$\mathbf{X}_2^a (\mathbf{X}_2^a)^T = \mathbb{E}[\delta \mathbf{x}_2^* (\delta \mathbf{x}_2^*)^T] = \mathbf{X}^* \mathbb{E}[\mathbf{w}^* (\mathbf{w}^*)^T] (\mathbf{X}^*)^T = \mathbf{X}^* \mathbf{D}^* (\mathbf{X}^*)^T \quad (2.139)$$

et

$$\mathbf{X}_2^a = \mathbf{X}^* (\mathbf{D}^*)^{1/2} = \mathbf{X}^* [\mathbf{I}_{m+m_q} + (\mathbf{H}^{(j)} \mathbf{X}^{(j)})^T \mathbf{R}^{-1} \mathbf{H}^{(j)} \mathbf{X}^{(j)}]^{-1/2}. \quad (2.140)$$

Similairement,

$$\mathbf{X}_1^s (\mathbf{X}_1^s)^T = \mathbb{E}[\delta \mathbf{x}_1^* (\delta \mathbf{x}_1^*)^T] = \mathbf{X}_1^f \mathbb{E}[\mathbf{u}^* (\mathbf{u}^*)^T] (\mathbf{X}_1^f)^T, \quad (2.141)$$

et donc

$$\mathbf{X}_1^s = \mathbf{X}_1^f (\mathbf{D}_{[1:m,1:m]}^*)^{1/2}, \quad (2.142)$$

avec $\mathbf{D}_{[1:m,1:m]}^*$ la sous-matrice de $\mathbf{D}^*$ formée des m premières lignes et colonnes de $\mathbf{D}^*$.

Les trois équations (2.135), (2.140) et (2.142) constituent l'IEnKF-Q.

2.5.4.2 Découplage de $\mathbf{u}$ et $\mathbf{v}$ quand l'opérateur d'observation est linéaire

On peut réécrire l'équation (2.137)

$$\begin{aligned} \mathbf{w}^{(j+1)} = & [\mathbf{I}_{m+m_q} + (\mathbf{H}^{(j)} \mathbf{X}^{(j)})^T \mathbf{R}^{-1} \mathbf{H}^{(j)} \mathbf{X}^{(j)}]^{-1} \\ & \times (\mathbf{H}^{(j)} \mathbf{X}^{(j)})^T \mathbf{R}^{-1} [\mathbf{y}_2 - \mathcal{H}(\mathbf{x}_2^{(j)}) + \mathbf{H}^{(j)} \mathbf{X}^{(j)} \mathbf{w}^{(j)}], \end{aligned} \quad (2.143)$$

et en utilisant la formule de SMW,

$$\begin{aligned} \mathbf{w}^{(j+1)} = & (\mathbf{H}^{(j)} \mathbf{X}^{(j)})^T [(\mathbf{H}^{(j)} \mathbf{X}^{(j)})^T \mathbf{H}^{(j)} \mathbf{X}^{(j)} + \mathbf{R}]^{-1} \\ & \times [\mathbf{y}_2 - \mathcal{H}(\mathbf{x}_2^{(j)}) + \mathbf{H}^{(j)} \mathbf{X}^{(j)} \mathbf{w}^{(j)}]. \end{aligned} \quad (2.144)$$

On peut déjà remarquer que dans le cas d'un $\mathcal{H}$ linéaire, et donc tel que $\mathcal{H} = \mathbf{H}^{(j)} = \mathbf{H}$, on a

$$\mathcal{H}(\mathbf{x}_2^{(j)}) = \mathcal{H}[\mathcal{M}(\mathbf{x}_1^{(j)}) + \mathbf{X}^q \mathbf{v}^{(j)}] = \mathcal{H} \circ \mathcal{M}(\bar{\mathbf{x}}_f + \mathbf{X}_1^f \mathbf{u}^{(j)}) + \mathbf{H} \mathbf{X}^q \mathbf{v}^{(j)}. \quad (2.145)$$

Ainsi, (2.144) devient, en séparant $\mathbf{w}$ en $\mathbf{u}$ et $\mathbf{v}$,

$$\begin{aligned} \mathbf{u}^{(j+1)} = & (\mathbf{H}^{(j)} \mathbf{M}^{(j)} \mathbf{X}_1^f)^T \left[(\mathbf{H}^{(j)} \mathbf{M}^{(j)} \mathbf{X}_1^f)^T \mathbf{H}^{(j)} \mathbf{M}^{(j)} \mathbf{X}_1^f + \mathbf{R}_u^{(j)} \right]^{-1} \\ & \times \left[\mathbf{y}_2 - \mathcal{H} \circ \mathcal{M}(\mathbf{x}_1^{(j)}) + \mathbf{H}^{(j)} \mathbf{M}^{(j)} \mathbf{X}_1^f \mathbf{u}^{(j)} \right], \end{aligned} \quad (2.146)$$

$$\begin{aligned} \mathbf{v}^{(j+1)} = & (\mathbf{H}^{(j)} \mathbf{X}^q)^T \left[(\mathbf{H}^{(j)} \mathbf{X}^q)^T \mathbf{H}^{(j)} \mathbf{X}^q + \mathbf{R}_v^{(j)} \right]^{-1} \\ & \times \left[\mathbf{y}_2 - \mathcal{H} \circ \mathcal{M}(\mathbf{x}_1^{(j)}) + \mathbf{H}^{(j)} \mathbf{M}^{(j)} \mathbf{X}_1^f \mathbf{u}^{(j)} \right], \end{aligned} \quad (2.147)$$

où

$$\mathbf{R}_u^{(j)} = \mathbf{H}^{(j)} \mathbf{X}^q (\mathbf{H}^{(j)} \mathbf{X}^q)^T + \mathbf{R} \quad (2.148)$$

$$\mathbf{R}_v^{(j)} = \mathbf{H}^{(j)} \mathbf{M}^{(j)} \mathbf{X}_1^f (\mathbf{H}^{(j)} \mathbf{M}^{(j)} \mathbf{X}_1^f)^T + \mathbf{R}. \quad (2.149)$$

On peut alors remarquer que (2.146) ne dépend plus de $\mathbf{v}$. Plus précisément, en la réécrivant

$$\begin{aligned} \mathbf{u}^{(j+1)} - \mathbf{u}^{(j)} = & \mathbf{D}_u^{(j)} \left\{ (\mathbf{H}^{(j)} \mathbf{M}^{(j)} \mathbf{X}_1^f)^T (\mathbf{R}_u^{(j)})^{-1} \right. \\ & \times \left. \left[\mathbf{y}_2 - \mathcal{H} \circ \mathcal{M}(\bar{\mathbf{x}}_f + \mathbf{X}_1^f \mathbf{u}^{(j)}) \right] - \mathbf{u}^{(j)} \right\}, \end{aligned} \quad (2.150)$$

avec

$$\mathbf{D}_u^{(j)} = \left[\mathbf{I}_{m+m_q} + (\mathbf{H}^{(j)} \mathbf{M}^{(j)} \mathbf{X}_1^f)^T (\mathbf{R}_u^{(j)})^{-1} \mathbf{H}^{(j)} \mathbf{M}^{(j)} \mathbf{X}_1^f \right]^{-1}, \quad (2.151)$$

on retrouve une formule similaire à (2.117), ce qui signifie que $\mathbf{u}$ peut être obtenu à partir d'un IEnKF standard (c'est-à-dire en supposant un modèle parfait), en substituant la matrice des erreurs d'observation par $\mathbf{R}_u^{(j)} = \mathbf{R} + \mathbf{H} \mathbf{Q} \mathbf{H}^T$.

À partir de là, $\mathbf{v}$ peut être obtenu en une unique itération. En effet, (2.150) nous donne $\mathbf{u}^* = (\mathbf{H} \mathbf{M}^* \mathbf{X}_1^f)^T (\mathbf{R} + \mathbf{H} \mathbf{Q} \mathbf{H}^T)^{-1} [\mathbf{y}_2 - \mathcal{H} \circ \mathcal{M}(\mathbf{x}_1^*)]$. Par conséquent, au minimum $\mathbf{v}^*$, l'équation (2.147) devient

$$\begin{aligned} \mathbf{v}^* = & (\mathbf{H} \mathbf{X}^q)^T \left[\mathbf{R} + \mathbf{H} \mathbf{Q} \mathbf{H}^T + (\mathbf{H} \mathbf{M}^* \mathbf{X}_1^f)(\mathbf{H} \mathbf{M}^* \mathbf{X}_1^f)^T \right]^{-1} \\ & \times \left[\mathbf{I} + (\mathbf{H} \mathbf{M}^* \mathbf{X}_1^f)(\mathbf{H} \mathbf{M}^* \mathbf{X}_1^f)^T (\mathbf{R} + \mathbf{H} \mathbf{Q} \mathbf{H}^T)^{-1} \right] [\mathbf{y}_2 - \mathcal{H} \circ \mathcal{M}(\mathbf{x}_1^*)] \end{aligned} \quad (2.152)$$

$$= (\mathbf{H} \mathbf{X}^q)^T (\mathbf{R} + \mathbf{H} \mathbf{Q} \mathbf{H}^T)^{-1} [\mathbf{y}_2 - \mathcal{H} \circ \mathcal{M}(\mathbf{x}_1^*)]. \quad (2.153)$$

On peut également atteindre cette conclusion en suivant un formalisme bayésien, ce qui permet de relier l'IEnKF-Q au filtre particulaire de Doucet *et al.* (2000), qui est une solution élégante à notre problème originel avec des applications en géosciences (Bocquet *et al.*, 2010; Snyder *et al.*, 2015; Slivinski et Snyder, 2016).

On a en effet la pdf de l'analyse $p(\mathbf{x}_1, \mathbf{x}_2 | \mathbf{y}_2)$ qui est reliée à la fonction de coût (2.123) de la sorte

$$\mathcal{J}(\mathbf{x}_1, \mathbf{x}_2) = -\ln p(\mathbf{x}_1, \mathbf{x}_2 | \mathbf{y}_2). \quad (2.154)$$

En appliquant la définition d'une probabilité conditionnelle, on transforme cette pdf en

$$p(\mathbf{x}_1, \mathbf{x}_2 | \mathbf{y}_2) = p(\mathbf{x}_2 | \mathbf{x}_1, \mathbf{y}_2) p(\mathbf{x}_1 | \mathbf{y}_2). \quad (2.155)$$

Il s'avère que dans le cas $\mathcal{H}$ linéaire, ces deux membres du produit ont une expression analytique. On peut ainsi montrer que

$$-2 \ln p(\mathbf{x}_1 | \mathbf{y}_2) = \|\mathbf{x}_1 - \bar{\mathbf{x}}_f\|_{\mathbf{P}_f}^2 + \|\mathbf{y}_2 - \mathcal{H} \circ \mathcal{M}(\mathbf{x}_1)\|_{\mathbf{R} + \mathbf{H}\mathbf{Q}\mathbf{H}^T}^2 + c_1 \quad (2.156)$$

et que

$$\begin{aligned} -2 \ln p(\mathbf{x}_2 | \mathbf{x}_1, \mathbf{y}_2) = & \|\mathbf{x}_2 - \mathcal{M}(\mathbf{x}_1) - \mathbf{Q}\mathbf{H}^T(\mathbf{R} + \mathbf{H}\mathbf{Q}\mathbf{H}^T)^{-1} \\ & \times [\mathbf{y}_2 - \mathcal{H} \circ \mathcal{M}(\mathbf{x}_1)]\|_{(\mathbf{Q}^{-1} + \mathbf{H}^T\mathbf{R}^{-1}\mathbf{H})^{-1}}^2 + c_2, \end{aligned} \quad (2.157)$$

où c_1 et c_2 sont des constantes indépendantes de $\mathbf{x}_1$ et $\mathbf{x}_2$. Nous précisons que $p(\mathbf{x}_1 | \mathbf{y}_2)$ n'est pas gaussienne, contrairement à $p(\mathbf{x}_2 | \mathbf{x}_1, \mathbf{y}_2)$ qui l'est grâce à la linéarité de $\mathbf{H}$.

Cette décomposition montre à nouveau le découplage de la minimisation de $\mathcal{J}(\mathbf{x}_1, \mathbf{x}_2)$ en deux étapes. On peut d'abord minimiser $-2 \ln p(\mathbf{x}_1 | \mathbf{y}_2)$ selon la variable $\mathbf{x}_1$ pour obtenir le maximum a posteriori (MAP) $\mathbf{x}_1^*$. Ensuite, cela nous permet d'obtenir directement le MAP de $-2 \ln p(\mathbf{x}_2 | \mathbf{x}_1^*, \mathbf{y}_2)$, donné par

$$\mathbf{x}_2^* = \mathcal{M}(\mathbf{x}_1^*) + \mathbf{Q}\mathbf{H}^T(\mathbf{R} + \mathbf{H}\mathbf{Q}\mathbf{H}^T)^{-1} [\mathbf{y}_2 - \mathcal{H} \circ \mathcal{M}(\mathbf{x}_1^*)]. \quad (2.158)$$

On retombe alors sur l'expression de $\mathbf{x}_2^*$ obtenue à partir de (2.153). On comprend donc d'un point de vue bayésien les raisons du découplage du calcul de l'optimum de l'IEnKF-Q. Cependant, ce découplage, valable pour le MAP, ne s'étend pas au calcul des anomalies.

2.5.4.3 Algorithmes

En se basant sur les équations (2.135), (2.140) et (2.142), on peut développer un algorithme pour l'IEnKF-Q. En l'occurrence, l'algorithme 8 présenté correspond à une version *transform*.

La ligne 5 de l'algorithme correspond à l'équation (2.125); la ligne 6 calcule la transformation de l'ensemble. Elle diffère de l'expression (2.142) car nous ne nous intéressons pas avec cet algorithme aux anomalies lissées, mais seulement à l'évaluation du filtrage. Les coefficients $\sqrt{m-1}$ des lignes 7, 9, 10, 20 et 21 permettent le passage des anomalies normalisées aux anomalies non-normalisées; dans le cas d'un $\mathcal{H}$ linéaire, la ligne 10 se simplifie en $\mathbf{H}\mathbf{X}^q = \mathcal{H}(\mathbf{X}^q)$.

La transformation de l'ensemble appliquée à la ligne 7 est un peu restrictive, même si elle est pratique. Son éventuelle sous-optimalité est évidente car la transformation s'applique correctement au modèle de propagation de l'ensemble mais pas à l'opérateur d'observation. Dans ce contexte, si on devait appliquer véritablement le principe de *transform* de Sakov *et al.* (2012), il faudrait appliquer la matrice de transformation $\mathbf{T} = \mathbf{D}^{1/2}$ à la matrice des anomalies jointes $[\mathbf{X}_1^f, \mathbf{X}^q]$, avant d'appliquer la fonction de l'espace de l'ensemble à celui des observations :

$$\mathbf{w} \rightarrow \mathcal{H}(\mathcal{M}(\mathbf{x}_1^f + \mathbf{X}_1^f \mathbf{w}_{1:m}) + \mathbf{X}^q \mathbf{w}_{m+1:m+m_q}) \quad (2.159)$$

Algorithme 8 Une version *transform* de l'IEnKF-Q, avec minimisation de Gauss-Newton. Les parties de l'algorithme en rouge correspondent aux changements par rapport à un IEnKF *transform* standard, en l'absence d'erreur modèle. “SR($\mathbf{X}, m$)” fait référence à une réduction de la taille de l'ensemble de $m + m_q$ à m .

Requis: $m, m_q, \mathbf{E}_1^f, \mathbf{X}^q, \mathbf{y}_2, \mathbf{R}, \mathcal{M}, \mathcal{H}, j_{\max}, e$

```

1:  $\mathbf{x}_1^f = \mathbf{E}_1^f \mathbf{1}/m$ 
2:  $\mathbf{X}_1^f = (\mathbf{E}_1^f - \mathbf{x}_1^f \mathbf{1}^T)/\sqrt{m-1}$ 
3:  $\mathbf{D} = \mathbf{I}_{m+m_q}, \quad \mathbf{w} = \mathbf{0}$ 
4: repeat
5:    $\mathbf{x}_1 = \mathbf{x}_1^f + \mathbf{X}_1^f \mathbf{w}_{1:m}$ 
6:    $\mathbf{T} = (\mathbf{D}_{1:m, 1:m})^{1/2}$ 
7:    $\mathbf{E}_1 = \mathbf{x}_1 \mathbf{1}^T + \mathbf{X}_1^f \mathbf{T} \sqrt{m-1}$ 
8:    $\mathbf{E}_2 = \mathcal{M}(\mathbf{E}_1)$ 
9:    $\mathbf{H}\mathbf{X}_2 = \mathcal{H}(\mathbf{E}_2)(\mathbf{I}_m - \mathbf{1}\mathbf{1}^T/m) \mathbf{T}^{-1}/\sqrt{m-1}$ 
10:   $\mathbf{H}\mathbf{X}^q = \mathcal{H}(\mathbf{E}_2 \mathbf{1}\mathbf{1}^T/m + \mathbf{X}^q \sqrt{m_q-1})(\mathbf{I}_{m_q} - \mathbf{1}\mathbf{1}^T/m_q)/\sqrt{m_q-1}$ 
11:   $\mathbf{H}\mathbf{X} = [\mathbf{H}\mathbf{X}_2, \mathbf{H}\mathbf{X}^q]$ 
12:   $\mathbf{x}_2 = \mathbf{E}_2 \mathbf{1}/m + \mathbf{X}^q \mathbf{w}_{m+1:m+m_q}$ 
13:   $\nabla \mathcal{J} = \mathbf{w} - (\mathbf{H}\mathbf{X})^T \mathbf{R}^{-1}[\mathbf{y}_2 - \mathcal{H}(\mathbf{x}_2)]$ 
14:   $\mathbf{D} = [\mathbf{I}_{m+m_q} + (\mathbf{H}\mathbf{X})^T \mathbf{R}^{-1} \mathbf{H}\mathbf{X}]^{-1}$ 
15:   $\Delta \mathbf{w} = -\mathbf{D} \nabla \mathcal{J}$ 
16:   $\mathbf{w} := \mathbf{w} + \Delta \mathbf{w}$ 
17:   $j := j + 1$ 
18: until  $j \geq j_{\max}$  or  $\|\Delta \mathbf{w}\| < e$ 
19:  $\mathbf{X}_2 = \mathbf{E}_2 (\mathbf{I}_m - \mathbf{1}\mathbf{1}^T/m) \mathbf{T}^{-1}$ 
20:  $\mathbf{X} = [\mathbf{X}_2/\sqrt{m-1}, \mathbf{X}^q] \mathbf{D}^{1/2}$ 
21:  $\mathbf{X}_2 = \text{SR}(\mathbf{X}, m)\sqrt{m-1}$ 
22:  $\mathbf{E}_2 = \mathbf{x}_2 \mathbf{1}^T + \mathbf{X}_2$ 

```

L'implémentation de cette transformation jointe est moins évidente que celle offerte par la ligne 7, qui revient à utiliser les anomalies lissées marginalisées à t_1 . Une autre possibilité est de choisir $\mathbf{T} = [\mathbf{D}^{1/2}]_{1:m,1:m}$ qui est encore une matrice définie positive. Nous avons vérifié que ces trois approches fournissent des résultats quantitatifs similaires. Cependant, nous supposons que la transformation jointe mentionnée ci-dessus pourrait fournir de meilleurs résultats en présence d'un opérateur d'observation fortement non-linéaire (expérience non testée).

Étant donné que l'IEEnKF-Q utilise des anomalies augmentées, la taille de l'ensemble est agrandie de m à $m + m_q$ à la ligne 11. Ainsi, pour retourner à la taille originelle de l'ensemble, il faut effectuer une réduction à la fin du cycle. Si la taille de l'ensemble m est strictement supérieure à la taille du modèle, cette réduction peut être effectuée sans perte d'information ; sinon, de l'information sera perdue. Cette réduction a lieu à la ligne 21 de l'algorithme.

Une première façon de réduire l'ensemble consiste à ne garder que les $m - 1$ premiers membres de $\mathbf{X}_2^a$ et conserver le degré de liberté restant pour centrer les anomalies autour de zéro. Même si cette démarche semble réductrice, elle s'avère suffisante dans le cas particulier d'un modèle linéaire et de matrices $\mathbf{R}$ et $\mathbf{Q}$ diagonales. Numériquement, quand cette réduction de l'ensemble a été appliquée, les résultats ne se sont avérés sous-optimaux qu'en présence de fortes non-linéarités.

Une autre façon de réduire l'ensemble consiste à appliquer une décomposition en valeurs singulières à $\mathbf{X}_2^a$, garder les $m - 1$ plus grandes directions de $\mathbf{X}_2^a$ et conserver le degré de liberté restant pour centrer les anomalies autour de zéro. En pratique, la norme des vecteurs engendrés par cette procédure est assez peu uniforme, ce qui peut détériorer les performances de l'algorithme. Une nette sous-performance a notamment été observée dans les expériences avec une faible erreur modèle. Par conséquent, on peut avoir à appliquer des perturbations stochastiques conservant la moyenne à l'ensemble pour le rendre plus gaussien.

Le coût de calcul supplémentaire de l'algorithme engendré par m_q est dû à l'application de l'opérateur d'observation aux m_q membres supplémentaires. Cependant, il peut être considéré comme négligeable par rapport au coût de propagation des m membres de l'ensemble, qui reste souvent le nœud du problème. Par ailleurs, un m_q de grande taille peut également augmenter le coût de l'optimisation de la fonctionnelle (2.130), du fait de l'inversion de matrices de taille $(m + m_q) \times (m + m_q)$. Dans des applications réelles, on pourra essayer de contenir la taille de m_q en pointant $\mathbf{X}^q$ sur les principales directions de l'erreur modèle, connues a priori, comme par exemple les forçages.

Une version localisée de cet algorithme peut également être développée, en s'inspirant de l'algorithme 7 de l'IEEnKS local qu'on aurait adapté en version *transform*. L'IEEnKF-Q localisé, qui n'est pas introduit dans Sakov *et al.* (2017, soumis), est ainsi présenté à l'algorithme 9. Nous rappelons que dans ce contexte local, la notation m_e est utilisée pour faire référence à la taille de l'ensemble et m désigne l'indice de chacune des M analyses locales. Notons que nous considérons un ensemble de M matrices d'erreur modèles différentes, dont nous avons les racines carrées $\{\mathbf{X}_m^q\}_{m \in [1, M]}$. Nous précisons également que cet algorithme, dans sa version *transform*, requiert l'inversion des M matrices $\mathbf{T}^m \in \mathbb{R}^{m_e \times m_e}$, ce qui est coûteux.

2.5.4.4 Discussion

La présence d'erreur modèle en assimilation de données bruite la transmission d'information dans le temps. Ainsi, les observations ont moins d'impact sur l'estimation d'un état loin en temps par rapport au cas sans erreur modèle ; réciproquement, les observations au même temps ont relativement plus d'impact. Concernant le deuxième point, les erreurs d'ébauche d'un EnKF sont augmentées de la matrice $\mathbf{Q}$, donnant ainsi plus d'importance aux observations. Concernant le premier point, il est justifié par le cas découplé, où l'analyse lissée est obtenue comme si le modèle

Algorithme 9 Une version *transform* de l'IEKF-Q localisé, avec minimisation de Gauss-Newton. “SR($\mathbf{X}, m_e$)” fait référence à une réduction de la taille de l'ensemble de $m_e + m_q$ à m_e , où m_e est la taille de l'ensemble.

Requis: $m_e, m_q, M, \mathbf{E}_1^f, \{\mathbf{X}_m^q\}_{m \in [1, M]}, \mathbf{y}_2, \mathbf{R}, \mathcal{M}, \mathcal{H}, j_{\max}, e$

```

1:  $\mathbf{x}_1^f = \mathbf{E}_1^f \mathbf{1} / m_e$ 
2:  $\mathbf{X}_1^f = (\mathbf{E}_1^f - \mathbf{x}_1^f \mathbf{1}^T) / \sqrt{m_e - 1}$ 
3: for  $m = 1 \dots M$  do
4:    $\mathbf{D}^m = \mathbf{I}_{m_e + m_q}, \quad \mathbf{w}^m = \mathbf{0}$ 
5: end for
6: repeat
7:   for  $m = 1 \dots M$  do
8:      $[\mathbf{x}_1]_m = [\mathbf{x}_1^f]_m + [\mathbf{X}_1^f \mathbf{w}_{1:m_e}^m]_m$ 
9:      $\mathbf{T}^m = ([\mathbf{D}^m]_{1:m_e, 1:m_e})^{1/2}$ 
10:     $[\mathbf{E}_1]_m = [\mathbf{x}_1 \mathbf{1}^T]_m + [\mathbf{X}_1^f \mathbf{T}^m]_m \sqrt{m_e - 1}$ 
11:  end for
12:   $\mathbf{E}_2 = \mathcal{M}(\mathbf{E}_1)$ 
13:  for  $m = 1 \dots M$  do
14:     $[\mathbf{x}_2]_m = [\mathbf{E}_2 \mathbf{1} / m_e]_m + [\mathbf{X}_m^q \mathbf{w}_{m_e+1:m_e+m_q}^m]_m$ 
15:  end for
16:  for  $m = 1 \dots M$  do
17:     $\mathbf{H}\mathbf{X}_2^m = \mathcal{H}(\mathbf{E}_2)(\mathbf{I}_{m_e} - \mathbf{1}\mathbf{1}^T / m_e)(\mathbf{T}^m)^{-1} / \sqrt{m_e - 1}$ 
18:     $\mathbf{H}\mathbf{X}_m^q = \mathcal{H}(\mathbf{E}_2 \mathbf{1}\mathbf{1}^T / m_e + \mathbf{X}_m^q \sqrt{m_q - 1})(\mathbf{I}_{m_q} - \mathbf{1}\mathbf{1}^T / m_q) / \sqrt{m_q - 1}$ 
19:     $\mathbf{H}\mathbf{X}^m = [\mathbf{H}\mathbf{X}_2^m, \mathbf{H}\mathbf{X}_m^q]$ 
20:     $\nabla \mathcal{J}^m = \mathbf{w}^m - (\mathbf{H}\mathbf{X}^m)^T (\mathbf{R}^m)^{-1} [\mathbf{y}_2 - \mathcal{H}(\mathbf{x}_2)]$ 
21:     $\mathbf{D}^m = [\mathbf{I}_{m_e + m_q} + (\mathbf{H}\mathbf{X}^m)^T (\mathbf{R}^m)^{-1} \mathbf{H}\mathbf{X}^m]^{-1}$ 
22:     $\Delta \mathbf{w}^m = -\mathbf{D}^m \nabla \mathcal{J}^m$ 
23:     $\mathbf{w}^m := \mathbf{w}^m + \Delta \mathbf{w}^m$ 
24:     $j := j + 1$ 
25:  end for
26: until  $j \geq j_{\max}$  or  $\sum_{m=1}^M \|\Delta \mathbf{w}^m\| < e \times M$ 
27: for  $m = 1 \dots M$  do
28:    $\mathbf{X}_2^m = \mathbf{E}_2 (\mathbf{I}_{m_e} - \mathbf{1}\mathbf{1}^T / m_e) (\mathbf{T}^m)^{-1}$ 
29:    $\mathbf{X}^m = [\mathbf{X}_2^m / \sqrt{m_e - 1}, \mathbf{X}_m^q] (\mathbf{D}^m)^{1/2}$ 
30:    $\mathbf{X}_2^m = \text{SR}(\mathbf{X}^m, m_e) \sqrt{m_e - 1}$ 
31:    $[\mathbf{E}_2]_m = [\mathbf{x}_2 \mathbf{1}^T]_m + [\mathbf{X}_2^m]_m$ 
32: end for

```

était parfait, mais avec cette fois-ci les erreurs d'observation augmentées de $\mathbf{H}\mathbf{Q}\mathbf{H}^T$.

En pratique, le concept d'erreur modèle additive est rarement applicable directement, pour deux raisons. Tout d'abord, la majorité des erreurs rencontrées comme les erreurs de forçage, de représentativité, de paramétrisation ou dans les équations de base, sont non-additives par nature. Ensuite, même si les erreurs sont additives, elles peuvent être difficiles à quantifier. Si l'erreur modèle ajoutée n'a pas la bonne amplitude, on risque de se retrouver avec une analyse sous-optimale. C'est pour cela que l'on a supposé ici que les statistiques de l'erreur modèle étaient connues. On pourrait malgré tout également les estimer par des méthodes comme celles de [Todling \(2015a\)](#).

Néanmoins, l'erreur modèle additive représente tout de même un aspect théorique important. Dans le 4D-Var par exemple, elle peut-être utilisée empiriquement pour régulariser la minimisation qui devient instable pour de très grandes fenêtres ([Blayo et al. \(2015\)](#), p. 451). L'erreur modèle pourrait être également définie comme un paramètre ajustable d'un système d'assimilation de données, au même titre que l'inflation par exemple. Ceci peut d'ailleurs s'avérer bénéfique dans certains cas d'applications de l'EnKF ([Whitaker et al. \(2008\)](#) par exemple). Les méthodes hybrides présentées à la partie 2.5.1 s'apparentent également à une utilisation d'erreur modèle additive.

2.6 Validation et indicateurs statistiques

Pour valider les résultats d'expériences de modélisation ou d'assimilation de données, il est nécessaire d'avoir des outils et indicateurs statistiques. Nous en présentons quelques uns que nous utiliserons par la suite.

2.6.1 Erreur quadratique : RMSE

La racine de l'erreur quadratique moyenne (RMSE) représente l'erreur quadratique entre un échantillon d'estimations fourni par un modèle ou une expérience d'assimilation de données et des valeurs observées ou réelles. Si on note $\mathbf{y} = (y_1, \dots, y_{N_t})$ une série temporelle d'observations à une station donnée et $\mathbf{x} = (x_1, \dots, x_{N_t})$ les valeurs estimées à cette station par un modèle ou une expérience d'assimilation de données, alors on définit la RMSE comme

$$\text{RMSE} = \sqrt{\frac{1}{N_t} \sum_{i=1}^{N_t} (x_i - y_i)^2}. \quad (2.160)$$

Si on a plusieurs stations d'observation à notre disposition, on peut alors à nouveau fournir la moyenne ou les centiles sur ces différentes stations.

Dans le cas d'une expérience jumelle, on pourra similairement définir une RMSE à chaque pas de temps, et la moyenniser sur l'ensemble des N_t cycles d'assimilation. Ainsi, si on note $\mathbf{x}_t^i \in \mathbb{R}^n$ la vérité à un temps i et $\mathbf{x}_a^i \in \mathbb{R}^n$ l'analyse à ce même temps, alors la RMSE moyenne est

$$\text{RMSE} = \frac{1}{N_t} \sum_{i=1}^{N_t} \sqrt{\frac{1}{n} \|\mathbf{x}_a^i - \mathbf{x}_t^i\|^2}, \quad (2.161)$$

où $\|\mathbf{x}\|^2 = \sum_{j=1}^n x_j^2$.

Enfin, dans le cas d'une expérience avec un lisseur, on pourra distinguer la RMSE de filtrage, calculée au temps présent, de la RMSE de lissage, calculée dans le passé, par exemple au début de la fenêtre. Si on illustre cela avec les mêmes notations et avec L la longueur de la fenêtre d'assimilation, on a alors

$$\text{RMSE}^{\text{filtrage}} = \frac{1}{N_t} \sum_{i=1}^{N_t} \sqrt{\frac{1}{n} \|\mathcal{M}_{L \leftarrow 0}(\mathbf{x}_a^i) - \mathbf{x}_t^{i+L}\|^2} \quad (2.162)$$

et

$$\text{RMSE}^{\text{lissage}} = \frac{1}{N_t} \sum_{i=1}^{N_t} \sqrt{\frac{1}{n} \|\mathbf{x}_a^i - \mathbf{x}_t^i\|^2}. \quad (2.163)$$

2.6.2 Corrélation de Pearson

A nouveau, si on note $\mathbf{y} = (y_1, \dots, y_{N_t})$ une série temporelle d'observations à une station donnée et $\mathbf{x} = (x_1, \dots, x_{N_t})$ les valeurs estimées à cette station par un modèle ou une expérience d'assimilation de données, alors on définit la corrélation de Pearson comme

$$C_r = \frac{\sum_{k=1}^{N_t} (x_k - \bar{x}) \times (y_k - \bar{y})}{\sqrt{\left(\sum_{k=1}^{N_t} (x_k - \bar{x})^2\right) \left(\sum_{k=1}^{N_t} (y_k - \bar{y})^2\right)}}, \quad (2.164)$$

où $\bar{x} = \frac{1}{N_t} \sum_{k=1}^{N_t} x_k$. Ce coefficient renseigne sur le degré de dépendance linéaire entre les deux variables $\mathbf{x}$ et $\mathbf{y}$.

2.6.3 Biais

Il existe plusieurs versions du biais, entre autres :

- le biais moyen : $MB = \frac{1}{N_t} \sum_{k=1}^{N_t} (x_k - y_k)$;
- l'erreur moyenne : $ME = \frac{1}{N_t} \sum_{k=1}^{N_t} |x_k - y_k|$;
- le biais fractionnel moyen : $MFB = \frac{1}{N_t} \sum_{k=1}^{N_t} \frac{(x_k - y_k)}{\left(\frac{x_k + y_k}{2}\right)}$;
- l'erreur fractionnelle moyenne : $MFE = \frac{1}{N_t} \sum_{k=1}^{N_t} \frac{|x_k - y_k|}{\left(\frac{x_k + y_k}{2}\right)}$.

2.6.4 Indicateurs statistiques χ^2 et RCRV

On a défini les variables $\mathbf{x}_f$, $\mathbf{y}$ et leurs erreurs associées $\boldsymbol{\varepsilon}_f$, $\boldsymbol{\varepsilon}_o$ de la sorte

$$\mathbf{x}_f = \mathbf{x}_t + \boldsymbol{\varepsilon}_f, \quad \text{avec} \quad \mathbb{E}[\boldsymbol{\varepsilon}_f \boldsymbol{\varepsilon}_f^T] = \mathbf{P}^f, \quad (2.165)$$

$$\mathbf{y} = \mathbf{H}\mathbf{x}_t + \boldsymbol{\varepsilon}_o, \quad \text{avec} \quad \mathbb{E}[\boldsymbol{\varepsilon}_o \boldsymbol{\varepsilon}_o^T] = \mathbf{R}, \quad (2.166)$$

où $\mathbb{E}$ désigne l'espérance. Ainsi, on a pour les innovations

$$\boldsymbol{\delta} = \mathbf{y} - \mathbf{H}\mathbf{x}_f = \boldsymbol{\varepsilon}_o - \mathbf{H}\boldsymbol{\varepsilon}_f, \quad (2.167)$$

et donc

$$\mathbb{E}[\boldsymbol{\delta} \boldsymbol{\delta}^T] = \mathbf{R} + \mathbf{H} \mathbf{P}^f \mathbf{H}^T, \quad (2.168)$$

puisque les erreurs d'ébauche et d'observation sont indépendantes, c'est-à-dire $\mathbb{E}[\boldsymbol{\varepsilon}_o \boldsymbol{\varepsilon}_f] = \mathbf{0}$.

La méthode du χ^2 consiste à s'assurer que les statistiques des innovations sont justes, ce qui nous donne une indication sur la qualité de la définition des matrices $\mathbf{R}$ et $\mathbf{P}^f$ conjointement. Il suffit pour cela de calculer

$$\chi^2 = \text{Tr} \left[(\mathbf{R} + \mathbf{H}\mathbf{P}^f\mathbf{H}^T)^{-1} \boldsymbol{\delta}\boldsymbol{\delta}^T \right], \quad (2.169)$$

et vérifier que χ^2 soit le plus proche possible de p , le nombre d'observations (Ménard *et al.*, 2000; Koohkan et Bocquet, 2012). Si $\chi^2 < p$ ($> p$), on surestime (sous-estime) les matrices $\mathbf{R}$ et/ou $\mathbf{P}^f$, et on peut revoir la définition de $\mathbf{R}$ ou l'inflation et la localisation appliquée à $\mathbf{P}^f$, voire la taille de l'ensemble. Nous précisons que le fait que l'on ait un unique critère pour déterminer deux paramètres rend le système sous-déterminé.

De plus, dans le contexte de l'EnKF, la matrice $\mathbf{P}^f$ est définie à partir de l'ensemble de taille réduit. On a vu qu'elle était donc approchée et nécessitait de recourir à de la localisation. Dans le cas de la localisation par produit de Schur, il suffit d'utiliser dans la formule la matrice localisée, à savoir $(\mathbf{R} + \mathbf{H}(\rho \circ \mathbf{P}^f)\mathbf{H}^T)$. Dans le cas d'une localisation par domaine, les statistiques sont un peu différentes. Si la localisation est effectuée avec un schéma *boxcar*, on peut appliquer la formule (2.169) dans le domaine réduit autour du point d'analyse. Si en revanche, une fonction de régularisation comme celle de GC est appliquée, elle transforme la matrice $\mathbf{R}$ en une matrice $\tilde{\mathbf{R}}$ locale, sachant que les statistiques de $[\boldsymbol{\delta}\boldsymbol{\delta}^T]$ ne changent pas. Aussi, la formule à appliquer localement devient

$$\chi^2 = \text{Tr} \left[(\tilde{\mathbf{R}} + \mathbf{H}\mathbf{P}^f\mathbf{H}^T)^{-1} (\boldsymbol{\delta}\boldsymbol{\delta}^T + \tilde{\mathbf{R}} - \mathbf{R}) \right], \quad (2.170)$$

qu'on peut réaliser en remplaçant l'innovation $\boldsymbol{\delta}$ par $\tilde{\boldsymbol{\delta}} = \boldsymbol{\delta} + \boldsymbol{\varepsilon}$ avec $\boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, \tilde{\mathbf{R}} - \mathbf{R})$.

À partir du même constat (2.168), on peut également définir le reduced centered random variable (RCRV), qui se concentre uniquement sur les termes diagonaux de la matrice $\mathbf{R} + \mathbf{H}\mathbf{P}^f\mathbf{H}^T$. Pour chaque observation y_i du vecteur $\mathbf{y}$ ($i \in [1, p]$), on peut calculer

$$\text{RCRV}_i = \frac{y_i - [\mathbf{H}\mathbf{x}_f]_i}{\sqrt{\sigma_o^2 + \sigma_f^2}}, \quad (2.171)$$

où σ_o^2 est le $i^{\text{ème}}$ terme de la diagonale de $\mathbf{R}$ et σ_f^2 est calculé avec chaque membre de l'ensemble $\mathbf{x}_f^j$ pour $j \in [1, m]$,

$$\sigma_f^2 = \frac{1}{m-1} \sum_{j=1}^M [\mathbf{H}\mathbf{x}_f^j]_i^2. \quad (2.172)$$

Une fois les p différents RCRV_i obtenus, on vérifie que leur moyenne est nulle et que leur variance est proche de 1.

Le diagnostic du χ^2 est présenté en tant que diagnostic sur les innovations dans Desroziers *et al.* (2005). Cet article démontre également plusieurs autres diagnostics similaires, fondés sur l'espérance des différences entre observations, ébauches et analyses, qui permettent d'obtenir séparément $\mathbf{R}$ ou $\mathbf{H}\mathbf{P}^f\mathbf{H}^T$. Dans le cas d'une expérience avec présence d'erreur modèle, on pourra également se reporter à Todling (2015a), qui montre comment estimer l'erreur modèle $\mathbf{H}\mathbf{Q}\mathbf{H}^T$ à partir des différences entre observations, ébauche, analyse et analyses lissées repropagées. On insistera également sur le fait que l'existence de diagnostics séparés pour l'estimation de $\mathbf{R}$, $\mathbf{H}\mathbf{P}^f\mathbf{H}^T$ ou $\mathbf{H}\mathbf{Q}\mathbf{H}^T$ ne prévaut en rien sur la sous-détermination du système, comme le présente la note Todling (2015b).

Développement de modèles jouets pour le transport et la chimie atmosphérique

Nous avons décrit au chapitre 1 les équations physiques régissant l'évolution de certaines variables météorologiques, les équations chimiques expliquant l'évolution des concentrations de l'ozone troposphérique et des polluants qui lui sont associés, et enfin comment coupler ces deux jeux d'équations avec l'équation du transport réactif (1.20). Nous avons vu pour chacun de ces jeux d'équations les problèmes qu'ils apportent : chaoticté, non-linéarité, recours à des schémas numériques d'intégration implicites et d'ordre élevé pour gérer la raideur des équations, approximations comme l'approximation des états quasi-stationnaires (QSSA) ou le *splitting*, nécessaires pour éviter une complexité numérique grandissante. Enfin, nous avons listé certaines des sources d'erreurs et d'incertitudes qui sont inhérentes à de tels modèles.

Nous avons décrit au chapitre 2 plusieurs méthodes d'assimilation de données. Pour chacune d'entre elles, nous avons décrit les hypothèses simplificatrices nécessaires à leurs dérivations et leurs contextes d'application optimaux. En outre, les avantages et défauts majeurs des méthodes variationnelles et d'ensemble, qui ont été listés en début de section 2.5, ont amené le développement des méthodes variationnelles d'ensemble qui sont encore en voie de développement à ce jour.

La complexité des modèles, aussi bien de météorologie que les modèle de chimie-transport (CTM) et encore plus les modèles couplés de chimie-météorologie (CCMMs), pose soucis pour tester de nouvelles méthodes d'assimilation de données. En effet, si une amélioration est observée lors de l'application de ces nouvelles méthodes sur un modèle complexe, il peut être difficile d'en expliquer la raison, car de très nombreux phénomènes peuvent intervenir, interagir, se compenser et s'amplifier mutuellement.

Ainsi, il est nécessaire de développer des modèles d'ordre réduit, dit modèles jouets, dont la physique est parfaitement comprise et maîtrisée, sur lesquels de nouvelles méthodes d'assimilation de données peuvent être testées. On veut par ce biais développer des bancs d'essai sur lesquels comparer les différentes méthodes et comprendre les gains en performance qu'elles peuvent apporter dans des contextes d'application variés. Par ailleurs, la taille réduite des modèles permet d'effectuer des simulations avec un temps de calcul limité et d'obtenir des statistiques stables grâce à un très grand nombre de cycles d'assimilation de données. On peut même tester des méthodes dont le coût de calcul sur des modèles complets serait prohibitif, ce qui permet d'avoir une performance de référence optimale vers laquelle on peut chercher à tendre avec des méthodes approchées mais numériquement moins coûteuses.

Nous débuterons donc ce chapitre en présentant le modèle Lorenz-95 (L95), qui modélise approximativement un phénomène météorologique, avant d'en introduire une première extension, le modèle L95-T, qui couple le modèle de Lorenz avec un tracer passif chimiquement. Enfin, nous ajouterons à ce modèle un module de chimie de l'ozone troposphérique, appelé *generic reaction set* (GRS), pour produire le modèle L95-GRS. Ce modèle, ainsi que sa validation en tant que banc d'essai pour des méthodes d'assimilation de données, ont fait l'objet d'une publication (Haussaire et Bocquet, 2016), et il est disponible sur le site <http://cerea.enpc.fr/195-grs/index.html>. Nous le

comparons ensuite succinctement au QG-Chem, qui a été développé avec les mêmes objectifs par Emili *et al.* (2016).

Sommaire

3.1	Un premier modèle simplifié de météorologie : L95	61
3.2	Un modèle avec transport d'un traceur inerte : L95-T	61
3.3	Introduction de la chimie atmosphérique : L95-GRS	63
3.3.1	Définition du L95-GRS	64
3.3.1.1	La chimie du GRS	64
3.3.1.2	Couplage avec le L95	65
3.3.1.3	Intégration numérique	66
3.3.1.4	Choix des paramètres	67
3.3.2	Validation du modèle	67
3.3.2.1	Évolution temporelle des concentrations	67
3.3.2.2	Isopleths	69
3.3.2.3	Transport transcontinental	69
3.3.2.4	Régimes météorologiques	70
3.3.3	Intérêt et nécessité d'un tel modèle	71
3.4	Un bref aperçu de QG-Chem	72
3.4.1	La météorologie avec un modèle quasi-géostrophique	72
3.4.2	La chimie atmosphérique avec RACM	72
3.4.3	Comparaison avec le L95-GRS	73

Les quatre modèles présentés dans ce chapitre sont issus, dans l'ordre, de

- E. N. LORENZ et K. A. EMANUEL : Optimal sites for supplementary weather observations: Simulation with a small model. *J. Atmos. Sci.*, 55:399–414, 1998,
- M. BOCQUET et P. SAKOV : Joint state and parameter estimation with an iterative ensemble Kalman smoother. *Nonlinear Process. Geophys.*, 20:803–818, 2013
- J.-M. HAUSSAIRE et M. BOCQUET : A low-order coupled chemistry meteorology model for testing online and offline data assimilation schemes: L95-GRS (v1. 0). *Geosci. Model Dev.*, 9:393–412, 2016
- E. EMILI, S. GÜROL et D. CARIOLLE : Accounting for model error in air quality forecasts: an application of 4DEnVar to the assimilation of atmospheric composition using QG-Chem 1.0. *Geosci. Model Dev.*, 9:3933–3959, 2016

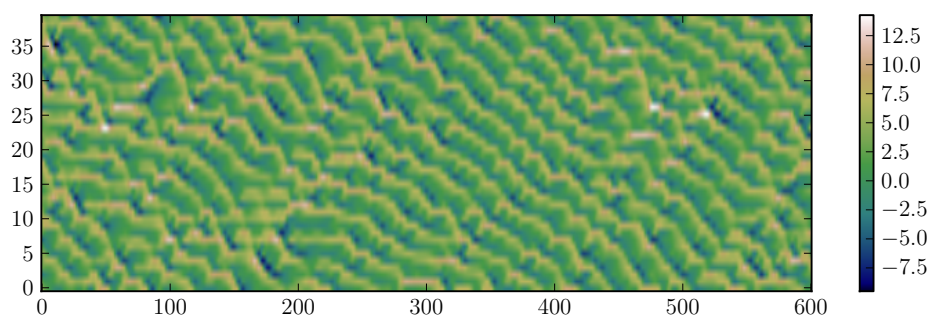


FIGURE 3.1 – Champs des 40 variables du modèle de Lorenz-95 en fonction du temps (en unité de $\Delta t = 0.05$).

3.1 Un premier modèle simplifié de météorologie : L95

Lorenz et Emanuel (1998) ont introduit un modèle chaotique devenu populaire en assimilation de données, que l'on notera Lorenz-95 (L95). C'est un modèle à une dimension dont les $n = 40$ variables s'étendent sur un cercle au niveau des latitudes moyennes. Les variables $\{x_i\}_{i \in [1, n]}$ suivent l'équation aux dérivées ordinaires

$$\frac{dx_i}{dt} = (x_{i+1} - x_{i-2})x_{i-1} - x_i + F. \quad (3.1)$$

Le domaine est périodique, de telle sorte que $x_{-1} = x_{39}$, $x_0 = x_{40}$ et $x_{41} = x_1$. On choisit le terme de forçage $F = 8$. Un pas de temps de $\Delta t = 0.05$ représente un intervalle de 6 h dans l'atmosphère (1 unité de temps de Lorenz correspond donc à 5 jours). Ce modèle est intégré numériquement avec un algorithme de Runge-Kutta d'ordre 4 avec ce pas de temps de $\Delta t = 0.05$. Les champs des variables du modèle sont représentés en fonction du temps sur la Fig. 3.1. Avec ces paramètres, la dynamique du modèle est chaotique avec un temps de doublement de 0.42 unité de temps et 13 exposants de Lyapunov positifs et 1 nul. Cela signifie que deux trajectoires du modèle initialisées avec une légère perturbation, vont diverger rapidement, comme l'illustre la Fig. 3.2.

Bien que ce modèle soit simple et phénoménologique plutôt qu'une dérivation rigoureuse des équations de l'atmosphère, son évolution ressemble à des ondes de Rossby que l'on retrouve dans la troposphère ; il a une taille limitée avec seulement 40 variables et ne nécessite qu'une intégration numérique avec un schéma explicite de type Runge-Kutta d'ordre 4 ; la physique de ce modèle est bien comprise, on peut en calculer les exposants de Lyapunov, calculer la vitesse de phase et de groupe des ondes. Il réunit donc toutes les qualités nécessaires pour servir de banc d'essai aux méthodes d'assimilation de données et s'est imposé comme modèle de choix parmi la communauté pour tester ces méthodes (par exemple, Anderson, 2001; Whitaker et Hamill, 2002; Ott *et al.*, 2004; Hunt *et al.*, 2004; van Leeuwen, 2010; Bocquet, 2011; Sakov *et al.*, 2012; Palatella *et al.*, 2013).

3.2 Un modèle avec transport d'un traceur inerte : L95-T

Le L95 s'apparente à un modèle de météorologie et ses variables peuvent représenter une direction et une amplitude de vents. Aussi peut-il être couplé à un traceur passif pour former un premier embryon de CCMM. C'est dans cet objectif que le L95-T a été développé dans Bocquet et Sakov (2013).

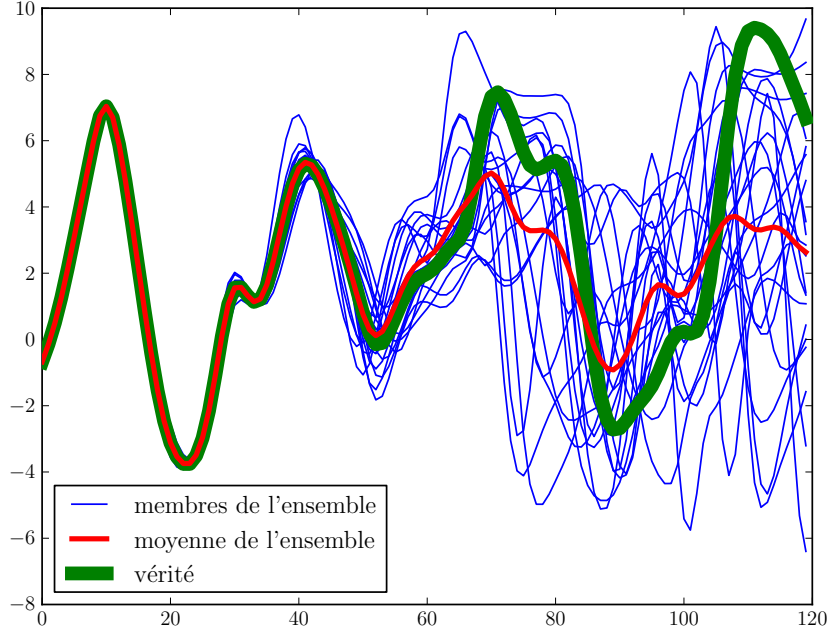


FIGURE 3.2 – Chaoticité du modèle illustrée : évolution d'une variable du Lorenz-95 en fonction du temps (en unité de Δt) et divergence rapide des modèles initialisés avec une légère perturbation.

Le champ du traceur est discrétisé selon 40 variables additionnelles, notées $c_{m+\frac{1}{2}}$ pour $m \in [1, 40]$, qui représentent les concentrations du traceur, réparties sur une grille-C d'Arakawa. Le traceur est advecté par le vent que représente le L95 avec un schéma *upwind*, comme défini dans le chapitre 1. Les équations du modèle couplé sont donc les suivantes :

$$\frac{dx_m}{dt} = (x_{m+1} - x_{m-2})x_{m-1} - x_m + F, \quad (3.2)$$

$$\frac{dc_{m+\frac{1}{2}}}{dt} = \psi_m - \psi_{m+1} - \lambda c_{m+\frac{1}{2}} + E_{m+\frac{1}{2}}, \quad (3.3)$$

$$\text{avec } \psi_m = x_m c_{m-\frac{1}{2}} \quad \text{si } x_m \geq 0, \quad (3.4)$$

$$\psi_m = x_m c_{m+\frac{1}{2}} \quad \text{si } x_m < 0. \quad (3.5)$$

Le traceur est émis tout autour du domaine et les flux sont notés $E_{m+\frac{1}{2}}$. Il est également déposé tout au long du domaine avec un coefficient de dépôt unique λ . Les valeurs de référence de ces paramètres choisies dans Bocquet et Sakov (2013) sont $E_{m+\frac{1}{2}} = E = 1$ et $\lambda = 0.1$. Nous noterons qu'un point stationnaire du modèle existe, avec $x_m = F$ et $c_{m+\frac{1}{2}} = E/\lambda$, ce qui donne une idée de l'ordre de grandeur des variables de vent et de concentration. Nous remarquons également que la condition de Courant-Frierichs-Lewy (CFL) est respectée (voir 1.4.1.2). En effet, le vent, en valeur absolue, ne dépasse que rarement 15 et on a $\Delta t = 0.05$ et $\Delta x = 1$ par construction.

Ce modèle est de nouveau intégré avec un schéma de Runge-Kutta d'ordre 4 et aucun *splitting* n'est appliqué. Les champs des 80 variables du modèle sont représentés en fonction du temps sur la Fig. 3.3.

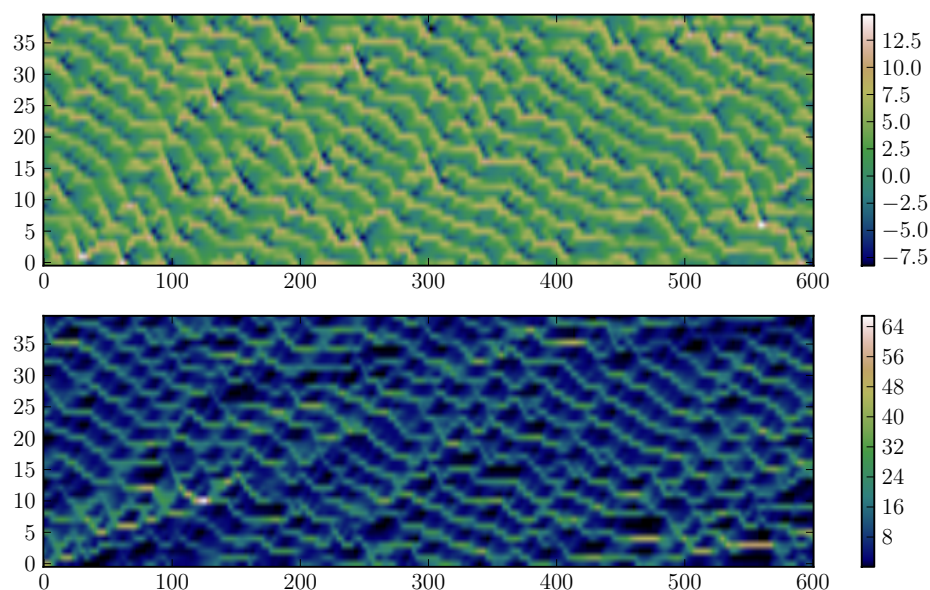


FIGURE 3.3 – Champs des 40 variables de vent (en haut) et des 40 variables de concentration (en bas) du modèle L95-T en fonction du temps (en unité de $\Delta t = 0.05$).

Comme la météorologie est simulée conjointement aux concentrations du traceur, il s'agit d'un modèle *online*, ou fortement couplé. Mais le couplage de ce modèle est à sens unique : les vents impactent les concentrations en les advectant, mais il n'y a pas d'influence physique des concentrations sur les vents. En revanche, l'assimilation de données fera échanger les informations dans les deux sens grâce aux covariances entre les sous-systèmes du traceur et de la météorologie. Notamment, des observations de concentrations pourront aider à estimer les variables météorologiques.

Ce modèle présente beaucoup de points communs avec des vrais CCMMs. En ce qui concerne la dynamique, on a avec ce modèle le couplage d'une météorologie chaotique avec un transport non chaotique, comme avec n'importe quel modèle de traceur réaliste. Les incertitudes sur la météorologie proviennent des erreurs sur la condition initiale qui grandissent en raison de la chaoticité. Les incertitudes sur le traceur viennent des champs d'émissions et des incertitudes sur le vent, mais comme la dynamique est stable, ces erreurs ne croissent pas.

3.3 Introduction de la chimie atmosphérique : L95-GRS

Le modèle L95-T est un premier embryon de CCMM. Cependant, pour avoir un véritable modèle d'ordre réduit représentatif d'un CCMM, il faut lui adjoindre une chimie non-linéaire. Le L95-GRS en résulte. En effet, ce modèle complète le L95-T en remplaçant le traceur, passif chimiquement, par un mécanisme simple de photochimie, nommé *generic reaction set* (GRS) (Azzi *et al.*, 1992). Nous allons présenter ce modèle dans cette section, en insistant sur ses caractéristiques et sa validation en tant que modèle simplifié, mais valide, de photochimie. Ce modèle est disponible en téléchargement sur le site <http://cerea.enpc.fr/l95-grs/index.html>.

3.3.1 Définition du L95-GRS

3.3.1.1 La chimie du GRS

Le module de chimie utilisé est le GRS de [Azzi *et al.* \(1992\)](#). Il consiste en 7 espèces chimiques supposées représenter la chimie de l’ozone troposphérique à partir de concentrations urbaines de composés organiques volatiles (COV) et de NO_x . Les réactions associées à ces espèces sont



où ROC représente les composés organiques réactifs (reactive organic compounds), RP sont les radicaux (radical pool), SGN sont les produits de réaction nitrogénés gazeux et SNGN sont les produits de réaction nitrogénés non-gazeux. Les ROC de ce module de chimie sont donc une forme d’espèce suppléante pour l’ensemble des COV d’une atmosphère réelle. Ce jeu d’équations englobe, de manière synthétique, l’ensemble des phénomènes complexes du cycle de Leighton présentés dans la section 1.2.1. Les constantes de cinétique chimique sont tirées de [Venkatram *et al.* \(1994\)](#) et valent

$$k_1 = 10\,000 \times e^{-\frac{4710}{T}} \times k_3 \quad (3.6)$$

$$k_2 = 5.482 \times e^{\frac{242}{T}} \text{ ppb}^{-1} \text{ min}^{-1}, \quad (3.7)$$

$$k_4 = 2.643 \times e^{-\frac{1370}{T}} \text{ ppb}^{-1} \text{ min}^{-1}, \quad (3.8)$$

$$k_5 = 10.2 \text{ ppb}^{-1} \text{ min}^{-1}, \quad (3.9)$$

$$k_6 = 0.12 \text{ ppb}^{-1} \text{ min}^{-1}, \quad (3.10)$$

$$k_7 = 0.12 \text{ ppb}^{-1} \text{ min}^{-1}, \quad (3.11)$$

où T est la température, choisie constante égale à 300 K par simplicité et k_3 est le coefficient de photolyse du NO_2 en min^{-1} , qui est une fonction de la luminosité.

Étant donné que $k_6 = k_7$, les réactions (R3.6) et (R3.7) sont similaires et on peut combiner les deux espèces en une espèce suppléante en utilisant la constante de réaction $0.24 \text{ ppb}^{-1} \text{ min}^{-1}$. Cette espèce suppléante représente la formation de tous les composés nitrogénés stables, qu’ils soient gazeux ou non. On a donc la réaction



Pour réduire davantage la taille du GRS et ainsi réduire sa complexité numérique, on utilise la QSSA (1.4.2.3) sur les radicaux RP, qui sont hautement réactifs et ont la plus courte durée de

vie parmi les espèces du GRS. Les radicaux sont alors considérés à l'équilibre dynamique, ce qui implique

$$0 = \frac{d[\text{RP}]}{dt} = k_1[\text{ROC}] - [\text{RP}] (k_2[\text{NO}] + 2k_6[\text{NO}_2] + k_5[\text{RP}]). \quad (3.12)$$

Après résolution de cette équation algébrique du second degré, on obtient la formule

$$[\text{RP}] = \frac{k_2[\text{NO}] + 2k_6[\text{NO}_2]}{2k_5} \left(\sqrt{1 + \frac{4k_1k_5[\text{ROC}]}{(k_2[\text{NO}] + 2k_6[\text{NO}_2])^2}} - 1 \right). \quad (3.13)$$

La validation de cette hypothèse a été faite a posteriori et il s'est avéré qu'elle n'avait que très peu d'impact sur les concentrations simulées, tout en permettant de fortement augmenter le pas de temps d'intégration moyen.

3.3.1.2 Couplage avec le L95

Le GRS est couplé au L95 de la même façon que le traceur du L95-T. Les variables du L95 sont considérées comme des vitesses et directions du vent qui advectent les espèces chimiques du GRS. Le but est alors de créer un modèle qui reproduit la variabilité temporelle et les ordres de grandeur de l'ozone troposphérique et de ses précurseurs. Après couplage, on a donc un total de 40 variables météorologiques et 200 variables de concentrations, à savoir les concentrations de ROC, NO, NO₂, O₃ et S(N)GN aux 40 points de grille définis selon la grille-C d'Arakawa. Rappelons que les concentrations de RP sont définies explicitement selon l'équation (3.13).

