



Estimation par maximum de vraisemblance dans des problèmes inverses non linéaires

Estelle Kuhn

► To cite this version:

Estelle Kuhn. Estimation par maximum de vraisemblance dans des problèmes inverses non linéaires. Mathématiques [math]. Université Paris Sud - Paris XI, 2003. Français. NNT: . tel-00008316

HAL Id: tel-00008316

<https://theses.hal.science/tel-00008316>

Submitted on 1 Feb 2005

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

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

ORSAY
N° D'ORDRE : 7408

UNIVERSITÉ DE PARIS SUD
U.F.R SCIENTIFIQUE D'ORSAY

THÈSE

présentée
pour obtenir

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

par

Estelle KUHN

Sujet : ESTIMATION PAR MAXIMUM DE VRAISEMBLANCE DANS
DES PROBLÈMES INVERSES NON LINÉAIRES

Rapporteurs : M. AZAIS Jean-Marc
M. MOULINES Eric

Soutenue le 12 décembre 2003 devant le jury composé de :

M.	AZAIS Jean-Marc	Rapporteur
M.	LAVIELLE Marc	Directeur de Thèse
M.	MASSART Pascal	Président
Mme.	MENTRÉ France	Examinatrice
M.	MOULINES Eric	Rapporteur
M.	ROBERT Christian	Examineur

Merci,

A Marc, Muchas Gracias pour avoir grand ouvert il y a maintenant plus de trois ans la porte que d'autres avaient entrebâillée, pour m'avoir encouragée à me lancer et pour m'avoir transmis tout au long de ces trois années son goût pour les applications.

A tous les membres du jury pour leurs présences aujourd'hui. Je tiens également à remercier Jean-Marc Azais et Eric Moulines, pour avoir accepté de rapporter ce travail, France Mentré pour sa collaboration dynamique et chaleureuse, Christian Robert, pour m'avoir accompagnée avec d'autres dans mes premiers pas d'enseignante à l'Université Dauphine, et Pascal Massart, pour m'avoir fait l'honneur d'accepter la présidence, ainsi que pour toutes les "cinq minutes" qu'il m'a accordées au cours de ces années passées à Orsay.

A Hermann Zeyen et Adeline Leclercq, pour m'avoir permis de réaliser de belles applications de la théorie dans leurs domaines respectifs.

Aux membres du laboratoire de mathématiques d'Orsay, plus particulièrement aux âmes du 430, au trio SOS informatique de notre équipe et à tout ceux qui un jour ou l'autre m'ont accordé un peu de leur temps pour une interrogation mathématique profonde ou pour une discussion moins conventionnelle, merci!

Aux membres du CEREMADE de l'Université Dauphine qui m'ont accueillie comme monitrice et comme ATER, plus particulièrement Martine Bellec, Claudine Dhuin et Anne-Marie Boussion avec qui j'ai pu partager mon goût pour l'enseignement.

A Jacques Chaumat, le principal entrebâilleur de porte, à Catherine Matias et Marie-Luce Taupin, pour leur aide précieuse et leur soutien acharné dans les moments difficiles (eh oui! il y en a eu...).

Aux thésards passés et présents de l'Ecole doctorale d'Orsay que j'ai cotoyés au quotidien pendant ces trois années et qui ont contribué par leur humour et leur enthousiasme à rendre ces moments plus qu'agréables... Antoine, pour m'avoir encouragée comme il l'a fait du début à la fin, Vincent, pas seulement pour son côté algébriste-contractant, Magalie, pour toutes les émotions que nous avons vécues quasi en parallèle depuis six ans maintenant, Violaine, pour les tête-à-tête nombreux et variés que nous avons partagés et pour ceux que nous partagerons, Béatrice et Réda, mes petits Suisses adorés, Catherine (bis!), pour sa présence constante mais discrète à mes côtés et pour sa franchise de chaque instant.

A mes parents, à mes grand-parents, à Guillaume et à Gene, pour leur sollicitude.

A Yann, pour m'avoir fait confiance dans ce que j'ai entrepris et pour avoir su me soutenir au bon moment,

A Salomé, ♡.

Table des matières

1	Coupling a stochastic approximation version of EM with an MCMC procedure	25
1.1	Introduction	26
1.2	Coupling the SAEM algorithm with MCMC	28
1.3	Convergence result	28
1.4	Applications	33
1.4.1	The deconvolution problem	33
1.4.2	The change-points problem	35
1.5	Appendix	38
2	Maximum likelihood estimation in nonlinear mixed effects models	41
2.1	Introduction	42
2.2	Algorithms proposed for maximum likelihood estimation	44
2.2.1	The EM algorithm	44
2.2.2	A stochastic version of the EM algorithm	45
2.3	Example of the orange trees	50
2.3.1	The model	50
2.3.2	The EM and SAEM algorithms	50
2.3.3	Extension to an heteroscedastic model	55
2.4	Comparisons with other methods used in PK/PD.	57
2.4.1	A first pharmacokinetic model	58
2.4.2	A pharmacodynamic model	61
2.4.3	A second pharmacokinetic model	62
2.5	Application to likelihood ratio test. ¹	65
2.5.1	Framework	65
2.5.2	The model	66
2.5.3	The likelihood ratio test	66
2.5.4	Estimation of the likelihood	66
2.5.5	Numerical application	67
3	Joint inversion of teleseismic delay times and gravity anomaly data: A new approach.²	73
3.1	Introduction	74

¹Joint work with France Mentré and Adeline Leclercq, INSERM.

²Joint work with Hermann Zeyen, Département des Sciences de la Terre, U.P.S., Orsay.

3.2	Method	75
3.2.1	Presentation of the probabilistic model	75
3.2.2	A stochastic algorithm for maximum likelihood estimation	75
3.2.3	Application to the Gaussian linear model	76
3.3	Application to the gravimetric and teleseismic problems	78
3.3.1	The gravimetric model	78
3.3.2	The teleseismic model	79
3.3.3	Joint inversion of gravimetric and teleseismic data	79
3.4	Synthetic model examples	81
3.5	Discussion	88
3.6	Appendix	89
4	Logspline density estimation by maximum likelihood in nonlinear inverse problem	91
4.1	Introduction	92
4.2	Logspline model	93
4.3	The SAEM algorithm	95
4.4	Convergence result	96
4.5	Applications	100
4.5.1	Application with missing data normally distributed	100
4.5.2	Application involving estimation of other parameters	104
4.5.3	Application with missing data drawn from a bimodal distribution	104
4.5.4	Application to the example of the orange trees	107
4.6	Appendix	107

Introduction

Sommaire

Cette thèse est consacrée à l'estimation par maximum de vraisemblance dans des problèmes inverses. Nous considérons des modèles statistiques à données manquantes, dans un cadre paramétrique au cours des trois premiers chapitres et dans un cadre non paramétrique dans le dernier chapitre. L'approche adoptée est basée sur une version stochastique de l'algorithme d'estimation EM (*Expectation Maximization*). Nous abordons à la fois le point de vue théorique de la convergence de l'algorithme vers l'estimateur du maximum de vraisemblance et les applications pratiques. Dans le cadre non paramétrique, nous considérons également le comportement asymptotique de l'estimateur.

Dans le premier chapitre, nous proposons une variante de l'algorithme SAEM (*Stochastic Approximation EM*) en le combinant à une méthode de Monte Carlo par chaînes de Markov. Les données manquantes ne sont plus simulées selon la loi a posteriori, mais selon la probabilité de transition d'une chaîne de Markov adéquate. Nous prouvons la convergence presque sûre de la suite (θ_k) générée par l'algorithme vers un maximum local de la vraisemblance des observations. Nous présentons également deux applications, la première à un problème de déconvolution, la seconde à un problème de détection de ruptures.

Dans le second chapitre, nous appliquons cette version de l'algorithme SAEM au cas particulier des modèles non linéaires à effets mixtes. Nous effectuons outre l'estimation des paramètres du modèle, des estimations de la vraisemblance du modèle et de l'information de Fisher. Nous illustrons sur des exemples choisis dans la littérature les performances de l'algorithme en effectuant des comparaisons avec d'autres méthodes. Nous nous intéressons plus particulièrement au domaine de la pharmacocinétique et de la pharmacodynamique.

Le troisième chapitre présente une application de l'algorithme SAEM à un problème inverse issu de la géophysique. Il s'agit d'effectuer deux inversions, la première entre les temps de parcours des ondes sismiques et leurs vitesses, la deuxième entre des mesures de gravimétrie effectuées en surface et les densités du sous-sol. Notre démarche est novatrice dans ce domaine car les paramètres du modèle sont estimés alors qu'ils étaient en général fixés arbitrairement. De plus, nous prenons en compte au cours de cette inversion jointe une relation linéaire entre les densités et les vitesses des ondes. Nous présentons des applications à des données simulées.

Le dernier chapitre propose une version non paramétrique de l'algorithme SAEM couplé à une méthode MCMC permettant d'estimer la densité π des données manquantes de façon non paramétrique. Nous exhibons un estimateur logspline $\pi_{\hat{\theta}}$ de cette densité qui maximise la vraisemblance des observations dans un modèle logspline de dimension J . Puis, nous nous intéressons à la convergence de cet estimateur $\pi_{\hat{\theta}}$ vers π lorsque la dimension J du modèle logspline et le nombre d'observations tendent vers l'infini. Nous présentons des applications à des données simulées et réelles.

Cadre et objectifs du problème paramétrique

On considère un problème statistique à données manquantes dans le cadre général suivant : les données observées sont notées y , les données non observées (ou manquantes) sont notées z et on note (y, z) les données complètes. On suppose que les données observées y et les données non observées z prennent respectivement leurs valeurs dans l'ensemble $\mathcal{Y}$, sous-ensemble de $\mathbb{R}^n$, et dans $\mathbb{R}^l$. Les dimensions n et l des vecteurs des observations et des données manquantes sont fixées au cours de ce travail. On suppose également que la vraisemblance des données complètes

(y, z) par rapport à la mesure de Lebesgue μ sur $\mathbb{R}^l$ (ou plus généralement une mesure borélienne σ -finie) appartient à une famille paramétrique $\{f(\cdot; \theta), \theta \in \Theta\}$ où Θ est un sous-ensemble de $\mathbb{R}^p$. La vraisemblance des observations y est donc égale à :

$$g(y; \theta) \triangleq \int_{\mathbb{R}^l} f(y, z; \theta) \mu(dz). \quad (1)$$

Pour estimer la valeur exacte du paramètre θ_e , on peut utiliser l'estimateur du maximum de vraisemblance, qui possède de bonnes propriétés asymptotiques (lorsque n tend vers l'infini) dans des modèles suffisamment réguliers. Notre but sera donc de déterminer une valeur $\hat{\theta}$ dans Θ qui maximise la log-vraisemblance des observations (égale à $\log g$) pour un jeu d'observations y fixé. Cette valeur n'est pas nécessairement unique, la log-vraisemblance pouvant avoir plusieurs maxima.

Ce type de modèle intervient dans de nombreux domaines comme par exemple en traitement du signal pour la déconvolution, la séparations de sources et la détection de ruptures, en tomographie, pour la tomographie par émission de positrons (*TEP*) en imagerie médicale ou pour la tomographie sismique en géophysique, illustrée par le Chapitre 3, ou en pharmacologie dont quelques applications sont présentées dans le Chapitre 2.

Une première approche : l'algorithme EM

Si les données z étaient observées, la valeur de θ qui nous intéresserait dans le cadre d'une estimation par maximum de vraisemblance serait $\arg \max_{\theta \in \Theta} \log f(y, z; \theta)$. Lorsque les données z ne sont pas observables, on peut considérer à la place de $\log f(y, z; \theta)$ son espérance conditionnelle par rapport aux données observées y . A partir de cette remarque, Dempster, Laird et Rubin (1977) ont proposé l'algorithme EM (*Expectation Maximization*).

Précisons les motivations et le fonctionnement de l'algorithme EM. Comme la valeur du paramètre θ est inconnue, on ne peut pas calculer directement l'espérance conditionnelle suivante $E[\log f(y, z; \theta) | y; \theta]$. On s'intéresse donc à la quantité :

$$Q(\theta' | \theta) \triangleq E[\log f(y, z; \theta') | y; \theta],$$

qui possède la propriété remarquable suivante :

Proposition. *Pour tout (θ, θ') dans Θ^2 , si $Q(\theta' | \theta) \geq Q(\theta | \theta)$, alors $l(\theta') \geq l(\theta)$, où l désigne la log-vraisemblance des observations.*

Remarque. *La dépendance des fonctions l et Q en les observations y n'est pas mentionnée dans les notations.*

La démonstration de cette proposition repose sur l'inégalité de Jensen (voir Lemme 1 et Théorème 1 de l'article de Dempster, Laird et Rubin (1977)). Cette propriété implique que tout accroissement de Q engendre un accroissement de l . Ainsi lorsque la maximisation de Q est plus simple que la maximisation directe de l , on peut procéder à des maximisations successives de Q pour espérer atteindre un maximum de l . Dempster, Laird et Rubin (1977) proposent donc l'algorithme itératif à deux étapes suivant : lors de la première étape notée E, on calcule l'espérance conditionnelle (sous la valeur courante du paramètre) de la log-vraisemblance complète; lors de la seconde étape notée M, on maximise cette quantité en θ . Ainsi, pour une

initialisation arbitraire θ_0 , la k -ième itération de l'algorithme consiste à évaluer lors de l'étape E la quantité :

$$Q(\theta|\theta_k) = E[\log f(y, z; \theta)|y; \theta_k], \quad (2)$$

où θ_k est la valeur courante du paramètre obtenue à la fin de l'itération précédente. Puis, lors de l'étape M, la valeur courante du paramètre est réactualisée de la façon suivante :

$$\theta_{k+1} = \arg \max_{\theta \in \Theta} Q(\theta|\theta_k). \quad (3)$$

La suite (θ_k) définie par l'algorithme EM induit donc une suite $(l(\theta_k))$ croissante. Plusieurs auteurs ont étudié la convergence de la suite (θ_k) . On peut citer de manière non exhaustive les travaux de Dempster, Laird et Rubin (1977), améliorés quelques années plus tard par ceux de Wu (1983). Ainsi, sous des hypothèses générales de régularité du modèle, Wu (1983) prouve la convergence de la suite (θ_k) vers un point stationnaire de la log-vraisemblance l des observations. La convergence de l'algorithme EM vers un maximum local de la log-vraisemblance l des observations est assurée par Wu (1983) sous des hypothèses assez contraignantes sur la régularité du modèle et sur l'ensemble des points stationnaires de l .

Nous présentons maintenant un résultat de convergence de l'algorithme EM proposé par Delyon, Lavielle et Moulines (1999), qui obtiennent pour le cadre particulier des modèles de type exponentiel des hypothèses plus simples que celles de Wu (1983). Dans la suite, nous noterons $p(z|y; \theta)$ la densité des données manquantes z conditionnellement aux données observées y :

$$p(z|y; \theta) \triangleq \begin{cases} f(y, z; \theta)/g(y; \theta) & \text{si } g(y; \theta) \neq 0 \\ 0 & \text{si } g(y; \theta) = 0 \end{cases}.$$

Delyon, Lavielle et Moulines (1999) supposent que le modèle vérifie les hypothèses suivantes :

- (M1) L'espace des paramètres Θ est un sous-ensemble ouvert de $\mathbb{R}^p$. La vraisemblance complète du modèle est donnée par l'expression suivante :

$$f(y, z; \theta) = \exp \left\{ -\psi(\theta) + \left\langle \tilde{S}(y, z), \phi(\theta) \right\rangle \right\}, \quad (4)$$

où $\langle \cdot, \cdot \rangle$ désigne le produit scalaire dans $\mathbb{R}^m$, $\tilde{S}$ est une fonction borélienne en la seconde variable sur $\mathbb{R}^l$ à valeurs dans un sous-ensemble ouvert $\mathcal{S}$ de $\mathbb{R}^m$. On désignera $\tilde{S}(y, z)$ comme une statistique exhaustive du modèle. De plus, on suppose que l'enveloppe convexe de $\tilde{S}(\mathbb{R}^l)$ est incluse dans $\mathcal{S}$, et que pour tout θ dans Θ ,

$$\int_{\mathbb{R}^l} |\tilde{S}(y, z)| p(z|y; \theta) \mu(dz) < \infty.$$

Remarque. Il n'y a pas unicité de la statistique exhaustive, qui est toujours définie à une constante près.

- (M2) Les fonctions ψ et ϕ sont deux fois continûment différentiables sur Θ .
- (M3) La fonction $\bar{s} : \Theta \rightarrow \mathcal{S}$ définie comme suit :

$$\bar{s}(\theta) \triangleq \int_{\mathbb{R}^l} \tilde{S}(y, z) p(z|y; \theta) \mu(dz)$$

est continûment différentiable sur Θ . Remarquons que la dépendance en y de la fonction $\bar{s}$ n'est pas notée.

- (M4) La log-vraisemblance des données observées l est continûment différentiable sur Θ et

$$\partial_\theta \int f(y, z; \theta) \mu(dz) = \int \partial_\theta f(y, z; \theta) \mu(dz),$$

où ∂_θ désigne l'opérateur de différentiation par rapport à θ .

- (M5) On définit la fonction $L : \mathcal{S} \times \Theta \rightarrow \mathbb{R}$ de la manière suivante :

$$L(s; \theta) \triangleq -\psi(\theta) + \langle s, \phi(\theta) \rangle,$$

alors il existe une fonction $\hat{\theta} : \mathcal{S} \rightarrow \Theta$, telle que :

$$\forall s \in \mathcal{S}, \quad \forall \theta \in \Theta, \quad L(s; \hat{\theta}(s)) \geq L(s; \theta).$$

De plus, la fonction $\hat{\theta}$ est continûment différentiable sur $\mathcal{S}$.

On introduit également l'ensemble $\mathcal{L}$ des points stationnaires de la log-vraisemblance l :

$$\mathcal{L} = \{\theta \in \Theta, \partial_\theta l(\theta) = 0\},$$

et on désigne par $d(x, A)$ la distance euclidienne du point x à l'ensemble fermé A . Le résultat de convergence de la suite (θ_k) s'énonce alors ainsi :

Théorème 1 (Delyon, Lavielle et Moulines (1999)). *Supposons vérifiées les hypothèses (M1)-(M5). Alors pour tout point initial θ_0 dans Θ , la suite $(l(\theta_k))$ obtenue par l'algorithme EM est croissante et la suite (θ_k) vérifie $\lim_{k \rightarrow \infty} d(\theta_k, \mathcal{L}) = 0$.*

L'utilisation pratique de cet algorithme pose plusieurs types de problèmes : d'une part, ceux relatifs aux limites obtenues pour la suite (θ_k) , d'autre part, ceux relatifs à la mise en oeuvre même.

La limite atteinte par la suite peut dépendre de la valeur initiale de la suite θ_0 . Ce problème est habituel dans la mesure où la plupart des algorithmes itératifs déterministes y sont confrontés. Supposons que les hypothèses suffisantes à la convergence de la suite (θ_k) vers un maximum local soit vérifiées. Dans le cas d'une vraisemblance unimodale, la limite sera bien le maximum global de la vraisemblance. Par contre, si on considère le cas d'une vraisemblance multimodale, la convergence aura lieu vers l'un des maxima locaux, rien n'assurant qu'il s'agisse du maximum global. Un des moyens de contourner ce problème consiste à faire varier le point initial θ_0 dans Θ pour atteindre différentes limites possibles et de comparer ensuite les valeurs de la vraisemblance obtenues en chacune de ces limites pour choisir finalement la limite en laquelle la vraisemblance est la plus grande (Wu (1983)).

Concernant les problèmes relatifs à la mise en oeuvre de l'algorithme, on en compte trois principaux dus au calcul de la quantité $Q(\theta|\theta_k)$, à la maximisation en θ de cette valeur et à la convergence relativement lente de l'algorithme. Lorsque la maximisation directe de $Q(\theta|\theta_k)$ est délicate, on peut procéder par des accroissements successifs en choisissant la valeur θ_{k+1} telle que $Q(\theta_{k+1}|\theta_k) \geq Q(\theta_k|\theta_k)$. Un algorithme du type gradient ou Newton-Raphson apporte donc une solution satisfaisante à ce problème comme le propose par exemple Lange (1995) qui utilise une itération d'un algorithme de Newton au cours de l'étape M. Un autre type de solution fut proposé par Meng et Rubin (1993) : l'étape M est remplacée par une étape CM constituée d'une suite de maximisations sous contraintes de la log-vraisemblance complète selon certaines

directions prédéfinies, inspirée de l'échantillonneur de Gibbs (Meng et Rubin (1992)). Cette modification apportée à l'algorithme EM, notée ECM (*Expectation Conditional Maximization*), préserve la propriété de croissance de la vraisemblance, ainsi que celle de la convergence de la suite (θ_k) moyennant des hypothèses sur la régularité du modèle et sur les contraintes choisies pour les maximisations (Meng (1994)).

Les problèmes persistants sont donc principalement liés aux difficultés de calcul de l'intégrale définissant $Q(\theta|\theta_k)$ et au temps de calcul nécessaire à la convergence de l'algorithme dans les applications pratiques. Des solutions à ces deux types de problèmes sont apparues par le biais d'une approche stochastique, alliant la simulation de variables aléatoires à l'algorithme EM.

Quelques variations stochastiques sur EM

Une des premières variantes stochastiques de l'algorithme EM fut la version SEM (*Simulated EM*) proposée par Celeux et Diebolt (1986) pour l'identification de mélanges finis de densités. A l'étape E s'ajoute une étape S de simulation : on calcule la densité conditionnelle $p(\cdot|y; \theta_k)$, puis on simule une réalisation z_k selon cette loi, qui remplace à l'itération k de l'algorithme les données manquantes. La maximisation en θ se fait donc sur la log-vraisemblance complète $\log f(y, z_k; \theta)$. Outre le fait de résoudre le problème lié au calcul de l'intégrale de l'étape E, le caractère stochastique du nouvel algorithme permet aussi de réduire la dépendance de la limite de la suite (θ_k) en les conditions initiales : en fait, la simulation de la réalisation z_k effectuée à chaque itération laisse à la suite (θ_k) une certaine souplesse en autorisant une exploration plus générale des modes de la vraisemblance dans le cas multimodal. Autrement dit, l'algorithme n'est plus contraint systématiquement dans un voisinage du mode le plus proche des conditions initiales. Les résultats de convergence obtenus par Celeux et Diebolt (1986) sont relativement limités dans la mesure où seul le comportement en moyenne de la suite (θ_k) est considéré. Des améliorations ont cependant été apportées par les auteurs un peu plus tard en considérant une version de type recuit simulé, combinant EM et la version SEM précédemment proposée (Celeux et Diebolt (1990), (1992)). Ils conservent le cadre des mélanges de densités et établissent la convergence presque sûre de la suite (θ_k) vers un maximum local de la vraisemblance.

Lorsque l'espérance conditionnelle de la log-vraisemblance complète ne peut être calculée, Wei et Tanner (1990) proposent une autre solution, notée MCEM (*Monte Carlo EM*), qui consiste lors de l'itération k de l'algorithme EM à approcher l'intégrale $Q(\theta|\theta_k)$ par une méthode de Monte Carlo : il s'agit donc de simuler un grand nombre T de variables aléatoires $(z^t)_{1 \leq t \leq T}$ à partir de la loi a posteriori p de z sachant y et de la valeur courante du paramètre θ_k , pour ensuite réaliser l'approximation suivante :

$$Q(\theta|\theta_k) \simeq \frac{1}{T} \sum_{t=1}^T \log f(y, z^t; \theta_k).$$

L'inconvénient notable de cette méthode est l'allongement du temps de calcul qu'elle engendre, vu le nombre important de simulations qu'il faut réaliser à chaque itération de l'algorithme. Dans l'article de Wei et Tanner (1990), l'algorithme MCEM est testé et illustré sur deux exemples, mais aucun résultat théorique de convergence ne vient appuyer ces simulations.

Pour éviter des simulations trop nombreuses tout en conservant de bonnes propriétés de convergence, Delyon, Lavielle et Moulines (1999) ont proposé une version intermédiaire utilisant la simulation via une approximation stochastique.

L'algorithme SAEM

Dans cette variante de l'algorithme EM notée SAEM (*Stochastic Approximation EM*), l'étape E est divisée en deux parties, d'une part une étape de simulation, notée S et d'autre part une étape d'approximation stochastique, notée A. Lors de l'itération k , l'étape S consiste à simuler une réalisation z_k des données manquantes z d'après la loi conditionnelle $p(\cdot|y; \theta_k)$. L'étape A introduit une suite décroissante de pas positifs (γ_k) et réalise l'approximation stochastique suivante

$$Q_{k+1}(\theta) = Q_k(\theta) + \gamma_k [\log f(y, z_k; \theta) - Q_k(\theta)], \quad (5)$$

à partir d'une initialisation déterministe $Q(\theta_0)$. L'étape M devient :

$$\theta_{k+1} = \arg \max_{\theta \in \Theta} Q_{k+1}(\theta). \quad (6)$$

En fait, si on suppose que le modèle est de type exponentiel (4), il suffit de réaliser l'approximation stochastique (5) sur une statistique exhaustive du modèle $\tilde{S}$. On construit alors une suite de points (s_k) , initialisée par s_0 , de la façon suivante :

$$s_{k+1} = s_k + \gamma_k [\tilde{S}(y, z_k) - s_k], \quad (7)$$

et l'étape M s'écrit avec les notations de l'hypothèse **(M5)** :

$$\theta_{k+1} = \hat{\theta}(s_{k+1}). \quad (8)$$

Pour motiver l'approximation stochastique effectuée en (7), on peut l'écrire de la manière suivante :

$$s_{k+1} = s_k + \gamma_k \left[\left(E[\tilde{S}(y, z)|y; \hat{\theta}(s_k)] - s_k \right) + \left(\tilde{S}(y, z_k) - E[\tilde{S}(y, z)|y; \hat{\theta}(s_k)] \right) \right].$$

En posant $h(s_k) = E[\tilde{S}(y, z)|y; \hat{\theta}(s_k)] - s_k$ et $e_k = \tilde{S}(y, z_k) - E[\tilde{S}(y, z)|y; \hat{\theta}(s_k)]$, on obtient :

$$s_{k+1} = s_k + \gamma_k [h(s_k) + e_k]. \quad (9)$$

La suite (e_k) représente une perturbation aléatoire. Pour une suite de pas (γ_k) bien choisie et pour des perturbations aléatoires (e_k) assez petites, on peut faire une analogie entre l'écriture (9) et l'équation différentielle ordinaire suivante :

$$ds(t)/dt = h(s(t)). \quad (10)$$

L'heuristique des approximations effectuées pour obtenir cette analogie est détaillée par Benveniste, Métivier et Priouret (1990). Ainsi, il existe un lien entre le comportement de la suite (s_k) et les solutions de l'équation différentielle ordinaire (10). Si cette dernière admet un point attracteur s^* , on peut en déduire, sous certaines hypothèses sur l'équation différentielle ordinaire (10) et sur la suite (γ_k) , une convergence presque sûre de la suite (s_k) vers le point s^* . Pour plus de détails sur les approximations stochastiques, on pourra se référer à l'ouvrage de Duflo (1996). Le résultat de convergence obtenu par Delyon, Lavielle et Moulines (1999) nécessite, outre les hypothèses de régularité **(M1)**-**(M5)** faites précédemment sur le modèle, des hypothèses relatives à la méthode de simulation des données manquantes z et à la suite de pas (γ_k) . Ainsi, on suppose que les variables aléatoires $z_1, z_2, \dots, z_k$ sont définies sur le même espace de probabilité $(\Omega, \mathcal{A}, P)$. On note $\mathcal{F} = \{\mathcal{F}_k\}_{k \geq 0}$ la famille croissante de tribus engendrées par les variables aléatoires $z_1, z_2, \dots, z_k$. De plus, on suppose vérifiées les hypothèses suivantes :

- **(SAEM1)** Pour tout k dans $\mathbb{N}$, $\gamma_k \in [0, 1]$, $\sum_{k=1}^{\infty} \gamma_k = \infty$ et $\sum_{k=1}^{\infty} \gamma_k^2 < \infty$.
- **(SAEM2)** $l : \Theta \rightarrow \mathbb{R}$ et $\hat{\theta} : \mathcal{S} \rightarrow \Theta$ sont m fois continûment différentiables, où m est l'entier tel que $\mathcal{S}$ est un ouvert de $\mathbb{R}^m$.
- **(SAEM3)**

- Pour toute fonction ϕ borélienne positive :

$$E[\phi(z_{k+1})|\mathcal{F}_k] = \int \phi(z)p(z|y; \theta_k)\mu(dz).$$

- Pour tout $\theta \in \Theta$, $\int \|\tilde{S}(y, z)\|^2 p(z|y; \theta)\mu(dz) < \infty$, et la fonction

$$\Gamma(\theta) \triangleq \text{Cov}[\tilde{S}(y, z)|y; \theta] \triangleq \int_{\mathbb{R}^l} (\tilde{S}(y, z))^2 p(z|y; \theta)\mu(dz) - \left[\int_{\mathbb{R}^l} \tilde{S}(y, z)p(z|y; \theta)\mu(dz) \right]^2$$

est continue en θ .

On rappelle que $d(x, A)$ est la distance euclidienne du point x à l'ensemble supposé fermé A . On peut alors énoncer le résultat de convergence de la suite (θ_k) :

Théorème 2 (Delyon, Lavielle et Moulines (1999)). *Supposons vérifiées les hypothèses (M1)-(M5) et (SAEM1)-(SAEM3), ainsi que l'hypothèse suivante notée (C) : la suite $(s_k)_{k \geq 0}$ est à valeurs dans un sous-ensemble compact de $\mathcal{S}$. Alors, presque sûrement, on a $\lim_{k \rightarrow +\infty} d(\theta_k, \mathcal{L}) = 0$.*

Remarque. *La condition de compacité (C) peut être relâchée dans la mesure où elle peut-être assurée par un moyen de construction imposé sur la suite (s_k) (voir la seconde remarque après le Théorème 1.1 du Chapitre 1).*

Cet algorithme est tout à fait satisfaisant en termes de rapidité de convergence dans la mesure où il améliore considérablement les performances en temps de l'algorithme MCEM. Cependant, alors que nous recherchons un maximum global ou de manière moins restrictive un maximum local, seule la convergence vers un point stationnaire de la vraisemblance des observations est assurée. Il se peut donc que la limite de la suite (θ_k) soit un minimum ou un point selle de la vraisemblance des observations. Pour éliminer ce type de limites non souhaitées, Delyon, Lavielle et Moulines (1999) ont proposé des hypothèses assurant la convergence presque sûre vers un maximum local de la vraisemblance observée. En fait, ils ont montré dans un premier temps que l'ensemble des points stationnaires de l coïncide avec l'ensemble des points fixes de l'application $T : \theta \rightarrow \hat{\theta}(\bar{s}(\theta))$, puis que les seuls points stationnaires de l stables pour T sont les maxima locaux de l , les autres types de points stationnaires étant tous instables ou hyperboliques. Ceci explique d'une part que la plupart des suites issues de l'algorithme EM ne convergent pas vers des points selle ou vers des minima de la vraisemblance observée, d'autre part que les suites issues de l'algorithme SAEM évitent en général ces points du fait du caractère stochastique de cet algorithme. Précisons ici quelles sont ces hypothèses suffisantes qui assurent la convergence vers un maximum de la vraisemblance observée.