L'équation du transport réactif pour une variable $[\text{C}_i]_{m+\frac{1}{2}}$, représentant la concentration de l'espèce C_i au point $m + \frac{1}{2}$, est alors

$$\frac{d[\text{C}_i]_{m+\frac{1}{2}}}{dt} = \psi_m^i - \psi_{m+1}^i - \lambda_i[\text{C}_i]_{m+\frac{1}{2}} + E_{m+\frac{1}{2}}^i + R_i([\text{C}_j]_{j=1,\dots,6}), \quad (3.14)$$

$$\text{avec } \psi_m^i = x_m[\text{C}_i]_{m-\frac{1}{2}} \quad \text{si } x_m \geq 0, \quad (3.15)$$

$$\psi_m^i = x_m[\text{C}_i]_{m+\frac{1}{2}} \quad \text{si } x_m < 0, \quad (3.16)$$

où λ_i est le coefficient de dépôt, $E_{m+\frac{1}{2}}^i$ est le taux d'émission et $R_i([\text{C}_j]_{j=1,\dots,6})$ est le terme de production et de destruction chimique pour l'espèce C_i . Précisons que ce dernier terme est nul pour le ROC, de telle sorte que cette espèce chimique se comporte comme le traceur du L95-T. Les équations d'évolution des différentes espèces sont plus précisément

$$\frac{d[\text{ROC}]_{m+\frac{1}{2}}}{dt} = \psi_m^{\text{ROC}} - \psi_{m+1}^{\text{ROC}} - \lambda[\text{ROC}]_{m+\frac{1}{2}} + E_{m+\frac{1}{2}}^{\text{ROC}}, \quad (3.17)$$

$$\begin{aligned} \frac{d[\text{RP}]_{m+\frac{1}{2}}}{dt} = & \psi_m^{\text{RP}} - \psi_{m+1}^{\text{RP}} - \lambda[\text{RP}]_{m+\frac{1}{2}} + k_1[\text{ROC}]_{m+\frac{1}{2}} \\ & - [\text{RP}]_{m+\frac{1}{2}} \left(k_2[\text{NO}]_{m+\frac{1}{2}} + 2k_6[\text{NO}_2]_{m+\frac{1}{2}} + k_5[\text{RP}]_{m+\frac{1}{2}} \right), \end{aligned} \quad (3.18)$$

$$\begin{aligned} \frac{d[\text{NO}]_{m+\frac{1}{2}}}{dt} = & \psi_m^{\text{NO}} - \psi_{m+1}^{\text{NO}} - \lambda[\text{NO}]_{m+\frac{1}{2}} + E_{m+\frac{1}{2}}^{\text{NO}} \\ & + k_3[\text{NO}_2]_{m+\frac{1}{2}} - [\text{NO}]_{m+\frac{1}{2}} \left(k_2[\text{RP}]_{m+\frac{1}{2}} + k_4[\text{O}_3]_{m+\frac{1}{2}} \right), \end{aligned} \quad (3.19)$$

$$\begin{aligned} \frac{d[\text{NO}_2]_{m+\frac{1}{2}}}{dt} = & \psi_m^{\text{NO}_2} - \psi_{m+1}^{\text{NO}_2} - \lambda[\text{NO}_2]_{m+\frac{1}{2}} + E_{m+\frac{1}{2}}^{\text{NO}_2} + k_4[\text{NO}]_{m+\frac{1}{2}}[\text{O}_3]_{m+\frac{1}{2}} \\ & + k_2[\text{NO}]_{m+\frac{1}{2}}[\text{RP}]_{m+\frac{1}{2}} - [\text{NO}_2]_{m+\frac{1}{2}} \left(k_3 + 2k_6[\text{RP}]_{m+\frac{1}{2}} \right), \end{aligned} \quad (3.20)$$

$$\frac{d[\text{O}_3]_{m+\frac{1}{2}}}{dt} = \psi_m^{\text{O}_3} - \psi_{m+1}^{\text{O}_3} - \lambda[\text{O}_3]_{m+\frac{1}{2}} + k_3[\text{NO}_2]_{m+\frac{1}{2}} - k_4[\text{NO}]_{m+\frac{1}{2}}[\text{O}_3]_{m+\frac{1}{2}}, \quad (3.21)$$

$$\frac{d[\text{S(N)GN}]_{m+\frac{1}{2}}}{dt} = \psi_m^{\text{S(N)GN}} - \psi_{m+1}^{\text{S(N)GN}} - \lambda[\text{S(N)GN}]_{m+\frac{1}{2}} + 2k_6[\text{NO}_2]_{m+\frac{1}{2}}[\text{RP}]_{m+\frac{1}{2}}. \quad (3.22)$$

L'équation 3.18 est notée à titre indicative dans le cas où la QSSA n'est pas utilisée.

Si on considère le L95-GRS comme un modèle global, alors le coefficient de photolyse k_3 , qui dépend de l'intensité lumineuse, doit varier tout autour du domaine et en fonction de la saison. Il devrait alors y avoir des points de grille où il fait jour, avec $k_3 \neq 0$, et d'autres où il fait nuit avec $k_3 = 0$. Pour simplifier, k_3 est par la suite fixé sur l'intégralité du domaine et il varie selon un cycle journalier prédéfini.

Étant donné que le ROC n'est pas consommé par la réaction (R3.1), il produira suffisamment de RP pour consommer l'intégralité du NO, NO₂ ou O₃ présent. Il a donc été nécessaire d'avoir un terme d'émission pour le ROC et les NO_x, ainsi qu'un coefficient de dépôt pour toutes les espèces. On a considéré les émissions comme constantes en temps et en espace, sauf dans le cas où une distinction entre continent et océan sera faite (voir 3.3.2.3).

3.3.1.3 Intégration numérique

Comme pour des modèles de chimie atmosphérique réalistes, il est utile ici d'appliquer un *splitting* entre la partie météorologie/transport et la partie chimie. Ainsi, comme pour le L95-T, la météorologie du L95, le transport, le dépôt et les émissions sont intégrés avec un schéma de Runge-Kutta d'ordre 4, tandis que la chimie nécessite le recours à un schéma de Rosenbrock d'ordre 2 pour les raisons mentionnées dans la section 1.4.2.4. Étant donné que la chimie est locale grâce au *splitting*, ce schéma peut être appliqué par bloc, avec un point de grille par bloc. Les systèmes linéaires à résoudre dans l'algorithme sont simplement de la taille du nombre d'espèces, qui est ici limité. Ces résolutions peuvent par ailleurs être parallélisées. Enfin, le tangent linéaire du modèle requis par le schéma de Rosenbrock peut être dérivé analytiquement et simplement grâce au nombre limité de réactions.

Pour finir, un pas de temps adaptatif, comme présenté en 1.4.2.5, a été mis en place, même s'il s'avère souvent inutile dans des simulations sans assimilation de données et avec la QSSA. On choisit un pas de temps de référence de $\delta t = 1$ h pour intégrer la chimie, tout comme pour la partie météorologie et transport du L95-GRS.

t	k_3	t	k_3	t	k_3	t	k_3
00 h	0	06 h	0.1972314	12 h	0.622824	18 h	0.00675528
01 h	0	07 h	0.3910734	13 h	0.611526	19 h	0
02 h	0	08 h	0.5074326	14 h	0.5755002	20 h	0
03 h	0	09 h	0.5755002	15 h	0.5074326	21 h	0
04 h	0	10 h	0.611526	16 h	0.3910734	22 h	0
05 h	0.00675528	11 h	0.622824	17 h	0.1972314	23 h	0

TABLEAU 3.1 – Valeurs horaires de k_3 en min^{-1} .

3.3.1.4 Choix des paramètres

Nous résumons ici l'ensemble des paramètres choisis pour le L95-GRS. Pour commencer, nous fixons le coefficient de dépôt λ constant sur l'intégralité du domaine, et choisissons la valeur $\lambda = 0.02 \text{ jour}^{-1}$ pour toutes les espèces. Cette valeur est la même que celle utilisée pour le L95-T, puisqu'une unité de temps de Lorenz correspond à 5 jours.

En ce qui concerne les émissions, nous avons choisi un coefficient pour les NO_x de $0.27 \text{ ppb} \cdot \text{jour}^{-1}$, dont 10% sont du NO_2 et 90% sont du NO. Ceci correspond à une émission de $3 \text{ kg an}^{-1} \text{ habitant}^{-1}$ pour 60 millions d'habitants dans un volume de $700 \text{ km} \times 700 \text{ km} \times 3 \text{ km}$. Cette valeur est donc proche de la situation d'un pays comme la France. Le dépôt et les émissions de NO_x étant fixés, nous avons pu déterminer les valeurs des émissions de ROC permettant d'obtenir des concentrations réalistes d' O_3 et de NO_x . Ainsi, la valeur retenue pour les émissions de ROC est $0.0235 \text{ ppb C jour}^{-1}$.

Il reste à déterminer le coefficient de photolyse k_3 du NO_2 , qui est une fonction du rayonnement solaire. Nous avons utilisé Fast-JX (Voulgarakis *et al.*, 2009) pour obtenir les valeurs de rayonnement. Nous avons choisi les valeurs à l'équateur du 21 mars sans atténuation, que nous répétons quotidiennement. Si k_3 est requis entre deux heures, on effectue une interpolation linéaire. Les valeurs exactes des coefficients horaires sont répertoriés dans le tableau 3.1.

Enfin, nous rappelons que le pas de temps choisi pour l'intégration du modèle est de $\delta t = 1 \text{ h}$. Nous récapitulons l'ensemble de ces paramètres dans le tableau 3.2.

3.3.2 Validation du modèle

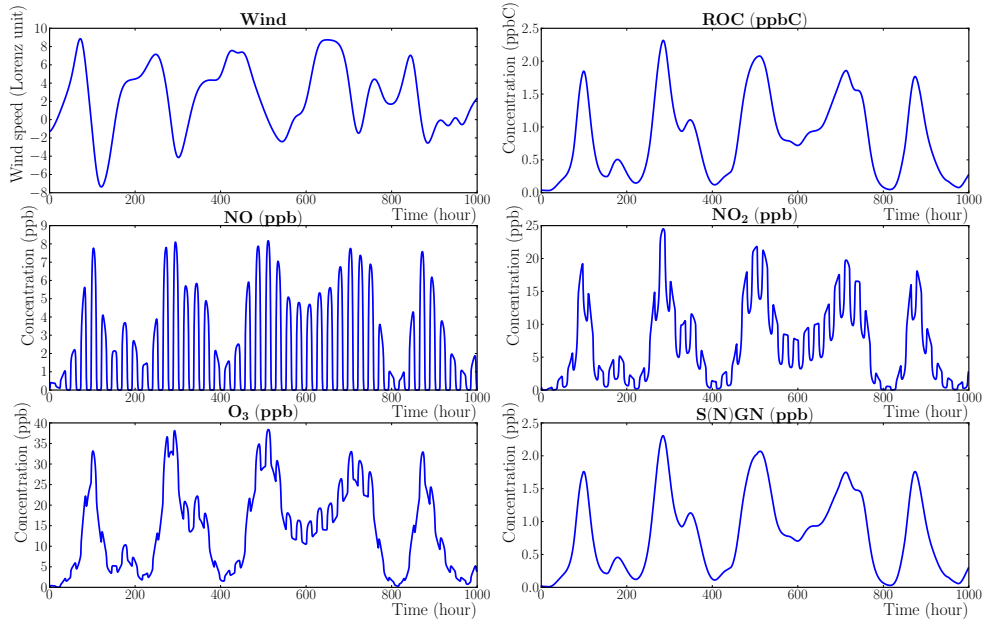
Il est nécessaire de valider la qualité physique de ce modèle pour déterminer s'il produit des concentrations de polluants réalistes et si certains des phénomènes observés en situation réelle peuvent être reproduits avec ce modèle. Notamment, on va dans un premier temps vérifier l'ordre de grandeur et l'évolution temporelle des concentrations, puis s'assurer que l'on retrouve les mêmes formes d'isopleths que celles présentées à la section 1.2.2. Enfin, on s'intéressera plus précisément à la dynamique du transport des polluants par la météorologie et les possibilités de différents régimes qu'offre le L95.

3.3.2.1 Évolution temporelle des concentrations

On trace l'évolution temporelle des concentrations des différents polluants en un point de grille, après *spin-up*, sur la Fig. 3.4. On peut y voir le cycle diurne sur les concentrations de NO, NO_2 et O_3 imposé par la variation du coefficient de photolyse k_3 , puisque ces trois variables en dépendent directement via les équations (R 3.3) et (R 3.4). Les vitesses et orientations du vent produisent quant à elles l'onde de période environ une semaine.

paramètre	valeur	unité
Nombre de points de grille	40	
Nombres total de variables	240	
Forçage du Lorenz F	8	unité de Lorenz
Dépôt λ	0.02	jour ⁻¹
Émissions de ROC E_{ROC}	0.0235	ppb C jour ⁻¹
Émissions de NO _x E_{NO_x}	0.27	ppb jour ⁻¹
Émissions de NO E_{NO}	$E_{\text{NO}_x} \times 0.9 = 0.243$	ppb jour ⁻¹
Émissions de NO ₂ E_{NO_2}	$E_{\text{NO}_x} \times 0.1 = 0.027$	ppb jour ⁻¹
Température T	300	K
Constante cinétique k_1	$0.00152 \times k_3$	min ⁻¹
Constante cinétique k_2	12.3	ppb ⁻¹ min ⁻¹
Constante cinétique k_3	Fast-JX le 21 mars à l'équateur	min ⁻¹
Constante cinétique k_4		ppb ⁻¹ min ⁻¹
Constante cinétique k_5	10.2	ppb ⁻¹ min ⁻¹
Constante cinétique k_6	0.12	ppb ⁻¹ min ⁻¹
Pas de temps d'intégration	1	h

TABLEAU 3.2 – Récapitulation des paramètres utilisés pour le L95-GRS.

FIGURE 3.4 – Évolution temporelle des variables du L95-GRS en un point de grille. Les variables du L95, notées *Wind*, sont montrées dans l'unité du Lorenz, alors que les concentrations sont en ppb (ppbC pour le ROC). Figure tirée de Haussaire et Bocquet (2016).

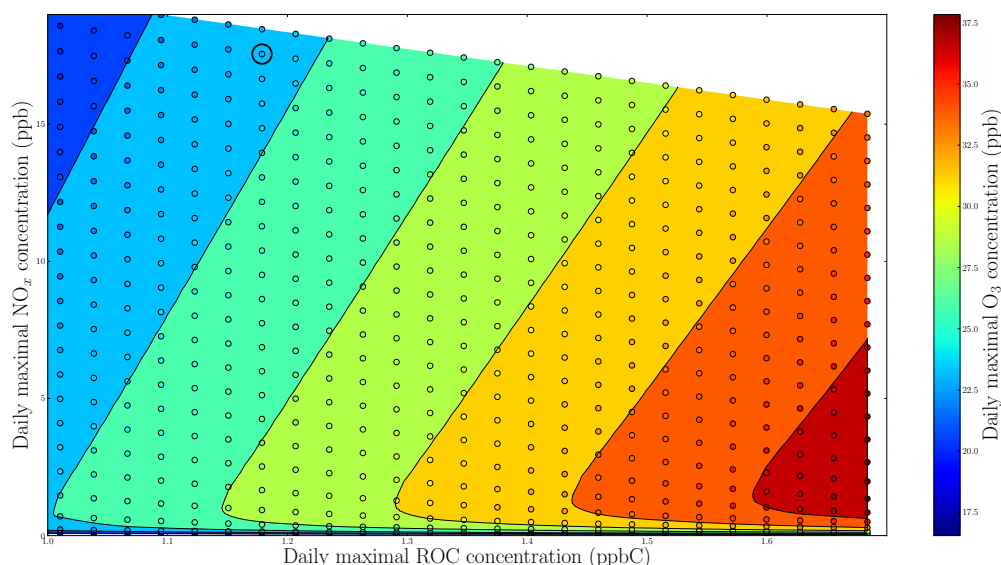


FIGURE 3.5 – Concentration maximale d’ozone (en ppb) moyennée sur le domaine en fonction des concentrations moyennes maximales de ROC (en ppbC) et de NO_x (en ppb). Chaque point correspond à une simulation avec un différent taux d’émission pour le ROC et les NO_x , générant une concentration moyenne maximale différente. La simulation de référence est entourée en noir. Figure tirée de Haussaire et Bocquet (2016).

Dans des situations réelles, les concentrations d’ozone et de NO peuvent tomber très bas pendant la nuit (Finlayson-Pitts et Pitts Jr, 1986). Avec le L95-GRS, seul le NO tombe à zéro la nuit, tandis que l’ozone reste à des niveaux élevés. L’inverse peut se produire si on diminue suffisamment les émissions de ROC ou augmente beaucoup celles de NO_x .

Ainsi, les niveaux de concentrations obtenus pour l’ozone et les NO_x sont dans l’ordre de grandeur de ce qui peut être observé en zone urbaine par exemple, avec des épisodes de forte pollution qui peuvent être provoqués par des conditions météorologiques favorables.

3.3.2.2 Isopleths

Ce modèle, même s’il n’est pas chaotique, est fortement non-linéaire. Un moyen de s’en convaincre est de tracer les isopleths d’ozone en fonction des concentrations moyennes (sur le domaine) de ROC et de NO_x . Ces isopleths sont tracées sur la Fig. 3.5. On y voit parfaitement les deux régimes chimiques distincts présentés dans la section 1.2.2. On rappelle qu’ici, nos ROC jouent le rôle des COV.

Étant donné que le cercle noir correspond à notre situation de référence, nous nous sommes placé dans un régime ROC-limité. Dans le cas NO_x -limité, les faibles concentrations de NO réduisent fortement l’amplitude du cycle diurne de l’ozone. On se retrouve ainsi avec une situation plus proche d’une simulation d’ozone de fond. Étant donné que le GRS a été proposé pour simuler des concentrations typiques d’un régime urbain, nous avons préféré rester dans le domaine ROC-limité.

3.3.2.3 Transport transcontinental

Pour mettre en exergue l’impact du transport, nous avons effectué une simulation où le domaine est coupé en deux : une partie, dite continentale, sur laquelle les émissions de ROC et NO_x sont non nulles ($E_i > 0$ pour $i \in [1, 20]$), et une partie, dite océanique, sur laquelle les émissions sont

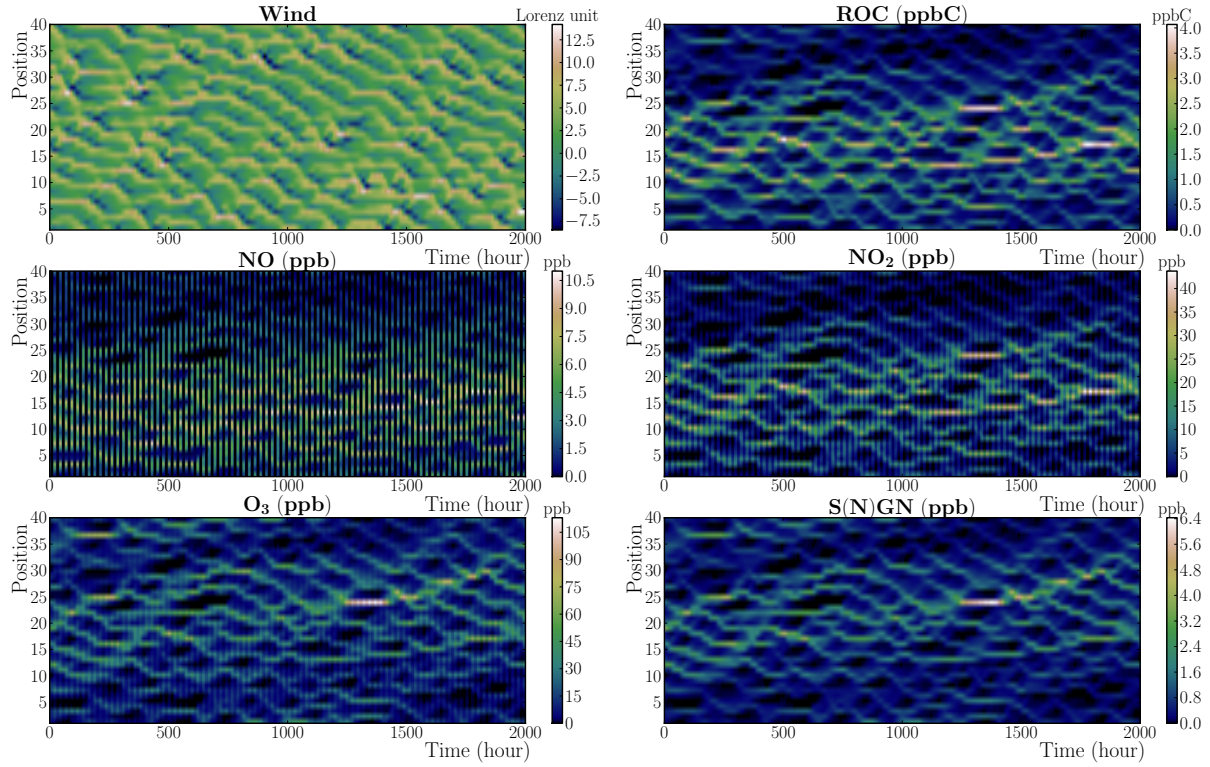


FIGURE 3.6 – Évolution temporelle des variables du L95-GRS sur l'intégralité du domaine dans le contexte d'une division continent/océan. Les variables du L95, notées *Wind*, sont montrées dans l'unité temporelle du Lorenz, alors que les concentrations sont en ppb (ppb C pour le ROC). Figure tirée de Haussaire et Bocquet (2016).

nulles ($E_i = 0$ pour $i \in [21, 40]$). L'évolution temporelle des concentrations sur l'ensemble des 40 points de grille du domaine est présentée sur la Fig. 3.6.

On constate que des bouffées d'ozone et de ses précurseurs peuvent traverser l'océan, comme on peut parfois l'observer au-dessus du Pacifique (Lin *et al.*, 2012) ou de l'Atlantique (Guerova *et al.*, 2006). De plus, on remarque que les concentrations d'ozone sont supérieures au dessus de l'océan, en l'absence d'émissions de NO. Par ailleurs, on notera que les bouffées voyagent vers l'est, en cohérence avec une vitesse de groupe du L95 positive, tandis que les crêtes du champ des variables de L95 voyagent vers l'ouest, étant donnée la vitesse de phase négative du L95 (Lorenz et Emanuel, 1998). Enfin, l'effet du rayonnement est visible avec les bandes noires sur le NO, quand les concentrations tombent à 0 la nuit. Dans le cas où un k_3 différent en chaque point de grille est utilisé, ces mêmes bandes sont obliques, sans impacter la variation des concentrations à grande échelle temporelle.

3.3.2.4 Régimes météorologiques

Jusqu'à maintenant, les seuls paramètres du L95-GRS qui ont été modulés concernent le GRS. Il est cependant possible de jouer sur le L95 pour changer l'échelle de temps d'évolution du L95. Dans la simulation de référence, la période des oscillations du L95 s'étendait sur plusieurs jours. Les concentrations étaient déterminées par cette cinétique du vent d'une part et un cycle journalier apparaissait d'autre part grâce à la photochimie. Pour changer cela, on peut remarquer que dans

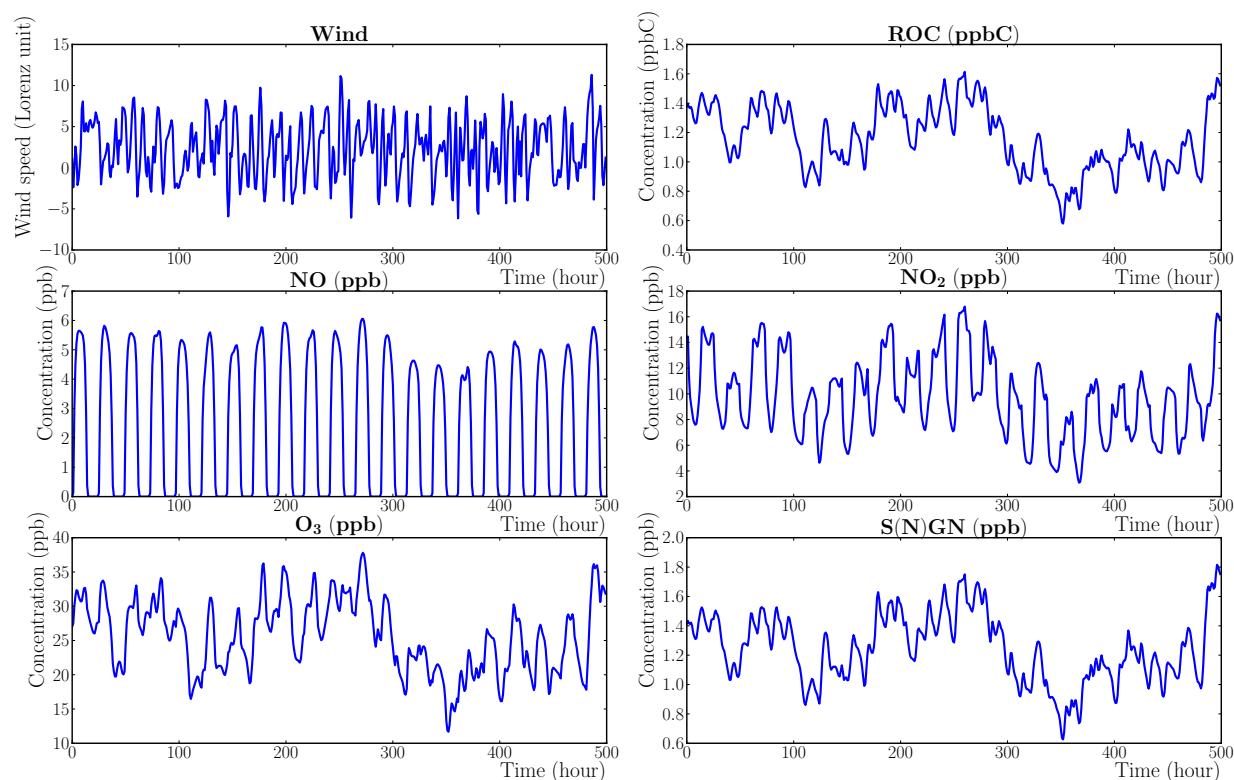


FIGURE 3.7 – Évolution temporelle des variables du L95-GRS en un point de grille, avec un changement d'échelle $\alpha = 20$ sur le L95. Les variables du L95, notées *Wind*, sont montrées dans l'unité temporelle du Lorenz, alors que les concentrations sont en ppb (ppb C pour le ROC). Figure tirée de Haussaire et Bocquet (2016).

l'équation du L95, si l'échelle de temps est modifiée par un facteur α et les variables de vent par β , l'équation devient

$$\frac{dx_m}{dt} = \alpha \left[\beta(x_{m+1} - x_{m-2})x_{m-1} - x_m + \frac{F}{\beta} \right]. \quad (3.23)$$

Par ce biais, la période d'oscillation du vent peut être réduite de telle sorte à correspondre à celle de la chimie et ainsi l'impact des variations du vent sur les concentrations peut être amplifié. Cela correspondrait alors à une modélisation à l'échelle régionale plutôt que continentale voire globale. La Fig. 3.7 montre ce changement d'échelle avec $\alpha = 20$ et $\beta = 1$.

On voit que le vent fluctue plus rapidement que les concentrations, à l'opposé du comportement de la simulation de référence (Fig. 3.4). On pourrait également ajuster β afin de réduire l'amplitude du vent pour avoir un transport similaire à celui en basse troposphère.

On précise qu' α et β sont simplement des paramètres de changement d'échelle qui n'impactent pas la nature de la dynamique du modèle, contrairement aux expériences effectuées dans Carrassi *et al.* (2008).

3.3.3 Intérêt et nécessité d'un tel modèle

La définition du modèle est simple, avec un schéma chimique dont le nombre de réactions et d'espèces est limité, permettant d'avoir un modèle avec peu de variables et d'en maîtriser parfaitement l'évolution. Cependant, le modèle fait preuve d'un comportement réaliste et a des propriétés fidèles

à l'évolution atmosphérique de la pollution à l'ozone. Par ailleurs, il possède un nombre conséquent de paramètres, notamment les émissions, le dépôt, le coefficient de photolyse ou l'échelle des vents, permettant de modifier le modèle pour produire un comportement donné souhaité. Ainsi, ce modèle est à la fois simple, réaliste et modulable. Il possède donc toutes les qualités requises pour étudier l'application de méthodes d'assimilation de données sur des modèles non-linéaires de chimie atmosphérique.

De tels modèles simplifiés sont nécessaires étant donnée la complexité croissante des modèles opérationnels ou universitaires de chimie atmosphérique d'une part et des méthodes d'assimilation de données d'autre part. Avec un modèle dont le comportement est parfaitement maîtrisé, l'accent peut être mis sur les effets de la méthode d'assimilation de données et son apport dans le contexte proposé par le modèle. Ceci explique notre démarche pour justifier la création du L95-GRS et c'est cette même motivation qui a conduit au QG-Chem de [Emili *et al.* \(2016\)](#) que nous présentons brièvement.

3.4 Un bref aperçu de QG-Chem

La similarité de la démarche proposée par [Emili *et al.* \(2016\)](#) nous pousse à succinctement présenter le modèle QG-Chem pour voir les différences entre ce modèle et le L95-GRS.

3.4.1 La météorologie avec un modèle quasi-géostrophique

La météorologie est tirée d'un modèle quasi-géostrophique (QG) à deux niveaux verticaux, qui représente de la dynamique aux latitudes moyennes à méso-échelle. Les équations du modèle pour les deux couches (indice 1 pour la couche du bas et 2 pour celle du haut) sont tirées de l'équation de conservation de la vorticité potentielle $\mathbf{q} = (q_1, q_2)$

$$\frac{dq_i}{dt} = 0, \quad (3.24)$$

qui s'écrit avec des variables adimensionnées

$$q_1 = \Delta\psi_1 - F_1(\psi_1 - \psi_2) + \beta y, \quad (3.25)$$

$$q_2 = \Delta\psi_2 - F_2(\psi_2 - \psi_1) + \beta y + R_s, \quad (3.26)$$

où Δ est le Laplacien, R_s représente l'orographie, β est la variation vers le nord (adimensionnée) du paramètre de Coriolis à une latitude fixée et F_1 et F_2 couplent les couches ensemble en étant une fonction de la force de Coriolis, la hauteur des couches et la gravité. Les flux $\psi = (\psi_1, \psi_2)$, dont les dérivés horizontales donnent les vents u_i, v_i , sont les vecteurs d'état du modèle.

QG peut produire des situations météorologiques variées et peut notamment développer des ondes de Rossby comme dans la configuration utilisée par [Emili *et al.* \(2016\)](#). Ce modèle a souvent été utilisé comme modèle jouet pour tester des méthodes d'assimilation de données destinées à la prévision numérique du temps ([Evensen \(1994\)](#); [Houtekamer et Mitchell \(1998\)](#) entre autres). Les champs de vent sur le domaine moyennés sur 24h sont montrés sur la Fig. 3.8.

3.4.2 La chimie atmosphérique avec RACM

Le schéma chimique choisi pour être couplé avec le QG est le Regional Atmospheric Chemistry Mechanism (RACM) de [Stockwell *et al.* \(1997\)](#), qui comprend 96 espèces chimiques et environ 300 réactions. Contrairement au GRS, ce schéma chimique est utilisé dans certains centres opérationnels

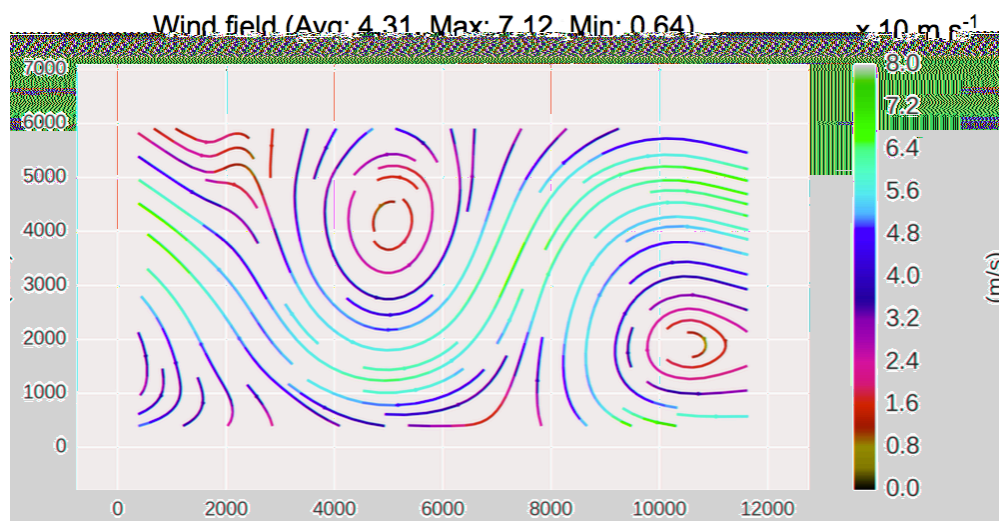


FIGURE 3.8 – Champs des vents moyennés sur 24h produits par le modèle QG. Figure tirée de Emili *et al.* (2016)

de modélisation de la qualité de l'air (Marécal *et al.*, 2015). Aussi, le choix de cette chimie rend le modèle plus détaillé, mais aussi plus complexe que le L95-GRS.

La définition des flux d'émission ou de dépôt est également plus détaillée que celle du L95-GRS, avec des émissions plus variées et des valeurs de dépôt différentes pour chaque espèce. Enfin, des conditions aux limites ont également été implémentées, ce qui rend le modèle plus réaliste mais de nouveau plus complexe. Cela permet dans le contexte de l'assimilation de données, de vérifier l'influence des conditions aux limites sur les résultats d'estimation.

3.4.3 Comparaison avec le L95-GRS

À la vue de cette description, il apparaît que les modèles de chimie atmosphérique et de météorologie utilisés dans le QG-Chem sont plus réalistes que ceux du L95-GRS. Le domaine est tridimensionnel contrairement au L95-GRS qui n'est qu'unidimensionnel, et des conditions aux limites sur deux des bords du domaine ont été implémentées. Les phénomènes d'émission et de dépôt sont définis avec plus de précision et un degré de spécification plus élevé. Il s'agit donc d'un modèle presque intermédiaire entre un modèle jouet comme le L95-GRS et un CTM complet. Un des avantages du QG-Chem est qu'il a été développé avec l'Object Oriented Prediction System (OOPS), ce qui lui permet d'être facilement pris en main par d'autres utilisateurs voulant tester de nouvelles méthodes d'assimilation de données.

Cependant, on peut se demander dans quelle mesure ce modèle n'est pas déjà trop complexe pour répondre à l'objectif recherché d'avoir un modèle d'ordre réduit simple pour tester des méthodes d'assimilation de données. En effet, le passage d'un modèle avec 240 variables à un modèle avec 12416 variables (16×8 points de grille et 97 espèces chimiques) pour le QG-Chem amplifie grandement le temps de calcul numérique du modèle. Ceci est d'autant plus vrai que le pas de temps d'intégration du L95-GRS a été fixé à 1 h, avec la possibilité d'adapter ce dernier pour le réduire jusqu'à 10 min, alors que le pas de temps d'intégration du QG-Chem est de 10 min avec la possibilité de le réduire jusqu'à 7 s. Ces différences expliquent notamment pourquoi les expériences d'assimilation de données effectuées avec le QG-Chem se réduisent à quelques cycles de 24 h, tandis qu'il est possible avec le L95-GRS d'effectuer 100 000 cycles d'assimilation (voir 4.4), permettant

d'étudier la stabilité et robustesse du système.

Même si la complexité du QG-Chem reste inférieure à celle d'un système opérationnel, elle est suffisante pour que les auteurs se soient placés dans des contextes simplifiés dans le cadre de leur étude. Notamment, le modèle n'est utilisé qu'en tant que CTM, avec une météorologie non perturbée, ce qui implique que seule de la chimie stable est modélisée. Cette stabilité impacte particulièrement le NO_2 dans les expériences sans erreur modèle. De plus, certaines expériences d'assimilation de données ont été effectuées sans avoir de transport par le vent, en se concentrant uniquement sur la chimie. Le modèle, initialement tridimensionnel, devient alors une collection de modèles 0D.

Le choix par les auteurs de ces simplifications montre la difficulté de créer des modèles jouets adaptés à l'évaluation de méthodes d'assimilation de données. Un équilibre entre réalisme, complexité numérique et complexité des phénomènes physiques intervenant dans le modèle doit être trouvé. Le choix assumé du QG-Chem s'est alors porté sur le réalisme du modèle de chimie atmosphérique choisi, quitte à réduire la complexité du couplage avec la météorologie dans leur étude. Cela démontre la modularité du QG-Chem qui offre la possibilité de se focaliser sur un phénomène physique, choisi comme centre d'intérêt d'une étude, en simplifiant certaines des autres caractéristiques du modèle.

Évaluation des méthodes EnVar 4D itératives sur des modèles jouets non-linéaires

Les modèles décrits dans le chapitre précédent ont été développés afin d'avoir des qualités physiques, telles que la chaotité ou la non-linéarité, faisant d'eux de parfaits bancs d'essai pour l'application de méthodes d'assimilation de données. Ainsi, nous présenterons dans ce chapitre l'évaluation de méthodes EnVar 4D, notamment l'itérative ensemble Kalman smoother (IEnKS) et l'IEnKF-Q, sur certains des précédents modèles.

La dynamique chaotique du Lorenz-95 (L95), d'abord sans puis avec erreur modèle, le couplage proposé par le L95-T et la chimie non-linéaire du L95-GRS présentent des difficultés variées qui seront abordées dans cet ordre. Au cours de ce chapitre, nous validerons ainsi numériquement l'algorithme IEnKF-Q. De plus, différentes stratégies d'application de l'IEnKS sur un modèle faiblement couplé de chimie-météorologie seront mises en avant et testées. Enfin, les qualités du L95-GRS, en tant que modèle jouet pour tester les méthodes d'assimilation de données dans un contexte avec chimie non-linéaire, seront vérifiées.

L'application de l'IEnKF-Q sur le L95, présentée dans la section 4.2, est tirée de [Sakov *et al.* \(2017, soumis\)](#) que nous avons soumis au cours de la thèse, et celle de l'IEnKS sur le L95-GRS, présentée dans la section 4.4, est quant à elle tirée de [Haussaire et Bocquet \(2016\)](#), article qui a été publié durant la thèse.

Sommaire

4.1	Expériences préliminaires sur le L95	77
4.1.1	Inter-comparaison 4D-Var, EnKS et IEnKS	77
4.1.2	Expérience avec estimation de paramètres	78
4.1.3	Expérience avec localisation	79
4.2	Expériences avec erreur modèle sur le L95	81
4.2.1	Base de comparaison	81
4.2.2	Impact de la non-linéarité	82
4.2.3	Impact de l'amplitude de l'erreur modèle	84
4.2.4	Impact de la taille de l'ensemble	85
4.2.5	Impact de la localisation	85
4.2.6	Conclusion des expériences avec erreur modèle	87
4.3	Résultats sur le L95-T	88
4.3.1	Estimation de paramètres	88
4.3.2	Régimes d'émission et de dépôt	89
4.3.3	Stratégies pour des modèles faiblement couplés	90
4.4	Assimilation de données sur le L95-GRS	93
4.4.1	Premières performances	94

4.4.2	Estimation de paramètres	96
4.4.3	Localisation	96
4.5	Conclusions	99

Nous nous appuierons pour la rédaction de ce chapitre sur

- M. ASCH, M. BOCQUET et M. NODÉ : Chapter 7: Ensemble variational methods. *In Data assimilation: methods, algorithms, and applications* Asch *et al.* (2016c), chap. 7, p. 195–216
- M. BOCQUET et P. SAKOV : Joint state and parameter estimation with an iterative ensemble Kalman smoother. *Nonlinear Process. Geophys.*, 20:803–818, 2013
- M. BOCQUET : Localization and the iterative ensemble Kalman smoother. *Q. J. R. Meteorol. Soc.*, 142:1075–1089, 2016
- J.-M. HAUSSAIRE et M. BOCQUET : A low-order coupled chemistry meteorology model for testing online and offline data assimilation schemes: L95-GRS (v1. 0). *Geosci. Model Dev.*, 9:393–412, 2016
- P. SAKOV, J.-M. HAUSSAIRE et M. BOCQUET : An iterative ensemble Kalman filter in presence of additive model error. *Q. J. R. Meteorol. Soc.*, 2017, soumis

4.1 Expériences préliminaires sur le L95

Nous présenterons dans cette section des résultats provenant de [Bocquet et Sakov \(2013\)](#) et [Bocquet \(2016\)](#), qui comparent tour à tour les performances de l'ensemble Kalman filter (EnKF) (dans son implémentation ensemble transform Kalman filter (ETKF)), de l'ensemble Kalman smoother (EnKS), du 4D-Var et de l'IEEnKS dans des contextes plus ou moins linéaires, en filtrage ou en lissage, avec ou sans localisation ainsi qu'avec ou sans estimation de paramètres. Ceci représente donc un large éventail de cas de figure qui sont tous permis par le L95, ce qui prouve encore une fois sa qualité en tant que modèle jouet. Chacune de ces expériences seront de type *observing system simulation experiment* (OSSE), autrement appelées expériences jumelles. On engendre une trajectoire avec un modèle, que l'on considère comme la vérité. On crée alors des observations vérifiant les statistiques de $\mathbf{R}$ en perturbant cette trajectoire. Enfin, on perturbe la condition initiale, en utilisant potentiellement également un modèle erroné, et on essaie d'obtenir une analyse la plus proche de la vérité possible en assimilant les observations précédemment engendrées. Pour estimer la qualité de l'analyse, on la compare à la vérité en calculant une racine de l'erreur quadratique moyenne (RMSE), de filtrage ou de lissage, selon les formules (2.162) et (2.163) respectivement. Dans les expériences que nous montrerons dans tout ce chapitre, 10^5 cycles d'assimilation de données ont été calculés, de telle sorte que les statistiques obtenues aient bien convergé.

4.1.1 Inter-comparaison 4D-Var, EnKS et IEEnKS

Dans cette expérience, le modèle est observé tous les $\Delta t = 0.05$ et toutes les $n = 40$ variables sont observées, avec une variance de 1, de telle sorte que la matrice des erreurs d'observation soit $\mathbf{R} = \mathbf{I}_{40}$. Ainsi, après assimilation de données, l'analyse doit fournir une RMSE inférieure à environ 1, sinon les observations sont considérées comme une meilleure estimation que l'analyse (du moins là où le modèle est observé). Par ailleurs, nous rappelons que le modèle est chaotique, avec 13 exposants de Lyapunov positifs et 1 nul, si bien que si l'on part d'une condition initiale erronée, on diverge rapidement de la vérité (voir Fig. 3.2). Les méthodes ensemblistes, EnKS ou IEEnKS, utilisent un ensemble de taille $m = 20$, de telle sorte qu'il n'est pas nécessaire de recourir à de la localisation mais de l'inflation est implicitement utilisée via le *finite-size* (voir 2.4.5.2).

Les RMSEs de l'expérience d'assimilation de données avec l'EnKS, le 4D-Var, l'IEEnKS *single data assimilation* (SDA) $S = 1$, l'IEEnKS SDA $S = L$ et l'IEEnKS *multiple data assimilation* (MDA) $S = 1$, en filtrage et en lissage, en fonction de la longueur de la fenêtre d'assimilation de données L (en unité de Δt), sont montrées sur la Fig. 4.1 ([Bocquet et Sakov, 2013](#)). On rappelle qu'en filtrage, l'EnKS donne un résultat strictement équivalent à l'EnKF (ici sous sa forme ETKF) qui ne dépend donc pas de la longueur de la fenêtre d'assimilation. La seule évolution que montre la courbe correspond à un bruit statistique dont l'amplitude n'est pas représentative d'une quelconque évolution. Par ailleurs, l'inflation des méthodes ensemblistes est optimisée pour fournir le meilleur résultat possible. La matrice des erreurs d'ébauche du 4D-Var est prise diagonale et la valeur des coefficients est optimisée pour chaque longueur de fenêtre afin d'obtenir le meilleur résultat. Pour le 4D-Var en lissage, la RMSE optimale au sein de la fenêtre est montrée. Elle n'est donc pas nécessairement calculée à l'extrémité de la fenêtre.

On voit que pour de courtes fenêtres d'assimilation ($L \leq 20$), l'EnKS bat le 4D-Var aussi bien en filtrage qu'en lissage. Ceci s'explique par le fait que l'EnKS propage les erreurs entre chaque analyse tandis que le 4D-Var se contente de propager un unique état optimal. En revanche, la tendance s'inverse pour de plus longues fenêtres ($L > 20$), grâce à l'analyse variationnelle du 4D-Var au sein de la fenêtre, où les non-linéarités s'amplifient avec la longueur.

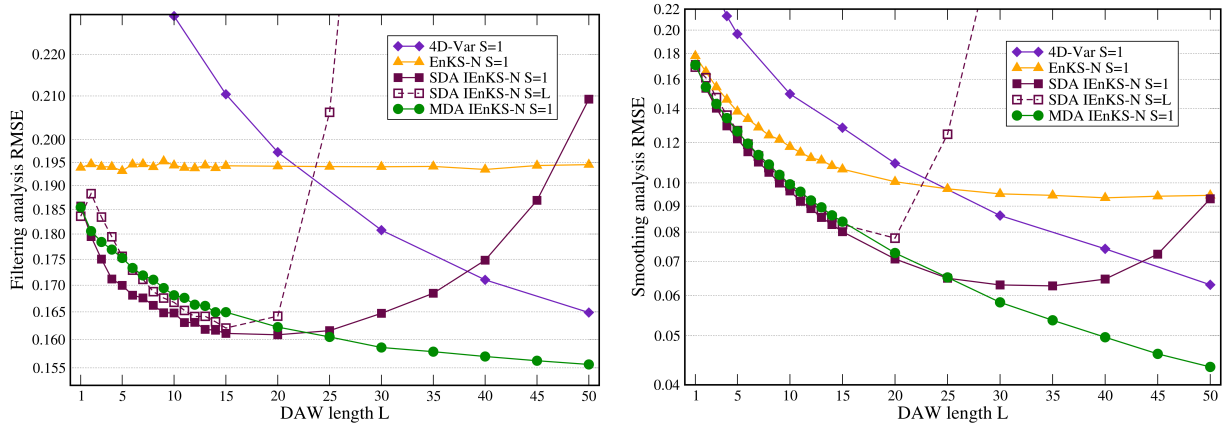


FIGURE 4.1 – RMSEs en fonction de la longueur de la fenêtre d’assimilation de données (en unité de $\Delta t = 0.05$), pour différentes méthodes d’assimilation. Le pas de temps entre deux observations est $\Delta t = 0.05$. À gauche, RMSEs de filtrage, à droite de lissage. Figure issue de [Bocquet et Sakov \(2013\)](#).

On remarque que l’IEnKS, aussi bien en SDA qu’en MDA, donne des résultats nettement meilleurs que l’EnKS ou le 4D-Var. Pour de courtes fenêtres, le SDA est meilleur que le MDA. Cependant, ce dernier devient nettement supérieur pour de longues fenêtres. Enfin, pour de courtes fenêtres, l’option SDA $S = L$ est sensiblement similaire au SDA $S = 1$, malgré un coup de calcul bien moindre. En effet, à partir d’une fenêtre de taille $15\Delta t = 0.75$, soit près de deux fois le temps de doublement, l’ébauche utilisée est issue d’une propagation sur une telle longueur qu’elle n’est pas de bonne qualité. Il est important de préciser que l’IEnKS n’a besoin en moyenne que de deux propagations de l’ensemble dans la fenêtre lors de la minimisation pour atteindre l’optimum. Ainsi, le coût de calcul qu’il représente n’est pas prohibitif dans ce cas là.

Lorsque l’on observe tous les $\Delta t = 0.05$, le modèle apparaît comme faiblement non-linéaire. Il est donc possible de jouer sur ce degré de non-linéarité en diminuant la fréquence des observations, c’est-à-dire en augmentant le temps entre deux jeux d’observations. Ainsi, dans l’expérience suivante, $\Delta t = 0.2$ a été choisi et les autres paramètres sont par ailleurs restés fixés. Les résultats sont montrés sur la Fig. 4.2. La supériorité du 4D-Var sur l’EnKS dans les cas non-linéaires est encore une fois mise en évidence et amplifiée ici. Les conclusions précédentes quant aux performances de l’IEnKS dans ses diverses configurations, relativement aux deux autres méthodes, tiennent toujours.

4.1.2 Expérience avec estimation de paramètres

Un des objectifs de l’assimilation de données, comme exprimé dans le chapitre 1, est d’estimer des paramètres méconnus des modèles. Un des paramètres importants du L95 est le forçage F , dont la valeur impacte fortement la dynamique du modèle. Aussi est-il intéressant d’essayer de l’estimer en même temps que les variables du modèle.

Une méthode simple pour estimer les paramètres d’un modèle avec des méthodes d’assimilation de données est d’utiliser un formalisme d’état augmenté. L’intérêt de cette méthode est synthétisé dans ([Ruiz et al., 2013](#)). Si on note $\theta \in \mathbb{R}^q$ le vecteur des paramètres à estimer conjointement au vecteur d’état $\mathbf{x} \in \mathbb{R}^n$, alors on augmente l’état pour former

$$\mathbf{z} = \begin{pmatrix} \mathbf{x} \\ \theta \end{pmatrix} \in \mathbb{R}^{n+q}, \quad (4.1)$$

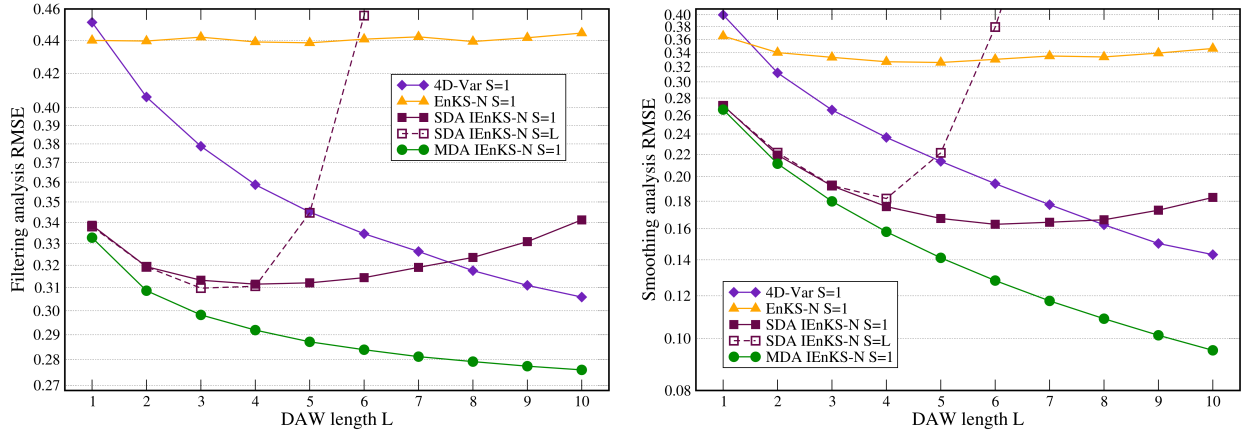


FIGURE 4.2 – RMSEs en fonction de la longueur de la fenêtre d’assimilation de données (en unité de $\Delta t = 0.2$), pour différentes méthodes d’assimilation. Le pas de temps entre deux observations est $\Delta t = 0.2$. À gauche, RMSEs de filtrage, à droite de lissage. Figure issue de [Bocquet et Sakov \(2013\)](#).

dans l’espace joint des états et des paramètres. Il est également nécessaire de définir un modèle de propagation pour les paramètres. Un premier choix simple peut alors être de prendre le modèle de persistance, c’est-à-dire $\theta_{k+1} = \theta_k$.

À nouveau, [Bocquet et Sakov \(2013\)](#) ont fait une OSSE dans ce contexte, où la vérité est engendrée avec une forçage $F = 8$ et où les membres de l’ensemble sont initialisés avec une valeur $7 + \varepsilon$ où $\varepsilon \sim \mathcal{N}(0, 0.1)$. Ainsi, la valeur moyenne de l’ensemble est biaisée et la méthode d’assimilation de données doit permettre de diagnostiquer la bonne valeur du paramètre, sachant que celui-ci n’est pas observé et qu’il n’agit que lors de la propagation du modèle. Notons que la distinction entre RMSE de filtrage et de lissage pour le paramètre ne s’applique pas ici, étant donné que l’on utilise le modèle de persistance et que le vrai paramètre est fixe. L’intégralité des autres paramètres de l’expérience sont gardés identiques au cas précédent et le pas de temps entre les observations est $\Delta t = 0.05$.

Les courbes des valeurs analysées du paramètre pour 5000 pas de temps (après spin-up), ainsi que les performances des différents algorithmes, sont tracées sur la Fig. 4.3. On y voit que l’IEnKS est toujours meilleur que l’EnKS et le 4D-Var dans ce cas. En ce qui concerne les performances des différents algorithmes sur l’état, elles ne sont que très peu impactées par l’estimation conjointe de la valeur du paramètre de forçage et ressemblent donc à la Fig. 4.1. Différentes configurations, où le paramètre évolue en fonction du temps, ont également été testées dans [Bocquet et Sakov \(2013\)](#) et ont mené aux mêmes conclusions, bien que le modèle de propagation pour les paramètres soit dans ces cas là erroné.

On a donc illustré avec cette expérience, d’une part la capacité du L95 à servir de modèle jouet pour l’estimation de paramètres, d’autre part l’utilité des méthodes d’assimilation de données pour l’estimation des paramètres et enfin la supériorité de l’IEnKS face au 4D-Var et à l’EnKS dans ce contexte également.

4.1.3 Expérience avec localisation

Si on diminue le nombre de membres de l’ensemble sous la taille de l’espace instable du L95, l’ensemble est trop petit pour empêcher la divergence du filtre. Il faut alors recourir à de la localisation (voir, par exemple, [Bocquet et Carrassi, 2017](#)).

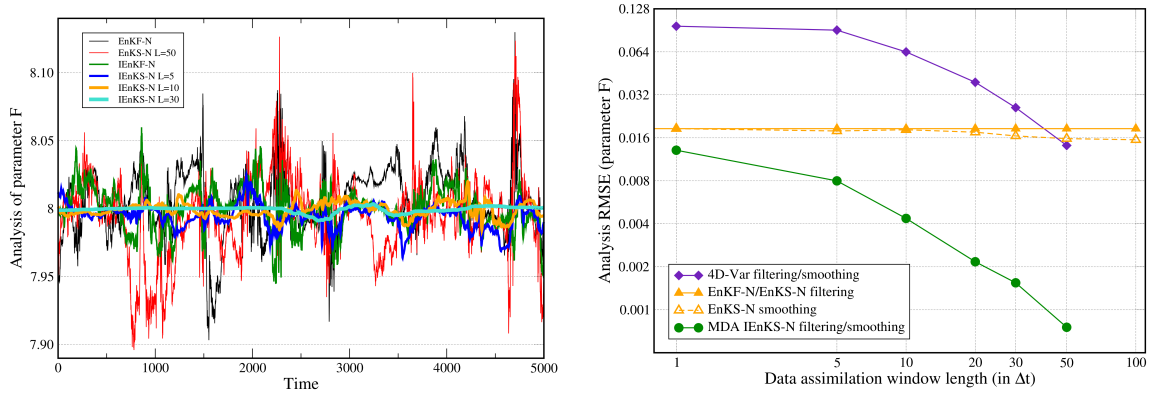


FIGURE 4.3 – À gauche : valeurs du paramètre en fonction du temps (en unité de Δt) pour différents algorithmes d’assimilation de données. À droite : RMSEs de filtrage et de lissage en fonction de la longueur de la fenêtre d’assimilation de données (en unité de $\Delta t = 0.05$), pour différentes méthodes d’assimilation. Le pas de temps entre deux observations est $\Delta t = 0.05$. Figure tirée de [Bocquet et Sakov \(2013\)](#)

Comme on l’a vu dans la section 2.5.2, l’application de la localisation sur les schémas de type EnVar 4D n’est pas évidente. Une première façon de procéder est de propager totalement les matrices de localisation, ce qui n’est pas du tout applicable en réalité car beaucoup trop coûteux et contre-productif. Cependant, l’intérêt d’utiliser un modèle d’ordre réduit est précisément de pouvoir tester ces méthodes qui fournissent alors une référence optimale vers laquelle tendre. Nous avons également à notre disposition, dans le cas d’une localisation par domaines, les algorithmes de localisation par domaines statiques et covariants présentés dans la section 2.5.3.3.

Dans cette expérience, tirée de [Bocquet \(2016\)](#), les méthodes avec localisation utilisent un ensemble de taille $m = 10$, tandis que 20 membres sont gardés quand aucune localisation n’est appliquée. Les résultats sont présentés sur la Fig. 4.4. Les notations utilisées pour les différentes courbes sont :

- R_1 : EnKS global ;
- R_2 : IEnKS global ;
- R_3 : IEnKS local avec propagation complète des matrices de localisation. Cas de référence non applicable sur des gros modèles ;
- DL_1 : IEnKS avec localisation par domaines statiques ;
- DL_2 : IEnKS avec localisation par domaines covariants ;
- DL_3 : EnKS avec localisation par domaines statiques (standard pour l’EnKS) ;
- CL_1 : IEnKS avec localisation des covariances en utilisant la propagation de 15 modes à travers la fenêtre. Le choix de 15 est tel que ce soit légèrement supérieur à la taille de l’espace instable.

On voit que la localisation permet de réduire la taille de l’ensemble tout en empêchant la divergence du filtre et en obtenant des scores similaires au cas global, comme la comparaison des courbes R_1 et DL_3 le montre. On voit également que l’approche naïve de localisation par domaines statiques DL_1 marche pour des courtes fenêtres, mais échoue pour de plus grandes, comme on pouvait le prévoir. Il en va de même de la localisation covariante avec un nombre limité de modes, pour laquelle la chaoticité du modèle pose problème. Enfin, on remarque que l’advection des domaines de localisation DL_2 fournit un résultat proche de la référence optimale R_3 .

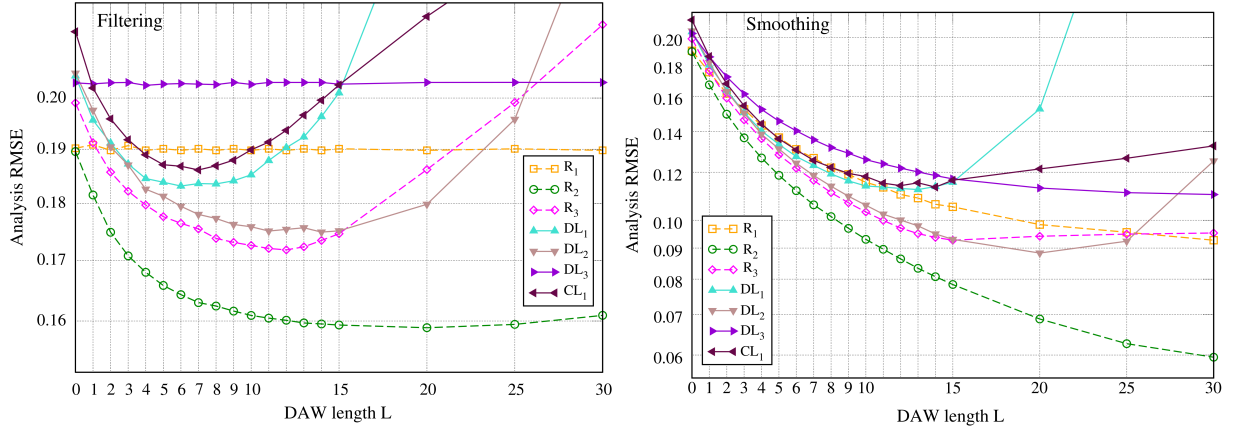


FIGURE 4.4 – RMSEs en fonction de la longueur de la fenêtre d’assimilation de données (en unité de $\Delta t = 0.05$), pour différentes méthodes d’assimilation globales et locales. Le pas de temps entre deux observations est $\Delta t = 0.05$. À gauche, RMSEs de filtrage, à droite de lissage. Figure tirée de Bocquet (2016).

4.2 Expériences avec erreur modèle sur le L95

À la suite du développement de l’IEnKF-Q, présenté dans la section 2.5.4, des expériences ont été menées dans Sakov *et al.* (2017, soumis) sur le L95 afin de valider numériquement ses performances. Nous rappelons que nous nous plaçons ici dans le contexte de l’iterative ensemble Kalman filter (IEnKF), à savoir que $L = 1$ dans cette section.

L’erreur modèle est incorporée dans la vérité par ajout d’un bruit stochastique de variance $\mathbf{Q}$ après la propagation sur l’intégralité de la fenêtre d’assimilation de données. De cette manière, l’erreur modèle utilisée dans les algorithmes va pouvoir être en accord avec celle effectivement appliquée sur la vérité.

Dans toutes les expériences de cette section, toutes les variables du L95 sont observées, avec une variance 1, de telle sorte que $\mathbf{R} = \mathbf{I}_{n=40}$. On moyenne à nouveau les RMSEs sur 10^5 cycles d’assimilation de données et on diagnostique, pour chaque expérience, l’inflation optimale parmi l’ensemble $\{1, 1.02, 1.05, 1.1, 1.15, 1.2, 1.25, 1.3, 1.4, 1.5, 1.75, 2, 2.5, 3, 4\}$. On n’utilise pas le *finite-size* dans ces expériences. On fixe le paramètre $m_q = 41$ pour l’IEnKF-Q, pour que $\mathbf{X}^q$ soit une racine exacte de $\mathbf{Q}$, éventuellement de rang plein, et afin de mettre en avant tout le potentiel de l’IEnKF-Q.

4.2.1 Base de comparaison

Avant de montrer les résultats de l’IEnKF-Q, nous présentons d’abord les méthodes auxquelles on va le comparer. Dans la section 2.4.5.3, nous avons présenté comment incorporer dans l’ébauche de l’inflation additive, de manière stochastique selon la formule (2.95) ou de manière déterministe selon la formule (2.102). Nous appliquerons ainsi ces deux méthodes sur un ETKF afin de définir respectivement l’EnKF-Rand et l’EnKF-Det. Ces mêmes incorporations de l’erreur modèle peuvent être appliquées avec l’IEnKF pour former les méthodes heuristiques IEnKF-Rand et IEnKF-Det. Nous précisons que dans ce cas, l’état $\mathbf{x}_f$ et les anomalies $\mathbf{X}_f$ à un instant t , qui sont utilisés dans ces formules, sont issus de la propagation de l’ensemble lissé à $t - 1$, obtenu par assimilation de l’observation au temps t .

Le découplage de l’analyse de l’IEnKF-Q suggère que l’IEnKF devrait utiliser la matrice $\mathbf{R} +$

$\mathbf{H}\mathbf{Q}\mathbf{H}^T$ à la place de $\mathbf{R}$ dans son analyse. Cependant, aucune différence de performance n'a été observée en appliquant cette modification.