- **(LOC1)** Les points stationnaires de l sont isolés : tout sous-ensemble compact de $\mathcal{X}$ contient un nombre fini de points stationnaires.
- **(LOC2)** Pour tout point stationnaire θ^* de $\mathcal{L}$, les matrices :

$$E_{\theta^*} \left[\partial_{\theta} L(\tilde{S}(y, Z); \theta^*)^t \partial_{\theta} L(\tilde{S}(y, Z); \theta^*) \right] \text{ et } - \partial_{\theta}^2 L(E_{\theta^*}[\tilde{S}(y, Z)]; \theta^*)$$

sont définies positives.

- **(LOC3)** : La plus petite valeur propre de la matrice de covariance :

$$R(\theta) = E_{\theta} \left[(\tilde{S}(y, Z) - \bar{s}(\theta))(\tilde{S}(y, Z) - \bar{s}(\theta))^t \right]$$

est strictement positive pour tout θ dans un sous-ensemble compact $\mathcal{K}$ de Θ .

Le résultat de convergence obtenu s'énonce alors comme suit :

Théorème 3 (Delyon, Lavielle et Moulines (1999)). *Supposons vérifiées les hypothèses (M1)-(M5), (SAEM1)-(SAEM3) et (LOC1)-(LOC3), ainsi que l'hypothèse suivante notée (C) : la suite $(s_k)_{k \geq 0}$ est à valeurs dans un sous-ensemble compact de $\mathcal{S}$. Alors, presque sûrement, $\lim_{k \rightarrow +\infty} \theta_k = \theta^*$ où θ^* est un maximum local de la vraisemblance observée l .*

La démonstration de ce résultat se fait en appliquant les résultats de Brandière et Duflo (1995). Les résultats de convergence obtenus pour l'algorithme SAEM sont donc satisfaisants, d'autant que Delyon, Lavielle et Moulines (1999) ont également démontré la normalité asymptotique pour une suite $(\hat{\theta}_k)$ déduite de la suite (θ_k) par une combinaison convexe.

Néanmoins, d'un point de vue pratique, la mise en oeuvre de l'algorithme SAEM ne peut pas toujours se faire car elle nécessite la connaissance de la loi a posteriori p . Plus exactement, elle nécessite de savoir simuler des réalisations de variables aléatoires suivant cette loi. Nous proposons donc dans le premier chapitre de cette thèse une méthode qui permet la mise en oeuvre de l'algorithme dans un cadre moins restrictif, basée sur les méthodes de Monte Carlo par chaînes de Markov.

Les méthodes de Monte Carlo par chaînes de Markov : MCMC

Pour une loi de probabilité donnée, on appelle algorithme de Monte Carlo par chaînes de Markov toute méthode produisant une chaîne de Markov ergodique ayant pour loi stationnaire la dite loi (Robert (1996)). L'existence de probabilités de transition générant des chaînes ayant les bonnes propriétés de convergence vers la loi stationnaire est assurée par exemple par l'algorithme de Hastings-Metropolis (Robert (1996)). Les méthodes MCMC sont apparues au début des années 1990. Les physiciens les utilisaient déjà sur des espaces finis, mais leur extension à des espaces quelconques n'a pu se faire qu'à partir du moment où les puissances de calcul considérables qu'elles requièrent furent disponibles. Cette limitation technique explique sans doute le fait que l'article précurseur de Hastings (1970) n'a pu avoir immédiatement les retombées conséquentes que l'on connaît actuellement. Les possibilités d'applications des méthodes MCMC pour l'estimation dans des modèles à données manquantes sont nombreuses : un certain nombre d'algorithmes stochastiques existants supposent la simulation de réalisations de variables aléatoires selon une distribution donnée. Mais il se peut que cette distribution ne soit pas connue, ou qu'il soit impossible de simuler facilement des variables aléatoires selon cette loi. Dans ce cas, on peut envisager de générer une chaîne de Markov ergodique dont la loi stationnaire sera cette loi. Les méthodes MCMC laissent donc entrevoir de possibles généralisations de ces algorithmes, permettant ainsi des applications moins restrictives, certaines ayant déjà été suggérées ou mises en pratique, sans justification théorique de convergence, comme dans les articles de Lavielle et Lebarbier (2001) et Cappé, Douc, Moulines et Robert (2002).

Nous présentons maintenant deux applications proches de notre problème pour illustrer les possibilités apportées par les méthodes MCMC. Les travaux de Fort et Moulines (2003) sur l'algorithme MCEM présenté ci-dessus sont très représentatifs des remarques précédentes, dans la

mesure où la simulation des données manquantes se fait à chaque itération suivant la probabilité de transition d'une chaîne de Markov ergodique de loi stationnaire la loi conditionnelle $p(\cdot|y; \theta)$. Supposant vérifiées des hypothèses d'ergodicité sur la chaîne, ainsi qu'un contrôle de l'accroissement de la suite (m_k) du nombre de simulations effectuées dans la moyenne de Monte Carlo à l'itération k de l'algorithme, les auteurs obtiennent des résultats de convergence presque sûre de la suite (θ_k) vers un point stationnaire de la vraisemblance, ainsi qu'une vitesse de convergence dépendant de la suite (m_k) choisie.

Par ailleurs, les travaux de Gu et Kong (2001) sont également illustratifs de l'application des méthodes MCMC. Les auteurs proposent un algorithme d'estimation adapté aux modèles à données manquantes qui s'inspire de la méthode de Newton-Raphson et associe le principe de l'approximation stochastique aux méthodes de Monte Carlo par chaînes de Markov. Leur algorithme permet la résolution en θ d'équations de la forme $H(y, z; \theta) = 0$ (H est égal au gradient de la vraisemblance dans le cas de l'estimation par maximum de vraisemblance). Comme z n'est pas observé, on se ramène à la résolution d'équation de la forme $h(y; \theta) = 0$ où h est égale à l'espérance conditionnelle de H . Partant de là, les auteurs génèrent de manière itérative une suite de couples (Γ_k, θ_k) comme suit :

$$\begin{aligned}\Gamma_k &= \Gamma_{k-1} + \gamma_k \left(\frac{1}{m} \sum_{i=1}^m I(\theta_{k-1}, z_{k,i}) - \Gamma_{k-1} \right) \\ \theta_k &= \theta_{k-1} + \gamma_k \Gamma_k^{-1} \left(\frac{1}{m} \sum_{i=1}^m H(\theta_{k-1}, z_{k,i}) - \theta_{k-1} \right),\end{aligned}$$

où les données manquantes z sont simulées à l'itération k selon une probabilité de transition Π_{θ_k} telle que pour tout θ , la chaîne issue de Π_{θ} ait pour loi stationnaire la loi conditionnelle $p(\cdot|y; \theta)$ de z en y . La fonction I est choisie de sorte que son espérance conditionnelle soit proche de $\partial_{\theta} h(\theta)$ au voisinage de la racine $\hat{\theta}$ de l'équation $h(\theta) = 0$. Sous certaines conditions de régularité du modèle et d'ergodicité de la chaîne, ils obtiennent la convergence du couple (Γ_k, θ_k) vers une solution stable d'une équation différentielle ordinaire déduite des fonctions H et I .

Utilisation d'une méthode MCMC dans l'algorithme SAEM

Décrivons à présent l'algorithme présenté dans le Chapitre 1 de cette thèse et appliqué dans le Chapitre 2. Comme la procédure de Gu et Kong (2001), la nôtre est basée sur l'association d'une approximation stochastique et d'une méthode de Monte Carlo par chaînes de Markov. Nous proposons de combiner judicieusement les itérations de l'algorithme SAEM à la procédure MCMC pour obtenir une convergence rapide de la suite (θ_k) . Lorsque la simulation de variables suivant la distribution a posteriori n'est pas réalisable, on considère une chaîne de Markov ergodique de probabilité de transition Π_{θ} ayant pour loi stationnaire la distribution a posteriori $p(\cdot|y; \theta)$. Pour éviter un allongement considérable du temps de calcul, on ne simule pas une chaîne à chaque itération pour obtenir une réalisation des données manquantes de loi asymptotiquement égale à la distribution a posteriori. L'idée principale est de faire évoluer les deux convergences en même temps, d'une part celle de la chaîne issue de la méthode de Monte Carlo vers la loi stationnaire $p(\cdot|y; \theta)$, d'autre part, celle de la suite (θ_k) .

Pour générer en pratique une telle chaîne, on se donne une probabilité de transition q appelée *loi instrumentale* (*proposal kernel*). A partir de la valeur z_{k-1} de la chaîne obtenue à l'itération

$k - 1$, on génère un candidat z' selon $q(z_{k-1}, \cdot)$. On choisit ensuite $z_k = z'$ avec une probabilité $\rho = \min(\frac{p(z'|y)}{p(z_{k-1}|y)} \frac{q(z', z_{k-1})}{q(z_{k-1}, z')}, 1)$ et $z_k = z_{k-1}$ avec une probabilité $1 - \rho$. Cette procédure accepte le candidat proposé chaque fois que le quotient $\frac{p(z'|y)}{q(z, z')}$ est supérieur au précédent. Pour notre algorithme, on choisira une transition Π_θ construite comme ci-dessus (voir Chapitre 1 Section 1.2 pour le détail de la procédure et Chapitre 2 Section 2.2.2 pour des applications pratiques). Ainsi notre version de l'algorithme diffère de la version de SAEM de Delyon et al. uniquement au niveau de l'étape S : lors de l'itération k , la réalisation z_k est simulée d'après la probabilité de transition Π_{θ_k} (au lieu de la loi conditionnelle $p(\cdot|y; \theta_k)$). Moyennant certaines hypothèses sur la chaîne de Markov de transition Π_θ , nous obtenons la convergence presque sûre de la suite (θ_k) vers un maximum local de la log-vraisemblance observée l . Comme pour la version de SAEM initiale, la démonstration est basée sur le théorème de Delyon (1996) concernant les approximations stochastiques. Cependant, la théorie des martingales utilisée par Delyon et al. ne s'applique plus dans notre cas, l'hypothèse **(SAEM3)** n'étant plus vérifiée. Nous utilisons donc à la place un résultat de Benveniste, Métivier et Priouret (1990) sur les perturbations aléatoires markoviennes.

A la fin du Chapitre 1, nous présentons deux applications pratiques, la première à un problème de déconvolution, la seconde à un problème de détection de ruptures. Ce second problème avait déjà été abordé par Lavielle et Lebarbier (2001), mais sans justification théorique de la convergence de l'algorithme qui utilise une méthode MCMC à la place de simulations indépendantes.

Applications aux modèles non linéaires à effets mixtes

Le deuxième chapitre de cette thèse présente un domaine d'application privilégié : les modèles à effets mixtes. Ce type de modélisation permet de prendre en compte, pour une série d'observations issues d'une expérience donnée, deux types de comportements : d'une part, un comportement global, d'autre part, un comportement individuel. Ces comportements sont matérialisés par deux types d'effets : les effets fixes et les effets aléatoires. Les premiers sont identiques pour l'ensemble des individus observés et sont donc des paramètres du modèle (dits aussi paramètres de population), alors que les seconds varient avec l'individu observé et sont donc aléatoires (désignés également comme paramètres individuels).

Bien qu'il y ait eu une recrudescence de l'intérêt porté à ce type de modèle ces dernières années, leur apparition n'est pas toute récente. Par exemple, ces modèles sont déjà mentionnés par Dempster, Laird et Rubin (1977). Les domaines d'application de ces modèles sont nombreux et variés, expliquant sans doute la multitude de travaux effectués sur le sujet. On peut citer à titre d'exemple l'économétrie, ainsi que la pharmacocinétique et la pharmacodynamique auxquelles on s'intéressera tout particulièrement dans les Sections 2.4 et 2.5, avec l'étude des courbes de croissance et de l'analyse de l'évolution de la concentration d'une substance dans un milieu au cours du temps. Détaillons le modèle probabiliste considéré :

$$y_{ij} = C(t_{ij}, \phi_i, \beta) + D(t_{ij}, \phi_i, \beta)\varepsilon_{ij} \quad \text{pour } 1 \leq i \leq n, \quad 1 \leq j \leq m_i, \quad (11)$$

où y_{ij} est la j -ième observation du i -ième individu, au temps t_{ij} . Le nombre d'individus est ici égal à n et le nombre d'observations de l'individu i à m_i . Les fonctions C et D sont à valeurs réelles. On dira que le modèle est non linéaire lorsque C ou D est non linéaire en l'effet aléatoire ϕ_i . Les erreurs (ε_{ij}) sont supposées indépendantes gaussiennes de moyenne nulle et de variance

σ^2 . De plus, on suppose les erreurs (ε_{ij}) indépendantes des effets aléatoires (ϕ_i) . Ces derniers sont modélisés de la façon suivante :

$$\phi_i = \mathbf{A}_i \boldsymbol{\mu} + \boldsymbol{\eta}_i \text{ avec } \boldsymbol{\eta}_i \sim_{i.i.d.} N(0, \boldsymbol{\Gamma}),$$

où $\boldsymbol{\mu}$ est un vecteur inconnu de paramètres. La matrice $\mathbf{A}_i$ est supposée connue alors que la matrice $\boldsymbol{\Gamma}$ est supposée inconnue. Le vecteur $\boldsymbol{\beta}$ désigne les autres paramètres du modèle qui interviennent en dehors des moyennes des effets aléatoires (on suppose parfois que les effets aléatoires sont de moyenne nulle, prenant en compte comme effets fixes leurs moyennes).

Notre but est d'estimer par maximum de vraisemblance les paramètres $\boldsymbol{\theta} = (\boldsymbol{\beta}, \boldsymbol{\mu}, \boldsymbol{\Gamma}, \sigma^2)$ du modèle.

Dans le cas d'un modèle linéaire, l'estimation par maximum de vraisemblance peut se faire soit par l'algorithme EM (Dempster et al. (1977)), soit par un algorithme de Newton-Raphson (Lindstrom et Bates (1988)). Dès lors que le modèle n'est plus linéaire, ces méthodes ne peuvent plus être appliquées. Par exemple, l'espérance conditionnelle intervenant dans l'étape E de l'algorithme EM n'admet plus d'expression analytique simple. Pourtant une modélisation non linéaire s'avère bien souvent plus adéquate d'un point de vue pratique. Cette limitation est donc très restrictive et des solutions très variées ont émergé pour permettre une estimation dans des modèles non linéaires. Des approches bayésiennes ont été proposées par Racine-Poon (1985), Wakefield et al. (1994) et Wakefield (1996). Parallèlement, d'autres approches basées sur une linéarisation de la vraisemblance se sont développées : Lindstrom et Bates (1990) propose une linéarisation basée sur un développement de Taylor, aboutissant à un estimateur sans propriété asymptotique connue. Vonesh (1996) propose lui une approximation de Laplace. L'estimateur fourni est consistant entre autres sous l'hypothèse contraignante et peu réaliste que les nombres d'observations par individu (m_i) tendent vers l'infini. Par ailleurs, des méthodes basées sur la simulation ont également fait leur apparition dans ce domaine : Walker (1996) applique l'algorithme MCEM aux modèles non linéaires à effets mixtes, alors que Concordet et Nunez (2002) proposent un estimateur du maximum de pseudo-vraisemblance consistant.

Nous proposons d'appliquer l'algorithme SAEM couplé à une méthode MCMC aux modèles non linéaires à effets mixtes dans le but d'approcher l'estimateur du maximum de vraisemblance. Les résultats de convergence présentés dans le premier chapitre de cette thèse assurent la convergence presque sûre de la suite produite par l'algorithme vers un maximum local de la vraisemblance sous des hypothèses réalistes, pour des nombres d'observations par individus (m_i) finis. Nous rappelons tout d'abord brièvement les étapes de l'algorithme SAEM couplé à une méthode MCMC et détaillons quelques points pratiques concernant la simulation des données manquantes et le choix des constantes de l'algorithme qui influent sur la qualité des résultats. Nous présentons également des méthodes d'estimation de la vraisemblance, par approximation stochastique et par intégration de Monte Carlo, ainsi qu'une méthode d'estimation de la matrice d'information de Fisher (Delyon et al. (1999)).

Nous étudions ensuite à travers quelques exemples les performances de notre algorithme SAEM couplé à une méthode MCMC en le comparant aux méthodes existantes. Nous avons donc sélectionné des données simulées et réelles dans la littérature. Le premier exemple est présenté de manière didactique et vise à comparer les algorithmes EM et SAEM : il s'agit en fait de données réelles issues de l'observation d'orangers, le modèle homoscedastique considéré comporte un effet aléatoire et deux effets fixes et est linéaire en l'effet aléatoire, permettant l'application de l'algorithme EM. Outre les paramètres, la vraisemblance et l'information de Fisher sont également estimées. Une extension à un modèle hétéroscedastique est traitée, moyennant une

estimation bayésienne des deux effets fixes. Les autres exemples sélectionnés sont issus soit de la pharmacocinétique, soit de la pharmacodynamique (*PK/PD*).

Le premier modèle pharmacocinétique considéré a été étudié par Concordet et Nunez (2002). Dans un premier temps, l'algorithme SAEM est mis en oeuvre pour obtenir une estimation des paramètres. Nous présentons à cette occasion diverses variantes sur le choix des constantes de l'algorithme, des noyaux de transition intervenant dans la méthode de Hastings-Metropolis, ainsi qu'une suite de pas (γ_k) adaptative. Les résultats obtenus sont comparés avec ceux de Concordet et Nunez, ainsi qu'avec ceux obtenus par la méthode FOCE, décrite dans la Section 2.5, et par la quadrature de Gauss. L'algorithme SAEM produit des estimations ayant un écart quadratique moyen globalement bien plus faible que les autres méthodes. Le deuxième exemple de cette section présente un modèle homoscedastique issu de la pharmacodynamique, traité par Walker (1996). Les résultats obtenus par l'algorithme MCEM de Walker et par notre algorithme sont du même ordre, un peu meilleurs que ceux fournis par les méthodes FOCE et LAPLACIAN du logiciel NONMEM, habituellement utilisé par les praticiens dans ce domaine. Cependant nous mettons ici en évidence la rapidité avec laquelle notre algorithme converge, par opposition à la convergence plus coûteuse en temps de l'algorithme MCEM (pour lequel un échantillonnage est réalisé à chaque itération). Le deuxième exemple issu de la pharmacocinétique est un modèle hétéroscedastique présenté par Mentré et Gomeni (1995). Il modélise l'évolution de la concentration du plasma au cours du temps après l'injection d'une dose de médicament. Ce modèle comporte également une covariable : l'âge du patient, qui intervient au sein de la moyenne d'un des trois effets aléatoires. Mentré et Gomeni appliquent l'algorithme P-Pharm décrit dans la Section 2.4.3. L'algorithme FOCE est également mis en oeuvre et les trois résultats obtenus sont comparés. Les biais relatifs moyens obtenus par notre algorithme sont sensiblement du même ordre, excepté pour l'une des variances des effets aléatoires où nous obtenons un biais relatif moyen cinq fois plus faible. Par contre, les écarts quadratiques moyens que nous obtenons sont trois fois plus petits que ceux obtenus par P-Pharm et cinq fois plus petits que ceux obtenus par FOCE.

Le dernier exemple présente une application à un test du rapport de vraisemblance, mis en place pour évaluer la pertinence de différentes trithérapies. Ce travail a été réalisé en collaboration avec France Mentré et Adeline Leclercq (INSERM, Hôpital Bichat, Paris).

Ding et Wu (1999) ont établi un lien entre l'effet des trithérapies et le taux de décroissance de la charge virale. Ainsi en choisissant une modélisation statistique bien adapté dans laquelle ce taux intervient comme paramètre, Ding et Wu (2001) proposent de tester l'efficacité du traitement par un test du rapport de vraisemblance. On considère donc deux traitements administrés chacun à un groupe différent. L'hypothèse nulle (**H0**) suppose que les deux groupes réagissent de manière identique aux traitements. L'alternative (**H1**) considère que tous les paramètres du modèle sont les mêmes, excepté le taux de décroissance de la charge virale qui est supposé différent. Dans un premier temps, les paramètres sont estimés par SAEM et par l'algorithme FOCE via la procédure NLME de S-Plus, puis les log-vraisemblances sous les hypothèses (**H0**) et (**H1**) sont estimées par une moyenne de Monte Carlo utilisant un échantillonnage pondéré (Robert (1996)). L'erreur de première espèce du test est approximée par la moyenne empirique effectuée sur 100 échantillons, pour effectuer la comparaison avec les résultats de Ding et Wu (2001) pour l'alternative (**H1**). Nos estimations des paramètres donnent lieu à des biais moyens relatifs plus faibles que ceux obtenus par NLME et des écarts quadratiques moyens relatifs du même ordre, excepté pour une des variances d'effets aléatoires qui est nettement améliorée.

Application à un problème pratique en géophysique

Le travail présenté dans le troisième chapitre a été réalisé en collaboration avec Hermann Zeyen (Département des Sciences de la Terre, Université Paris Sud, Orsay).

Dans ce chapitre, nous proposons une application de l'algorithme SAEM en géophysique. Nous disposons de deux types d'observations, d'une part, des observations de temps de parcours d'ondes sismiques, notées dt , d'autre part, des observations gravimétriques de surface notées dg . Ces données sont reliées par des relations physiques respectivement à la vitesse des ondes (*velocity*), plus précisément à son inverse (*slowness*) notée ds , et à la densité du sous-sol notée $d\rho$. Ces relations sont linéaires ou supposées linéarisées et donnent lieu à deux problèmes inverses :

$$\begin{aligned} dt &= A_t ds + \varepsilon_t \\ dg &= A_g d\rho + \varepsilon_g, \end{aligned}$$

où A_t et A_g sont des matrices connues et ε_t et ε_g des termes d'erreur aléatoires. D'un point de vue statistique, la nouveauté apportée par cette approche réside dans le fait que les paramètres du modèles sont estimés au lieu d'être fixés arbitrairement. Par ailleurs, nous prenons en compte une relation linéaire entre l'inverse de la vitesse des ondes et la densité du sous-sol :

$$ds = b d\rho + \eta,$$

où b est un paramètre inconnu pouvant dépendre de la profondeur et η un terme d'erreur aléatoire. Nous présentons dans un premier temps l'aspect géophysique du problème, les données et les relations physiques utilisées. Puis nous choisissons un modèle gaussien pour les différentes variables mises en oeuvre, dont nous estimons les paramètres. Nous reconstruisons ensuite les données non observées par le maximum a posteriori (*MAP*). Nous comparons nos résultats à ceux obtenus par l'algorithme de type bayésien de Zeyen et Achauer (1997). Nous présentons des applications sur des données simulées, inspirées d'un site réel : le rift kenyan. Nous réalisons d'abord l'inversion simple des temps de parcours téléseismiques, puis nous traitons l'inversion simultanée des deux problèmes (téléseismiques et gravimétriques), pour un modèle physique stratifié en trois couches et un modèle physique stratifié en six couches. Enfin, nous réalisons une inversion simultanée sur un modèle comportant une irrégularité des densités au niveau de la première couche. Outre l'estimation des paramètres du modèle, l'application de l'algorithme SAEM permet une nette amélioration des résultats obtenus par Zeyen et Achauer pour la reconstruction des données non observées, tout particulièrement pour les inverses des vitesses ds .

Bien que les résultats obtenus par ce couplage de l'algorithme SAEM avec une méthode MCMC soient encourageants aussi bien en théorie que dans les applications, il serait néanmoins souhaitable de s'affranchir d'un choix arbitraire d'un a priori paramétrique (en général gaussien) sur la loi des données manquantes. Dans ce sens, nous présentons dans le quatrième chapitre une version non paramétrique de l'algorithme SAEM couplé avec une méthode MCMC.

Cadre et objectifs du problème non paramétrique

On considère dans le quatrième chapitre un problème statistique à données manquantes dans le cadre non paramétrique suivant : les données observées sont notées $(y_1, \dots, y_n)$ et les données

non observées (ou manquantes) sont notées $(z_1, \dots, z_n)$. On suppose que chaque observation y_i et chaque donnée non observée z_i pour i compris entre 1 et n est à valeurs respectivement dans un intervalle $\mathcal{Y}$ de $\mathbb{R}$ et dans un intervalle compact $\mathcal{I} = [a, b]$ de $\mathbb{R}$. On suppose que les observations $(y_1, \dots, y_n)$ sont indépendantes et on note, pour tout i compris entre 1 et n , h_i la densité de y_i conditionnellement à z_i . Les données non observées $(z_1, \dots, z_n)$ sont indépendantes et identiquement distribuées selon une loi non paramétrique de densité π par rapport à la mesure de Lebesgue sur $\mathcal{I}$. Notre but dans ce dernier chapitre est de proposer un estimateur du maximum de vraisemblance pour la densité non paramétrique π et un outil algorithmique pour l'approcher, puis d'étudier les propriétés asymptotiques de l'estimateur lorsque le nombre d'observations n tend vers l'infini. Pour cela, nous nous plaçons dans un premier temps dans un modèle logspline de dimension fixée, dans lequel nous considérons l'estimateur du maximum de vraisemblance paramétrique pour ce modèle. Nous approchons cet estimateur par l'algorithme SAEM couplé à la méthode MCMC. Puis nous faisons tendre la dimension du modèle logspline ainsi que le nombre d'observations vers l'infini pour obtenir des propriétés de convergence de notre estimateur vers la densité π .

Estimation de la densité des données manquantes par maximum de vraisemblance dans un modèle logspline

Précisons le modèle logspline considéré. Soit $\mathcal{I} = [a, b]$ un intervalle compact de $\mathbb{R}$. Pour tout entier q supérieur ou égal à 1, pour tout entier K supérieur ou égal à 1 et pour toute subdivision $\tau = (t_l)_{1 \leq l \leq K+1}$ de $[a, b]$ telle que $a = t_1$ et $b = t_{K+1}$, l'espace des fonctions polynomiales par morceaux sur $\mathcal{I}$, appelées spline, d'ordre q est de dimension $J = q + K - 1$ et admet une base de B-splines notée $(B_1, \dots, B_J)$. On renvoie à la lecture de de Boor (1978) pour plus de détails sur la construction des B-splines. On définit alors pour tout vecteur $\theta = (\theta_1, \dots, \theta_J)$ de $\mathbb{R}^J$ et pour toute valeur z de $\mathcal{I}$ les quantités suivantes :

$$s(z; \theta) = \sum_{j=1}^J \theta_j B_j(z) \quad \text{et} \quad c(\theta) = \log \left[\int_{\mathcal{I}} \exp[s(z; \theta)] dz \right],$$

de sorte que la fonction π_θ définie par

$$\pi_\theta(z) = \exp[s(z; \theta) - c(\theta)]$$

est une densité pour tout θ de $\mathbb{R}^J$. On définit pour tout entier q supérieur ou égal à 1, pour tout entier K supérieur ou égal à 1 et pour toute subdivision $\tau = (t_l)_{1 \leq l \leq K+1}$ de $[a, b]$ telle que $a = t_1$ et $b = t_{K+1}$, l'ensemble des fonctions logspline associé noté $\mathcal{M}_{q,\tau} = \{\pi_\theta, \theta \in \mathbb{R}^J\}$. Ainsi définies, les fonctions logspline sont systématiquement positives et ont une masse totale égale à 1. Elles sont donc particulièrement adaptées à l'estimation de densité. De plus, elles possèdent de bonnes propriétés d'approximations fonctionnelles. Ainsi, pour une fonction f telle que $\log f$ est supposé de classe $\mathcal{C}^m$, on a :

$$\inf_{\theta \in \mathbb{R}^J} \|\log f - \log \pi_\theta\|_\infty = O(J^{-m}).$$

Ces propriétés suggèrent qu'un modèle logspline de dimension J suffisamment grande sera assez riche pour fournir un bon estimateur de la densité non paramétrique π . Dans la suite, on considère

uniquement des subdivisions équiréparties de l'intervalle $\mathcal{I} = [a, b]$. Ainsi pour tout entier q supérieur ou égal à 1 et pour tout entier K supérieur ou égal à 1, on définit le modèle logspline paramétrique $\mathcal{M}_{q,K}$ pour la densité des données manquantes. On note $(h_i)_{1 \leq i \leq n}$ les densités conditionnelles de y_i par rapport à z_i pour $1 \leq i \leq n$. La log-vraisemblance l_n des observations $(y_1, \dots, y_n)$ dans le modèle logspline $\mathcal{M}_{q,K}$ s'écrit alors pour θ dans $\mathbb{R}^J$:

$$l_n(\theta) = \frac{1}{n} \sum_{i=1}^n \log \int_{\mathcal{I}} h_i(y_i|z_i) \pi_{\theta}(z_i) dz_i.$$

L'estimateur $\pi_{\hat{\theta}_{n,J}}$ du maximum de vraisemblance dans ce modèle logspline est défini par :

$$\hat{\theta}_{n,J} = \arg \max_{\theta \in \mathbb{R}^J} l_n(\theta).$$

Pour simplifier les notations, on résume la dépendance en q et en K en notant uniquement la dimension J en indice.

Pour approcher cet estimateur, nous proposons d'appliquer l'algorithme SAEM couplé avec une méthode MCMC au modèle $\mathcal{M}_{q,K}$. La log-vraisemblance complète dans ce modèle logspline s'écrit alors pour θ dans $\mathbb{R}^J$:

$$\begin{aligned} l_n^{com}(\theta) &= \frac{1}{n} \sum_{i=1}^n \log h_i(y_i|z_i) + \frac{1}{n} \sum_{i=1}^n \log \pi_{\theta}(z_i) \\ &= \frac{1}{n} \sum_{i=1}^n \log h_i(y_i|z_i) + \sum_{j=1}^J \theta_j \left[\frac{1}{n} \sum_{i=1}^n B_j(z_i) \right] - c(\theta). \end{aligned}$$

Pour J fixé, ce modèle est exponentiel et permet donc une mise en oeuvre aisée de l'algorithme SAEM. Dans notre cas, l'itération k s'effectue de la manière suivante : lors de l'étape S, on utilise comme loi instrumentale pour la méthode MCMC la densité π_{θ_k} . On effectue l'approximation stochastique sur la statistique exhaustive $(\frac{1}{n} \sum_{i=1}^n B_j(z_i), 1 \leq j \leq J)$, puis on réactualise le paramètre selon (8). D'après les résultats du Chapitre 1, l'algorithme SAEM couplé à la méthode MCMC génère sous des hypothèses de régularité du modèle et de convergence de la méthode MCMC une suite (θ_k) qui converge presque sûrement vers un maximum local de la vraisemblance des observations l_n et permet donc d'approcher notre estimateur $\pi_{\hat{\theta}_{n,J}}$ du maximum de vraisemblance.