Par ailleurs, l'IEnKF-Rand conduit à une matrice de covariance des erreurs d'analyse suivant

$$\mathbf{P}^a \simeq \mathbb{E}[\mathbf{M}\mathbf{P}^s\mathbf{M}^T + \mathbf{Q}] \quad (4.2)$$

$$\begin{aligned} &\simeq \mathbf{Q} + \mathbf{M}\mathbf{P}^f\mathbf{M}^T \\ &\quad - \mathbf{M}\mathbf{P}^f\mathbf{M}^T\mathbf{H}^T \left[\mathbf{R} + \mathbf{H}(\mathbf{Q} + \mathbf{M}\mathbf{P}^f\mathbf{M}^T)\mathbf{H}^T \right]^{-1} \mathbf{H}\mathbf{M}\mathbf{P}^f\mathbf{M}^T, \end{aligned} \quad (4.3)$$

tandis que l'IEnKF-Q conduit aux erreurs

$$\begin{aligned} \mathbf{P}^a &= \mathbf{Q} + \mathbf{M}\mathbf{P}^f\mathbf{M}^T \\ &\quad - (\mathbf{Q} + \mathbf{M}\mathbf{P}^f\mathbf{M}^T)\mathbf{H}^T \left[\mathbf{R} + \mathbf{H}(\mathbf{Q} + \mathbf{M}\mathbf{P}^f\mathbf{M}^T)\mathbf{H}^T \right]^{-1} \mathbf{H}(\mathbf{Q} + \mathbf{M}\mathbf{P}^f\mathbf{M}^T). \end{aligned} \quad (4.4)$$

Ces deux expressions, bien que différentes, sont très similaires. Ainsi, l'IEnKF-Rand peut être considéré comme une bonne approximation de l'IEnKF-Q. Par ailleurs, l'équation (4.2) fournit l'expression exacte de $\mathbf{P}^a$ pour l'IEnKF-Det si la taille de l'ensemble m est strictement plus grande que la taille de l'espace des états n . Cela explique que, sous cette hypothèse, les résultats de l'IEnKF-Det et de l'IEnKF-Q puissent être similaires.

4.2.2 Impact de la non-linéarité

Dans cette première expérience, la taille de la fenêtre d'assimilation de données est variée, pour vérifier l'impact de la non-linéarité sur les performances des différentes méthodes. On fixe la taille de l'ensemble à $m = 20$. La matrice de covariance de l'erreur modèle vaut $\mathbf{Q} = 0.01T\mathbf{I}_n$, où T est l'intervalle entre deux jeux d'observations. Elle est donc proportionnelle à la longueur de la fenêtre d'assimilation de données. Étant donné que l'erreur modèle est incorporée à l'issue de l'intégration sur l'ensemble de la longueur de la fenêtre, cela signifie que l'erreur modèle par unité de temps reste constante. Même si cette valeur de $\mathbf{Q}$ est deux ordres de grandeur plus faible que celle de $\mathbf{R}$ quand $T = 1$, elle est raisonnable. En effet, l'écart type du bruit ajouté à chaque pas de temps est de 0.1, à mettre en regard d'une RMSE d'environ 0.2 pour les résultats présentés pour l'ETKF à la Fig. 4.1.

La Fig. 4.5 compare les performances des EnKFs, des IEnKFs et de l'IEnKF-Q en fonction de T , le temps entre deux jeux d'observations. À partir de $T = 3$, les méthodes itératives battent nettement les EnKFs, du fait de la non-linéarité croissante. Pour toutes les valeurs de T , l'IEnKF-Q fait mieux que les IEnKF-Rand et IEnKF-Det, y compris pour $T = 1$. La raison est que l'IEnKF-Q utilise un ensemble augmenté de taille $m + m_q$ durant la minimisation, alors que les autres méthodes n'ont accès qu'à un ensemble de taille m . Cet avantage décroît logiquement quand la taille de l'ensemble m est augmentée (voir 4.2.4).

La Fig. 4.6 reproduit avec les mêmes jeux de paramètres l'expérience de la Fig. 4 de Raanes *et al.* (2015). Notre EnKF-Rand et EnKF-Det correspondent respectivement à leurs Add-Q et SQRT-CORE. La seule différence est que Raanes *et al.* (2015) ajoutent de l'erreur modèle de variance $[\mathbf{Q}]_{i,j} = 0.05 \left(\exp[-d(i-j)^2/30] + 0.1\delta_{i,j} \right)$ après chaque intégration du modèle, plutôt que T fois cette valeur ajoutée une unique fois à la fin de l'intégration sur la totalité du cycle. On a noté ici $d(i,j) = \min(|i-j|, 40 - |i-j|)$. Comparée à la Fig. 4.5, la taille de l'ensemble a été augmentée à 30, et l'erreur modèle est plus importante et spatialement corrélée. Cette hausse de l'erreur modèle accentue l'avantage de l'IEnKF-Q sur les autres méthodes. Les filtres non-itératifs présentent des résultats qui paraissent relativement meilleurs grâce à l'augmentation de la taille de l'ensemble.

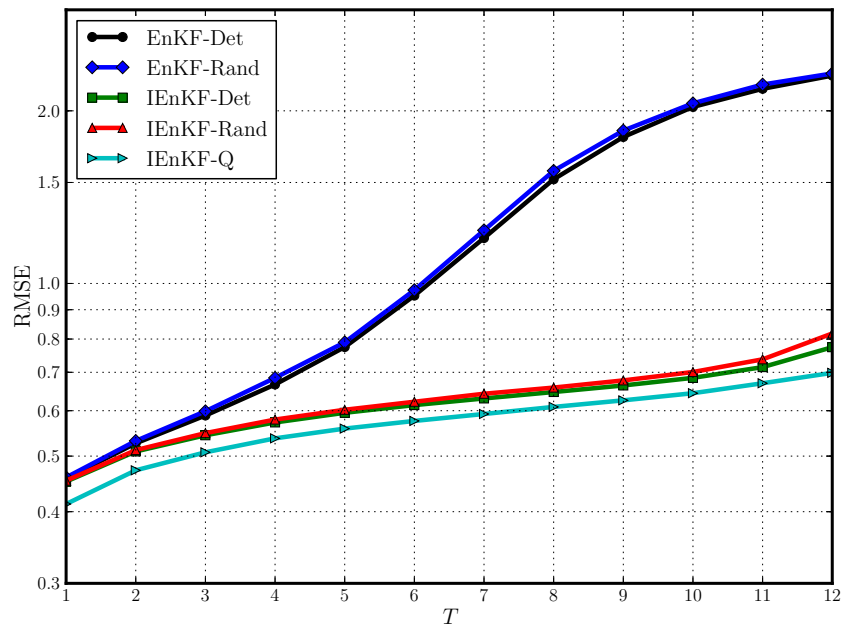


FIGURE 4.5 – RMSEs de filtrage en fonction de T , l'intervalle de temps entre les observations. $m = 20$, $\mathbf{Q} = 0.01T\mathbf{I}_n$. Figure tirée de Sakov *et al.* (2017, soumis).

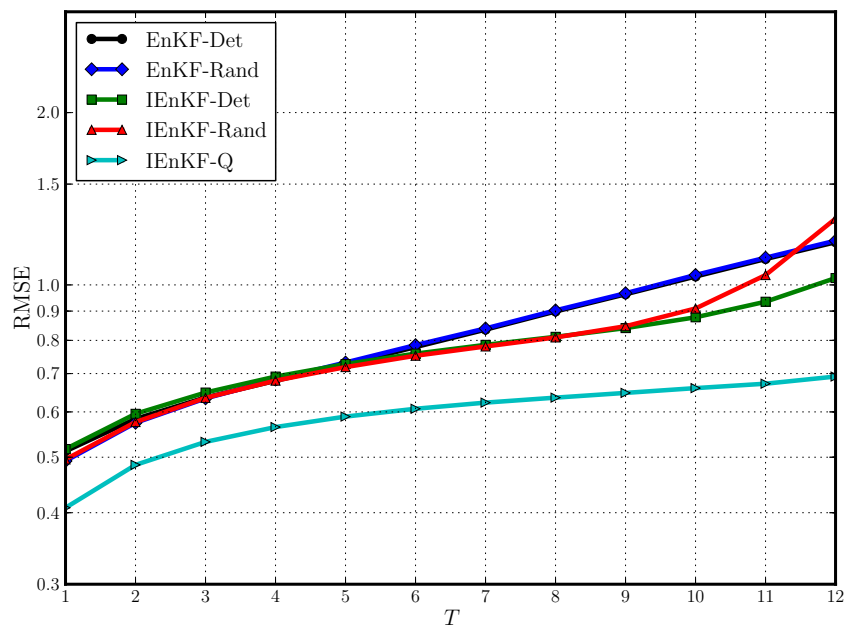


FIGURE 4.6 – RMSEs de filtrage en fonction de T , l'intervalle de temps entre les observations, avec des paramètres similaires à la Fig. 4 de Raanes *et al.* (2015). $m = 30$, $[\mathbf{Q}]_{i,j} = 0.05T (\exp[-d(i-j)^2/30] + 0.1\delta_{i,j})$, où $\delta_{i,j}$ est le symbole de Kronecker. Figure tirée de Sakov *et al.* (2017, soumis).

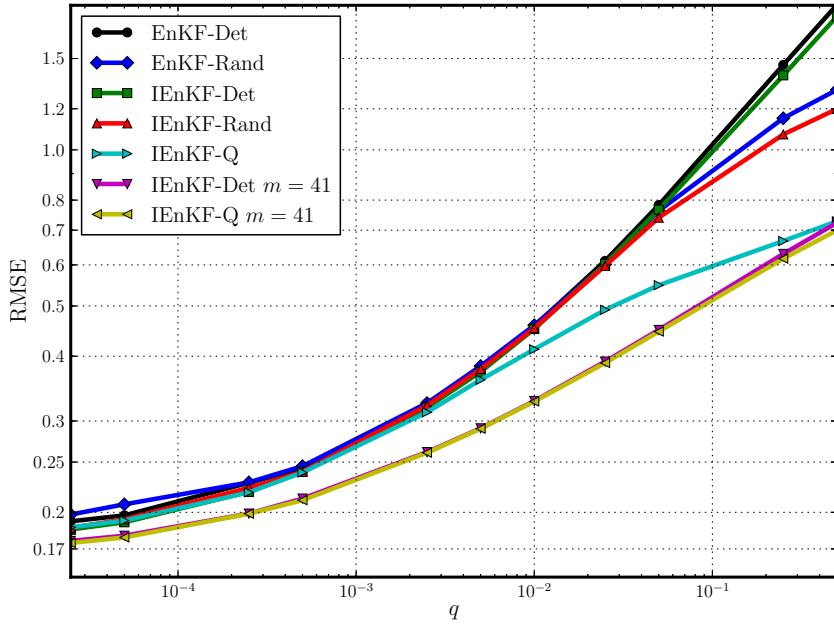


FIGURE 4.7 – RMSEs de filtrage en fonction de q , l'amplitude de l'erreur modèle dans un cas faiblement non-linéaire. $m = 20$, $T = 1$, $\mathbf{Q} = qT\mathbf{I}_n$. Figure tirée de Sakov *et al.* (2017, soumis).

4.2.3 Impact de l'amplitude de l'erreur modèle

Dans cette expérience, l'amplitude de l'erreur modèle q , telle que $\mathbf{Q} = qT\mathbf{I}_n$, est variée sur une large plage de valeurs dans un cas faiblement non-linéaire et fortement non-linéaire. Tous les tests utilisent un ensemble de taille $m = 20$ et deux tests complémentaires avec $m = 41$, afin que les matrices soient de rang plein, sont montrés pour l'IEnKF-Det et l'IEnKF-Q.

La Fig. 4.7 montre les performances des différents algorithmes dans le cas faiblement non-linéaire, c'est-à-dire avec $T = 1$. Pour de petites erreurs modèle, $q \lesssim 3 \cdot 10^{-3}$, tous les schémas sont presque indiscernables. Pour de plus grandes valeurs de q , l'IEnKF-Q commence à dépasser les autres méthodes et pour $q \gtrsim 5 \cdot 10^{-2}$, les méthodes itératives IEnKF-Rand et IEnKF-Det dépassent leurs équivalents non-itératifs. Étonnamment, l'IEnKF-Det et l'IEnKF-Q avec $m = 41$ donnent des résultats très similaires, avec un très léger avantage de l'IEnKF-Q pour les très grandes valeurs de q .

Bien qu'étonnant, ce résultat est conforme à ce que prédit la théorie. En effet, le découplage montre que l'analyse en lissage des deux méthodes est strictement équivalente dans ce cas. Le filtrage quant à lui est différent, même si très proche. Celui de l'IEnKF-Q est fourni par l'équation (2.158) et celui de l'IEnKF-Det correspond simplement au premier terme de cette expression. En ce qui concerne les anomalies analysées, on a vu avec les équations (4.2) et (4.4) qu'elles étaient très proches, sauf peut-être précisément quand $\mathbf{Q}$ est grand.

Dans le cas fortement non-linéaire de la Fig. 4.8, les algorithmes non-itératifs ne sont plus capables de contraindre le modèle. Les algorithmes itératifs quant à eux donnent des résultats similaires à l'IEnKF-Q jusque des valeurs $q \lesssim 2 \cdot 10^{-3}$ avec un ensemble de taille $m = 20$ et $q \lesssim 2 \cdot 10^{-2}$ avec $m = 41$. Cependant, pour de fortes valeurs de q , ils perdent en stabilité et ne réussissent pas à achever les 10^5 cycles d'assimilation de données requis. On remarquera également que pour une très forte erreur modèle $q = 0.5$ ($\mathbf{Q} = 5\mathbf{I}_n$), l'IEnKF-Q produit des résultats similaires que l'on utilise $m = 20$ ou $m = 41$, avec une RMSE d'environ 0.94, nettement inférieure à l'amplitude de

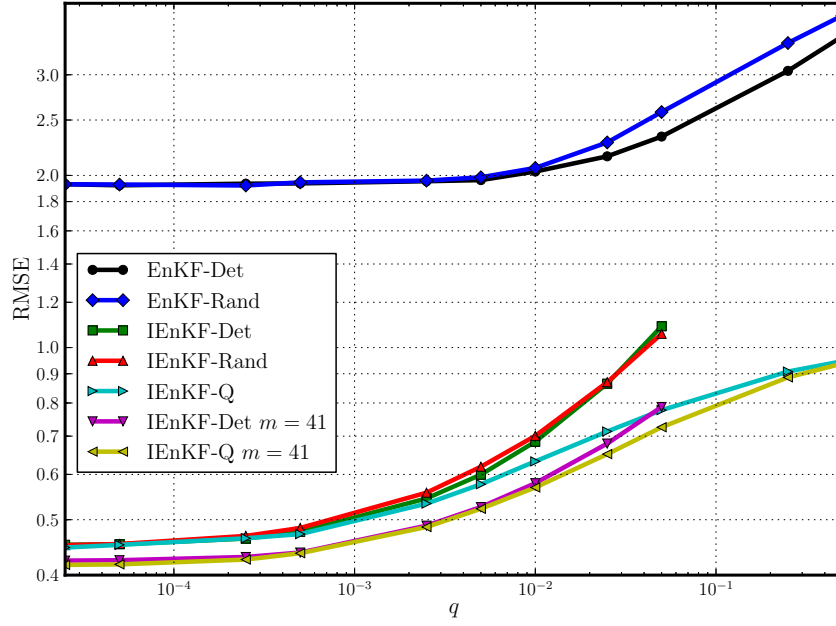


FIGURE 4.8 – RMSEs de filtrage en fonction de q , l’amplitude de l’erreur modèle dans un cas fortement non-linéaire. $m = 20$, $T = 10$, $\mathbf{Q} = q\mathbf{I}_n$. Figure tirée de Sakov *et al.* (2017, soumis).

l’erreur modèle d’environ 2.2, et même inférieure à l’erreur d’observation. Ceci est probablement dû à l’analyse variationnelle qui explore l’ensemble de l’espace des états étant donné que $m_q = 41$ et que, dans ce régime, il y a assez peu de dépendance en $\mathbf{X}^f$.

4.2.4 Impact de la taille de l’ensemble

Dans cette expérience, la taille de l’ensemble est variée dans un cas faiblement non-linéaire ($T = 1$, Fig. 4.9) et fortement non-linéaire ($T = 10$, Fig. 4.10). L’erreur modèle est à nouveau fixée à $\mathbf{Q} = 0.01\mathbf{I}_n$. Conformément aux résultats des expériences de la section précédente, on remarque que les méthodes non-itératives ne fonctionnent pas correctement dans les cas fortement non-linéaires. L’IEnKF-Q bat les autres méthodes quand la taille de l’ensemble est faible, mais l’IEnKF-Det est à nouveau très similaire lorsque l’ensemble est grand. L’avantage de l’IEnKF-Q est probablement dû au fait qu’il recherche l’analyse optimale dans un sous-espace plus grand.

4.2.5 Impact de la localisation

Dans ces expériences, qui ne sont pas publiées dans Sakov *et al.* (2017, soumis), nous validons l’algorithme 9 en observant l’impact de la localisation quand la taille de l’ensemble est réduite. Nous utilisons ici la localisation par domaines covariants. Cette dernière ne présente d’intérêt pour ces expériences que lorsque la non-linéarité est grande.

Nous reprenons dans un premier temps la configuration utilisée pour la Fig. 4.5, dont nous ne conservons que la courbe de l’IEnKF-Q global qui utilise $m = 20$ membres, et nous ajoutons une courbe pour l’IEnKF-Q local avec $m = 10$ membres. Les résultats sont tracés sur la Fig. 4.11. Nous remarquons que dans tous les régimes de non-linéarités, l’IEnKF-Q localisé avec un plus petit ensemble produit de meilleurs résultats que son homologue global avec un plus grand ensemble.

Nous reproduisons également la configuration utilisée pour la Fig. 4.7, dont nous ne conservons

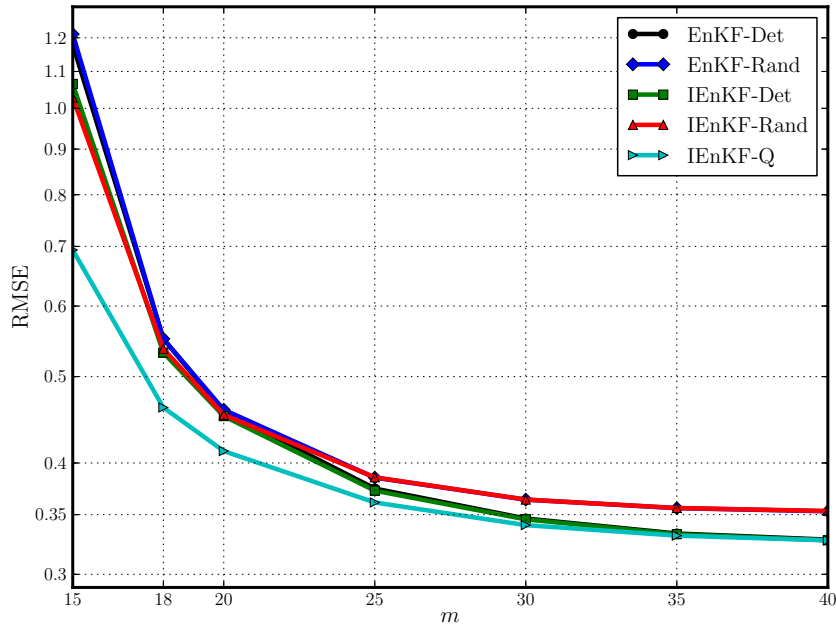


FIGURE 4.9 – RMSEs de filtrage en fonction de m , la taille de l'ensemble, dans un cas faiblement non-linéaire. $T = 1$, $\mathbf{Q} = 0.01T\mathbf{I}_n$. Figure tirée de Sakov *et al.* (2017, soumis).

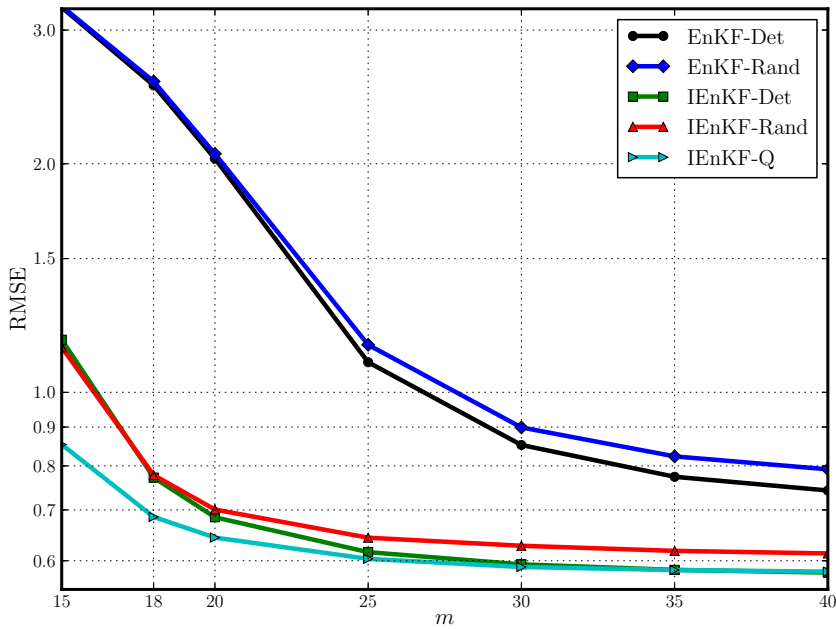


FIGURE 4.10 – RMSEs de filtrage en fonction de m , la taille de l'ensemble, dans un cas fortement non-linéaire. $T = 10$, $\mathbf{Q} = 0.01T\mathbf{I}_n$. Figure tirée de Sakov *et al.* (2017, soumis).

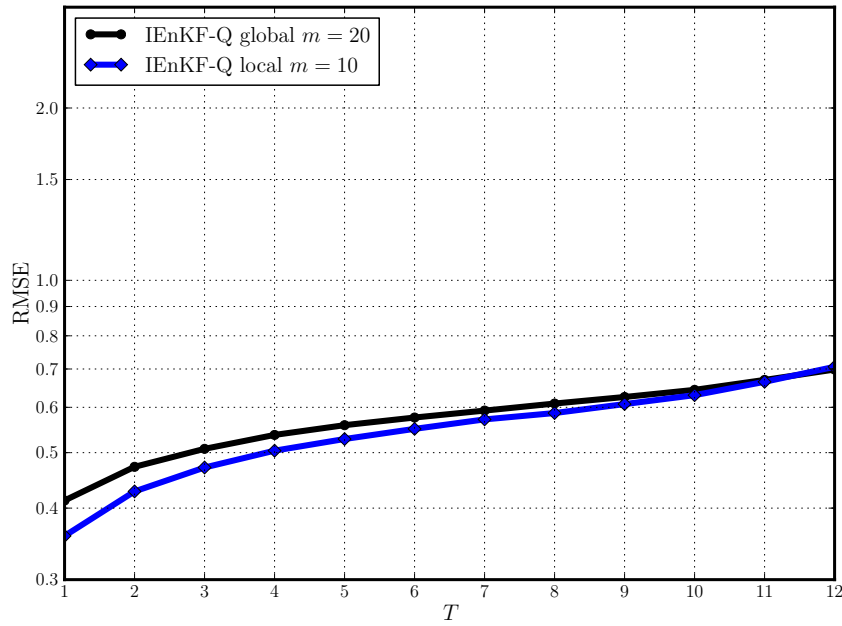


FIGURE 4.11 – RMSEs de filtrage en fonction de T , l'intervalle de temps entre les observations, pour l'IEnKF-Q global ($m = 20$) et local ($m = 10$). $\mathbf{Q} = 0.01T\mathbf{I}_n$.

que les courbes de l'IEnKF-Q global qui utilisent $m = 20$ et $m = 41$ membres, et nous ajoutons une courbe pour l'IEnKF-Q local avec $m = 10$ membres. Les résultats sont tracés sur la Fig. 4.12. Les performances de l'IEnKF-Q localisé sont comprises entre celles de l'IEnKF-Q global avec $m = 20$ et $m = 41$, sauf pour de très petites amplitudes d'erreur modèle où l'algorithme localisé est légèrement moins efficace.

Ces deux expériences valident donc la qualité de la localisation appliquée à l'IEnKF-Q. L'algorithme utilisé, qui a été dérivé empiriquement en combinant les algorithmes de l'IEnKS local et de l'IEnKF-Q global, s'avère efficace pour réduire la taille de l'ensemble sans impacter les performances.

4.2.6 Conclusion des expériences avec erreur modèle

Dans toutes les expériences présentées ici, les résultats de l'IEnKF-Q sont meilleurs que ceux des EnKFs et des IEnKFs adaptés pour tenir compte de l'erreur modèle, à l'exception de l'IEnKF-Det qui donne des résultats équivalents dans des régimes faiblement non-linéaires ou lorsque l'ensemble est suffisamment grand. La supériorité de l'IEnKF-Q dans les cas faiblement non-linéaires est étonnante étant donné que le minimum de la fonction de coût est atteint en une unique itération ce qui est censé limiter l'intérêt des méthodes itératives. Cela est dû au recours aux anomalies d'ébauche augmentées dans l'analyse.

Nous rappelons que l'erreur modèle considérée ici était additive. En fonction de la nature de l'erreur modèle, d'autres méthodes peuvent être envisagées. Par exemple, si la nature de l'erreur modèle se manifeste sous la forme d'un biais systématique, [Emili et al. \(2016\)](#) présentent un algorithme, équivalent à celui de [Piccolo et Cullen \(2016\)](#), qui permet de l'estimer et ainsi débiaiser les prévisions.

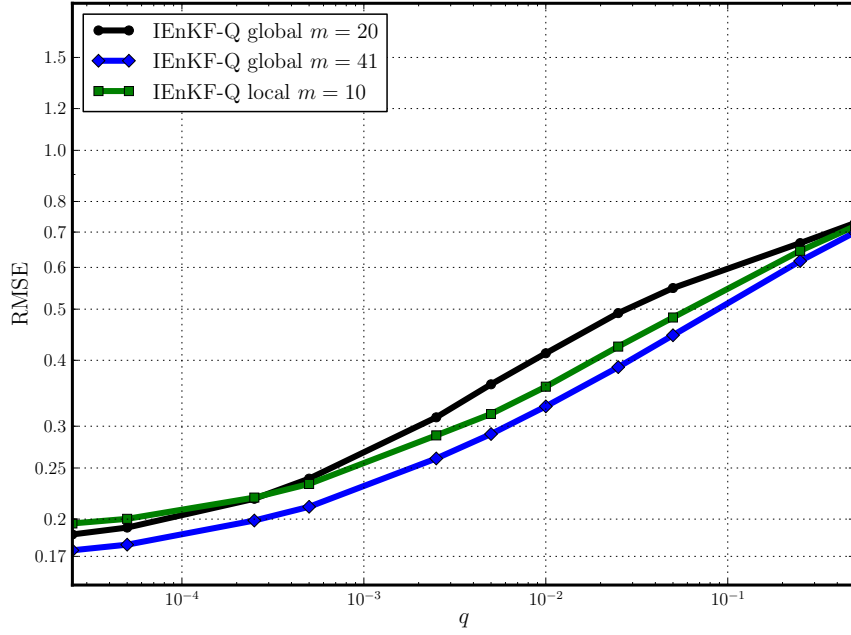


FIGURE 4.12 – RMSEs de filtrage en fonction de q , l’amplitude de l’erreur modèle dans un cas faiblement non-linéaire, pour l’IEEnKF-Q global ($m = 20$ et $m = 41$) et local ($m = 10$). $T = 1$, $\mathbf{Q} = qT\mathbf{I}_n$.

4.3 Résultats sur le L95-T

Dans cette section, nous évaluerons l’IEEnKS sur le L95-T. Nous commencerons par présenter les résultats d’expériences avec estimation de paramètres issues de [Bocquet et Sakov \(2013\)](#). L’intérêt de l’adjonction d’un traceur au L95 est toutefois de permettre de tester le comportement des méthodes d’assimilation de données face au transport des concentrations par la météorologie. À cette fin, nous varierons dans un premier temps l’amplitude du transport, puis nous nous placerons dans un contexte *offline* et étudierons l’impact du passage d’un couplage fort à un couplage faible sur les performances de l’IEEnKS.

Toutes les expériences présentées dans cette section utilisent $S = 1$.

4.3.1 Estimation de paramètres

Des expériences d’assimilation de données similaires à celles réalisées sur le L95 ont été effectuées sur le L95-T, à nouveau dans [Bocquet et Sakov \(2013\)](#). L’intégralité des 80 variables est cette fois-ci assimilée, avec la même variance $\mathbf{R} = \mathbf{I}_{80}$, tous les $\Delta t = 0.05$. Les valeurs de référence d’émission $E = 1$ et de dépôt $\lambda = 0.1$ sont utilisées. Le *finite-size* est appliqué pour ne pas avoir à optimiser l’inflation. De plus, les paramètres de forçage F et d’émission E ont été estimés conjointement. Un problème avec ce modèle vient de la nécessité de garder $c_{m+\frac{1}{2}}$, E et F positifs. Une solution de facilité est de seuiller les concentrations et paramètres négatifs après l’assimilation en les mettant à 0. Malheureusement, cette méthode est sous-optimale car elle rompt la gaussianité et induit un biais positif. Une autre solution est d’appliquer une anamorphose en utilisant le logarithme des concentrations et paramètres, qui sont alors définis sur $\mathbb{R}$. [Bocquet et Sakov \(2013\)](#) remarquent cependant que l’anamorphose sur les concentrations n’est pas nécessaire et se contentent de considérer

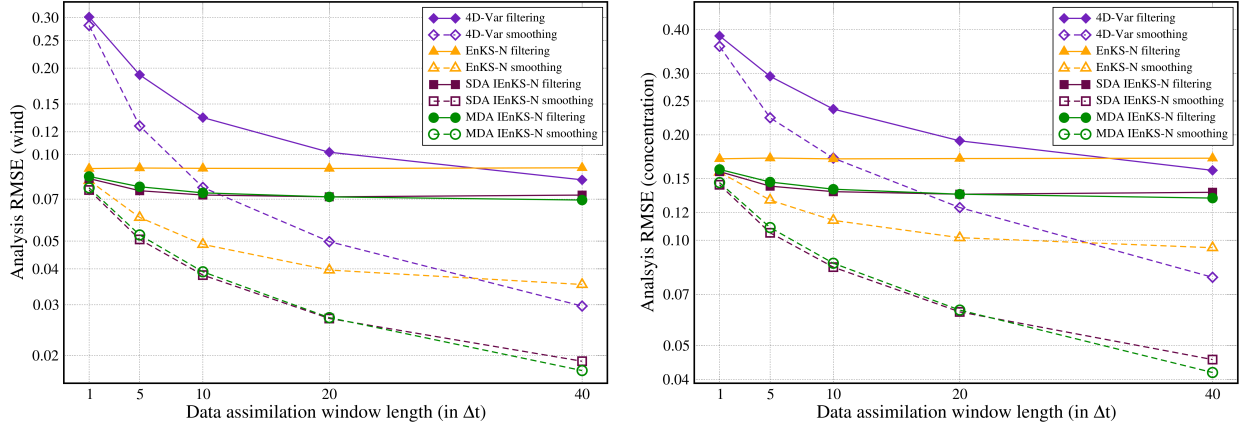


FIGURE 4.13 – RMSEs de filtrage et de lissage en fonction de la longueur de la fenêtre d’assimilation de données (en unité de $\Delta t = 0.05$), pour le vent (à gauche) et pour les concentrations (à droite). Figure tirée de Bocquet et Sakov (2013).

le vecteur d’état

$$[x_1, \dots, x_{39}, c_{\frac{1}{2}}, \dots, c_{39-\frac{1}{2}}, \ln F, \ln E]^T. \quad (4.5)$$

Ils précisent que l’anamorphose sur F n’est pas non plus nécessaire mais qu’elle permet de définir des statistiques homogènes sur les paramètres, notamment une RMSE dont la formule peut alors être

$$\text{RMSE} = \sqrt{\frac{1}{2}(\ln F^a - \ln F^t)^2 + \frac{1}{2}(\ln E^a - \ln E^t)^2}, \quad (4.6)$$

où $F^t = 8$ et $E^t = 1$.

Les résultats en lissage et en filtrage, pour différents algorithmes d’assimilation de données, et pour le vent et les concentrations, sont montrés sur la Fig. 4.13. On remarque dans un premier temps que les scores sur le vent sont meilleurs dans ce contexte que dans l’expérience sur le L95 seul, malgré l’estimation de paramètres. Cela montre l’impact des observations des concentrations pour l’estimation des vents et l’intérêt du couplage. Par ailleurs, le nombre plus grand d’observations peut expliquer la plus faible différence que l’on remarque entre l’IEnKS SDA et MDA.

Les résultats sur les paramètres estimés par différents algorithmes sont représentés sur la Fig. 4.14. On y voit une fois de plus la supériorité de l’IEnKS sous ses différentes formes par rapport au 4D-Var et à l’EnKS.

4.3.2 Régimes d’émission et de dépôt

Les coefficients de dépôt λ et d’émission E régissent la masse totale du traceur dans le domaine. Leur ratio peut donner différents régimes dans la physique du modèle. Dans le cas de référence précédent ($E = 1, \lambda = 0.1$), des bouffées de traceur peuvent traverser de longues distances avant d’être déposées. Ainsi les observations des concentrations de traceur en un point de grille donné contiennent de l’information sur la force et la direction du vent en un point de grille différent, plusieurs pas de temps dans le passé. Haussaire et Bocquet (2016) présentent une expérience pour illustrer ce phénomène. L’ordre de grandeur du coefficient de dépôt λ est varié, en même temps que le coefficient d’émission E , de telle sorte à garder le ratio E/λ constant. De ce fait, les ordres de grandeur des concentrations restent inchangés, et garder une matrice des erreurs d’observation constante au cours de l’expérience ne change alors pas l’erreur relative des observations des

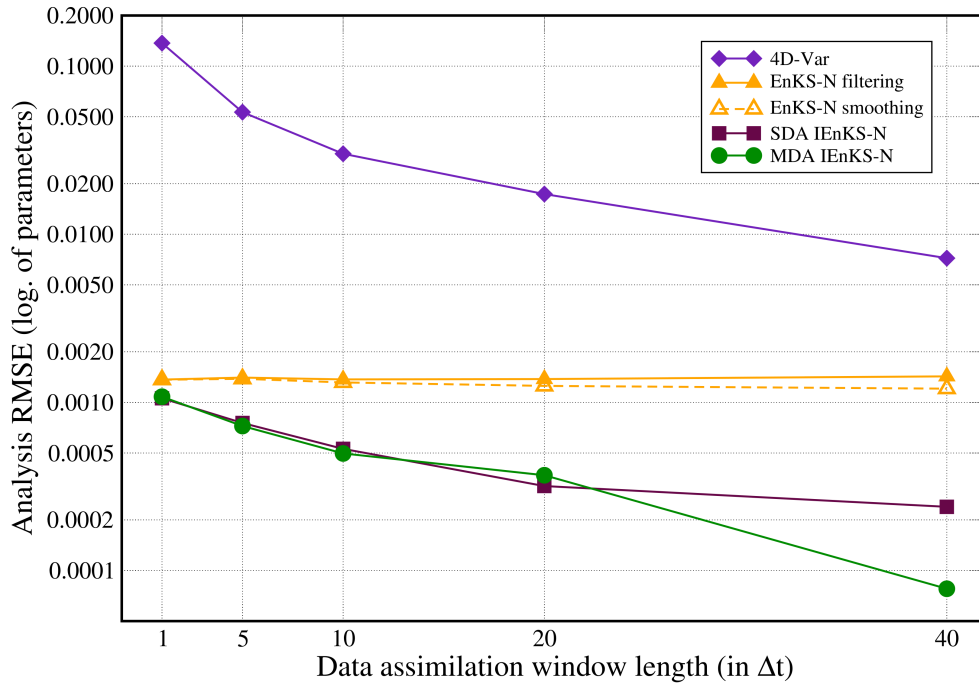


FIGURE 4.14 – RMSEs de filtrage et de lissage en fonction de la longueur de la fenêtre d’assimilation de données (en unité de $\Delta t = 0.05$), pour les paramètres estimés. Figure tirée de [Bocquet et Sakov \(2013\)](#).

concentrations.

Les RMSEs de filtrage des concentrations et du vent sont tracées sur la Fig. 4.15, en fonction du coefficient de dépôt. Dans cette expérience, on a utilisé l’IEnKS SDA avec $L = 0$ (équivalent à un ETKF, dans sa formulation maximum likelihood ensemble filter (MLEF)) et $L = 15$.

Pour des petites valeurs du dépôt, la RMSE reste relativement constante pour le vent et les concentrations, proche du cas de référence ($\lambda = 0.1$; $E = 1$). Cependant, dès que $\lambda > 1$, le vent n’a plus le temps de transporter le traceur sur des distances suffisantes. Les observations de concentration contiennent alors moins d’information sur le vent et la RMSE sur le vent augmente en conséquence. Parallèlement, l’absence de transport du traceur réduit l’impact négatif de la dispersion par le vent et rend l’estimation des concentrations plus facile, notamment car le traceur est stable. Par ailleurs, l’intérêt de la grande fenêtre d’assimilation de données $L = 15$ est amplifié avec des petits coefficients de dépôt car le traceur est advecté plus loin dans l’espace. Comme déjà mis en avant, les observations des concentrations contiennent de l’information sur le vent dans le passé, avant que le traceur ne soit advecté par ces vents. On a ainsi vérifié l’avantage des méthodes variationnelles développées sur une longue fenêtre par rapport à l’EnKF dans ce contexte.

4.3.3 Stratégies pour des modèles faiblement couplés

Nous rappelons dans un premier temps la distinction entre un modèle *offline* et un modèle *online*. Pour un modèle comme le L95-T tel qu’il a été présenté jusque maintenant, la météorologie et le transport du traceur sont intégrés numériquement en même temps. Il s’agit alors d’un modèle couplé *online*, ou modèle couplé de chimie-météorologie (CCMM). Toutefois, le L95 pourrait être intégré séparément. Les champs de vent seraient fournis en tant que données d’entrée au traceur pour effectuer le transport. Il s’agirait alors d’un modèle *offline*, tel un modèle de chimie-transport

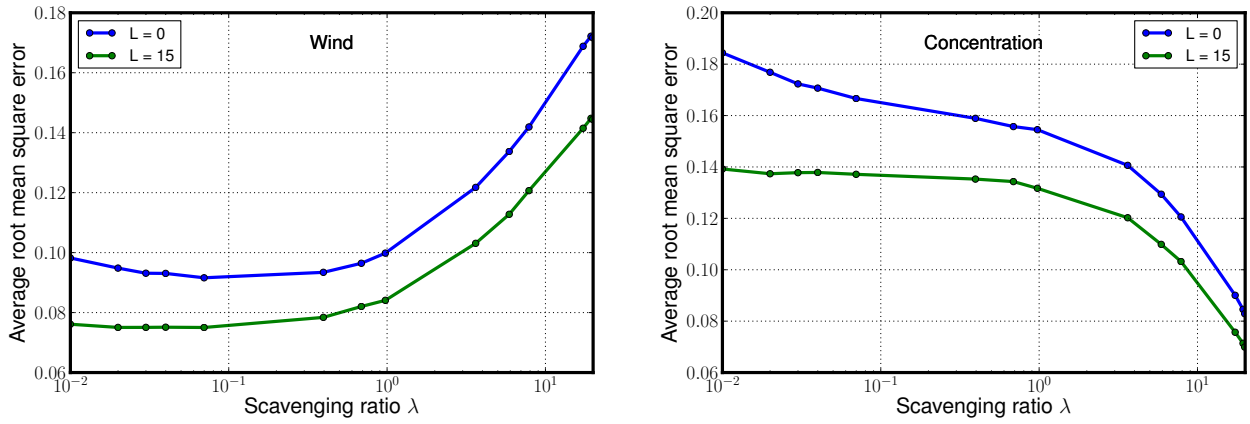


FIGURE 4.15 – RMSEs de filtrage pour l’IEnKS SDA avec $L=0$ et $L=15$, en fonction du coefficient de dépôt, pour le vent (*Wind*, à gauche) et pour les concentrations (à droite). Figure tirée de Haussaire et Bocquet (2016).

(CTM). Pour les modèles *offline* et *online*, nous avons donc deux sous-systèmes qui communiquent, à savoir le transport du traceur d’un côté, la météorologie de l’autre. Dans le premier cas, on parle de couplage faible, car il n’y a pas de rétroaction possible du traceur vers la météorologie. Par opposition, on parle de couplage fort pour les modèles *online*. Si on applique de l’assimilation de données sur un modèle *offline*, les informations ne pouvant pas communiquer entre les deux sous-systèmes, les performances vont être impactées.

Le L95-T fournit un cadre d’application simple pour tester les méthodes d’assimilation de données dans des contextes *offline* et *online*. Le cas *online* correspond aux expériences présentées aux sections précédentes. Pour le cas *offline*, on applique un IEnKS sur la partie L95 du modèle avec une longueur d’assimilation L_w . Ceci nous fournit une estimation erronée des vents que nous utilisons pour l’advection du traceur. Nous appliquons également un IEnKS avec une longueur de fenêtre d’assimilation de données L_c sur les concentrations. Ainsi, les observations de concentrations du traceur ne pourront pas aider à améliorer l’estimation du vent. De plus, les incertitudes sur le vent constituent une source réaliste et tangible d’erreur modèle pour le sous-système du traceur.

On présente ici une série d’OSSEs ayant été effectuées sur le L95-T *offline*. Les paramètres sont identiques à ceux du paragraphe 4.3.1, où on fixe les émissions et le dépôt aux valeurs de référence. On précise que dans le cas *offline*, de l’erreur modèle est introduite par le découplage. Le *finite-size* permet de corriger les erreurs d’échantillonnage en diagnostiquant une inflation optimale, mais cette inflation ne corrige pas l’erreur modèle. Aussi est-il nécessaire d’ajouter de l’inflation multiplicative, cette dernière ayant été optimisée pour minimiser la RMSE. L’intérêt de conserver le *finite-size* en plus de l’inflation multiplicative permet de quantifier séparément les deux type d’inflations, celle pour les erreurs d’échantillonnage et celle pour l’erreur modèle.

Nous avons plusieurs systèmes *offline* présentés. Dans le premier, noté Offline 1a, on fournit le vent analysé moyen à l’IEnKS sur le sous-système du traceur, aussi bien pour l’analyse que pour la propagation. Dans cette référence, on choisit $L_c = L_w$. Dans une deuxième variante, nommée Offline 1b, les vents moyens issus d’un ETKF sont fournis, c’est-à-dire $L_w = 0$, et on fait varier L_c . La troisième possibilité (Offline 1c) prend les vents moyens issus d’un IEnKS avec un L_w donné et un ETKF est appliqué sur le traceur, c’est-à-dire $L_c = 0$ et L_w varie.

Dans une dernière variante *offline*, appelée Offline 2, les vents moyens issus d’un IEnKS sont utilisés pour l’étape d’analyse de l’assimilation de données sur le traceur, tandis que l’ensemble des

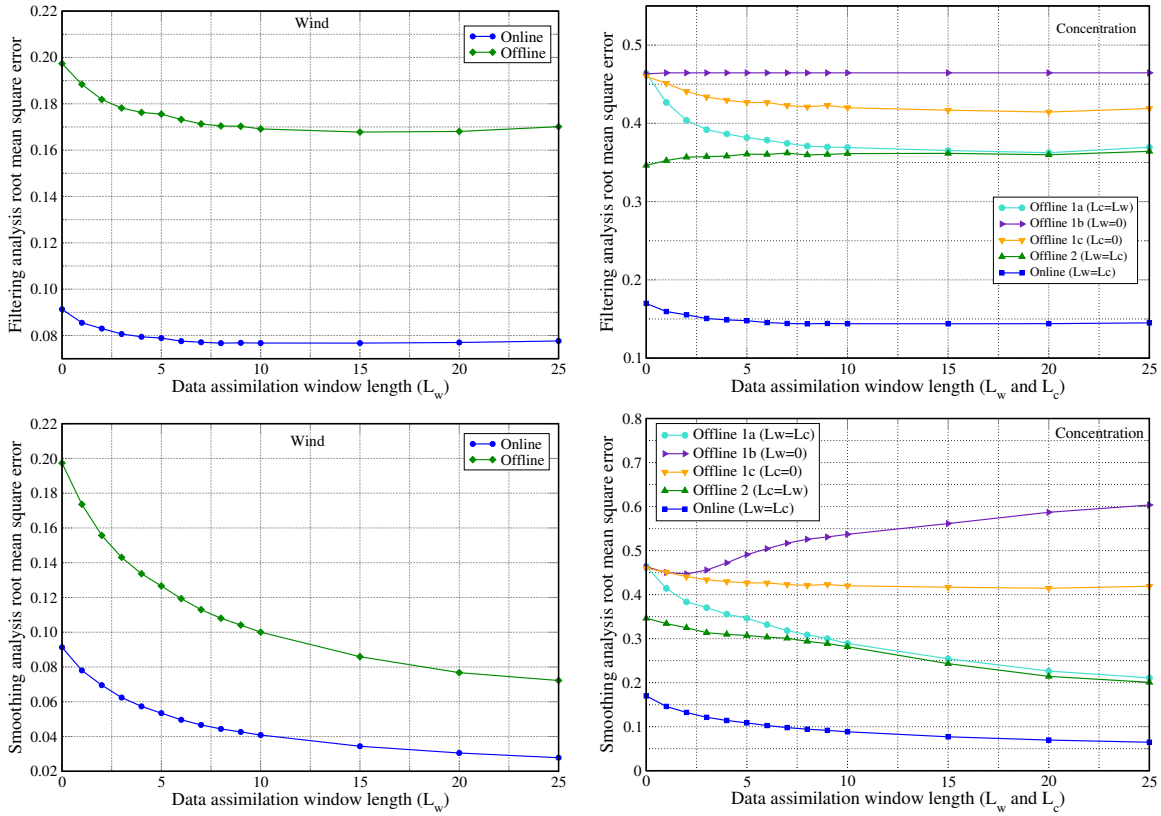


FIGURE 4.16 – RMSEs des expériences d’assimilation de données avec l’IEnKS sur le L95-T, en fonction de la longueur de la fenêtre d’assimilation de données (en unité de $\Delta t = 0.05$). Les figures du haut représentent les RMSE de filtrage, celles du bas de lissage. Les scores sur les variables du vent sont à gauche (*Wind*), tandis que les scores sur les concentrations sont à droite. Le cas $L = 0$ correspond à l’ETKF. Figure tirée de [Haussaire et Bocquet \(2016\)](#).

vents analysés est utilisé pour l’étape de propagation. Si la dispersion de l’ensemble des vents est représentative de l’incertitude sur la météorologie, on peut alors espérer que l’erreur modèle soit proprement comptabilisée lors de l’étape de propagation. Dans cette variante, on prend à nouveau $L_c = L_w$.

Les résultats en fonction de la longueur de la fenêtre d’assimilation de données, sur les vents et les concentrations du traceur, en filtrage et en lissage, pour les différentes expériences *offline*, sont montrés sur la Fig. 4.16. À titre de comparaison, une expérience *online* a également été effectuée et est tracée sur la figure. Tout d’abord, on voit que le système *online* a un net avantage sur chacun des systèmes *offline* grâce au transfert d’information entre les deux sous-systèmes. Ce résultat est en accord avec [Semane et al. \(2009\)](#) qui ont assimilé des observations réelles d’ozone stratosphérique. Pour toutes les expériences *offline*, le vent ne peut pas tirer partie des observations d’ozone. Nous nous concentrons désormais sur le traceur.

Nous insistons à nouveau sur le besoin de tenir compte de l’erreur modèle, soit en implémentant une inflation multiplicative (en plus du *finite-size*), soit par le biais choisi par l’expérience Offline 2. Si ce n’est pas fait, le filtre diverge et une RMSE d’environ 0.65 est obtenue, ce qui correspond à une expérience sans assimilation de données (erreur climatologique).

Dans les cas avec inflation multiplicative Offline 1a, 1b, 1c, le premier fournit une référence, dont la RMSE diminue progressivement quand la longueur de la fenêtre augmente. Cependant, la

RMSE reste loin du cas *online*, car le vent utilisé est tout de même moins bon.

Dans la configuration Offline 1b, où le vent moyen est issu d'un ETKF ($L_w = 0$), la RMSE du traceur en filtrage ne bénéficie pas de l'extension de la fenêtre d'assimilation, tandis que le lissage en profite pour de petites fenêtres (jusqu'à $L_c = 2$) avant de se dégrader. Cela est probablement dû à la dynamique stable et linéaire du transport, ainsi qu'à l'erreur modèle.

Dans le cas Offline 1c, l'assimilation de données sur le traceur consiste en un ETKF, tandis que le vent s'améliore avec la longueur de la fenêtre L_w . Ainsi, l'amélioration que l'on observe provient de l'erreur réduite sur le vent. Dans cette configuration, le filtrage et le lissage sont identiques car $L_c = 0$.

À la vue des deux résultats précédents, il apparaît que l'amélioration que présentait le traceur dans l'expérience Offline 1a provenait d'une meilleure estimation des vents au sein de la fenêtre d'assimilation. Quand $L_c = L_w$ grandit, l'analyse du traceur au sein de la fenêtre bénéficie de meilleurs vents lissés, ce qui explique l'amélioration plus importante que dans le cas Offline 1b, qui n'utilise que les vents filtrés. On conclut donc qu'une longue fenêtre d'assimilation de données est inutile pour le transport *offline* étant donnée la dynamique linéaire et stable du système.

Dans l'expérience Offline 2, on prend en compte l'erreur modèle d'une part avec de l'inflation mais surtout en ayant un ensemble de vents utilisé dans l'étape de propagation, ensemble supposé représenter l'incertitude sur le champ de vent. Ceci est similaire à de la paramétrisation stochastique où les paramètres d'entrée des modèles sont différents pour chaque membre de l'ensemble de CTMs, comme dans Constantinescu *et al.* (2007b) ou Wu *et al.* (2008). Cette expérience fournit de bien meilleurs résultats, obtenus sans inflation pour le lissage et avec une inflation croissant avec L_c (allant jusqu'à 1.06 pour $L_c = 25$) pour le filtrage.

Cette expérience montre l'intérêt du L95-T pour tester des méthodes d'assimilation de données *offline* et de paramétrisation de l'erreur modèle. Elle montre également les limites du recours à des méthodes variationnelles sur des longues fenêtres d'assimilation lorsque la dynamique est stable et linéaire.

4.4 Assimilation de données sur le L95-GRS

L'intérêt du L95-GRS par rapport au L95-T est d'introduire de la chimie non-linéaire et d'évaluer les performances des méthodes d'assimilation de données dans ce contexte. Des expériences d'assimilation avec le L95-GRS ont donc été présentées dans Haussaire et Bocquet (2016). Nous nous plaçons avec les paramètres décrits dans la section 3.3.1.4. Plus précisément, nous rappelons que nous sommes en régime ROC-limité, même si quelques expériences ont également été réalisées en régime NO_x -limité. Parmi ces expériences en régime NO_x -limité, nous avons remarqué que l'assimilation de données dans ce contexte est moins efficace pour estimer les concentrations de ROC, et le contraire est vrai pour le NO_2 . Ceci s'explique probablement par le fait que les concentrations de NO_2 , dans ce cas, contrôlent totalement le modèle alors qu'une erreur d'estimation des ROC n'a que peu d'impact sur les autres espèces.

Le but de ces expériences est de prouver la richesse du L95-GRS en tant que modèle jouet développé pour l'assimilation de données. Dans ces OSSEs, des observations sont tirées de la vérité toutes les $\Delta t = 6$ h et perturbées selon un bruit gaussien. L'écart type de ce bruit est choisi afin de correspondre à environ 10% de la concentration moyenne sur le domaine. Concrètement, nous avons choisi une matrice des erreurs d'observation de la forme $\mathbf{R} = \sigma^2 \mathbf{I}$ avec $\sigma^2 = 1$ pour le vent, en unité de Lorenz, $\sigma^2 = 0.01 \text{ (ppb C)}^2$ pour le ROC, $\sigma^2 = 0.16 \text{ ppb}^2$ pour le NO, $\sigma^2 = 1 \text{ ppb}^2$ pour le NO_2 , $\sigma^2 = 4 \text{ ppb}^2$ pour le O_3 , et $\sigma^2 = 0.01 \text{ ppb}^2$ pour le S(N)GN. Quand une RMSE sera

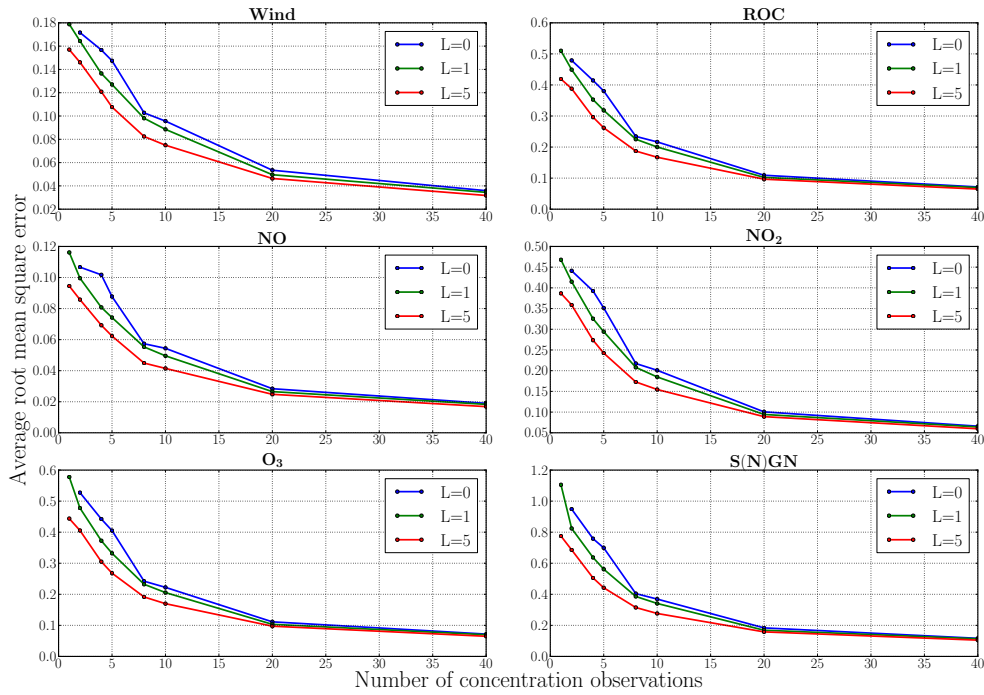


FIGURE 4.17 – RMSEs de filtrage des variables du L95-GRS en fonction du nombre d’observations de concentrations pour l’IEnKS avec 3 longueurs de fenêtre d’assimilation de données. Le cas $L = 0$ correspond à l’ETKF. Les RMSEs sont normalisées par l’écart type de l’erreur d’observation. Les variables du L95, notées *Wind*, sont montrées dans l’unité du Lorenz, les concentrations sont en ppb (ppbC pour le ROC). Figure tirée de Haussaire et Bocquet (2016).

tracée, elle aura été normalisée par l’écart type de l’erreur d’observation. La taille de l’ensemble est maintenue à $m = 20$.

4.4.1 Premières performances

On évalue la performance de l’IEnKS, avec différentes longueurs de fenêtre, quand on n’assimile qu’une partie des observations de concentration. Dans cette expérience, l’intégralité des 40 points de grille est observée pour le vent, tandis que seule une fraction d’entre eux est gardée pour observer les concentrations. On choisit de répartir les observations des concentrations uniformément sur le domaine, de telle sorte que le nombre d’observations soit un diviseur de 40, à savoir $\{1, 2, 4, 5, 8, 10, 20, 40\}$.

Les RMSEs de filtrage sur les concentrations et le vent en fonction du nombre d’observations sont tracées sur la Fig. 4.17. On voit que si on observe trop peu le système, l’état estimé est imprécis, au point que l’ETKF diverge quand on n’observe les concentrations qu’en un point de grille. Même l’IEnKF (avec $L = 1$) ne réussit pas à estimer le S(N)GN mieux que l’erreur d’observation, en moyenne sur le domaine.

Pour être plus réaliste, nous ne garderons dans la suite des expériences que 8 observations, soit 1 tous les 5 points de grille. Dans ce contexte, on a fait varier plus largement la longueur de la fenêtre d’assimilation de données et nous présentons les RMSEs de filtrage et de lissage du vent et des concentrations sur la Fig. 4.18. On voit donc que même si le modèle est non-linéaire, l’IEnKS peut prendre en compte ces non-linéarités grâce à l’analyse variationnelle à l’intérieur de la fenêtre.

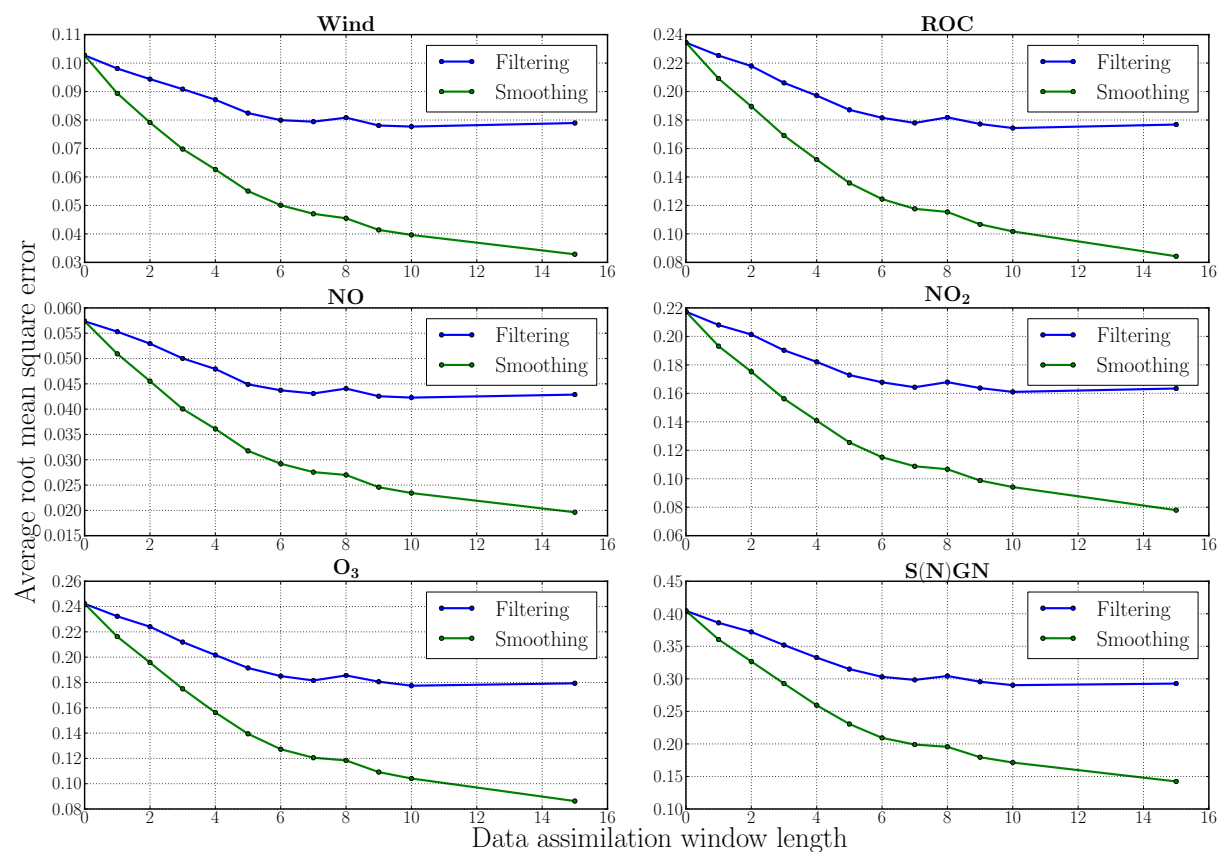


FIGURE 4.18 – RMSEs de filtrage et de lissage pour les variables du L95-GRS, en fonction de la longueur de la fenêtre d'assimilation de données (en unité de $\Delta t = 6$ h). Le cas $L = 0$ correspond à l'ETKF. Les RMSEs sont normalisées par l'écart type de l'erreur d'observation. Les variables du L95, notées *Wind*, sont montrées dans l'unité du Lorenz, les concentrations sont en ppb (ppb C pour le ROC). Figure tirée de Haussaire et Bocquet (2016).

Ainsi, les scores s'améliorent avec L en filtrage comme en lissage. Le S(N)GN est l'espèce la moins bien estimée, probablement car elle est assez peu corrélée avec les autres espèces, à part le ROC, et une erreur d'estimation du S(N)GN n'impacte pas les autres espèces.

4.4.2 Estimation de paramètres

Comme il a été expliqué au chapitre introductif, il y a de nombreuses sources d'incertitudes au sein des CTMs et CCMs, comme les conditions aux limites (Roustan *et al.*, 2010), les champs météorologiques (Dawson *et al.*, 2007) ou les émissions de polluants et de leurs précurseurs (Cohan *et al.*, 2005). Le L95-GRS étant un modèle avec un domaine cyclique, il n'y a pas de conditions aux limites pouvant être estimées, mais les deux autres paramètres peuvent faire l'objet d'une telle étude.

Nous utilisons un système augmenté, comme avec le L95-T, auquel on applique l'IEnKS pour estimer conjointement les émissions de NO_x , de ROC ainsi que le forçage F du L95. L'état augmenté contient donc 243 variables dans cette expérience. Les paramètres sont initialisés de la sorte :

- pour les émissions : l'ensemble est initialisé autour de la vérité par l'ajout d'un bruit gaussien centré d'écart type 10% de la vérité ;
- pour F : l'ensemble est initialisé autour de la valeur $F = 7$ par l'ajout d'un bruit gaussien centré d'écart type 10% de la vérité.

Dans cette expérience, on a eu recours à l'IEnKS MDA.

Comme pour le L95-T, l'estimation des variables n'est pas affectée par l'estimation conjointe des paramètres. Ainsi, les RMSEs montrées sur la Fig. 4.18 tiennent toujours dans ce contexte. Les valeurs estimées des paramètres en fonction du temps sont tracées sur la Fig. 4.19. On voit que seulement quelques jours sont nécessaires à la convergence vers la valeur du paramètre F , tandis que quelques dizaines de jours sont requis pour que les émissions se stabilisent. Cela est dû à la forte sensibilité des variables de vent et de concentration vis-à-vis du paramètre F , qui contrôle la chaoticité du modèle, ainsi qu'au fait qu'il y ait un biais sur la valeur initiale de F . À la fin d'une longue expérience d'assimilation de données, l'algorithme converge vers la bonne valeur du paramètre avec une précision de l'ordre du pourcent. Le recours à une longue fenêtre d'assimilation permet d'améliorer l'estimation des paramètres et la régularité de la solution. Le cas $L = 1$ montre que la méthode peut parfois converger vers une valeur biaisée pour une assez longue période.

Il pourrait être envisagé d'estimer les coefficients de réaction chimique, par exemple la photolyse du ROC. Cependant, nos tests ont montré la divergence du filtre dans ce cas là. Cela arrive probablement car le taux de photolyse est égal à 0 la nuit. Ainsi, il est nécessaire de définir une évolution a priori pour ce type de paramètres et éviter la divergence.

4.4.3 Localisation

Nous avons voulu vérifier que le L95-GRS se prêtait à l'utilisation de méthodes de localisation. Pour cela, nous avons utilisé le deterministic EnKF (DEnKF), présenté en section 2.4.2.1, en lui appliquant une localisation par produit de Schur. L'intérêt est que le DEnKF est efficace en terme de complexité et la localisation par produit de Schur est facilement implémentable. Les RMSEs en fonction de la taille de l'ensemble, avec et sans localisation, avec inflation et rayon de localisation optimaux, sont montrés sur la Fig. 4.20. On voit que la localisation nous permet de réduire la taille de l'ensemble en deçà de 18 sans provoquer la divergence du filtre, même si les scores se dégradent pour de très petits ensembles ($N < 10$).

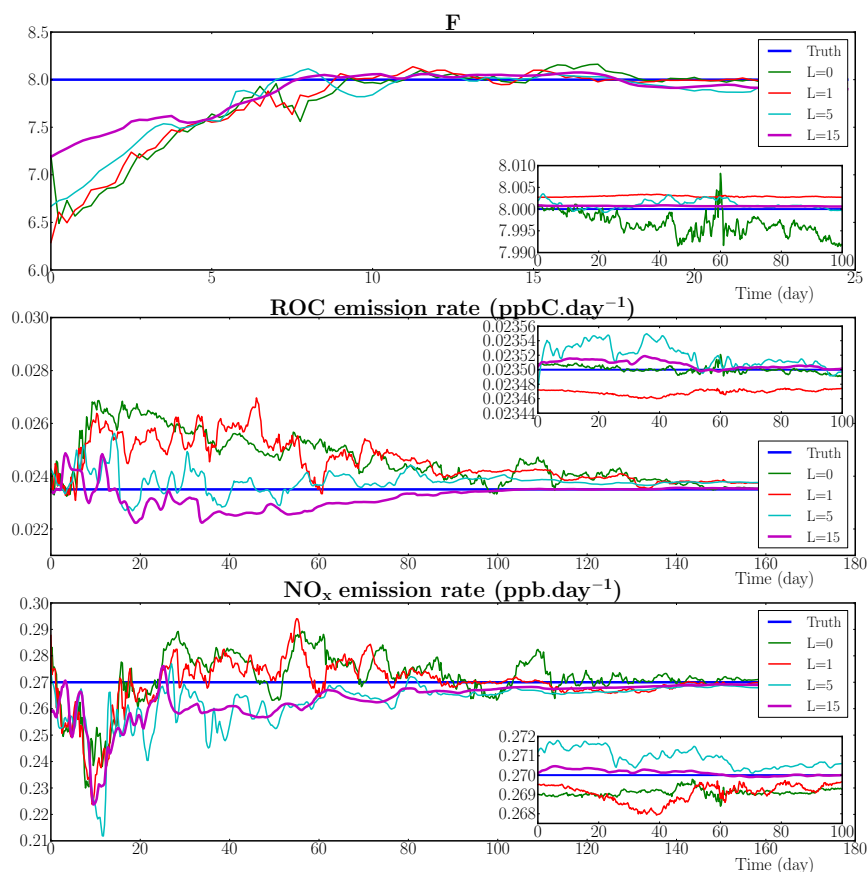


FIGURE 4.19 – Évolution temporelle (en jour) de la valeur des paramètres pour différentes longueurs de fenêtre d'assimilation de données sans spin-up (image principale) ou après un long temps (image insérée). Le cas $L = 0$ correspond à l'ETKF. F est montré en unité de Lorenz, tandis que les émissions sont en ppbC jour⁻¹ ou ppb jour⁻¹. Figure tirée de Haussaire et Bocquet (2016).

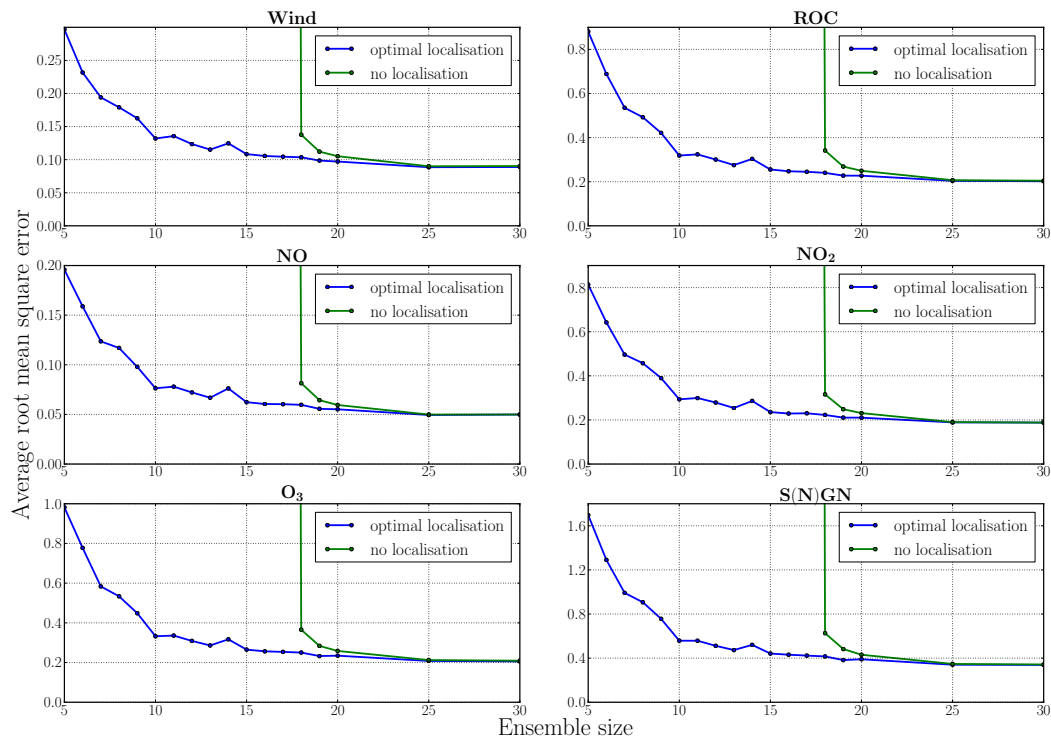


FIGURE 4.20 – RMSEs de filtrage pour les variables du L95-GRS, en fonction de la taille de l'ensemble pour le DEnKF sans localisation ou avec une localisation optimale. Les variables du L95, notées *Wind*, sont montrées dans l'unité du Lorenz, les concentrations sont en ppb (ppbC pour le ROC). Figure tirée de Haussaire et Bocquet (2016).

4.5 Conclusions

Les modèles présentés au chapitre 3 ont été développés pour servir de bancs d'essai aux méthodes d'assimilation de données. Au cours du présent chapitre, ils se sont avérés efficaces à ce titre en fournissant des contextes d'application variés, offrant des difficultés diverses. Le L95 permet en effet de moduler le degré de non-linéarité et de proposer un contexte d'estimation de paramètres. De plus, malgré sa petite dimension, il permet de tester et valider les méthodes de localisation. Cette petite dimension offre également la possibilité de tester des expériences en contrainte faible, où la matrice de covariance de l'erreur modèle peut être définie explicitement, y compris non-diagonale. Outre la possibilité de tester l'estimation de paramètres, le L95-T permet particulièrement de vérifier l'impact du transport et du découplage du modèle sur l'assimilation de données. Enfin, la performance de l'IEnKS sur un modèle non-linéaire de chimie atmosphérique peut être validée sur le L95-GRS, avant d'essayer de l'appliquer à un CTM complet.

Dans tous les cas testés, l'IEnKS s'est montré supérieur au 4D-Var ou à l'ETKF. L'analyse variationnelle itérative effectuée dans l'algorithme permet de bien tenir compte de la non-linéarité du modèle ce qui lui donne un net avantage face à l'ETKF. De plus, la propagation des erreurs d'ébauche grâce à l'ensemble lui permet d'atteindre des performances supérieures au 4D-Var. L'IEnKF-Q, la formulation en contrainte faible de l'IEnKS, s'est également avéré être efficace pour prendre en compte l'erreur modèle.

Néanmoins, ces modèles jouets permettent également de montrer les limites des algorithmes EnVar 4D, dont l'IEnKS est un archétype. Tout d'abord, la chaotité du L95 empêche le recours à un IEnKS SDA avec $S = L$ pour de trop grandes fenêtres d'assimilation de données. Ce sont pourtant des algorithmes similaires, comme le 4D-EnVar, qui sont implémentés à l'heure actuelle dans les centres opérationnels de Météo-France ou du Met Office. Ensuite, lors de l'application de la localisation, le L95 a permis de montrer la nécessité de recourir à une localisation adaptée des covariances 4D, en exhibant la sous-performance de la localisation par domaines statiques et l'utilité des domaines covariants. Enfin, le L95-T a permis de mitiger l'impact de l'IEnKS dans un contexte *offline* avec un modèle de traceur stable et linéaire faiblement couplé à de la météorologie. Ce premier résultat, corroboré par des tests effectués sur le L95-GRS faiblement couplé (non présentés dans ce chapitre), peut tempérer l'intérêt de l'application d'un algorithme EnVar 4D sur un CTM opérationnel. On rappellera cependant que des résultats encourageants ont été obtenus sur le L95-GRS fortement couplé. De plus, des différences majeures existent entre ce modèle simplifié et un modèle opérationnel. Tout d'abord, la non-linéarité du CTM sera vraisemblablement plus importante, ce qui pourra rendre utile le recours à l'IEnKS. Enfin, l'hypothèse de modèle parfait ne tiendra plus du fait de l'erreur modèle associée à tout système réaliste.

Assimilation de données avec un modèle de chimie-transport

L'objectif des expériences d'assimilation de données sur des modèles jouets, comme présentées au chapitre précédent, est de tester dans un contexte simplifié de nouvelles méthodes avant de les évaluer sur des modèles plus réalistes. Les résultats encourageants qui ont ainsi été obtenus sur le L95-GRS nous amènent à vouloir tester l'itérative ensemble Kalman smoother (IEnKS) sur un modèle complet de chimie atmosphérique.

Pour ce faire, nous avons utilisé le modèle de chimie-transport (CTM) Polair3D de la plateforme Polyphemus (Mallet *et al.*, 2007) pour simuler les concentrations d'ozone sur un domaine européen. Nous allons donc tout d'abord décrire la simulation de référence qui a été utilisée et ses performances par rapport aux observations d'ozone in situ disponibles. Ensuite, nous définirons les paramètres optimaux pour l'application d'un localized ensemble transform Kalman filter (LETKF). Les résultats fournis par le LETKF nous serviront de base de comparaison. Ils mettent en avant un problème d'erreur de représentativité des stations d'observation utilisées, erreur qui sera évaluée grâce à de l'estimation de paramètres. Enfin, nous concluons par des résultats d'expériences avec l'IEnKS. Des premiers résultats exploratoires sont obtenus lors d'expériences synthétiques, puis étendus au cas avec observations réelles. Nous nous sommes alors heurtés à des complications à cause de la forte erreur modèle présente dans le système d'assimilation de données, que même l'IEnKF-Q n'arrive pas à prendre en compte.

Les résultats présentés dans ce chapitre sont des travaux originaux et exploratoires, qui ont fait l'objet d'un poster à l'EGU General Assembly 2016 et au Colloque National d'Assimilation de données 2016.

Sommaire

5.1	Simulation de référence avec Polair3D	102
5.1.1	La configuration du modèle	102
5.1.2	Évaluation en regard des observations	102
5.2	Assimilation de données de référence avec le LETKF	105
5.2.1	Calibrage de $\mathbf{P}^f$ avec EMEP	107
5.2.2	Calibrage de $\mathbf{R}$ pour AirBase	109
5.2.3	Résumé des paramètres et des performances	109
5.3	Estimation des erreurs de représentativité des stations AirBase	111
5.4	Application de l'IEnKS sur Polair3D	115
5.4.1	Expériences synthétiques	116
5.4.2	Expériences avec observations réelles	119
5.4.3	Assimilation des observations débiaisées	122
5.4.4	Assimilation avec l'IEnKF-Q	122
5.4.5	Conclusions de l'application de l'IEnKS sur un CTM	126

5.1 Simulation de référence avec Polair3D

5.1.1 La configuration du modèle

Pour notre simulation, nous avons utilisé le modèle Polair3D de la plateforme Polyphemus (Mallet *et al.*, 2007) sur l'année 2009 et simulons les concentrations d'ozone sur un domaine européen, qui s'étend sur ($9^{\circ}\text{O} - 24^{\circ}\text{E}$, $36^{\circ}\text{N} - 59^{\circ}\text{N}$). Le modèle est eulérien et la résolution horizontale du maillage est fixée à 0.5° selon la latitude et la longitude ; 11 niveaux verticaux sont choisis, dont les limites sont (0 m, 40 m, 90 m, 180 m, 320 m, 600 m, 1000 m, 1400 m, 1900 m, 2400 m, 3000 m, 5000 m).

En ce qui concerne l'intégration numérique de l'équation du transport réactif par le modèle, un *splitting* d'ordre 1 est effectué. Le transport par advection est résolu numériquement avec un schéma *direct space time* (DST) d'ordre 3 et un limiteur de flux de Koren-Sweby. Quant à l'intégration de la partie transport par diffusion turbulente et de la partie chimique, elles sont résolues avec un Rosenbrock d'ordre 2 avec un pas de temps adaptatif de 10 min.

Pour chacune des configurations présentées au chapitre 1, des tests de performances ont été effectués, tests qui nous ont permis de converger vers les paramètres présentés ici. Nous utilisons donc la chimie de Carbon Bond 2005 (CB05) (Yarwood *et al.*, 2005), sans le module d'aérosol, pour diminuer le temps de calcul. Les champs météorologiques sont fournis par le modèle opérationnel du Centre européen pour les prévisions météorologiques à moyen terme (CEPMMT), avec une fréquence de 3 h sur un domaine de résolution 0.25° . Les conditions initiales et aux limites sont extraites des simulations globales du Model for OZone And Related chemical Tracers (MOZART), dans sa version 2.0 (Horowitz *et al.*, 2003). Pour les émissions, l'inventaire European Monitoring and Evaluation Programme (EMEP) est sélectionné (Vestreng *et al.*, 2004). Il est complété par des émissions biogéniques calculées avec le Model of Emissions of Gases and Aerosols from Nature (MEGAN) (Guenther *et al.*, 2006), grâce aux données d'occupation des sols de la Global Land Cover Facility (GLCF). On fixe le coefficient de diffusion horizontale à $100\,000\text{ m}^2\text{ s}^{-1}$ et le coefficient de diffusion verticale est déterminé en utilisant Louis (1979). La modélisation du dépôt utilise Zhang *et al.* (2003).

Cependant, les résultats de simulation avec l'ensemble de ces paramètres se sont avérés biaisés. En effet, pendant que les observations d'ozone au sol issues d'AirBase, marquées comme rurales et de fond, enregistraient une concentration moyenne sur le domaine d'environ $68\,\mu\text{g m}^{-3}$ au cours du mois d'août, notre simulation enregistre $75\,\mu\text{g m}^{-3}$. Nous avons donc étendu notre étude de sensibilité en appliquant des facteurs correctifs à certains des champs d'entrée utilisés jusqu'ici. Ainsi, les conditions aux limites de MOZART ont été multipliées empiriquement par un facteur $68/75$. Une autre source d'incertitude provenant des champs d'émission et en particulier des émissions de NO_x , ces dernières ont été multipliées par un facteur 1.2. Enfin, le dépôt d'ozone a également été multiplié par un facteur 1.7. Cette combinaison de modifications a conduit aux meilleures performances du modèle. Tous ces paramètres et données d'entrée sont consignés dans le tableau 5.1.

5.1.2 Évaluation en regard des observations

Pour évaluer nos résultats de simulation, nous les comparons aux observations horaires d'ozone issues des stations d'AirBase marquées rurales et de fond sur les mois estivaux de juin-juillet-août. Nous ne gardons que cette fraction des observations étant donnée la résolution spatiale de notre modèle, qui ne peut alors pas permettre de distinguer des phénomènes plus locaux. La réglementation européenne imposant des concentrations maximales sur un nombre d'heures ou de jours consécutifs fixés, il est alors important de se concentrer sur la capacité de notre modèle à détecter et prévoir ces pics de pollution pour mieux pouvoir les diagnostiquer et les éviter. Cela explique notre choix

Paramètre	Choix
Année	2009
Résolution	0.5°
Domaine	9°O – 24°E ($N_x = 67$) 36°N – 59°N ($N_y = 47$)
Niveaux verticaux	$N_z = 11$
Chimie	CB05 ($N_s = 52$)
Aérosol	non
Météorologie	CEPMMT
Conditions initiales	MOZART 2.0
Conditions aux limites	MOZART 2.0 * $\frac{68}{75}$
Émissions anthropiques	EMEP, avec $\text{NO}_x * 1.2$
Émissions biogéniques	MEGAN
Diffusion verticale	Louis (1979)
Diffusion horizontale	100 000 m ² s ⁻¹
Dépôt	Zhang <i>et al.</i> (2003) * 1.7
Limiteur de flux	Koren-Sweby

TABEAU 5.1 – Compilation des paramètres et données utilisés dans notre simulation de référence du modèle Polair3D.

de se concentrer sur la période estivale pour évaluer notre modèle, d'autant que les niveaux élevés de concentration sont mieux modélisés.

Les stations d'observation mesurent en moyenne sur cette période estivale une concentration d'ozone de $69.5 \mu\text{g m}^{-3}$. Notre modèle, interpolé aux stations, diagnostique quant à lui $65.2 \mu\text{g m}^{-3}$. Il semblerait donc que nous ayons trop fortement débiaisé notre simulation. Cependant, cette configuration reste celle qui donne les meilleurs statistiques parmi celles que nous avons testées. La moyenne des racines de l'erreur quadratique moyenne (RMSEs) sur chaque station est de $25.0 \mu\text{g m}^{-3}$ avec une corrélation de 57.5%. Nous précisons que ces statistiques sont calculées à partir de l'ensemble des observations horaires sur la période. Ces premiers résultats, bien que modérément satisfaisants, restent globalement de l'ordre de grandeur de ce qui est attendu d'un CTM (Bessagnet *et al.*, 2016).

Nous avons tracé sur la Fig. 5.1 l'évolution en temps des concentrations moyennes simulées et observées aux stations. En vis-à-vis, est tracée l'erreur moyenne en fonction du temps, calculée à partir de la formule présentée à la section 2.6.3 où $N_t = p$, le nombre d'observations. Elle correspond à une RMSE calculée avec (2.160) où $N_t = 1$. Nous nous référons à cette statistique indépendamment en tant que RMSE ou erreur moyenne en fonction du temps. On remarque que le modèle tend à sous-estimer les concentrations, surtout au mois de juillet, en accord avec les moyennes citées précédemment. Par ailleurs, la variabilité du modèle est nettement inférieure à celle des observations, qui présentent une plus forte amplitude. Le biais moyen quant à lui oscille entre $10 \mu\text{g m}^{-3}$ et $35 \mu\text{g m}^{-3}$ au cours de l'été. Nous avons également tracé le profil journalier moyen sur la Fig. 5.2 qui fournit la même conclusion, à savoir une amplitude réduite de la simulation et un léger biais négatif.

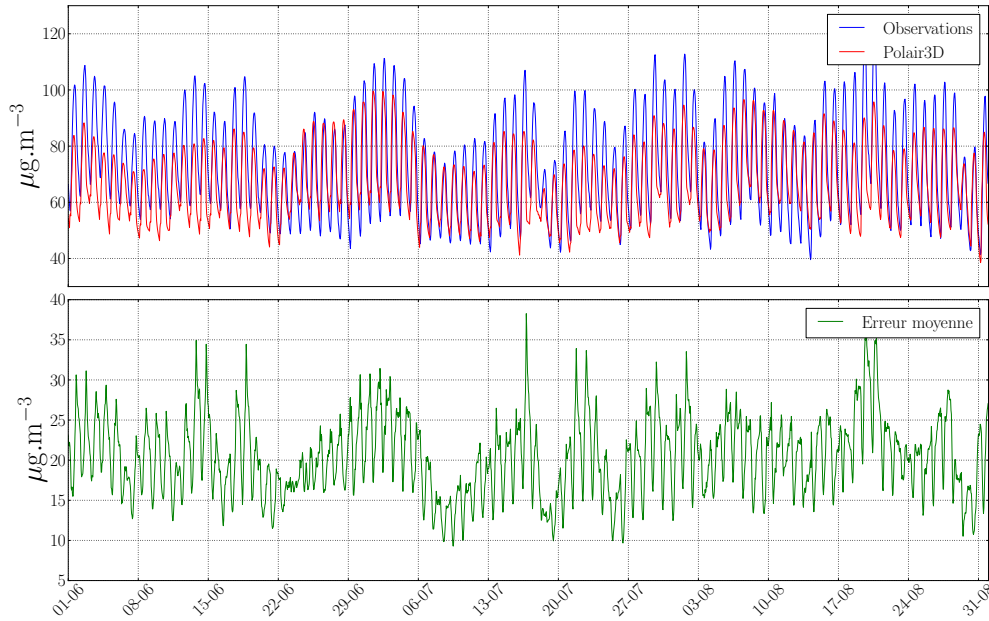


FIGURE 5.1 – Profil horaire d’ozone (en $\mu\text{g m}^{-3}$), moyenné sur l’ensemble des stations d’AirBase rurales et de fond, du 01/06/2009 au 31/08/2009. Les concentrations d’ozone calculées à l’aide du modèle Polair3D sont bilinéairement interpolées sur la position de chacune des stations d’observations. L’erreur moyenne est également indiquée en $\mu\text{g m}^{-3}$.

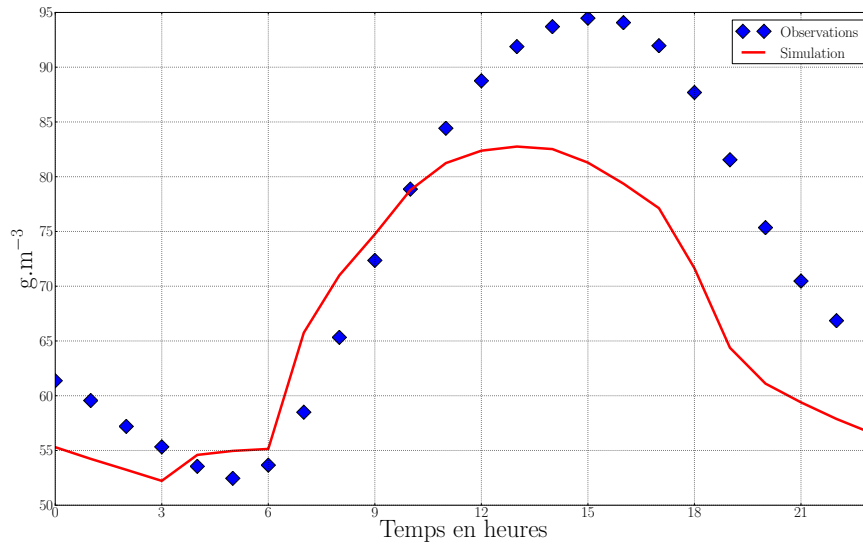


FIGURE 5.2 – Profil journalier d’ozone (en $\mu\text{g m}^{-3}$), moyenné sur l’ensemble des stations d’AirBase rurales et de fond, du 01/06/2009 au 31/08/2009. Les concentrations d’ozone calculées à l’aide du modèle Polair3D sont bilinéairement interpolées sur la position de chacune des stations d’observations. Le 0 correspond à minuit UTC.

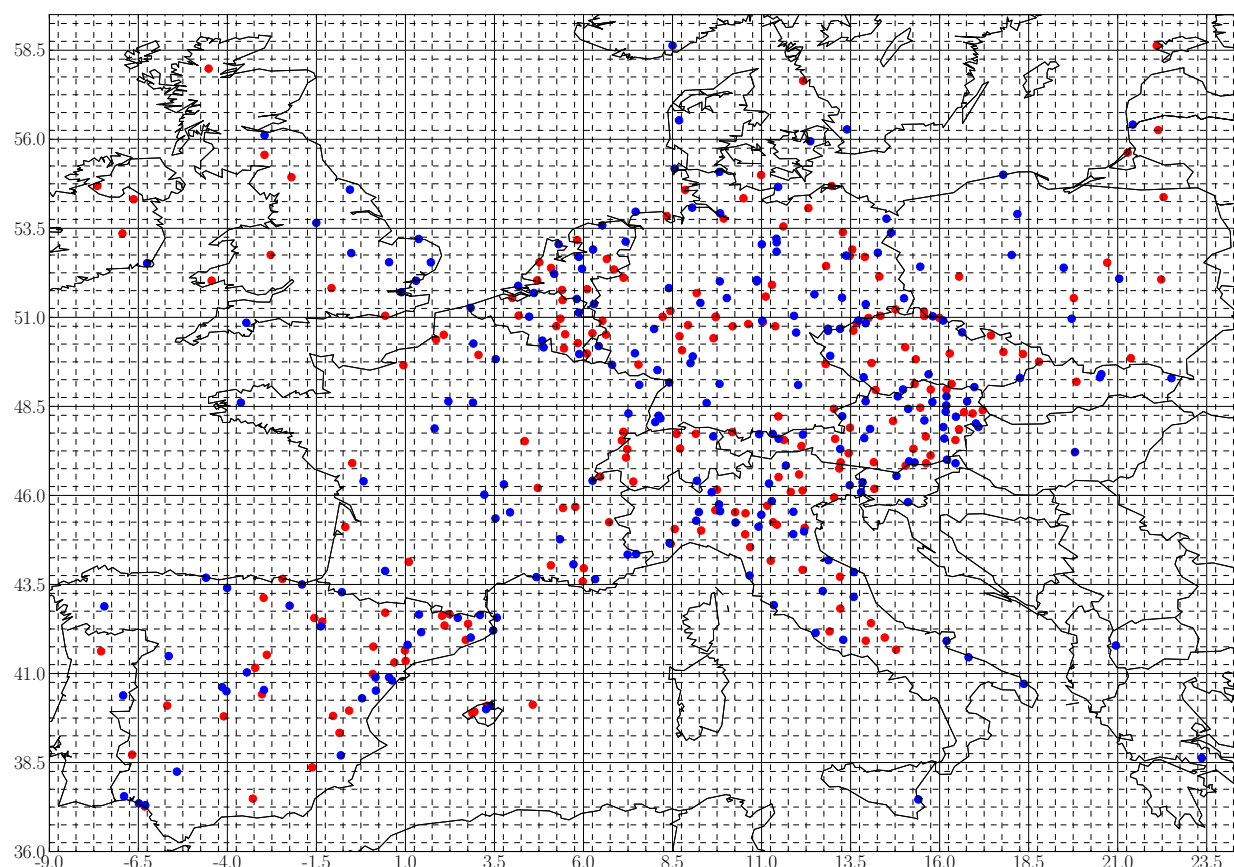


FIGURE 5.3 – Sites de mesure d’ozone du réseau AirBase notés ruraux et de fond, divisés en deux groupes. Les stations du jeu d’observations assimilées sont en rouge tandis que les stations du jeu de validation sont en bleu. Le quadrillage en pointillé a pour résolution 0.5° . Le quadrillage en traits pleins représente les blocs de taille 5.

5.2 Assimilation de données de référence avec le LETKF

Une fois notre simulation de référence définie, nous pouvons assimiler les observations d’ozone d’AirBase avec un LETKF qui utilise donc la localisation dite par domaine (voir 2.4.4). Ces observations sont fournies avec une fréquence horaire.

L’assimilation a lieu sur 72 h, à partir du 14 août 2009 à minuit UTC, et est suivie de 48 h de propagation sans assimilation. Cela nous permet de pouvoir se comparer partiellement aux résultats de Gaubert *et al.* (2014), qui font courir leur expérience sur 10 jours à partir de cette même date. Cette période correspond à une situation anticyclonique en Europe centrale.

Parmi les observations d’ozone au sol d’AirBase, nous ne gardons à nouveau que les observations dites rurales et de fond. Cela représente, sur notre domaine spatial et temporel, 409 stations. Nous les divisons aléatoirement en deux groupes. Le premier groupe fournira le jeu d’observations assimilées. Le deuxième servira de jeu d’observations de contrôle ou de validation. Il permettra de faire une validation croisée, pour évaluer la performance de notre expérience d’assimilation. La répartition des observations assimilées et de validation est présentée sur la Fig. 5.3. On y voit également la résolution du domaine, en quadrillage pointillé.

Nous avons dû définir un certain nombre de paramètres pour notre système d’assimilation, comme la taille de l’ensemble, la variance de l’erreur d’observation, l’amplitude et la corrélation

spatiale de l'inflation additive ajoutée entre chaque cycle ou encore les paramètres associés à la localisation. Nous présentons ici les valeurs retenues pour ces paramètres.

La littérature semble s'accorder à considérer qu'un ensemble de taille comprise entre 20 et 50 membres est raisonnable lorsqu'un algorithme local est appliqué en chimie atmosphérique (par exemple, Hanea *et al.*, 2004; Sandu *et al.*, 2005; Constantinescu *et al.*, 2007a; Wu *et al.*, 2008; Carmichael *et al.*, 2008; Sandu *et al.*, 2011; Boynard *et al.*, 2011; Pagowski et Grell, 2012; Miyazaki *et al.*, 2012; Gaubert *et al.*, 2014). Certaines de ces études explorent l'impact de la taille de l'ensemble sur les résultats d'assimilation. Les améliorations constatées lorsqu'elle est augmentée restent modérées et ne justifient pas le coût de calcul supplémentaire induit (Hanea *et al.*, 2004; Constantinescu *et al.*, 2007a; Miyazaki *et al.*, 2012). Ainsi, nous nous contenterons de 20 membres pour notre ensemble.

Pour initialiser notre ensemble, nous perturbons les concentrations initiales d'ozone de notre modèle en ajoutant un bruit stochastique corrélé en espace de variance $1000 (\mu\text{g m}^{-3})^2$. Cela permet d'avoir un ensemble très dispersé au départ. La nature de la corrélation spatiale est la même que celle utilisée dans l'inflation additive qui est explicitée ultérieurement (équation (5.1)).

Même si cet ensemble est de faible taille, nous avons décidé de garder dans notre vecteur de contrôle l'intégralité des espèces et des niveaux verticaux. Cela permet que des corrélations entre espèces puissent se développer au cours de l'assimilation et ainsi contraindre certaines espèces même si elles ne sont pas observées. Ce choix est avalisé par Miyazaki *et al.* (2012), qui présentent les corrélations entre toutes les espèces chimiques d'un CTM, obtenues avec un LETKF ayant un nombre limité de membres. Cependant, ils ont tout de même appliqué de la localisation en fixant à 0 les corrélations entre espèces faiblement corrélées, car certaines d'entre elles peuvent être erronées et dues aux erreurs d'échantillonnage.

Nous précisons que si jamais une concentration estimée après l'assimilation est négative, cette dernière est seuillée à 0. Il s'avère par ailleurs que c'est le cas également pour l'intégration numérique du modèle, qui fixe à 0 les concentrations négatives après intégration sur un pas de temps. Nous ajoutons que même si avant assimilation, les concentrations des membres de l'ensemble et les observations sont positives, cela n'assure en rien que le résultat de l'assimilation soit positif, d'où le recours au seuillage.

En théorie, lorsqu'un LETKF est utilisé, une analyse par point de grille est effectuée. En pratique, il est possible de regrouper ces points de grille par blocs de taille prédéfinie, et ainsi faire une analyse par bloc. On diminue de la sorte le nombre d'analyses à effectuer, ce qui représente un gain en coût de calcul. Nous avons décidé de prendre des blocs carrés et avons vérifié si les résultats se dégradaient lorsque différentes tailles de bloc étaient appliquées. Le passage d'un bloc de taille 1 (une analyse par point de grille) à un bloc de taille 5 ne semble pas avoir d'impact important sur les résultats d'assimilation, tout en permettant de diviser par près de 25 le nombre d'analyses à effectuer. C'est donc cette taille de bloc qui a été retenue. La position des blocs est représentée sur la Fig. 5.3, en quadrillage plein. On voit que nombre d'entre eux n'ont pas d'observation à l'intérieur.

Pour la localisation par domaines, les erreurs d'observation sont modifiées en fonction de leur distance au point de grille où a lieu l'analyse, de sorte qu'une observation proche du point d'analyse ait plus de poids qu'une très distante. Nous utilisons pour cela la fonction de Gaspari-Cohn (GC), et désignons le paramètre c de la fonction (2.82) comme la longueur de localisation. Dans le cas où des blocs de taille 5 ont été utilisés, la distance est calculée entre l'observation et le centre du bloc où a lieu l'analyse.

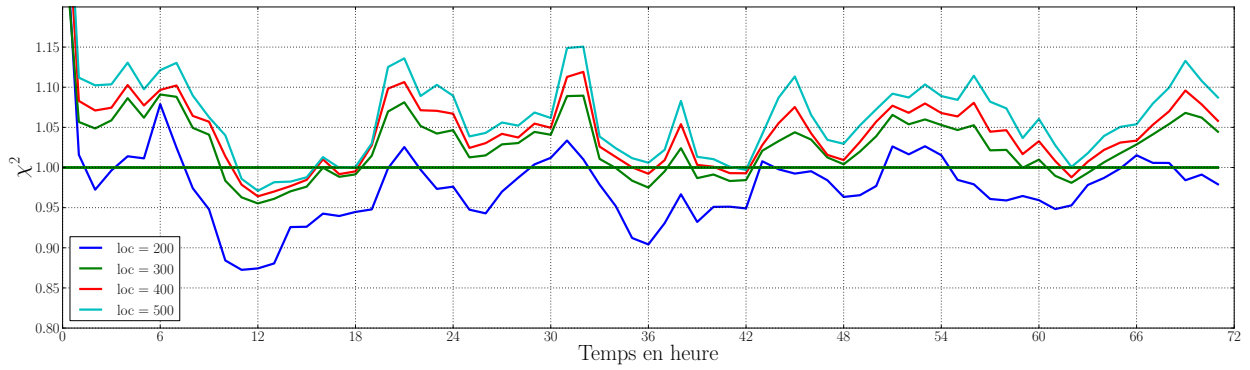


FIGURE 5.4 – Évolution du χ^2 en fonction du temps (en h), pour différentes longueurs de localisation (en km).

5.2.1 Calibrage de $\mathbf{P}^f$ avec EMEP

Les méthodes du χ^2 et du reduced centered random variable (RCRV) présentées dans la section 2.6.4 permettent d'estimer conjointement les amplitudes de $\mathbf{R}$ et $\mathbf{P}^f$. Dans notre cas, $\mathbf{P}^f$ étant fourni par l'ensemble, il s'agit ici d'estimer l'inflation additive qui est appliquée. En effet, entre chaque cycle d'assimilation de données, un bruit stochastique corrélé en espace est ajouté aux champs d'ozone. L'évaluation conjointe de deux hyperparamètres avec le χ^2 conduit à une sous-détermination. Il faut donc être en mesure de fixer l'un d'entre eux pour n'évaluer que l'autre.

Dans un premier temps, nous avons fixé $\mathbf{R} = 10^2 \mathbf{I}_p (\mu\text{g m}^{-3})^2$. Cette valeur est celle que nombre de cas d'étude utilisent (Gaubert *et al.*, 2014; Wu *et al.*, 2008; Hanea *et al.*, 2004; Blond *et al.*, 2003). Chai *et al.* (2007) ont également diagnostiqué la même erreur pour le jour, mais l'augmentent à $13^2 (\mu\text{g m}^{-3})^2$ la nuit. À ce stade, nous avons préféré utiliser le réseau d'observation EMEP, car il est plus fiable. On doit alors fixer la longueur de localisation adaptée, pour ensuite se concentrer sur l'estimation de l'inflation additive grâce au χ^2 .

Cette inflation est incorporée dans les résultats d'analyse à la manière d'un EnKF-Rand, selon l'équation (2.95). Toutefois, nous avons jugé préférable de l'incorporer avant la propagation par le modèle plutôt qu'après, car, ce dernier étant stable, la propagation va permettre de ramener partiellement les membres de l'ensemble perturbé sur l'attracteur. Cela permet également à la perturbation de se propager sur toutes les autres variables chimiques. Nos membres de l'ensemble seront ainsi probablement plus physiques lors de l'analyse. Nous précisons que le bruit additif ajouté entre chaque cycle d'assimilation est corrélé en espace. On applique en effet une fonction de Balgovind (Balgovind *et al.*, 1983), telle que la covariance de l'erreur entre deux points soit

$$f(d_h) = e \times (1 + d_h/L_h) \exp(-d_h/L_h), \quad (5.1)$$

où d_h est la distance entre les deux points, L_h est la longueur de corrélation et e est l'amplitude de l'erreur modèle. En ce qui concerne la longueur de corrélation L_h , nous avons choisi 200 km conformément à Gaubert *et al.* (2014); Boynard *et al.* (2011); Coman *et al.* (2012) et similairement à Chai *et al.* (2007); Constantinescu *et al.* (2007b); Frydendall *et al.* (2009), qui utilisent 270 km. Nous perturbons ainsi les niveaux verticaux jusqu'à 1000 m d'altitude, ce qui correspond en moyenne à la hauteur de la couche limite, et nous corrélons également les perturbations des niveaux verticaux entre elles de la même manière, avec une longueur de corrélation de 3000 m.

Nous avons tout d'abord fixé $e = 15^2 (\mu\text{g m}^{-3})^2$ et fait varier la longueur de localisation. La Fig. 5.4 montre que la longueur optimale semble être légèrement supérieure à 200 km. Nous fixons

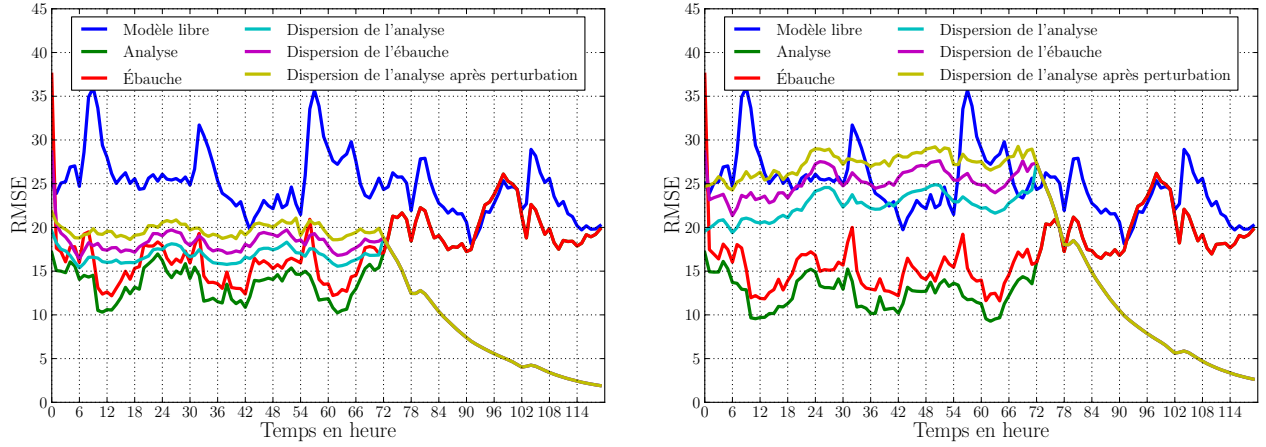


FIGURE 5.5 – Erreur moyenne et dispersion de l'ensemble en fonction du temps, pour une expérience d'assimilation de données avec $e = 10^2 (\mu\text{g m}^{-3})^2$ à gauche et $e = 15^2 (\mu\text{g m}^{-3})^2$ à droite. L'erreur moyenne de la simulation sans analyse est marquée en bleu, l'analyse est en verte, l'ébauche, ou la prévision à 1 h, est en rouge. La dispersion de l'ensemble après analyse est en cyan, après ajout de l'inflation additive en magenta, et après propagation en jaune.

ainsi la longueur de localisation à 200 km, ce qui est proche des 250 km utilisés par [Gaubert et al. \(2014\)](#) par exemple. Cette valeur est par ailleurs cohérente avec le fait que la durée de vie de l'ozone dans l'atmosphère est de quelques jours, ce qui lui laisse le temps de se disperser sur de telles distances.

Il reste donc à estimer la valeur de e , de sorte que le χ^2 soit satisfait et que la dispersion de l'ensemble soit du même ordre de grandeur que l'erreur d'analyse. Nous avons donc testé, avec cette longueur de localisation, le comportement du χ^2 avec $e = 10^2 (\mu\text{g m}^{-3})^2$. Ce dernier est étonnamment peu affecté par ce changement, bien qu'un peu détérioré. En revanche, la dispersion de l'ensemble est en meilleur accord avec l'erreur moyenne montrée sur la Fig. 5.5, sans impacter les résultats en prévision. Cette erreur moyenne est calculée en se comparant avec le même jeu d'observation que celui utilisé dans l'analyse.

Les trois courbes de dispersion de l'ensemble valident le fait que, d'une part, la propagation tend à réduire la dispersion de l'ensemble, conformément à la stabilité du modèle ; d'autre part, l'analyse, elle aussi, réduit la dispersion de l'ensemble. Ces deux phénomènes nous imposent logiquement l'application d'inflation, pour éviter la sous-dispersion de l'ensemble.

Cette expérience nous montre donc que choisir une inflation additive de variance $10^2 (\mu\text{g m}^{-3})^2$ semble adapté. Cette valeur est par ailleurs du même ordre de grandeur que celle prise par [Constantinescu et al. \(2007a\)](#), qui choisissent $12^2 (\mu\text{g m}^{-3})^2$, ou même avec le profil journalier de [Gaubert et al. \(2014\)](#), qui a été validé par le RCRV et des méthodes de Desroziers ([Desroziers et al., 2005](#)).

Ainsi, une fois l'amplitude de l'erreur d'observation fixée et la longueur de la localisation à appliquer aux observations déterminée, nous avons pu estimer l'amplitude de l'inflation additive à appliquer à notre système d'assimilation de données. Cette valeur va donc pouvoir être conservée quand le réseau d'observation sera changé d'EMEP à AirBase. Nous pourrions ainsi nous concentrer sur les valeurs de l'erreur d'observation et la longueur de localisation à appliquer.

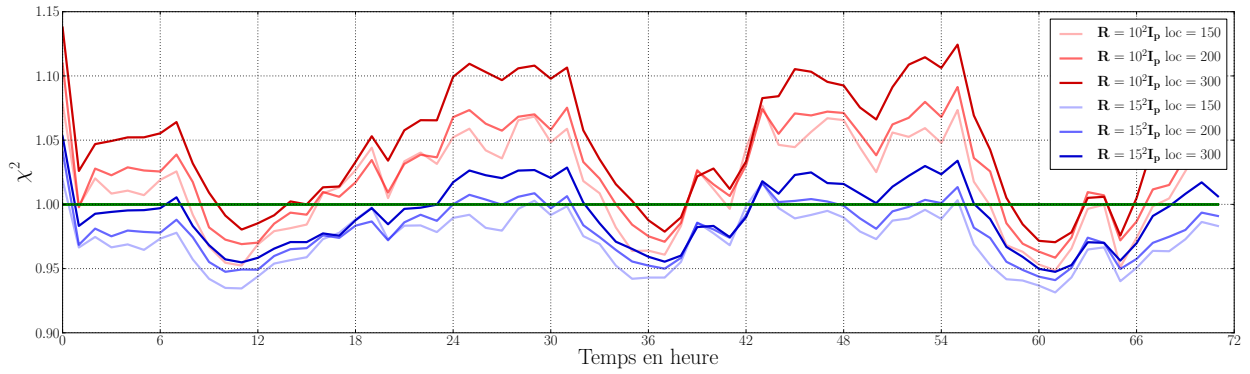


FIGURE 5.6 – Évolution du χ^2 en fonction du temps (en h), pour différentes longueurs de localisation et erreurs d'observation, lorsque les observations d'AirBase rurales et de fond sont assimilées.

5.2.2 Calibrage de $\mathbf{R}$ pour AirBase

Nous conservons donc ici l'écart type de l'inflation additive de $10 \mu\text{g m}^{-3}$, diagnostiquée à la section précédente, et cherchons à évaluer les paramètres de $\mathbf{R}$, à savoir la variance de l'erreur d'observation et la longueur de localisation à appliquer.

Nous considérons désormais les observations issues des stations rurales et de fond du réseau d'observation AirBase, qui sont celles que nous utiliserons in fine dans nos expériences d'assimilation de données. En effet, ce réseau de stations est plus dense que celui d'EMEP, même si probablement moins fiable, car plus entaché d'erreur de représentativité. Notamment, à notre résolution grossière de 0.5° , il peut y avoir plusieurs observations contradictoires au sein d'une même maille du modèle.

Nous faisons donc varier ces deux paramètres conjointement en regardant leur impact sur le χ^2 . Les résultats de la Fig. 5.6 montrent qu'une erreur plus importante est nécessaire pour AirBase que celle choisie initialement avec EMEP. En effet, à moins d'utiliser un écart type de $15 \mu\text{g m}^{-3}$, il est impossible d'avoir un χ^2 inférieur à 1, montrant que notre erreur est effectivement sous-estimée. En prenant un écart type de $15 \mu\text{g m}^{-3}$, en revanche cela devient possible et la longueur de corrélation de 200 km utilisée avec EMEP semble toujours convenir. Cette erreur est donc nettement supérieure à ce qui est habituellement utilisé dans la littérature. Par exemple, Gaubert *et al.* (2014) utilisaient un écart type de $10 \mu\text{g m}^{-3}$ avec les observations d'AirBase dont ils ont requalifié la nature. Nous précisons qu'avec un bloc de taille 5, la distance maximale entre le centre du bloc et le coin de la maille la plus éloignée peut aller au-delà de 200 km. Ainsi, nous choisirons une longueur de localisation de 220 km.

Il s'avère donc que les observations d'AirBase sont effectivement moins fiables que celles d'EMEP et nécessitent d'être retraitées comme Gaubert *et al.* (2014) l'ont fait. Nous expliquons cela par une forte erreur de représentativité de ces observations qui peuvent être influencées par des phénomènes sous-maille que notre modèle n'est pas capable de discerner à résolution grossière. Il est donc intéressant d'essayer d'estimer cette erreur de représentativité pour débiaiser les observations d'AirBase. Ceci est l'objectif de la section 5.3.

5.2.3 Résumé des paramètres et des performances

Nous avons décrit au cours de cette section 5.2 notre système d'assimilation de données de référence. Il utilise les paramètres répertoriés dans le tableau 5.2. Ces paramètres ont été validés par l'application du χ^2 . Les courbes de performance de notre système d'assimilation, avec ces paramètres, sont présentées sur la Fig. 5.7.

Paramètre	Choix
Algorithme	LETKF
Date	14/08/2009
Durée d'assimilation	72 h
Durée de propagation après assimilation	48 h
Nombre de stations d'observation assimilées	198
Nombre de stations d'observation de validation	211
Taille de l'ensemble	20
Taille des blocs	5
Erreur d'observation	$\mathbf{R} = 15^2 \mathbf{I}_p (\mu\text{g m}^{-3})^2$
Localisation	Gaspari et Cohn (1999), 220 km
Variance de l'inflation additive	$10^2 (\mu\text{g m}^{-3})^2$
Corrélation spatiale de l'inflation additive	Balgovind <i>et al.</i> (1983), 200 km

TABLEAU 5.2 – Compilation des paramètres utilisés dans notre système d'assimilation de données de référence.

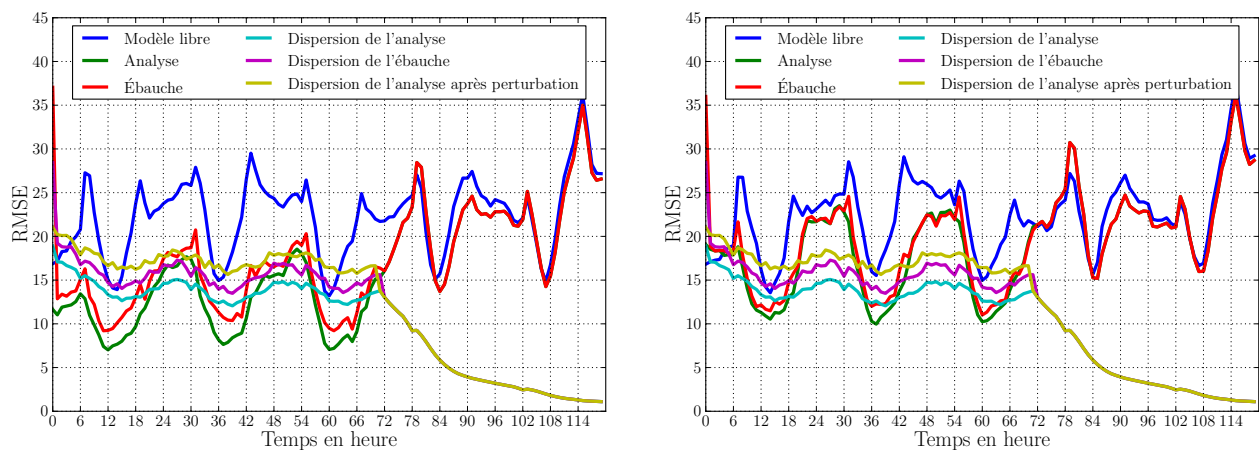


FIGURE 5.7 – Erreur moyenne et dispersion de l'ensemble en fonction du temps, pour notre expérience d'assimilation de données de référence, quand on se compare au jeu d'observations assimilées (à gauche) et au jeu de validation (à droite). L'erreur moyenne de la simulation sans analyse est marquée en bleu, l'analyse est en verte, l'ébauche, ou la prévision à 1 h, est en rouge. La dispersion de l'ensemble après analyse est en cyan, après ajout de l'inflation additive en magenta, et après propagation en jaune.

Les statistiques de RMSE et de corrélation pour cette même expérience, en se comparant au jeu d'observations assimilées et au jeu d'observations de contrôle, sont condensées dans le tableau 5.3. Elles sont calculées séparément pour chaque station puis moyennées sur l'ensemble des stations. On y voit que par rapport à la simulation de référence, notre assimilation de référence se rapproche fortement des observations, aussi bien celles assimilées que celles du jeu de contrôle. On gagne ainsi réciproquement 40% et 20% sur ces différents jeux d'observations. Les corrélations sont également fortement améliorées, gagnant réciproquement 30 et 20 points.

RMSE Corrélation	Observations assimilées		Observations de contrôle	
	Assimilation	Propagation	Assimilation	Propagation
Simulation de référence	25.92 53.6	27.95 48.9	25.84 54.8	28.19 51.8
Assimilation de référence	15.13 85.4	26.72 44.1	20.28 74.4	27.90 45.8

TABLEAU 5.3 – RMSE (en haut) et corrélation (en bas) du modèle seul et de notre assimilation de référence sur la période d'assimilation et de propagation, comparés aux observations assimilées et aux observations de contrôle. Les statistiques sont calculées pour chacune des stations d'observation et moyennées sur l'ensemble des stations.

En revanche, les performances de prévision sur les 48 h suivant l'assimilation ne bénéficient pas de l'assimilation, au contraire. Même si la RMSE semble être meilleure, d'environ 1 point sur les observations assimilées, la corrélation perd quant à elle 4 points. Si on regarde la RMSE sur des périodes plus courtes après la période d'assimilation, on peut voir que celle-ci est surtout améliorée sur les premières 24 h (tableau 5.4).

RMSE	Observations assimilées				
	0 – 6 h	0 – 12 h	0 – 24 h	24 – 48 h	0 – 48 h
Simulation de référence	23.89	25.13	26.15	28.70	27.95
Assimilation de référence	20.31	23.73	24.29	27.85	26.72

TABLEAU 5.4 – RMSE du modèle seul et de notre assimilation de référence sur les 6, 12 et 24 premières heures de la période de propagation ainsi que sur les 24 heures suivantes et la totalité de la période de propagation, comparés aux observations assimilées. Les statistiques sont calculées pour chacune des stations d'observation et moyennées sur l'ensemble des stations.

5.3 Estimation des erreurs de représentativité des stations AirBase

Nous avons diagnostiqué à la section précédente une erreur d'observation élevée pour les mesures d'AirBase. Nous allons ici estimer un biais pour chacune des stations de ce réseau d'observation. Cette démarche a déjà été entreprise avec succès par Koohkan et Bocquet (2012), qui ont inversé les flux de monoxyde de carbone en France en assimilant avec un 4D-Var les observations de CO de la Base de Données de la Qualité de l'Air (BDQA).

Nous utilisons donc un état augmenté comme présenté dans la section 4.1.2. Le vecteur d'état

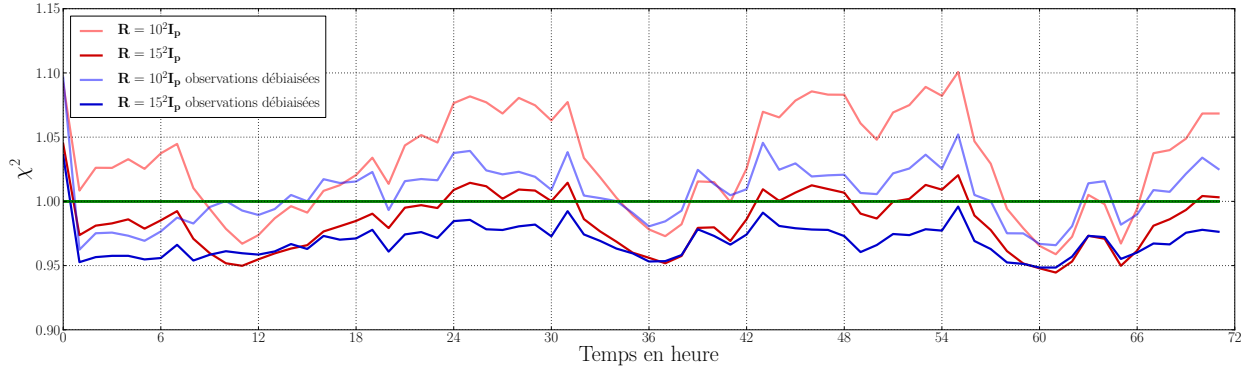


FIGURE 5.8 – Évolution du χ^2 en fonction du temps (en h), pour des écarts types de l'erreur d'observation de $10 \mu\text{g m}^{-3}$ et $15 \mu\text{g m}^{-3}$, lorsque les observations sont débiaisées ou non.

est augmenté de $\mathbf{x} \in \mathbb{R}^n$ à

$$\mathbf{z} = \begin{pmatrix} \mathbf{x} \\ \mathbf{b} \end{pmatrix} \in \mathbb{R}^{n+p}, \quad (5.2)$$

où $\mathbf{b} \in \mathbb{R}^p$ sont les biais associés à chaque station d'observation, avec $p = 198$ le nombre de stations d'observation assimilées, et où n vaut ici $N_x \times N_y \times N_z \times N_s = 67 \times 47 \times 11 \times 52 = 1801\,228$. La matrice d'observation $\mathbf{H} \in \mathbb{R}^{p \times n}$ est également modifiée en $\tilde{\mathbf{H}} \in \mathbb{R}^{p \times (n+p)}$ et est telle que

$$\tilde{\mathbf{H}} = (\mathbf{H}, \mathbf{I}_p), \text{ afin que } \tilde{\mathbf{H}}\mathbf{z} = \mathbf{H}\mathbf{x} + \mathbf{b}. \quad (5.3)$$

Les observations sont ainsi débiaisées quand les innovations sont calculées

$$\mathbf{y} - \tilde{\mathbf{H}}\mathbf{z} = (\mathbf{y} - \mathbf{b}) - \mathbf{H}\mathbf{x}. \quad (5.4)$$

Nous précisons que cette démarche peut nous imposer d'assimiler une observation négative si jamais le biais estimé est supérieur à une observation mesurée à un instant donné. Cependant, ceci n'arrive que très rarement avec notre système d'assimilation. Le modèle choisi pour la propagation des biais est le modèle de persistance. L'ensemble des biais est initialisé avec un bruit gaussien centré en 0 et de variance 10.

Avec cette estimation des biais des stations d'observation, le χ^2 est très fortement impacté. Cette remarque est en accord avec [Koohkan et Bocquet \(2012\)](#), qui ont développé un processus itératif pour optimiser conjointement la variance des erreurs d'observation et la valeur des biais. Pour des biais donnés, il détermine la variance de l'erreur d'observation satisfaisant le χ^2 , avant de mettre à jour ces biais et d'itérer ainsi jusqu'à convergence.

Les valeurs du χ^2 en fonction du temps pour des expériences avec et sans estimation des biais, avec un écart type de l'erreur d'observation de $10 \mu\text{g m}^{-3}$ et $15 \mu\text{g m}^{-3}$ sont tracées sur la Fig. 5.8. L'estimation des biais sur les stations d'observation diminue fortement le χ^2 , de telle sorte que la valeur de l'écart type des erreurs d'observation de $15 \mu\text{g m}^{-3}$ précédemment définie mène à un $\chi^2 < 1$. Si au contraire, on revient à une valeur de $10 \mu\text{g m}^{-3}$ comme initialement utilisée, le χ^2 est bien meilleur qu'avant et semble même optimal. Cela montre que lorsque l'on débiaise les observations, l'erreur d'observation est naturellement altérée, de telle sorte que l'on puisse utiliser un écart type d'environ $10 \mu\text{g m}^{-3}$ pour l'erreur d'observation.

Avec cet écart type de l'erreur d'observation, l'évolution du biais estimé pour chacune des 198 stations observées est illustrée sur la Fig. 5.9. On y voit que les biais convergent rapidement vers

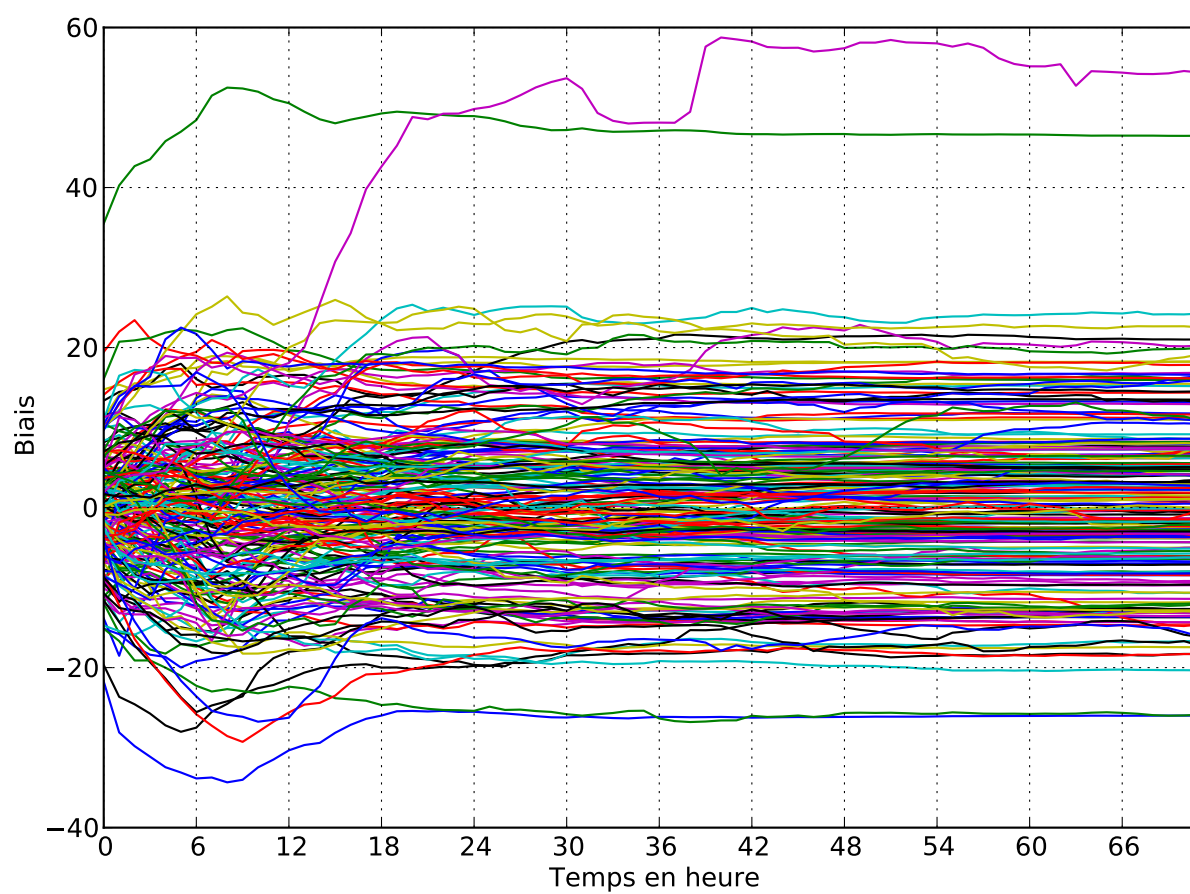


FIGURE 5.9 – Valeurs des biais (en $\mu\text{g m}^{-3}$) estimés pour chaque station, en fonction du temps.