Parmi les travaux proposant une estimation de la densité des données manquantes dans le cadre des problèmes inverses via les logspline, on peut citer l'article de Koo et Park (1996) qui présente une estimation par logspline dans le cas du problème inverse linéaire de la déconvolution, basée sur l'application de l'algorithme EM, sans fournir de propriétés asymptotiques de l'estimateur du maximum de vraisemblance obtenu dans le modèle logspline. Par ailleurs, Koo et Chung (1998) adoptent une approche du type *standard value decomposition (SVD)*, alors que Koo (1999) propose un estimateur de la densité π maximisant la vraisemblance dite indirecte dans le cadre de la déconvolution. Dans les deux cas, les estimateurs fournis possèdent des propriétés de convergence vers la densité π , en norme L_2 et en divergence de Kullbak-Leiber, sous certaines hypothèses de régularité de la densité π et de l'opérateur intégral considéré. D'autres approches du type *SVD* existent pour ce type de problèmes à opérateur intégral, on peut citer par exemple de manière non exhaustive les travaux de Cavalier et Tsybakov (2002) et de Donoho

(1995). D'autres auteurs ont proposé des estimateurs issus de méthodes itératives calquées sur l'algorithme EM. On peut citer les travaux de Silverman, Jones, Nychka et Wilson (1990), qui s'appuient sur ceux de Vardi, Sheep et Kaufman (1985) dans le cadre particulier de la tomographie par émission de positrons. On peut également citer les travaux de Schumitzky (1991) qui propose une version non paramétrique où l'estimateur de la densité π est obtenu comme moyenne des densités a posteriori $(p(\cdot|y_i; \theta))$ pour les observations $(y_1, \dots, y_n)$. Cette procédure préserve la propriété remarquable de l'algorithme EM d'accroissement de la vraisemblance des observations. Mais aucun autre résultat de convergence n'est fourni. Par ailleurs, Eggermont et Lariccia (1995) et Eggermont (1999) proposent un lissage par noyau et obtiennent la convergence presque sûre de leur estimateur du maximum de vraisemblance vers la densité π en norme L_1 lorsque n tend vers l'infini.

Nous nous intéressons à la convergence de l'estimateur $\pi_{\hat{\theta}_{n,J}}$ vers la densité π lorsque la dimension n des observations et la dimension J du modèle logspline considéré tendent vers l'infini. Dans ce but, on définit $\Lambda_n(\theta) = \frac{1}{n} \sum_{i=1}^n E[\log g_{i,\theta}(Y_i)]$ où $g_{i,\theta}(y) = \int_{\mathcal{I}} h_i(y|z) \pi_{\theta}(z) dz$ pour tout y de $\mathcal{Y}$. Puis on définit :

$$\theta_{n,J}^* = \arg \max_{\theta \in \mathbb{R}^J} \Lambda_n(\theta). \quad (12)$$

Si la densité π appartient à un ensemble $\mathcal{M}_J$ pour un certain J , alors $\pi_{\theta_{n,J}^*} = \pi$. Ainsi pour une densité π quelconque, on s'intéresse à la convergence de $\pi_{\hat{\theta}_{n,J}}$ vers $\pi_{\theta_{n,J}^*}$ pour n assez grand et celle de $\pi_{\theta_{n,J}^*}$ vers π pour J assez grand. Nous avons démontré, sous certaines hypothèses de régularité du modèle et d'existence de $\theta_{n,J}^*$, l'existence de $\pi_{\hat{\theta}_{n,J}}$ sur un ensemble de probabilité tendant vers 1 lorsque n tend vers l'infini et lorsque J ne tend pas trop vite vers l'infini par rapport à n , ainsi que la convergence en probabilité vers 0 de $\|\log \pi_{\hat{\theta}_{n,J}} - \log \pi_{\theta_{n,J}^*}\|_2$. Pour la démonstration, nous nous sommes appuyés sur les travaux de Stone (1986) et Stone (1990) effectués dans le cadre du problème d'estimation directe, ainsi que sur l'article de Barron et Sheu (1991) pour certaines inégalités reliant la divergence de Kullback-Leiber à la norme L_2 .

Il reste à établir la convergence de $\pi_{\theta_{n,J}^*}$ vers π , en utilisant les bonnes qualités d'approximation des modèles logspline. Nous la démontrons dans le cas particulier suivant :

$$Y = \Phi(Z) + \varepsilon,$$

où Φ est une fonction injective de $\mathbb{R}$ dans $\mathbb{R}$ et ε un terme d'erreur aléatoire de loi ayant une fonction caractéristique ne s'annulant pas (ce qui est par exemple vérifié pour une loi gaussienne). On suppose que toutes les densités conditionnelles h_i sont égales et on note h cette densité conditionnelle. On considère une suite de variables aléatoires (Z_J) de densités $(\pi_{\theta_J^*})$ ($\theta_{n,J}^*$ ne dépend plus de n sous nos hypothèses). Alors la suite (Z_J) converge en loi vers la variable aléatoire Z de densité π sous une hypothèse de régularité de la densité conditionnelle h de Y sachant Z .

Nous présentons maintenant quatre applications. Considérons tout d'abord le modèle suivant :

$$y_i = \alpha_i \exp(-z_i t) + \varepsilon_i \quad \text{pour } 1 \leq i \leq n. \quad (13)$$

Les observations $(y_i)_{1 \leq i \leq n}$ sont obtenues de la manière suivante : les coefficients (α_i) sont indépendants et identiquement distribués de loi gaussienne de moyenne m_α et de variance η_α^2 .

Les données manquantes sont indépendantes et identiquement distribuées de loi gaussienne de moyenne μ et de variance γ^2 . Les erreurs (ε_i) sont indépendantes et identiquement distribuées suivant une loi gaussienne centrée de variance σ^2 . De plus, les coefficients (α_i) , les données manquantes (z_i) et les erreurs (ε_i) sont mutuellement indépendants.

Dans un premier temps, on suppose que les coefficients (α_i) et la variance σ^2 sont connus. On estime alors la densité des données manquantes dans le modèle logspline par l'algorithme SAEM. Nous mettons en évidence les effets du choix de l'ordre q des splines et de la dimension J du modèle logspline sur la convergence de la suite (θ_k) . L'algorithme SAEM couplé à une méthode MCMC fournit une suite π_{θ_k} qui approche bien la densité π .

On considère ensuite que seuls les coefficients (α_i) sont connus. On estime alors la densité des données manquantes dans le modèle logspline et le paramètre σ^2 . L'algorithme SAEM couplé à une méthode MCMC fournit, outre une bonne approximation de la densité π , une estimation correcte de la variance σ^2 . Puis, on considère que les coefficients (α_i) sont également inconnus. On suppose qu'ils suivent une loi gaussienne de moyenne m et de variance η^2 . On estime alors la densité des données manquantes dans le modèle logspline, ainsi que les paramètres m , η et σ^2 . L'algorithme SAEM couplé à une méthode MCMC fournit toujours une bonne approximation de la densité π , ainsi qu'une bonne estimation de la variance σ^2 . Les estimations de m et η sont plus éloignées des vraies valeurs des paramètres.

On considère toujours le modèle (13), mais les données manquantes ne sont plus simulées selon une loi gaussienne, mais selon une distribution bimodale, en l'occurrence un mélange de deux gaussiennes. L'algorithme SAEM couplé à une méthode MCMC donne un estimateur π_{θ_k} qui est proche de la densité π , dans la mesure où il se situe correctement par rapport aux deux modes.

Enfin, pour finir, nous reconsidérons l'exemple des orangers présenté au Chapitre 2 dans la Section 2.3. Dans le cadre paramétrique, nous avons choisi un a priori gaussien pour l'effet aléatoire et nous effectuons désormais une estimation non paramétrique de sa densité, tout en estimant les autres paramètres du modèle. Les résultats obtenus fournissent des estimations du même ordre pour les autres paramètres et une estimation de la densité π de l'effet aléatoire proche de celle obtenue dans le cas paramétrique.

L'algorithme SAEM associé à une méthode MCMC permet donc d'obtenir une approximation pratique de l'estimateur $\pi_{\hat{\theta}_{n,J}}$ du maximum de vraisemblance en fournissant une suite (θ_k) qui converge presque sûrement vers un maximum local de la log-vraisemblance l_n . Les résultats obtenus lors des applications nous font espérer l'obtention de propriétés plus générales pour la convergence de $\pi_{\hat{\theta}_{n,J}}$ vers π . En particulier, nous espérons exhiber des hypothèses assurant l'existence de $\theta_{n,J}^*$ (au moins pour une valeur assez grande de J), obtenir la convergence en probabilité de $\pi_{\hat{\theta}_{n,J}}$ vers $\pi_{\theta_{n,J}^*}$ lorsque n tend vers l'infini et prouver, dans un cadre moins restrictif, la convergence de $\pi_{\theta_J^*}$ vers π lorsque J tend vers l'infini.

Chapitre 1

Coupling a stochastic approximation version of EM with an MCMC procedure

Summary

The stochastic approximation version of EM (SAEM) proposed by Delyon, Lavielle, and Moulines (1999) is a powerful alternative to EM when the E-step is intractable. Convergence of SAEM toward a maximum of the observed likelihood is established when the non observed data are simulated at each iteration under the conditional distribution. We show that this very restrictive assumption can be weakened. Indeed, the results of Benveniste *et al.* for stochastic approximation with Markovian perturbations are used to establish the convergence of SAEM when it is coupled with a Markov chain Monte-Carlo procedure. This result is very useful for many practical applications. Applications to the convolution model and the change-points model are presented to illustrate the proposed method.

Ce travail a donné lieu à la rédaction d'un article co-signé par Estelle Kuhn et Marc Lavielle, accepté par ESAIM PS, actuellement en révision finale.

1.1 Introduction

A wide class of statistical problems involves observed and unobserved data. We can consider, for example, inverse problems such that deconvolution, source separation, change-points detection, etc. Linear and nonlinear mixed effects models can also be considered as incomplete-data models. The purpose is to get the best information about the unknown parameters of the model and about the unobserved data (also called missing data) from the observed data. In particular, the estimation of the parameters of these models, for example by maximum likelihood estimate, is a difficult challenge, mainly because the likelihood of the observations cannot usually be maximized in closed form. Concerning the missing data, they can for example be approached using these parameters estimations by some characteristic of their a posteriori distributions like its maximum or its mean (usually denoted *MAP* for maximum a posteriori or *EPM* for expected posterior mean).

We define here the standard incomplete-data scheme as follow: we consider that the observable incomplete data y having the parametric distribution $g(y; \theta)$ results from partial observation of complete data (y, z) having the parametric distribution $f(y, z; \theta)$, where g and f are some known density functions, z referring to the missing data. Maximum likelihood estimation of θ consists then in computing the value of θ that maximizes the observed likelihood g , considering the observation y are given. If the likelihood of the observation g can not be maximized directly, it is possible to do iterative maximization steps in order to approach the maximum. For example, the expectation-maximization (EM) algorithm, proposed by Dempster, Laird, and Rubin in year 1977, is a broadly applicable approach for the iterative computation of maximum likelihood estimates, useful in a variety of incomplete-data (or partially-observed-data) statistical problems. We recall here briefly the principle of the EM algorithm, which is detailed in the introduction. The EM algorithm is a two-step iterative procedure. One iteration is composed by the E-step and the M-step. The E-step computes $Q(\theta|\theta_k) = E[\log f(y, z; \theta)|y; \theta_k]$ and the M-step determines θ_{k+1} as maximizing $Q(\theta|\theta_k)$. Then, as presented in the proposition in the introduction, the observed-data likelihood sequence $(g(y; \theta_k))$ is nondecreasing along any EM sequence (see Dempster et al. (1977) and Wu (1983) for more details). The convergence to a stationary point of the observed likelihood occurs under some regularity assumptions and to a saddle point or a local maximum provided that checking some additional assumptions. Stochastic versions of EM have been introduced from different perspectives to deal with situations where the E-step is infeasible in closed form. This is often the case, because of the expectation which has no analytical expression. So Monte Carlo EM (MCEM) replaces the E-step by a Monte Carlo approximation of the expectation based on a large number of independent simulations of the missing data, see Meng and Rubin (1993). An other way to get round the difficulty of the computation of the expectation is also to use simulation, not for approximating some integral deriving from an expectation, but for getting some numerical plausible value standing for the missing data. One proposition was to simulate the missing data from the a posteriori distribution with the current values of the parameters (see Diebolt and Celeux (1993)). To be more efficient, the SAEM algorithm, proposed by Delyon, Lavielle, and Moulines (1999), replaces the E-step not only by the simulation of the missing data, but by stochastic approximation involving these simulated data. At iteration k , SAEM generates $m(k)$ realizations $z_k(j)$ ($1 \leq j \leq m(k)$) from

the a posteriori distribution denoted $p(z|y; \theta_k)$ and updates $Q_{k-1}(\theta)$ according to

$$Q_k(\theta) = Q_{k-1}(\theta) + \gamma_k \left(\frac{1}{m(k)} \sum_{j=1}^{m(k)} \log f(y, z_k(j); \theta) - Q_{k-1}(\theta) \right),$$

where (γ_k) is a sequence of positive step-sizes decreasing to 0. It is possible to select $m(k) = 1$ for all k when simulation of z_k is heavy. The M-step consists in determining θ_{k+1} which maximizes $Q_k(\theta)$. Precise results of convergence of SAEM are presented in Delyon et al. (1999) in the case where $f(y, z; \theta)$ belongs to a regular curved exponential family. In this case, SAEM can be written in terms of the complete-data sufficient statistics. This leads to a general Robbins-Monro type scheme and the almost sure convergence of the sequence (θ_k) to a local maximum of the likelihood is proved under very general assumptions by Delyon et al. (1999) like it is recalled in the introduction of this thesis. At the same time as the estimation of the parameters, the likelihood and the Fisher Information matrix of the model can also be estimated by stochastic approximation.

Unfortunately, for most nonlinear models or non-Gaussian models, the unobserved data cannot be simulated exactly under the conditional distribution. A well-known alternative consists in using a Monte Carlo Markov Chain (MCMC) method: that means to introduce an ergodic Markov chain which transition probability has as unique invariant distribution the conditional distribution p we want to simulate (see for example Robert (1996)). In this situation, the assumptions of Delyon et al. (1999) that ensure the convergence of SAEM are no longer satisfied. The aim of this chapter is to show that these assumptions can be weakened such that the algorithm SAEM still converges when it is coupled with a Markov chain Monte Carlo procedure. The results of Benveniste et al. (1990) for stochastic approximation with little Markovian perturbations are used to establish the convergence of this algorithm instead of results of the martingale theory which were used by Delyon, Lavielle, and Moulines (1999) to prove their convergence theorem.

Anyway, the use of simulated data for estimating parameters in missing data statistical problems is a powerful approach that tends to become popular since the 90's. In his paper published in 2000, Yao defines and studies an online stochastic approximation scheme. In this situation, the number of observations goes to infinity and the sequence of estimate converges to the true value of the parameter. In their paper published in 1998, Gu and Kong propose a stochastic version of a Newton-Raphson algorithm for incomplete data estimation. The algorithm proposed by Gu and Zhu (2001) combines an MCMC procedure with stochastic approximation for spatial models estimation. In these two papers, the Fisher information matrix is also estimated by stochastic approximation. In 2002, Concordet and Nunez propose a pseudo simulated maximum likelihood method for estimating the parameters of a nonlinear mixed effects model. By the way, Fort and Moulines (2003) realize an extension of the MCEM algorithm of Wei and Tanner (1990), drawing the unobserved data from a transition probability of an MCMC procedure.

Applications of SAEM to the change-points model was proposed in Lavielle and Lebarbier (2001) and to the convolution model in Lavielle and Moulines (1997), but no results of convergence of the algorithm were given. We show that the convergence results obtained for SAEM coupled with an MCMC procedure are very general and apply it to these both examples. A Monte Carlo experiment illustrates the performances of the proposed procedure with these both models.

1.2 Coupling the SAEM algorithm with MCMC

Assuming we are not able to simulate the missing data z_k at iteration k from the posterior distribution $p(\cdot|y; \theta_k)$, we propose to introduce an MCMC procedure to get the simulated value z_k . So in practice only the simulation step of the SAEM algorithm is changed. Given some objective distribution to reach, a MCMC procedure consists in generating an ergodic Markov chain having it as limiting distribution. Assume for the time being that the conditional distribution $p(\cdot|y; \theta)$ is the unique limiting distribution of some transition probability Π_θ , for any $\theta \in \Theta$. We will see later that there exists some method to ensure the existence of such probability transition. We could use this tool for generating a Markov chain at *each* iteration k of the algorithm from the transition probability Π_{θ_k} in order to get some realization z_k of the missing data having asymptotically the density function $p(\cdot; \theta_k)$. But this will be very expensive in time cost. So we propose to combine both convergences, on one hand this of the chain generating by the MCMC procedure, on the other hand this of the sequence (θ_k) . To that purpose, we consider one inhomogeneous Markov chain generating as follow: at iteration k , we draw the realization z_k from the previous value z_{k-1} from the transition probability Π_{θ_k} . So the k -th iteration of the proposed algorithm can be summarized in three steps as follows:

- **S-step:** using z_{k-1} , generate a realization z_k from the transition probability $\Pi_{\theta_k}(z_{k-1}, \cdot)$.
- **A-step:** update s_{k-1} according to (7).
- **M-step:** update θ_k according to (8).

Usually, Π_θ will be defined as the succession of M iterations of a MCMC procedure, such as the Metropolis-Hastings algorithm (see Chapter 2 of this thesis for some applications and for example Robert (1996) for more details). Then, the simulation step of iteration k consists in simulating z_k with the transition probability $\Pi_{\theta_k}(z_{k-1}, dz_k) = P_{\theta_k}^M(z_{k-1}, dz_k)$, where

$$P_{\theta_k}(z, dz') = q_{\theta_k}(z, z') \min \left\{ \frac{p(z'|y; \theta_k) q_{\theta_k}(z', z)}{p(z|y; \theta_k) q_{\theta_k}(z, z')}, 1 \right\} dz' \quad (1.1)$$

for $z' \neq z$ and $P_{\theta_k}(z, \{z\}) = 1 - \int_{z' \neq z} P_{\theta_k}(z, dz')$, where $q_\theta(z, z')$ is any aperiodic recurrent transition density, called instrumental distribution or proposal distribution. This procedure always accepts the new value z' if the likelihood ratio $\frac{p(z'|y; \theta_k)}{q_{\theta_k}(z, z')}$ is upper than the previous ratio.

For example, we can use the marginal distribution π of z_k as a proposal distribution. Then, writing $f(y, z; \theta) = \pi(z; \theta) h(y|z; \theta)$, the acceptance probability defined in (1.1) only depends on the conditional distribution h of the observation y which is often easy to evaluate:

$$P_{\theta_k}(z, dz') = \pi(z'; \theta_k) \min \left\{ \frac{h(y|z'; \theta_k)}{h(y|z; \theta_k)}, 1 \right\} dz'. \quad (1.2)$$

1.3 Convergence result

To prove the convergence of the proposed method which combines the SAEM algorithm with a MCMC procedure, we will use some results that work for general Robbins-Monro type stochastic approximation procedures. So if we write the recursion (7) into the general form of such algorithms, it becomes:

$$s_k = s_{k-1} + \gamma_k h(s_{k-1}) + \gamma_k e_k \quad (1.3)$$

where h stands for the mean field of the algorithm and e_k is a random perturbation. In our case, we obtain the following expressions for the function h and the sequence (e_k) :

$$\begin{aligned} h(s) &= E_{\hat{\theta}(s)}(\tilde{S}(y, z)) - s = \bar{s}(\hat{\theta}(s)) - s \\ e_k &= \tilde{S}(z_k) - E[\tilde{S}(z_k) | \mathcal{F}_{k-1}] = \tilde{S}(z_k) - \bar{s}(\hat{\theta}(s_{k-1})) \end{aligned}$$

where $E_{\theta}[\Phi(z)] \triangleq \int \Phi(z) p(z|y; \theta) \mu(dz)$.

The first point of assumption **(SAEM3)** detailed in the introduction means that, given $\theta_0, \dots, \theta_k$, the random variables $z_0, \dots, z_k$ are independent. Coupled with the second point of assumption **(SAEM3)**, it is sufficient to prove that the series $\sum \gamma_k e_k$ converges thanks to the martingale theory. This property plays a key role in the proof of the convergence of the SAEM algorithm in Delyon et al. (1999). In this new version of the algorithm, we allow Markovian dependence between z_k and z_{k+1} : we suppose that z_{k+1} is obtained from z_k due to a Markovian transition depending on θ_{k+1} , so that we will need some technical tools presented in Benveniste et al. (1990) to show that the series $\sum \gamma_k e_k$ still converges.

As announced above, the assumption **(SAEM3)** can be weakened. It is assumed that the random variables $s_0, z_1, z_2, \dots, z_k, \dots$ are defined on the same probability space $(\Omega, \mathcal{A}, P)$. We denote $\mathcal{F} = \{\mathcal{F}_k\}_{k \geq 0}$ the increasing family of σ -algebras generated by the random variables $s_0, z_1, z_2, \dots, z_k$. Since we introduce the transition probability Π_{θ} in order to approach the conditional distribution, we have to make some assumptions on it:

- **(SAEM3')**

- The chain $(z_k)_{k \geq 0}$ takes its values in a compact subset $\mathcal{E}$ of $\mathbb{R}^l$.
- For any compact subset V of Θ , there exists a real constant L such that for any (θ, θ') in V^2

$$\sup_{(x,y) \in \mathcal{E}^2} |\Pi_{\theta}(x, y) - \Pi_{\theta'}(x, y)| \leq L|\theta - \theta'|.$$

- The transition probability Π_{θ} generates a uniformly ergodic chain whose invariant probability is the conditional distribution $p(\cdot|y; \theta)$:

$$\exists K_{\theta} \in \mathbb{R}^+ \quad \exists \rho_{\theta} \in]0, 1[\quad | \quad \forall z \in \mathcal{E} \quad \forall k \in \mathbb{N} \quad \|\Pi_{\theta}^k(z, \cdot) - p(\cdot|y; \theta)\|_{TV} \leq K_{\theta} \rho_{\theta}^k,$$

where $\|\cdot\|_{TV}$ denotes the total variation norm. We suppose also that:

$$K \triangleq \sup_{\theta} K_{\theta} < +\infty \quad \text{and} \quad \rho \triangleq \sup_{\theta} \rho_{\theta} < 1.$$

- The function $\tilde{S}$ is bounded on $\mathcal{E}$.

We also weaken the assumption **(SAEM1)** and replace it by assumption **(SAEM1')**:

- **(SAEM1')** For all k in $\mathbb{N}$, $\gamma_k \in [0, 1]$, $\sum_{k=1}^{\infty} \gamma_k = \infty$ and there exists λ in $]\frac{1}{2}, 1]$ such that $\sum_{k=1}^{\infty} \gamma_k^{1+\lambda} < \infty$.

With these modified assumptions, we obtain the following convergence result:

Theorem 1.1. *Assume that assumptions **(M1)-(M5)**, **(SAEM1')**, **(SAEM2)** and **(SAEM3')** hold. Assume in addition the assumption **(C)**: the sequence $(s_k)_{k \geq 0}$ takes its values in a compact subset of $\mathcal{S}$. Then, w.p. 1, $\lim_{k \rightarrow +\infty} d(\theta_k, \mathcal{L}) = 0$ where $d(x, A)$ denotes the distance of x to the closed subset A and $\mathcal{L} = \{\theta \in \Theta, \partial_{\theta} l(y; \theta) = 0\}$ is the set of stationary points of l .*

Remark. For checking the third point of assumption **(SAEM3')**, it is possible to verify some minorization condition or Doeblin's condition for the transition probability Π_θ (see Chapter 16 of Meyn and Tweedie (1993)). Otherwise, we have to consider each case individually. Consider for example an independent Metropolis-Hastings algorithm: the transition density q_θ mentioned in equation (1.1) defines an independence sampler:

$$\forall (z, z') \in \mathcal{E}^2 \quad q_\theta(z, z') = q_\theta(z').$$

Then the uniform ergodicity is ensured if the transition q_θ satisfies the following inequality (see Theorem 2.1 in Mengersen and Tweedie (1996)):

$$\exists \beta \in \mathbb{R}^+ \quad | \quad \forall z \in \mathcal{E} \quad q_\theta(z) \geq \beta p(z|y; \theta).$$

Remark. In cases where the compactness condition **(C)** is not satisfied or is difficult to check, it is possible to stabilize the algorithm by using the method of dynamic bounds proposed by Chen, Lei, and Gao (1988) and already used in this context in Delyon et al. (1999). Consider a sequence of compact subsets (K_k) such that $K_k \subset \text{int}(K_{k+1})$ and recovering all the set $\mathcal{S}$. If s_k is outside the given compact set K_k of $\mathcal{S}$, it is reinitialized in a specific compact set K_0 .

Proof of Theorem 1.1:

Theorem 1.1 is an application of Theorem 2 presented in Delyon et al. (1999), which gives a general result about convergence of Robbins-Monro type stochastic approximation procedures of the form (1.3):

Theorem (Delyon et al. (1999)). Assume that:

- **(SA0)** w.p.1, $\forall k \geq 0, s_k \in \mathcal{S}$.
- **(SA1)** $(\gamma_k)_{k \geq 0}$ is a decreasing sequence of positive number such that $\sum_{k=1}^{\infty} \gamma_k = \infty$.
- **(SA2)** The vector field h is continuous on $\mathcal{S}$ and there exists a continuously differentiable function $V : \mathcal{S} \rightarrow \mathbb{R}$ such that:

- $\forall s \in \mathcal{S} \quad F(s) = \langle \partial_s V(s), h(s) \rangle \leq 0$.
- $\text{int}(V(\mathcal{L}_0)) = \emptyset$, where $\mathcal{L}_0 \triangleq \{s \in \mathcal{S}, F(s) = 0\}$.

- **(SA3)** w.p.1, the closure of $(\{s_k\}_{k \geq 0})$ is a compact subset of $\mathcal{S}$.
- **(SA4)** w.p.1, $\lim_{n \rightarrow \infty} \sum_{k=0}^n \gamma_k e_k$ exists and is finite.

Then, w.p.1, $\limsup d(s_k, \mathcal{L}) = 0$.

The assumptions **(SA0)**-**(SA3)** don't use the dependence structure of the sequence (z_k) and are satisfied under **(M1-M5)**, **(SAEM1')** and **(SAEM2)** like in the case where the missing data are exactly simulated under the a posteriori distribution (see Delyon et al. (1999)). Under the assumption **(SAEM3)**, the condition **(SA4)** is satisfied because the sequence of the partial sum of $\sum \gamma_k e_k$ is a convergent martingale. For checking **(SA4)** under **(SAEM3')**, we use a result presented in Proposition 7 of Chapter 1 of Benveniste et al. (1990). The general model of the algorithm considered by Benveniste *et al.* is of the form:

$$s_k = s_{k-1} + \gamma_k H(s_{k-1}, z_k),$$

where the sequence $(s_k)_{k \geq 0}$ evolves in $\mathbb{R}^m$ and the sequence $(z_k)_{k \geq 0}$ in $\mathbb{R}^l$. In our algorithm, the function H is equal to:

$$H(s, z) = \tilde{S}(y, z) - s,$$

where the dependence of H in y is not written since the observation y are considered as constant.

Proposition (Benveniste et al. (1990)). *Assume the following assumptions:*

- **(A1)** $(\gamma_k)_{k \geq 0}$ is a decreasing sequence of positive real numbers such that $\sum \gamma_k = +\infty$.
- **(A2)** There exists a family $\{\Pi_{\hat{\theta}(s)}, s \in \mathbb{R}^m\}$ of transition probabilities on $\mathbb{R}^l$ such that for any Borel subset A of $\mathbb{R}^l$, we have

$$P(z_k \in A | \mathcal{F}_{k-1}) = \Pi_{\hat{\theta}(s_{k-1})}(z_{k-1}, A).$$

- **(A3)** For any compact subset Q of $\mathcal{S}$, there exists a constant C_1 depending on Q such that for all s in Q and all z in $\mathcal{E}$ we have

$$|H(s, z)| \leq C_1.$$

- **(A4)** There exists a function h on $\mathcal{S}$ and for each s in $\mathcal{S}$ a function ν_s on $\mathbb{R}^l$ such that
 - h is locally Lipschitz on $\mathcal{S}$: for all s in $\mathcal{S}$, there exist a neighborhood $\mathcal{V}$ of s and a real constant C such that

$$\forall (s', s'') \in \mathcal{V}^2 \quad |h(s') - h(s'')| \leq C|s' - s''|.$$

- $(I - \Pi_{\hat{\theta}(s)})\nu_s = H_s - h(s)$ for all s in $\mathcal{S}$, where for all z in $\mathcal{E}$, $\Pi_{\hat{\theta}(s)}\nu_s(z) \triangleq \int \nu_s(z') \Pi_{\hat{\theta}(s)}(z, dz')$.
- For any compact subset Q of $\mathcal{S}$, there exist real constants C_2, C_3 and λ in $]\frac{1}{2}, 1]$, such that for all s and s' in Q and for all z in $\mathcal{E}$

$$|\nu_s(z)| \leq C_2$$

$$|\Pi_{\hat{\theta}(s)}\nu_s(z) - \Pi_{\hat{\theta}(s')}\nu_{s'}(z)| \leq C_3|s - s'|^\lambda.$$

- **(A5)** For any compact subset Q of $\mathcal{S}$ and any positive real q , there exists a real number $\mu_q(Q)$ such that, for all n , for all z_0 in $\mathbb{R}^l$ and all s_0 in $\mathbb{R}^m$

$$E_{z_0, s_0} \left((1 + |z_k|^q) \mathbb{1}(s_{k-1} \in Q, k \leq n) \right) \leq \mu_q(Q)$$

where E_{z_0, s_0} denotes the expectation under the distribution of $(z_k, s_k)_{k \geq 0}$ for the initial conditions (z_0, s_0) .

Then, for any compact subset Q of $\mathcal{S}$, denoting $\tau(Q) = \inf\{k, s_k \notin Q\}$, if the constant λ from **(A4)** verifies $\sum \gamma_k^{1+\lambda} < +\infty$, then, on $\{\tau(Q) = \infty\}$, the series $\sum_k \gamma_k e_k$ converges a.s. and in L^2 .

Remark. The proposition of Benveniste et al. requires $\sum \gamma_k^{1+\lambda} \leq 1$, but we only need $\sum \gamma_k^{1+\lambda} < +\infty$ to obtain the convergence of the series $\sum_k \gamma_k e_k$.

In order to apply this proposition in our case, we have to check the assumptions **(A1)-(A5)**:

- **(A1)** is implied by **(SAEM1')**.
- **(A2)** is verified with the transition probability $\Pi_{\hat{\theta}(s)}$.
- **(A3)** is verified since $\tilde{S}$ is bounded on $\mathcal{E}$.
- **(A5)** is verified since the sequence $(z_k)_{k \geq 0}$ takes its values in a compact subset $\mathcal{E}$ of $\mathbb{R}^l$.