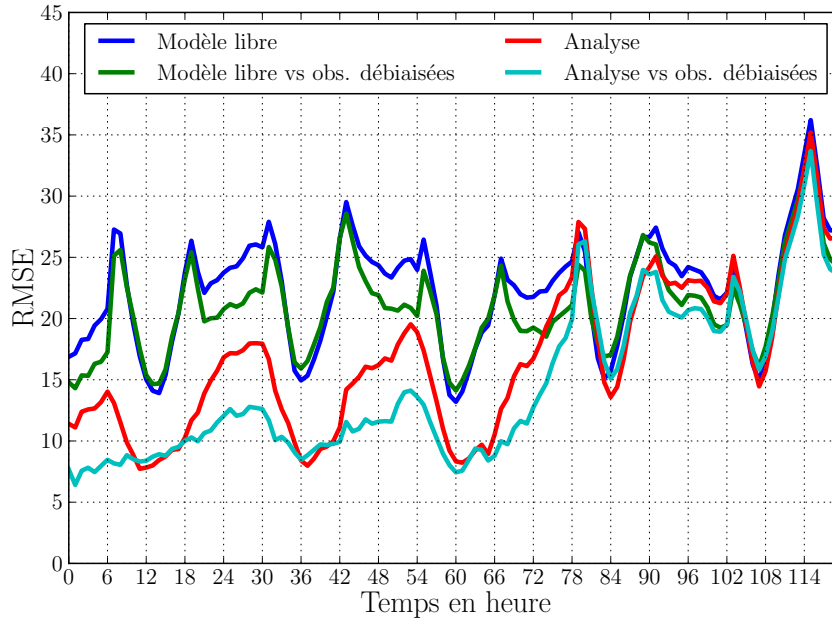


FIGURE 5.10 – Erreur moyenne en fonction du temps, pour notre expérience d’assimilation de données avec estimation des biais d’observation, quand on se compare au jeu originel d’observations assimilées et au jeu d’observations débiaisées. L’erreur moyenne de la simulation sans analyse, comparée au jeu originel, est marquée en bleu, l’analyse, comparée au jeu originel, est en rouge. L’erreur moyenne de la simulation sans analyse, comparée au jeu débiaisé, est marquée en verte, l’analyse, comparée au jeu débiaisé, est en cyan.

des valeurs très variées, allant d’environ $-20 \mu\text{g m}^{-3}$ à jusque $+20 \mu\text{g m}^{-3}$, à part deux stations qui ont des énormes biais. Ces deux stations sont celles qui enregistrent les plus fortes concentrations sur la période d’assimilation, ce qui peut expliquer ces valeurs.

La convergence rapide pourrait s’expliquer par le fait que l’on ait utilisé le modèle de persistance pour la propagation des biais. Cependant, si on ajoute un bruit stochastique (d’écart type $1 \mu\text{g m}^{-3} \text{h}^{-1}$) lors de la propagation, on observe toujours la même convergence, juste légèrement bruitée. Cela montre que la convergence n’est pas fortuite ou causée par une sous-dispersion de l’ensemble, mais bien un résultat physique.

L’erreur moyenne en fonction du temps avec estimation des biais, est tracée sur la Fig. 5.10. Dans ce contexte, deux choix s’offrent à nous. On peut comparer notre simulation de référence et notre résultat d’assimilation aux observations d’AirBase originelles, ou se comparer aux observations que l’on débiaise grâce au biais estimé lors de l’assimilation. En ce qui concerne la période de propagation, nous n’avons aucun biais sur les observations à notre disposition. Comme le biais semble avoir convergé vers une valeur fixe à l’issue de la période d’assimilation, la moyenne de la valeur du biais sur les dernières 24 h d’assimilation est utilisée pour la période de propagation.

Ces résultats montrent que nous arrivons effectivement correctement à débiaiser les observations, de sorte que notre simulation de référence tout comme notre résultat d’assimilation soient bien plus en accord avec les observations débiaisées. Ceci est vrai sur la période d’assimilation aussi bien que sur la période de propagation. Les RMSEs et corrélations pour la simulation de référence et notre assimilation, comparés aux observations assimilées débiaisées ou non, sont consignées dans le tableau (5.5).

RMSE Corrélation	Observations assimilées		Observations assimilées débiaisées	
	Assimilation	Propagation	Assimilation	Propagation
Simulation de référence	25.92 53.6	27.95 48.9	24.47 53.8	26.72 48.8
Assimilation	15.63 86.7	26.89 44.0	12.43 87.5	25.60 43.9

TABLEAU 5.5 – RMSE (en haut) et corrélation (en bas) du modèle seul et de notre assimilation de référence sur la période d'assimilation et de propagation, comparés aux observations assimilées débiaisées ou non. Les statistiques sont calculées pour chacune des stations d'observation et moyennées sur l'ensemble des stations.

En revanche, les résultats en se comparant au jeu de validation restent à ce stade inchangés. En effet, nous n'avons pas à notre disposition de valeur de biais pour ces observations de validation, qui sont de toute évidence elles-mêmes biaisées. Nous avons donc effectué une expérience où nous assimilons les observations issues de l'ensemble des 409 stations d'AirBase rurales et de fonds, ce qui nous permet de définir un biais pour les stations de validation. Si on compare alors notre simulation de référence et notre assimilation, où seulement la moitié des observations ont été assimilées, aux observations de validation débiaisées, on peut observer un net gain. Les RMSEs et corrélations sont consignées dans le tableau (5.6).

RMSE Corrélation	Observations de contrôle		Observations de contrôle débiaisées	
	Assimilation	Propagation	Assimilation	Propagation
Simulation de référence	25.84 54.8	28.19 51.8	24.91 54.8	26.80 51.8
Assimilation	20.82 74.1	27.85 46.1	19.40 74.3	26.30 46.1

TABLEAU 5.6 – RMSE (en haut) et corrélation (en bas) du modèle seul et de notre assimilation de référence sur la période d'assimilation et de propagation, comparés aux observations de contrôle, débiaisées ou non. Les statistiques sont calculées pour chacune des stations d'observation et moyennées sur l'ensemble des stations.

Des raffinements supplémentaires ont également été testés sans apporter de résultats complémentaires. On peut en effet par exemple imaginer que le biais d'une station ne soit pas fixe, mais présente un cycle journalier similaire à ce qui est visible sur la Fig. 5.2. Notamment, les concentrations pendant la journée sont nettement supérieures à celles de la nuit. On pourrait donc estimer deux biais distincts, l'un pour le jour et l'autre pour la nuit. Une autre approche consisterait à remplacer le modèle de persistance pour la propagation des biais. On pourrait imaginer appliquer un modèle cyclique dont on estimerait l'amplitude. Enfin, une dernière possibilité serait de considérer le biais comme proportionnel à la valeur du modèle interpolé à la station. On estimerait alors un biais relatif, plutôt qu'un biais absolu.

5.4 Application de l'IEEnKS sur Polair3D

Ces premières difficultés rencontrées sur l'application du LETKF avec Polair3D et des observations réelles nous poussent à prendre des précautions avant d'utiliser une méthode EnVar 4D

Paramètre	Simulation de référence	Vérité de l'OSSE
Résolution	0.5°	0.5°
Domaine	9°O – 24°E ($N_x = 67$) 36°N – 59°N ($N_y = 47$)	9°O – 24°E ($N_x = 67$) 36°N – 59°N ($N_y = 47$)
Niveaux verticaux	$N_z = 11$	$N_z = 11$
Chimie	CB05 ($N_s = 52$)	CB05 ($N_s = 52$)
Aérosol	non	non
Météorologie	CEPMMT	CEPMMT
Conditions initiales	MOZART 2.0	MOZART 2.0
Émissions biogéniques	MEGAN	MEGAN
Diffusion verticale	Louis (1979)	Louis (1979)
Diffusion horizontale	100 000 m ² s ⁻¹	100 000 m ² s ⁻¹
Limiteur de flux	Koren-Sweby	Koren-Sweby
Conditions aux limites	MOZART 2.0 * ⁶⁸ / ₇₅	MOZART 2.0 *1
Émissions anthropiques	EMEP, avec NO _x *1.2	EMEP *1
Dépôt	Zhang <i>et al.</i> (2003)*1.7	Zhang <i>et al.</i> (2003)*1

TABLEAU 5.7 – Compilation des paramètres et données utilisés dans notre simulation de référence et dans la simulation considérée comme la vérité. Les différences sont mises en rouge et à la fin du tableau.

comme l'IEnKS. Nous allons donc tout d'abord l'appliquer à Polair3D dans un contexte maîtrisé. Nous nous plaçons dans un cadre d'*observing system simulation experiment* (OSSE), pour parfaitement maîtriser l'erreur modèle et ne pas avoir d'observations biaisées comme celles d'AirBase. Nous verrons ensuite seulement les résultats d'application de l'IEnKS avec des observations réelles.

5.4.1 Expériences synthétiques

Nous nous plaçons ici dans un contexte d'OSSE. Nous allons donc engendrer une vérité à partir d'une simulation, créer des observations en bruitant cette simulation et utiliser ces observations dans une expérience d'assimilation de données. Cependant, le modèle étant stable, si la même simulation était utilisée pour la vérité et pour notre expérience d'assimilation de données, on aurait une convergence rapide de notre état estimé vers la vérité, que l'on assimile des observations ou non. Il est donc nécessaire d'avoir deux modèles différents, un pour engendrer la vérité et les observations, et un autre utilisé dans notre système d'assimilation de données. Il s'agit alors non plus d'une expérience jumelle, mais d'une expérience dite fraterne. Cette situation est plus réaliste, car elle inclut de l'erreur modèle.

Pour simuler la vérité, nous utilisons donc toujours Polair3D avec des paramètres similaires à notre situation de référence, mais avec des conditions aux limites, des émissions de NO_x et un dépôt d'ozone modifié. Nous gardons en revanche la simulation de référence inchangée pour notre système d'assimilation de données. Nous recopions dans le tableau 5.7 les paramètres de la simulation de référence et de la simulation utilisée pour engendrer la vérité. La RMSE entre ces deux simulations sur les 120 h d'expérience est de 13.66 µg m⁻³, ce qui est du même ordre de grandeur que l'erreur modèle que l'on a diagnostiquée dans notre assimilation de référence. Aussi, gardons-nous l'inflation additive inchangée.

Pour créer les observations, nous interpolons la vérité au niveau des stations d'AirBase et bruitons la valeur avec un bruit gaussien d'écart type 15 µg m⁻³. Notre matrice de covariance d'er-

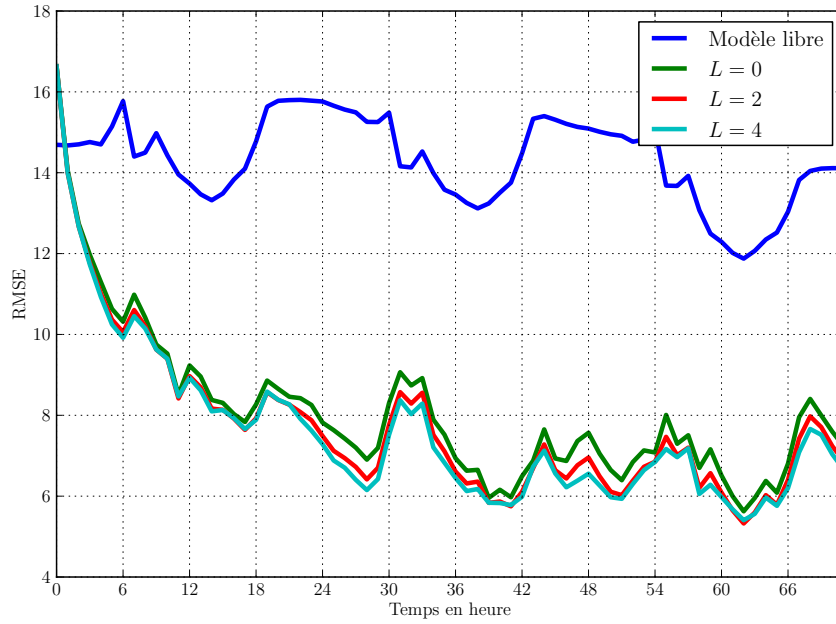


FIGURE 5.11 – RMSEs de filtrage en fonction du temps pour le modèle sans assimilation (en bleu) et les résultats d’assimilation avec différentes longueurs de fenêtre d’assimilation de données.

reur d’observation est donc $\mathbf{R} = 15^2 (\mu\text{g m}^{-3})^2$, comme pour notre assimilation de référence. Si une observation ainsi engendrée devait être négative, on considère que la station n’a pas produit d’observation à cette heure là. Ceci arrive fréquemment dans les données d’AirBase, qui peuvent présenter ainsi des trous de plusieurs heures, voire même plusieurs jours.

Dans ces expériences, nous n’avons pas advecté dans l’espace les observations avec un modèle de substitution comme présenté dans la section 2.5.3.3. Nous avons donc utilisé ce que nous avons appelé la localisation par domaines statiques. En effet, dans ce contexte idéalisé, et avec des petites longueurs de fenêtres d’assimilation, l’utilisation des domaines covariants semble détériorer les résultats. Il se peut que cette détérioration soit causée par l’erreur modèle. Une étude plus approfondie de l’impact de la localisation par domaines covariants devrait néanmoins être menée.

Nous avons donc fait varier la longueur de la fenêtre d’assimilation de données L avec l’IEEnKS local, avec $S = 1$. En théorie, on s’attend à ce que l’augmentation de la fenêtre d’assimilation de données permette de mieux prendre en compte la non-linéarité du modèle. Dans une certaine mesure, les RMSEs de filtrage et de lissage devraient donc s’améliorer quand L augmente. C’est ce que nous remarquons dans un premier temps sur les RMSEs de filtrage en fonction du temps, montrées sur la Fig. 5.11 pour la période d’assimilation. Cependant, le gain reste très modéré, permettant de passer, sur ces 72 h d’assimilation, d’une RMSE de $8.08 \mu\text{g m}^{-3}$ quand $L = 0$ à $7.61 \mu\text{g m}^{-3}$ quand $L = 4$. Nous précisons qu’il s’agit ici de RMSEs calculées entre la vérité et l’assimilation à tous les points de grille, et non d’une erreur moyenne aux stations.

Ces résultats encourageants ne se vérifient en revanche pas sur la RMSE de lissage. En effet, cette dernière, tracée sur la Fig. 5.12, se dégrade très nettement quand la fenêtre d’assimilation est augmentée. On remarque par ailleurs que les courbes semblent suivre la même évolution temporelle, amplifiée et décalée dans le temps de L heures.

Ces conclusions sont en accord avec celles obtenues sur le L95-T faiblement couplé (section 4.3.3), pour lequel le filtrage était constant quand L était varié et le lissage se dégradait rapidement. De

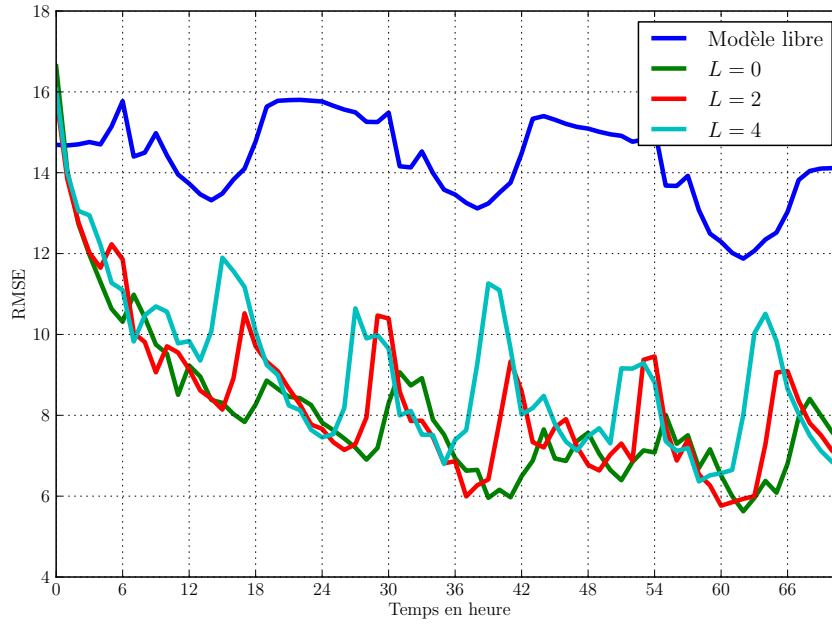


FIGURE 5.12 – RMSEs de lissage en fonction du temps pour le modèle sans assimilation (en bleu) et les résultats d’assimilation avec différentes longueurs de fenêtre d’assimilation de données. La courbe pour $L = 0$ est la même que celle de la figure précédente.

plus, le comportement de la RMSE de lissage n’est pas anodin. Nous l’expliquons par la forte stabilité dynamique du modèle Polair3D. Le schéma Fig. 5.13 illustre le phénomène à l’œuvre. On y voit l’évolution d’une variable de la vérité en rouge. La trajectoire du modèle erroné utilisé par le système d’assimilation de données est en bleu. Si on part de deux conditions initiales différentes, la stabilité de ce modèle fait converger les deux trajectoires issues de ces conditions initiales. On a une observation, marquée en vert, au temps t_L .

L’assimilation de données va déterminer l’incrément à imposer à l’état du modèle pour être en accord avec cette observation. Dans le cas où $L = 0$, on a un petit incrément qui nous décale sur l’observation directement. En revanche, quand L augmente, l’incrément à appliquer au modèle dans le passé grandit, du fait de la stabilité du modèle. Ainsi, l’incrément que l’on doit appliquer à t_0 pour que le modèle soit sur l’observation à t_L est grand et tel que si on devait calculer la RMSE de lissage, cette dernière serait mauvaise.

Nous voyons donc ici les limites de l’application aux modèles de chimie atmosphérique d’une méthode EnVar 4D, dont l’IEnKS est un archétype. En effet, un point qui fait la force de cette méthode sur les modèles jouets fondés sur le Lorenz-95 (L95), qui est chaotique, est qu’elle place l’analyse sur l’espace instable du modèle (Bocquet et Carrassi, 2017). L’espace instable étant inexistant ici, l’intérêt est alors effectivement limité.

Pour limiter les effets de la stabilité que l’on a mis en avant ici, nous pouvons assimiler plusieurs observations dans la fenêtre. Notamment, on peut fixer $S = L + 1$ et ainsi assimiler toutes les observations contenues dans la fenêtre, à la manière d’un 4D-LETKF (Hunt *et al.*, 2007), mais en gardant plusieurs itérations de minimisation de la fonction de coût. En comparaison, nous avons jusqu’ici $S = 1$ quel qu’était L . Pour passer au cycle suivant, on propage d’autant de pas de temps que d’observations ainsi assimilées, à savoir S . On peut alors espérer que les observations au début de la fenêtre vont contraindre l’état estimé pour empêcher d’imposer à t_0 un incrément comme celui

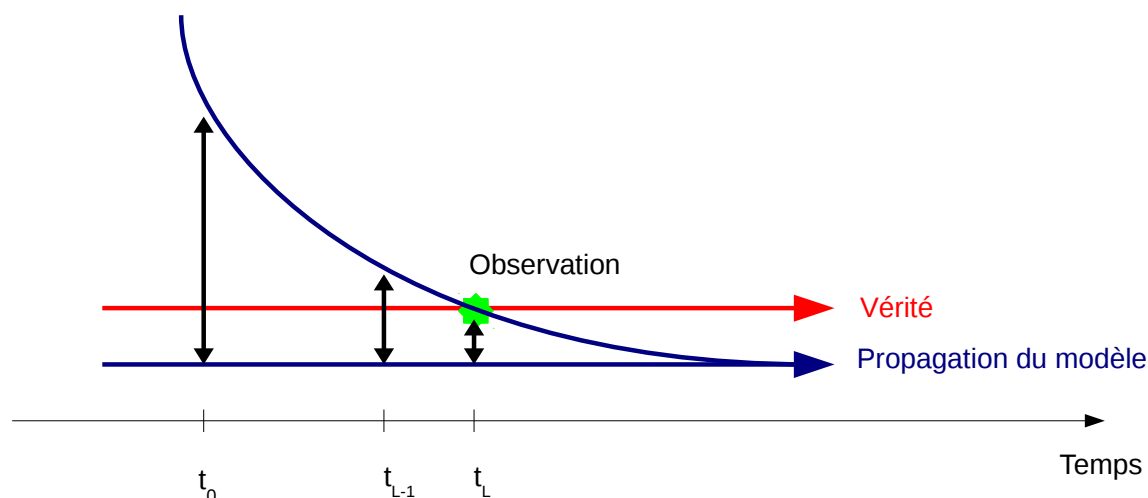


FIGURE 5.13 – Représentation schématique de l'évolution d'une variable de la vérité en rouge et du modèle erroné stable en bleu. Pour ce dernier, si on part d'une condition initiale différente, on converge avec la stabilité vers l'attracteur du modèle. L'observation est en vert. Les incréments à appliquer au modèle pour atteindre l'observation au temps t_L sont les flèches noires.

représenté sur la Fig. 5.13.

Les résultats d'assimilation avec cette configuration, pour plusieurs valeurs de S (et de L par la même occasion), sont montrés sur la Fig. 5.14. On ne peut présenter que des résultats en lissage dans ce contexte, puisqu'à part le bout de la fenêtre, toutes les estimations utilisent des observations futures.

On voit qu'on améliore ainsi très fortement la qualité de l'estimation en lissage par rapport au cas où $S = 1$. Cependant, l'estimation en lissage n'apporte quasiment aucun gain par rapport au filtrage précédent. Quand la RMSE de filtrage était $7.61 \mu\text{g m}^{-3}$ pour $L = 4$, elle est désormais de 7.51, ce qui est trop proche pour conclure sur une différence statistiquement valide. En revanche, le nombre de propagations requises dans ce contexte est nettement inférieur. En effet, étant donné qu'on assimile S observations par cycle, on effectue S fois moins de cycles, qui chacun requiert $m \times L$ propagations par itération de minimisation, avec m la taille de l'ensemble.

5.4.2 Expériences avec observations réelles

Dans cette section, les paramètres seront identiques à ceux résumés dans le tableau 5.2. Nous présentons dans un premier temps les résultats de l'IEnKS lorsque nous assimilons de nouveau les observations d'AirBase notées rurales et de fond. Il s'agit donc simplement d'étendre les résultats de la section 5.2.3 au cas avec $L > 0$ et $S \geq 1$.

Dans cette expérience, nous faisons varier L avec $S = 1$ et vérifions la qualité de l'analyse en filtrage et en lissage, comparée aux observations analysées ou de contrôle. Cela représente une première différence avec le cas synthétique où l'on pouvait se comparer directement à la vérité. De plus, contrairement au cas synthétique, nous avons appliqué ici de la localisation par domaines covariants, même si les résultats sont quasiment identiques à ceux obtenus avec de la localisation par domaines statiques. La modification de la position des stations d'observation par le transport avec un modèle de substitution peut s'apparenter à une correction du second ordre quand on connaît

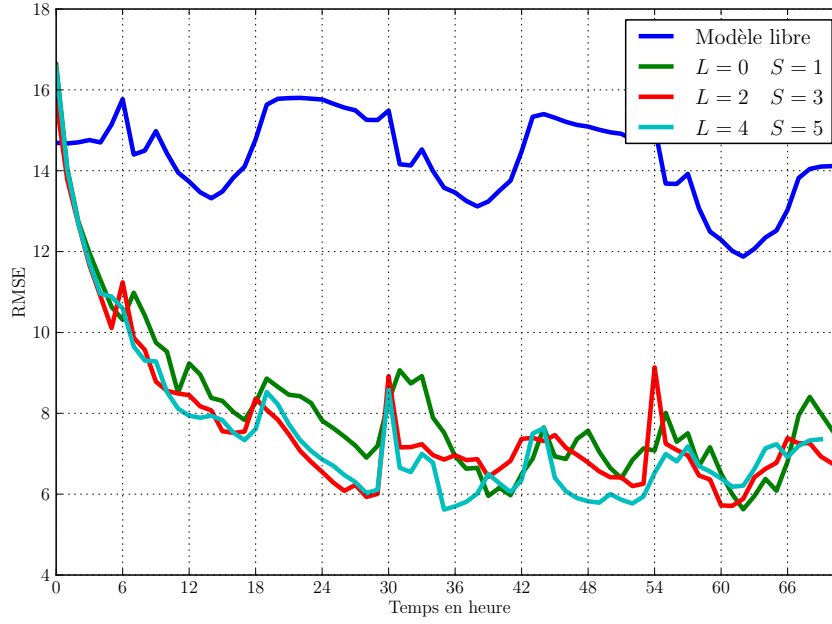


FIGURE 5.14 – RMSEs de lissage en fonction du temps pour le modèle sans assimilation (en bleu) et les résultats d’assimilation avec un IEnKS avec $S = L + 1$, pour différentes longueurs de fenêtre d’assimilation de données. La courbe pour $L = 0$, $S = 1$ est la même que celle des figures précédentes.

les biais possibles qu’ont les observations. De plus, nous sommes dans une situation météorologique anticyclonique, donc avec assez peu de transport par le vent, ce qui peut expliquer le peu d’impact. On pourrait espérer que les résultats soient similaires au cas synthétique. Ce n’est malheureusement pas le cas.

Les courbes mentionnées précédemment sont rassemblées sur la Fig. 5.15. Les conclusions que l’on peut tirer de ces courbes sont les mêmes quel que soit le jeu d’observations auquel on se compare. Tout d’abord, similairement au cas avec des observations synthétiques, la RMSE de lissage est dégradée quand la longueur de la fenêtre d’assimilation augmente. Cependant, l’évolution des courbes quand L varie n’est pas aussi caractéristique que ce qui était observé dans le cas synthétique. De plus, contrairement au cas avec observations synthétiques, la RMSE de filtrage se dégrade également à mesure que L augmente. Ce résultat est surprenant, puisqu’on attendait de l’utilisation d’une méthode EnVar 4D itérative qu’elle puisse améliorer les prévisions en prenant en compte la non-linéarité du modèle grâce à l’analyse variationnelle au sein de la fenêtre. Enfin, les performances en prédiction ne semblent presque pas impactées par le recours à l’IEnKS.

Ces premiers résultats semblent décourageants, surtout étant données les réussites, même modérées, qui ont été obtenues avec des observations synthétiques. De toute évidence, la forte baisse de performance en lissage quand L grandit peut de nouveau être attribuée à la stabilité du modèle. Cela nous pousse à vouloir vérifier si le cas $S = L + 1$ donne de meilleurs résultats. Ce n’est de nouveau pas le cas, comme le montre la Fig. 5.16. Alors que l’assimilation simultanée de plusieurs observations permettait, dans le cas synthétique, de contraindre le lissage, cela n’est plus le cas avec des observations réelles.

On peut potentiellement expliquer ces résultats de plusieurs manières. Tout d’abord, la qualité des observations peut être mise en cause, comme nous l’avons déjà montré dans la section 5.3. Aussi

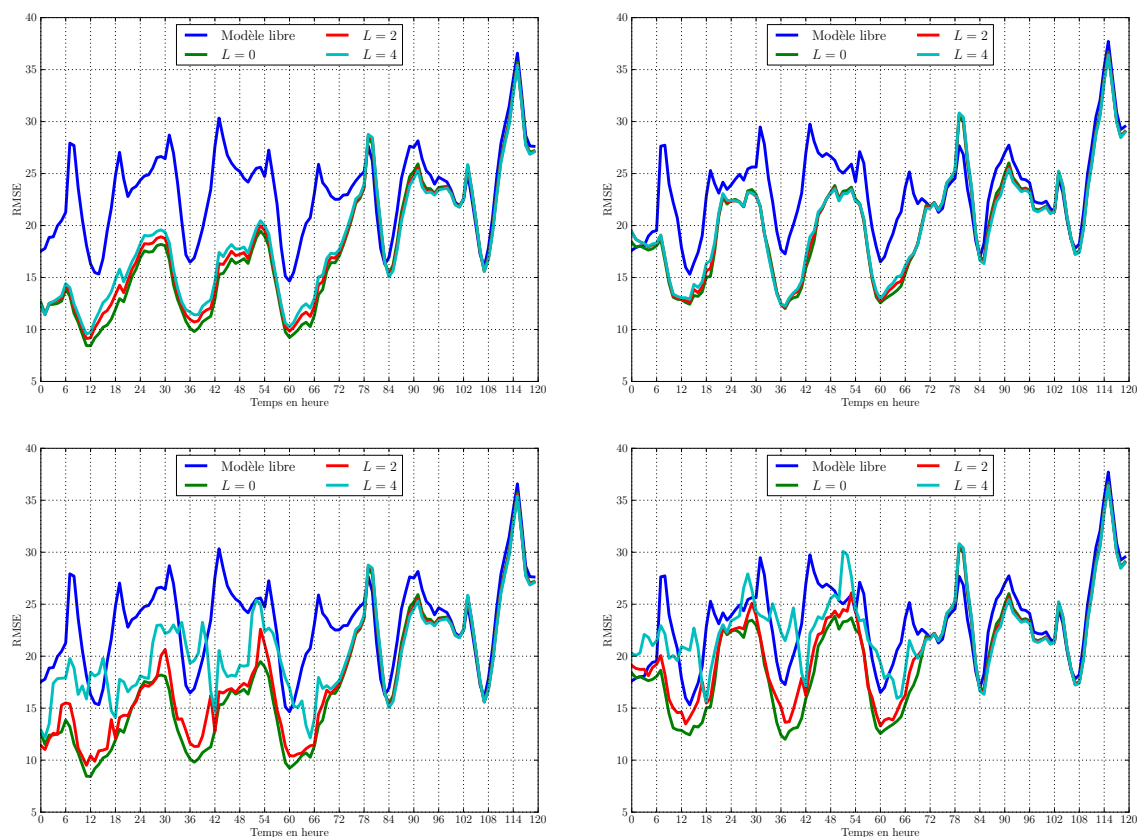


FIGURE 5.15 – Erreurs moyennes de filtrage (en haut) et de lissage (en bas), en fonction du temps, pour le modèle sans assimilation (en bleu) et les résultats d'assimilation avec différentes longueurs de fenêtre d'assimilation de données. Ces erreurs sont calculées en se comparant au jeu d'observations analysées (à gauche) ou aux observations de contrôle (à droite).

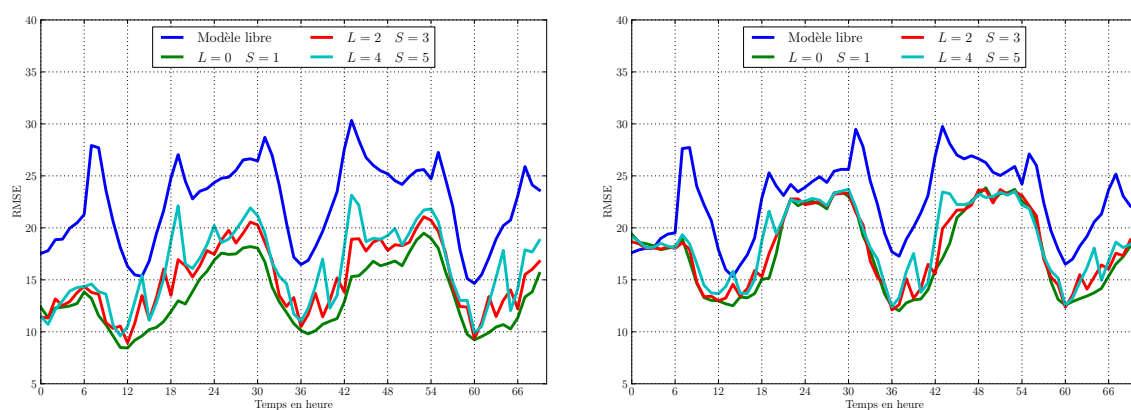


FIGURE 5.16 – Erreurs moyennes de lissage en fonction du temps, pour le modèle sans assimilation (en bleu) et les résultats d'assimilation avec différentes longueurs de fenêtre d'assimilation de données, pour l'IEnKS avec $S = L + 1$. Ces erreurs moyennes sont calculées en se comparant au jeu d'observations analysées (à gauche) ou aux observations de contrôle (à droite).

semble-t-il nécessaire de remédier à la faible précision des observations en assimilant les observations débiaisées de la valeur diagnostiquée à cette même section. Nous présentons ces résultats dans la section suivante. Ensuite, contrairement au cas synthétique, nous avons une très forte erreur modèle dont nous ne connaissons pas la nature. Nous avons défini l'inflation additive ajoutée entre chaque cycle de sorte à ce qu'elle représente une erreur climatologique de notre modèle et son amplitude et sa corrélation spatiale a été ajustée, mais elle reste définie de manière empirique. À défaut d'une meilleure connaissance, nous pouvons utiliser cette définition de l'erreur modèle avec un IEnKF-Q. La section 5.4.4 présente des expériences dans ce contexte.

5.4.3 Assimilation des observations débiaisées

Dans la section 5.3, un biais sur chacune des 409 stations d'observation a été diagnostiqué. Nous pouvons donc soustraire ce biais des observations et assimiler les observations ainsi débiaisées. Nous avons vu dans ce contexte qu'une matrice d'erreur d'observation réduite à $\mathbf{R} = 10^2 \mathbf{I}_p (\mu\text{g m}^{-3})^2$ était suffisante. Les paramètres de l'assimilation de données présentée ici seront donc identiques au cas de référence, à l'exception de la variance de l'erreur d'observation. Nous avons utilisé la localisation par domaines statiques, qui donne des résultats très similaires à de la localisation par domaines covariants.

Les résultats en filtrage et en lissage, comparés aux observations analysées ou de contrôle, sont présentés sur la Fig. 5.17. Les conclusions obtenues avec les observations d'AirBase non modifiées tiennent toujours. Le fait d'utiliser dans cette expérience des observations débiaisées ne permet pas à l'IEEnKS de s'améliorer relativement au LETKF.

Même si l'erreur modèle est partiellement prise en compte par le débiaisage des observations, qui réduit l'erreur de représentativité, il semblerait qu'elle soit trop importante pour qu'une méthode s'étendant sur une fenêtre de temps puisse présenter un quelconque intérêt. En effet, une des hypothèses de la majorité des méthodes d'assimilation de données, dont les méthodes EnVar 4D et à ce titre l'IEEnKS, est d'avoir un modèle parfait, ce qui n'est de toute évidence pas le cas ici. La très forte inflation additive, dont l'écart-type correspond à environ 15% de la concentration moyenne sur le domaine, en est la preuve. Nous allons donc appliquer l'IEEnKF-Q pour pallier cela.

5.4.4 Assimilation avec l'IEEnKF-Q

Jusqu'ici, outre le débiaisage des observations, l'erreur modèle était prise en compte par l'implémentation d'une inflation additive sur une partie des champs d'ozone, de variance $10^2 (\mu\text{g m}^{-3})^2$, corrélée en espace. Le schéma correspondait alors à un EnKF-Rand ou IEnKS-Rand, présentés à la section 4.2. Les champs jusqu'à 1000 m sont perturbés, mais notons que si seuls les champs d'ozone au sol étaient perturbés, les résultats obtenus seraient similaires. Nous avons donc à notre disposition une matrice $\mathbf{Q}$ que l'on peut considérer comme une approximation de la matrice de l'erreur modèle additive. Nous pouvons ainsi utiliser l'algorithme 9 de l'IEEnKF-Q local avec cette matrice $\mathbf{Q}$.

Étant données les dimensions de notre domaine ($N_x = 67$; $N_y = 47$), notre matrice d'erreur est de taille 3149^2 . On devrait alors choisir $m_q = 3150$ pour l'IEEnKF-Q afin d'avoir $\mathbf{X}^q$ une racine exacte. Une telle dimension risque d'être prohibitive pour l'application de notre algorithme. Cependant, dans la mesure où l'analyse est locale, on peut se contenter de définir un $\mathbf{Q}_m$ différent pour chacun des M points d'analyse, de taille réduite et portant sur les incertitudes dans le sous-domaine de l'analyse. En pratique, nous effectuons une analyse par bloc de taille 5. Nous pouvons alors ne sélectionner que les lignes et les colonnes de la matrice $\mathbf{Q}$ correspondant aux points de grille dans le

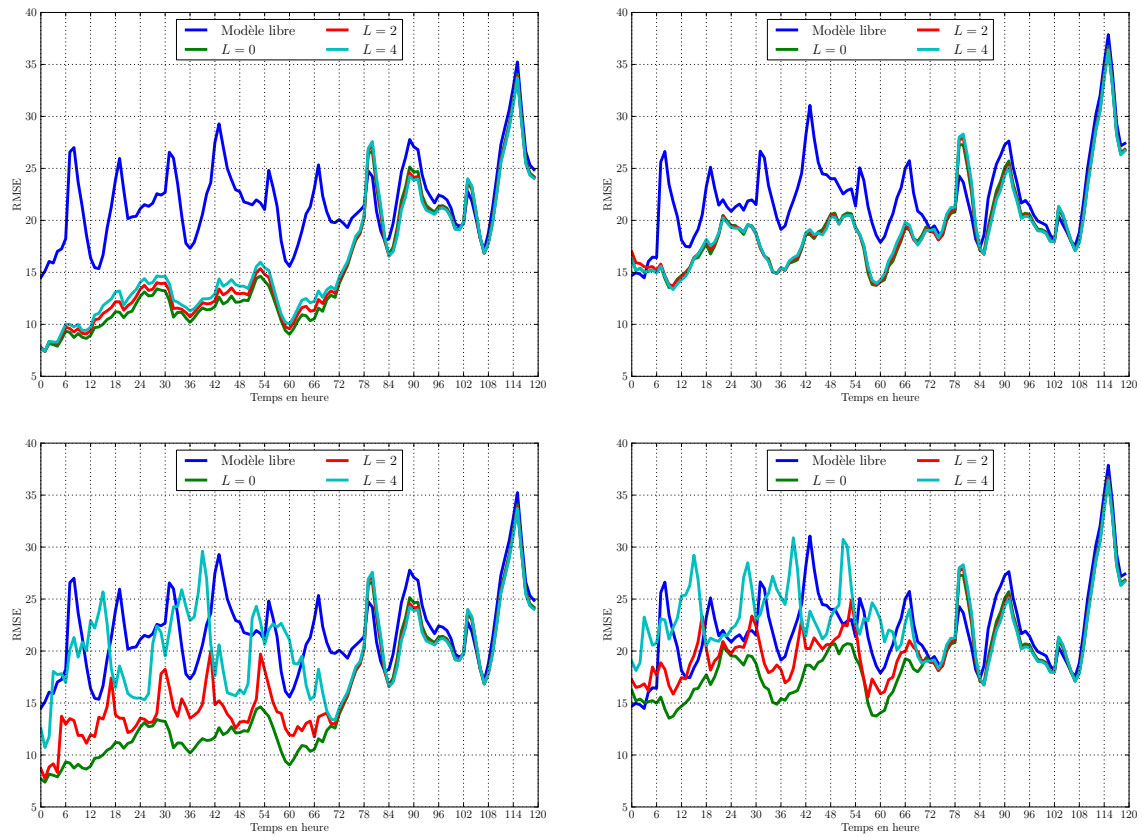


FIGURE 5.17 – Erreurs moyennes de filtrage (en haut) et de lissage (en bas), en fonction du temps, pour le modèle sans assimilation (en bleu) et les résultats d'assimilation avec différentes longueurs de fenêtre d'assimilation de données. Les observations analysées et de contrôle ont été débiaisées a priori. Ces erreurs moyennes sont calculées en se comparant au jeu d'observations analysées (à gauche) ou aux observations de contrôle (à droite).

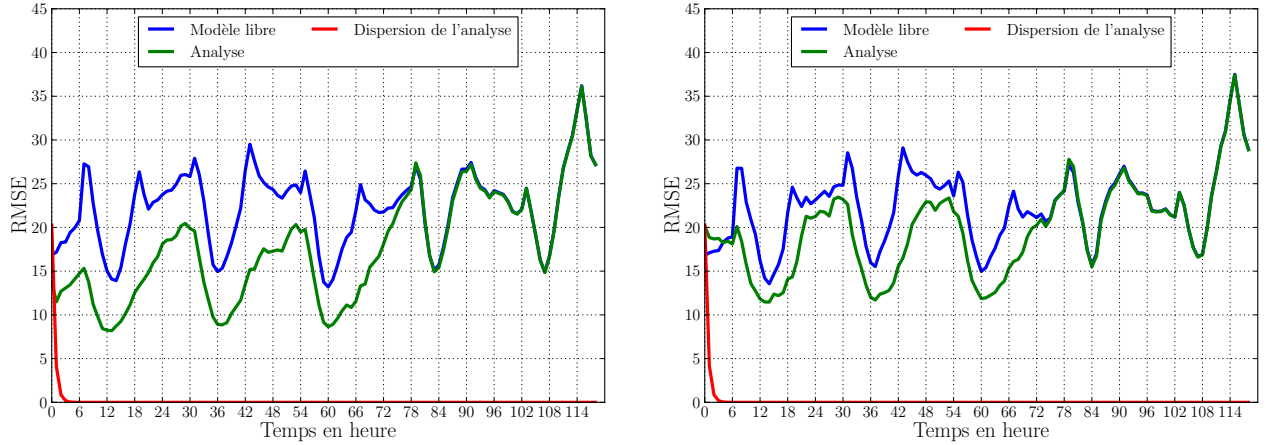


FIGURE 5.18 – Erreur moyenne et dispersion de l'ensemble en fonction du temps, pour une expérience d'assimilation de données avec l'IEnKF-Q sans inflation additive, quand on se compare au jeu d'observations assimilées (à gauche) et au jeu d'observations de contrôle (à droite). L'erreur moyenne de la simulation sans analyse est marquée en bleu, le filtrage est en vert et est confondu avec le lissage. La dispersion de l'ensemble après analyse est en rouge.

bloc, quitte à rajouter une petite marge supplémentaire. Dans notre cas, nous avons donc conservé les points de grille dans le bloc, plus une bande autour du bloc, soit un totale de 7×7 valeurs. On peut donc se contenter de $m_q = 50$ et on définit pour chaque point d'assimilation une matrice $\mathbf{X}_m^q$ comme la racine carrée, de taille $(N_y N_x \times m_q)$, de $\mathbf{Q}_m$.

Pour l'étape $\text{SR}(\mathbf{X}, m_e)$ de la ligne 25 de l'algorithme, on souhaite éviter de faire une décomposition en valeurs singulières par point de grille, même pour des matrices avec uniquement $m_e + m_q$ colonnes. De plus, dans un contexte local, il n'est pas possible d'effectuer des rotations aléatoires de l'ensemble. En effet, les algorithmes locaux fonctionnent sur le principe qu'entre deux points d'analyse proches, les ensembles analysés sont similaires et peuvent être recollés continûment. Cette propriété est perdue si une rotation aléatoire de l'ensemble est appliquée en chaque point de l'analyse. Nous nous contentons donc de ne conserver que les m_e premiers membres des matrices $\mathbf{X}^m$. Dans les expériences synthétiques sur le L95, cette démarche donnait des résultats similaires à ceux montrés dans la section 4.2, à part une légère sous-optimalité quand $\mathbf{Q}$ était grand.

Nous avons donc appliqué l'IEnKF-Q local sur notre CTM avec les observations d'AirBase. En théorie, si la matrice des erreurs modèle est parfaitement définie, on pourrait s'attendre à ce que nous n'ayons plus besoin d'inflation. Par conséquent, nous avons dans un premier temps enlevé toute forme d'inflation additive. Les erreurs moyennes de filtrage et de lissage en fonction du temps, ainsi que la dispersion de l'ensemble, comparées aux observations analysées et de contrôle, sont tracées sur la Fig. 5.18. On remarque que l'ensemble s'effondre totalement, avec une dispersion qui tombe très rapidement à 0, bien plus rapidement que si une simple propagation de l'ensemble avait été effectuée. Cela n'empêche pas pour autant les résultats de filtrage d'être sensiblement les mêmes que ceux obtenus lors de l'application d'un LETKF.

Ces résultats ne sont que modérément étonnants. Pour comprendre, reportons nous à ceux obtenus lors de l'application de l'IEnKF-Q sur le L95. Tout d'abord, rappelons que de l'inflation multiplicative avait été implémentée et était nécessaire pour obtenir les scores optimaux. Cependant, dans ces expériences, la méthode $\text{SR}(\mathbf{X}, m_e)$ appliquait une décomposition en valeurs singulières à $\mathbf{X}_2^a$, tandis qu'ici nous nous contentons de ne conserver que les m_e premiers membres de l'ensemble.

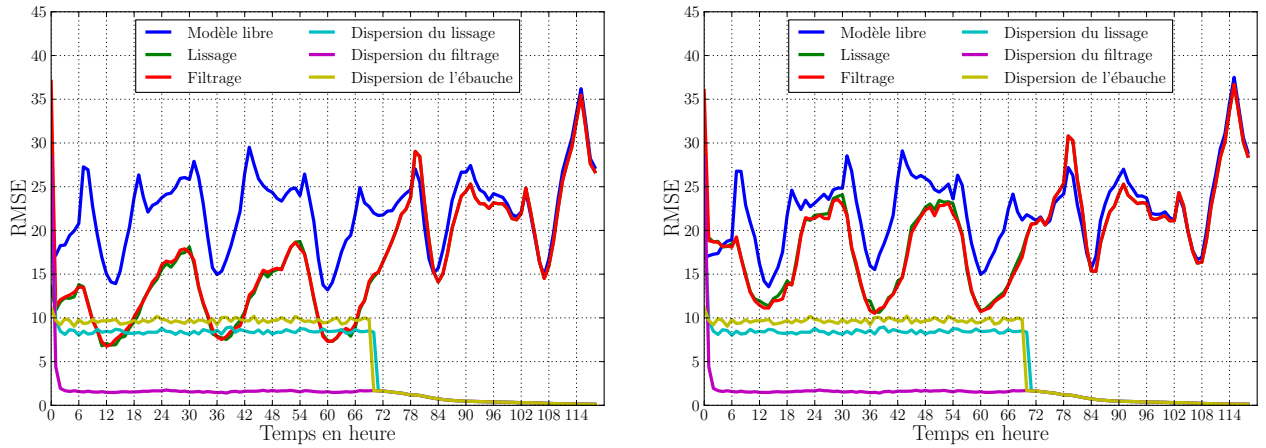


FIGURE 5.19 – Erreur moyenne et dispersion de l'ensemble en fonction du temps, pour une expérience d'assimilation de données avec l'IEEnKF-Q et inflation additive, quand on se compare au jeu d'observations assimilées (à gauche) et au jeu d'observations de validation (à droite). L'erreur moyenne de la simulation sans analyse est marquée en bleu, le filtrage est en rouge, le lissage en vert. La dispersion de l'ensemble en lissage est en cyan, en filtrage en magenta et après ajout de l'inflation additive en jaune.

Or, sur le L95, pour de très grandes valeurs de $\mathbf{Q}$ et lorsque la méthode $\text{SR}(\mathbf{X}, m_e)$ se contentait de conserver les m_e premiers membres de l'ensemble, l'absence d'inflation correspondait au cas optimal (résultats non montrés). Dans ce cas, la dispersion de l'ensemble tombait également à 0. L'algorithme ayant pour autant toujours recours aux anomalies d'erreur modèle $\mathbf{X}^q$, qui ne peuvent pas être sous-dispersées, les scores étaient quasiment identiques à ceux montrés sur les Figs. 4.7 et 4.8. Ainsi, on voit comment, sur le L95, il était possible d'avoir un ensemble sous-dispersé avec des scores de performance raisonnables.

Ce résultat est en accord avec le fait que la taille de l'ensemble importe peu quand $\mathbf{Q}$ est très grand. En effet, on voit sur ces mêmes courbes Figs. 4.7 et 4.8 que le passage de $m_e = 41$ à $m_e = 20$ n'impacte que peu les scores, et il s'avère que l'on peut aisément descendre en dessous de la taille de l'espace instable sans provoquer la divergence du filtre. Tous ces résultats expérimentaux obtenus sur le L95 avec l'IEEnKF-Q permettent de comprendre les résultats obtenus dans cette expérience avec Polair3D, à savoir d'une part l'effondrement de l'ensemble et d'autre part la qualité relative de l'analyse. On précise par ailleurs que l'hypothèse de dire que $\mathbf{Q}$ est grand est vérifiée, étant donnée l'amplitude relative du $\mathbf{Q}$ choisie par rapport aux concentrations moyennes d'une part (de l'ordre de 15%), et à l'évolution temporelle des concentrations d'autre part, qui ne varient guère de plus de $30 \mu\text{g m}^{-3}$ en une heure.

Nous avons donc voulu vérifier le comportement de l'algorithme si une inflation additive était conservée. Afin d'obtenir une dispersion de l'ordre de l'erreur d'analyse, la totalité de l'inflation additive utilisée jusqu'ici est appliquée. Les résultats sont affichés sur la Fig. 5.19. On remarque tout d'abord que l'analyse contracte énormément l'ensemble en filtrage, mais assez peu en lissage. Cependant, cette dispersion accrue de l'ensemble ne semble pas impacter les résultats qui sont très similaires à ceux obtenus sans inflation additive.

On remarque donc que l'IEEnKF-Q ne permet pas dans ce contexte d'améliorer les scores d'estimation des concentrations de l'ozone par rapport à un LETKF. Plusieurs hypothèses peuvent être avancées pour expliquer ce résultat. Tout d'abord, l'erreur modèle utilisée peut être mal dé-

finie dans notre expérience. Étant donné que seul l'espace des anomalies de modèle semble être utilisé pour trouver l'analyse, celles-ci se doivent d'être mieux définies pour améliorer l'analyse. La deuxième hypothèse, qui y est reliée, est que les erreurs modèle auxquelles nous sommes confrontés, sont trop importantes et potentiellement non-additives. Nos conditions aux limites, nos champs d'émissions, notre modélisation du dépôt, sont autant de paramètres que nous avons déjà modifiés empiriquement pour améliorer la concordance de notre modèle avec les observations. Les champs météorologiques sont une source importante d'erreur modèle supplémentaire que nous n'avons pas changée. De nombreuses études sont effectuées en utilisant pour chaque membre de l'ensemble une configuration différente (Constantinescu *et al.*, 2007a; Wu *et al.*, 2008; Boynard *et al.*, 2011). Nous avons déjà vu dans la section 4.3.3 que cette méthode conduisait aux meilleurs résultats sur le L95-T dans un contexte *offline* idéalisé. Elle permet une configuration adaptée de l'erreur modèle et nous aurions pu y recourir. À défaut, une méthode d'estimation de l'erreur modèle comme celles de Piccolo et Cullen (2016); Emili *et al.* (2016) pourrait permettre d'améliorer les performances en prévision et empêcher temporairement les membres de l'ensemble de converger vers l'attracteur stable du modèle.

5.4.5 Conclusions de l'application de l'IEEnKS sur un CTM

Nous avons présenté dans ce chapitre l'ensemble des expériences exploratoires qui ont été effectuées concernant l'application d'une méthode EnVar 4D, à savoir l'IEEnKS, sur le CTM Polair3D. Les résultats préliminaires issus de l'application de l'IEEnKS sur le L95-T avaient montré qu'un algorithme variationnel s'étendant sur une longue fenêtre risquait de ne pas s'avérer utile dans un contexte *offline* comme celui d'un CTM. Cependant, pour des modèles non-linéaires, la possibilité d'avoir une analyse variationnelle sur une longue fenêtre d'assimilation de données présente un intérêt. Les expériences avec des observations synthétiques ont laissé entrevoir cet avantage, avec des RMSEs de filtrage qui s'amélioraient quand la fenêtre était augmentée. Malheureusement, alors que le lissage sur le L95-T s'améliorait sur de courtes fenêtres avant de se détériorer, les résultats en lissage avec Polair3D se dégradent dans tous les cas de figure quand L grandit. Nous attribuons ce comportement à la très forte stabilité du modèle combinée à la présence d'erreur modèle.

Malgré ces premiers résultats encourageants, le passage à des observations réelles, même débiaisées, n'a pas permis de retrouver un tel comportement. L'utilisation de l'IEEnKS, en filtrage comme en lissage, dégrade la prévision de l'état, aussi bien par rapport à des observations de contrôle qu'en prévision. La très forte erreur modèle associée à notre CTM explique probablement ces résultats. Cependant, le recours à l'IEEnKF-Q, censé prendre en compte cette erreur modèle, n'a pas suffi pour améliorer les prévisions. Cela est probablement dû au fait qu'elle est mal définie via notre simple matrice d'inflation corrélée en espace d'une part et qu'elle n'est pas uniquement additive d'autre part.

Ces résultats, même si pessimistes, ne sont pas définitifs et n'excluent pas la possibilité d'améliorer notre estimation des champs d'ozone grâce à une méthode 4D Envar. Le point critique de cette expérience semble être la forte erreur modèle associée à notre CTM. Étant donné ce constat, plusieurs voies d'amélioration s'offrent à nous. Nous pouvons tout d'abord chercher à améliorer notre simulation de base. Cela permettrait de réduire certaines sources d'erreur pour se retrouver dans une situation où les méthodes supposant un modèle parfait sont applicables. Par exemple, nous pourrions augmenter la résolution de notre domaine, ce qui permettrait de réduire l'erreur de représentativité des observations. De plus, les conditions aux limites semblent être une forte source d'incertitude. Notre utilisation de MOZART, dans ses versions 2 ou 4, ne semble pas optimale. Une solution pourrait être d'utiliser un modèle global par exemple. Les champs d'émission d'EMEP qui

ont été utilisés, peuvent aussi être améliorés, en tentant de faire de la modélisation inverse et estimer ces flux (Liu *et al.*, 2017, soumis, et les références qui y figurent). Par ailleurs, d'autres inventaires, comme celui de GEMS-TNO, peuvent être utilisés. Enfin, la météorologie est une forte source d'incertitude. Les champs de réanalyses auraient pu être utilisés plutôt que le modèle opérationnel du CEPMMT. De plus, la modélisation d'une inflation implémentée par la perturbation des modèles de chacun des membres de l'ensemble aurait pu permettre d'améliorer les résultats (Constantinescu *et al.*, 2007a; Wu *et al.*, 2008; Boynard *et al.*, 2011). Cependant, il aurait surtout été préférable d'avoir à notre disposition un modèle couplé de chimie-météorologie (CCMM), comme WRF-Chem par exemple, pour que les deux sous-systèmes du modèle puissent interagir lors de l'assimilation, comme pour le L95-T ou le L95-GRS. Des expériences d'assimilation de données sur des modèles fortement couplés ont déjà prouvé l'intérêt de cette démarche (Semane *et al.*, 2009; Milewski et Bourqui, 2011, 2013).

L'autre voie possible d'amélioration peut être le recours à des méthodes plus avancées pour tenir compte de l'erreur modèle. Tout d'abord, nous notons que seules de courtes fenêtres ont été testées pour l'IEEnKS, car moins coûteuses numériquement à appliquer. Une fenêtre plus longue, notamment de 24 h, aurait pu permettre au système d'avoir à sa disposition un cycle journalier complet, et ainsi améliorer les prévisions. Ensuite, bien que nos expériences avec l'IEEnKF-Q se soient soldées par un échec provisoire, il existe d'autres méthodes pour tenir compte de l'erreur modèle, y compris dans un contexte variationnel d'ensemble comme le nôtre. Nous pensons par exemple au récent algorithme de Mandel *et al.* (2016) ou encore ceux de Piccolo et Cullen (2016); Emili *et al.* (2016), qui auraient pu être appliqués pour améliorer les performances en prévision.

Nos expériences ont donc le mérite de montrer en premier les limites de l'application des méthodes variationnelles d'ensemble sur des CTMs avec des observations réelles. En effet, nous pensons qu'il n'y a que peu d'utilité pour ces modèles stables de recourir à une longue fenêtre d'assimilation de données, du moins tant qu'une méthode adaptée de prise en compte de l'erreur modèle n'est pas implémentée. Ainsi, de nombreuses voies d'amélioration sont encore à explorer. Ces travaux seront donc poursuivis et feront l'objet d'une publication.

Méthodes de quantification des incertitudes liées à l'estimation d'une source de polluants

L'erreur modèle associée aux expériences présentées au chapitre 5 a fortement impacté les résultats des expériences d'assimilation de données. Il est donc nécessaire d'être en mesure de la diagnostiquer correctement. En particulier, les champs d'émission sont une forte source d'erreur des modèles de chimie-transport. Ils sont généralement issus d'inventaires d'émissions. Cependant, de tels inventaires n'existent pas lors de l'émission accidentelle de polluants, comme lors d'accidents nucléaires. Or il est nécessaire d'avoir une connaissance la plus détaillée possible de ces rejets pour pouvoir évacuer les zones dangereuses et évaluer l'impact sanitaire de l'accident. Il faut alors recourir, pour estimer la quantité de polluants rejetée, à des méthodes de modélisation inverse. Ces méthodes doivent répondre à deux défis : d'une part, améliorer l'estimation des champs d'émission pour réduire l'erreur modèle, et d'autre part, déterminer l'incertitude associée à cette estimation.

L'application des méthodes de modélisation inverse pour l'estimation de sources de polluants a significativement progressé au cours des deux dernières décennies. Cependant, malgré des estimations relativement fiables, ces dernières ne sont que rarement accompagnées d'une évaluation de leurs incertitudes, sauf quand des statistiques gaussiennes sont utilisées. Dans ce chapitre, nous testons des techniques rigoureuses d'estimation de l'incertitude dans le contexte de la modélisation inverse du terme source d'un traceur émis ponctuellement. Des statistiques log-normales sont utilisées pour le terme source et éventuellement les observations, ce qui interdit le recours à des techniques basées sur des statistiques gaussiennes.

Dans un premier temps, via l'approche bayésienne empirique, les paramètres des statistiques d'erreur, appelés hyperparamètres, sont estimés grâce à l'algorithme espérance-maximisation. Cela permet d'estimer un terme source. Puis les incertitudes attachées à la masse totale de polluant émise, ainsi que les taux d'émission au cours du temps, sont estimés par quatre méthodes de Monte-Carlo : (i) un échantillonnage préférentiel fondé sur l'approximation de Laplace ; (ii) un échantillonnage avec un *randomize-then-optimize* (RTO) naïf ; (iii) un échantillonnage avec un RTO sans biais ; (iv) une simulation avec une méthode basique de Monte-Carlo par chaînes de Markov (MCMC). Dans un deuxième temps, ces méthodes sont comparées à une approche bayésienne hiérarchique complète, en utilisant une méthode de MCMC fondée sur une représentation transdimensionnelle du terme source pour réduire le coût de calcul.

Ces méthodes sont appliquées à l'estimation des termes sources de césium-137 des accidents des centrales nucléaires de Tchernobyl en avril et mai 1986 et de Fukushima Daiichi en mars 2011. Ces résultats ont fait l'objet d'une publication (Liu *et al.*, 2017, [soumis](#)). Il s'agit de la première quantification rigoureuse de l'incertitude des estimations des termes sources de ces accidents.

Sommaire

6.1 Introduction	131
-----------------------------------	------------

6.1.1	Modélisation inverse et quantification d'incertitude	131
6.1.2	Modélisation inverse pour le rejet accidentel de radionucléides	131
6.1.3	Quantification d'incertitude et objectifs de cette étude	131
6.2	Les rejets de radionucléides des accidents de Tchernobyl et Fukushima . .	132
6.2.1	Contexte	132
6.2.2	Estimation du terme source pour les accidents de Tchernobyl et de Fukushima Daiichi	133
6.2.3	Base de données d'observations	134
6.2.4	Modélisation	134
6.3	Modélisation inverse	134
6.3.1	Statistiques d'erreur a priori log-normales	134
6.3.2	Solution bayésienne théorique	136
6.3.3	Les solutions par Bayes hiérarchique et Bayes empirique	137
6.3.4	Estimation des hyperparamètres avec l'algorithme espérance-maximisation . . .	138
6.3.5	A priori sur les hyperparamètres	140
6.3.6	Application aux accidents de Tchernobyl et Fukushima Daiichi	141
6.4	Quantification des incertitudes	144
6.4.1	Échantillonnage préférentiel fondé sur l'approximation de Laplace	144
6.4.2	Échantillonnage par <i>randomize-then-optimize</i> naïf et sans biais	146
6.4.3	Méthode basique de Monte-Carlo par chaînes de Markov	147
6.4.4	Application aux accidents de Tchernobyl et Fukushima Daiichi	147
6.5	Hiérarchie bayésienne complète	148
6.5.1	Méthode hiérarchique de Monte-Carlo par chaînes de Markov	149
6.5.2	Analyse transdimensionnelle	149
6.5.3	Application aux accidents de Tchernobyl et Fukushima Daiichi	150
6.6	Résumé et conclusions	153

6.1 Introduction

6.1.1 Modélisation inverse et quantification d'incertitude

La modélisation inverse des champs d'émission de polluants est un défi en chimie atmosphérique, particulièrement en ce qui concerne les gaz à effet de serre (par exemple, [Gurney *et al.*, 2002](#); [Kaminski *et al.*, 2001](#); [Chevallier *et al.*, 2006](#); [Bergamaschi *et al.*, 2009](#)) et la qualité de l'air (par exemple, [Quélo *et al.*, 2005](#); [Elbern *et al.*, 2007](#); [Koohkan et Bocquet, 2012](#); [Koohkan *et al.*, 2013](#)). Il s'agit d'assimiler des observations de concentration de polluants et d'estimer les flux de polluants en utilisant un modèle de chimie-transport (CTM) avec une représentation statistique correcte des erreurs (voir, par exemple, la section 3 de [Zhang *et al.*, 2012](#)). Cette approche, considérée comme *top-down*, est complémentaire à celle consistant à construire un inventaire d'émission, qui est considérée *bottom-up*. Comme pour les méthodes d'assimilation de données utilisées en géophysique, la majorité des études de modélisation inverse repose sur des principes bayésiens. En particulier, l'a priori de l'approche bayésienne, qui synthétise l'ensemble des informations disponibles avant l'assimilation des observations, est fourni par un inventaire et si possible une estimation a priori de l'incertitude intrinsèque de ce dernier.

6.1.2 Modélisation inverse pour le rejet accidentel de radionucléides

La modélisation inverse a également été appliquée à l'estimation de la quantité de radionucléides libérée dans l'atmosphère lors d'un accident nucléaire. Les taux d'émission à la centrale nucléaire sur une période de temps pour un ou plusieurs radionucléides se nomment le terme source. Le transport, la microphysique ou la décroissance radioactive des radionucléides peuvent souvent être vus comme linéaires, c'est-à-dire que la relation source-récepteur, qui relie les mesures de concentration au terme source, est linéaire. En comparaison des autres problèmes de modélisation inverse, en particulier ceux mentionnés plus haut, la modélisation numérique est donc considérablement plus simple. La position exacte de l'émission est souvent connue, aussi les seules variables à reconstruire sont les taux horaires d'émission sur plusieurs jours, pour des accidents majeurs comme ceux de Tchernobyl ou Fukushima.