Consider now the mean field of the algorithm defined by:

$$h(s) \triangleq \bar{s}(\hat{\theta}(s)) - s.$$

We will show that this function h satisfies **(A4)**. The assumptions **(M3)** and **(M5)** imply that the function h is continuously differentiable on $\mathcal{S}$. So h is locally Lipschitz on $\mathcal{S}$ and assumption **(A4)1** is satisfied.

We define now:

$$\begin{aligned}\nu_s(z) &\triangleq \sum_{k \geq 0} \Pi_{\hat{\theta}(s)}^k (H(s, z) - h(s)) \\ &= \sum_{k \geq 0} \left(\Pi_{\hat{\theta}(s)}^k \tilde{S}(y, z) - \bar{s}(\hat{\theta}(s)) \right) \\ &= \sum_{k \geq 0} \left(\Pi_{\hat{\theta}(s)}^k \tilde{S}(y, z) - p_{\hat{\theta}(s)} \tilde{S} \right),\end{aligned}$$

where $p_{\hat{\theta}(s)} \tilde{S} \triangleq \int \tilde{S}(y, z) p_{\hat{\theta}(s)}(z|y) dz$.

Since $\Pi_{\hat{\theta}(s)}$ is uniformly ergodic, there exist $K_{\hat{\theta}(s)} \in \mathbb{R}^+$ and $\rho_{\hat{\theta}(s)} \in]0, 1[$, such that, for any $k \in \mathbb{N}$, for any z in $\mathcal{E}$

$$\sup_{\|u\| \leq 1} \left| \Pi_{\hat{\theta}(s)}^k u(z) - p_{\hat{\theta}(s)} u \right| \leq K_{\hat{\theta}(s)} \rho_{\hat{\theta}(s)}^k.$$

Since $\tilde{S}$ is bounded on $\mathcal{E}$, the series defining ν_s is convergent. Moreover, we have for all z in $\mathcal{E}$:

$$\begin{aligned}(I - \Pi_{\hat{\theta}(s)}) \nu_s(z) &= (I - \Pi_{\hat{\theta}(s)}) \sum_{k \geq 0} \left(\Pi_{\hat{\theta}(s)}^k \tilde{S}(y, z) - p_{\hat{\theta}(s)} \tilde{S} \right) \\ &= \sum_{k \geq 0} \left(\Pi_{\hat{\theta}(s)}^k \tilde{S}(y, z) - \Pi_{\hat{\theta}(s)}^{k+1} \tilde{S}(y, z) \right) \\ &= \tilde{S}(y, z) - p_{\hat{\theta}(s)} \tilde{S},\end{aligned}$$

which proves **(A4)2**.

Under the third point of assumption **(SAEM3')** and since $\tilde{S}$ is bounded on $\mathcal{E}$, we obtain the first inequality of **(A4)3**. We will now prove the second inequality of **(A4)3**. Consider the following decomposition:

$$\begin{aligned}\Pi_{\hat{\theta}(s)} \nu_s(z) - \Pi_{\hat{\theta}(s')} \nu_{s'}(z) &= \sum_{k \geq 1} \left(\Pi_{\hat{\theta}(s)}^k \tilde{S}(y, z) - p_{\hat{\theta}(s)} \tilde{S} \right) - \sum_{k \geq 1} \left(\Pi_{\hat{\theta}(s')}^k \tilde{S}(y, z) - p_{\hat{\theta}(s')} \tilde{S} \right) \\ &= \nu_s(z) - \nu_{s'}(z) + p_{\hat{\theta}(s')} \tilde{S} - p_{\hat{\theta}(s)} \tilde{S},\end{aligned}$$

so we have

$$|\Pi_{\hat{\theta}(s)} \nu_s(z) - \Pi_{\hat{\theta}(s')} \nu_{s'}(z)| \leq |\nu_s(z) - \nu_{s'}(z)| + |p_{\hat{\theta}(s)} \tilde{S} - p_{\hat{\theta}(s')} \tilde{S}|.$$

We use the following technical lemma proved in the appendix to prove the second inequality of **(A4)3**:

Lemma 1.2. *If we assume **(SAEM2)**, **(SAEM3')** and **(C)**, then for any compact subset $\mathcal{Q}$ of $\mathcal{S}$, there exist K_1 and K_2 in $\mathbb{R}^+$ such that, for any $\alpha \in]0, 1[$, for any $(s, s', z) \in \mathcal{Q}^2 \times \mathcal{E}$,*

$$|p_{\hat{\theta}(s)} \tilde{S} - p_{\hat{\theta}(s')} \tilde{S}| \leq K_1 |s - s'|^\alpha \quad \text{and} \quad |\nu_s(z) - \nu_{s'}(z)| \leq K_2 |s - s'|^\alpha.$$

This result implies that for any compact subset $\mathcal{Q}$ of $\mathcal{S}$, there exist K in $\mathbb{R}^+$ and λ in $]\frac{1}{2}, 1]$ such that for any $(s, s', z) \in \mathcal{Q}^2 \times \mathcal{E}$:

$$|\Pi_{\hat{\theta}(s)} \nu_s(z) - \Pi_{\hat{\theta}(s')} \nu_{s'}(z)| \leq K|s - s'|^\alpha,$$

concluding the proof of **(A4)3**.

All the assumptions of the proposition 7 of Benveniste et al. (1990) are satisfied so it can be applied to prove **(SA4)**. Then Theorem 2 of Delyon et al. (1999) allows us to prove the a.s. convergence of the sequence $(s_k)_{k \geq 0}$ to a point of the set $\mathcal{L}_0$. To conclude, we use the Lemma 2 of Delyon et al. (1999) to prove that the sequence $(\theta_k)_{k \geq 0}$ converges a.s. to a point of the set $\mathcal{L}$ since the function $\hat{\theta}$ is continuous. This closes the proof of Theorem 1.1 which is equivalent to Theorem 2 of Delyon, Lavielle, and Moulines (1999). $\square$

We also obtain in the case where the missing data are not simulated under the posterior distribution, but from an MCMC procedure, a convergence result to a proper maximizer for some additional assumptions **(LOC1)**-**(LOC3)** presented in the introduction:

Theorem 1.3. *Assume the following assumptions **(M1)**-**(M5)**, **(SAEM1')**, **(SAEM2)**, **(SAEM3')** and **(LOC1)**-**(LOC3)** are satisfied, as well as **(C)** : the sequence $(s_k)_{k \geq 0}$ takes its values in a compact subset of $\mathcal{S}$. Then, w.p. 1 $\lim_{k \rightarrow +\infty} \theta_k = \theta^*$ where θ^* is a proper maximizer of the observed log-likelihood l .*

It means that the algorithm SAEM a.s. avoids traps, i.e., it can only converge to a proper local maximizer of the likelihood. The proof is the same as this done by Delyon, Lavielle, and Moulines (1999). It is based on the result of Brandière and Duflo (1995) presented in paragraph II.3. about little Markovian perturbations which gives some sufficient conditions for avoiding the local minima or saddle points of l .

1.4 Applications

1.4.1 The deconvolution problem

In a convolution model, the observation $\mathbf{y} = (y_{L+1}, \dots, y_n)$ is the linear convolution of an unobserved input sequence $\mathbf{z} = (z_1, \dots, z_n)$ with additive noise ε , e.g.,

$$y_t = \sum_{l=0}^L \varphi_l z_{t-l} + \sigma \varepsilon_t, \quad L+1 \leq t \leq n \quad (1.4)$$

where $\varphi = (\varphi_0, \dots, \varphi_L)$ is the convolution filter and σ a positive parameter corresponding to the standard deviation of the additive noise.

This kind of model is commonly used in seismic deconvolution, fMRI data analysis and statistical signal analysis.

We shall make the following assumptions, concerning the random sequences $\mathbf{z}$ and ε : **(i)** $(z_t, 1 \leq t \leq n)$ is a sequence of independent and identically distributed random variables, distributed according to some distribution function π , and taking their values in some compact subset of $\mathbb{R}$. **(ii)** $(\varepsilon_t, L+1 \leq t \leq n)$ is a sequence of independent standardized Gaussian variables, **(iii)** z_t and $\varepsilon_{t'}$ are independent for all pairs of time indexes $1 \leq t \leq n$ and $L+1 \leq t' \leq n$.

These assumptions together with equation (1.4) specify completely the log-likelihood of the observed data samples. Let $\mathcal{M}(\mathbf{z})$ be the $(n-L) \times (L+1)$ matrix such that $\mathcal{M}_{ij}(\mathbf{z}) = z_{L+1+i-j}$. So the distribution of $\mathbf{y}$ conditionally to $\mathbf{z}$, denoted h , is a Gaussian distribution with mean $\mathcal{M}(\mathbf{z})$ and variance $\sigma^2 I_{n-L}$, where I_m denotes the identity matrix of size m . Then, the complete log-likelihood is, up to a constant,

$$\begin{aligned} \log f(\mathbf{y}, \mathbf{z}; \theta) &= \log h(\mathbf{y}|\mathbf{z}; \theta) + \log \pi(\mathbf{z}) \\ &= C - \frac{n-L}{2} \log \sigma^2 - \frac{1}{2\sigma^2} \|\mathbf{y} - \mathcal{M}(\mathbf{z}) \boldsymbol{\varphi}\|^2 + \log \pi(\mathbf{z}), \end{aligned} \quad (1.5)$$

where C is a constant.

Deconvolution consists in recovering the input sequence $\mathbf{z}$ from the observation $\mathbf{y}$. Of course, deconvolution requires an accurate estimation of the convolution filter $\boldsymbol{\varphi}$ and of the noise variance σ^2 . SAEM will be very useful for estimating these parameters. Furthermore, the marginal distribution π of the input sequence $\mathbf{z}$ can also be estimated, whenever it belongs to the exponential family and depends on an unknown parameter ψ :

$$\pi(\mathbf{z}; \psi) = C(\psi) \exp \left\{ - \left\langle \tilde{S}_\pi(\mathbf{z}), \psi \right\rangle \right\}. \quad (1.6)$$

Here, the vector of parameters of the model is $\theta = (\psi, \boldsymbol{\varphi}, \sigma^2)$. For estimating them by the SAEM algorithm, we define the minimal sufficient statistics of the model according to (4). Since the log-likelihood of the complete data is expressed by:

$$\log f(\mathbf{y}, \mathbf{z}; \theta) = C' - \frac{n-L}{2} \log \sigma^2 - \frac{1}{2\sigma^2} \|\mathbf{y} - \mathcal{M}(\mathbf{z}) \boldsymbol{\varphi}\|^2 - \left\langle \tilde{S}_\pi(\mathbf{z}), \psi \right\rangle, \quad (1.7)$$

we choose as minimal sufficient statistics $\tilde{S}(\mathbf{y}, \mathbf{z}) = (\mathcal{M}(\mathbf{z})^t \mathcal{M}(\mathbf{z}), \mathcal{M}(\mathbf{z})^t \mathbf{y}, \tilde{S}_\pi(\mathbf{z}))$. At step k , we used the following procedure for generating $\mathbf{z}_k$ from $\mathbf{z}_{k-1}$, using $\theta_k = (\psi_k, \boldsymbol{\varphi}_k, \sigma_k^2)$:

- i) a permutation p_k of $\{1, 2, \dots, n\}$ is randomly chosen,
- ii) for $i = 1, 2, \dots, n$:
 1. Let $j = p_k(i)$. Set $\tilde{z}_t = z_{k-1,t}$ for any $t \neq j$ and generate $\tilde{z}_j \sim \pi(\cdot; \psi_k)$.
 2. Compute

$$\alpha = \frac{1}{2\sigma_k^2} (\|\mathbf{y} - \mathcal{M}(\tilde{\mathbf{z}}) \boldsymbol{\varphi}_k\|^2 - \|\mathbf{y} - \mathcal{M}(\mathbf{z}) \boldsymbol{\varphi}_k\|^2). \quad (1.8)$$

- 3. Generate $u \sim \text{Exp}(1)$. Set $\mathbf{z}_k = \tilde{\mathbf{z}}$ if $\alpha < u$ and $\mathbf{z}_k = \mathbf{z}_{k-1}$ otherwise.

We can easily show that $\Pi_{\theta_k}(\mathbf{z}_{k-1}, \cdot)$ is the transition probability of an ergodic Markov chain, that converges uniformly to the conditional distribution $p(\cdot|\mathbf{y}; \theta_k)$. We define then a sequence $(s_{k,1}, s_{k,2}, s_{k,3})$ according to (7) :

$$\begin{aligned} s_{k,1} &= s_{k-1,1} + \gamma_k (\mathcal{M}(\mathbf{z}_k)^t \mathcal{M}(\mathbf{z}_k) - s_{k-1,1}) \\ s_{k,2} &= s_{k-1,2} + \gamma_k (\mathcal{M}(\mathbf{z}_k)^t \mathbf{y} - s_{k-1,2}) \\ s_{k,3} &= s_{k-1,3} + \gamma_k (\tilde{S}_\pi(\mathbf{z}_k) - s_{k-1,3}). \end{aligned}$$

Then, the maximization step yields directly

$$\begin{aligned} \boldsymbol{\varphi}_{k+1} &= (s_{k,1})^{-1} s_{k,2} \\ \sigma_{k+1}^2 &= \frac{1}{n-L} \left(\mathbf{y}^t \mathbf{y} - (s_{k,2})^t \boldsymbol{\varphi}_{k+1} \right) \\ \psi_{k+1} &= \text{Argmax } C(\psi) e^{-\langle s_{k,3}, \psi \rangle}. \end{aligned}$$

All the assumptions of Theorem 1.1 are satisfied whenever π has bounded support, and the SAEM algorithm converges to a (local) maximum of the observed likelihood. To assume that the input variables $\mathbf{z}$ are bounded is not a restrictive assumption from a practical point of view. Indeed, any nonbounded distribution is truncated in practice.

We will now proceed to a numerical implementation and used the convolution model described in (1.4) for simulating an observed series $\mathbf{y}$ of length $n = 500$. In this example, the input sequence $\mathbf{z}$ are independent $Beta(a, b)$ random variables on $[0, 1]$ with $a = b = 3$. So the statistic $\tilde{S}_\pi(\mathbf{z})$ is $(S_1(\mathbf{z}), S_2(\mathbf{z})) = (\sum \log(z_j), \sum \log(1 - z_j))$ since the density function of a random variable $Beta(a, b)$ is $z^{a-1}(1 - z)^{b-1}\mathbb{1}_{[0,1]}(z)/B(a, b)$ where $B(a, b) = \Gamma(a)\Gamma(b)/\Gamma(a + b)$ and $\Gamma(a) = \int_0^{+\infty} e^{-x}x^{a-1}dx$. The convolution filter is $\varphi = (1, -3, 2, 6, 2, -3, 1)$. We choose $\sigma^2 = 0.2286$ in order to ensure a Signal to Noise ratio equal to 10dB, *i.e.* $\text{Var}(\mathcal{M}(\mathbf{z})\varphi) = 10\text{Var}(\sigma\varepsilon)$.

Table 1.1 gives the estimation of $\theta = (a, b, \varphi, \sigma^2)$. A Monte-Carlo experiment based on 100 replications was used for estimating the mean and the standard deviation of two estimators. First, the maximum likelihood estimator $\hat{\theta}_f$, which maximizes the complete likelihood $f(\mathbf{y}, \mathbf{z}; \theta)$ assuming that the input series $\mathbf{z}$ is known, was computed as follows by direct maximization of (1.7):

$$\begin{aligned}\hat{\varphi}_f &= (\mathcal{M}(\mathbf{z})^t \mathcal{M}(\mathbf{z}))^{-1} \mathcal{M}(\mathbf{z})^t \mathbf{y} \\ \hat{\sigma}_f^2 &= \frac{1}{n - L} \left(\mathbf{y}^t \mathbf{y} - \mathbf{y}^t \mathcal{M}(\mathbf{z}) \hat{\varphi}_f \right) \\ (\hat{a}_f, \hat{b}_f) &= \text{Arg max}_{a,b} \frac{1}{B(a, b)} e^{-(a-1)S_1(\mathbf{z}) - (b-1)S_2(\mathbf{z})}\end{aligned}$$

where $B(a, b) = \Gamma(a)\Gamma(b)/\Gamma(a + b)$ and $\Gamma(a) = \int_0^{+\infty} e^{-x}x^{a-1}dx$.

On the other hand, the estimator $\hat{\theta}_g$ maximizes the incomplete likelihood $g(\mathbf{y}; \theta)$, considering that the input series $\mathbf{z}$ is unknown.

We computed $\hat{\theta}_g$ with 100 iterations of SAEM, using $\gamma_k = 1$ for $1 \leq k \leq 30$ and $\gamma_k = 1/(k - 29)$ for $k \geq 31$.

The initialization chosen arbitrarily is the uniform distribution for $\mathbf{z}$ ($a_0 = b_0 = 0$). The initial guess for the filter is a spike located at $j = 4$. That ensures that the algorithm recovers the good phase of the convolution filter. For a different initialization, the algorithm can converge to a local maximum of the likelihood that cannot be compared with the true filter φ^* . For example, using as initial guess a spike at $j = 2$, a simulation gives $\hat{\varphi} = (0.73, 6.65, 1.42, -3.51, -0.58, 2.54, -1.32)$. We remark that the phase of the true filter is not recovered, but the estimated filter and the true filter have both almost the same transfer function. The problem of convergence to the global maximum of the likelihood is beyond the scope of this paper (see Lavielle and Moulines (1997) for a simulated annealing version of this algorithm).

The results presented in Table 1.1 confirm that $\hat{\theta}_f$ is a more accurate estimate of θ than $\hat{\theta}_g$. Nevertheless, we can remark that, when $\mathbf{z}$ is not observed, SAEM provides a good estimation of θ .

1.4.2 The change-points problem

We use the model considered in the paper of Lavielle and Lebarbier (2001). We observe a real sequence $\mathbf{y} = (y_i, 1 \leq i \leq n)$, such that, for any $1 \leq i \leq n$,

$$y_i = f(t_i) + \sigma\varepsilon_i, \quad (1.9)$$

TAB. 1.1: Estimation of $\theta = (a, b, \varphi, \sigma^2)$: θ^* is the true value of θ , θ_0 is the initialization, $\hat{\theta}_f$ is the estimation obtained by maximizing $f(\mathbf{y}, \mathbf{z}; \theta)$ and $\hat{\theta}_g$ is the estimation obtained by maximizing $g(\mathbf{y}; \theta)$.

	θ^*	θ_0	E $\hat{\theta}_f$	std $\hat{\theta}_f$	E $\hat{\theta}_g$	std $\hat{\theta}_g$
a	3	0	3.0075	0.1846	2.8615	0.4014
b	3	0	3.0220	0.1838	2.6396	0.3712
φ_1	1	0	1.0002	0.1020	1.0277	0.5820
φ_2	-3	0	-3.0058	0.0997	-2.6564	0.4612
φ_3	2	0	2.0100	0.1046	1.8809	0.8416
φ_4	6	1	5.9810	0.0963	5.5637	0.4024
φ_5	2	0	2.0034	0.1128	1.8963	0.7713
φ_6	-3	0	-2.9985	0.0955	-2.7920	0.5766
φ_7	1	0	1.0009	0.1047	0.8361	0.5523
σ^2	0.2286	1	0.2266	0.0157	0.2615	0.0593

where $(t_i, 1 \leq i \leq n)$ is a sequence of known observation times, σ is an unknown positive parameter and $(\varepsilon_i, 1 \leq i \leq n)$ is a sequence of independent zero-mean Gaussian variables with unit variance. The function f to recover is piecewise constant. Thus, there exists a sequence of instants $(\tau_j, j \geq 0)$ among the sequence $(t_i, 1 \leq i \leq n)$ and a sequence $(m_j, j \geq 1)$ such that, for any $j \geq 1$,

$$f(t) = m_j \text{ for all } \tau_{j-1} < t \leq \tau_j, \quad (1.10)$$

with the convention $\tau_0 = 0$.

We introduce a latent sequence of independent identically distributed Bernoulli random variables $(z_i, 1 \leq i \leq n-1)$ that take the value 1 at the change instants and 0 between two changes :

$$z_i = \begin{cases} 1 & \text{if there exists } j \text{ such that } t_i = \tau_j \\ 0 & \text{otherwise} \end{cases}. \quad (1.11)$$

Let λ be the unknown parameter of the Bernoulli and for any $\mathbf{z} = (z_i, 1 \leq i \leq n-1)$ in $\Omega = \{0, 1\}^{n-1}$, let $K_{\mathbf{z}} = \sum_{i=1}^{n-1} z_i + 1$ be the number of segments (*i.e.* $K_{\mathbf{z}} - 1$ is the number of change-points) defined by $\mathbf{z}$. Then,

$$\pi(\mathbf{z}; \lambda) = \lambda^{K_{\mathbf{z}}-1} (1 - \lambda)^{n-K_{\mathbf{z}}}. \quad (1.12)$$

Conditionally to the change-points sequence, the vector $\mathbf{m} = (m_j, 1 \leq j \leq K_{\mathbf{z}})$ is assumed to be Gaussian with the following density function:

$$\pi(\mathbf{m}|\mathbf{z}; \mu, V) = \prod_{j=1}^{K_{\mathbf{z}}} \left(\frac{2\pi V}{n_j} \right)^{-\frac{1}{2}} \exp \left\{ -\frac{n_j}{2V} (m_j - \mu)^2 \right\} \quad (1.13)$$

where $n_j = \sum_{i=1}^n \mathbb{1}_{[\tau_{j-1}, \tau_j]}(t_i)$ is the length of segment j for $1 \leq j \leq K_{\mathbf{z}}$. The parameters μ and V are unknown and such that each element of the mean of the vector $\mathbf{m}$ conditionally to $\mathbf{z}$ is

equal to μ and the variance of the vector $\mathbf{m}$ conditionally to $\mathbf{z}$ is equal to V multiply by the identity matrix of size $K_{\mathbf{z}}$.

Moreover, $(\varepsilon_i, 1 \leq i \leq n)$ is assumed to be a sequence of independent zero-mean and unit variance Gaussian random variables. Thus, the conditional distribution of the observations is defined by:

$$h(\mathbf{y}|\mathbf{z}, \mathbf{m}; \sigma^2) = (2\pi\sigma^2)^{-\frac{n}{2}} \exp \left\{ -\frac{1}{2\sigma^2} \sum_{j=1}^{K_{\mathbf{z}}} \sum_{l=N_{j-1}+1}^{N_j} (y_l - m_j)^2 \right\}, \quad (1.14)$$

where $N_j = \sum_{l=1}^j n_l$ for $1 \leq j \leq K_{\mathbf{z}}$ and $N_0 = 0$.

Now let $\theta = (\mu, \lambda, V, \sigma^2)$ be the set of parameters of the model and for any configuration of changes $\mathbf{z}$, let $\bar{y}_j = n_j^{-1} \sum_{l=N_{j-1}+1}^{N_j} y_l$, $\bar{y} = n^{-1} \sum_{l=1}^n y_l$ and $C_{\mathbf{z}} = \sum_{j=1}^{K_{\mathbf{z}}} \sum_{l=N_{j-1}+1}^{N_j} (y_l - \bar{y}_j)^2$. Then, after some calculation, it can be shown (see Lavielle and Lebarbier (2001)) that the complete log-likelihood of the model is defined by

$$\begin{aligned} \log f(\mathbf{y}, \mathbf{z}; \theta) &= C - \frac{n}{2} \log(\sigma^2) - \frac{K_{\mathbf{z}}}{2} \log \left(\frac{\sigma^2 + V}{\sigma^2} \right) + (K_{\mathbf{z}} - 1) \log \lambda + (n - K_{\mathbf{z}}) \log(1 - \lambda) \\ &\quad \times - \left\{ \frac{1}{2(V + \sigma^2)} \left(\sum_{i=1}^n (y_i - \mu)^2 + \frac{V}{\sigma^2} C_{\mathbf{z}} \right) \right\}. \end{aligned} \quad (1.15)$$

So according to (4), the minimal sufficient statistics to be approximated is $\tilde{S}(\mathbf{y}, \mathbf{z}) = (K_{\mathbf{z}}, C_{\mathbf{z}})$ and we define then according to (7) the following sequence:

$$\begin{aligned} s_{k,1} &= s_{k-1,1} + \gamma_k (K_{\mathbf{z}} - s_{k-1,1}) \\ s_{k,2} &= s_{k-1,2} + \gamma_k (C_{\mathbf{z}} - s_{k-1,2}). \end{aligned}$$

Then, the maximization step yields directly

$$\begin{aligned} \mu_k &= \bar{y} \\ \lambda_k &= \frac{s_{k,1} - 1}{n - 1} \\ \sigma_k^2 &= \frac{s_{k,2}}{n - s_{k,1}} \\ V_k &= \frac{\sum_{i=1}^n (y_i - \bar{y})^2 - s_{k,2}}{s_{k,1}} - \sigma_k^2, \end{aligned}$$

where $\bar{y} = \frac{1}{n} \sum_{i=1}^n y_i$.

Here, $\mathbf{z}$ takes a finite number of values. Then, any irreducible proposal kernel q can be used to generate a geometrically ergodic kernel. In this application, we used alternatively the three following kernels: *i*) a new vector $\tilde{\mathbf{z}}$ is drawn independently of the current value $\mathbf{z}_{k-1}$ with the marginal distribution $\pi(\cdot; \psi_{k-1})$, *ii*) a new change-point is created, or an existing change-point is removed, *iii*) an existing change-point is shifted (see Lavielle and Lebarbier (2001) for more details concerning the MCMC procedure).

TAB. 1.2: Estimation of $\theta = (\lambda, V, \sigma^2)$: θ^* is the true value of θ , θ_0 is the initialization, $\hat{\theta}_f$ is the estimation obtained by maximizing $f(\mathbf{y}, \mathbf{z}; \theta)$ and $\hat{\theta}_g$ is the estimation obtained by maximizing $g(\mathbf{y}; \theta)$.

	θ^*	θ_0	E $\hat{\theta}_f$	std $\hat{\theta}_f$	E $\hat{\theta}_g$	std $\hat{\theta}_g$
λ	0.02	0.05	0.0201	0.0051	0.0235	0.0105
V	40	10	38.1	15.8	32.5	19.0
σ^2	1	5	1.01	0.07	1.01	0.08

The maximum likelihood estimate of μ is $\bar{y}$. For estimating the others parameters, the SAEM algorithm can be used since it is easy to check that the assumptions of Theorem 1.1 are still satisfied.

We now propose some numerical application and used the change-point model described above for simulating an observed series $\mathbf{y}$ of length $n = 500$. We set the values of the parameters to $\lambda^* = 0.02$, $V^* = 40$ and $\sigma^{2*} = 1$.

As in the previous example, we compute two estimators: on one hand, $\hat{\theta}_f$, which maximizes the complete likelihood $f(\mathbf{y}, \mathbf{z}; \theta)$ assuming $\mathbf{y}$ and $\mathbf{z}$ are observed, on the other hand, $\hat{\theta}_g$ that maximizes the incomplete likelihood $g(\mathbf{y}; \theta)$. Table 1.2 gives the estimation of $\theta = (\lambda, V, \sigma^2)$. A Monte-Carlo experiment based on 100 replications was used for estimating the mean and the standard deviation of $\hat{\theta}_f$ and of $\hat{\theta}_g$. We computed $\hat{\theta}_g$ with 100 iterations of SAEM, using $\gamma_k = 1$ for $1 \leq k \leq 30$ and $\gamma_k = 1/(k - 29)$ for $k \geq 31$.

In this example, SAEM still produces a good estimation of θ . In particular, the noise variance σ^2 is estimated with almost the same accuracy, when the change-points are known and when they are unknown.

1.5 Appendix

Proof of theorem 1.1:

We first prove a technical lemma which will be necessary to prove Lemma 1.2.

Lemma 1.4. *If we assume (SAEM2), (SAEM3') and (C), then for any compact subset $\mathcal{Q}$ of $\mathcal{S}$, there exists a real constant M such that:*

$$\forall (s, s', y, z) \in \mathcal{Q}^2 \times \mathcal{Y} \times \mathcal{E} \quad \forall k \in \mathbb{N} \quad |\Pi_{\hat{\theta}(s)}^k \tilde{S}(y, z) - \Pi_{\hat{\theta}(s')}^k \tilde{S}(y, z)| \leq Mk|s - s'|.$$

Proof of lemma 1.4: $\forall (s, s') \in \mathcal{Q}^2 \quad \forall (y, z) \in \mathcal{Y} \times \mathcal{E} \quad \forall k \in \mathbb{N},$

$$\begin{aligned}
& |\Pi_{\hat{\theta}(s)}^k \tilde{S}(y, z) - \Pi_{\hat{\theta}(s')}^k \tilde{S}(y, z)| \\
& \leq \sum_{i=0}^{k-1} \left| \Pi_{\hat{\theta}(s)}^{k-i} \Pi_{\hat{\theta}(s')}^i \tilde{S}(y, z) - \Pi_{\hat{\theta}(s)}^{k-1-i} \Pi_{\hat{\theta}(s')}^{i+1} \tilde{S}(y, z) \right| \\
& = \sum_{i=0}^{k-1} \left| \Pi_{\hat{\theta}(s)}^{k-1-i} \left(\Pi_{\hat{\theta}(s)} - \Pi_{\hat{\theta}(s')} \right) \Pi_{\hat{\theta}(s')}^i \tilde{S}(y, z) \right| \\
& = \sum_{i=0}^{k-1} \left| \int \int \int \Pi_{\hat{\theta}(s')}^i(z, u) \left(\Pi_{\hat{\theta}(s)}(u, v) - \Pi_{\hat{\theta}(s')}^i(u, v) \right) \Pi_{\hat{\theta}(s)}^{k-1-i}(v, w) \tilde{S}(y, w) du dv dw \right| \\
& \leq \|\tilde{S}\|_{\infty} \sum_{i=0}^{k-1} \int \int \int \Pi_{\hat{\theta}(s')}^i(z, u) \left| \Pi_{\hat{\theta}(s)}(u, v) - \Pi_{\hat{\theta}(s')}^i(u, v) \right| \Pi_{\hat{\theta}(s)}^{k-1-i}(v, w) du dv dw.
\end{aligned}$$

The assumption **(SAEM2)** ensures that the set $\hat{\theta}(\mathcal{Q})$ is compact, so that the second point of assumption **(SAEM3')** ensures the existence of real constants L and $\tilde{L}$ such that:

$$\begin{aligned}
|\Pi_{\hat{\theta}(s)}^k \tilde{S}(y, z) - \Pi_{\hat{\theta}(s')}^k \tilde{S}(y, z)| & \leq \|\tilde{S}\|_{\infty} L |\hat{\theta}(s) - \hat{\theta}(s')| \sum_{i=0}^{k-1} \int \int \int \Pi_{\hat{\theta}(s')}^i(z, u) \Pi_{\hat{\theta}(s)}^{k-1-i}(v, w) du dv dw \\
& \leq \|\tilde{S}\|_{\infty} \tilde{L} |s - s'| \sum_{i=0}^{k-1} \int \int \int \Pi_{\hat{\theta}(s')}^i(z, u) \Pi_{\hat{\theta}(s)}^{k-1-i}(v, w) du dv dw,
\end{aligned}$$

since $\hat{\theta}$ is continuously differentiable. Moreover $\mathcal{E}$ is compact, so we obtain:

$$|\Pi_{\hat{\theta}(s)}^k \tilde{S}(y, z) - \Pi_{\hat{\theta}(s')}^k \tilde{S}(y, z)| \leq \|\tilde{S}\|_{\infty} \tilde{L} M(\mathcal{E}) k |s - s'|,$$

where $M(A)$ denotes the Lebesgue measure of the set A . □

Proof of lemma 1.2:

The first and third points of assumption **(SAEM3')** imply the existence of constants $K \in \mathbb{R}^+$ and $\rho \in]0, 1[$ such that for all (s, s') in $\mathcal{Q}^2$, for all (y, z) in $\mathcal{Y} \times \mathcal{E}$ and for all k in $\mathbb{N}$, we have

$$\begin{aligned}
|p_{\hat{\theta}(s)} \tilde{S} - p_{\hat{\theta}(s')} \tilde{S}| & \leq |p_{\hat{\theta}(s)} \tilde{S} - \Pi_{\hat{\theta}(s)}^k \tilde{S}(y, z)| + |\Pi_{\hat{\theta}(s)}^k \tilde{S}(y, z) - \Pi_{\hat{\theta}(s')}^k \tilde{S}(y, z)| + |\Pi_{\hat{\theta}(s')}^k \tilde{S}(y, z) - p_{\hat{\theta}(s')} \tilde{S}| \\
& \leq 2\|\tilde{S}\|_{\infty} K \rho^k + M k |s - s'|.
\end{aligned}$$

We choose $k \approx \log(|s - s'|) / \log(\rho)$ so that the two terms have approximatively the same weight. Then, there exists a constant K_1 in $\mathbb{R}^+$ such that

$$\forall (s, s') \in \mathcal{Q}^2 \quad \text{s.t.} \quad |s - s'| < 1, \quad |p_{\hat{\theta}(s)} \tilde{S} - p_{\hat{\theta}(s')} \tilde{S}| \leq \frac{K_1}{\log(\rho)} \log(|s - s'|) |s - s'|.$$

Since we have

$$\forall \alpha \in]0, 1[\quad \lim_{|s-s'| \rightarrow 0} \log(|s - s'|) |s - s'|^{1-\alpha} = 0,$$

we can deduce that $\forall \alpha \in]0, 1[, \exists \eta > 0$ s.t.

$$\forall (s, s') \in \mathcal{Q}^2 \quad \text{s.t.} \quad |s - s'| \leq \eta, \quad |p_{\hat{\theta}(s)} \tilde{S} - p_{\hat{\theta}(s')} \tilde{S}| \leq K_1 |s - s'|^\alpha, \quad (1.16)$$

which proves the first inequality of Lemma 1.2 for all pair (s, s') in a compact $\mathcal{Q}$ such that $|s - s'| \leq \eta$.

Concerning the second inequality of Lemma 1.2, let us define

$$a_{k,s,s'}(y, z) = \left| \left(\Pi_{\hat{\theta}(s)}^k \tilde{S}(y, z) - p_{\hat{\theta}(s)} \tilde{S} \right) - \left(\Pi_{\hat{\theta}(s')}^k \tilde{S}(y, z) - p_{\hat{\theta}(s')} \tilde{S} \right) \right|.$$

On one hand, since $\tilde{S}$ is bounded and $\Pi_{\hat{\theta}(s)}$ is uniformly ergodic, there exist $K \in \mathbb{R}^+$ and $\rho \in]0, 1[$ such that, for any $k \in \mathbb{N}$ and any $(s, s', y, z) \in \mathcal{Q}^2 \times \mathcal{Y} \times \mathcal{E}$,

$$\begin{aligned} a_{k,s,s'}(y, z) &\leq |\Pi_{\hat{\theta}(s)}^k \tilde{S}(y, z) - p_{\hat{\theta}(s)} \tilde{S}| + |\Pi_{\hat{\theta}(s')}^k \tilde{S}(y, z) - p_{\hat{\theta}(s')} \tilde{S}| \\ &\leq 2 \|\tilde{S}\|_\infty K \rho^k. \end{aligned} \quad (1.17)$$

On the other hand, for any $\beta \in]0, 1[$, there exists $\eta > 0$ such that, for any $k \in \mathbb{N}$ and any $(s, s', y, z) \in \mathcal{Q}^2 \times \mathcal{Y} \times \mathcal{E}$ such that $|s - s'| \leq \eta$,

$$\begin{aligned} a_{k,s,s'}(y, z) &\leq |\Pi_{\hat{\theta}(s)}^k \tilde{S}(y, z) - \Pi_{\hat{\theta}(s')}^k \tilde{S}(y, z)| + |p_{\hat{\theta}(s)} \tilde{S} - p_{\hat{\theta}(s')} \tilde{S}| \\ &\leq Mk |s - s'| + K_2 |s - s'|^\beta \\ &\leq \max(M, K_2)(k + 1) |s - s'|^\beta. \end{aligned} \quad (1.18)$$

Then we use the following convexity inequality to combine both upper bounds obtained in (1.17) and (1.18):

$$\forall (x_1, x_2) \in \mathbb{R}^{+2} \quad \forall a \in]0, 1[\quad \min(x_1, x_2) \leq x_1^a x_2^{1-a}.$$

So there exist $L_1 \in \mathbb{R}^+$, $L_2 \in \mathbb{R}^+$ and $L \in \mathbb{R}^+$ such that, for any $a \in]0, 1[$, for any $\beta \in]0, 1[$, there exists $\eta > 0$ such that, for any $k \in \mathbb{N}$ and any $(s, s', y, z) \in \mathcal{Q}^2 \times \mathcal{Y} \times \mathcal{E}$ such that $|s - s'| \leq \eta$,

$$\begin{aligned} a_{k,s,s'}(y, z) &\leq \min(L_1 \rho^k, L_2(k + 1) |s - s'|^\beta) \\ &\leq L \rho^{(1-a)k} |s - s'|^{\beta a} (k + 1)^a. \end{aligned}$$

Thus, for any $\beta \in]0, 1[$, for any $a \in]0, 1[$, there exists $\eta > 0$ such that, for any $k \in \mathbb{N}$ and any $(s, s', y, z) \in \mathcal{Q}^2 \times \mathcal{Y} \times \mathcal{E}$ such that $|s - s'| \leq \eta$,

$$\begin{aligned} |\nu_s(z) - \nu_{s'}(z)| &\leq \sum_{k \geq 0} a_{k,s,s'}(y, z) \\ &\leq L \left(\sum_{k \geq 0} \rho^{(1-a)k} (k + 1)^a \right) |s - s'|^{a\beta}. \end{aligned}$$

So we obtain the following conclusion: for any $a \in]0, 1[$, for any $\alpha \in]0, a[$, there exist real constants K_2 and η such that for all z in $\mathcal{E}$

$$\forall (s, s') \in \mathcal{Q}^2 \quad \text{s.t.} \quad |s - s'| < \eta, \quad |\nu_s(z) - \nu_{s'}(z)| \leq K_2 |s - s'|^\alpha. \quad (1.19)$$

Since $\mathcal{Q}$ is compact, we can recover it with a finite number N of balls of diameter η and thus obtain the inequalities (1.16) and (1.19) with the same constants multiplied by N for any pair (s, s') in $\mathcal{Q}^2$. $\square$

Chapitre 2

Maximum likelihood estimation in nonlinear mixed effects models

Summary

A stochastic approximation version of EM for maximum likelihood estimation of a wide class of nonlinear mixed effects models is proposed. The main advantage of this algorithm is its ability to provide an estimator close to the MLE in very few iterations. The likelihood of the observations as well as the Fisher Information matrix can also be estimated by stochastic approximations. Numerical experiments realized in the field of pharmacokinetic and pharmacodynamic allow to highlight the very good performances of the proposed method. The last application to pharmacodynamic consists in the implementation of a likelihood ratio test.

Le travail présenté dans ce chapitre, excepté la dernière section, a donné lieu à la rédaction d'un article co-signé par Estelle Kuhn et Marc Lavielle, actuellement soumis.

2.1 Introduction

Although the interest for mixed effects models was increasing since few years, there were appeared long time ago. For example, Dempster, Laird, and Rubin (1977) already mentioned mixed effects models in their paper dealing with the EM algorithm, but in a different form, reminding some Bayesian point of view, without any interpretation as fixed and random effects: there were two levels of parameters, namely these called parameters, which were supposed to follow some given distribution, depending on some hyperparameters, supposed to be constant. This fact highlights that the notion of mixed effects models is linked with the notion of missing data problems. More recently, the mixed effects models were introduced mainly for modeling responses of individuals that have the same global behavior with individual variations (see the book of Pinheiro and Bates (2000) and the many references therein, for example), implying in particular lots of applications in the study of population behavior like in pharmacokinetic and pharmacodynamic. In fact, we consider that all the responses follow a common known functional form that depends on unknown effects. Some of them are fixed (i.e. the same for all the individuals), the others are random, so they depend on the individuals (or on sub-groups of the population). Then, the model has two types of parameters: global parameters that correspond to the fixed effects, which are also designed as population parameters, and parameters which vary among the population that correspond to the random effects, which are also designed as individual parameters. We can here come back on missing data problems and draw the parallel between its and mixed effects models: the random effects correspond to the missing data (or non observed data) whereas the fixed effects are some parameters of the model. Anyway this kind of observations are usually the result of repeated measurements: for example some individuals are repeatedly observed under different experimental conditions or at different times. Note that it is not necessary that all the individuals have the same number of observations. This approach seems also to be adapted to some other situations besides pharmacokinetic and pharmacodynamic fields, for example in econometry. We will now present more precisely the general framework of mixed effects models.

Let us consider here the following general nonlinear mixed effects model:

$$y_{ij} = C(t_{ij}, \phi_i, \beta) + D(t_{ij}, \phi_i, \beta)\varepsilon_{ij} \quad \text{for } 1 \leq i \leq n, \quad 1 \leq j \leq m_i, \quad (2.1)$$

where y_{ij} is the j th observation of the i th individual, at some known instant t_{ij} . Here, n is the number of individuals and m_i is the number of observations of the individual i . The functions C and D take real values and we say that the model is nonlinear if C or D are nonlinear functions of the random vector ϕ_i , modeled by:

$$\phi_i = A_i \mu + \eta_i \quad \text{with } \eta_i \sim_{i.i.d.} N(0, \Gamma),$$

where μ is an unknown vector of population parameters. The individual matrix A_i is supposed to be known whereas the covariance matrix Γ is supposed to be unknown. The vector β denotes also unknown population parameters, which do not appear in the random effects (ϕ_i). The within-group errors (ε_{ij}) are supposed to be *i.i.d.* Gaussian random variables with mean zero and unknown variance σ^2 . We suppose that the (ε_{ij}) and the (η_i) are mutually independent. We highlight here the fact that in pharmacology, the considered random effects are always centered, namely the mean of there distributions are equal to zero. So may be a rewriting of the model including the possible non zero means in the fixed effects and taking account only of centered

random effects shall be done. Note that the parameterization proposed above for the mixed effects models allows more parameters model as just the means and the variances of the random effects and of the errors.

Our purpose is to compute the maximum likelihood estimator of the unknown parameter vector of the model $\theta = (\beta, \mu, \Gamma, \sigma^2)$ from the observations. In the case of a linear model, the estimation of the unknown parameters can be treated with the usual EM algorithm (see Dempster, Laird, and Rubin (1977)), or with a Newton-Raphson algorithm (see Lindstrom and Bates (1988)). But a nonlinear function is often more suitable for modeling faithfully the physical beginning problem. However this requires a specific approach for estimating the parameters because the methods mentioned above failed, involving usually the computation of some integral which becomes intractable in nonlinear models. So different methods were suggested for dealing with nonlinear models. Bayesian approaches were proposed by Racine-Poon (1985), Wakefield et al. (1994) and Wakefield (1996). The most usual were based on linearization of the model. The common methods used by the professional in pharmacology, implemented in the soft-ware NONMEM, are first-order method (FO) and first-order conditional estimates (FOCE) (see Lindstrom and Bates (1990) and section 2.5 for more details), based on a first order Taylor expansion of the response about the fixed effects. Vonesh (1996) propose a Laplace approximation, but its estimator is consistant under restrictive assumptions: he supposes that the numbers of observations per individual tend to infinity. There exist also other two-step iterative algorithms like the algorithm P-Pharm of Mentré and Gomeni (1995) (see section 2.4.3 for more details). Using simulation, Walker (1996) uses a Monte-Carlo EM algorithm, whereas a simulated pseudo maximum likelihood estimator for these specific models is developed by Concordet and Nunez (2002). Their estimator is consistant when the number of observations tends to infinity.

In this chapter, we show that the SAEM algorithm (Stochastic Approximation version of EM) is very efficient for computing the Maximum Likelihood Estimate of θ in mixed effects models in general, but particularly in nonlinear models. This iterative procedure consists at each iteration, in successively simulating the random effects with the conditional distribution, and updating the unknown parameters of the model. This algorithm was shown to converge under very general conditions by Delyon, Lavielle, and Moulines (1999). If it is not possible to simulate the random effects from the conditional distribution, this algorithm can be coupled with a MCMC procedure for the simulation step, allowing to treat more generally models. In this case, the convergence of the algorithm toward the Maximum Likelihood Estimate is also established in the first chapter of this thesis. Furthermore, the observed likelihood and the Fisher Information matrix can also be estimated, by using also a stochastic approximation procedure. This method has the very nice advantage to converge very quickly to a neighborhood of the Maximum Likelihood Estimate. Then, only a few seconds are required for computing a MLE confidence interval, in any of the usual models used in the practice. The SAEM algorithm can be used for estimating homoscedastic models as well as heteroscedastic models. For the latter, the parameters related to the fixed effects may be estimated in a Bayesian context in term of their expectations.

The second section presents succinctly the EM and the SAEM algorithms. We briefly recall the main convergence results, which are detailed in the introduction of this thesis. We illustrate the proposed method in Section 2.3 with the very simple example of orange data, reported in Pinheiro and Bates (2000). The model used for this data is linear with respect to the random effects ϕ_i , so that the exact MLE can be computed with EM, and compared to the results provided by the SAEM algorithm. Section 2.4 is dedicated to the comparison of the SAEM

algorithm performances with these of some other popular methods of estimation through three numerical examples from the pharmacological field. Section 2.5 presents an application to a likelihood ratio test on a pharmacodynamic model.

2.2 Algorithms proposed for maximum likelihood estimation

Since the EM algorithm is devoted to missing data problems, we begin by specifying the remark mentioned above, namely that any mixed effects model can be seen as an usual missing data problem. Indeed, the observed data are the (y_{ij}) , for $1 \leq i \leq n$ and $1 \leq j \leq m_i$, whereas the random effects (ϕ_i) , for $1 \leq i \leq n$, are the non observed data. Then, the complete data of the model are $(y_{ij}, \phi_i)_{1 \leq i \leq n, 1 \leq j \leq m_i}$. In the sequel, we will make the following assumption concerning the model:

- **(M0)** For any $1 \leq i \leq n$ and any $1 \leq j \leq m_i$,

$$y_{ij} = C(t_{ij}, \phi_i, \beta) + D(t_{ij}, \phi_i, \beta)\varepsilon_{ij}, \quad (2.2)$$

where C and D can be decomposed as:

$$C(t_{ij}, \phi_i, \beta) = C_1(t_{ij}, \phi_i)C_2(t_{ij}, \beta) \quad (2.3)$$

$$D(t_{ij}, \phi_i, \beta) = D_1(t_{ij}, \phi_i)D_2(t_{ij}, \beta). \quad (2.4)$$

Assumption **(M0)** is equivalent to suppose that the complete data likelihood $f(\mathbf{y}, \phi; \theta)$ belongs to the curved exponential family and can be written:

$$f(\mathbf{y}, \phi; \theta) = \exp \left\{ -\Psi(\theta) + \left\langle \tilde{S}(\mathbf{y}, \phi), \Phi(\theta) \right\rangle \right\}, \quad (2.5)$$

where $\langle \cdot, \cdot \rangle$ denotes the scalar product on $\mathbb{R}^m$, Ψ and Φ two functions of θ taking values respectively in $\mathbb{R}$ and $\mathbb{R}^m$ and $\tilde{S}(\mathbf{y}, \phi)$ a vector of $\mathbb{R}^m$ known as the minimal sufficient statistics of the complete model.

Remark. If the model admits no parameter β , then the condition **(M0)** is always satisfied. This is for example the case in the most common models used in pharmacology, where the parameters of the model consist only in the means and in the variances of the Gaussian random effects.

2.2.1 The EM algorithm

We recall here briefly the iterative procedure of the EM algorithm. Assuming the exponential form (2.5) of the complete likelihood, the k th iteration reduces to the following two steps:

- **E-step:** evaluate the quantity $\mathbf{s}_{k+1} = E \left[\tilde{S}(\mathbf{y}, \phi) | \mathbf{y}; \theta_k \right]$.
- **M-step:** compute $\theta_{k+1} = \arg \max \{ -\Psi(\theta) + \langle \mathbf{s}_{k+1}, \Phi(\theta) \rangle \}$.

If the function D does not depend on the random effects ϕ , the complete likelihood f has the form

$$f(\mathbf{y}, \phi; \theta) = \mathcal{C} |\Gamma|^{-\frac{n}{2}} \prod_{i,j} (\sigma^2 D^2(\beta, t_{ij}))^{-\frac{1}{2}} \\ \times \exp \left\{ -\frac{1}{2} \sum_i (\phi_i - \mathbf{A}_i \mu)^t \Gamma^{-1} (\phi_i - \mathbf{A}_i \mu) - \frac{1}{2\sigma^2} \sum_{i,j} \left(\frac{y_{ij} - C(\beta, \phi_i, t_{ij})}{D(\beta, t_{ij})} \right)^2 \right\}, \quad (2.6)$$

where $\mathcal{C}$ is a normalizing constant. If the function C depends linearly on ϕ , this joint distribution is Gaussian, and the E-step can be performed in a close form. Following Dempster et al. (1977), Wu (1983) and Delyon, Lavielle, and Moulines (1999), convergence of the EM algorithm requires the following regularity assumptions on the model corresponding to the assumptions (M1)-(M5) presented in the introduction:

- (H1) The parameter vector θ belongs to a parameter space Θ , which is an open subset of $\mathbb{R}^p$.
- (H2) The functions C and D are twice continuously differentiable with respect to β .
- (H3) There exists a unique minimum for the function Λ defined by:

$$\Lambda(\beta) = \frac{1}{2\sigma^2} \sum_{ij} \left(\frac{y_{ij} - C(\beta, \phi_i, t_{ij})}{D(\beta, t_{ij})} \right)^2 + \sum_{ij} \log |D(\beta, t_{ij})|.$$

We denote this minimum $\hat{\beta}(\tilde{S}(\mathbf{y}, \phi))$ and suppose also that the function $\hat{\beta}$ is continuously differentiable on $\mathcal{S}$, which is the space where the function $\tilde{S}$ takes its values.

So we obtain the following theorem in the particular case of mixed effects models:

Theorem 2.1. *Assume that assumptions (H1)-(H3) hold. Then, the sequence (θ_k) obtained from the EM algorithm converges with probability 1 to a stationary point of the log-likelihood l of the observations $\mathbf{y}$ (i.e a point where the derivative of l is 0).*

In cases where the function D depends on the random effects or where the function C does not depend linearly on the random effects, we propose an alternative which is a stochastic version of the EM algorithm.

2.2.2 A stochastic version of the EM algorithm

Description and general convergence result of the algorithm

The stochastic approximation version of the standard EM algorithm, proposed by Delyon et al. (1999) consists in replacing the usual E-step of EM by a stochastic procedure composed of two steps: first a simulation step of the missing data under the conditional distribution, second a stochastic approximation step:

- **S-step:** draw $\phi^{(k+1)}$ from the conditional distribution $p(\cdot|\mathbf{y}; \theta_k)$.
- **A-step:** update $\mathbf{s}_k$ according to

$$\mathbf{s}_{k+1} = \mathbf{s}_k + \gamma_k \left(\tilde{S}(\mathbf{y}, \phi^{(k+1)}) - \mathbf{s}_k \right). \quad (2.7)$$

- **M-step:** update θ_k according to

$$\theta_{k+1} = \arg \max \{ -\Psi(\theta) + \langle \mathbf{s}_{k+1}, \Phi(\theta) \rangle \}.$$

The assumptions they need and their convergence result are presented in the paper cited above and are also recalled in the introduction. When the simulation step cannot be directly performed, we propose to combine this algorithm with a MCMC procedure: the sequence $(\phi^{(k)})$ becomes a Markov Chain with transition kernels (Π_{θ_k}) . So the simulation step involves a transition probability $\Pi_{\theta_k}(\cdot, \cdot)$ from which the missing data $\phi^{(k+1)}$ were simulated using the previous value $\phi^{(k)}$ (see chapter 1). We have shown that SAEM converges to a maximum (local or global) of the log-likelihood of the observations l , under very general conditions. We recall here briefly the theorem presented in the first chapter of this thesis.

- **(SAEM1')** For all k in $\mathbb{N}$, $\gamma_k \in [0, 1]$, $\sum_{k=1}^{\infty} \gamma_k = \infty$ and there exists λ in $]\frac{1}{2}, 1]$ such that $\sum_{k=1}^{\infty} \gamma_k^{1+\lambda} < \infty$.
- **(SAEM2)** The functions $l : \Theta \rightarrow \mathbb{R}$ and $\hat{\theta} : \mathcal{S} \rightarrow \Theta$ are m times differentiable, where m is the integer such that $\mathcal{S}$ is an open subset of $\mathbb{R}^m$.
- **(SAEM3')**
 - The chain $(z_k)_{k \geq 0}$ takes its values in a compact subset $\mathcal{E}$ of $\mathbb{R}^l$.
 - For any compact subset V of Θ , there exists a real constant L such that for any (θ, θ') in V^2

$$\sup_{(x,y) \in \mathcal{E}^2} |\Pi_{\theta}(x, y) - \Pi_{\theta'}(x, y)| \leq L|\theta - \theta'|.$$

- The transition probability Π_{θ} generates a uniformly ergodic chain which invariant probability is the conditional distribution $p(\cdot | y; \theta)$:

$$\exists K_{\theta} \in \mathbb{R}^+ \quad \exists \rho_{\theta} \in]0, 1[\quad | \quad \forall z \in \mathcal{E} \quad \forall k \in \mathbb{N} \quad \|\Pi_{\theta}^k(z, \cdot) - p(\cdot | y; \theta)\|_{TV} \leq K_{\theta} \rho_{\theta}^k,$$

where $\|\cdot\|_{TV}$ denotes the total variation norm. We suppose also that:

$$K \triangleq \sup_{\theta} K_{\theta} < +\infty \quad \text{and} \quad \rho \triangleq \sup_{\theta} \rho_{\theta} < 1.$$

- The function $\tilde{S}$ is bounded on $\mathcal{E}$.

Theorem 2.2. Assume that assumptions **(M0)**, **(H1)**-**(H3)**, **(SAEM1')**, **(SAEM2)**, **(SAEM3')** and **(LOC1)**-**(LOC3)** hold. Then, the sequence (θ_k) obtained from the SAEM algorithm converges with probability 1 to a maximum (local or global) of the log-likelihood l of the observations.

Remark. One of the technical conditions required to prove the convergence of SAEM is the compactness of the support of $p(\cdot | y; \theta)$. In the case of the nonlinear mixed effects models we have chosen to consider here, the non observed data ϕ follow a Gaussian distribution and theoretically take their values over an infinite set which is not compact. From a practical point of view, this assumption is not a restriction, since in the practice, any Gaussian random variable takes values in a (very large) compact.

Remark. Consider the homoscedastic model where the function D defined in (2.2) is constant and equal to 1. In this particular case, the assumptions are quite simpler. Indeed, convergence of SAEM is ensured if the function $\Lambda_1(\beta) = \|\mathbf{y} - C(\beta, \phi, x)\|^2$ possesses a unique minimum.

In the case of an heteroscedastic model such that $C = D$, convergence of SAEM requires that the following function possesses a unique minimum:

$$\Lambda_2(\beta) = \frac{1}{2\sigma^2} \sum_{i,j} \left(\frac{y_{ij}}{C(\beta, \phi_i, t_{ij})} - 1 \right)^2 + \sum_{ij} \log |C_1(\beta, t_{ij})|.$$

Remark. The SAEM is useful for fitting models that belong to the exponential family. If this assumption is not satisfied, it is always possible to consider the vector of fixed parameters β as a random vector. Then, β is estimated in a Bayesian context in term of its expectation. An example of such extension is proposed section 2.3.3.

Simulation of the missing data

When $\phi^{(k+1)}$ cannot be exactly drawn from the conditional distribution $p(\cdot|\mathbf{y}; \boldsymbol{\theta}_k)$, we have to choose judiciously a transition $\Pi_{\boldsymbol{\theta}}$ that converges to the target distribution $p(\cdot|\mathbf{y}; \boldsymbol{\theta})$. The Metropolis-Hastings algorithm provides a solution in general cases. Usually, $\Pi_{\boldsymbol{\theta}}$ will be defined as the succession of M iterations of a MCMC procedure and the simulation-step of iteration k consists in simulating $\phi^{(k+1)}$ with the transition probability $\Pi_{\boldsymbol{\theta}_k}(\phi^{(k)}, d\phi^{(k+1)}) = P_{\boldsymbol{\theta}_k}^M(\phi^{(k)}, d\phi^{(k+1)})$, where

$$P_{\boldsymbol{\theta}_k}(\phi, d\phi') = r_{\boldsymbol{\theta}_k}(\phi, \phi') \min \left\{ \frac{p(\phi'|\mathbf{y}; \boldsymbol{\theta}_k) r_{\boldsymbol{\theta}_k}(\phi', \phi)}{p(\phi|\mathbf{y}; \boldsymbol{\theta}_k) r_{\boldsymbol{\theta}_k}(\phi, \phi')}, 1 \right\} d\phi', \quad (2.8)$$

for $\phi' \neq \phi$ and $P_{\boldsymbol{\theta}_k}(\phi, \{\phi\}) = 1 - \int_{\phi' \neq \phi} P_{\boldsymbol{\theta}_k}(\phi, d\phi')$, where $r_{\boldsymbol{\theta}}(\phi, \phi')$ is any aperiodic recurrent transition density.

We set $M = 1$ in the sequel, in order to avoid intricate notations. Extension to any value of M is straightforward. We can use the population distribution π as an instrumental distribution. Then, writing $f(\mathbf{y}, \phi; \boldsymbol{\theta}) = h(\mathbf{y}|\phi; \boldsymbol{\theta})\pi(\phi; \boldsymbol{\theta})$, the acceptance probability only depends on the conditional distribution h of the observation $\mathbf{y}$:

$$P_{\boldsymbol{\theta}_k}(\phi, d\phi') = \pi(\phi'; \boldsymbol{\theta}_k) \min \left\{ \frac{h(\mathbf{y}|\phi'; \boldsymbol{\theta}_k)}{h(\mathbf{y}|\phi; \boldsymbol{\theta}_k)}, 1 \right\} d\phi'.$$

In the case of the nonlinear mixed effects models, the k th step of this MCMC procedure consists in:

1. Draw $\phi' = (\phi'_1, \dots, \phi'_n)'$ i.i.d. with the prior distribution $\mathcal{N}(\boldsymbol{\mu}_k, \boldsymbol{\Gamma}_k)$ and $\mathbf{u} = (u_1, \dots, u_n)$ i.i.d. with the uniform distribution $\mathcal{U}([0, 1])$.
2. For $i = 1 \dots n$, compute

$$\Delta_i = \sum_{j=1}^{m_i} \left[\log \left(\frac{D_2(\phi'_i, t_{ij})}{D_2(\phi_i^{(k)}, t_{ij})} \right) + \frac{1}{2\sigma^2} \left(\frac{y_{ij} - C(\boldsymbol{\beta}_k, \phi'_i, t_{ij})}{D(\boldsymbol{\beta}_k, \phi'_i, t_{ij})} \right)^2 - \frac{1}{2\sigma^2} \left(\frac{y_{ij} - C(\boldsymbol{\beta}_k, \phi_i^{(k)}, t_{ij})}{D(\boldsymbol{\beta}_k, \phi_i^{(k)}, t_{ij})} \right)^2 \right].$$

3. For $i = 1 \dots n$, set

$$\begin{aligned} \phi_i^{(k+1)} &= \phi'_i \quad \text{if } \Delta_i \leq \log(u_i) \\ \phi_i^{(k+1)} &= \phi_i^{(k)} \quad \text{elsewhere.} \end{aligned}$$

Another way to simulate the Markov chain $(\phi^{(k)})$ is to generate the new possible value of $\phi^{(k+1)}$ by adding to the previous value $\phi^{(k)}$ a random perturbation. The acceptance probability shall then be evaluated from the formula (2.8). Other methods are detailed in practice in section 2.4.1.

Some improvements of the algorithm

Some parameters of the algorithm may be well chosen to improve the results, like the general number of iterations, the number of iterations in the MCMC procedure for generating each element of the Markov chain $(\phi^{(k)})$, the step size $(\gamma_k)_{k \geq 0}$, possibly the number of Markov chain generated at the same time. More these choices are adapted at each case, better the results are. For example, we have to choose the sequence of step size $(\gamma_k)_{k \geq 0}$ such that the assumption

(SAEM1') is satisfied: each γ_k must belong to $[0, 1]$, the series $\sum \gamma_k$ must diverge, and the series $\sum \gamma_k^{1+\lambda}$ must converge for some λ in $]\frac{1}{2}, 1]$. We must also take into account the influence of the step size on the convergence speed of the algorithm. In the practice, it is useful to choose the first step size equal to 1, in order to allow more flexibility during the first iterations: in fact, the initial guess θ_0 may be far from the maximum likelihood value we are looking for and the first iterations will require big variations of the sequence (θ_k) . After converging to a neighborhood of the MLE, it is interesting to choose smaller step size in order to refine the estimation near the objective value and ensure the almost sure convergence of the algorithm. Finally, we recommend to set $\gamma_k = 1$ for $1 \leq k \leq K$ and $\gamma_k = (k - K)^{-1}$ for $k \geq K + 1$. In the practice K can be chosen between 50 and 100.