Toutefois, le fait que la source soit ponctuelle peut provoquer des phénomènes de double pénalité (voir, dans ce contexte, [Farchi *et al.*, 2016](#)) près des forts gradients de panache de polluants. De plus, il n'y a dans ce cas précis pas d'inventaire à proprement parler, même si des physiciens nucléaires peuvent évaluer ce qui pourrait être émis de la centrale, en fonction du scénario de l'accident. Ainsi, on se doit de supposer que l'a priori sur les émissions a une distribution très large et non-informative.

Comme nous le verrons dans la section 6.2, la modélisation inverse a été appliquée à l'estimation des termes sources des accidents de la centrale nucléaire de Tchernobyl (CNPP, pour Chernobyl nuclear power plant en anglais) et plus récemment de la centrale nucléaire de Fukushima Daiichi (FDNPP, pour Fukushima Daiichi nuclear power plant en anglais).

6.1.3 Quantification d'incertitude et objectifs de cette étude

Dans ce contexte accidentel, la plupart des inversions publiées n'ont pas été obtenues avec une quantification de l'incertitude (UQ, pour uncertainty quantification en anglais) de la meilleure estimation du terme source. Toutefois, [Winiarek *et al.* \(2014\)](#) s'est appuyé sur une méthode de Monte-Carlo (décrite plus tard comme RTO naïf) pour fournir un intervalle de confiance pour le terme source de césium-137 de la FDNPP.

D'une manière plus générale, on observe, en modélisation inverse pour la chimie atmosphérique, une tendance récente à essayer de quantifier l'incertitude des émissions estimées. Cette démarche est naturelle en assimilation de données météorologiques, où elle repose sur une ébauche réaliste et souvent gaussienne. Cela n'est pas aussi évident pour les problèmes de modélisation inverse car une telle ébauche n'est souvent pas disponible ou peu fiable. Ce problème est amplifié par le fait que les données d'observation étant limitées, l'ébauche joue un rôle primordial dans l'inversion. De plus, l'UQ est mathématiquement aisée si les erreurs peuvent être considérées comme gaussiennes. Ce n'est cependant pas le cas des émissions qui sont des variables aléatoires positives et de facto intrinsèquement non-gaussiennes.

Le but de ce chapitre est donc de présenter, tester, comparer et critiquer les méthodes objectives d'estimation de l'incertitude du terme source reconstruit par modélisation inverse. Certaines des méthodes répertoriées sont présentées avec de nouvelles variations. La plupart d'entre elles n'ont jamais été testées dans ce domaine.

À la section 6.3, nous introduirons le cadre théorique, reposant sur des méthodes bayésiennes, qui tient compte de la positivité du terme source et utilise des statistiques log-normales pour les erreurs d'ébauche. Il s'appuie sur une approximation, appelée Bayes empirique (BE), qui sera implémentée au travers de l'algorithme espérance-maximisation (EM). À la section 6.4, la quantification des incertitudes des estimations obtenues par l'approche BE sera étudiée avec les méthodes suivantes : un échantillonnage préférentiel fondé sur l'approximation de Laplace, le RTO naïf, le RTO sans biais, et une méthode MCMC. À la section 6.5, une approche bayésienne hiérarchique est développée pour résoudre conjointement le problème inverse ainsi que l'UQ. Pour simplifier les calculs numériques, le maillage est intégré à la hiérarchie bayésienne, ce qui est nommé analyse transdimensionnelle. Un résumé et des conclusions sont fournies dans la section 6.6.

Pour chacune des sections principales 6.3, 6.4 et 6.5, nous illustrons les méthodes en les appliquant à l'inversion des termes sources des accidents des CNPP et FDNPP, en utilisant des observations réelles de concentration. Nous fournissons ainsi la première UQ rigoureuse de ces inversions. Sans chercher à être exhaustifs, nous souhaitons mettre en exergue les éléments clés de ces méthodes et les résultats marquants. L'erreur modèle, y compris celle provenant des champs météorologiques, n'est pas considérée dans cette étude. Cependant, elle sera partiellement prise en compte par l'estimation des paramètres statistiques du *a priori*.

Étant donné le choix de ces illustrations, nous rappelons le contexte des accidents des CNPP et FDNPP, et décrivons leur modélisation dans la section 6.2, avant de décrire les inversions et les méthodes d'UQ.

6.2 Les rejets de radionucléides des accidents de Tchernobyl et Fukushima

Dans cette section, nous rappelons les caractéristiques des accidents de la CNPP et la FDNPP et leurs conséquences en terme de dispersion atmosphérique de radionucléides. Les différentes études utilisant la modélisation inverse pour estimer le terme source sont également listées. Enfin, nous décrivons les jeux d'observations et les modèles que nous utiliserons lors de l'inversion des termes sources et du calcul de leurs incertitudes.

6.2.1 Contexte

Le 25 avril 1986, à 21h30 UTC, la CNPP subit un accident majeur avec l'explosion du quatrième réacteur, un important feu et la fusion du réacteur. En conséquence, des radionucléides ont

été répandus dans l'atmosphère, y compris des espèces résilientes comme l'iode-131, le césium-134 et également le césium-137, d'une demi-vie de 30 ans. L'émission principale de radionucléides a duré environ 10 jours. Les radionucléides se sont déposés principalement en Ukraine, en Biélorussie et en Scandinavie. L'United Nations Scientific Committee on the Effects of Atomic Radiation (UNSCEAR) estime les quantités émises d'iode-131 et de césium-137 à environ respectivement 1760 PBq et 85 PBq (United Nations, 2000).

À la suite d'un tremblement de terre sous-marin de magnitude 9.0 et du tsunami qui en a découlé le 11 mars 2011, une fusion partielle du réacteur et des explosions d'hydrogène dans plusieurs réacteurs de la FDNPP ont provoqué le rejet massif de radionucléides dans l'atmosphère, sans compter le déversement d'eaux contaminées dans l'océan Pacifique. Des radionucléides, en particulier du césium-137, ont été déposés dans la préfecture de Fukushima, au Nord-Ouest de la FDNPP. L'émission principale a eu lieu sur une période de trois semaines à partir du 11 mars. Les analyses s'accordent sur une quantité émise comprise entre 100 et 500 PBq pour l'iode-131 et entre 6 et 20 PBq pour le césium-137 (United Nations, 2016).

6.2.2 Estimation du terme source pour les accidents de Tchernobyl et de Fukushima Daiichi

Lors d'accidents nucléaires graves, l'évaluation de l'impact environnemental dépend fortement de l'évolution temporelle du taux de rejet dans l'atmosphère et de sa distribution entre les différents radionucléides. C'est ce que l'on nomme terme source. Une méthode d'estimation de ce terme source consiste à faire l'inventaire des quantités de radionucléides présents dans le réacteur et leur comportement en cas de fusion du cœur du réacteur. Les autres méthodes consistent à combiner un modèle numérique de dispersion atmosphérique et des observations environnementales. Les méthodes de modélisation inverse en font partie. Elles ont vu le jour après l'accident de la CNPP. Elles s'appuient sur un formalisme mathématique rigoureux et se sont avérées efficaces pour reconstituer le terme source (Gudiksen *et al.*, 1989; Davoine et Bocquet, 2007; Krysta et Bocquet, 2007; Bocquet, 2012).

Six ans après l'accident de la FDNPP, la plupart des estimations du terme source ont été obtenues en couplant des observations et des modèles de dispersion atmosphérique. Deux types d'approche peuvent être distingués : des méthodes directes, qui comparent simplement les données observées et mesurées (Chino *et al.*, 2011; Mathieu *et al.*, 2012; Terada *et al.*, 2012; Katata *et al.*, 2012; Hirao *et al.*, 2013; Kobayashi *et al.*, 2013; Katata *et al.*, 2015), et des méthodes de modélisation inverse (Winiarek *et al.*, 2012; Stohl *et al.*, 2012; Saunier *et al.*, 2013; Winiarek *et al.*, 2014; Yumimoto *et al.*, 2016). Ces estimations ont fourni les quantités totales de radionucléides rejetés qui ont été mentionnées précédemment. Cependant, les termes sources reconstruits sont sensibles au type de mesures utilisé, à la qualité des champs météorologiques et à celle du modèle de dispersion atmosphérique. Les estimations récentes (Katata *et al.*, 2015; Yumimoto *et al.*, 2016) sont toujours incertaines et diffèrent grandement sur les taux de rejet. Les comparaisons entre les résultats de simulations et les mesures de concentration de particules de césium-137 dans l'air (Tsuruta *et al.*, 2014; Oura *et al.*, 2015) ont mis en avant le besoin d'améliorer la reconstruction du terme source (Nakajima *et al.*, 2017). L'assimilation de ces mêmes mesures de concentration de césium-137 dans l'air peut le permettre (Saunier *et al.*, 2016).

6.2.3 Base de données d'observations

La reconstruction du terme source de l'accident de Tchernobyl repose sur des observations récoltées dans la base de données Radioactivity Environmental Monitoring (REM) du Joint Research Centre. Les mesures de concentration de césium-137 dans l'air effectuées en Europe de l'Ouest par 92 stations sont utilisées (1293 observations disponibles). Étant donné que cette base de données a déjà permis plusieurs inversions du terme source (Davoine et Bocquet, 2007; Bocquet, 2012), elle semble adaptée pour étudier les méthodologies développées dans ce chapitre. Cependant, une étude du terme source à proprement parler bénéficierait probablement des efforts récents apportés à l'enrichissement de la base de données REM (Evangeliou *et al.*, 2016).

Les mesures de concentration de césium-137 dans l'air effectuées dans un rayon de 200 km autour de la FDNPP sont utilisées pour reconstruire le terme source dans notre étude. Cela correspond aux stations situées dans la région de Tohoku et du Nord de la région de Kanto. Nous nous défaussons des stations situées plus au Sud. Le choix de cette zone géographique restreinte permet de limiter l'impact de l'erreur modèle sur la reconstruction du terme source. La base de données originelle est constituée de 6451 concentrations mesurées, dont nous ne gardons que les 3294 situées dans un rayon de 200 km de la FDNPP. Cela inclut les données décrites et utilisées dans Winiarek *et al.* (2012) et les données rendues disponibles suite à l'immense travail crucial effectué par Oura *et al.* (2015).

6.2.4 Modélisation

La dispersion atmosphérique des panaches de césium-137 issus des CNPP et FDNPP est simulée par le modèle eulérien de la plate-forme Polyphemus, déjà utilisé au chapitre 5 (Mallet *et al.*, 2007; Quélo *et al.*, 2007). Les paramètres utilisés sont choisis sur la base de précédents travaux (Bocquet, 2012; Winiarek *et al.*, 2014; Quérel *et al.*, 2015). Les principales caractéristiques de ces configurations sont rappelées dans le tableau 6.1. Les simulations démarrent le 25 avril 1986 à 00h00 UTC et courent sur 14 jours pour l'accident de la CNPP et démarrent le 11 mars 2011 à 00h00 UTC et courent sur 24 jours pour l'accident de la FDNPP.

Nous notons que K_h a été fixé à 0 pour le cas de la CNPP, puisque la résolution est grossière et implique déjà une forte diffusion numérique. Au contraire, la configuration pour l'accident de la FDNPP a une résolution spatiale bien plus fine qui par conséquent peut nécessiter un fort K_h . Nous avons choisi $K_h = 25\,000\text{ m}^2\text{s}^{-1}$, conformément à Winiarek *et al.* (2014), afin de mitiger grâce à la diffusion l'erreur modèle et l'effet négatif de la double pénalité qui en résulte (Farchi *et al.*, 2016).

6.3 Modélisation inverse

Dans cette section, nous présentons le cadre méthodologique pour les inversions effectuées dans ce chapitre, qui repose sur des techniques bayésiennes approchées ou exactes.

6.3.1 Statistiques d'erreur a priori log-normales

Pour s'assurer de la positivité de la source estimée par modélisation inverse dans un cadre bayésien, une densité de probabilité positive des taux d'émission doit être utilisée. Par exemple, Winiarek *et al.* (2012, 2014); Tichý *et al.* (2016) ont choisi des statistiques gaussiennes tronquées. Nous choisissons ici des distributions log-normales. Une des raisons de ce choix est la large plage de valeurs de concentration des polluants et de leurs erreurs, qui pose problème pour les rejets de

TABLEAU 6.1 – Principales caractéristiques des configurations des simulations pour les accidents de la CNPP et la FDNPP.

	Tchernobyl	Fukushima Daiichi
extension du domaine	[10.96°E-68.915°W] et [30.57°N-72.195°N]	[137.53°W-143.48°W] et [34.72°N-40.67°N]
résolution angulaire	1.125°×1.125°	0.05°×0.05°
résolution verticale	15 niveaux parallèles au sol (de 0 à 8 073 m)	
météorologie	réanalyses ERA-40 du CEPMMT (Uppala <i>et al.</i> , 2005)	simulation WRF (Winiarek <i>et al.</i> , 2014)
mélange vertical	coefficient vertical de diffusion turbulente, K_z , suivant Louis (1979)	
mélange horizontal	coefficient horizontal de diffusion turbulente constant K_h $K_h = 0 \text{ m}^2 \text{ s}^{-1}$	$K_h = 25\,000 \text{ m}^2 \text{ s}^{-1}$
lessivage dans le nuage	Pudykiewicz (1989) avec $\Lambda^a = 1.5 \cdot 10^{-5} \text{ s}^{-1}$	Roselle et Binkowski (1999)
dépôt humide	aucun	Baklanov et Sørensen (2001)
dépôt sec	vélocité constante $v_d = 1.45 \cdot 10^{-3} \text{ m s}^{-1}$	$v_d = 1.5 \cdot 10^{-3} \text{ m s}^{-1}$ sur la terre $v_d = 0.1 \cdot 10^{-3} \text{ m s}^{-1}$ sur l'eau

^a Λ est la constante produisant la valeur maximale de coefficient de lessivage.

radionucléides. En effet, on observe des écarts d'un facteur 10^5 entre les niveaux de radioactivité possibles. Cette remarque vaut aussi bien pour les observations que pour les taux de rejet de la source. Ce problème est significativement amoindri par le recours à des statistiques log-normales. En pratique, cela donne plus de poids aux faibles observations.

Avec des statistiques log-normales, le terme source $\mathbf{x}$, composé de M taux d'émission et vu comme une variable aléatoire vectorielle dans $\mathbb{R}_+^M$, a la forme

$$\mathbf{x} = \bar{\mathbf{x}}e^{\mathbf{z}}, \text{ avec } \mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{B}), \quad (6.1)$$

qui signifie que $\mathbf{z}$ est échantillonné à partir de $\mathcal{N}(\mathbf{0}, \mathbf{B})$, la distribution gaussienne multivariée de moyenne $\mathbf{0}$ et de matrice de covariance $\mathbf{B}$, et a pour fonction de densité de probabilité (pdf)

$$p(\mathbf{z}) = \frac{e^{-\frac{1}{2}\|\mathbf{z}\|_{\mathbf{B}}^2}}{\sqrt{(2\pi)^M |\mathbf{B}|}}, \quad (6.2)$$

où on rappelle que $\|\mathbf{z}\|_{\mathbf{B}}^2 = \mathbf{z}^T \mathbf{B}^{-1} \mathbf{z}$ et $|\mathbf{B}|$ est le déterminant de $\mathbf{B}$. Le vecteur $\bar{\mathbf{x}}$ dans l'équation (6.1) donne un ordre de grandeur des taux d'émission et pourrait être une ébauche du terme source. À partir du changement de variable défini à l'équation (6.1), on déduit que la pdf de $\mathbf{x}$ est (voir, par exemple, Fletcher et Zupanski, 2006)

$$p(\mathbf{x}) = \frac{e^{-\sum_{m=1}^M \ln x_m - \frac{1}{2}\|\ln \mathbf{x} - \ln \bar{\mathbf{x}}\|_{\mathbf{B}}^2}}{\sqrt{(2\pi)^M |\mathbf{B}|}}. \quad (6.3)$$

De la même manière, les P observations, considérées comme des variables aléatoires dans $\mathbb{R}_+^P$, sont supposées avoir des statistiques log-normales :

$$\mathbf{y} = \boldsymbol{\mu}e^{\boldsymbol{\varepsilon}}, \text{ avec } \boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{R}). \quad (6.4)$$

Par conséquent, la pdf des observations s'écrit

$$p(\mathbf{y}) = \frac{e^{-\sum_{p=1}^P \ln y_p - \frac{1}{2}\|\ln \mathbf{y} - \ln \boldsymbol{\mu}\|_{\mathbf{R}}^2}}{\sqrt{(2\pi)^P |\mathbf{R}|}}. \quad (6.5)$$

L'opportunité qu'offrent des statistiques log-normales de traiter des observations d'ordres de grandeur significativement différents est un grand avantage. Le défaut de cette approche est toutefois la trop grande précision implicitement associée aux observations quasi-nulles, dont la fiabilité peut être remise en question (par exemple, [Farchi et al., 2016](#)). Plusieurs solutions sont envisageables pour mitiger cet effet. Premièrement, on peut se contenter de ne considérer que les observations qui sont au-dessus d'un certain seuil y_t . Ce seuil peut être choisi comme valant par exemple la précision de mesure de radioactivité atmosphérique ou l'activité de fond pour des mesures d'activités déposées.

Une solution plus élégante consiste à modifier la distribution log-normale et l'amender pour les petites valeurs d'observation. Par exemple, on peut considérer la distribution suivante pour les observations,

$$p(\mathbf{y}) \propto e^{\sum_{p=1}^P \ln|\zeta'(y_p)| - \frac{1}{2} \|\zeta(\mathbf{y}) - \zeta(\boldsymbol{\mu})\|_{\mathbf{R}}^2}, \quad (6.6)$$

où la fonction $\zeta : \mathbb{R}_+ \rightarrow \mathbb{R}$ s'applique composante par composante à $\mathbf{y}$ et $\boldsymbol{\mu}$, et ζ' est sa dérivée. On peut la définir suivant $\zeta(y_p) = \ln y_p$ si $y_p \geq y_t$ et $\zeta(y_p) = \ln y_t - 1 + y_p/y_t$ si $0 \leq y_p \leq y_t$. La fonction a été choisie de telle sorte à ce qu'elle ait un comportement log-normal pour la majorité des observations ($y_p \geq y_t$) et qu'elle soit linéaire pour les faibles valeurs, en plus d'être continue et différentiable en y_t . Cela tempère l'impact des observations de faible concentration. Une troisième possibilité intermédiaire est de choisir simplement $y_p \mapsto \zeta(y_p) = \ln(y_p + y_t)$. Cette dernière a le mérite de se passer d'un test conditionnel, de faciliter le calcul analytique des vraisemblances et de leurs dérivées, et d'éviter de se défausser de certaines observations. Nous avons vérifié que les trois solutions conduisent à des inversions similaires quand y_t est choisi de telle sorte à avoir un sens physique. Par conséquent, nous avons sélectionné la dernière solution. La valeur du seuil est fixée à $y_t = 10^{-3} \text{ Bq.m}^{-3}$ pour l'étude de l'accident de la CNPP et à $y_t = 10^{-1} \text{ Bq.m}^{-3}$ pour celui de la FDNPP.

Enfin, à titre de comparaison, nous considérerons également le cas d'observations gaussiennes, c'est-à-dire telle que $\mathbf{y} = \boldsymbol{\mu} + \boldsymbol{\varepsilon}$ avec $\boldsymbol{\varepsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{R})$, tout en gardant les erreurs sur l'ébauche log-normales.

6.3.2 Solution bayésienne théorique

La distribution des statistiques des erreurs a priori pour le terme source et les observations étant choisie, nous pouvons désormais présenter les principes de l'inversion appliquée dans ce chapitre. Le problème inverse est résolu formellement par la formule de Bayes, qui s'écrit pour deux variables $\mathbf{z}$ et $\mathbf{y}$,

$$p(\mathbf{z}|\mathbf{y}) = \frac{p(\mathbf{y}|\mathbf{z})p(\mathbf{z})}{p(\mathbf{y})}, \text{ où } p(\mathbf{y}) = \int_{\mathbb{R}^M} d\mathbf{z} p(\mathbf{y}|\mathbf{z})p(\mathbf{z}). \quad (6.7)$$

La pdf $p(\mathbf{y}|\mathbf{z})$ est la vraisemblance des observations, issue des erreurs d'observation a priori; la pdf $p(\mathbf{z})$ est l'ébauche du terme source; et la pdf $p(\mathbf{y})$, constante selon $\mathbf{z}$, est l'évidence.

Les observations (non bruitées) et le terme source sont reliés par la résolvante $\mathbf{H}$ du modèle de chimie-transport

$$\boldsymbol{\mu} = \mathbf{H}\mathbf{x} = \mathbf{H}\bar{\mathbf{x}}e^{\mathbf{z}}, \quad (6.8)$$

où $\mathbf{H}$ est souvent appelée matrice jacobienne du modèle et définit la relation source-récepteur. Notons que le modèle physique est ici supposé linéaire. L'opposé du logarithme de la vraisemblance de $\mathbf{z}$, étant donné l'ensemble des paramètres $\boldsymbol{\theta} \equiv \bar{\mathbf{x}}, \mathbf{B}, \mathbf{R}$, c'est-à-dire $\mathcal{L}(\mathbf{z}; \boldsymbol{\theta}) = -\ln p(\mathbf{z}|\mathbf{y}, \boldsymbol{\theta})$, est obtenu par l'application de la loi de Bayes :

$$\mathcal{L}(\mathbf{z}; \boldsymbol{\theta}) = \frac{1}{2} \|\ln \mathbf{y} - \ln \mathbf{H}\bar{\mathbf{x}}e^{\mathbf{z}}\|_{\mathbf{R}}^2 + \frac{1}{2} \|\mathbf{z}\|_{\mathbf{B}}^2 + \xi(\boldsymbol{\theta}), \quad (6.9)$$

où $\xi(\boldsymbol{\theta})$ représente les termes qui ne dépendent pas de $\mathbf{z}$. Quand nous cherchons le maximum a posteriori (MAP), c'est-à-dire le minimum de $\mathcal{L}(\mathbf{z}; \boldsymbol{\theta})$ en fonction de $\mathbf{z}$, il est préférable d'appliquer le changement de variable $\mathbf{z} = \mathbf{B}^{\frac{1}{2}} \tilde{\mathbf{z}}$, avec $\mathbf{B}^{\frac{1}{2}}$ la matrice racine carrée de $\mathbf{B}$, qui produit une fonction de coût mieux conditionnée :

$$\tilde{\mathcal{L}}(\tilde{\mathbf{z}}; \boldsymbol{\theta}) = \frac{1}{2} \left\| \ln \mathbf{y} - \ln \mathbf{H} \bar{\mathbf{x}} e^{\mathbf{B}^{\frac{1}{2}} \tilde{\mathbf{z}}} \right\|_{\mathbf{R}}^2 + \frac{1}{2} \|\tilde{\mathbf{z}}\|^2 + \tilde{\xi}(\boldsymbol{\theta}), \quad (6.10)$$

où on utilise la norme $\|\mathbf{v}\| = \sqrt{\mathbf{v}^T \mathbf{v}}$. Ce préconditionnement est nécessaire quand les valeurs propres de $\mathbf{B}$ approchent 0 et impactent la convergence de la minimisation de (6.9). Ceci étant, par la suite, les développements théoriques utiliseront la formulation de (6.9), plus conventionnelle.

6.3.3 Les solutions par Bayes hiérarchique et Bayes empirique

Si les statistiques des erreurs sont effectivement log-normales et que les paramètres $\boldsymbol{\theta}$ sont précis, la méthode présentée plus haut devrait fournir une estimation fiable de la pdf a posteriori du problème inverse. Cependant, $\boldsymbol{\theta}$, souvent appelés hyperparamètres, sont incertains, tout particulièrement en ce qui concerne l'inversion de sources de polluants. En statistique bayésienne, il est par conséquent conseillé de considérer $\boldsymbol{\theta}$ comme un vecteur aléatoire et de l'incorporer dans la formule de Bayes habituelle

$$p(\mathbf{x}|\mathbf{y}) = \frac{p(\mathbf{y}|\mathbf{x})p(\mathbf{x})}{p(\mathbf{y})}, \quad (6.11)$$

ce qui donne la hiérarchie d'inférences suivante :

$$p(\mathbf{x}, \boldsymbol{\theta}|\mathbf{y}) = \frac{p(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta})p(\mathbf{x}, \boldsymbol{\theta})}{p(\mathbf{y})} = \frac{p(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta})p(\mathbf{x}|\boldsymbol{\theta})p(\boldsymbol{\theta})}{p(\mathbf{y})}, \quad (6.12)$$

où $p(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta})$ est la vraisemblance des observations conditionnée par $\mathbf{x}$ et $\boldsymbol{\theta}$ et $p(\boldsymbol{\theta})$ est l'ébauche de ces hyperparamètres. Il s'agit de l'approche hiérarchique bayésienne complète (BC). La pdf a posteriori initiale $p(\mathbf{x}|\mathbf{y})$, sans référence aux hyperparamètres $\boldsymbol{\theta}$, peut ensuite être obtenue comme une marginalisation sur $\boldsymbol{\theta}$

$$p(\mathbf{x}|\mathbf{y}) = \int d\boldsymbol{\theta} p(\mathbf{x}, \boldsymbol{\theta}|\mathbf{y}). \quad (6.13)$$

Une approximation consiste à faire apparaître $p(\mathbf{x}|\mathbf{y}, \boldsymbol{\theta})$ dans cette marginalisation

$$p(\mathbf{x}|\mathbf{y}) = \int d\boldsymbol{\theta} p(\mathbf{x}|\mathbf{y}, \boldsymbol{\theta})p(\boldsymbol{\theta}|\mathbf{y}) \quad (6.14)$$

et supposer que $p(\boldsymbol{\theta}|\mathbf{y})$ est piquée autour d'une valeur $\boldsymbol{\theta}^*$ et concentrée sur un voisinage de $\boldsymbol{\theta}^*$ où la variation de $p(\mathbf{x}|\mathbf{y}, \boldsymbol{\theta})$ est faible. Cela amène à une marginalisation approchée

$$p(\mathbf{x}|\mathbf{y}) \approx p(\mathbf{x}|\mathbf{y}, \boldsymbol{\theta}^*) \int d\boldsymbol{\theta} p(\boldsymbol{\theta}|\mathbf{y}) = p(\mathbf{x}|\mathbf{y}, \boldsymbol{\theta}^*). \quad (6.15)$$

Cette approche est nommée Bayes empirique (BE). Même si cette démarche n'est qu'approchée, elle peut représenter avec fidélité la hiérarchie bayésienne, à supposer que $p(\boldsymbol{\theta}|\mathbf{y})$ soit suffisamment piquée autour de $\boldsymbol{\theta}^*$. Elle est bien connue en statistiques environnementales et a été prônée dans des revues en géosciences (Michalak *et al.*, 2005; Wu *et al.*, 2013; Bocquet, 2015). Elle a été utilisée en modélisation inverse pour le terme source de polluants dangereux dans Davoine et Bocquet (2007); Winiarek *et al.* (2012, 2014) et en chimie atmosphérique (par exemple, Koohkan *et al.*, 2013; Berchet *et al.*, 2013).

Même si elle est répandue en statistiques environnementales, l'approche hiérarchique a rarement été utilisée en géosciences avec des modèles complexes car elle nécessite d'échantillonner selon une succession de pdf, ce qui n'est pas une tâche numériquement aisée. Entre autres exceptions, elle a été utilisée en assimilation de données pour mieux tenir compte des incertitudes dans les méthodes de filtre de Kalman d'ensemble (Bocquet, 2011). Il s'agit de l'EnKF-N décrit à la section 2.4.5.2. On peut citer également Ganesan *et al.* (2014) qui l'a utilisée dans l'inversion de gaz à effet de serre. Étant donné que le nombre de degrés de liberté d'une source ponctuelle avec quelques taux d'émissions est nettement moindre, l'approche bayésienne complète a été rendue possible pour la modélisation inverse du terme source de polluants dangereux dans Delle Monache *et al.* (2008); Yee (2008). Enfin, nous précisons qu'une comparaison plus approfondie entre les inférences bayésiennes hiérarchiques et empiriques est présentée dans Malinverno et Briggs (2004).

6.3.4 Estimation des hyperparamètres avec l'algorithme espérance-maximisation

Dans le cadre du BE, plusieurs estimateurs peuvent être utilisés pour les hyperparamètres, comme le maximum de vraisemblance, la validation croisée généralisée, la L-curve, l'estimateur de risque prédictif sans biais, etc. (voir, par exemple, Vogel, 2002; Hansen, 2010; Doicu *et al.*, 2010, pour des monographies avancées sur le sujet). La plupart d'entre eux ont été utilisés dans les études de modélisation inverse déjà mentionnées. Étant donnée la justification du BE fournie plus haut, le maximum de vraisemblance est souvent employé. Cependant, en l'absence d'expression analytique pour cet estimateur à l'exception du cas de statistiques gaussiennes, sa solution requiert des algorithmes numériques non-triviaux. Une approche rigoureuse, bien connue dans le domaine des statistiques mais assez peu dans celui de l'assimilation de données en géophysique, est celle de l'espérance-maximisation (EM) (Dempster *et al.*, 1977), qui sera utilisée dans ce chapitre pour estimer les hyperparamètres. Le schéma est intéressant car il est facile à implémenter. Son fondement et sa justification sont en revanche sophistiqués, ce qui a pu empêcher dans une certaine mesure sa diffusion aux géosciences.

L'algorithme EM est censé trouver le maximum local de la vraisemblance

$$l(\boldsymbol{\theta}) \equiv p(\mathbf{y}|\boldsymbol{\theta}) = \int_{\mathbb{R}^M} d\mathbf{x} p(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta}) p(\mathbf{x}|\boldsymbol{\theta}), \quad (6.16)$$

quand $p(\mathbf{y}|\boldsymbol{\theta})$ n'est pas nécessairement connue analytiquement car l'intégrale n'est pas calculable. L'algorithme EM est itératif, jusqu'à convergence, et alterne entre les deux étapes :

- **Étape E (espérance)** : Étant donnée la valeur estimée des paramètres $\boldsymbol{\theta}^{(i)}$ à l'itération i , on calcule la fonction $L(\boldsymbol{\theta}|\boldsymbol{\theta}^{(i)}) = \mathbb{E}_{\mathbf{x}|\mathbf{y}, \boldsymbol{\theta}^{(i)}} [\ln p(\mathbf{x}, \mathbf{y}|\boldsymbol{\theta})]$, où l'indice $\mathbf{x}|\mathbf{y}, \boldsymbol{\theta}^{(i)}$ signifie que l'espérance $\mathbb{E}$ est prise sur la distribution de $\mathbf{x}$ sachant $\mathbf{y}$ mais en supposant la valeur fixe $\boldsymbol{\theta}^{(i)}$ pour $\boldsymbol{\theta}$.
- **Étape M (maximisation)** : Étant donné $L(\boldsymbol{\theta}|\boldsymbol{\theta}^{(i)})$ calculée à l'étape d'espérance, nous cherchons son maximum $\boldsymbol{\theta}^{(i+1)} = \underset{\boldsymbol{\theta}}{\operatorname{argmax}} L(\boldsymbol{\theta}|\boldsymbol{\theta}^{(i)})$.

Pour une condition initiale donnée $\boldsymbol{\theta}_0$, on répète les itérations jusqu'à convergence des paramètres. La convergence de l'algorithme EM vers le maximum local dont le bassin d'attraction contient $\boldsymbol{\theta}_0$ est garantie. Cet algorithme est connu pour être efficace pour l'estimation de paramètres de modèles de mélange gaussien (voir par exemple, Bishop, 2006), qui sont souvent utilisés en assimilation de données non-linéaire. Cependant, une expression analytique de $L(\boldsymbol{\theta}|\boldsymbol{\theta}^{(i)})$ n'est bien souvent pas disponible, ce qui est le cas de notre application. Une solution consiste à estimer $L(\boldsymbol{\theta}|\boldsymbol{\theta}^{(i)})$ par une

méthode de Monte-Carlo (Wei et Tanner, 1990), c'est-à-dire obtenir N échantillons de la densité de probabilité $p(\mathbf{x}|\mathbf{y}, \boldsymbol{\theta}^{(i)})$, de telle sorte que

$$L(\boldsymbol{\theta}|\boldsymbol{\theta}^{(i)}) \approx \frac{1}{N} \sum_{n=1}^N \ln p(\mathbf{x}_n^{(i)}, \mathbf{y}|\boldsymbol{\theta}^{(i)}), \quad (6.17)$$

où $\mathbf{x}_n^{(i)} \sim \mathbf{x}|\mathbf{y}, \boldsymbol{\theta}^{(i)}$. En géosciences, cette approche a été mise en avant par Ueno et Nakamura (2014); Tandeo *et al.* (2015); Ueno et Nakamura (2016), où les matrices d'erreur d'observation, d'ébauche et de modèle sont estimées. Leurs modèles physiques sont non-linéaires (contrairement au notre qui est linéaire) et leurs statistiques d'erreur sont gaussiennes, ou un mélange de gaussiennes, (par opposition à log-normales dans notre cas). Cela a de plus été utilisé avec succès par Pulido *et al.* (2017).

Ici, nous appliquons l'algorithme EM pour estimer les hyperparamètres $\boldsymbol{\theta} = \mathbf{R}, \mathbf{B}$ et considérons $\bar{\mathbf{x}}$ comme donné et fixé. À l'itération i , nous avons donc les estimations $\mathbf{R}^{(i)}, \mathbf{B}^{(i)}$ de $\mathbf{R}, \mathbf{B}$. Nous voulons donc estimer la fonctionnelle suivante, redéfinie comme deux fois l'opposé du L précédent :

$$\begin{aligned} L(\mathbf{R}, \mathbf{B}|\mathbf{R}^{(i)}, \mathbf{B}^{(i)}) &= -2\mathbb{E}_{\mathbf{z}|\mathbf{y}, \mathbf{R}^{(i)}, \mathbf{B}^{(i)}} [\ln p(\mathbf{z}, \mathbf{y}|\mathbf{R}, \mathbf{B})] \\ &= -2\mathbb{E}_{\mathbf{z}|\mathbf{y}, \mathbf{R}^{(i)}, \mathbf{B}^{(i)}} [\ln p(\mathbf{y}|\mathbf{z}, \mathbf{R}, \mathbf{B}) + \ln p(\mathbf{z}|\mathbf{R}, \mathbf{B})] \\ &= \mathbb{E}_{\mathbf{z}|\mathbf{y}, \mathbf{R}^{(i)}, \mathbf{B}^{(i)}} \left[\|\ln \mathbf{y} - \ln \mathbf{H}\bar{\mathbf{x}}e^{\mathbf{z}}\|_{\mathbf{R}}^2 \right] + \mathbb{E}_{\mathbf{z}|\mathbf{y}, \mathbf{R}^{(i)}, \mathbf{B}^{(i)}} \left[\|\mathbf{z}\|_{\mathbf{B}}^2 \right] + \ln |\mathbf{R}| + \ln |\mathbf{B}| + (M + P) \ln(2\pi), \end{aligned} \quad (6.18)$$

qui ne peut pas totalement être calculée analytiquement. Nous allons donc, comme annoncé, tirer N échantillons de la distribution a posteriori étant donnés $\mathbf{R}^{(i)}, \mathbf{B}^{(i)}$, c'est-à-dire, $\mathbf{z}_n^{(i)} \sim \mathbf{z}|\mathbf{y}, \mathbf{R}^{(i)}, \mathbf{B}^{(i)}$. Un tel échantillonnage n'est pas trivial et fera l'objet de la section 6.4. Chacune des méthodes d'échantillonnage qui y sont proposées peut fonctionner. Notons cependant que pour atteindre la convergence de l'algorithme EM, le nombre d'échantillons est nettement inférieur à ceux requis par l'UQ exposée dans la section 6.4.

Définissons de manière empirique les moments d'ordre 2

$$\mathbf{S}_o^{(i)} = \frac{1}{N} \sum_{n=1}^N \left(\ln \mathbf{y} - \ln \mathbf{H}\bar{\mathbf{x}}e^{\mathbf{z}_n^{(i)}} \right) \left(\ln \mathbf{y} - \ln \mathbf{H}\bar{\mathbf{x}}e^{\mathbf{z}_n^{(i)}} \right)^T, \quad (6.19)$$

$$\mathbf{S}_b^{(i)} = \frac{1}{N} \sum_{n=1}^N \mathbf{z}_n^{(i)} \left(\mathbf{z}_n^{(i)} \right)^T. \quad (6.20)$$

Avec ces définitions, on peut simplifier l'expression de L en

$$L(\mathbf{R}, \mathbf{B}|\mathbf{R}^{(i)}, \mathbf{B}^{(i)}) \approx \text{Tr} \left(\mathbf{R}^{-1} \mathbf{S}_o^{(i)} \right) + \ln |\mathbf{R}| + \text{Tr} \left(\mathbf{B}^{-1} \mathbf{S}_b^{(i)} \right) + \ln |\mathbf{B}| + (M + P) \ln(2\pi). \quad (6.21)$$

L'étape de maximisation consiste à minimiser $L(\mathbf{R}, \mathbf{B}|\mathbf{R}^{(i)}, \mathbf{B}^{(i)})$ selon $\mathbf{R}$ et $\mathbf{B}$ en annulant les gradients

$$\nabla_{\mathbf{R}} L(\mathbf{R}, \mathbf{B}|\mathbf{R}^{(i)}, \mathbf{B}^{(i)}) = \mathbf{R}^{-1} (\mathbf{R} - \mathbf{S}_o^{(i)}) \mathbf{R}^{-1}, \quad (6.22)$$

$$\nabla_{\mathbf{B}} L(\mathbf{R}, \mathbf{B}|\mathbf{R}^{(i)}, \mathbf{B}^{(i)}) = \mathbf{B}^{-1} (\mathbf{B} - \mathbf{S}_b^{(i)}) \mathbf{B}^{-1}, \quad (6.23)$$

ce qui donne les estimations $\mathbf{R}^{(i+1)}, \mathbf{B}^{(i+1)}$ de $\mathbf{R}, \mathbf{B}$ à l'itération $i + 1$

$$\mathbf{R}^{(i+1)} = \mathbf{S}_o^{(i)}, \quad \mathbf{B}^{(i+1)} = \mathbf{S}_b^{(i)}. \quad (6.24)$$

Deux formes simples de $\mathbf{R}$ et $\mathbf{B}$ sont considérées. Dans un premier temps, nous les choisissons proportionnelles à la matrice identité : $\mathbf{R} = r\mathbf{I}_P$ dans $\mathbb{R}^{P \times P}$ et $\mathbf{B} = b\mathbf{I}_M$ dans $\mathbb{R}^{M \times M}$. Même si cette approximation est quelque peu irréaliste, elle fournit un ordre de grandeur pour ces covariances. Nous nommerons cette configuration pour les hyperparamètres $\mathcal{D}_u$ (pour diagonal, paramètre unique). Dans ce cas, les équations (6.19) et (6.20) se simplifient en équations scalaires

$$s_o^{(i)} = \frac{1}{NP} \sum_{n=1}^N \left\| \ln \mathbf{y} - \ln \mathbf{H} \bar{\mathbf{x}} e^{\mathbf{z}_n^{(i)}} \right\|^2, \quad (6.25)$$

$$s_b^{(i)} = \frac{1}{NM} \sum_{n=1}^N \left\| \mathbf{z}_n^{(i)} \right\|^2, \quad (6.26)$$

et l'équation (6.24) devient

$$r^{(i+1)} = s_o^{(i)}, \quad b^{(i+1)} = s_b^{(i)}. \quad (6.27)$$

Dans un second temps, nous continuons de considérer une forme diagonale pour les matrices de covariances, cette fois-ci en prenant la configuration $\mathbf{R} = r\mathbf{I}_P$ et $\mathbf{B} = \text{Diag}(b_1, b_2, \dots, b_M)$. Dans cette configuration, nous choisissons autant de paramètres de variance que de taux d'émission, étant donné que les différents taux d'émission et leurs statistiques d'erreur sont vraisemblablement indépendants ou assez peu corrélés en temps. Nous nommerons cette configuration des hyperparamètres $\mathcal{D}_m$ (pour diagonal, paramètres multiples). Dans ce cas, les itérations sur $\mathbf{B}$ sont modifiées en

$$s_{b,m}^{(i)} = \frac{1}{N} \sum_{n=1}^N \left(z_{n,m}^{(i)} \right)^2, \quad b_m^{(i+1)} = s_{b,m}^{(i)}. \quad (6.28)$$

pour $m = 1, \dots, M$. Notons que l'algorithme EM peut aussi s'appliquer au cas où les observations ont des erreurs gaussiennes de manière évidente : l'équation (6.25) deviendrait $s_o^{(i)} = \frac{1}{NP} \sum_{n=1}^N \left\| \mathbf{y} - \mathbf{H} \bar{\mathbf{x}} e^{\mathbf{z}_n^{(i)}} \right\|^2$.

6.3.5 A priori sur les hyperparamètres

Dans les cas réels étudiés dans ce chapitre, avec des jeux d'observations limités, les hyperparamètres sont souvent assez peu contraints. Il peut alors être important de définir ce à quoi correspond l'a priori sur ces hyperparamètres. Comme il a été vu dans la section 6.3.4, l'algorithme EM fournit un maximum de $p(\mathbf{y}|\boldsymbol{\theta})$, la vraisemblance des observations. Cependant, dans l'approche bayésienne empirique, nous nous intéressons plutôt au maximum de $p(\boldsymbol{\theta}|\mathbf{y}) = p(\mathbf{y}|\boldsymbol{\theta})p(\boldsymbol{\theta})/p(\mathbf{y})$. Ainsi, appliquer l'EM tel qu'il a été présenté jusqu'ici reviendrait à supposer une distribution uniforme pour $p(\boldsymbol{\theta})$. On pourrait préférer une pdf moins informative, telle l'a priori de Jeffreys (Jeffreys, 1961). Pour les deux formes supposées de $\mathbf{R}$ et $\mathbf{B}$ à la section 6.3.4, cet a priori sur les hyperparamètres r , b ou les b_m s'écrit

$$p(r) \propto 1/\sqrt{r}, \quad p(b) \propto 1/\sqrt{b}, \quad p(b_m) \propto 1/\sqrt{b_m}. \quad (6.29)$$

Dans ce cas, il est aisé de vérifier que les équations (6.25), (6.26) et (6.28) doivent être modifiées

TABLEAU 6.2 – Valeurs clés de l’inversion : les activités totales des termes sources inversés (TRRAs) sont en PBq ; les paramètres r^* sont en $\text{Bq}^2.\text{m}^{-6}$ pour le cas gaussien ; toutes les autres valeurs sont sans unité. Dans la configuration $\mathcal{D}_m$, b^* correspond à la moyenne de tous les b_m^* . Le biais (en $\text{Bq}.\text{m}^{-3}$) est la différence moyenne entre les observations et les observations simulées ; ρ est le coefficient de corrélation de Pearson ; la NMSE est l’erreur moyenne normalisée (sans unité).

		Gaussien						Log-normal					
		TRRA	r^*	b^*	biais	ρ	NMSE	TRRA	r^*	b^*	biais	ρ	NMSE
CNPP	$\mathcal{D}_u$	82.5	1.51	0.49	0.05	0.71	1.81	82.0	1.32	3.99	0.05	0.70	1.87
	$\mathcal{D}_m$	83.5	1.50	0.71	0.05	0.71	1.81	87.2	1.35	2.26	0.05	0.70	1.86
FDNPP	$\mathcal{D}_u$	12.3	564	1.69	2.85	0.57	17.4	24.1	2.02	5.15	0.61	0.21	37.0
	$\mathcal{D}_m$	22.4	549	1.64	2.70	0.58	16.4	20.4	2.05	2.64	0.42	0.22	36.8

en

$$s_o^{(i)} = \frac{1}{N(P+1)} \sum_{n=1}^N \left\| \ln \mathbf{y} - \ln \mathbf{H} \bar{\mathbf{x}} e^{\mathbf{z}_n^{(i)}} \right\|^2, \quad (6.30)$$

$$s_b^{(i)} = \frac{1}{N(M+1)} \sum_{n=1}^N \left\| \mathbf{z}_n^{(i)} \right\|^2, \quad (6.31)$$

$$s_{b,m}^{(i)} = \frac{1}{2N} \sum_{n=1}^N \left(\mathbf{z}_{n,m}^{(i)} \right)^2, \quad (6.32)$$

pour $m = 1, \dots, M$. Les changements $P \rightarrow P+1$, $M \rightarrow M+1$ et $N \rightarrow 2N$ dans ces espérances empiriques proviennent de l’a priori de Jeffreys. Dans le reste du chapitre, nous optons pour cet a priori non-informatif dès que $p(\boldsymbol{\theta})$ est utilisé, ainsi que dans cette extension de la méthode EM.

6.3.6 Application aux accidents de Tchernobyl et Fukushima Daiichi

Avec ce premier jeu d’illustrations, nous appliquons l’approche BE développée au cours de cette section pour estimer le terme source de césium-137 des CNPP et FDNPP, à l’aide de la méthode EM (utilisant l’a priori de Jeffreys) pour estimer les hyperparamètres. Pour les deux cas CNPP et FDNPP, nous implémentons les configurations $\mathcal{D}_u$ et $\mathcal{D}_m$ pour les hyperparamètres. Les termes sources inversés sont montrés sur la Fig. 6.1. Les statistiques gaussiennes et log-normales ont été utilisées pour les observations. Dans les deux cas, les taux d’émission correspondent à des durées de 6 heures. Il y a ainsi 40 taux d’émission (s’étalant sur 10 jours) et 80 taux d’émission (s’étalant sur 20 jours) respectivement dans les cas de la CNPP et la FDNPP. Pour le cas de la CNPP, l’ébauche est le terme source d’UNSCEAR (United Nations, 2000), tandis que pour la FDNPP, elle est choisie constante et telle que la masse totale émise soit de 1 PBq. Le tableau 6.2 donne les valeurs clés des hyperparamètres optimaux, les quantités d’activité totales reconstruites et plusieurs indicateurs statistiques associés à la meilleure estimation. Le nombre d’échantillons pour l’étape d’espérance de la méthode EM est $N = 1000$. Ces échantillons sont obtenus avec un RTO naïf, qui sera introduit dans la section 6.4 et qui n’est qu’une approximation. Nous remarquons que la convergence de l’EM pour r ne nécessite que quelques itérations tandis que la convergence pour b et les b_m n’en nécessite que quelques dizaines tout au plus.

Dans le cas de la CNPP, les statistiques d’erreur d’observation gaussiennes et log-normales conduisent à des résultats similaires et en particulier la même activité totale du terme source inversé (TRRA, pour total retrieved released activity en anglais). Les résultats sont en accord avec l’ordre

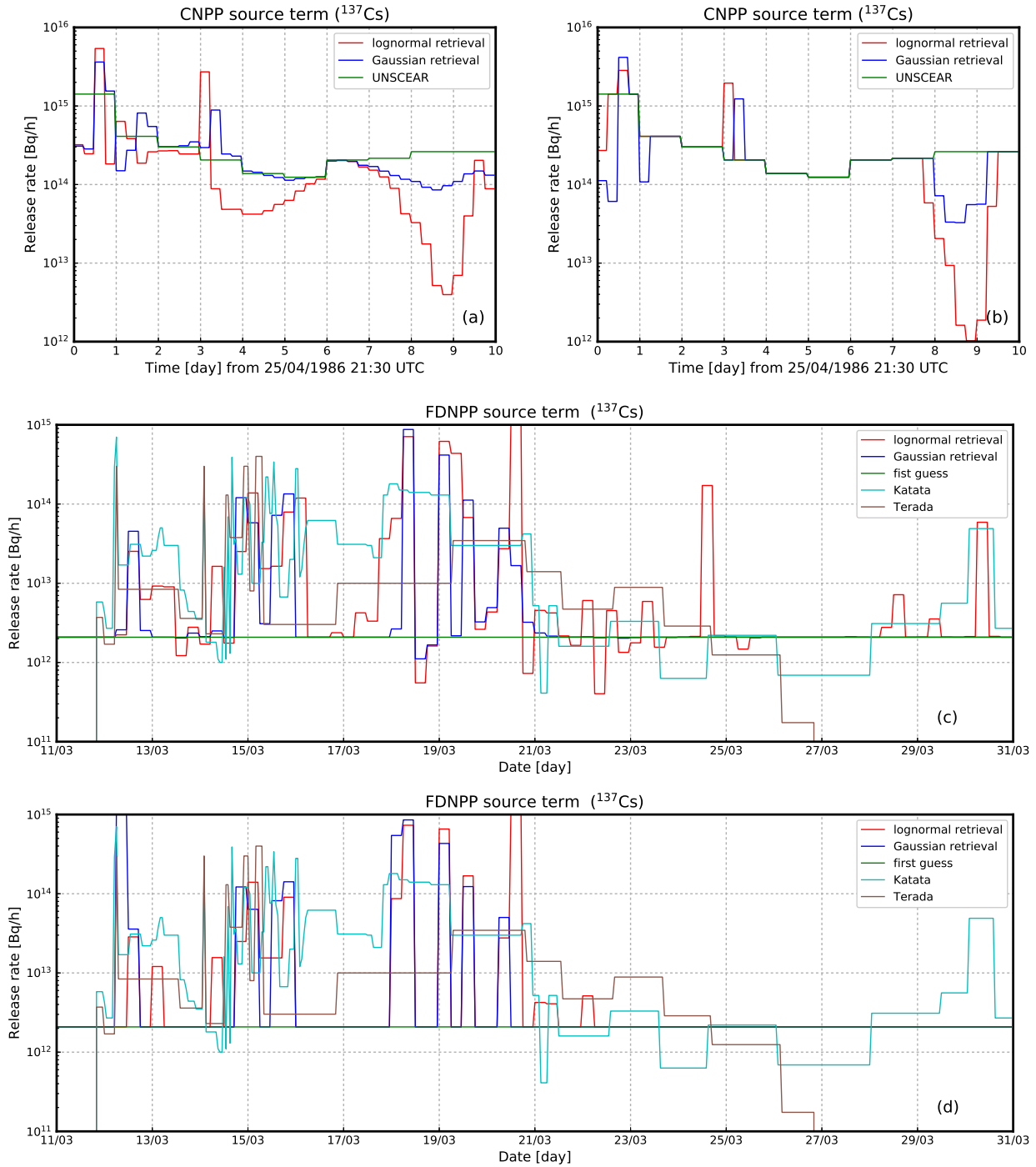


FIGURE 6.1 – Terme source de césium-137 reconstruit pour le cas de l’accident de la CNPP (images du haut, a et b) et pour l’accident de la FDNPP (images du bas, c et d). Les images a et c utilisent la configuration $\mathcal{D}_u$ tandis que les images b et d utilisent la configuration $\mathcal{D}_m$. Les statistiques gaussiennes aussi bien que log-normales sont considérées. Le terme source d’UNSCEAR est également tracé pour le cas de la CNPP. L’ébauche ainsi que les termes sources de Katata et Terada sont également tracés dans le cas de la FDNPP à titre de comparaison. Figure tirée de Liu *et al.* (2017, soumis).

de grandeur 70–90 PBq admis pour le césium-137 (Talerko, 2005) et ceux obtenus par modélisation inverse dans Davoine et Bocquet (2007); Bocquet (2012). Le paramètre estimé r^* nous indique un écart type de l'erreur d'observation additive de $\sqrt{r^*} \approx 1.23 \text{ Bq.m}^{-3}$ (cas gaussien/ $\mathcal{D}_u$) et une erreur relative typique de $\sqrt{r^*} \approx 1.14$ (cas log-normal/ $\mathcal{D}_u$). Nous rappelons que les erreurs d'observation comptabilisent partiellement l'erreur modèle, y compris les erreurs de représentativité. Ainsi, r^* est un indicateur important.

Le cas $\mathcal{D}_m$ choisi pour $\mathbf{B}$ impose de fixer plusieurs des b_m à 0 ou une très faible valeur. Comme on pouvait s'y attendre, cela arrive pour des taux d'émission qui ne sont pas suffisamment informés par les observations. Cela revient à faire une confiance aveugle à l'ébauche sur ces intervalles de temps. L'écart entre le cas gaussien et log-normal est plus petit dans la configuration $\mathcal{D}_m$ que dans la $\mathcal{D}_u$. Cela peut être dû aux degrés de liberté supplémentaires dans les b_m , qui apportent plus de flexibilité pour se rabattre sur l'ébauche durant plusieurs périodes en l'absence d'information suffisante.

Dans le cas FDNPP, les statistiques d'erreur d'observation gaussiennes et log-normales donnent des résultats similaires qualitativement mais quantitativement différents. La TRRA avec l'hypothèse log-normale et la configuration $\mathcal{D}_u$ est deux fois plus importante que celle obtenue avec l'hypothèse gaussienne. Cela est dû à des observations significativement élevées et faibles où les mesures et le modèle sont en désaccord : les deux approche statistiques apportent des réponses différentes puisque nous savons qu'elles traitent différemment ces valeurs élevées et faibles de concentration. Les TRRAs gaussiennes et log-normales sont plus proches l'une de l'autre dans la configuration $\mathcal{D}_m$. Ces résultats sont dans la fourchette 1–3 PBq, en accord avec les estimations passées (Chino *et al.*, 2011; Winiarek *et al.*, 2012; Terada *et al.*, 2012; Saunier *et al.*, 2013; Winiarek *et al.*, 2014; Katata *et al.*, 2015). Le paramètre estimé r^* nous indique un écart type de l'erreur d'observation additive de $\sqrt{r^*} \approx 24 \text{ Bq.m}^{-3}$ (cas gaussien/ $\mathcal{D}_u$) et une erreur relative typique de $\sqrt{r^*} \approx 1.41$ (cas log-normal/ $\mathcal{D}_u$). Les erreurs additives sont nettement plus importantes que dans le cas de la CNPP, car les stations de mesure utilisées sont dans un rayon de 200 km de la FDNPP, alors que dans le cas de la CNPP, elles se situaient à plusieurs milliers de kilomètres de la centrale, conduisant à des concentrations d'activité mesurées plusieurs ordres de grandeur plus faibles. Cependant, les erreurs multiplicatives montrent que, indépendamment de leur valeur absolue, les erreurs d'observation sont plus grandes que dans le cas de la CNPP. Cela trahit les problèmes significatifs d'erreur modèle liés à la modélisation numérique de l'accident de la FDNPP. Les inversions fondées sur les configurations $\mathcal{D}_u$ et $\mathcal{D}_m$ montrent des différences notables entre le 23 et le 31 mars (dernière semaine de la période d'étude). Tandis que la configuration $\mathcal{D}_m$ fixe les taux reconstruits à l'ébauche sur cette période, la configuration $\mathcal{D}_u$ conduit à la reconstruction de plusieurs pics d'émission, certains d'entre eux (25 mars) en cohérence avec Winiarek *et al.* (2014), d'autres (29-31 mars) en accord avec Katata *et al.* (2015). En effet, le b^* optimal dans la configuration $\mathcal{D}_u$ étant plus important et étant applicable à toute la durée de l'inversion, il intervient de façon permissive dans la dernière semaine. Cela contraste avec la configuration $\mathcal{D}_m$ qui fixe les b_m à de petites valeurs sur la dernière semaine, considérant probablement les informations comme trop peu fiables.

Comme on peut le voir sur le tableau 6.2, pour les deux cas de la CNPP et la FDNPP ainsi que les configurations $\mathcal{D}_u$ et $\mathcal{D}_m$, le b^* optimal et la moyenne des b_m^* optimaux du cas log-normal sont toujours nettement supérieurs à ceux du cas gaussien. Cela peut être expliqué par le fait que l'erreur d'observation gaussienne est significativement plus contraignante que les statistiques log-normales, de telle sorte que le terme d'ébauche devrait également l'être (avec des b^* et b_m^* plus petits) s'il existe un équilibre entre les observations et l'ébauche dans la fonction de coût.

Enfin, le tableau 6.2 présente les indicateurs statistiques pour les activités et les simulations des panaches reconstruits en utilisant les hyperparamètres optimaux. Dans le cas de la CNPP, ces

indicateurs sont en parfait accord avec ceux obtenus dans Bocquet (2012, 2015). Les statistiques gaussiennes et log-normales conduisent à des indicateurs statistiques quasiment identiques. Le cas de la FDNPP est différent, les statistiques gaussiennes produisant un grand biais mais une bonne corrélation de Pearson, tandis que les statistiques log-normales conduisent à un biais bien plus équilibré mais une corrélation bien moindre. Cela indique à nouveau les difficultés de la modélisation de la dispersion atmosphérique de l'accident de la FDNPP.

6.4 Quantification des incertitudes

À terme, nous sommes intéressés par l'estimation rigoureuse de la distribution a posteriori $p(\mathbf{x}|\mathbf{y})$ ou de certains de ses moments statistiques. Dans l'approche BC, la quantification des incertitudes est permise par l'échantillonnage hiérarchique d'un ensemble d'hyperparamètres $\boldsymbol{\theta}_n$ et des variables $\mathbf{x}_n$, avec $n = 1, \dots, N$. De ces nombreux échantillons $\mathbf{x}_n, \boldsymbol{\theta}_n$, on peut estimer les moments empiriques de $p(\mathbf{x}|\mathbf{y})$ et quantifier l'incertitude des estimations.

Si l'approche BC est trop coûteuse numériquement, on peut accessoirement recourir à l'approche BE, estimer $\boldsymbol{\theta}^*$ et ainsi réduire la dimension du problème. Il reste nécessaire de savoir échantillonner rigoureusement selon $p(\mathbf{x}|\mathbf{y}, \boldsymbol{\theta}^*)$. Cette tâche est le but de cette section. Nous présentons, testons et comparons quatre méthodes.

6.4.1 Échantillonnage préférentiel fondé sur l'approximation de Laplace

La première méthode s'appuie sur le principe d'échantillonnage préférentiel (par exemple, McKay, 2003). On considère $\mathbf{x} \mapsto \phi(\mathbf{x})$ une fonction de test valide du vecteur de contrôle $\mathbf{x}$ à valeur dans un espace vectoriel et définissons l'espérance a posteriori

$$\xi_\phi \equiv \mathbb{E}_{\mathbf{x}|\mathbf{y}}[\phi] = \int_{\mathbb{R}^M} d\mathbf{x} p(\mathbf{x}|\mathbf{y}) \phi(\mathbf{x}), \quad (6.33)$$

où l'indice $\mathbf{x}|\mathbf{y}$ signifie que l'espérance $\mathbb{E}$ est prise sur la distribution a posteriori de pdf $p(\mathbf{x}|\mathbf{y})$.

Nous introduisons une pdf $q(\mathbf{x})$ que l'on juge plus facile à échantillonner que $p(\mathbf{x}|\mathbf{y})$. Il est souhaitable que $q(\mathbf{x})$ soit similaire à $p(\mathbf{x}|\mathbf{y})$ mais la seule condition mathématique nécessaire est que le support de $q(\mathbf{x})$ inclut celui de $p(\mathbf{x}|\mathbf{y})$. Nous avons alors

$$\xi_\phi = \int_{\mathbb{R}^M} d\mathbf{x} q(\mathbf{x}) \frac{p(\mathbf{x}|\mathbf{y})}{q(\mathbf{x})} \phi(\mathbf{x}). \quad (6.34)$$

Étant donnée cette condition, l'intégrale peut être définie de manière équivalente sur le domaine où $q(\mathbf{x}) \neq 0$. Afin de calculer numériquement ξ_ϕ , l'équation (6.34) suggère de tirer un échantillon $\mathbf{x}_n$ de la pdf préférentielle $q(\mathbf{x})$ et d'estimer

$$\xi_\phi \approx \frac{\sum_{n=1}^N \omega_n \phi(\mathbf{x}_n)}{\sum_{n=1}^N \omega_n} \quad \text{où} \quad \omega_n = \frac{p(\mathbf{x}_n|\mathbf{y})}{q(\mathbf{x}_n)}. \quad (6.35)$$

Nous souhaitons implémenter cette stratégie d'échantillonnage préférentiel pour calculer la distribution a posteriori $p(\mathbf{z}|\mathbf{y}) \propto p(\mathbf{y}|\mathbf{z})p(\mathbf{z}) \propto e^{-\mathcal{L}(\mathbf{z})}$ où $\mathcal{L}(\mathbf{z})$ est définie à l'équation (6.9). Un choix simple de distribution préférentielle pourrait être l'a priori $q(\mathbf{z}) \equiv p(\mathbf{z})$. Les poids de l'équation (6.35) sont alors simplement $\omega_n = p(\mathbf{y}|\mathbf{z}_n)$. Cependant, cet échantillonnage s'est avéré inefficace au-delà d'un terme source de quelques variables pour notre application, car cette distribution est trop éloignée de celle que l'on cherche à estimer. Ainsi, nous avons utilisé une pdf préférentielle

plus sophistiquée, fournie par l'approximation de Laplace de $p(\mathbf{z}|\mathbf{y})$. Il s'agit simplement de la distribution gaussienne définie selon

$$q(\mathbf{z}) \propto e^{-\frac{1}{2}\|\mathbf{z}-\mathbf{z}^*\|_{\mathbf{G}}^{-1}} \quad (6.36)$$

où $\mathbf{z}^* = \operatorname{argmin} \mathcal{L}(\mathbf{z})$ et $[\mathbf{G}]_{ml} = \partial_{z_m} \partial_{z_l} \mathcal{L}(\mathbf{z})|_{\mathbf{z}^*}$. Puisque q est une distribution gaussienne multivariée, un échantillon $\mathbf{z}_n$ peut être tiré selon elle en tirant $\boldsymbol{\eta}_n \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_M)$ et en choisissant

$$\mathbf{z}_n = \mathbf{z}^* + \mathbf{G}^{-\frac{1}{2}} \boldsymbol{\eta}_n. \quad (6.37)$$

Le poids attaché à $\mathbf{z}_n$ est

$$\omega_n = e^{\frac{1}{2}\|\boldsymbol{\eta}_n\|^2 - \mathcal{L}(\mathbf{z}_n)}. \quad (6.38)$$

La Hessienne $\mathbf{G}$ requise par l'échantillonnage peut être calculée analytiquement dans notre problème à partir de l'équation (6.9). De plus, seul un unique MAP $\mathbf{z}^*$ a besoin d'être calculé par minimisation de $\mathcal{L}$. Ainsi les échantillons peuvent être obtenus de manière numériquement efficace, ce qui permet de calculer un grand nombre d'entre eux en parallèle.