It is possible to reduce the variations of the sequence (θ_k) around the MLE, by averaging the sequence (s_k) after iteration K as follows: *i*) run SAEM with $\gamma_k = (k - K)^{-\alpha}$ for $k \geq K + 1$, where $0.5 < \alpha < 1$ *ii*) compute $\bar{s}$ as an average of the sequence (s_k) , $k \geq K + 1$ and compute $\hat{\theta}(\bar{s})$. Some theoretical results concerning averaging of SAEM are given by Delyon et al. (1999). From a practical point of view, the improvement is slight. We think that it is not really useful for computing an estimate that possesses a variance much bigger than the improvement we can expect.

Concerning the MCMC procedure, we have remarked that it improves considerably the results if we choose a big value for M during the first iterations of the algorithm, in order to adjust rapidly the chain of good values for the missing data. We usually use $M = 100$ or $M = 200$ for the iterations 1 to 20 or 50, and M between 10 and 100 after.

The MCEM (Monte-Carlo EM) was proposed by Wei and Tanner (1990), and used by Walker (1996) for nonlinear mixed effects models (see section 2.4.2 for some applied results). This algorithm approximates the conditional expectation $Q_k(\theta) = E(\log f(y, \phi; \theta) | y; \theta_k)$ in the *E-step* thanks to a Monte Carlo method. It requires to draw a sequence $(\phi^{(k+1, \ell)}), 1 \leq \ell \leq L$ at iteration k and to compute s_{k+1} as an average of $(\tilde{S}(y, \phi^{(k+1, \ell)}))_{1 \leq \ell \leq L}$. A good approximation requires a large number L of simulations at each iteration. We can combine MCEM with SAEM by adapting the approximation step (2.7) as follows:

$$s_{k+1} = s_k + \gamma_k \left(\frac{1}{L} \sum_{\ell=1}^L \tilde{S}(y, \phi^{(k+1, \ell)}) - s_k \right). \quad (2.9)$$

With this new version of the algorithm, a small value of L (smaller than 10 in the practice) is enough to ensure very satisfactory results.

Estimation of the likelihood

The likelihood of the observations can be approximated via a Monte Carlo integration method (see Walker (1996) for example), using the following importance sampling identity

$$q(y; \theta) = \int h(y | \phi; \theta) \pi(\phi; \theta) d\phi \simeq \frac{1}{T} \sum_{t=1}^T h(y | \phi^{(t)}; \theta), \quad (2.10)$$

where the $\phi^{(t)}$ are drawn independently with the prior distribution $\pi(\cdot; \theta)$.

It can be useful if we want to study the behavior of the likelihood at each iteration of the EM or SAEM algorithm. At iteration k of the algorithm, a sequence $\phi^{(k1)}, \phi^{(k2)}, \dots, \phi^{(kT)}$ is drawn with the prior $\pi(\cdot; \theta_k)$ and $q(y; \theta_k)$ is estimated using (2.10).

If the $(\phi^{(k,t)})$ are drawn independently at each iteration, the estimated likelihood sequence $(q(\mathbf{y}; \boldsymbol{\theta}_k))$ will be quite perturbed if T is not chosen large enough. A smooth curve can be obtained by using the same random numbers at each iteration: for example, if $\phi \sim \mathcal{N}(\boldsymbol{\mu}, \boldsymbol{\tau})$, then we set $\phi^{(k,t)} = \boldsymbol{\mu}_k + \boldsymbol{\tau}_k Z^{(t)}$ for any k , and where the $Z^{(t)}$ are independent $\mathcal{N}(0, 1)$.

Another estimator of the likelihood is obtained by stochastic approximation, setting

$$q_k(\mathbf{y}; \boldsymbol{\theta}_k) = q_{k-1}(\mathbf{y}; \boldsymbol{\theta}_{k-1}) + \gamma_k \left(\frac{1}{T} \sum_{t=1}^T h(\mathbf{y} | \phi^{(k,t)}; \boldsymbol{\theta}_k) - q_{k-1}(\mathbf{y}; \boldsymbol{\theta}_{k-1}) \right), \quad (2.11)$$

where the $\phi^{(k,t)}$ are independent. Indeed, when the sequence $(\boldsymbol{\theta}_k)$ converges almost surely to $\boldsymbol{\theta}^*$, this sequence $(q_k(\mathbf{y}; \boldsymbol{\theta}_k))$ converges almost surely to $q(\mathbf{y}; \boldsymbol{\theta}^*)$.

Estimation of the variance of the estimates

Thanks to the maximum likelihood estimator obtained with the proposed algorithm, it is possible to obtain simultaneously an estimation of the Fisher Information matrix. Delyon et al. (1999) propose a method to estimate this matrix by using the fact that the gradient (the Fisher score function) and the Hessian (observed Fisher Information) of l can be obtained almost directly from the simulated missing data ϕ . Using the so-called Fisher identity, the Jacobian of the log-likelihood of the observed data l is equal to the conditional expectation of the complete data likelihood:

$$\partial_{\theta} l(\boldsymbol{\theta}) \triangleq E_{\theta}[\partial_{\theta} \log f(\mathbf{y}, \phi; \boldsymbol{\theta}) | \mathbf{y}; \boldsymbol{\theta}],$$

where ∂_{θ} denotes the differential with respect to $\boldsymbol{\theta}$. By analogy with the implementation of the SAEM algorithm, the following approximation scheme is proposed:

$$\Delta_k = \Delta_{k-1} + \gamma_k \left[\partial_{\theta} \log f(\mathbf{y}, \phi^{(k)}; \boldsymbol{\theta}_k) - \Delta_{k-1} \right].$$

Using the Louis's missing information principle (Louis (1982)), the Hessian of l at $\boldsymbol{\theta}$, is the observed Fisher information matrix $\partial_{\theta}^2 l(\boldsymbol{\theta})$ that may be expressed as

$$\partial_{\theta}^2 l(\boldsymbol{\theta}) = E_{\theta}[\partial_{\theta}^2 \log f(\mathbf{y}, \phi; \boldsymbol{\theta})] + \text{Cov}_{\theta}[\partial_{\theta} \log f(\mathbf{y}, \phi; \boldsymbol{\theta})].$$

where $\text{Cov}_{\theta}(\psi(\phi)) \triangleq E_{\theta}[(\psi(\phi) - E_{\theta}(\psi(\phi)))(\psi(\phi) - E_{\theta}(\psi(\phi)))^t]$. Using this expression, it is possible to derive the following stochastic approximation procedure to approximate $\partial_{\theta}^2 l(\boldsymbol{\theta})$:

$$\begin{aligned} \mathbf{G}_k &= \mathbf{G}_{k-1} + \gamma_k \left(\partial_{\theta}^2 \log f(\mathbf{y}, \phi^{(k)}; \boldsymbol{\theta}_k) + \partial_{\theta} \log f(\mathbf{y}, \phi^{(k)}; \boldsymbol{\theta}_k) \partial_{\theta} \log f(\mathbf{y}, \phi^{(k)}; \boldsymbol{\theta}_k)^t - \mathbf{G}_{k-1} \right) \\ \mathbf{H}_k &= \mathbf{G}_k - \Delta_k \Delta_k^t. \end{aligned}$$

Knowing that the algorithm proposed above converges to a limiting value $\boldsymbol{\theta}^*$ and that l is regular enough, $(-\mathbf{H}_k)$ converges to $-\partial_{\theta}^2 l(\boldsymbol{\theta}^*)$. When l is an incomplete data log-likelihood function sufficiently smooth, the maximum likelihood estimator is asymptotically normal and the inverse of the observed Fisher information matrix $-\partial_{\theta}^2 l(\boldsymbol{\theta}^*)$ converges to the asymptotic covariance of the estimators.

2.3 Example of the orange trees

2.3.1 The model

We choose the example of the orange trees to illustrate our algorithm for two reasons: the first one is that the model chosen for this study is a mixed effects model with only one random effect, moreover linear in this random effect, so we are able to implement the EM algorithm at the same time as the SAEM algorithm, in order to get some reliable estimations of the maximum likelihood with which we can compare these provided by the SAEM algorithm. Since the model is relatively simple, the implementation involves only very soft calculation. The second one is that these data, available for example on S-plus, were studied by Pinheiro and Bates (2000) and other authors, so allow us some comparisons with results obtained by others methods, in order to assess the performance of the SAEM algorithm.

The data, shown in Figure 2.1, consist in seven measurements of the trunk circumference of each of five orange trees. Inspection of Figure 2.1 suggests that a “tree effect” is present. So a mixed effects model seems to be very adapted to the analyze of these data.

Following Pinheiro and Bates (2000), a logistic curve is used for modeling the trunk circumference y_{ij} of tree i at age t_j :

$$y_{ij} = C(t_j, \phi_i; \beta_1, \beta_2) + \varepsilon_{ij} \text{ for } 1 \leq i \leq n, 1 \leq j \leq m, \quad (2.12)$$

$$C(t_j, \phi_i; \beta_1, \beta_2) = \frac{\phi_i}{1 + \exp\left(-\frac{t_j - \beta_1}{\beta_2}\right)}. \quad (2.13)$$

We suppose here that the (ε_{ij}) are independent Gaussian error terms with variance σ^2 . On the one hand, the age at which the tree attains half of its asymptotic trunk circumference β_1 , and the grow scale β_2 are treated as two fixed effects. On the other hand, the asymptotic trunk circumference ϕ_i is treated as a random effect and is assumed to be Gaussian with mean μ and variance τ^2 . We suppose also that the random effects (ϕ_i) and the error terms (ε_{ij}) are mutually independent.

Setting

$$\alpha_j(\beta_1, \beta_2) = \frac{1}{1 + \exp\left(-\frac{t_j - \beta_1}{\beta_2}\right)}, \quad (2.14)$$

the likelihood of the complete model as the form:

$$f(\mathbf{y}, \boldsymbol{\phi}; \boldsymbol{\theta}) = (2\pi\sigma^2)^{-\frac{nm}{2}} (2\pi\tau^2)^{-\frac{m}{2}} \times \exp\left[-\frac{1}{2\sigma^2} \sum_{i,j} (y_{ij} - \alpha_j(\beta_1, \beta_2)\phi_i)^2 - \frac{1}{2\tau^2} \sum_i (\phi_i - \mu)^2\right], \quad (2.15)$$

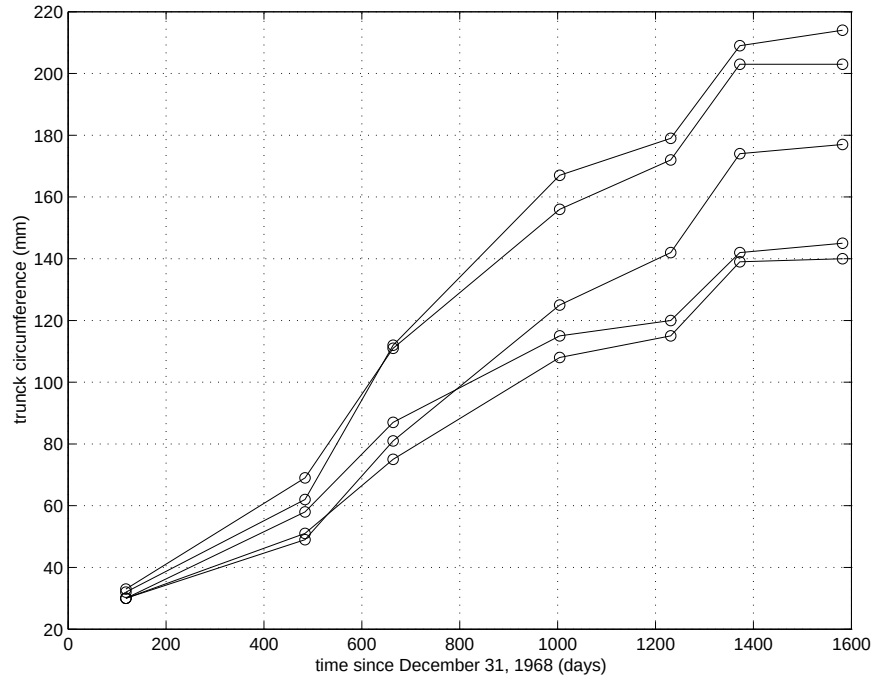
where $\boldsymbol{\theta} = (\beta_1, \beta_2, \mu, \tau^2, \sigma^2)$ are all the parameters of the model which we will estimate from the observations (y_{ij}) .

2.3.2 The EM and SAEM algorithms

The non observed data is the sequence of independent random effects $(\phi_i)_{1 \leq i \leq n}$. From (2.15), we deduce by identification that the conditional distribution $p(\phi_i | \mathbf{y}; \boldsymbol{\theta})$ is Gaussian:

$$\phi_i | \mathbf{y}; \boldsymbol{\theta} \sim \mathcal{N}(u_i, V),$$

FIG. 2.1: Circumference of the five orange trees.



where

$$u_i = V \left(\frac{1}{\sigma^2} \sum_{j=1}^n y_{ij} \alpha_j(\beta_1, \beta_2) + \frac{\mu}{\tau^2} \right) \quad \text{and} \quad V = \left(\frac{1}{\sigma^2} \sum_{j=1}^n \alpha_j^2(\beta_1, \beta_2) + \frac{1}{\tau^2} \right)^{-1}. \quad (2.16)$$

Thus, the Expectation-step of the EM algorithm as well as the Simulation-step of the SAEM algorithm can be easily done.

By the way, equation (2.15) implies that we can choose as minimal sufficient statistics the function $\tilde{S}$ equal to:

$$\tilde{S}(\mathbf{y}, \phi) = \left(\tilde{S}_{\mathbf{y}^2}, \tilde{S}_{\phi}, \tilde{S}_{\phi^2}, (\tilde{S}_{\mathbf{y}\phi}(j))_{1 \leq j \leq m} \right), \quad (2.17)$$

where

$$\tilde{S}_{\mathbf{y}^2} = \sum_{i=1}^n \sum_{j=1}^m y_{ij}^2, \quad \tilde{S}_{\phi} = \sum_{i=1}^n \phi_i, \quad \tilde{S}_{\phi^2} = \sum_{i=1}^n \phi_i^2, \quad \tilde{S}_{\mathbf{y}\phi}(j) = \sum_{i=1}^n y_{ij} \phi_i. \quad (2.18)$$

TAB. 2.1: Comparison of EM and SAEM estimates after 100 and 1000 iterations.

Parameters	β_1	β_2	μ	τ^2	σ^2
θ_0	650	250	100	50	10
θ_{100}^{EM}	727.89	348.06	192.05	1001.45	61.51
θ_{1000}^{EM}	727.91	348.07	192.05	1001.49	61.51
θ_{100}^{SAEM}	725.34	346.11	191.46	1020.26	60.56
θ_{1000}^{SAEM}	727.36	347.67	191.93	1003.76	61.52

Then the maximum likelihood estimate of θ can be expressed as a function of $\tilde{S}(\mathbf{y}, \phi)$:

$$(\hat{\beta}_1, \hat{\beta}_2) = \text{Arg min}_{\beta_1, \beta_2} \left\{ -2 \sum_{j=1}^m \alpha_j(\beta_1, \beta_2) \tilde{S}_{\mathbf{y}\phi}(j) + \tilde{S}_{\phi^2} \sum_{j=1}^m \alpha_j^2(\beta_1, \beta_2) \right\} \quad (2.19)$$

$$\hat{\mu} = \frac{\tilde{S}_{\phi}}{n} \quad (2.20)$$

$$\hat{\tau}^2 = \frac{\tilde{S}_{\phi^2}}{n} - \left(\frac{\tilde{S}_{\phi}}{n} \right)^2 \quad (2.21)$$

$$\hat{\sigma}^2 = \frac{1}{nm} \left(\tilde{S}_{\mathbf{y}^2} - 2 \sum_{j=1}^m \alpha_j(\hat{\beta}_1, \hat{\beta}_2) \tilde{S}_{\mathbf{y}\phi}(j) + \tilde{S}_{\phi^2} \sum_{j=1}^m \alpha_j^2(\hat{\beta}_1, \hat{\beta}_2) \right). \quad (2.22)$$

The M-step of both algorithms requires the use of a Newton-Raphson algorithm for computing $(\hat{\beta}_1, \hat{\beta}_2)$ since the direct maximization fails.

The sequence of estimates (θ_k^{EM}) and (θ_k^{SAEM}) are displayed Figure 2.2 and Figure 2.3, together with the log-likelihood sequences $\log q(\mathbf{y}, \theta_k^{EM})$ and $\log q(\mathbf{y}, \theta_k^{SAEM})$ that can easily be computed in a close form:

$$q(\mathbf{y}; \theta) = (2\pi\sigma^2)^{-\frac{nm}{2}} \left(\frac{V}{\tau^2} \right)^{\frac{n}{2}} \exp \left(-\frac{1}{2\sigma^2} \sum_{ij} y_{ij}^2 - \frac{m\mu^2}{2\tau^2} + \frac{V}{2} \sum_i u_i^2 \right), \quad (2.23)$$

by integration of (2.15) since $q(\mathbf{y}; \theta) = \int f(\mathbf{y}, \phi; \theta) d\phi$.

The estimation of the parameters after 100 and 1000 iterations with EM and SAEM are displayed Table 2.1.

We can remark that EM has almost converged after 100 iterations. Thus, the value obtained with this algorithm can be considered as the maximum likelihood estimate of θ . In this example, the step size sequence (γ_k) used for SAEM was: $\gamma_k = 1$ for $1 \leq k \leq 100$ and $\gamma_k = (k - 99)^{-1}$ for $k \geq 100$. After some iterations, the SAEM algorithm has converged to a neighborhood of the MLE of θ . Since $\gamma_k = 1$ during the first iterations, no stochastic approximation are still performed, and $\mathbf{s}_k = \tilde{S}(\mathbf{y}, \phi^{(k)})$. Thus, the behavior of θ_k^{SAEM} remains quite perturbed until iteration 100. After that, the introduction of a decreasing step size allows the almost sure convergence of the sequence θ_k^{SAEM} to $\hat{\theta}^{MLE}$.

The estimation of the observed log-likelihood is displayed Figure 2.4. In this very simple example, the estimator, proposed in Section 2.2.2 can be compared to the true log-likelihood,

FIG. 2.2: Estimation of θ using EM. A logarithmic scale is used for the x-axis. (a) (β_{1k}) , (b) (β_{2k}) , (c) (μ_k) , (d) (τ_k^2) , (e) (σ_k^2) , (f) $\log q(\mathbf{y}, \theta_k^{EM})$.

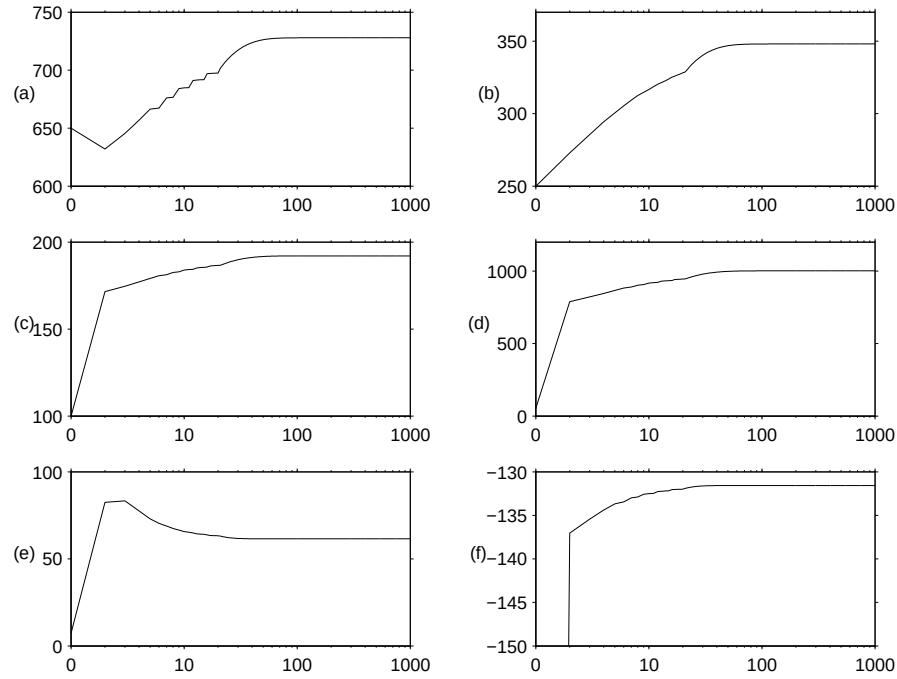


FIG. 2.3: Estimation of θ using SAEM. A logarithmic scale is used for the x-axis. (a) (β_{1k}) , (b) (β_{2k}) , (c) (μ_k) , (d) (τ_k^2) , (e) (σ_k^2) , (f) $\log q(\mathbf{y}, \theta_k^{SAEM})$.

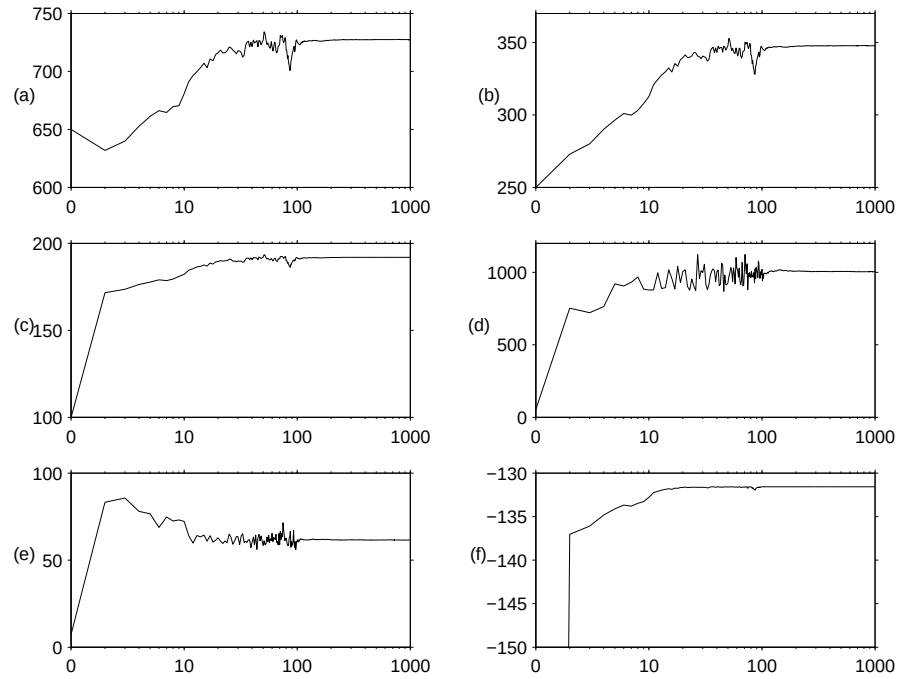
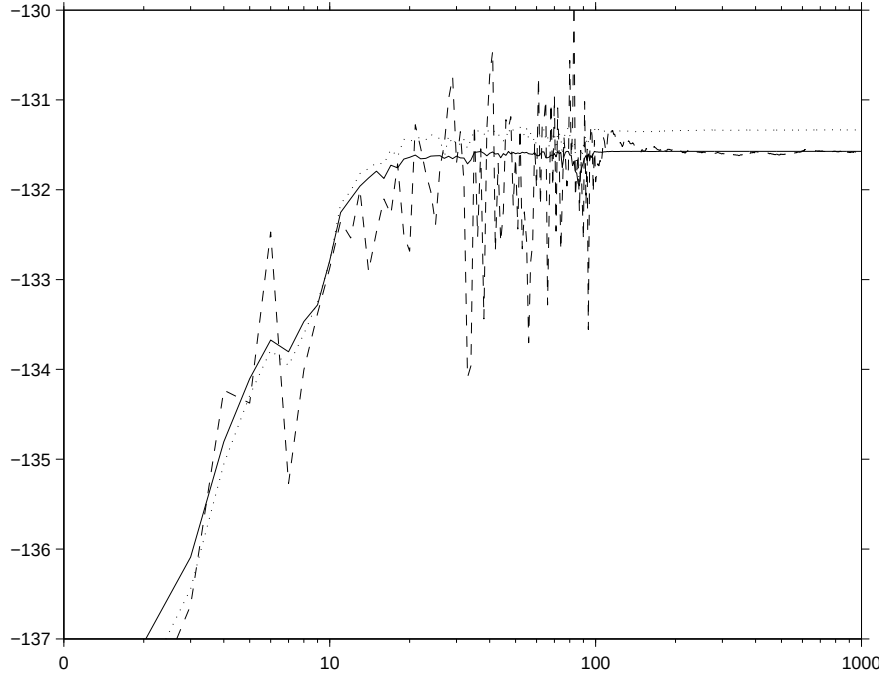


FIG. 2.4: Estimation of the log-likelihood of the observations. The exact log-likelihood is in solid line. The estimated log-likelihood, computed by using the same random sequence at each iteration, is displayed in dotted line. The stochastic approximation of the log-likelihood is in dashed line.



since this one can be computed by using (2.23). When the same random Gaussian sequence is used at each iteration, for drawing $\phi^{(k,1)}, \phi^{(k,2)}, \dots, \phi^{(k,T)}$, with $T = 100$, we can see that the estimated log-likelihood tends to increase at each iteration. We also remark that this sequence converges to an erroneous value. On the other hand, the stochastic approximation proposed in (2.11) converges to the true log-likelihood. In this example, the step size sequence (γ_k) is the same sequence used for estimating the parameters.

The Fisher information of the MLE can also be estimated by using the stochastic approximation scheme presented in Section 2.2.2. We present in Table 2.2 the estimated standard deviation of each component of θ^{EM} and θ^{SAEM} , obtained after 100 and 1000 iterations. We remark again that SAEM gives a good estimation in few iterations. Furthermore, 100 iterations of SAEM, including the estimation of the likelihood and the Fisher information matrix, only require about 1.5s with a Pentium IV.

Of course, we cannot conclude about the performance of SAEM with only one realization of the algorithms. Indeed, the trajectories of EM and SAEM strongly depend on the value of θ_0 . We ran 50 times the SAEM algorithm, with different random initial guess. Each value of θ_0 was independently drawn with a Gaussian distribution of mean $\hat{\theta}^{MLE}$ and standard deviation $0.6\hat{\theta}^{MLE}$. Then, we ran SAEM, and estimated the expected relative deviation between θ_k^{SAEM} and $\hat{\theta}^{MLE}$. For example, the expected relative deviation for the parameter μ , at iteration k , is

TAB. 2.2: Estimation of the standard deviation of θ^{EM} and θ^{SAEM} obtained after 100 and 1 000 iterations.

Parameters	β_1	β_2	μ	τ^2	σ^2
$\hat{\sigma}(\theta_{100}^{EM})$	13.51	13.04	14.15	633.39	14.70
$\hat{\sigma}(\theta_{1000}^{EM})$	13.51	13.04	14.15	633.40	14.70
$\hat{\sigma}(\theta_{100}^{SAEM})$	12.89	12.51	13.83	604.53	13.41
$\hat{\sigma}(\theta_{1000}^{SAEM})$	13.51	13.04	14.15	633.40	14.70

approximated by

$$\text{RDEV}_k(\mu) = \frac{1}{50} \sum_{\ell=1}^{50} \left| \frac{\mu_k^{(\ell)} - \hat{\mu}^{MLE}}{\hat{\mu}^{MLE}} \right|.$$

We also estimated the expected absolute deviation between the log-likelihood $\log q(\mathbf{y}; \theta_k^{SAEM})$ and the maximum log-likelihood $q(\mathbf{y}; \hat{\theta}^{MLE})$ by

$$\text{ADEV}_k(\log q) = \frac{1}{50} \sum_{\ell=1}^{50} (\log q(\mathbf{y}; \theta_k^{(\ell)}) - \log q(\mathbf{y}; \hat{\theta}^{MLE})).$$

The results are displayed Figure 2.5. We clearly see that a good estimation is expected after only 100 iterations. Indeed, only 100 iterations are enough if a relative precision of about 5% is required in the estimation of θ . Also the estimation of the log-likelihood provides good results.

2.3.3 Extension to an heteroscedastic model

Let us consider now the following heteroscedastic model for this same example:

$$y_{ij} = \frac{\phi_i}{1 + \exp\left(-\frac{t_j - \beta_1}{\beta_2}\right)} (1 + \varepsilon_{ij}) \quad \text{for } 1 \leq i \leq n, 1 \leq j \leq m. \quad (2.24)$$

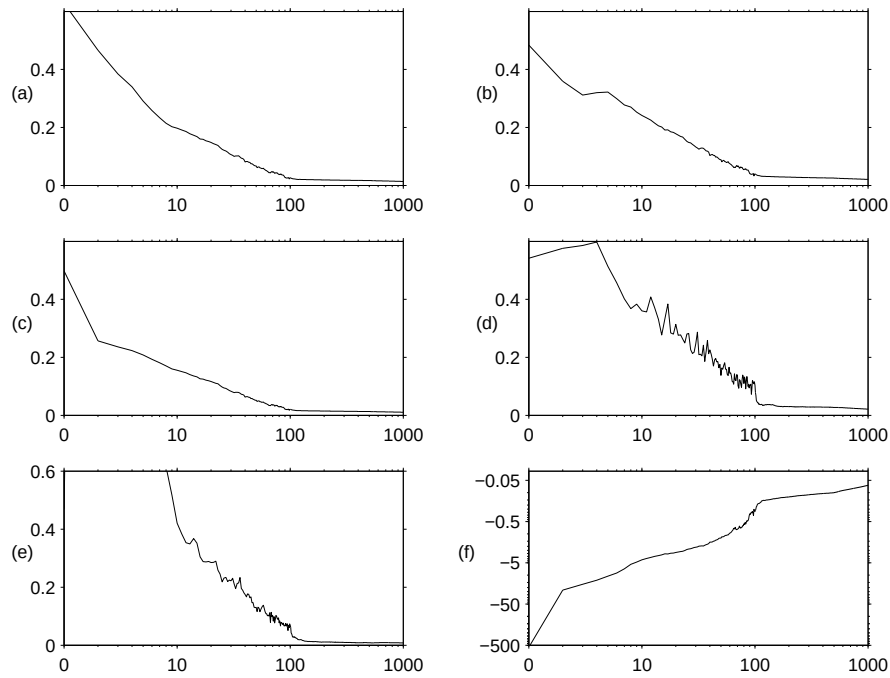
We are outside the scope of the exponential model and SAEM cannot be used as before. The solution consists in regarding the fixed parameters (β_1, β_2) as the realization of a Gaussian random vector of mean (μ_1, μ_2) and diagonal covariance matrix with diagonal terms (τ_1^2, τ_2^2) . As before, $\phi = (\phi_i)$ is a sequence of *i.i.d.* Gaussian random variables of mean μ and variance τ^2 .

It is important to notice that we do not change the model by doing that. The fixed effects remain fixed effects, since we still consider only one vector (β_1, β_2) for the all population.

The complete likelihood of this model is:

$$\begin{aligned} f(\mathbf{y}, \phi; \theta) = & 2\pi\tau_1\tau_2(2\pi\sigma^2)^{-\frac{nm}{2}}(2\pi\tau^2)^{-\frac{n}{2}} \exp \left[-\frac{1}{2\sigma^2} \sum_{i,j} \left(\frac{y_{ij}}{C(t_j, \phi_i; \beta_1, \beta_2)} - 1 \right)^2 \right. \\ & \left. - \sum_{i,j} \log(C(t_j, \phi_i; \beta_1, \beta_2)) - \sum_{i=1}^n \frac{(\phi_i - \mu)^2}{2\tau^2} - \frac{(\beta_1 - \mu_1)^2}{2\tau_1^2} - \frac{(\beta_2 - \mu_2)^2}{2\tau_2^2} \right]. \end{aligned} \quad (2.25)$$

FIG. 2.5: (a-e) The five components of the mean deviation between θ_k^{SAEM} and $\hat{\theta}^{MLE}$, (f) The mean deviation between the log-likelihood $\log q(\mathbf{y}; \theta_k^{SAEM})$ and the maximum log-likelihood $\log q(\mathbf{y}; \hat{\theta}^{MLE})$.



TAB. 2.3: Heteroscedastic model: mean value and standard deviation of the 50 estimations of θ , after 1000 iterations of SAEM.

	Parameters	β_1	β_2	μ	τ^2	σ^2
$\tau_1^2 = \tau_2^2 = 10$	Mean value	757.29	378.78	197.50	722.48	$8.5 \cdot 10^{-3}$
	Standard deviation	11.80	4.96	2.18	17.61	$3 \cdot 10^{-5}$
$\tau_1^2 = \tau_2^2 = 1000$	Mean value	756.43	379.45	195.65	724.28	$8.9 \cdot 10^{-3}$
	Standard deviation	13.59	5.18	2.59	19.13	$4 \cdot 10^{-5}$

where $\theta = (\mu_1, \mu_2, \tau_1^2, \tau_2^2, \mu, \tau^2, \sigma^2)$ is the new vector of parameters and $C(t_j, \phi_i; \beta_1, \beta_2) = \frac{\phi_i}{1 + \exp\left(-\frac{t_j - \beta_1}{\beta_2}\right)}$. This model is now clearly exponential and we are able to apply the SAEM algorithm for estimating $(\mu_1, \mu_2, \mu, \tau^2, \sigma^2)$. Then we can use the estimation of (μ_1, μ_2) as estimation of (β_1, β_2) .

We ran 50 times the SAEM algorithm (1000 iterations), with different random initial guess. The mean value and the standard deviation of the 50 estimations of $(\mu_1, \mu_2, \mu, \tau^2, \sigma^2)$ are displayed in Table 2.3. It is interesting to remark that the estimation of the parameters is not influenced at all by the choice of (τ_1^2, τ_2^2) . Indeed, the same results are obtained with $\tau_1^2 = \tau_2^2 = 10$, or $\tau_1^2 = \tau_2^2 = 1000$. The results of Table 2.3 clearly also show that convergence of SAEM does not depend on the initial value. Furthermore, the estimated values of (μ, β_1, β_2) are not very different from the values obtained with an additive model.

2.4 Comparisons with other methods used in PK/PD.

This section is devoted to some examples coming from pharmacokinetic and pharmacodynamic modeling. This field motivates lots of works about nonlinear mixed effects models. For examples, we can cite the growth-curves analysis like this of the orange trees, the analysis of dose-response curves (see Grizzle and Allen (1979), Mentré and Gomeni (1995)) like the examples proposed in sections 2.4.1, 2.4.2 and 2.4.3 or the analysis of viral decay after some treatment (see Ding and Wu (2001)) like in section 2.5.

We consider now the mixed effects models from the point of view of PK/PD field. In fact, in those kinds of models, the observations are generally some repeated measurements done of a group of individuals. So in most cases, there are two types of parameters which correspond respectively to the fixed effects and to the random effects, namely the population parameters which are the same for all the individuals and the individual parameters which are varying with the individual. Usually the individual parameters and the with-in group errors are supposed to be normally distributed with mean zero. Our purpose consists then in estimating the population parameters and the variances of all the Gaussian distributions involved by the random effects and the error terms, from the data of a sample of individuals, each individual having one or more measurements.

The first mixed effects models involved in PK/PD were linear (see for example Grizzle and Allen (1979)). Since the models are linear, the EM algorithm allows the estimation of the population parameters and of the unknown variances of the model. Some applications are presented in Laird and Ware (1982) and in Lindstrom and Bates (1988). Nevertheless, modeling

practical problems in PK/PD only with linear mixed effects models is very limited. So nonlinear mixed effects models were introduced, involving practical computing problems, because the likelihood is expressed due to an integral which has no analytical expression. Some solutions were proposed, for example, Racine-Poon (1985) uses a Bayesian method assuming some prior distribution on the parameters. But the most usual way to get round the previous difficulty is to do some kind of linearization on the model function in order to obtain some linear approached model. The common methods used by the professional in this field are implemented in the software NONMEM, namely first-order method (FO) and first-order conditional estimates (FOCE) (see Lindstrom and Bates (1990) and section 2.5 for more details). Both are based on a first order Taylor expansion of the response about the fixed effects. Other two-step iterative algorithms were proposed more recently, for example the algorithm P-Pharm of Mentré and Gomeni (1995) (see section 2.4.3 for more details). A common characteristic is that the first step usually provides an estimation of the individuals parameters and that the second step provides estimations of the population parameters using the chosen linearization. Some problem appearing with those method is that the quality of the approximation realized by the linearization is difficult to quantify, so that it is often impossible to get any theoretical result of the convergence. For all these reasons, the application of the SAEM algorithm in this field seems to be very promising, because theoretical result exists and the assumptions done can be easily satisfied in practice for those kinds of models.

2.4.1 A first pharmacokinetic model

This example was proposed by Concordet and Nunez (2002). It is a kinetic population homoscedastic model, used for example for analyzing the concentration obtained after a constant drug diffusion. The data were simulated according to:

$$y_{ij} = \phi_{i1}(1 - \exp[-\phi_{i2}t_j]) + \varepsilon_{ij} \text{ for } 1 \leq i \leq n, 1 \leq j \leq m,$$

where y_{ij} denotes the concentration on time $t_j = j$ for the individual i with $n = 30$ and $m = 7$. The random effects $(\phi_i) = ((\phi_{i1}, \phi_{i2}))$ are independent and identically normally distributed with mean $\mu = (20, 0.5)$ and covariance matrix Γ

$$\text{with } \Gamma = \begin{pmatrix} 4 & 0.0574 \\ 0.0574 & 0.00328 \end{pmatrix}.$$

The errors (ε_{ij}) are independent and follow a Gaussian distribution with mean zero and variance $\sigma^2 = 16$. We suppose also that the random effects (ϕ_i) and the errors (ε_{ij}) are mutually independent. The parameters of the model are then equal to $\theta = (\mu, \Gamma, \sigma^2)$.

Following Concordet and Nunez (2002), the starting values of the parameters are chosen equal to their true value. Indeed, they do not study the optimization method, but the performances of the different estimates.

We propose to highlight on this example how the different methods presented in section 2.2.2 for improving the performances of the algorithm work in practice. We first run the algorithm with 150 iterations using one chain simulated at iteration k with the following transition kernel which is repeated 200 times at each simulation:

1. For $i = 1 \cdots n$, draw $\phi'_i = (\phi'_{i1}, \phi'_{i2})$ i.i.d. with the prior distribution $\mathcal{N}(\mu_k, \Gamma_k)$ and $\mathbf{u} = (u_1, \dots, u_n)$ i.i.d. with the uniform distribution $\mathcal{U}([0, 1])$.

2. For $i = 1 \cdots n$, compute

$$\Delta_i = \frac{h(y_i | \phi'_i; \sigma_k^2)}{h(y_i | \phi_i; \sigma_k^2)} = \exp \left[\frac{1}{2\sigma_k^2} \sum_{j=1}^m \left[(y_{ij} - \phi_{i1}(1 - \exp[-\phi_{i2}t_j]))^2 - (y_{ij} - \phi'_{i1}(1 - \exp[-\phi'_{i2}t_j]))^2 \right] \right],$$

where h designs the posterior distribution of y_i for all $1 \leq i \leq n$.

3. For $i = 1 \cdots n$, set

$$\begin{aligned} \phi_i^{(k+1)} &= \phi'_i & \text{if } \Delta_i \geq u_i \\ \phi_i^{(k+1)} &= \phi_i^{(k)} & \text{elsewhere.} \end{aligned}$$

All the results presented in the two next tables are based on 20 simulations. With this method noted 1, we obtain the results presented in the first column in Table 2.4 for the bias and Table 2.5. The method 2 consists in simulating not just one chain but 10 chains in parallel, without changing anything else. These results are presented in the second columns of Tables 2.4 and 2.5. The method 3 introduces another transition kernel for simulating the chains: since the model function is linear in the random effect ϕ_{i1} , we can easily evaluate the density function $\tilde{p}$ of ϕ_{i2} conditionally to y_i by integration in ϕ_{i1} . Since the acceptance probability associated is also easy to evaluate, we propose to use as proposal distribution the prior distribution $\tilde{\pi}$ of ϕ_{i2} to draw a new candidate ϕ'_{i2} . After evaluating the corresponding acceptance probability which leads to accept or not this new value ϕ'_{i2} , we finally draw ϕ_{i1} from the conditional distribution knowing y_i and the re-actualized value of ϕ_{i2} . So we use the following scheme for simulating the chains $\phi^{(k)}$:

1. For $i = 1 \cdots n$, draw ϕ'_{i2} i.i.d. with the prior distribution $\tilde{\pi}$ equal to $\mathcal{N}(\mu_{k2}, \Gamma_{k22})$ where $\mu_k = (\mu_{k1}, \mu_{k2})$ and $\Gamma_k = \begin{pmatrix} \Gamma_{k11} & \Gamma_{k12} \\ \Gamma_{k12} & \Gamma_{k22} \end{pmatrix}$ and $\mathbf{u} = (u_1, \dots, u_n)$ i.i.d. with the uniform distribution $\mathcal{U}([0, 1])$.

2. For $i = 1 \cdots n$, compute

$$\Delta_i = \frac{\tilde{p}(\phi'_{i2} | y_i; \theta_k) \tilde{\pi}(\phi_{i2}; \theta_k)}{\tilde{p}(\phi_{i2} | y_i; \theta_k) \tilde{\pi}(\phi'_{i2}; \theta_k)}.$$

3. For $i = 1 \cdots n$, set

$$\begin{aligned} \phi_{i2}^{(k+1)} &= \phi'_{i2} & \text{if } \Delta_i \geq u_i \\ \phi_{i2}^{(k+1)} &= \phi_{i2}^{(k)} & \text{elsewhere.} \end{aligned}$$

4. For $i = 1 \cdots n$, draw $\phi_{i1}^{(k+1)}$ from the conditional distribution of ϕ_{i1} conditionally to $(\phi_{i2}^{(k+1)}, y_i)$ which is also Gaussian and can be easily evaluated using Bayes's formula.

We alternate this simulation method with the previous and repeated both only 100 times to ensure the same number of iterations in one MCMC procedure. The method 4 adds adaptive step size, since it was deterministic before. Our purpose is to reduce the step size if the stochastic perturbation $\tilde{S}(y, \phi^{(k+1)}) - s_k$ added to the sequence of minimal sufficient statistics has changed of direction, concerning the optimization problem, so that the sequence will not go further in a bad direction. Indeed, we evaluate the scalar product between the current stochastic perturbation $\tilde{S}(y, \phi^{(k+1)}) - s_k$ and the previous one $\tilde{S}(y, \phi^{(k)}) - s_{k-1}$. If it is negative, we reduce the step size

TAB. 2.4: First pharmacokinetic model: comparison of the bias obtained by SAEM with the different methods, based on 20 simulations.

Parameters	1	2	3	4	5
μ_1	0.027	0.014	0.052	0.018	0.005
μ_2	-0.002	-0.001	-0.002	0.001	0.001
Γ_{11}	3.654	3.386	1.243	0.667	-0.671
$\Gamma_{12} \times 100$	4.717	4.611	1.945	1.469	-0.069
$\Gamma_{22} \times 1000$	2.977	2.576	2.163	0.686	0.120
σ^2	-3.545	-3.128	-1.169	-0.533	0.459

TAB. 2.5: First pharmacokinetic model: comparison of the square roots mean squared errors obtained by SAEM with the different methods, based on 20 simulations.

Parameters	1	2	3	4	5
μ_1	0.604	0.569	0.569	0.518	0.525
μ_2	0.025	0.025	0.024	0.011	0.013
Γ_{11}	3.952	3.834	2.698	1.679	1.640
$\Gamma_{12} \times 100$	6.491	6.260	5.390	3.808	3.675
$\Gamma_{22} \times 1000$	3.256	2.933	2.697	1.186	1.110
σ^2	4.023	3.661	2.712	1.918	1.779

for example by updating it according to $\gamma_{k+1} = (\frac{\gamma_k+1}{\gamma_k})^{-1}$. If it is positive, we set $\gamma_{k+1} = \gamma_k$. Finally, the method 5 uses only the second transition kernel defined above with the adaptative step size.

Concerning the results we obtained, the increase of the number of Markov chains generated in parallel reduces in the same time the bias and the MSE's, but the order are roughly the same. When we change the simulation kernel used in the MCMC procedure, we improve globally the results, for the bias as much as for the MSE's. Likewise, the implementation of the adaptative step size contributes considerably to the improvement of both and we obtain finally the best results using only the second transition kernel and the adaptative step size. We will now compare these square roots of the MSE's with these obtained by Concordet and Nunez (2002).

In their article, Concordet and Nunez (2002) estimate the parameters of the model by maximizing a simulated pseudo-likelihood. Indeed, they consider a penalised least square criterion equal to this one which would be issu from a Gaussian model where the means and the variances were estimated by a Monte Carlo method, implying the denomination *simulated pseudo-likelihood*. We reproduce in Table 2.6 their simulation results, comparing FOCE (First-Order Conditional Estimation), Gaussian quadrature and their SPML (Simulated Pseudo Maximum Likelihood) method. These results are compared with the SAEM algorithm. Here, the number of iterations in SAEM was also taken equal to 150.

For each algorithm, the square roots of the MSE's were estimated, based on 20 simulations. We clearly see that the maximum likelihood estimate, obtained with SAEM has the smallest MSE's, for the mean parameters but also for the variance-covariance matrix. Indeed, Concordet and Nunez (2002) have shown that the SPML estimator is consistent. Their method provides a

TAB. 2.6: First pharmacokinetic model: comparison of parameters estimates. The estimated square roots of the MSE's, based on 20 simulations.

Parameters	FOCE	Gauss	SPML	SAEM
μ_1	1.72	0.58	0.82	0.52
μ_2	0.09	0.04	0.06	0.01
Γ_{11}	7.03	4.78	2.96	1.64
$\Gamma_{12} \times 100$	42.01	12.19	15.40	3.67
$\Gamma_{22} \times 1000$	24.70	6.20	11.84	1.11
σ^2	1.70	2.05	1.73	1.78

good estimation for large values of n , but not for a small number of individuals.

2.4.2 A pharmacodynamic model

We consider in this section the nonlinear population pharmacodynamic model used by Walker (1996).

Simulated data are given by

$$y_{ij} = \phi_{i1} - \frac{\phi_{i2}t_j}{\phi_{i3} + t_j} + \varepsilon_{ij} \quad \text{for } 1 \leq i \leq n, 1 \leq j \leq m, \quad (2.26)$$

with $n = 30$, $m = 6$, $t_1 = 0$, $t_2 = 5$, $t_3 = 10$, $t_4 = 20$, $t_5 = 40$ and $t_6 = 80$. The random effects $(\phi_i) = ((\phi_{i1}, \phi_{i2}, \phi_{i3}))$ are supposed to be independent and are simulated with Gaussian distribution with mean $\boldsymbol{\mu} = (\mu_1, \mu_2, \mu_3)$ and diagonal covariance matrix $\boldsymbol{\gamma}^2 = (\gamma_1^2, \gamma_2^2, \gamma_3^2)$. The additive noise (ε_{ij}) are also supposed independent and normally distributed with mean zero and variance equal to σ^2 . Moreover, we suppose that the random effects (ϕ_i) and the noise (ε_{ij}) are mutually independent. The parameters of the model are then equal to $\boldsymbol{\theta} = (\boldsymbol{\mu}, \boldsymbol{\gamma}^2, \sigma^2)$ and the values chosen for the simulation are the following:

$$\varepsilon_{ij} \sim_{iid} \mathcal{N}(0, 4), \quad \phi_{i1} \sim_{iid} \mathcal{N}(105, 64), \quad \phi_{i2} \sim_{iid} \mathcal{N}(12, 36), \quad \phi_{i3} \sim_{iid} \mathcal{N}(10, 12.25).$$

According to Sheiner, Hashimoto, and Beal (1991) and Walker (1996), this model can be used for the analysis of blood pressure $\mathbf{y}$ as a function of the dose of an anti-hypertensive drug from a longitudinal study.

Walker (1996) compares different popular methods of estimation, such as FOCE (First-Order Conditional Estimation) and LAPLACIAN methods of NONMEM. He computes the standard errors (as the square roots of the MSE's) for these different estimators based on 50 simulations. Table 2.7 reproduces these values, with also the standard errors obtained with SAEM. The results concerning the parameter σ^2 are not given since Walker does not produced them in his paper. We see that the MCEM algorithm of Walker and SAEM give similar results, but it is important here to remark that only 300 iterations of SAEM are performed for each single simulated data set. Then, computing time for SAEM is very much reduced, in comparison of the Monte-Carlo EM that requires to sample 10,000 random variates at each iteration and converges very slowly.

Table 2.7 also gives in the last column the estimation of the standard deviation of the MLE, using the approach proposed in Section 2.2.2. This method seems to be very accurate, since these values are close to the empirical standard deviation computed from the 50 simulations.

TAB. 2.7: Pharmacodynamic model: comparison of parameters estimates. The estimated square roots of the MSE's, based on 50 simulations.

Parameters	FOCE	LAPLACIAN	EM	SAEM	$\hat{\sigma}(\hat{\theta}^{ML})$
μ_1	1.8	1.6	1.7	1.5	1.4
μ_2	1.2	1.3	1.3	1.3	1.0
μ_3	2.8	2.4	1.6	0.9	0.6
γ_1^2	20.9	20.2	20.1	16.6	14.5
γ_2^2	11.0	11.9	10.7	10.8	7.6
γ_3^2	7.6	6.1	2.9	3.0	2.8

2.4.3 A second pharmacokinetic model

We consider here the pharmacokinetic model used by Mentré and Gomeni (1995), which is close to the one used by Hashimoto and Sheiner (1991). The aim is to analyze the evolution in time of the plasma concentration after an injection of a dose D of some drug at a given time. The data were simulated following an open one-compartment model with first-order absorption, which involves three individuals parameters : the absorption rate-constant Ka , the volume of distribution V and the total clearance of the drug Cl . The bioavailability was considered fixed to 1. The predicted plasma concentration noted C_p at a time t after a dose D is given by:

$$C_p(Ka, V, Cl, t) = D \frac{Ka}{Cl - V.Ka} \left[\exp(-Ka.t) - \exp\left(-\frac{Cl.t}{V}\right) \right].$$

So let us denote as before the random effect $\phi = (Ka, V, Cl)$. We consider an heteroscedastic error, such that the observation y_{ij} is given by:

$$y_{ij} = C_p(\phi_i, t_j)(1 + \varepsilon_{ij}), \text{ for } 1 \leq i \leq n, \quad 1 \leq j \leq m_i,$$

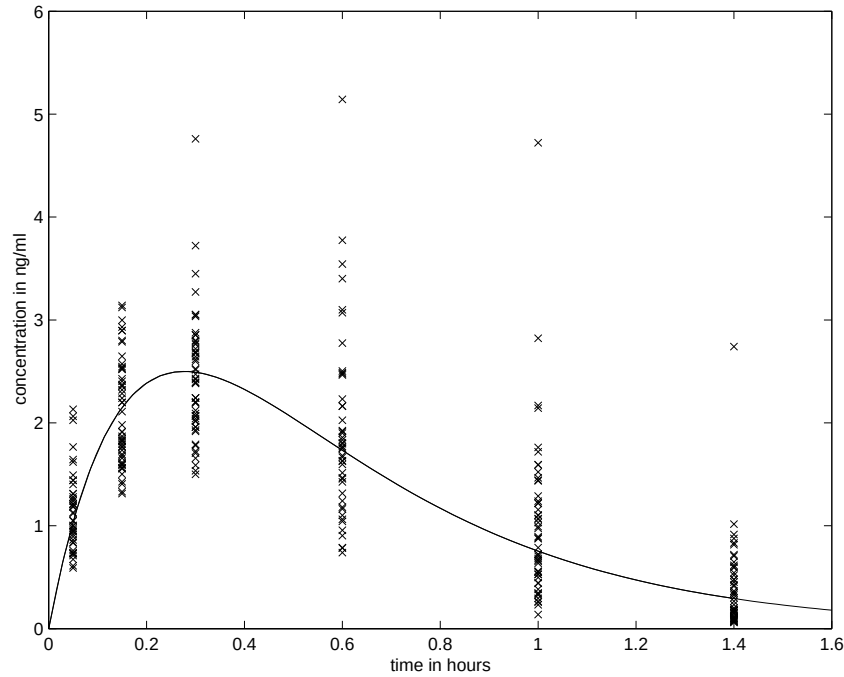
where (ε_{ij}) are independent normally distributed with mean zero and variance σ^2 . We suppose also that the random (ε_{ij}) and the individual parameters (ϕ_i) are mutually independent. By the way, the three individuals parameters $(\phi_i = (Ka_i, V_i, Cl_i))$ are supposed to be independent and normally distributed. The age of the individuals, noted z , was considered as a covariate of the model. For the simulation, we draw the values of the covariable z from a Gaussian distribution with mean 50 and variance 20^2 . This covariate only affects the individual parameter Cl : its mean is supposed to be equal to $m_{Cl} - m_{Cov}z$ and its variance to w_{Cl}^2 . The two others individuals parameters Ka and V have respectively means and variances equal to m_{Ka} and m_V and w_{Ka}^2 and w_V^2 :

$$\varepsilon_{ij} \sim_{iid} \mathcal{N}(0, \sigma^2), \quad Ka_i \sim_{iid} \mathcal{N}(m_{Ka}, w_{Ka}^2),$$

$$V_i \sim_{iid} \mathcal{N}(m_V, w_V^2), \quad Cl_i \sim_{iid} \mathcal{N}(m_{Cl} - m_{Cov}z_i, w_{Cl}^2).$$

So the parameters θ of the model are equal to $(m_{Ka}, m_V, m_{Cl}, m_{Cov}, w_{Ka}^2, w_V^2, w_{Cl}^2, \sigma^2)$. For each individual, three times of measurement were randomly chosen among six times given in hours after the injection of the dose D , namely $\{0.05, 0.15, 0.3, 0.6, 1.0, 1.4\}$. We only kept the observations which were upper than 0.05, so that some individuals may have less than three observations, may be zero observation. This was done to mimic a real situation, namely the

FIG. 2.6: Simulated concentrations of one data set of 100 individuals. The concentration-time curve obtained with the exact mean parameters is plotted in solid line.

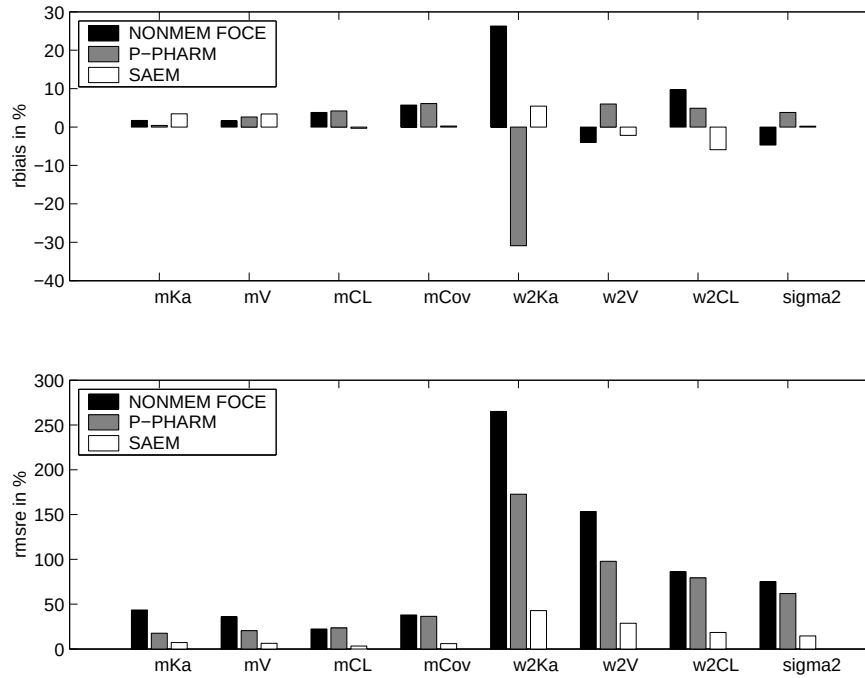


fact that very low concentrations can not be measured because of assay detection limit. The simulation presented below was done with the same values of the parameters as in the paper of Mentré and Gomeni (1995), namely $m_{Ka} = 5$, $m_V = 2$, $m_{Cl} = 10$, $m_{Cov} = 0.1$, $w_{Ka}^2 = 1$, $w_V^2 = 0.16$, $w_{Cl}^2 = 1$ and $\sigma^2 = 0.15$. For example, Figure 2.6 displays the concentrations obtained after simulation for one data set of 100 individuals. Here we kept only 298 observations upper than 0.05. For the study, we simulate as Mentré and Gomeni thirty data sets and use the same starting values for the algorithm, namely $m_{Ka,0} = 6$, $m_{V,0} = 4$, $m_{Cl,0} = 4$, $m_{Cov,0} = 0.0001$, $w_{Ka,0}^2 = 2$, $w_{V,0}^2 = 1$, $w_{Cl,0}^2 = 2$ and $\sigma_0^2 = 0.04$.

The article of Mentré and Gomeni (1995) compares the results obtained with the algorithm P-Pharm to these produced with the usual methods used in pharmacokinetic FO and FOCE proposed in the software NONMEM. The algorithm P-Pharm is an EM like iterative algorithm, which proceed in two steps : during the first step, the individuals parameters are approached by the maximum of their posterior distribution given the current population parameters. The second step consists in the estimation by maximization of the likelihood given the current estimates of the individuals parameters, obtained after a first order expansion of the model function about it.

Mentré and Gomeni (1995) propose an evaluation on simulated data to assess the performance of their algorithm. We present in Tables 2.8 and 2.9 the results obtained by FOCE and P-Pharm versus these obtained by the SAEM algorithm (we do 100 iterations each containing 100 MCMC steps generating 5 chains). Figure 2.7 displays a graphic representation of these results. The mean relative bias obtained by SAEM are globally of the same order than these obtained by P-Pharm or FOCE, excepted for the parameter w_{Ka}^2 for which it is five times lower.

FIG. 2.7: Mean relative bias and root mean squared relative error in % obtained for each parameters from the 30 data sets.



TAB. 2.8: Second pharmacokinetic model: comparison of the mean relative bias in % of the parameters estimates, based on 30 simulations.

Parameters	FOCE	P-Pharm	SAEM
m_{Ka}	1.7	0.4	3.4
m_V	1.6	2.6	3.4
m_{Cl}	3.8	4.2	-0.3
m_{Cov}	5.7	6.1	0.3
w_{Ka}^2	26.3	-30.9	5.4
w_V^2	-4.0	6.0	-2.2
w_{Cl}^2	9.7	4.9	-5.9
σ^2	-4.7	3.8	0.2

TAB. 2.9: Second pharmacokinetic model: comparison of the root mean squared relative error of the parameters estimates in %, based on 30 simulations.

Parameters	FOCE	P-Pharm	SAEM
m_{K_a}	43.5	17.5	7.1
m_V	35.9	20.5	6.3
m_{C_l}	22.3	23.4	3.2
$m_{C_{ov}}$	37.9	36.4	6.0
$w_{K_a}^2$	265.1	172.7	42.8
w_V^2	153.3	97.8	28.7
$w_{C_l}^2$	86.3	79.3	18.4
σ^2	75.1	61.9	14.6

By the way, the root mean squared error provided by SAEM are three times lower than these obtained by P-Pharm and five times lower than these obtained by FOCE.

2.5 Application to likelihood ratio test.¹

2.5.1 Framework

The application treated in this section is based on this proposed in the paper of Ding and Wu (2001). The purpose is to assess anti-viral potency of anti-HIV therapies. In fact, it was shown in practice by clinical data that this potency can be well evaluated by some virological marker, the viral load (number of HIV RNA copies). Until recently, evaluation were done only based on the endpoints, which are in fact good markers for evaluation of long-term treatment effects. But these endpoints are also affected by drug resistance, pharmacokinetics, toxicity and other long-term clinical factors. So they may not represent the potency of the therapy faithfully. A possibility to improve the evaluation was to consider some dynamics study of the viral load. At the beginning, the viral dynamic models including treatment effects were very complicated, involving differential equations. Some simplified versions were proposed recently, in particular a two-phase viral decay model by Ding and Wu (1999), who establishes a relationship between viral decay rates and treatment effects. Ding and Wu (2001) believe that viral decay rates are good markers for drug potency, because it can separate evaluation of anti-viral potency from some others characteristics of new anti-HIV therapies. So we are interested in their two-phase viral decay rates which seems to reflect well the treatment effects. To analyze the treatment effects, they use a likelihood ratio test. Nevertheless some problems appear during its implementation, particularly concerning the type I error when the likelihood estimations were computed using estimations of the parameters obtained with others methods. Our purpose was to apply the SAEM algorithm and to assess the results.

¹Joint work with France Mentré and Adeline Leclercq, INSERM.

2.5.2 The model

The study involves a nonlinear biexponential mixed effects model. The observations (y_{ij}) correspond to the decimal logarithm of the viral load and can be written as follow:

$$y_{ij} = \log_{10}[P_{1i} \exp(-d_{1i}t_j) + P_{2i} \exp(-d_{2i}t_j)] + \varepsilon_{ij} \text{ for } 1 \leq i \leq n, \quad 1 \leq j \leq m,$$

where $(P_{1i}, P_{2i}, d_{1i}, d_{2i})$ are individuals parameters which constitute the random effects of the model, (t_j) are the observation times and (ε_{ij}) the with-in group errors. We consider n individuals and suppose each individual is observed m times. Parameters d_{1i} and d_{2i} are the decay rates of the two phases of plasma virus. Parameters P_{1i} and P_{2i} are reparameterized positive parameters depending on some concentrations values. In order to be sure to obtain positive estimations for these random effects, we choose to express its as an exponential, so the individuals parameters become $\phi_i = (\log P_{1i}, \log P_{2i}, \log d_{1i}, \log d_{2i})$ and the model:

$$y_{ij} = \log_{10}[\exp(\log P_{1i}) \exp(-\exp(\log d_{1i})t_j) + \exp(\log P_{2i}) \exp(-\exp(\log d_{2i})t_j)] + \varepsilon_{ij}. \quad (2.27)$$

We consider that the random effects (ϕ_i) are independent and follow a Gaussian distribution with mean $\beta = (\log P_1, \log P_2, \log d_1, \log d_2)$ respectively and diagonal covariance matrix equal to $w^2 = (w_1^2, w_2^2, w_3^2, w_4^2)$ respectively. The errors (ε_{ij}) are also supposed to be independent and normally distributed with mean zero and variance σ^2 . Moreover the random effects (ϕ_i) and the errors (ε_{ij}) are supposed mutually independent. So all the parameters of the model are equal to $\theta = (\beta, w^2, \sigma^2) = (\log P_1, \log P_2, \log d_1, \log d_2, w_1^2, w_2^2, w_3^2, w_4^2, \sigma^2)$.