Quand le système devient significativement non-gaussien, le MAP et l'inverse de la Hessienne deviennent de piètres estimations de la moyenne et de la matrice de covariance. Dans ce cas, nous proposons d'estimer la moyenne et la matrice de covariance au fur et à mesure que les échantillons sont calculés. Nous proposons de mettre à jour $\bar{\mathbf{z}}_n$ selon

$$\bar{\mathbf{z}}_n = \bar{\mathbf{z}}_{n-1} + \frac{\omega_n}{\sum_{i=1}^n \omega_i} (\mathbf{z}_n - \bar{\mathbf{z}}_{n-1}), \quad (6.39)$$

et la somme $\mathbf{S}_n = \sum_{i=1}^n \omega_i (\mathbf{z}_i - \bar{\mathbf{z}}_n)(\mathbf{z}_i - \bar{\mathbf{z}}_n)^T$ selon

$$\mathbf{S}_n = \mathbf{S}_{n-1} + \omega_n (\mathbf{z}_n - \bar{\mathbf{z}}_n)(\mathbf{z}_n - \bar{\mathbf{z}}_{n-1})^T. \quad (6.40)$$

Un rapide calcul montre que l'estimateur sans biais de la matrice de covariance des erreurs, $\mathbf{C}_n$, est relié à $\mathbf{S}_n$ par

$$\mathbf{C}_n = \frac{\sum_{i=1}^n \omega_i}{(\sum_{i=1}^n \omega_i)^2 - \sum_{i=1}^n \omega_i^2} \mathbf{S}_n, \quad (6.41)$$

pour des poids normalisés ou non. Les deux équations de mise à jour (6.39) et (6.40) sont cohérentes dans la mesure où $\bar{\mathbf{z}}_n$ converge vers la moyenne et $\mathbf{C}_n$ converge vers la matrice de covariance. De plus, ces deux formules ont été développées afin d'éviter les effets d'annulation catastrophique (Welford, 1962) qui risquent fortement d'arriver avec des formules naïves de mise à jour.

En notant $\mathbf{L}_n$ la factorisation de Cholesky de $\mathbf{C}_n$, c'est-à-dire, $\mathbf{C}_n = \mathbf{L}_n \mathbf{L}_n^T$, l'équation (6.37) devient

$$\mathbf{z}_n = \bar{\mathbf{z}}_{n-1} + \mathbf{L}_{n-1} \boldsymbol{\eta}_n. \quad (6.42)$$

En pratique, le MAP et la Hessienne sont toujours utilisés pour initialiser la récursion en tant qu'approximations grossières de la moyenne et de la matrice de covariance. Nous avons constaté que cette méthode adaptative était supérieure à son équivalent non-adaptatif en situation modérément non-gaussienne. De toute évidence, pour que la méthode soit efficace, la pdf a posteriori se doit de ne pas être trop éloignée d'une gaussienne. Du fait de l'étape séquentielle de mise à jour, les échantillons sont obtenus séquentiellement. Néanmoins, en pratique, cela n'est pas trop contraignant puisque plusieurs chaînes séquentielles peuvent être engendrées en parallèle de façon indépendante.

6.4.2 Échantillonnage par *randomize-then-optimize* naïf et sans biais

Nous considérons ici deux autres méthodes non-linéaires pour échantillonner selon la distribution a posteriori de la source. Un problème inverse avec des statistiques d'a priori gaussiennes et un modèle physique non-linéaire a une pdf a posteriori $p(\mathbf{z}|\mathbf{y}) \propto e^{-J(\mathbf{z})}$ avec

$$J(\mathbf{z}) = \frac{1}{2} \|\mathbf{y} - h(\mathbf{z})\|_{\mathbf{R}}^2 + \frac{1}{2} \|\mathbf{z} - \mathbf{z}_b\|_{\mathbf{B}}^2 \quad (6.43)$$

où h est le modèle non-linéaire de l'espace des états à l'espace des observations. Notre problème, fondé sur des statistiques d'erreur log-normales et un modèle linéaire d'après l'équation (6.9), peut aisément être transformé avec le formalisme en effectuant les substitutions : $\mathbf{0} \leftarrow \mathbf{z}_b$, $\mathbf{y} \leftarrow \ln \mathbf{y}$, et $h(\mathbf{z}) \leftarrow \ln \mathbf{H}\bar{\mathbf{x}}e^{\mathbf{z}}$. Définissons

$$\mathbf{w} = \begin{pmatrix} \mathbf{R}^{-\frac{1}{2}}\mathbf{y} \\ \mathbf{B}^{-\frac{1}{2}}\mathbf{z}_b \end{pmatrix} \quad \text{et} \quad f(\mathbf{z}) = \begin{pmatrix} \mathbf{R}^{-\frac{1}{2}}h(\mathbf{z}) \\ \mathbf{B}^{-\frac{1}{2}}\mathbf{z} \end{pmatrix}. \quad (6.44)$$

$\mathbf{w}$ et $f(\mathbf{z})$ sont tous deux des vecteurs de $\mathbb{R}^{M+P}$. Avec ces définitions, J peut être condensée en $J(\mathbf{z}) = \frac{1}{2} \|\mathbf{w} - f(\mathbf{z})\|^2$.

Une méthode intuitive pour échantillonner selon $p(\mathbf{z}|\mathbf{w})$ est de perturber $\mathbf{w}$ selon ses statistiques gaussiennes a priori $\mathcal{N}(\mathbf{0}, \mathbf{I}_{M+P})$ et de minimiser $J(\mathbf{z})$ pour obtenir un échantillon. Autant d'échantillons que nécessaires peuvent être ainsi obtenus en parallèle. Cette méthode produit un échantillonnage sans biais de la distribution a posteriori sous réserve que f soit linéaire. Cependant, comme il a été remarqué par Bardsley *et al.* (2014), cet échantillonnage n'est pas exact dans le cas général où $\mathbf{z} \mapsto f(\mathbf{z})$ est non-linéaire, même s'il peut s'agir d'une bonne approximation. Cela n'empêche pas cette approche naïve d'être au cœur des méthodes d'assimilation de données connues sous le nom d'ensemble of data assimilation (EDA), qui sont utilisées opérationnellement dans les centres de modélisation numérique du temps (par exemple, Raynaud *et al.*, 2011; Bonavita *et al.*, 2012, pour Météo-France et le Centre européen pour les prévisions météorologiques à moyen terme (CEPMMT)). Elle a également été utilisée par Winiarek *et al.* (2014) pour estimer l'incertitude a posteriori de l'inversion du terme source de césium-137 de la FDNPP. Cette démarche approchée d'échantillonnage non-linéaire sera nommée RTO naïf par la suite.

Bardsley *et al.* (2014) offrent un schéma original qui imite cette approche naïve mais effectue un échantillonnage non-linéaire sans biais de la distribution a posteriori $p(\mathbf{z}|\mathbf{w})$. Sans biais signifie qu'un échantillon produit par cette méthode est véritablement tiré selon la pdf a posteriori. Elle s'appuie sur les méthodes de maximum de vraisemblance aléatoire de Oliver *et al.* (1996), qui furent les premiers à proposer un échantillonnage par RTO, potentiellement suivi d'une étape de correction (en utilisant un critère de Metropolis-Hastings) dans le cas non-linéaire.

Définissons le MAP du problème : $\mathbf{z}^* = \operatorname{argmin} J(\mathbf{z})$. Soit $\mathbf{f}_{\mathbf{z}^*} \in \mathbb{R}^{(M+P) \times M}$ le tangent linéaire de f au point $\mathbf{z}^*$. Nous effectuons une décomposition QR de $\mathbf{f}_{\mathbf{z}^*}$: $\mathbf{f}_{\mathbf{z}^*} = \mathbf{Q}\mathbf{U}$ où $\mathbf{Q}$ est une matrice orthogonale de $\mathbb{R}^{(M+P) \times M}$ et $\mathbf{U}$ est une matrice triangulaire supérieure de $\mathbb{R}^{M \times M}$. Soit $\tilde{\mathbf{Q}}$ une matrice orthogonale dans $\mathbb{R}^{(M+P) \times P}$ qui complète la base orthonormale formée par les colonnes de $\mathbf{Q}$. Pour obtenir un échantillon $\mathbf{z}_n$ de la distribution a posteriori, on échantillonne tout d'abord $\boldsymbol{\varepsilon}_n \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_{M+P})$ et calculons

$$\mathbf{z}_n = \operatorname{argmin}_{\mathbf{z}} \|\mathbf{Q}^T(\mathbf{w} + \boldsymbol{\varepsilon}_n - f(\mathbf{z}))\|. \quad (6.45)$$

Cependant, il y a un poids attaché à cet échantillon,

$$\omega_n = \frac{e^{-\frac{1}{2}\|\tilde{\mathbf{Q}}^T(\mathbf{w} - f(\mathbf{z}_n))\|^2}}{|\mathbf{Q}^T \mathbf{f}_{\mathbf{z}^*}|}, \quad (6.46)$$

qui doit être calculé pour obtenir un ensemble sans biais d'échantillons $\{\mathbf{z}_n, \omega_n\}_{1 \leq n \leq N}$. Bardsley *et al.* (2014) ont nommé ce schéma d'échantillonnage *randomize-then-optimize* (RTO). Nous nous référons à leur article pour une démonstration complète du schéma, ou à l'appendice B de Liu *et al.* (2017, soumis) pour une démonstration heuristique mais compacte.

Les deux méthodes d'échantillonnage RTO, naïf et sans biais, sont coûteuses car nécessitent autant de minimisations que d'échantillons, même si celles-ci sont indépendantes et peuvent être effectuées en parallèle. Contrairement au RTO naïf, le RTO sans biais peut être impacté par une variance significative des poids, qui provoque des erreurs d'échantillonnage et invite donc à tirer plus d'échantillons. Cependant, les deux RTO ont besoin de moins d'échantillons que la méthode de Laplace présentée dans la section précédente, car chacun des échantillons produit est plus pertinent.

6.4.3 Méthode basique de Monte-Carlo par chaînes de Markov

Des échantillons de la pdf équation (6.7) peuvent être engendrés numériquement par les techniques versatiles de Monte-Carlo par chaînes de Markov (MCMC, Metropolis *et al.*, 1953; Hastings, 1970; Green, 1995). Le MCMC fournit une séquence d'échantillons dans l'espace de l'ensemble dont la distribution invariante est la pdf que nous souhaitons échantillonner, à savoir $p(\mathbf{x}|\mathbf{y})$. La transition de cette chaîne, ou déplacement, consiste à tirer $\mathbf{x}'$ selon $q(\mathbf{x}'|\mathbf{x})$, connaissant la valeur actuelle $\mathbf{x}$. Par conséquent, cette séquence est une chaîne de Markov, puisque la pdf de transition q pour $\mathbf{x}'$ ne dépend que de $\mathbf{x}$. Un tel déplacement $\mathbf{x}'$ est accepté si et seulement si

$$u \leq \frac{p(\mathbf{x}'|\mathbf{y}) q(\mathbf{x}|\mathbf{x}')}{p(\mathbf{x}|\mathbf{y}) q(\mathbf{x}'|\mathbf{x})}, \quad (6.47)$$

où u est tiré selon $\mathcal{U}[0, 1]$, la distribution uniforme sur $[0, 1]$; sinon $\mathbf{x}$ est inchangé. Cette règle de sélection (6.47), de type Metropolis-Hastings, assure que la chaîne engendre des échantillons de la distribution invariante, c'est-à-dire $p(\mathbf{x}|\mathbf{y})$. De la chaîne $\{\mathbf{x}_n\}_{n=1, \dots, N}$, on peut estimer les moments de la pdf a posteriori $p(\mathbf{x}|\mathbf{y})$. Dans l'équation (6.47), le terme $p(\mathbf{x}'|\mathbf{y})/p(\mathbf{x}|\mathbf{y})$ peut de manière équivalente être remplacé par $\{p(\mathbf{y}|\mathbf{x}')p(\mathbf{x}')\} / \{p(\mathbf{y}|\mathbf{x})p(\mathbf{x})\}$, grâce à l'annulation des constantes de normalisation. Ce ratio peut être calculé numériquement. La convergence de la chaîne est garantie sous des hypothèses faibles, telles l'irréductibilité de la chaîne de Markov. On entend par convergence le fait que les échantillons produits fournissent une approximation de Monte-Carlo fiable de la pdf $p(\mathbf{x}|\mathbf{y})$. Cependant, cette convergence peut être longue, tout particulièrement lorsque la dimension de l'espace des états est grande et que le déplacement est inefficace. La transition que l'on utilise est l'application d'un bruit gaussien à un taux d'émission choisi au hasard.

Cette technique peut être appliquée à l'approche BE, et donc à l'estimation de $p(\mathbf{x}|\mathbf{y}, \boldsymbol{\theta}^*)$, qui représente un échantillonnage approché de la pdf $p(\mathbf{x}, \boldsymbol{\theta}|\mathbf{y}) \approx p(\mathbf{x}|\mathbf{y}, \boldsymbol{\theta}^*)\delta(\boldsymbol{\theta} - \boldsymbol{\theta}^*)$.

6.4.4 Application aux accidents de Tchernobyl et Fukushima Daiichi

Dans ce second jeu d'illustrations, nous appliquons ces méthodes d'UQ, l'échantillonnage préférentiel de Laplace, le RTO naïf, le RTO sans biais et le MCMC, à l'estimation de la pdf de la TRRAS de la CNPP et la FDNPP. Nous utilisons uniquement des statistiques d'erreur log-normales. Les pdfs sont tracées sur la Fig. 6.2.

L'approximation de Laplace utilise l'estimation adaptative de la covariance, qui améliore considérablement la convergence dans ce cas. Le RTO sans biais utilise la variante Metropolis-Hastings décrite dans Bardsley *et al.* (2014), qui incorpore une étape de Metropolis-Hastings pour éviter de calculer le poids des échantillons, même si nous n'avons pas remarqué d'amélioration en termes de

convergence, comparée à la variante originelle. Pour atteindre un compromis entre précision et coût de calcul, nous choisissons $M = 40$ taux d'émissions pour la CNPP, chacun d'entre eux d'une durée de 6 h. Avec ces paramètres, la configuration $\mathcal{D}_m$ de l'EM donne $r^* = 1.35$ et $b^* = 2.26$ en moyenne sur les 40 b_m^* . Les courbes issues des échantillonnages de Laplace et MCMC sont confondues, comme on pouvait l'espérer, étant donné qu'il s'agit de méthodes exactes d'échantillonnage (pour un jeu fixé d'hyperparamètres). Le RTO naïf donne une distribution différente qui est clairement biaisée en direction de plus grosses TRRAs. Le RTO sans biais effectue un échantillonnage exact du problème BE. Cependant, un léger écart situé au niveau du MAP de la distribution est observé, par rapport à la référence fournie par les approches de Laplace et MCMC. Le RTO sans biais agit comme s'il avait du mal à ajuster les échantillons proches du MAP, qui devraient être les mêmes que ceux du RTO naïf, mais avec des poids. La méthode de Laplace utilise 2.4×10^8 échantillons ; le MCMC utilise 2.4×10^7 échantillons ; les RTO naïf et sans biais produisent 4.8×10^6 échantillons.

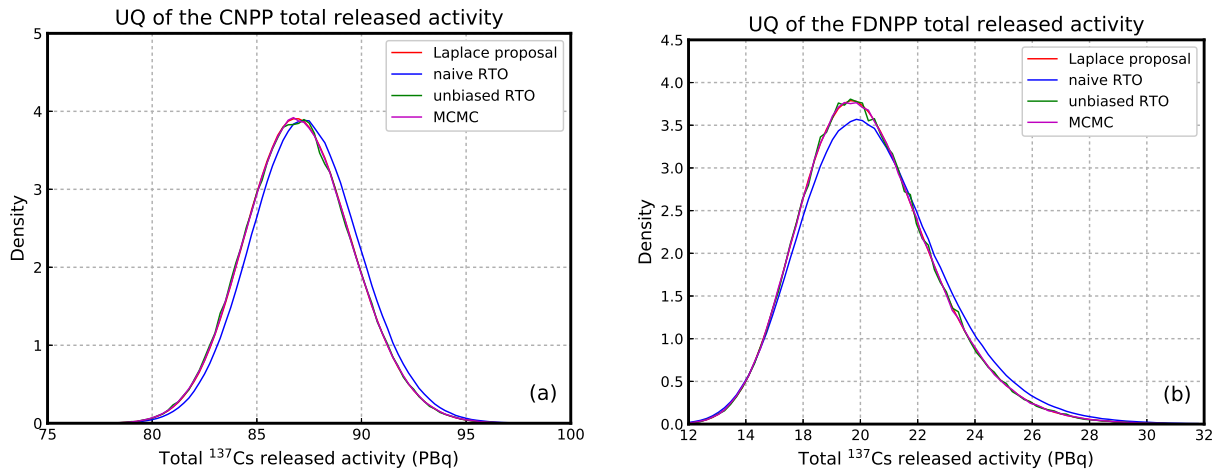


FIGURE 6.2 – Fonctions de densité de probabilité des TRRAs de la CNPP (à gauche) et de la FDNPP (à droite). L'échantillonnage de Laplace, le RTO naïf, le RTO sans biais et le MCMC sont comparés. Les courbes des solutions par les méthodes de Laplace, du RTO sans biais et du MCMC sont presque confondues. Figure tirée de [Liu et al. \(2017, soumis\)](#).

Dans le cas de la FDNPP, nous choisissons $M = 80$ taux d'émission, chacun d'entre eux d'une durée de 6 h. Avec ces paramètres, la configuration $\mathcal{D}_m$ de l'EM donne $r^* = 2.05$ et $b^* = 2.64$ en moyenne sur les 80 b_m^* . Les résultats obtenus sont similaires à ceux de la CNPP, même si les distributions sont apparemment moins gaussiennes, avec une queue de distribution plus épaisse en direction des fortes valeurs de TRRA. Le MAP est à 20.2 PBq, la moyenne à 20.9 PBq tandis que la médiane est à 20.7 PBq. Ces différences prouvent la présence de cette queue de distribution. Comme dans le cas de la CNPP, le RTO sans biais n'a pas parfaitement convergé et présente un bruit d'échantillonnage résiduel.

6.5 Hiérarchie bayésienne complète

Dans l'approche BE, l'échantillonnage est exact une fois que θ^* est déterminé. Dans cette section, nous considérons cette fois l'approche bayésienne complète et l'échantillonnage exact de $p(\mathbf{x}, \theta | \mathbf{y})$. Les résultats peuvent servir de référence pour se comparer à la section 6.4. En effet, nous avons vu que si $p(\theta | \mathbf{y})$ est piquée autour de θ^* , la marginalisation de $p(\mathbf{x}, \theta | \mathbf{y})$ sur θ est approchée

par $p(\mathbf{x}|\boldsymbol{\theta}^*, \mathbf{y})$.

6.5.1 Méthode hiérarchique de Monte-Carlo par chaînes de Markov

Comme il a été vu dans la section 6.3.3, il est possible d'échantillonner selon la distribution a posteriori et prendre en compte les incertitudes sur $\boldsymbol{\theta}$ par la hiérarchie présentée à l'équation (6.12). Les échantillons issus d'une telle hiérarchie peuvent être engendrés avec un MCMC dont la distribution invariante est la pdf conditionnelle $p(\mathbf{x}, \boldsymbol{\theta}|\mathbf{y})$. La transition pour cette chaîne consiste à échantillonner $\mathbf{x}', \boldsymbol{\theta}'$ de $q(\mathbf{x}', \boldsymbol{\theta}'|\mathbf{x}, \boldsymbol{\theta})$, connaissant $\mathbf{x}, \boldsymbol{\theta}$. Une telle transition $\mathbf{x}', \boldsymbol{\theta}'$ est acceptée si et seulement si

$$u \leq \frac{p(\mathbf{x}', \boldsymbol{\theta}'|\mathbf{y}) q(\mathbf{x}, \boldsymbol{\theta}|\mathbf{x}', \boldsymbol{\theta}')}{p(\mathbf{x}, \boldsymbol{\theta}|\mathbf{y}) q(\mathbf{x}', \boldsymbol{\theta}'|\mathbf{x}, \boldsymbol{\theta})}, \quad (6.48)$$

où u est tiré selon $\mathcal{U}[0, 1]$; sinon, $\mathbf{x}, \boldsymbol{\theta}$ sont conservés. Dans l'équation (6.48), le terme $p(\mathbf{x}', \boldsymbol{\theta}'|\mathbf{y})/p(\mathbf{x}, \boldsymbol{\theta}|\mathbf{y})$ peut de manière équivalente être remplacé par le ratio de pdfs hiérarchiques

$$\{p(\mathbf{y}|\mathbf{x}', \boldsymbol{\theta}')p(\mathbf{x}'|\boldsymbol{\theta}')p(\boldsymbol{\theta}')\} / \{p(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta})p(\mathbf{x}|\boldsymbol{\theta})p(\boldsymbol{\theta})\}, \quad (6.49)$$

grâce à l'annulation des constantes de normalisation. Ce ratio peut être calculé numériquement. Ainsi, un échantillonnage par méthode MCMC existe aussi bien pour l'approche BE que BC.

Les transitions que l'on considère pour notre application sont : (i) l'application d'un bruit gaussien à un taux d'émission choisi au hasard, (ii) l'application d'un bruit gaussien à un hyperparamètre choisi au hasard. La variance de ces distributions gaussiennes est prise de telle sorte que le taux d'acceptation soit autour d'1/3.

6.5.2 Analyse transdimensionnelle

Le but principal d'une analyse transdimensionnelle (Sambridge *et al.*, 2013, et les références qui y figurent) est de déterminer le maillage le plus pertinent pour les variables de contrôle utilisées pour effectuer l'inversion. Ainsi, il s'agit d'un outil adaptatif permettant de réduire le temps de calcul, particulièrement lors du recours à un MCMC pour échantillonner selon la distribution a posteriori. Plus précisément, le maillage est lui-même considéré comme une variable faisant partie de l'inversion. Cela revient à étendre la hiérarchie de l'équation (6.12). Si $\mathbf{g}$ est le vecteur des paramètres décrivant le maillage permettant la discrétisation du problème, à savoir la position des points de grille et leur nombre M , alors la nouvelle hiérarchie bayésienne s'écrit

$$\begin{aligned} p(\mathbf{x}, \boldsymbol{\theta}, \mathbf{g}|\mathbf{y}) &= \frac{p(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta}, \mathbf{g})p(\mathbf{x}, \boldsymbol{\theta}, \mathbf{g})}{p(\mathbf{y})} \\ &= \frac{p(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta}, \mathbf{g})p(\mathbf{x}|\boldsymbol{\theta}, \mathbf{g})p(\boldsymbol{\theta}, \mathbf{g})}{p(\mathbf{y})} \\ &= \frac{p(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta}, \mathbf{g})p(\mathbf{x}|\boldsymbol{\theta}, \mathbf{g})p(\boldsymbol{\theta}|\mathbf{g})p(\mathbf{g})}{p(\mathbf{y})}, \end{aligned} \quad (6.50)$$

où $p(\boldsymbol{\theta}|\mathbf{g})$ est la distribution a priori des hyperparamètres connaissant le maillage et $p(\mathbf{g})$ est l'a priori pour ce maillage.

Le recours à une méthode de MCMC pour échantillonner selon $p(\mathbf{x}, \boldsymbol{\theta}, \mathbf{g}|\mathbf{y})$ nécessite des modifications mentionnées dans Green (1995, 2003), pour créer ce qu'on appelle un MCMC avec saut réversible. Il est toujours possible d'utiliser une règle de Metropolis-Hastings comme dans l'équation (6.48), mais il faut l'amender d'un facteur correctif, une Jacobienne $\mathbf{J}$ qui tient compte du

changement de dimension des espaces des états. La règle devient donc

$$u \leq \frac{p(\mathbf{x}', \boldsymbol{\theta}', \mathbf{g}' | \mathbf{y})}{p(\mathbf{x}, \boldsymbol{\theta}, \mathbf{g} | \mathbf{y})} \frac{q(\mathbf{x}, \boldsymbol{\theta}, \mathbf{g} | \mathbf{x}', \boldsymbol{\theta}', \mathbf{g}')}{q(\mathbf{x}', \boldsymbol{\theta}', \mathbf{g}' | \mathbf{x}, \boldsymbol{\theta}, \mathbf{g})} |\mathbf{J}|. \quad (6.51)$$

Cette Jacobienne peut être calculée en utilisant une extension stochastique du déplacement transdimensionnel pour en faire un difféomorphisme avec une Jacobienne correctement définie. Les détails théoriques peuvent être trouvés dans [Green \(2003\)](#) et un exemple clair est présenté dans [Bodin et Sambridge \(2009\)](#). Nous notons également que dans le contexte de modélisation inverse d'un rejet accidentel de polluants dans l'atmosphère, [Yee \(2008\)](#) a utilisé un MCMC avec saut réversible afin d'estimer plusieurs sources ponctuelles. Dans le contexte de la modélisation inverse des émissions de gaz à effet de serre, cette méthode a été utilisée par [Lunt et al. \(2016\)](#) pour l'estimation d'un champ bidimensionnel d'émission de méthane au Nord-Ouest de l'Europe.

Dans le cas de nos illustrations, nous considérons un maillage adaptatif avec $M + 1$ nœuds, deux d'entre eux étant respectivement les dates de début et de fin de la période d'émission. Cela correspond donc à M taux d'émission définissant chacun une émission uniforme entre deux nœuds. Cela coïncide avec un pavage de Voronoi dont l'implémentation est aisée en comparaison du cas bidimensionnel de [Lunt et al. \(2016\)](#). Les déplacements que l'on considère sont : les déplacements (i) et (ii) de la section 6.5.1, (iii) le déplacement aléatoire sur une distance gaussienne d'un nœud choisi aléatoirement. Cela est suffisant si on choisit de fixer M . Dans ce cas précis, la Jacobienne est $|\mathbf{J}| = 1$ car le nombre de nœuds est inchangé. Si M est variable, on peut alors également considérer (iv) le processus de naissance d'un nœud, (v) le processus de mort d'un nœud. Il est possible, comme dans l'appendice B de [Bodin et Sambridge \(2009\)](#), de construire ces processus de telle sorte que la Jacobienne reste 1, ce qui n'est en général pas le cas pour ces deux derniers processus.

6.5.3 Application aux accidents de Tchernobyl et Fukushima Daiichi

Dans ce troisième jeu d'illustrations, une approche BC est implémentée, en utilisant des statistiques d'erreur log-normales. Cette approche se fonde sur une analyse transdimensionnelle et utilise la hiérarchie sur les hyperparamètres et le maillage. Par simplicité, nous choisissons un maillage adaptatif avec un nombre fixe de points de grille, que l'on limite à $M = 20$ pour la CNPP et $M = 40$ pour la FDNPP, ce qui correspond à la moitié des expériences précédentes. En particulier, cela implique que la Jacobienne de l'équation (6.51) est 1. La solution du problème est obtenue numériquement par un MCMC. Nous nous concentrons sur l'hypothèse $\mathcal{D}_m$. La configuration plus grossière $\mathcal{D}_u$ donnerait des distributions encore plus larges pour les TRRAs.

Dans les deux cas d'étude, nous récupérons, de la simulation MCMC, la moyenne et la médiane des activités reconstruites, ainsi que la moyenne plus ou moins l'écart type. Cela quantifie partiellement l'incertitude liée au terme source. Nous récupérons également la pdf des hyperparamètres, tout particulièrement celle de r , ainsi que la distribution de la TRRA. Ces deux pdfs sont des marginalisations qui peuvent être estimées de la chaîne d'échantillons.

Toutes ces estimations pour le terme source de la CNPP sont tracées sur la Fig. 6.3a. On y voit de nettes différences entre la moyenne et la médiane, ce qui prouve que la distribution a posteriori est non-gaussienne. La moyenne, la médiane et le MAP de la TRRA sont respectivement 100.2 PBq, 95.3 PBq, et 89.0 PBq. Malgré les $M = 20$ points de grille, le terme source paraît continu par morceaux car le maillage est adaptatif, permettant des intervalles de temps localement aussi petits qu'une heure. Les discontinuités du terme source reconstruit proviennent des discontinuités de l'ébauche d'UNSCEAR utilisée. Les écarts types au-dessus et en dessous de la moyenne montrent les périodes plus ou moins fiables du terme source. Tandis que la borne supérieure contraint bien

la solution, la borne inférieure est très faible pour certains évènements (impact du feu du troisième jour ou à la fin du terme source).

La distribution de la TRRA, tracée sur la Fig. 6.3b a une queue nettement plus large que dans l'approche BE. De toute évidence, cela est dû à l'approche hiérarchique du BC où les variations des hyperparamètres trahissent le fait que l'a priori de l'ébauche est peu informatif. Cela souligne le risque d'obtenir une TRRA plus grande que ce que l'approche BE prédirait.

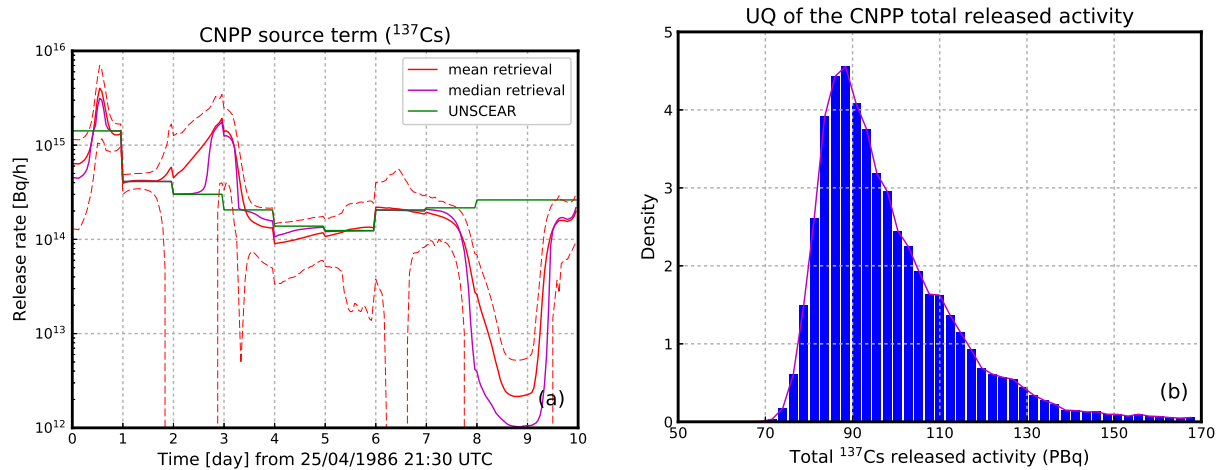


FIGURE 6.3 – À gauche : terme source de césium-137 pour l'accident de la CNPP reconstruit en utilisant la hiérarchie transdimensionnelle du BC sur les hyperparamètres et le maillage (avec $M = 20$ fixé). Les lignes en pointillés indiquent la moyenne de la source plus ou moins l'écart type. À droite : pdf de la TRRA de césium-137. Figure tirée de Liu *et al.* (2017, soumis).

La Fig. 6.4a montre la pdf de l'hyperparamètre r qui quantifie la matrice des erreurs d'observation. La distribution est approximativement gaussienne et centrée autour de $r^* = 1.35$. Cette valeur coïncide avec celle obtenue en appliquant l'approche BE avec une résolution fixe de $M = 40$ mailles uniformément espacées (voir tableau 6.2). La convergence rapide de l'algorithme EM pour estimer r^* est cohérente avec la distribution piquée observée avec l'approche BC.

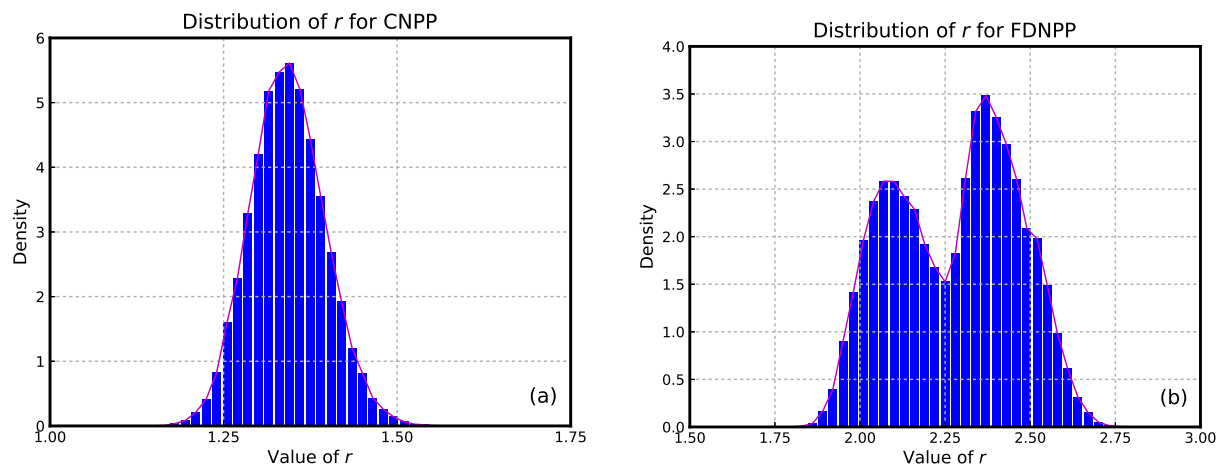


FIGURE 6.4 – Distribution de la variance des erreurs d'observation (hyperparamètre r) pour le cas de la CNPP (à gauche) et de la FDNPP (à droite). Figure tirée de Liu *et al.* (2017, soumis).

Les résultats principaux, à savoir la moyenne et la médiane du terme source estimé pour la FDNPP, sont tracés sur la Fig. 6.5. Les distributions sont significativement moins gaussiennes que dans le cas de la CNPP. En effet, la médiane et la moyenne du terme source ne sont pas nécessairement proches. Elles sont proches ou ont des variations similaires pour les taux d'émission qui sont bien informés par les observations et quand les termes sources de Terada et Katata sont en accord. Au contraire, dans la période après le 23 mars, la médiane s'effondre sur l'ébauche tandis que la moyenne prédit des taux d'émission non négligeables, mais avec une très forte incertitude attachée.

La distribution de la TRRA est multimodale, comme le montre la Fig. 6.6. Elle est fortement impactée par la présence ou l'absence d'émissions après le 23 mars. La moyenne, la médiane et le MAP de la TRRA sont respectivement 30.3 PBq, 17.0 PBq, et 7.9 PBq. Le MAP secondaire se situe à 19.4 PBq. Cette multimodalité est également trahie par la distribution de l'hyperparamètre r , comme le montre la Fig. 6.4b. Cela est d'autant plus remarquable que l'estimation de r est plutôt robuste, avec seulement deux itérations de l'EM requises pour atteindre la convergence, comme mentionné auparavant. Le premier maximum se situe à $r^* = 2.05$, qui coïncide avec la valeur obtenue en appliquant l'approche BE avec une résolution fixe de $M = 80$ mailles uniformément espacées (voir tableau 6.2). Le deuxième maximum est à $r^* = 2.38$, qui peut être obtenu avec l'EM et une approche BE mais avec une résolution fixe de $M = 40$ mailles uniformément espacées. En règle générale, r^* décroît quand M augmente, car l'erreur de représentativité diminue. Pour un maillage plus grossier et un r^* plus important, la TRRA se situe autour de 10 PBq, tandis que pour un maillage plus fin et un r^* plus petit, la TRRA est autour de 20 PBq. Cela est en accord avec la structure bimodale de la pdf de la TRRA. Cela peut probablement s'expliquer par des flux de courte durée permettant d'expliquer des observations d'activité élevées, mais qui ne sont pas retenus car la résolution grossière va provoquer un désaccord avec d'autres observations. Cette multimodalité démontre également que dans certains cas complexes, l'approche BE devrait être utilisée avec précaution, car elle risque de ne capturer qu'un seul mode.

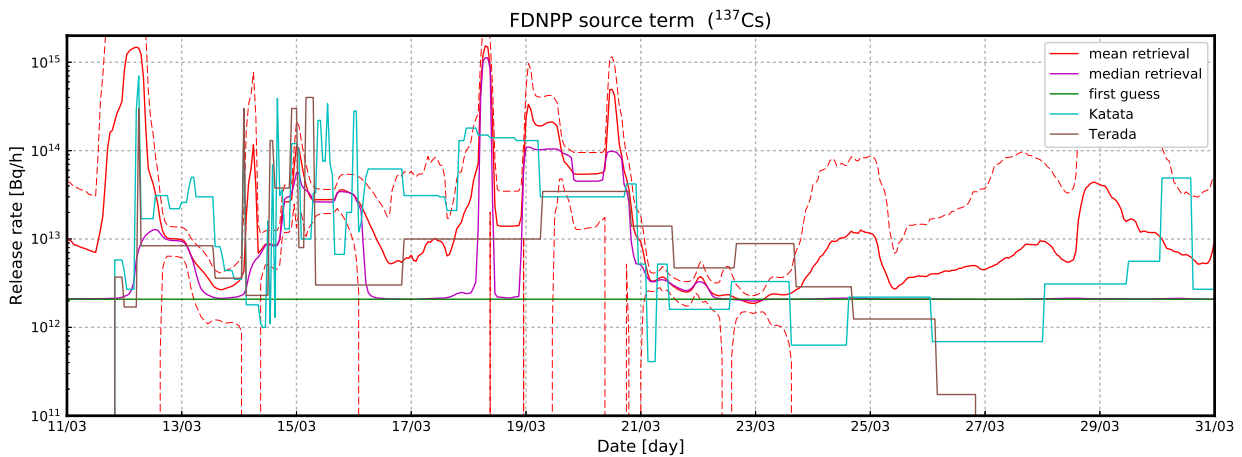


FIGURE 6.5 – Terme source de césium-137 pour l'accident de la FDNPP reconstruit en utilisant la hiérarchie transdimensionnelle du BC sur les hyperparamètres et le maillage (avec $M = 40$ fixé). Les lignes en pointillés indiquent la moyenne de la source plus ou moins l'écart type. Figure tirée de Liu *et al.* (2017, soumis).

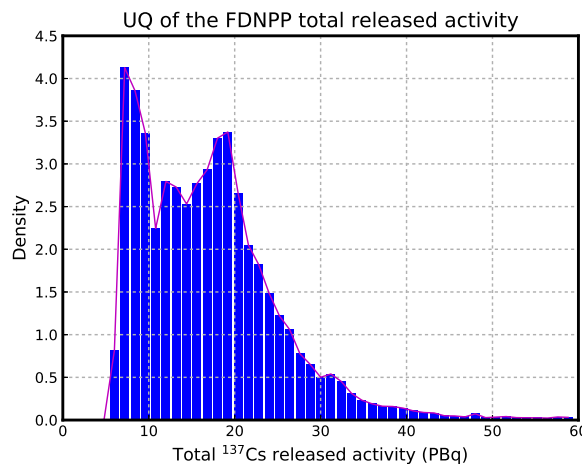


FIGURE 6.6 – Pdf de la TRRA de césium-137 pour l’accident de la FDNPP en utilisant la hiérarchie transdimensionnelle du BC sur les hyperparamètres et le maillage (avec $M = 40$ fixé). Figure tirée de Liu *et al.* (2017, soumis).

6.6 Résumé et conclusions

Dans ce chapitre, nous avons étudié la quantification d’incertitude (UQ) liée à une estimation, ayant pour objet la reconstruction du terme source d’un rejet dans l’atmosphère de polluants dangereux. Même si la relation source-récepteur était linéaire dans notre cas, les statistiques des erreurs sont fondamentalement non-gaussiennes car les taux d’émission et les observations sont des variables positives. Ainsi, il s’agit d’un problème d’inversion non-linéaire et d’échantillonnage non-gaussien.

Nous avons décrit et implémenté deux voies pour le problème de l’UQ. Premièrement, le problème peut être approché par le recours à une démarche de type Bayes empirique (BE), qui détermine la valeur optimale des hyperparamètres θ^* , avant d’échantillonner $p(\mathbf{x}|\mathbf{y}, \theta^*)$ pour évaluer l’incertitude sur $\mathbf{x}$. Deuxièmement, le problème peut être résolu numériquement sans approximation en échantillonnant la hiérarchie complète de $p(\mathbf{x}, \theta|\mathbf{y})$, voire en ajoutant le maillage du terme source aux variables de contrôle et en échantillonnant $p(\mathbf{x}, \theta, \mathbf{g}|\mathbf{y})$. Dans l’approche BE, les hyperparamètres ont été obtenus avec l’algorithme espérance-maximisation (EM) pour des statistiques gaussiennes et log-normales pour les erreurs d’observation. Nous avons échantillonné $p(\mathbf{x}|\mathbf{y}, \theta^*)$ en utilisant une méthode approchée, le *randomize-then-optimize* (RTO) naïf, et trois méthodes exactes : un échantillonnage préférentiel fondé sur l’approximation de Laplace, le RTO sans biais et une méthode de Monte-Carlo par chaînes de Markov (MCMC). Dans l’approche hiérarchique bayésienne complète (BC), un MCMC est utilisé avec un maillage adaptatif, ou analyse transdimensionnelle, pour accélérer la convergence. La plupart des méthodes utilisées ne l’ont jamais été dans ce contexte auparavant. Certaines d’entre elles ont été étendues ou améliorées dans cette étude (ajout d’un a priori sur les hyperparamètres estimés avec l’EM ; utilisation de l’EM pour des statistiques d’erreur log-normales ; échantillonnage préférentiel de Laplace adaptatif).

Ces méthodes ont été testées et comparées via les applications à l’estimation des rejets de césium-137 lors des accidents des centrales nucléaires de Tchernobyl et Fukushima Daiichi (CNPP et FDNPP), en utilisant des observations réelles de concentration de radioactivité. L’étude de l’UQ a fourni des fourchettes de 70 – 130 PBq pour l’accident de la CNPP et 7 – 35 PBq pour l’accident

de la FDNPP, en accord avec la littérature récente sur le sujet. Ces résultats constituent la première analyse objective de l'UQ des termes sources des accidents de la CNPP et de la FDNPP.

Nous avons confirmé que le RTO naïf était biaisé et qu'il peut s'avérer très différent de la véritable distribution des erreurs, tout particulièrement dans un cas non-gaussien comme celui de la FDNPP. Cette conclusion peut avoir un impact quant à l'utilisation du RTO naïf en assimilation de données pour les géosciences, tout particulièrement pour l'ensemble of data assimilation (EDA), où le schéma est souvent considéré à tort comme un échantillonnage sans biais.

Malgré son élégance, le RTO sans biais s'est avéré moins efficace que les autres méthodes d'UQ quand le nombre de variables était de l'ordre de quelques dizaines. Cela est dû au coût individuel de chaque échantillon. Ce coût est partagé avec le RTO naïf, mais la forte variabilité des poids correctifs du RTO sans biais ralentit la convergence. C'est pourquoi ce schéma s'est révélé être numériquement le moins efficace dans l'application à la CNPP et plus encore à la FDNPP.

Ainsi, nous concluons de ces expériences que la démarche BE la plus satisfaisante consiste à estimer les hyperparamètres θ^* avec l'EM, puis d'échantillonner $p(\mathbf{x}|\mathbf{y}, \theta^*)$ avec un MCMC.

De plus, nous avons montré que l'approche BE peut n'apporter qu'une solution partielle, tout particulièrement si la distribution des hyperparamètres est multimodale, car elle ne suppose qu'un unique mode pour ces hyperparamètres. Il est probable que cela soit provoqué par une incohérence entre le modèle et la base de données d'observations ; ce problème ne peut être résolu qu'en prenant mieux en compte l'erreur modèle.

La solution BC avec une représentation transdimensionnelle du maillage nous a permis de mettre en avant ce risque en diagnostiquant une distribution multimodale dans le cas de la FDNPP. Le schéma est plus lent à converger qu'un MCMC avec un maillage fixe et un nombre équivalent de variables, en raison du coût supplémentaire imposé par le maillage adaptatif et du plus faible taux d'acceptation de l'étape de Metropolis-Hastings. Cependant, les résultats obtenus sont très satisfaisants, produisant une solution quasiment continue par morceaux.

Pour limiter la taille cette étude, nous n'avons présenté qu'une sélection des résultats d'UQ pour la CNPP et la FDNPP. En particulier, nous n'avons étudié que le césium-137, nous avons réduit la base de données d'observations et nous n'avons pas tenu compte de l'erreur modèle au-delà de l'estimation de r . De plus, des inversions et des UQ utilisant un paramètre d'erreur d'observation par station de mesure dans le cas de la FDNPP ont été effectuées mais n'ont pas été présentées dans ce chapitre.

Conclusions

L'assimilation de données pour les modèles climatiques, de prévision numérique du temps ou de chimie atmosphérique, permet d'améliorer l'estimation de l'état de l'atmosphère et la qualité des prévisions. À ce titre, elle est utilisée dans des centres opérationnels de météorologie, comme Météo-France ou le Centre européen pour les prévisions météorologiques à moyen terme (CEPMMT). La présence accrue de données satellitaires denses permet aujourd'hui d'assimiler dans ces centres des millions d'observations quotidiennement. Des méthodes d'assimilation de données adaptées aux problèmes liées à ces modèles et à ces observations ont donc vu le jour.

En effet, depuis les années 60 et le développement du filtre de Kalman, les méthodes d'assimilation de données ont grandement évolué. Le filtre de Kalman nécessite que le modèle utilisé soit linéaire et parfaitement connu, que des matrices explicites de covariance des erreurs soient calculées et propagées dans le temps, et que ces mêmes erreurs soient gaussiennes. L'ensemble de ces hypothèses a peu à peu été assoupli. La non-linéarité des modèles a pu être prise en compte grâce aux méthodes variationnelles comme le 3D-Var et le 4D-Var, apparues dès les années 80. Elles définissent une fonction de coût non-linéaire qui doit être minimisée en recourant à l'une des méthodes offertes par la théorie de l'optimisation. Parallèlement, pour éviter la définition et la propagation de matrices de covariance d'erreur d'ébauche de grandes dimensions, des méthodes ensemblistes, nommées ensemble Kalman filter (EnKF), sont apparues. Les erreurs d'ébauche sont ainsi définies statistiquement, grâce à un ensemble de membres dont la dispersion est censée être représentative de l'erreur. La taille de l'ensemble étant généralement plus faible que celle des états, la propagation des erreurs est permise par la simple propagation des membres de l'ensemble. La faible taille de l'ensemble n'est cependant possible que par l'hypothèse de gaussianité des erreurs. On peut noter que des méthodes de filtres particulières existent également pour s'abstraire de cette dernière hypothèse et permettent de reconstituer entièrement la fonction de densité de probabilité associée à l'état estimé. Enfin, une branche de l'assimilation de données a été développée pour s'abstraire de l'hypothèse de modèles parfaits.

Le domaine d'application auquel nous nous sommes intéressés est celui de la chimie atmosphérique et plus précisément la chimie de l'ozone troposphérique. L'ozone, en tant qu'oxydant, est un polluant néfaste pour la santé et est sous le coup de législations européennes visant à limiter sa concentration. Les phénomènes chimiques responsables de sa formation font également intervenir les oxydes d'azote et les composés organiques volatils, mais la relation entre ces précurseurs et les concentrations d'ozone résultantes est non-linéaire. De plus, les modèles numériques utilisés pour prédire les concentrations d'ozone présentent de nombreuses sources d'incertitudes. L'enjeu sanitaire lié à l'estimation de la concentration d'ozone, ainsi que la complexité des modèles numériques rend nécessaire le recours à des méthodes avancées d'assimilation de données.

Parmi ces méthodes, le 4D-Var et l'EnKF, dans sa formulation localized ensemble transform Kalman filter (LETKF), sont majoritairement utilisées pour des modèles de météorologie ou de chimie atmosphérique. Cependant, elles ont chacune leur lot de défauts. La première nécessite d'avoir à sa disposition l'adjoint du tangent linéaire du modèle et de définir des matrices de covariance des erreurs d'ébauche fixes au cours du temps. La deuxième quant à elle ne permet pas de tenir compte de la non-linéarité des modèles au cours de l'analyse et nécessite de recourir à des méthodes empiriques, l'inflation et la localisation, pour pallier les erreurs d'échantillonnage provoquées par la taille réduite de l'ensemble. Face aux limites du 4D-Var et de l'EnKF, des méthodes variationnelles d'ensemble, ou EnVar 4D, ont vu le jour. Elles visent à tirer parti des deux approches pour définir

des méthodes qui prennent en compte les non-linéarités des modèles, sans recourir à leur tangent linéaire, et propagent les incertitudes au cours du temps sans imposer d'important coût de calcul. À ce titre, l'iterative ensemble Kalman smoother (IEnKS) est un parfait archétype des méthodes EnVar 4D. Il englobe en son sein les méthodes telles le LETKF, le maximum likelihood ensemble filter (MLEF), l'iterative ensemble Kalman filter (IEnKF), le 4D-LETKF et, dans une certaine mesure, le 4D-EnVar. Cependant, à ce jour, aucune des méthodes EnVar 4D n'a été appliqué dans un contexte réel avec un modèle de chimie atmosphérique.

Les méthodes d'assimilation de données sont donc variées et en constante évolution pour s'adapter aux différentes applications. Notre travail s'est inscrit dans cet effort combiné de développement de nouvelles méthodes et d'application à la chimie atmosphérique.

Tout d'abord, l'IEnKF a initialement été développé en supposant un modèle parfait, sans prendre en compte l'erreur modèle. Nous avons donc étendu cet algorithme au cas où une erreur modèle additive existe. L'IEnKF-Q a ainsi été introduit (Sakov *et al.*, 2017, soumis). Cet algorithme a été testé numériquement sur un modèle jouet de météorologie. Il a été comparé à un EnKF avec prise en compte de l'erreur modèle. Un IEnKF bénéficiant de la même prise en compte a également été testé, ce qui correspond à une implémentation naturelle, naïve et sous-optimal de l'erreur modèle. En effet, l'IEnKF-Q a fourni de meilleures performances que ces deux bases de comparaison dans tous les régimes testés.

Ensuite, nous avons appliqué les méthodes EnVar 4D sur un modèle de chimie atmosphérique. Étant donnée la complexité de ces modèles, la quantité d'incertitudes qui y sont attachées ainsi que la difficulté que représente l'assimilation d'observations réelles, nous avons testé ces méthodes dans un contexte idéalisé dans un premier temps. Pour ce faire, imitant le L95-T, nous avons développé un modèle d'ordre réduit de chimie atmosphérique, nommé L95-GRS (Haussaire et Bocquet, 2016). De faible dimension, issu du couplage d'une météorologie simpliste (le modèle Lorenz-95 (L95)) et d'un module de chimie de l'ozone troposphérique simplifié (le modèle *generic reaction set* (GRS)), ce modèle s'est avéré être un parfait banc d'essai pour l'application des méthodes EnVar 4D sur de la chimie atmosphérique. En effet, il reproduit les phénomènes principaux de la chimie de l'ozone troposphérique, malgré un coût de calcul réduit et une relative simplicité des phénomènes modélisés. Il permet donc de tester les méthodes d'assimilation de données dans un contexte maîtrisé mais présentant certaines des difficultés de la chimie atmosphérique.

Les résultats de l'application de l'IEnKS sur le L95-GRS et le L95-T se sont montrés encourageants, l'IEnKS démontrant sa supériorité face à l'ensemble transform Kalman filter (ETKF) dans divers contextes et permettant de bien tenir compte de la non-linéarité du modèle. Cependant, des premières limites ont aussi été soulevées. En effet, bien que le recours à une méthode variationnelle s'étendant sur une longue fenêtre d'assimilation de données sur un modèle fortement couplé apporte des résultats plus que satisfaisants, une telle méthode s'avère inutile si elle est appliquée sur un modèle faiblement couplé de traceur stable et linéaire. Ainsi, dans le cas *offline*, ces méthodes permettent d'améliorer les champs de météorologie, ce qui va avoir un effet positif sur les champs de traceur estimés, mais à qualité de météorologie constante, l'utilisation d'une longue fenêtre d'assimilation variationnelle n'impacte pas les performances d'estimation du traceur. En revanche, la présence de l'ensemble des champs météorologiques fourni par une méthode ensembliste, dont la dispersion est représentative de l'incertitude sur la météorologie, permet de tenir compte de l'erreur modèle due au couplage faible et d'améliorer grandement l'estimation des champs de traceur.

Fort de l'expérience de cette application de l'IEnKS sur ce modèle jouet, nous avons cherché à étendre ces résultats sur un modèle de chimie-transport (CTM) plus précis. Ces derniers se sont alors confirmés. Tout d'abord, l'assimilation d'observations synthétiques a validé la capacité de l'IEnKS à supplanter l'ETKF en filtrage, grâce à la meilleure prise en compte des non-linéarités du modèle. Les

résultats obtenus sur le modèle jouet en lissage se sont également vérifiés et ces derniers se dégradent quand la fenêtre d'assimilation est étendue. La stabilité des modèles de chimie atmosphérique est incriminée pour justifier ce résultat. Elle limite l'intérêt de méthodes variationnelles s'étendant sur une longue fenêtre d'assimilation pour la réanalyse de l'état du système.

Le passage d'observations synthétiques à des observations réelles soulève en revanche de nouvelles problématiques. En effet, l'IEnKS, et donc a fortiori les méthodes EnVar 4D, ne semble pas en mesure d'améliorer l'estimation de l'état en filtrage, en lissage ou en prédiction. Les raisons de cet échec provisoire peuvent être attribuées à la forte erreur du modèle utilisé, que même le recours à l'IEnKF-Q ne parvient à correctement prendre en considération. En effet, cette dernière n'est pas nécessairement additive et est également probablement mal paramétrée. Ainsi, il semblerait que tant que l'erreur modèle n'est pas réduite, mieux connue et paramétrisée, ou mieux prise en compte par la méthode d'assimilation de données, les méthodes EnVar 4D ne permettront pas d'améliorer l'estimation des concentrations de polluants dans l'atmosphère.

Parmi les sources d'erreur modèle possibles, on pourra par exemple relever les fortes incertitudes sur les champs d'émission de polluants utilisés, qui sont issus d'inventaires d'émission. L'estimation de telles incertitudes a fait l'objet de notre dernier axe de recherche. Plus précisément, la dernière partie du travail de thèse s'est portée sur les méthodes statistiques d'estimation de l'erreur associée à la reconstruction du terme source des accidents nucléaires de Tchernobyl et Fukushima Daiichi. Depuis longtemps, l'application de méthodes de modélisation inverse a permis de reconstruire le terme source de ces accidents, avec une certaine cohérence entre les différentes expériences de modélisation inverse. Cependant, peu de ces expériences ont fourni une estimation de l'incertitude à ces termes sources. Liu *et al.* (2017, soumis) présentent des méthodes permettant d'obtenir cette incertitude et de reconstruire la distribution de probabilité du terme source. Certaines d'entre elles ont fait l'objet d'améliorations, comme l'algorithme espérance-maximisation (EM), appliqué sur des statistiques log-normales, que nous avons étendu en utilisant un a priori de Jeffreys pour les hyperparamètres. La plupart des méthodes n'avaient jamais été utilisées dans ce contexte. C'est le cas en particulier des méthodes de Monte-Carlo par chaînes de Markov (MCMC), avec analyse transdimensionnelle, pour résoudre la hiérarchie d'inférences bayésiennes complète du problème d'estimation du terme source. De telles méthodes peuvent être utiles pour quantifier nombre des sources d'incertitudes attachées aux modèles de chimie atmosphérique, qui pénalisent l'application des méthodes d'assimilation de données. Par ailleurs, l'application de ces méthodes a montré que la distribution de probabilité reconstruite peut ne pas être gaussienne, voire même être multimodale. Un tel diagnostic remet en question l'hypothèse de gaussianité faite par la majorité des méthodes d'assimilation de données, autre que les filtres particuliers.

L'ensemble de ce travail a eu le mérite d'ouvrir la voie à l'assimilation de données variationnelle d'ensemble sur des modèles non-linéaires de chimie atmosphérique. Il aura permis d'une part de montrer les points forts et les limites de l'application de ces méthodes sur de tels modèles, et d'autre part de pointer vers des voies possibles d'amélioration pour rendre ces méthodes efficaces dans ce contexte.

Une première voie d'exploration possible consiste à tester ces méthodes EnVar 4D sur des modèles couplés de chimie-météorologie (CCMMs) plutôt que des CTMs. La réduction de l'erreur modèle associée aux champs de vent, les interactions entre les deux sous-systèmes de météorologie et de chimie lors de la propagation du modèle, et plus encore lors de l'assimilation de données, permises par le développement de covariances entre les variables de ces sous-systèmes, sont autant d'avantages fournis par l'utilisation d'un CCMM. Ces avantages qu'apporte un couplage fort ont été mis en avant sur un modèle jouet et leur absence s'est ressentie lors de l'application de l'IEnKS sur un CTM complexe. Cependant, l'application d'une méthode EnVar 4D sur un CCMM apportera

son propre lot de complications. L'espace des états étant augmenté, la taille de l'ensemble devra l'être également en concordance et de la localisation entre les variables météorologiques et chimiques doit être définie et implémentée.

La deuxième voie d'exploration consiste à développer et utiliser des méthodes d'assimilation de données adaptées à la présence d'erreur modèle. L'IEEnKF-Q en est un exemple, mais se contente de tenir compte d'une erreur modèle additive. D'autres moyens existent pour s'adapter à la présence d'erreur modèle. On aurait pu par exemple appliquer une perturbation stochastique sur les champs météorologiques ou d'émission de chacun des différents modèles de l'ensemble utilisé dans l'assimilation de données.

Cependant, pour appliquer efficacement de telles méthodes, il faut connaître avec précision la source, la nature et l'amplitude de cette erreur modèle. Il faut donc être en mesure de l'estimer convenablement et d'évaluer l'incertitude liée à cette estimation. À cette fin, l'application de méthodes statistiques fiables pour estimer les incertitudes liées notamment aux émissions ou aux champs météorologiques utilisés par ces modèles est un travail qui a déjà été engagé, et les méthodes EnVar 4D peuvent s'avérer utiles dans ce contexte.

Enfin, notre quantification d'incertitude liée à l'estimation du terme source des accidents de Tchernobyl et de Fukushima Daiichi ayant montré que la distribution de certaines des quantités considérées n'était pas gaussienne, un effort de développement doit être également porté sur des méthodes non-gaussiennes d'assimilation de données. Aussi, mon travail d'application des méthodes EnVar 4D sur des modèles de chimie atmosphérique sera continué et complété pour le comparer à l'utilisation de nouveaux algorithmes de filtres particuliers localisés, qui permettront entre autres de vérifier le bien fondé de l'hypothèse de gaussianité.

Bibliographie

- J. L. ANDERSON : An ensemble adjustment Kalman filter for data assimilation. *Mon. Weather Rev.*, 129:2884–2903, 2001. (Cité en page 61.)
- M. ASCH, M. BOCQUET et M. NODET : Chapter 6 : The ensemble Kalman filter. *In Data assimilation : methods, algorithms, and applications Asch et al. (2016c)*, chap. 6, p. 153–193. (Cité en page 30.)
- M. ASCH, M. BOCQUET et M. NODET : Chapter 7 : Ensemble variational methods. *In Data assimilation : methods, algorithms, and applications Asch et al. (2016c)*, chap. 7, p. 195–216. (Cité en page 39.)
- M. ASCH, M. BOCQUET et M. NODET : *Data assimilation : methods, algorithms, and applications*. Fundamentals of Algorithms. SIAM, 2016c. (Cité en page 159.)
- M. AZZI, G. M. JOHNSON et M. COPE : An introduction to the generic reaction set photochemical smog mechanism. *Proc. 11th Int. Clean Air Conf. 4th Regional IUAPPA Conf. Brisbane, Australia*, July 1992. (Cité en pages 63 et 64.)
- A. BAKLANOV et J. H. SØRENSEN : Parameterisation of radionuclide deposition in atmospheric long-range transport modelling. *Phys. Chem. Earth (B)*, 26:787–799, 2001. (Cité en page 135.)
- R. BALGOVIND, A. DALCHER, M. GHIL et E. KALNAY : A stochastic-dynamic model for the spatial structure of forecast error statistics. *Mon. Weather Rev.*, 111:701–722, 1983. (Cité en pages 107 et 110.)
- J. M. BARDSLEY, A. SOLONEN, H. HAARIO et M. LAINE : Randomize-then-optimize : A method for sampling from posterior distributions in nonlinear inverse problems. *SIAM J. Sci. Comput.*, 36:A1895–A1910, 2014. (Cité en pages 146 et 147.)
- A. BERCHET, I. PISON, F. CHEVALLIER, P. BOUSQUET, S. CONIL, M. GEEVER, T. LAURILA, J. LAVRIČ, M. LOPEZ, J. MONCRIEFF, J. NECKI, M. RAMONET, M. SCHMIDT, M. STEINBACHER et J. TARNIEWICZ : Towards better error statistics for atmospheric inversions of methane surface fluxes. *Atmos. Chem. Phys.*, 13:7115–7132, 2013. (Cité en page 137.)
- P. BERGAMASCHI, C. FRANKENBERG, J. F. MEIRINK, M. KROL, M. G. VILLANI, S. HOUWELING, F. DENTENER, E. J. DLUGOKENCKY, J. B. MILLER, L. V. GATTI, A. ENGEL et I. LEVIN : Inverse modeling of global and regional CH₄ emissions using SCIAMACHY satellite retrievals. *J. Geophys. Res.*, 114, 2009. (Cité en page 131.)
- B. BESSAGNET, G. PIROVANO, M. MIRCEA, C. CUVELIER, A. AULINGER, G. CALORI, G. CIARELLI, A. MANDERS, R. STERN, S. TSYRO, M. GARCÍA VIVANCO, P. THUNIS, M.-T. PAY, A. COLETTE, F. COUVIDAT, F. MELEUX, L. ROUİL, A. UNG, S. AKSOYOGLU, J. M. BALDASANO, J. BIESER, G. BRIGANTI, A. CAPPELLETTI, M. D’ISIDORO, S. FINARDI, R. KRANENBURG, C. SILIBELLO, C. CARNEVALE, W. AAS, J.-C. DUPONT, H. FAGERLI, L. GONZALEZ, L. MENUT, A. S. H. PRÉVÔT, P. ROBERTS et L. WHITE : Presentation of the EURODELTA III intercomparison exercise – evaluation of the chemistry transport models’ performance on criteria pollutants and joint analysis with meteorology. *Atmos. Chem. Phys.*, 16:12667–12701, 2016. (Cité en page 103.)

- C. M. BISHOP, éd. *Pattern Recognition and Machine Learning*. Springer-Verlag New-York Inc, 2006. (Cit  en page 138.)
- C. H. BISHOP, B. J. ETHERTON et S. J. MAJUMDAR : Adaptive sampling with the ensemble transform Kalman filter. Part I : Theoretical aspects. *Mon. Weather Rev.*, 129:420–436, 2001. (Cit  en page 26.)
-  . BLAYO, M. BOCQUET, E. COSME et L. F. CUGLIANDOLO,  ds. *Advanced Data Assimilation for Geosciences*. Lecture Notes of the Les Houches School of Physics : Special Issue, June 2012. Oxford University Press, Oxford, 2015. (Cit  en pages 55 et 160.)
- N. BLOND, L. BEL et R. VAUTARD : Three-dimensional ozone data analysis with an air quality model over the Paris area. *J. Geophys. Res.*, 108, 2003. ISSN 2156-2202. (Cit  en page 107.)
- M. BOCQUET : Parameter field estimation for atmospheric dispersion : Application to the Chernobyl accident using 4D-Var. *Q. J. R. Meteorol. Soc.*, 138:664–681, 2012. (Cit  en pages 133, 134, 143 et 144.)
- M. BOCQUET, C. A. PIRES et L. WU : Beyond Gaussian statistical modeling in geophysical data assimilation. *Mon. Weather Rev.*, 138:2997–3023, 2010. (Cit  en page 51.)
- M. BOCQUET, P. N. RAANES et A. HANNART : Expanding the validity of the ensemble Kalman filter without the intrinsic need for inflation. *Nonlinear Process. Geophys.*, 22:645–662, 2015. (Cit  en pages 37 et 38.)
- M. BOCQUET et P. SAKOV : An iterative ensemble Kalman smoother. *Q. J. R. Meteorol. Soc.*, 140:1521–1535, 2014. (Cit  en pages 39, 42, 43 et 45.)
- M. BOCQUET : Ensemble Kalman filtering without the intrinsic need for inflation. *Nonlinear Process. Geophys.*, 18:735–750, 2011. (Cit  en pages 37, 61 et 138.)
- M. BOCQUET : Introduction to the principles and methods of data assimilation in the geosciences. Notes de cours,  cole des Ponts ParisTech, 2014. URL <http://cerea.enpc.fr/HomePages/bocquet/teaching/assim-mb-en.pdf>. Cours du Master2 OACOS & WAPE. (Non cit .)
- M. BOCQUET : An introduction to inverse modelling and parameter estimation for atmosphere and ocean sciences. In *Blayo et al. (2015)*, p. 461–493. (Cit  en pages 137 et 144.)
- M. BOCQUET : Localization and the iterative ensemble Kalman smoother. *Q. J. R. Meteorol. Soc.*, 142:1075–1089, 2016. (Cit  en pages 39, 41, 45, 77, 80 et 81.)
- M. BOCQUET et A. CARRASSI : Four-dimensional ensemble variational data assimilation and the unstable subspace. *Tellus A*, 69:1304504, 2017. (Cit  en pages 79 et 118.)
- M. BOCQUET et P. SAKOV : Combining inflation-free and iterative ensemble Kalman filters for strongly nonlinear systems. *Nonlinear Process. Geophys.*, 19:383–399, 2012. (Cit  en pages 37 et 42.)
- M. BOCQUET et P. SAKOV : Joint state and parameter estimation with an iterative ensemble Kalman smoother. *Nonlinear Process. Geophys.*, 20:803–818, 2013. (Cit  en pages 43, 61, 62, 77, 78, 79, 80, 88, 89 et 90.)

- T. BODIN et M. SAMBRIDGE : Seismic tomography with the reversible jump algorithm. *Geophysical Journal International*, 178:1411–1436, 2009. (Cité en page 150.)
- M. BONAVIDA, L. ISAKSEN et E. HÓLM : On the use of EDA background error variances in the ECMWF 4D-Var. *Q. J. R. Meteorol. Soc.*, 138:1540–1559, 2012. (Cité en page 146.)
- A. BOYNARD, M. BEEKMANN, G. FORET, A. UNG, S. SZOPA, C. SCHMECHTIG et A. COMAN : An ensemble assessment of regional ozone model uncertainty with an explicit error representation. *Atmos. Environ.*, 45:784 – 793, 2011. ISSN 1352-2310. (Cité en pages 106, 107, 126 et 127.)
- M. BUEHNER, P. HOUTEKAMER, C. CHARETTE, H. L. MITCHELL et B. HE : Intercomparison of variational data assimilation and the ensemble Kalman filter for global deterministic NWP. Part I : Description and single-observation experiments. *Mon. Weather Rev.*, 138:1550–1566, 2010. (Cité en page 40.)
- M. BUEHNER, R. McTAGGART-COWAN, A. BEAULNE, C. CHARETTE, L. GARAND, S. HEILLIETTE, E. LAPALME, S. LAROCHE, S. R. MACPHERSON, J. MORNEAU et A. ZADRA : Implementation of deterministic weather forecasting systems based on ensemble–variational data assimilation at environment canada. Part I : The global system. *Mon. Weather Rev.*, 143:2532–2559, 2015. (Cité en page 40.)
- G. BURGERS, P. J. van LEEUWEN et G. EVENSEN : Analysis scheme in the ensemble Kalman filter. *Mon. Weather Rev.*, 126:1719–1724, 1998. (Cité en page 24.)
- G. R. CARMICHAEL, A. SANDU, T. CHAI, D. N. DAESCU, E. M. CONSTANTINESCU et Y. TANG : Predicting air quality : Improvements through advanced methods to integrate models and measurements. *J. Comput. Phys.*, 227:3540 – 3571, 2008. (Cité en page 106.)
- A. CARRASSI, S. VANNITSEM et C. NICOLIS : Model error and sequential data assimilation : A deterministic formulation. *Q. J. R. Meteorol. Soc.*, 134:1297–1313, 2008. (Cité en page 71.)
- A. CARRASSI, S. VANNITSEM, D. ZUPANSKI et M. ZUPANSKI : The maximum likelihood ensemble filter performances in chaotic systems. *Tellus A*, 61:587–600, 2009. (Cité en page 30.)
- T. CHAI, G. R. CARMICHAEL, Y. TANG, A. SANDU, M. HARDESTY, P. PILEWSKIE, S. WHITLOW, E. V. BROWELL, M. A. AVERY, P. NÉDÉLEC, J. T. MERRILL, A. M. THOMPSON et E. WILLIAMS : Four-dimensional data assimilation experiments with International Consortium for Atmospheric Research on Transport and Transformation ozone measurements. *J. Geophys. Res.*, 112, 2007. D12S15. (Cité en page 107.)
- F. CHEVALLIER, N. VIOVY, M. REICHSTEIN et P. CIAIS : On the assignment of prior errors in Bayesian inversion of CO₂ surface fluxes. *Geophys. Res. Lett.*, 33:L13802, 2006. (Cité en page 131.)
- M. CHINO, H. NAKAYAMA, H. NAGAI, H. TERADA, G. KATATA et H. YAMAZAWA : Preliminary estimation of release amounts of ¹³¹I and ¹³⁷Cs accidentally discharged from the Fukushima Daiichi nuclear power plant into the atmosphere. *J. Nucl. Sci. Technol.*, 48:1129–1134, 2011. (Cité en pages 133 et 143.)
- D. S. COHAN, A. HAKAMI, Y. HU et A. G. RUSSELL : Nonlinear response of ozone to emissions : Source apportionment and sensitivity analysis. *Environ. Sci. Technol.*, 39:6739–6748, 2005. (Cité en page 96.)

- S. E. COHN : An introduction to estimation theory. *J. Meteorol. Soc. Jpn.*, 75:257–288, 1997. (Cité en page 19.)
- A. COMAN, G. FORET, M. BEEKMANN, M. EREMENKO, G. DUFOUR, B. GAUBERT, A. UNG, C. SCHMECHTIG, J.-M. FLAUD et G. BERGAMETTI : Assimilation of IASI partial tropospheric columns with an ensemble Kalman filter over Europe. *Atmos. Chem. Phys.*, 12:2513–2532, 2012. (Cité en page 107.)
- E. M. CONSTANTINESCU, A. SANDU, T. CHAI et G. R. CARMICHAEL : Ensemble-based chemical data assimilation. I : General approach. *Q. J. R. Meteorol. Soc.*, 133:1229–1243, 2007a. (Cité en pages 106, 108, 126 et 127.)
- E. M. CONSTANTINESCU, A. SANDU, T. CHAI et G. R. CARMICHAEL : Ensemble-based chemical data assimilation. II : Covariance localization. *Q. J. R. Meteorol. Soc.*, 133:1245–1256, 2007b. (Cité en pages 93 et 107.)
- E. COSME, J.-M. BRANKART, J. VERRON, P. BRASSEUR et M. KRYSTA : Implementation of a reduced rank square-root smoother for high resolution ocean data assimilation. *Ocean Modelling*, 33:87 – 100, 2010. (Cité en page 31.)
- E. COSME, J. VERRON, P. BRASSEUR, J. BLUM et D. AUROUX : Smoothing problems in a Bayesian framework and their linear Gaussian solutions. *Mon. Weather Rev.*, 140:683–695, 2012. (Cité en page 34.)
- X. DAVOINE et M. BOCQUET : Inverse modelling-based reconstruction of the Chernobyl source term available for long-range transport. *Atmos. Chem. Phys.*, 7:1549–1564, 2007. (Cité en pages 133, 134, 137 et 143.)
- J. P. DAWSON, P. J. ADAMS et S. N. PANDIS : Sensitivity of PM_{2.5} to climate in the Eastern US : a modeling case study. *Atmos. Chem. Phys.*, 7:4295–4309, 2007. (Cité en page 96.)
- L. DELLE MONACHE, J. K. LUNDQUIST, B. KOSOVIĆ, G. JOHANNESSON, K. M. DYER, R. D. AINES, F. K. CHOW, R. D. BELLES, W. G. HANLEY, S. C. LARSEN, G. A. LOOSMORE, J. J. NITAO, G. A. SUGIYAMA et P. J. VOGT : Bayesian inference and Markov chain Monte Carlo sampling to reconstruct a contaminant source on a continental scale. *Journal of Applied Meteorology and Climatology*, 47:2600–2613, 2008. (Cité en page 138.)
- A. P. DEMPSTER, N. M. LAIRD et D. B. RUBIN : Maximum likelihood from incomplete data via the EM algorithm. *Journal of the Royal Statistical Society. Series B*, p. 1–38, 1977. (Cité en page 138.)
- G. DESROZIER, L. BERRE, B. CHAPNIK et P. POLI : Diagnosis of observation, background and analysis-error statistics in observation space. *Q. J. R. Meteorol. Soc.*, 131:3385–3396, 2005. (Cité en pages 57 et 108.)
- G. DESROZIER, J.-T. CAMINO et L. BERRE : 4DEnVar : link with 4D state formulation of variational assimilation and different possible implementations. *Q. J. R. Meteorol. Soc.*, 140:2097–2110, 2014. (Cité en pages 39 et 40.)
- G. DESROZIER, É. ARBOGAST et L. BERRE : Improving spatial localisation in 4DEnVar. *Q. J. R. Meteorol. Soc.*, 142:3171–3185, 2016. (Cité en page 41.)

- A. DOICU, T. TRAUTMANN et F. SCHREIER : *Numerical Regularization for Atmospheric Inverse Problems*. Springer-Verlag Berlin Heidelberg, 2010. (Cité en page 138.)
- A. DOUCET, S. GODSILL et C. ANDRIEU : On sequential Monte Carlo sampling methods for Bayesian filtering. *Stat. Comput.*, 10:197–208, 2000. (Cité en page 51.)
- J. DUDHIA, D. GILL, K. MANNING, W. WANG, C. BRUYERE, S. KELLY et K. LACKEY : PSU/NCAR mesoscale modeling system tutorial class notes and user’s guide : MM5 modeling system version 3. Rap. tech., NCAR, 2005. (Cité en page 10.)
- H. ELBERN, A. STRUNK, H. SCHMIDT et O. TALAGRAND : Emission rate and chemical state estimation by 4-dimensional variational inversion. *Atmos. Chem. Phys.*, 7:3749–3769, 2007. (Cité en page 131.)
- E. EMILI, S. GÜROL et D. CARIOLLE : Accounting for model error in air quality forecasts : an application of 4DEnVar to the assimilation of atmospheric composition using QG-Chem 1.0. *Geosci. Model Dev.*, 9:3933–3959, 2016. (Cité en pages 60, 72, 73, 87, 126 et 127.)
- L. K. EMMONS, S. WALTERS, P. G. HESS, J.-F. LAMARQUE, G. G. PFISTER, D. FILLMORE, C. GRANIER, A. GUENTHER, D. KINNISON, T. LAEPPE, J. ORLANDO, X. TIE, G. TYNDALL, C. WIEDINMYER, S. L. BAUGHUM et S. KLOSTER : Description and evaluation of the Model for Ozone and Related chemical Tracers, version 4 (MOZART-4). *Geosci. Model Dev.*, 3:43–67, 2010. (Cité en page 9.)
- N. EVANGELIOU, T. HAMBURGER, N. TALERKO, S. ZIBTSEV, Y. BONDAR, A. STOHL, Y. BALKANSKI, T. A. MOUSSEAU et A. P. MØLLER : Reconstructing the Chernobyl Nuclear Power Plant (CNPP) accident 30 years after. A unique database of air concentration and deposition measurements over Europe. *Environ. Pollut.*, 216:408–418, 2016. (Cité en page 134.)
- G. EVENSEN : Sequential data assimilation with a nonlinear quasi-geostrophic model using Monte Carlo methods to forecast error statistics. *J. Geophys. Res.*, 99:10143–10162, 1994. (Cité en pages 24 et 72.)
- G. EVENSEN : The ensemble Kalman filter : theoretical formulation and practical implementation. *Ocean dynamics*, 53:343–367, 2003. (Cité en page 31.)
- G. EVENSEN et P. J. van LEEUWEN : An ensemble Kalman smoother for nonlinear dynamics. *Mon. Weather Rev.*, 128:1852–1867, 2000. (Cité en pages 31 et 48.)
- D. FAIRBAIRN, S. R. PRING, A. C. LORENC et I. ROULSTONE : A comparison of 4DVar with ensemble data assimilation methods. *Q. J. R. Meteorol. Soc.*, 140:281–294, 2014. (Cité en page 41.)
- A. FARCHI, M. BOCQUET, Y. ROUSTAN, A. MATHIEU et A. QUÉREL : Using the Wasserstein distance to compare fields of pollutants : application to the radionuclide atmospheric dispersion of the Fukushima-Daiichi accident. *Tellus B*, 68:31682, 2016. (Cité en pages 131, 134 et 136.)
- B. J. FINLAYSON-PITTS et J. N. PITTS JR : *Atmospheric chemistry : Fundamentals and experimental techniques*. John Wiley and Sons, New York, NY, 1986. (Cité en page 69.)
- S. J. FLETCHER et M. ZUPANSKI : A data assimilation method for log-normally distributed observational errors. *Q. J. R. Meteorol. Soc.*, 132:2505–2519, 2006. (Cité en page 135.)

- J. FRYDENDALL, J. BRANDT et J. H. CHRISTENSEN : Implementation and testing of a simple data assimilation algorithm in the regional air pollution forecast model, DEOM. *Atmos. Chem. Phys.*, 9:5475–5488, 2009. (Cité en page 107.)
- A. L. GANESAN, M. RIGBY, A. ZAMMIT-MANGION, A. J. MANNING, R. G. PRINN, P. J. FRASER, C. M. HARTH, K.-R. KIM, P. B. KRUMMEL, S. LI, J. MÜHLE, S. J. O'DOHERTY, S. PARK, P. K. SALAMEH, L. P. STEELE et R. F. WEISS : Characterization of uncertainties in atmospheric trace gas inversions using hierarchical Bayesian methods. *Atmos. Chem. Phys.*, 14:3855–3864, 2014. (Cité en page 138.)
- G. GASPARI et S. E. COHN : Construction of correlation functions in two and three dimensions. *Q. J. R. Meteorol. Soc.*, 125:723–757, 1999. (Cité en pages 36 et 110.)
- B. GAUBERT, A. COMAN, G. FORET, F. MELEUX, A. UNG, L. ROUIL, A. IONESCU, Y. CANDAU et M. BEEKMANN : Regional scale ozone data assimilation using an ensemble Kalman filter and the CHIMERE chemical transport model. *Geoscientific Model Development*, 7:283–302, 2014. (Cité en pages 19, 105, 106, 107, 108 et 109.)
- B. GAUBERT : *Assimilation des observations pour la modélisation de la qualité de l'air*. Thèse de doctorat, Paris 7, 2013. (Cité en page 19.)
- W. S. GOLIFF, W. R. STOCKWELL et C. V. LAWSON : The regional atmospheric chemistry mechanism, version 2. *Atmos. Environ.*, 68:174–185, 2013. (Cité en page 13.)
- P. J. GREEN : Reversible jump Markov chain Monte Carlo computation and Bayesian model determination. *Biometrika*, 82:711–732, 1995. (Cité en pages 147 et 149.)
- P. J. GREEN : Trans-dimensional Markov chain Monte Carlo. In P. J. GREEN, N. L. HJORT et S. RICHARDSON, eds : *Highly Structured Stochastic Systems*, chap. 6, p. 179–198. Oxford University Press, Oxford, 2003. (Cité en pages 149 et 150.)
- Y. GU et D. S. OLIVER : An iterative ensemble Kalman filter for multiphase fluid flow data assimilation. *SPE Journal*, 12:438–446, 2007. (Cité en page 48.)
- P. H. GUDIKNEN, T. F. HARVEY et R. LANGE : Chernobyl source term, atmospheric dispersion, and dose estimation. *Health Phys.*, 57:697–706, 1989. (Cité en page 133.)
- A. GUENTHER, T. KARL, P. HARLEY, C. WIEDINMYER, P. I. PALMER et C. GERON : Estimates of global terrestrial isoprene emissions using MEGAN (Model of Emissions of Gases and Aerosols from Nature). *Atmos. Chem. Phys.*, 6:3181–3210, 2006. (Cité en pages 9 et 102.)
- G. GUEROVA, I. BEY, J.-L. ATTÍÉ, R. V. MARTIN, J. CUI et M. SPRENGER : Impact of transatlantic transport episodes on summertime ozone in Europe. *Atmos. Chem. Phys.*, 6:2057–2072, 2006. (Cité en page 70.)
- K. R. GURNEY, R. M. LAW, A. S. DENNING, P. J. RAYNER, D. BAKER, P. BOUSQUET, L. BRUH-WILER, Y.-H. CHEN, P. CIAIS, S. FAN *et al.* : Towards robust regional estimates of CO₂ sources and sinks using atmospheric transport models. *Nature*, 415:626–630, 2002. (Cité en page 131.)
- T. M. HAMILL et C. SNYDER : A hybrid ensemble Kalman filter-3D variational analysis scheme. *Mon. Weather Rev.*, 128:2905–2919, 2000. (Cité en page 40.)

- R. G. HANEA, G. J. M. VELDEERS et A. HEEMINK : Data assimilation of ground-level ozone in Europe with a Kalman filter and chemistry transport model. *J. Geophys. Res.*, 109, 2004. D10302. (Cité en pages 106 et 107.)
- P. C. HANSEN : *Discrete Inverse Problems : Insight and Algorithms*. Fundamentals of Algorithms. SIAM, Philadelphia, 2010. (Cité en page 138.)
- W. K. HASTINGS : Monte Carlo sampling methods using Markov chains and their applications. *Biometrika*, 57:97–109, 1970. (Cité en page 147.)
- J.-M. HAUSSAIRE et M. BOCQUET : A low-order coupled chemistry meteorology model for testing online and offline data assimilation schemes : L95-GRS (v1. 0). *Geosci. Model Dev.*, 9:393–412, 2016. (Cité en pages 43, 59, 68, 69, 70, 71, 75, 89, 91, 92, 93, 94, 95, 97, 98 et 156.)
- S. HIRAO, H. YAMAZAWA et T. NAGAE : Estimation of release rate of iodine-131 and cesium-137 from the Fukushima Daiichi nuclear power plant. *J. Nucl. Sci. Technol.*, 50:139–147, 2013. (Cité en page 133.)
- A.-G. HJELLBREKKE, A. M. FÆJRAA et S. SOLBER : Ozone measurements 2009. Rap. tech., EMEP, 2011. (Cité en page 19.)
- L. W. HOROWITZ, S. WALTERS, D. L. MAUZERALL, L. K. EMMONS, P. J. RASCH, C. GRANIER, X. TIE, J.-F. LAMARQUE, M. G. SCHULTZ, G. S. TYNDALL, J. J. ORLANDO et G. P. BRASSEUR : A global simulation of tropospheric ozone and related tracers : Description and evaluation of MOZART, version 2. *J. Geophys. Res.*, 108, 2003. 4784. (Cité en pages 9 et 102.)
- P. L. HOUTEKAMER et H. L. MITCHELL : Data assimilation using an ensemble Kalman filter technique. *Mon. Weather Rev.*, 126:796–811, 1998. (Cité en page 72.)
- W. HUNSDORFER et J. G. VERWER : *Numerical Solution of Time-Dependent Advection-Diffusion-Reaction Equations*, vol. 33 de *Springer Series in Computational Mathematics*. Springer-Verlag Berlin Heidelberg, 2003. (Cité en page 13.)
- B. R. HUNT, E. J. KOSTELICH et I. SZUNYOGH : Efficient data assimilation for spatiotemporal chaos : A local ensemble transform Kalman filter. *Physica D*, 230:112–126, 2007. (Cité en pages 26, 34, 35, 41, 42 et 118.)
- B. R. HUNT, E. KALNAY, E. J. KOSTELICH, E. OTT, D. PATIL, T. SAUER, I. SZUNYOGH, J. YORKE et A. ZIMIN : Four-dimensional ensemble Kalman filtering. *Tellus A*, 56:273–277, 2004. (Cité en pages 34, 41, 48 et 61.)
- H. JEFFREYS : *Theory of Probability*. Oxford University Press, troisième édn, 1961. (Cité en page 140.)
- R. E. KALMAN : A new approach to linear filtering and prediction problems. *J. Basic. Eng.*, 82:35–45, 1960. (Cité en page 20.)
- T. KAMINSKI, P. J. RAYNER, M. HEIMANN et I. G. ENTING : On aggregation errors in atmospheric transport inversions. *J. Geophys. Res.*, 106:4703–4715, 2001. (Cité en page 131.)

- G. KATATA, M. CHINO, T. KOBAYASHI, H. TERADA, M. OTA, H. NAGAI, M. KAJINO, R. DRAXLER, M. C. HORT, A. MALO, T. TORII et Y. SANADA : Detailed source term estimation of the atmospheric release for the Fukushima Daiichi Nuclear Power Station accident by coupling simulations of an atmospheric dispersion model with an improved deposition scheme and oceanic dispersion model. *Atmos. Chem. Phys.*, 15:1029–1070, 2015. (Cité en pages 133 et 143.)
- G. KATATA, M. OTA, H. TERADA, M. CHINO et H. NAGAI : Atmospheric discharge and dispersion of radionuclides during the Fukushima Dai-ichi Nuclear Power Plant accident, Part I : Source term estimation and local-scale atmospheric dispersion in early phase of the accident. *J. Environ. Radioact.*, 109:103–113, 2012. (Cité en page 133.)
- T. KOBAYASHI, H. NAGAI, M. CHINO et H. KAWAMURA : Source term estimation of atmospheric release due to the Fukushima Dai-ichi Nuclear power Plant accident by atmospheric and oceanic dispersion simulations. *J. Nucl. Sci. Technol.*, 50:255–264, 2013. (Cité en page 133.)
- M. R. KOOHKAN et M. BOCQUET : Accounting for representativeness errors in the inversion of atmospheric constituent emissions : Application to the retrieval of regional carbon monoxide fluxes. *Tellus B*, 64:19047, 2012. (Cité en pages 57, 111, 112 et 131.)
- M. R. KOOHKAN, M. BOCQUET, Y. ROUSTAN, Y. KIM et C. SEIGNEUR : Estimation of volatile organic compound emissions for Europe using data assimilation. *Atmos. Chem. Phys.*, 13:5887–5905, 2013. (Cité en pages 9, 131 et 137.)
- M. KRISTA et M. BOCQUET : Source reconstruction of an accidental radionuclide release at European scale. *Q. J. R. Meteorol. Soc.*, 133:529–544, 2007. (Cité en page 133.)
- F. LE DIMET et J. BLUM : Assimilation de données pour les fluides géophysiques. *Matapli*, (67), p. 33–55, 2002. (Cité en page 17.)
- M. LIN, A. M. FIORE, L. W. HOROWITZ, O. R. COOPER, V. NAIK, J. HOLLOWAY, B. J. JOHNSON, A. M. MIDDLEBROOK, S. J. OLTMANS, I. B. POLLACK, T. B. RYERSON, J. X. WARNER, C. WIEDINMYER, J. WILSON et B. WYMAN : Transport of Asian ozone pollution into surface air over the western United States in spring. *J. Geophys. Res.*, 117, 2012. D00V07. (Cité en page 70.)
- C. LIU, Q. XIAO et B. WANG : An ensemble-based four-dimensional variational data assimilation scheme. Part I : Technical formulation and preliminary test. *Mon. Weather Rev.*, 136:3363–3373, 2008. (Cité en page 40.)
- Y. LIU, J.-M. HAUSSAIRE, M. BOCQUET, Y. ROUSTAN, O. SAUNIER et A. MATHIEU : Uncertainty quantification of pollutant source retrieval : comparison of Bayesian methods with application to the Chernobyl and Fukushima Daiichi accidental releases of radionuclides. *Q. J. R. Meteorol. Soc.*, 2017, soumis. (Cité en pages 127, 129, 142, 147, 148, 151, 152, 153 et 157.)
- A. C. LORENC : Recommended nomenclature for EnVar data assimilation methods. *Research Activities in Atmospheric and Oceanic Modeling, WGNE*, 2013. (Cité en page 39.)
- A. C. LORENC, N. E. BOWLER, A. M. CLAYTON, S. R. PRING et D. FAIRBAIRN : Comparison of hybrid-4DVar and hybrid-4DVar data assimilation methods for global NWP. *Mon. Weather Rev.*, 143:212–229, 2015. (Cité en page 40.)

- E. N. LORENZ et K. A. EMANUEL : Optimal sites for supplementary weather observations : Simulation with a small model. *J. Atmos. Sci.*, 55:399–414, 1998. (Cité en pages 61 et 70.)
- J.-F. LOUIS : A parametric model of vertical eddy fluxes in the atmosphere. *Bound.-Layer Meteorol.*, 17:187–202, 1979. (Cité en pages 6, 102, 103, 116 et 135.)
- M. F. LUNT, M. RIGBY, A. L. GANESAN et A. J. MANNING : Estimation of trace gas fluxes with objectively determined basis functions using reversible-jump Markov chain Monte Carlo. *Geosci. Model Dev.*, 9:3213–3229, 2016. (Cité en page 150.)
- D. J. C. MACKAY : *Information Theory, Inference and Learning Algorithms*. Cambridge University Press, Cambridge, 2003. (Cité en page 144.)
- A. MALINVERNO et V. A. BRIGGS : Expanded uncertainty quantification in inverse problems : Hierarchical Bayes and empirical Bayes. *Geophysics*, 69:1005–1016, 2004. (Cité en page 138.)
- V. MALLET : Introduction aux modèles de chimie-transport pour la qualité de l'air. Notes de cours, École des Ponts ParisTech, Janvier 2008. URL <http://vivienmallet.net/static/publication/mallet08introduction.pdf>. Cours du Master SGE-AQA. (Non cité.)
- V. MALLET, D. QUÉLO, B. SPORTISSE, M. Ahmed de BIASI, É. DEBRY, I. KORSAKISSOK, L. WU, Y. ROUSTAN, K. SARTELET, M. TOMBETTE et H. FOU DHIL : Technical note : The air quality modeling system Polyphemus. *Atmos. Chem. Phys.*, 7:5479–5487, 2007. (Cité en pages 10, 13, 101, 102 et 134.)
- J. MANDEL, E. BERGOU, S. GÜROL, S. GRATTON et I. KASANICKÝ : Hybrid Levenberg-Marquardt and weak-constraint ensemble Kalman smoother method. *Nonlinear Process. Geophys.*, 23:59–73, 2016. (Cité en pages 48 et 127.)
- V. MARÉCAL, V.-H. PEUCH, C. ANDERSSON, S. ANDERSSON, J. ARTETA, M. BEEKMANN, A. BENEDICTOW, R. BERGSTRÖM, B. BESSAGNET, A. CANSADO, F. CHÉROUX, A. COLETTE, A. COMAN, R. L. CURIER, H. A. C. Denier van der GON, A. DROUIN, H. ELBERN, E. EMILI, R. J. ENGELEN, H. J. ESKES, G. FORET, E. FRIESE, M. GAUSS, C. GIANNAROS, J. GUTH, M. JOLY, E. JAUMOILLÉ, B. JOSSE, N. KADYGROV, J. W. KAISER, K. KRAJSEK, J. KUENEN, U. KUMAR, N. LIORA, E. LOPEZ, L. MALHERBE, I. MARTINEZ, D. MELAS, F. MELEUX, L. MENUT, P. MOINAT, T. MORALES, J. PARMENTIER, A. PIACENTINI, M. PLU, A. POUPKOU, S. QUEGUINER, L. ROBERTSON, L. ROUÏL, M. SCHAAP, A. SEGERS, M. SOFIEV, L. TARASSON, M. THOMAS, R. TIMMERMAN, A. VALDEBENITO, P. van VELTHOVEN, R. van VERSEDAAL, J. VIRA et A. UNG : A regional air quality forecasting system over Europe : the MACC-II daily ensemble production. *Geosci. Model Dev.*, 8:2777–2813, 2015. (Cité en page 73.)
- A. MATHIEU, I. KORSAKISSOK, D. QUÉLO, J. GROËLL, M. TOMBETTE, D. DIDIER, E. QUENTRIC, O. SAUNIER, J.-P. BENOIT et O. ISNARD : Atmospheric dispersion and deposition of radionuclides from the Fukushima Daiichi nuclear power plant accident. *Elements*, 8:195–200, 2012. (Cité en page 133.)
- R. MÉNARD, S. E. COHN, L.-P. CHANG et P. M. Lyster : Assimilation of stratospheric chemical tracer observations using a Kalman filter. Part I : Formulation. *Mon. Weather Rev.*, 128:2654–2671, 2000. (Cité en page 57.)

- L. MENUT, B. BESSAGNET, D. KHVOROSTYANOV, M. BEEKMANN, N. BLOND, A. COLETTE, I. COLL, G. CURCI, G. FORET, A. HODZIC, S. MAILLER, F. MELEUX, J.-L. MONGE, I. PISON, G. SIOUR, S. TURQUETY, M. VALARI, R. VAUTARD et M. G. VIVANCO : CHIMERE 2013 : a model for regional atmospheric composition modelling. *Geosci. Model Dev.*, 6(4):981–1028, 2013. (Non cité.)
- N. METROPOLIS, A. W. ROSENBLUTH, M. N. ROSENBLUTH, A. H. TELLER et E. TELLER : Equation of state calculations by fast computing machines. *J. Chem. Phys.*, 21:1087–1092, 1953. (Cité en page 147.)
- A. M. MICHALAK, A. HIRSCH, L. BRUHWILER, K. R. GURNEY, W. PETERS et P. P. TANS : Maximum likelihood estimation of covariance parameters for Bayesian atmospheric trace gas surface flux inversions. *J. Geophys. Res.*, 110:D24107, 2005. (Cité en page 137.)
- T. MILEWSKI et M. S. BOURQUI : Assimilation of stratospheric temperature and ozone with an ensemble Kalman filter in a chemistry–climate model. *Mon. Weather Rev.*, 139:3389–3404, 2011. (Cité en page 127.)
- T. MILEWSKI et M. S. BOURQUI : Potential of an ensemble Kalman smoother for stratospheric chemical-dynamical data assimilation. *Tellus A*, 65:18541, 2013. (Cité en page 127.)
- K. MIYAZAKI, H. J. ESKES, K. SUDO, M. TAKIGAWA, M. van WEELE et K. F. BOERSMA : Simultaneous assimilation of satellite NO₂, O₃, CO, and HNO₃ data for the analysis of tropospheric chemical composition and emissions. *Atmos. Chem. Phys.*, 12:9545–9579, 2012. (Cité en page 106.)
- T. NAKAJIMA, S. MISAWA, Y. MORINO, H. TSURUTA, D. GOTO, J. UCHIDA, T. TAKEMURA, T. OHARA, Y. OURA, M. EBIHARA et M. SATOH : Model depiction of the atmospheric flows of radioactive cesium emitted from the Fukushima Daiichi Nuclear Power Station accident. *Progress in Earth and Planetary Science*, 4:2, 2017. (Cité en page 133.)
- J. NOCEDAL et S. J. WRIGHT : *Numerical Optimization*. Springer Series in Operations Research and Financial Engineering. Springer-Verlag New York, 2006. (Cité en page 22.)
- D. S. OLIVER, N. HE et A. C. REYNOLDS : Conditioning permeability fields to pressure data. In *ECMOR V-5th European Conference on the Mathematics of Oil Recovery*, p. 259–269, 1996. (Cité en page 146.)
- E. OTT, B. R. HUNT, I. SZUNYOGH, A. V. ZIMIN, E. J. KOSTELICH, M. CORAZZA, E. KALNAY, D. J. PATIL et J. A. YORKE : A local ensemble Kalman filter for atmospheric data assimilation. *Tellus A*, 56:415–428, 2004. (Cité en pages 26 et 61.)
- Y. OURA, M. EBIHARA, H. TSURUTA, T. NAKAJIMA, T. OHARA, M. ISHIMOTO, H. SAWAHATA, Y. KATSUMURA et W. NITTA : A database of hourly atmospheric concentrations of radiocesium (¹³⁴Cs and ¹³⁷Cs) in suspended particulate matter collected in March 2011 at 99 air pollution monitoring stations in Eastern Japan. *J. Nucl. Radiochem. Sci.*, 15:1–12, 2015. (Cité en pages 133 et 134.)
- M. PAGOWSKI et G. A. GRELL : Experiments with the assimilation of fine aerosols using an ensemble Kalman filter. *J. Geophys. Res.*, 117, 2012. D21302. (Cité en page 106.)

- L. PALATELLA, A. CARRASSI et A. TREVISAN : Lyapunov vectors and assimilation in the unstable subspace : theory and applications. *Journal of Physics A : Mathematical and Theoretical*, 46:254020, 2013. (Cité en page 61.)
- C. PICCOLO et M. CULLEN : Ensemble data assimilation using a unified representation of model error. *Mon. Weather Rev.*, 144:213–224, 2016. (Cité en pages 87, 126 et 127.)
- J. PUDYKIEWICZ : Simulation of the Chernobyl dispersion with a 3-D hemispheric tracer model. *Tellus B*, 41:391–412, 1989. (Cité en page 135.)
- M. PULIDO, P. TANDEO, M. BOCQUET, A. CARRASSI et M. LUCINI : Parameter estimation in stochastic dynamical systems using expectation-maximization and Newton-Raphson maximum likelihood. *Phys. Rev. E*, 0:0–0, 2017. submitted. (Cité en page 139.)
- D. QUÉLO, M. KRISTA, M. BOCQUET, O. ISNARD, Y. MINIER et B. SPORTISSE : Validation of the Polyphemus platform on the ETEX, Chernobyl and Algeciras cases. *Atmos. Environ.*, 41:5300–5315, 2007. (Cité en page 134.)
- D. QUÉLO, V. MALLET et B. SPORTISSE : Inverse modeling of NO_x emissions at regional scale over northern France : Preliminary investigation of the second-order sensitivity. *J. Geophys. Res.*, 110:D24310, 2005. (Cité en page 131.)
- A. QUÉREL, Y. ROUSTAN, D. QUÉLO et J.-P. BENOIT : Hints to discriminate the choice of wet deposition models applied to an accidental radioactive release. *International Journal of Environment and Pollution*, 58:268–279, 2015. (Cité en page 134.)
- P. N. RAANES : On the ensemble Rauch-Tung-Striebel smoother and its equivalence to the ensemble Kalman smoother. *Q. J. R. Meteorol. Soc.*, 142:1259–1264, 2016. (Cité en page 27.)
- P. N. RAANES, A. CARRASSI et L. BERTINO : Extending the square root method to account for additive forecast noise in ensemble methods. *Mon. Weather Rev.*, 143:3857–3873, 2015. (Cité en pages 39, 82 et 83.)
- L. RAYNAUD, L. BERRE et G. DESROZIER : An extended specification of flow-dependent background error variances in the Météo-France global 4D-Var system. *Q. J. R. Meteorol. Soc.*, 137:607–619, 2011. (Cité en page 146.)
- S. J. ROSELLE et F. S. BINKOWSKI : Cloud dynamics and chemistry. In D. W. BYUN et C. J. K. S., eds : *Science Algorithms of the EPA Models-3 Community Multiscale Air Quality (CMAQ) Modeling System*. U.S. Environmental Protection Agency, 1999. (Cité en page 135.)
- Y. ROUSTAN, K. SARTELET, M. TOMBETTE, É. DEBRY et B. SPORTISSE : Simulation of aerosols and gas-phase species over Europe with the POLYPHEMUS system. part II : Model sensitivity analysis for 2001. *Atmos. Environ.*, 44:4219–4229, 2010. (Cité en page 96.)
- G. A. RUGGIERO, E. COSME, J.-M. BRANKART, J. L. SOMMER et C. UBELMANN : An efficient way to account for observation error correlations in the assimilation of data from the future SWOT high-resolution altimeter mission. *J. Atmos. Oceanic Technol.*, 33:2755–2768, 2016. (Cité en page 19.)
- J. J. RUIZ, M. PULIDO et T. MIYOSHI : Estimating model parameters with ensemble-based data assimilation : A review. *J. Meteorol. Soc. Jpn.*, 91:79–99, 2013. (Cité en page 78.)

- P. SAKOV et P. R. OKE : A deterministic formulation of the ensemble Kalman filter : an alternative to ensemble square root filters. *Tellus A*, 60:361–371, 2008. (Cité en page 30.)
- P. SAKOV : EnKF lectures. Notes de cours, NERSC, 2011. URL <http://enkf.nersc.no/Presentations/EnKFlecturesbyPavelSakov/>. School on Data assimilation. (Non cité.)
- P. SAKOV et L. BERTINO : Relation between two common localisation methods for the EnKF. *Comput. Geosci.*, 15:225–237, 2011. (Cité en page 36.)
- P. SAKOV, G. EVENSEN et L. BERTINO : Asynchronous data assimilation with the EnKF. *Tellus A*, 62:24–29, 2010. (Cité en page 48.)
- P. SAKOV, J.-M. HAUSSAIRE et M. BOCQUET : An iterative ensemble Kalman filter in presence of additive model error. *Q. J. R. Meteorol. Soc.*, 2017, soumis. (Cité en pages 17, 48, 53, 75, 81, 83, 84, 85, 86 et 156.)
- P. SAKOV, D. S. OLIVER et L. BERTINO : An iterative EnKF for strongly nonlinear systems. *Mon. Weather Rev.*, 140:1988–2004, 2012. (Cité en pages 30, 42, 51 et 61.)
- M. SAMBRIDGE, T. BODIN, K. GALLAGHER et H. TKALČIĆ : Transdimensional inference in the geosciences. *Phil. Trans. R. Soc. A*, 371:20110547, 2013. (Cité en page 149.)
- A. SANDU, E. CONSTANTINESCU, G. R. CARMICHAEL, T. CHAI, D. DAESCU et J. H. SEINFELD : Ensemble methods for dynamic data assimilation of chemical observations in atmospheric models. *Journal of Algorithms & Computational Technology*, 5:667–692, 2011. (Cité en page 106.)
- A. SANDU, E. M. CONSTANTINESCU, W. LIAO, G. R. CARMICHAEL, T. CHAI, J. H. SEINFELD et D. DAESCU : Ensemble-based data assimilation for atmospheric chemical transport models. In V. S. SUNDERAM, G. D. van ALBADA, P. M. A. SLOOT et J. J. DONGARRA, eds : *Computational Science – ICCS 2005 : 5th International Conference, Atlanta, GA, USA, May 22-25, 2005. Proceedings, Part II*, p. 648–655, Berlin, Heidelberg, 2005. Springer Berlin Heidelberg. (Cité en page 106.)
- O. SAUNIER, A. MATHIEU, D. DIDIER, M. TOMBETTE, D. QUÉLO, V. WINIAREK et M. BOCQUET : An inverse modeling method to assess the source term of the Fukushima Nuclear Power Plant accident using gamma dose rate observations. *Atmos. Chem. Phys.*, 13:11403–11421, 2013. (Cité en pages 133 et 143.)
- O. SAUNIER, A. MATHIEU, T. SEKIYAMA, M. KAJINO, K. ADACHI, M. BOCQUET, Y. IGARASHI, T. MAKI et D. DIDIER : A new perspective on the Fukushima releases brought by newly available ^{137}Cs air concentration observations and reliable meteorological fields. In *17th International Conference on Harmonisation within Atmospheric Dispersion Modelling for Regulatory Purposes, 9-12 May 2016, Budapest, Hungary*, 2016. (Cité en page 133.)
- V. SAVCENCO, W. HUNSDORFER et J. VERWER : A multirate time stepping strategy for stiff ordinary differential equations. *BIT Numerical Mathematics*, 47:137–155, 2007. (Cité en page 14.)
- C. SEIGNEUR : Environnement atmosphérique et qualité de l’air. Notes de cours, École des Ponts ParisTech, 2012. URL http://cerea.enpc.fr/fr/cours-web/POLU1_2012-2013.html. Cours de deuxième année. (Non cité.)

- N. SEMANE, V.-H. PEUCH, S. PRADIER, G. DESROZIER, L. EL AMRAOUI, P. BROUSSEAU, S. MASSART, B. CHAPNIK et A. PEUCH : On the extraction of wind information from the assimilation of ozone profiles in Météo-France 4-D-Var operational NWP suite. *Atmos. Chem. Phys.*, 9:4855–4867, 2009. (Cité en pages 92 et 127.)
- D. SIMEOENS : AirBase version 7 data products on EEA data service. Rap. tech., European Environment Agency, February 2013. URL <http://www.eea.europa.eu/data-and-maps/data/airbase-the-european-air-quality-database-7>. (Cité en page 19.)
- D. SIMPSON, W. WINIWARTER, G. BÖRJESSON, S. CINDERBY, A. FERREIRO, A. GUENTHER, C. N. HEWITT, R. JANSON, M. A. K. KHALIL, S. OWEN, T. E. PIERCE, H. PUXBAUM, M. SHEARER, U. SKIBA, R. STEINBRECHER, L. TARRASÓN et M. G. ÖQUIST : Inventorying emissions from nature in Europe. *J. Geophys. Res.*, 104:8113–8152, 1999. (Cité en page 9.)
- W. C. SKAMAROCK, J. B. KLEMP, J. DUDHIA, D. O. GILL, D. M. BARKER, M. G. DUDA, X.-Y. HUANG, W. WANG et J. G. POWERS : A description of the advanced research WRF version 3. Rap. tech., NCAR, 2008. (Cité en page 10.)
- L. SLIVINSKI et C. SNYDER : Exploring practical estimates of the ensemble size necessary for particle filters. *Mon. Weather Rev.*, 144:861–875, 2016. (Cité en page 51.)
- C. SNYDER, T. BENGTTSSON et M. MORZFELD : Performance bounds for particle filters using the optimal proposal. *Mon. Weather Rev.*, 143:4750–4761, 2015. (Cité en page 51.)
- B. SPORTISSE : *Pollution atmosphérique : des processus à la modélisation*. Ingénierie et développement durable. Springer-Verlag Paris, 2008. (Non cité.)
- B. SPORTISSE et V. MALLET : Calcul scientifique pour l’environnement. Notes de cours, ENSTA ParisTech, 2005. URL <http://vivienmallet.net/static/publication/sportisse07calcul.pdf>. Cours de deuxième année. (Cité en page 11.)
- W. R. STOCKWELL, F. KIRCHNER, M. KUHN et S. SEEFELD : A new mechanism for regional atmospheric chemistry modeling. *J. Geophys. Res.*, 102:25847–25879, doi :10.1029/97JD00849, 1997. (Cité en page 72.)
- A. STOHL, P. SEIBERT, G. WOTAWA, D. ARNOLD, J. F. BURKHART, S. ECKHARDT, C. TAPIA, A. VARGAS et T. J. YASUNARI : Xenon-133 and caesium-137 releases into the atmosphere from the Fukushima Dai-ichi nuclear power plant : determination of the source term, atmospheric dispersion, and deposition. *Atmos. Chem. Phys.*, 12:2313–2343, 2012. (Cité en page 133.)
- G. STRANG : On the construction and comparison of difference schemes. *SIAM J. Numer. Anal.*, 5:506–517, 1968. (Cité en page 14.)
- N. TALERKO : Mesoscale modelling of radioactive contamination formation in Ukraine caused by the Chernobyl accident. *J. Environ. Radioactivity*, 78:311–329, 2005. (Cité en page 143.)
- P. TANDEO, M. PULIDO et F. LOTT : Offline parameter estimation using EnKF and maximum likelihood error covariance estimates : Application to a subgrid-scale orography parametrization. *Q. J. R. Meteorol. Soc.*, 141:383–395, 2015. (Cité en page 139.)

- H. TERADA, G. KATATA, M. CHINO et H. NAGAI : Atmospheric discharge and dispersion of radionuclides during the Fukushima Dai-ichi Nuclear Power Plant accident. Part II : verification of the source term and analysis of regional-scale atmospheric dispersion. *J. Environ. Radioact.*, 112:141–154, 2012. (Cité en pages 133 et 143.)
- O. TICHÝ, V. ŠMÍDL, R. HOFMAN et A. STOHL : LS-APC v1.0 : a tuning-free method for the linear inverse problem and its application to source-term determination. *Geosci. Model Dev.*, 9:4297–4311, 2016. (Cité en page 134.)
- M. K. TIPPETT, J. L. ANDERSON, C. H. BISHOP, T. M. HAMILL et J. S. WHITAKER : Ensemble square root filters. *Mon. Weather Rev.*, 131:1485–1490, 2003. (Cité en page 26.)
- R. TODLING : A lag-1 smoother approach to system-error estimation : sequential method. *Q. J. R. Meteorol. Soc.*, 141:1502–1513, 2015a. (Cité en pages 55 et 57.)
- R. TODLING : A complementary note to ‘A lag-1 smoother approach to system-error estimation’ : the intrinsic limitations of residual diagnostics. *Q. J. R. Meteorol. Soc.*, 141:2917–2922, 2015b. (Cité en page 57.)
- Y. TRÉMOLET : Accounting for an imperfect model in 4D-Var. *Q. J. R. Meteorol. Soc.*, 132:2483–2504, 2006. (Cité en page 48.)
- I. B. TROEN et L. MAHRT : A simple model of the atmospheric boundary layer ; sensitivity to surface evaporation. *Bound.-Layer Meteorol.*, 37:129–148, 1986. (Cité en page 6.)
- H. TSURUTA, Y. OURA, M. EBIHARA, T. OHARA et T. NAKAJIMA : First retrieval of hourly atmospheric radionuclides just after the Fukushima accident by analyzing filter-tapes of operational air pollution monitoring stations. *Sci. Rep.*, 4:6717, 2014. (Cité en page 133.)
- G. UENO et N. NAKAMURA : Iterative algorithm for maximum-likelihood estimation of the observation-error covariance matrix for ensemble-based filters. *Q. J. R. Meteorol. Soc.*, 140:295–315, 2014. (Cité en page 139.)
- G. UENO et N. NAKAMURA : Bayesian estimation of the observation-error covariance matrix in ensemble-based filters. *Q. J. R. Meteorol. Soc.*, 142:2055–2080, 2016. (Cité en page 139.)
- UNITED NATIONS : Sources and effects of ionizing radiation. United Nations scientific committee on the effects of atomic radiation. Report to general assembly. Rap. tech., United Nations, New-York, 2000. (Cité en pages 133 et 141.)
- UNITED NATIONS : Developments since the 2013 UNSCEAR report on the levels and effects of radiation exposure due to the nuclear accident following the great east-Japan earthquake and tsunami. Rap. tech., United Nations, New-York, 2016. URL http://www.unscear.org/docs/publications/2016/UNSCEAR_WP_2016.pdf. (Cité en page 133.)
- S. M. UPPALA, P. W. KÅLLBERG, A. J. SIMMONS, U. ANDRAE, V. D. C. BECHTOLD, M. FIORINO, J. K. GIBSON, J. HASELER, A. HERNANDEZ, G. A. KELLY, X. LI, K. ONOGI, S. SAARINEN, N. SOKKA, R. P. ALLAN, E. ANDERSSON, K. ARPE, M. A. BALMASEDA, A. C. M. BELJAARS, L. V. D. BERG, J. BIDLOT, N. BORMANN, S. CAIRES, F. CHEVALLIER, A. DETHOF, M. DRAGOSAVAC, M. FISHER, M. FUENTES, S. HAGEMANN, E. HÓLM, B. J. HOSKINS, L. ISAKSEN, P. A. E. M. JANSSEN, R. JENNE, A. P. MCNALLY, J.-F. MAHFOUF, J.-J. MORCRETTE, N. A. RAYNER, R. W. SAUNDERS, P. SIMON, A. STERL, K. E. TRENBERTH, A. UNTCH, D. VASILJEVIC,

- P. VITERBO et J. WOOLLEN : The ERA-40 re-analysis. *Q. J. R. Meteorol. Soc.*, 131:2961–3012, 2005. (Cité en page 135.)
- P. J. van LEEUWEN : Nonlinear data assimilation in geosciences : an extremely efficient particle filter. *Q. J. R. Meteorol. Soc.*, 136:1991–1999, 2010. (Cité en page 61.)
- P. J. van LEEUWEN et G. EVENSEN : Data assimilation and inverse methods in terms of a probabilistic formulation. *Mon. Weather Rev.*, 124:2898–2913, 1996. (Cité en page 34.)
- A. VENKATRAM, P. KARAMCHANDANI, P. PAI et R. GOLDSTEIN : The development and application of a simplified ozone modeling system (SOMS). *Atmos. Environ.*, 28:3665–3678, 1994. (Cité en page 64.)
- V. VESTRENG, M. ADAMS et J. GOODWIN : Inventory review 2004. emission data reported to CLRTAP and the NEC Directive. Rap. tech., EMEP/EEA, 2004. ISSN 0804-2446. (Cité en pages 9 et 102.)
- A. VISSCHEDIJK, P. ZANDVELD et H. Denier van der GON : A high resolution gridded European emission database for the EU integrated project gems. *TNO, Apeldoorn, Netherlands, TNO-report*, 2007. (Cité en page 9.)
- C. R. VOGEL : *Computational Methods for Inverse Problems*. SIAM, Frontiers in Applied Mathematics, 2002. (Cité en page 138.)
- A. VOULGARAKIS, N. SAVAGE, O. WILD, G. CARVER, K. CLEMITSHAW et J. PYLE : Upgrading photolysis in the p-TOMCAT CTM : model evaluation and assessment of the role of clouds. *Geosci. Model Dev.*, 2:59–72, 2009. (Cité en page 67.)
- G. C. G. WEI et M. A. TANNER : A Monte Carlo implementation of the EM algorithm and the poor man’s data augmentation algorithms. *Journal of the American Statistical Association*, 85:699–704, 1990. (Cité en page 139.)
- B. P. WELFORD : Note on a method for calculating corrected sums of squares and products. *Technometrics*, 4:419–420, 1962. (Cité en page 145.)
- M. L. WESELY : Parameterization of surface resistances to gaseous dry deposition in regional-scale numerical models. *Atmos. Environ.*, 23:1293–1304, 1989. (Cité en page 6.)
- J. S. WHITAKER et T. M. HAMILL : Ensemble data assimilation without perturbed observations. *Mon. Weather Rev.*, 130:1913–1924, 2002. (Cité en page 61.)
- J. S. WHITAKER et T. M. HAMILL : Evaluating methods to account for system errors in ensemble data assimilation. *Mon. Weather Rev.*, 140:3078–3089, 2012. (Cité en page 37.)
- J. S. WHITAKER, T. M. HAMILL, X. WEI, Y. SONG et Z. TOTH : Ensemble data assimilation with the NCEP global forecast system. *Mon. Weather Rev.*, 136:463–482, February 2008. (Cité en page 55.)
- V. WINIAREK, M. BOCQUET, N. DUHANYAN, Y. ROUSTAN, O. SAUNIER et A. MATHIEU : Estimation of the caesium-137 source term from the Fukushima Daiichi nuclear power plant using a consistent joint assimilation of air concentration and deposition observations. *Atmos. Environ.*, 82:268–279, 2014. (Cité en pages 131, 133, 134, 135, 137, 143 et 146.)

- V. WINIAREK, M. BOCQUET, O. SAUNIER et A. MATHIEU : Estimation of errors in the inverse modeling of accidental release of atmospheric pollutant : Application to the reconstruction of the cesium-137 and iodine-131 source terms from the Fukushima Daiichi power plant. *J. Geophys. Res.*, 117:D05122, 2012. (Cité en pages 133, 134, 137 et 143.)
- L. WU, M. BOCQUET, F. CHEVALLIER, T. LAUVAUX et K. DAVIS : Hyperparameter estimation for uncertainty quantification in mesoscale carbon dioxide inversions. *Tellus B*, 65:20894, 2013. (Cité en page 137.)
- L. WU, V. MALLET, M. BOCQUET et B. SPORTISSE : A comparison study of data assimilation algorithms for ozone forecasts. *J. Geophys. Res.*, 113:D20310, 2008. (Cité en pages 93, 106, 107, 126 et 127.)
- G. YARWOOD, S. RAO, M. YOCKE et G. WHITTEN : Updates to the carbon bond chemical mechanism : CB05. *Final report to the US EPA, RT-0400675*, 8, 2005. (Cité en pages 13 et 102.)
- E. YEE : Theory for reconstruction of an unknown number of contaminant sources using probabilistic inference. *Boundary-Layer Meteorology*, 127:359–394, 2008. (Cité en pages 138 et 150.)
- K. YUMIMOTO, Y. MORINO, T. OHARA, Y. OURA, M. EBIHARA, H. TSURUTA et T. NAKAJIMA : Inverse modeling of the ^{137}Cs source term of the Fukushima Dai-ichi Nuclear Power Plant accident constrained by a deposition map monitored by aircraft. *J. Environ. Radioact.*, 164:1–12, 2016. (Cité en page 133.)
- L. ZHANG, J. R. BROOK et R. VET : A revised parameterization for gaseous dry deposition in air-quality models. *Atmos. Chem. Phys.*, 3:2067–2082, 2003. (Cité en pages 6, 102, 103 et 116.)
- Y. ZHANG, M. BOCQUET, V. MALLET, C. SEIGNEUR et A. BAKLANOV : Real-time air quality forecasting, part II : State of the science, current research needs, and future prospects. *Atmos. Environ.*, 60:656–676, 2012. (Cité en page 131.)
- M. ZUPANSKI : Maximum likelihood ensemble filter : Theoretical aspects. *Mon. Weather Rev.*, 133:1710–1726, 2005. (Cité en page 30.)

Summary

Data assimilation methods are constantly evolving to adapt to the various application domains. In atmospheric sciences, each new algorithm has first been implemented on numerical weather prediction models before being ported to atmospheric chemistry models. It has been the case for 4D variational methods and ensemble Kalman filters for instance. The new 4D ensemble variational methods (4D EnVar) are no exception. They were developed to take advantage of both variational and ensemble approaches and they are starting to be used in operational weather prediction centers, but have yet to be tested on operational atmospheric chemistry models.

The validation of new data assimilation methods on these models is indeed difficult because of the complexity of such models. It is hence necessary to have at our disposal low-order models capable of synthetically reproducing key physical phenomena from operational models while limiting some of their hardships. Such a model, called L95-GRS, has therefore been developed. It combines the simple meteorology from the Lorenz-95 model to a tropospheric ozone chemistry module with 7 chemical species. Even though it is of low dimension, it reproduces some of the physical and chemical phenomena observable in real situations. A data assimilation method, the iterative ensemble Kalman smoother (IEnKS), has been applied to this model. It is an iterative 4D EnVar method which solves the full non-linear variational problem. This application validates 4D EnVar methods in the context of non-linear atmospheric chemistry, but also raises the first limits of such methods, most noticeably when they are applied to weakly coupled stable models.

After this experiment, results have been extended to a realistic atmospheric pollution prediction model. 4D EnVar methods, via the IEnKS, have once again shown their potential to take into account the non-linearity of the chemistry model in a controlled environment, with synthetic observations. However, the assimilation of real tropospheric ozone concentrations mitigates these results and shows how hard atmospheric chemistry data assimilation is. A strong model error is indeed attached to these models, stemming from multiple uncertainty sources. Two steps must be taken to tackle this issue.

First of all, the data assimilation method used must be able to efficiently take into account the model error. However, most methods are developed under the assumption of a perfect model. To avoid this hypothesis, a new method has then been developed. Called IEnKF-Q, it expands the IEnKS to the model error framework. It has been validated on a low-order model, proving its superiority over data assimilation methods naively adapted to take into account model error.

Nevertheless, such methods need to know the exact nature and amplitude of the model error which needs to be accounted for. Therefore, the second step is to use statistical tools to quantify this model error. The expectation-maximization algorithm, the naive and unbiased randomize-then-optimize algorithms, an importance sampling based on a Laplace proposal, and a Markov chain Monte Carlo simulation, potentially transdimensional, have been assessed, expanded, and compared to estimate the uncertainty on the retrieval of the source term of the Chernobyl and Fukushima-Daiichi nuclear power plant accidents.

This thesis therefore improves the domain of 4D EnVar data assimilation by its methodological input and by paving the way to applying these methods on atmospheric chemistry models.

Résumé

Les méthodes d'assimilation de données sont en constante évolution pour s'adapter aux problèmes à résoudre dans les multiples domaines d'application. En sciences de l'atmosphère, chaque nouvel algorithme a d'abord été implémenté sur des modèles de prévision numérique du temps avant d'être porté sur des modèles de chimie atmosphérique. Ce fut le cas des méthodes variationnelles 4D et des filtres de Kalman d'ensemble par exemple. La nouvelle génération d'algorithmes variationnels d'ensemble quadridimensionnels (EnVar 4D) ne fait pas exception. Elle a été développée pour tirer partie des deux approches variationnelle et ensembliste et commence à être appliquée au sein des centres opérationnels de prévision numérique du temps, mais n'a à ce jour pas été testée sur des modèles opérationnels de chimie atmosphérique.

En effet, la complexité de ces modèles rend difficile la validation de nouvelles méthodes d'assimilation. Il est ainsi nécessaire d'avoir à disposition des modèles d'ordre réduit, qui doivent être en mesure de synthétiser les phénomènes physiques à l'œuvre dans les modèles opérationnels tout en limitant certaines des difficultés liées à ces derniers. Un tel modèle, nommé L95-GRS, a donc été développé. Il associe la météorologie simpliste du modèle de Lorenz-95 à un module de chimie de l'ozone troposphérique avec 7 espèces chimiques. Bien que de faible dimension, il reproduit des phénomènes physiques et chimiques observables en situation réelle. Une méthode d'assimilation de donnée, le lisseur de Kalman d'ensemble itératif (IEnKS), a été appliquée sur ce modèle. Il s'agit d'une méthode EnVar 4D itérative qui résout le problème non-linéaire variationnel complet. Cette application a permis de valider les méthodes EnVar 4D dans un contexte de chimie atmosphérique non-linéaire, mais aussi de soulever les premières limites de telles méthodes, notamment lorsqu'elles sont appliquées sur un modèle faiblement couplé stable.

Fort de cette expérience, les résultats ont été étendus au cas d'un modèle réaliste de prévision de pollution atmosphérique. Les méthodes EnVar 4D, via l'IEnKS, ont montré leur potentiel pour tenir compte de la non-linéarité du modèle de chimie dans un contexte maîtrisé, avec des observations synthétiques. Cependant, le passage à des observations réelles d'ozone troposphérique mitige ces résultats et montre la difficulté que représente l'assimilation de données en chimie atmosphérique. En effet, une très forte erreur est associée à ces modèles, provenant de sources d'incertitudes variées. Deux démarches doivent alors être entreprises pour pallier ce problème.

Tout d'abord, la méthode d'assimilation doit être en mesure de tenir compte efficacement de l'erreur modèle. Cependant, la majorité des méthodes sont développées en supposant au contraire un modèle parfait. Pour se passer de cette hypothèse, une nouvelle méthode a donc été développée. Nommée IEnKF-Q, elle étend l'IEnKS au cas avec erreur modèle. Elle a été validée sur un modèle jouet, démontrant sa supériorité par rapport à des méthodes d'assimilation adaptées naïvement pour tenir compte de l'erreur modèle.

Toutefois, une telle méthode nécessite de connaître la nature et l'amplitude exacte de l'erreur modèle qu'elle doit prendre en compte. Aussi, la deuxième démarche consiste à recourir à des outils statistiques pour quantifier cette erreur modèle. Les algorithmes d'espérance-maximisation, de *randomize-then-optimize* naïf et sans biais, un échantillonnage préférentiel fondé sur l'approximation de Laplace, ainsi qu'un échantillonnage avec une méthode de Monte-Carlo par chaînes de Markov, y compris transdimensionnelle, ont ainsi été évalués, étendus et comparés pour estimer l'incertitude liée à la reconstruction du terme source des accidents des centrales nucléaires de Tchernobyl et Fukushima-Daiichi.

Cette thèse a donc enrichi le domaine de l'assimilation de données EnVar 4D par ses apports méthodologiques et en ouvrant la voie à l'application de ces méthodes sur les modèles de chimie atmosphérique.