2.5.3 The likelihood ratio test

Knowing that the two decay rates of plasma virus are monotone functions of the treatment potency of anti-viral therapy (see Ding and Wu (1999)), the anti-viral effects can be assessed by comparing the two decay rates between two treatments or two patient groups using the relationships between the decay rates and the treatment effects. The natural null hypothesis (**H0**) chosen by Ding and Wu (2001) is that all the population parameters are the same for groups A and B , i.e. a patient has the same virological response to treatments A and B . The alternative hypothesis (**H1**) also chosen by Ding and Wu (2001) is that all the population parameters are the same for groups A and B , excepted for the parameter $\log d_1$ which is supposed to be different for the two treatments. So the likelihood ratio test consists in three steps: first, fit the data under (**H0**) by maximum likelihood and compute the value of the log-likelihood function l_0 , second, fit the data under (**H1**) by maximum likelihood and compute the value of the log-likelihood function l_1 , third, compute the test statistic $\Delta = -2(l_1 - l_0)$ and compare it with the quantile of the χ_ν^2 distribution where ν is the difference of degrees of freedom between the alternative hypothesis and the null hypothesis. The alternative model we choose is one of this proposed by Ding and Wu (2001). We add a parameter β to the first decay rates as follows, so that the mean of the random effect $\log d_{1i}$ is equal to $\log d_1 + \beta$. So we have one degree of freedom for our likelihood ratio test since the model under (**H1**) has one parameter more than the model under (**H0**).

2.5.4 Estimation of the likelihood

First, the parameters were estimated by the SAEM algorithm. Then we use these estimations for approximating by a Monte Carlo method the likelihood value under (**H0**) on the one hand

and under **(H1)** on the other hand. To improve this integral approximation, we implement some importance sample like describe below.

Consider the log-likelihood l of the model:

$$l(\theta) = \sum_{i=1}^n \log g(y_i; \theta),$$

where g denote the density function of $y_i = (y_{i1}, \dots, y_{im})$ for $1 \leq i \leq n$. We introduce now the distribution h of y_i conditionally to ϕ_i and the prior distribution π of ϕ_i for $1 \leq i \leq n$, so that we have

$$g(y_i; \theta) = \int h(y_i | \phi_i; \theta) \pi(\phi_i; \theta) d\phi_i,$$

which allows us to use a Monte Carlo method to approximate the value of this integral. The basic method consists in sampling $(\phi_i^l)_{1 \leq l \leq T}$ for a big value T from the distribution $\pi(\cdot; \theta)$ and to compute $\frac{1}{T} \sum_{l=1}^T h(y_i | \phi_i^l; \theta)$ which approaches as better $g(y_i; \theta)$ as T is great. To improve this approximation, we propose to use some importance sampling which will favor specifically values for ϕ_i , for example the posterior mean of ϕ_i , by using some well chosen instrumental distribution $\tilde{\pi}$ instead of π to sample the random $(\phi_i^l)_{1 \leq l \leq T}$. So we obtain

$$g(y_i; \theta) = \int h(y_i | \phi_i; \theta) \frac{\pi(\phi_i; \theta)}{\tilde{\pi}(\phi_i; \theta)} \tilde{\pi}(\phi_i; \theta) d\phi_i$$

and we approach this integral by $\frac{1}{T} \sum_{l=1}^T h(y_i | \phi_i^l; \theta) \frac{\pi(\phi_i^l; \theta)}{\tilde{\pi}(\phi_i^l; \theta)}$. We choose for the instrumental distribution $\tilde{\pi}$ the Gaussian distribution with mean $\tilde{m}$ and variance $\tilde{w}^2$ respectively equal to the posterior mean and variance of ϕ_i . These value are obtained by stochastic approximation by the SAEM algorithm during the same run which provides the estimation of the parameters of the model.

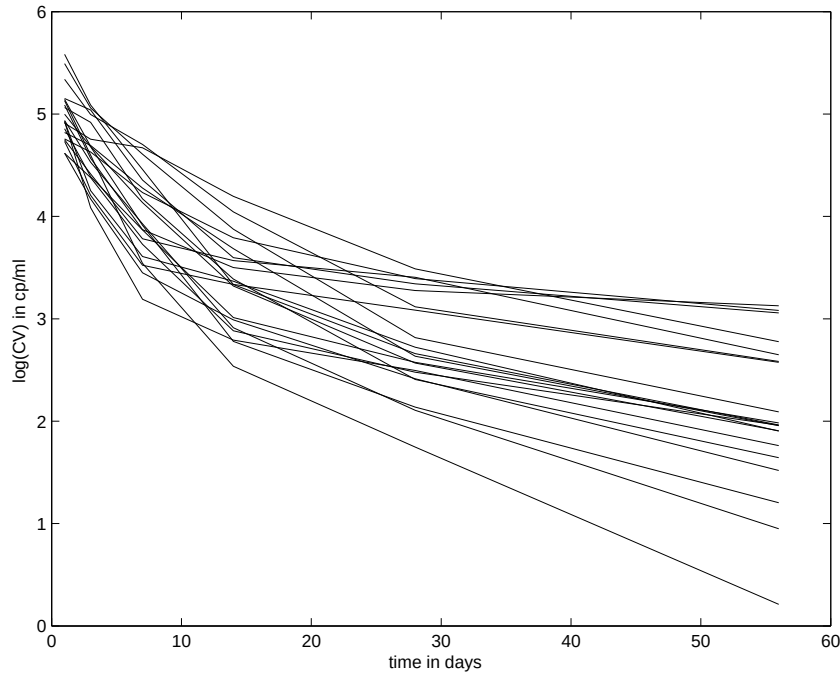
2.5.5 Numerical application

The simulation concerning the NLME method were done in S-Plus by Adeline Leclercq (INSERM).

The simulated data

To estimate empirically the type I error of the likelihood ratio test, we simulate about 100 samples of 40 subjects under **(H0)**, i.e. with the same population parameters for the 20 subjects of group A that for the 20 subjects of group B . We suppose that six measurements were done for each subject at days 1, 3, 7, 14, 28 and 56. The values chosen for the parameters for the simulations of the data were the following: $\log P_{1,0} = 12$, $\log P_{2,0} = 8$, $\log d_{1,0} = \log(0.5)$, $\log d_{2,0} = \log(0.05)$. All the variances of the random effects were chosen equal to 0.3. The variance σ^2 of the error was chosen equal to 0.004 in order to have a variability coefficient of 15%. Figure 2.8 presents the 20 observations simulated which composed one of the two groups.

FIG. 2.8: Simulated values for 20 subjects.



Estimation of the parameters of the model

We compare two estimation methods, on the one hand the usual NLME function implemented for example in *R* using the first-order conditional estimates (FOCE) method and on the other hand the SAEM algorithm. As precise above, the FOCE algorithm is an iterative algorithm. Each iteration of this algorithm proceed in two steps. Knowing the values obtained at the previous iteration for the population parameters and for the individual parameters, the first step consists in the reactualization of the individual parameters by a Bayesian estimation, namely the maximum a posteriori involving the Bayes's formula. Some calculation show that this is equivalent to the minimization of some penalized nonlinear least squares (PNLS) like explain in Pinheiro and Bates (1995). This minimization problem is done with some iterations of a Newton-Raphson algorithm. The second step consists in a linearization by a first order Taylor expansion of the decimal logarithmic function in (2.27) about the new current values of the individual parameters providing an approximation of the likelihood. The population parameters are updated by maximization of this pseudo likelihood. This maximization problem also involves some iterations of a Newton-Raphson method. The NLME function implemented in *R* allows at the most 50 iterations for the FOCE algorithms with at the most 7 Newton-Raphson iterations for the PNLS step and 50 Newton-Raphson iterations for the likelihood maximization. If the algorithm does not converge after all these iterations, i.e. the convergence criterion was not reached (difference between two consecutive estimations lower than 10^{-6}), the run failed and some error message is given.

The SAEM was implemented with 200 iterations of the algorithm and 5 chains updated at each iteration using 50 iterations for the MCMC procedure.

Table 2.10 shows that the quality of the estimation is globally the same with SAEM than with

TAB. 2.10: Relative bias and relative root mean squared error in % obtained under H0.

Parameters	Bias NLME	Bias SAEM	Rmsre NLME	Rmsre SAEM
$\log P_1$	0.19	0.16	0.82	0.80
$\log P_2$	0.53	0.01	1.35	1.20
$\log \lambda_1$	-3.48	-2.36	14.65	14.22
$\log \lambda_2$	-0.93	-0.02	2.79	3.85
w_1^2	-3.25	-3.10	38.56	27.72
w_2^2	9.48	-2.04	66.38	30.96
w_3^2	-1.33	1.45	23.73	23.25
w_4^2	-4.09	3.82	21.73	24.22
σ^2	2.41	0.47	16.23	15.79

TAB. 2.11: Relative bias and relative root mean squared error in % (excepted for the parameter β) obtained under H1.

Parameters	Bias NLME	Bias SAEM	Rmsre NLME	Rmsre SAEM
$\log P_1$	0.19	0.16	0.82	0.80
$\log P_2$	0.45	0.03	1.28	1.22
$\log \lambda_1$	-4.19	-3.40	19.53	19.30
$\log \lambda_2$	-0.90	-0.02	3.82	3.86
w_1^2	-3.50	-3.05	27.78	28.01
w_2^2	1.13	-1.99	49.73	30.37
w_3^2	-2.64	-1.00	23.89	23.45
w_4^2	-4.07	3.87	21.73	24.18
σ^2	2.42	0.67	16.21	15.61
β	-0.12	-0.01	0.17	0.03

FIG. 2.9: Relative bias and relative root mean squared error in % obtained for the parameters under H_0 . On the x-axis, from left to right, $\log P_1, \log P_2, \log d_1, \log d_2, w_1^2, w_2^2, w_3^2, w_4^2$ and σ^2 .

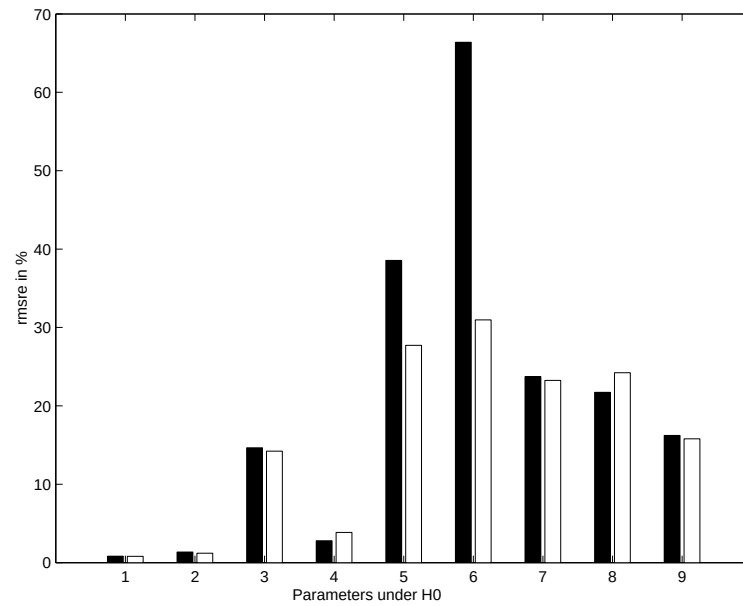
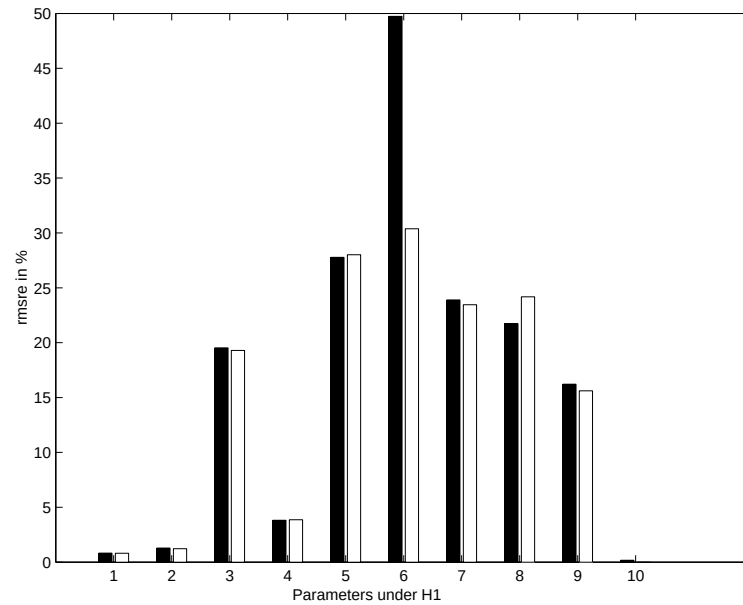


FIG. 2.10: Relative bias and relative root mean squared error in % (excepted for the parameter β) obtained under H_1 . On the x-axis, from left to right, $\log P_1, \log P_2, \log d_1, \log d_2, w_1^2, w_2^2, w_3^2, w_4^2, \sigma^2$ and β .



NLME. Nevertheless, the relative bias obtained with SAEM are far lower than those obtained with NLME. Concerning the relative root mean squared error, there are globally the same, excepted for the variances of w_1^2 and w_2^2 which are quite improved.

Estimation of the type I error of the likelihood ratio test

For each of the 100 samples, we compute using the estimations of the parameters obtained on one hand by NLME and on the other hand by SAEM an estimation of the log-likelihood using the method described in section 2.5.4, one time the log-likelihood l_0 in the model under **(H0)** and one time the log-likelihood l_1 in the model under **(H1)**. Then we evaluate the empirical estimator $\hat{\alpha}$ of the type I error as follow:

$$\hat{\alpha} = \frac{1}{100} \sum_{i=1}^{100} \mathbb{1}(\Delta_i > 3.84),$$

where 3.84 correspond to the 5% quantile of the χ^2 distribution with one degree of freedom.

We obtain the following results:

$$\hat{\alpha}_{NLME} = 0.13 \text{ and } \hat{\alpha}_{SAEM} = 0.07,$$

which seems to be a better estimation since the theoretical value for the type I error is asymptotically equal to 0.05 when n goes to infinity.

Chapitre 3

Joint inversion of teleseismic delay times and gravity anomaly data: A new approach.¹

Summary

We present a new method for joint inversion of gravity and teleseismic delay time data. It is based on the assumption of a linear but unknown relationship between density and velocity variations which may vary with depth. The main novelty of this method is that it allows to estimate the model hyperparameters (in particular variances of the different data sets and of the physical properties) instead of fixing them arbitrarily. We use the stochastic algorithm SAEM for obtaining an estimation of the variances and of others unknown hyperparameters. Then we use them for obtaining an estimation of the velocity and density variations. The algorithm has been applied to different synthetic data sets in order to show its performance.

Ce travail a donné lieu à la rédaction d'un article co-signé par Estelle Kuhn, Marc Lavielle et Hermann Zeyen, actuellement soumis.

¹Joint work with Hermann Zeyen, Département des Sciences de la Terre, U.P.S., Orsay.

3.1 Introduction

Three-dimensional teleseismic tomography has become in the last decade an established method for high resolution investigation of the upper mantle. A mobile array of seismic receivers is generally installed in an area of several 100 km diameter, with the principal interest to record seismic events from distances larger than some 20° . The measured arrival times for each seismic event are stripped of the theoretical arrival time for a global mean model in order to obtain the delay times. For the interpretation of these delay times, different inversion techniques are being employed, which can be summarized into two groups: exclusive inversion of delay times and joint inversion of delay times and other data sets, mainly gravity and geoid data.

Initially, most inversion methods of the first group were based on the ACH algorithm (see Aki, Christofferson, and Husebye (1977)) and later modifications (see Evans and Achauer (1993), Granet and Cara (1988)). These methods subdivide a three dimensional space into rectangular prisms of variable size and perform a least squares inversion regularized by smoothing, usually through singular value decomposition. More recently, Bayesian inversion algorithms were developed that allow a better, more localized treatment of the uncertainties of data and model parameters (e.g. use of a priori information). Examples of these inversion methods include Humphreys and Dueker (1994) and Ritter, Jordan, Christensen, and Achauer (2001). One of the important problems of these inversion schemes concerns "smearing", the numerical propagation of the effects of strong anomalies into regions that are in reality void of velocity anomalies. Smearing in teleseismic tomography inversion is due to reduced resolution along the ray paths, paths that are generally near vertical, inside a cone of $40 - 60^\circ$ opening.

Partly in order to reduce the non-uniqueness of teleseismic tomography, several authors have proposed to use the empirically observed relationship between density and velocity variations (see Glaznev, Raevsky, and Skopenko (1996), Ludwig, Nafe, and Drake (1970)) to constrain the seismic models. Lines, Schultz, and Treitel (1988) introduced the concept of "cooperative inversion", that consists of an iterative procedure, doing delay-time and gravity inversion in separated steps but using the result of one iteration step as starting model for the next one, through a linear velocity-density relation. Later in year 1991, Lees and VanDecar used Birch's law (see Birch (1961)) for a joint inversion of delay time and Bouguer anomaly data, connected through a fixed linear velocity-density relationship. Zeyen and Achauer (1997) introduced an inversion scheme that allows to invert gravity and delay time data assuming a linear but unknown velocity-density relation. Their aim was not only to reduce the non-uniqueness of the single inversions but also to investigate variations in the correlation coefficient between velocity and density variations that may be due to mineralogical or thermal variations.

This article presents a new inversion algorithm. The general ideas and aims of this method are the same as for the approach by Zeyen and Achauer, i.e. a joint inversion of gravity (Bouguer or free air anomalies) and teleseismic delay times under the assumption of a linear relationship between velocity and density variations. The coefficients of this linear relationship are supposed to be unknown and may change in different areas of the model namely in the vertical. The algorithm is based on a Bayesian inversion scheme, however, in contrast to the former algorithm, we use here the maximum likelihood approach for estimating the statistical hyperparameters of the model; we will first introduce the theoretical background and the model definition before we show results of the inversion of data produced with synthetic models with different degrees of complexity. The results will be compared with those of earlier methods in order to demonstrate the advantages of the new method.

3.2 Method

The inversion scheme is based on a model consisting of horizontal layers with variable thickness, each of which is composed of rectangular prisms of variable size. In principal, a slowness (reciprocal of seismic velocity) and a density are attributed to each prism. However, in general, not all prisms of a layer are crossed by a sufficient number of seismic rays. Therefore, one has to account for the possibility that for certain prisms only one parameter, the density, is inverted. In order to minimize boundary effects, we work with relative density and velocity variations.

For each layer, we assume a constant but unknown linear velocity-density relation, simulating mainly a temperature dependence of this relationship (Sobolev, Zeyen, Granet, Stoll, Achauer, Bauer, Werling, Altherr, and Fuchs (1997)). Although this is the application in our models, it is possible, if geological evidence exists, to split each layer into regions with different relationships or to join several layers into one single region.

3.2.1 Presentation of the probabilistic model

The chosen framework is the one of incomplete data. We denote the observed data y (gravity and delay time data) and the non observed data x (density and velocity variations). We consider the following model:

$$y = h(x) + \varepsilon,$$

where h is a known function and ε represents the noise due inter alia to the errors of measurement.

Our purpose is to recover the non observed data x from the observed data y . To do this we will adopt a probabilistic approach and make the following assumptions:

- **(H1)** The non observed data x is a random vector, drawn from a prior distribution π that depends on some set of unknown hyperparameters ϕ .
- **(H2)** The noise ε is a Gaussian white noise with unknown variance σ^2 . Moreover, the non observed data x and the noise ε are mutually independent.

Since the set of hyperparameters $\theta = (\phi, \sigma^2)$ is usually unknown, we propose to estimate them directly from the data and to use this estimation for recovering x . The proposed method can be summarized as follows:

- 1) Compute θ_{ML} , the maximum likelihood estimate of θ , that is the value of θ that maximizes the likelihood $g(y, \theta)$ of the observations.
- 2) Compute the posterior distribution $p(x|y; \theta_{ML})$ of the non observed data x , with the maximum likelihood estimate of θ .

Then, different estimates of x can be defined. For example, the so-called MAP (Maximum a Posteriori) estimate of x maximizes the posterior distribution $p(x|y; \theta_{ML})$. The EPM (Expected Posterior Mean) is the mean of $p(x|y; \theta_{ML})$.

3.2.2 A stochastic algorithm for maximum likelihood estimation

The expectation-maximization (EM) algorithm, proposed by Dempster, Laird, and Rubin (1977), is a well known iterative algorithm for computation of maximum likelihood estimates, useful in a variety of partially-observed-data problems.

The E-step of the EM algorithm computes $Q(\theta|\theta_k) = E(\log f(x, y; \theta)|y; \theta_k)$ where $E(.|y; \theta)$ denotes the conditional expectation and $f(x, y; \theta)$ is the likelihood of the complete data (x, y) .

So we have:

$$E(\log f(x, y; \theta) | y; \theta_k) = \int \log f(x, y; \theta) p(x | y; \theta_k) dx.$$

The M-step determines θ_{k+1} as maximizing $Q(\theta | \theta_k)$. All the statistical models considered here belong to the exponential family. That means that there exist a constant C , two functions Ψ and Φ of θ and a function S of (x, y) such that the joint log-likelihood has the form

$$\log f(x, y; \theta) = C + \Psi(\theta) - \langle S(x, y), \Phi(\theta) \rangle, \quad (3.1)$$

where $\langle \cdot, \cdot \rangle$ denotes the scalar product. The vector $S(x, y)$ is called the sufficient statistics of the model (see paragraph 2.3. for a concrete application).

In this case, the k -th iteration consists in two steps:

- **E-step:** compute

$$S_k = E(S(x, y) | y; \theta_k).$$

- **M-step:** determine θ_{k+1} by maximizing $\Psi(\theta) - \langle S_k, \Phi(\theta) \rangle$.

This procedure is repeated till the convergence, which is ensured under very general assumptions, becomes apparent.

When the E-step is infeasible in a closed form, the SAEM (stochastic approximation of EM) algorithm, proposed by Delyon, Lavielle, and Moulines (1999), introduces a simulation step at iteration k to obtain a stochastic approximation of S_k :

- **S-step:** using the current value θ_k of θ , draw a realization $x^{(k)}$ of the missing data under the posterior distribution $p(x | y; \theta_k)$.
- **A-step:** update the approximation of S_k according to

$$s_k = s_{k-1} + \gamma_k (S(x^{(k)}, y) - s_{k-1}), \quad (3.2)$$

where $\{\gamma_k\}_{k \geq 0}$ is a sequence of positive step-sizes such that $\sum \gamma_k = +\infty$ and $\sum \gamma_k^2 < +\infty$ (for example, $\gamma_k = 1/k$).

- **M-step:** determine θ_{k+1} by maximizing $\Psi(\theta) - \langle s_k, \Phi(\theta) \rangle$.

The convergence of this algorithm is proved under very general conditions, see Delyon et al. (1999) for more details.

3.2.3 Application to the Gaussian linear model

We suppose that the non observed data take their values in $\mathbb{R}^p$ and that the observed data take their values in $\mathbb{R}^n$. We consider the basic case where the function h is linear, so there exists some matrix A such that:

$$y = Ax + \epsilon.$$

We consider that x is a sequence of independent Gaussian random variables with mean μ and variance γ^2 :

$$x \sim N(\mu \mathbb{1}_p, \gamma^2 I_p),$$

where N indicates "normal distribution", $\mathbb{1}_p$ is the unit vector of $\mathbb{R}^p$ and I_p is the identity matrix of size p . It corresponds to the prior distribution π mentioned in assumption **(H1)**. The set of hyperparameters to estimate is $\theta = (\mu, \gamma^2, \sigma^2)$ and the joint likelihood of (x, y) has the form:

$$f(x, y; \theta) = \frac{1}{\sqrt{2\pi\gamma^2}^p} \exp\left(-\frac{\|x - \mu \mathbb{1}_p\|^2}{2\gamma^2}\right) * \frac{1}{\sqrt{2\pi\sigma^2}^n} \exp\left(-\frac{\|y - Ax\|^2}{2\sigma^2}\right).$$

If we develop this expression, we obtain for the log-likelihood:

$$\begin{aligned} \log f(x, y; \theta) = & -\frac{n+p}{2} \log(2\pi) - \frac{p}{2} \log(\gamma^2) - \frac{n}{2} \log(\sigma^2) \\ & - \frac{\mu^2 p}{2\gamma^2} - \left(-\frac{\mu}{\gamma^2} \sum_{i=1}^p x_i + \frac{1}{2\gamma^2} \sum_{i=1}^p x_i^2 + \frac{1}{2\sigma^2} \|y - Ax\|^2 \right). \end{aligned} \quad (3.3)$$

So we have an expression like the one presented in (3.1), with:

$$\begin{aligned} C &= -\frac{n+p}{2} \log(2\pi) \\ \Psi(\theta) &= -\frac{p}{2} \log(\gamma^2) - \frac{n}{2} \log(\sigma^2) - \frac{\mu^2 p}{2\gamma^2} \\ S(x, y) &= \left(\sum_{i=1}^p x_i, \sum_{i=1}^p x_i^2, \|y - Ax\|^2 \right) \\ \Phi(\theta) &= \left(-\frac{\mu}{\gamma^2}, \frac{1}{2\gamma^2}, \frac{1}{2\sigma^2} \right). \end{aligned}$$

Using the Bayes's formula and (3.3), we can easily compute the posterior distribution of x :

$$x|y; \theta \sim N(m(\theta), \Sigma(\theta)),$$

where

$$m(\theta) = \left(\frac{1}{\sigma^2} A^T A + \frac{1}{\gamma^2} I_p \right)^{-1} \left(\frac{1}{\sigma^2} A^T y + \frac{\mu}{\gamma^2} \mathbb{1}_p \right) \quad (3.4)$$

$$\Sigma(\theta) = \left(\frac{1}{\sigma^2} A^T A + \frac{1}{\gamma^2} I_p \right)^{-1}. \quad (3.5)$$

Then, iteration k of the SAEM algorithm reduces to the three following steps:

- **S-step:** using $\theta_k = (\mu_k, \gamma_k, \sigma_k^2)$, compute $m(\theta_k)$ and $\Sigma(\theta_k)$ from (3.4,3.5), and draw $x^{(k)}$ as a Gaussian random vector with mean $m(\theta_k)$ and variance $\Sigma(\theta_k)$.
- **A-step:** update the sufficient statistics $s_k = (s_{k,1}, s_{k,2}, s_{k,3})$ of the complete model as follows:

$$s_{k,1} = s_{k-1,1} + \gamma_k \left(\sum_{i=1}^p x_i^{(k)} - s_{k-1,1} \right) \quad (3.6)$$

$$s_{k,2} = s_{k-1,2} + \gamma_k \left(\sum_{i=1}^p \left(x_i^{(k)} \right)^2 - s_{k-1,2} \right) \quad (3.7)$$

$$s_{k,3} = s_{k-1,3} + \gamma_k \left(\|y - Ax^{(k)}\|^2 - s_{k-1,3} \right). \quad (3.8)$$

- **M-step:** compute $\theta_{k+1} = (\mu_{k+1}, \gamma_{k+1}, \sigma_{k+1}^2)$:

$$\mu_{k+1} = \frac{s_{k,1}}{p} \quad (3.9)$$

$$\gamma_{k+1}^2 = \frac{s_{k,2}}{p} - \mu_{k+1}^2 \quad (3.10)$$

$$\sigma_{k+1}^2 = \frac{s_{k,3}}{n}. \quad (3.11)$$

$$(3.12)$$

Initial values for θ_0 and s_0 have to be chosen (for example an initial model containing prior information for θ_0 and zero for s_0). This algorithm converges to a value $\theta_{ML} = (\mu_{ML}, \gamma_{ML}^2, \sigma_{ML}^2)$ that will be used for the estimation of the non observed data sequence x . In the Gaussian case, the posterior mean and the posterior mode coincide. Thus, we estimate x by the posterior mean $m(\theta_{ML})$, computed using (3.4) with θ_{ML} .

3.3 Application to the gravimetric and teleseismic problems

3.3.1 The gravimetric model

The intensity of the gravity field at the earth's surface is a function of the density distribution underneath the surface, which is approximated by horizontal layers, subdivided into rectangular blocks. There exists then a linear function F such that

$$g = F(\rho),$$

where g is a vector containing gravity values and ρ is a vector containing the densities of each block. This equation can be written in matrix form. Since our aim is not the reconstruction of absolute densities but of density variations, we subtract the average from the gravity data and denote the residual field by dg_{th} (theoretical values) and the density variations by $d\rho$. This results in the following equation, where A_g is the matrix $F'(\rho)$:

$$dg_{th} = A_g d\rho. \quad (3.13)$$

The measurement errors (instrument, reading etc.) in the observed data dg are modeled by Gaussian white noise ε_g , with zero mean and an unknown variance $\sigma_g^2 I_{n_g}$, where n_g is the size of the vector dg (number of measured gravity data). As a first approximation, we obtain therefore

$$dg = A_g d\rho + \varepsilon_g. \quad (3.14)$$

Since the random variables ε_g and $d\rho$ are independent, we can not consider that this white noise describes also the modeling errors which are not independent from the random variable $d\rho$. We do, however, not model those errors since this would complicate too much the entire inversion process and we restrict the errors to the measurement noise ε_g .

We assume that the prior distribution π of the vector $d\rho$ is a Gaussian distribution with an unknown mean μ_ρ and an unknown variance $\sigma_\rho^2 I_{n_\rho}$, where n_ρ denotes the size of the vector $d\rho$ (number of blocks for which the density variations are inverted):

$$d\rho \sim N(\mu_\rho \mathbb{1}_{n_\rho}, \sigma_\rho^2 I_{n_\rho}).$$

The SAEM algorithm presented in section 3.2.3 can be used for estimating $\theta_g = (\mu_\rho, \sigma_\rho^2, \sigma_g^2)$ and for recovering the unknown density sequence $d\rho$.

3.3.2 The teleseismic model

The travel time of a seismic ray is the sum of the travel times in the different blocks crossed. For simplicity we suppose that the velocity inside a block is constant and therefore the rays are straight. In this case, a linear relationship exists between the travel time within a block and the inverse of the velocity, called slowness. The travel time is the product of slowness and length of the ray within a block. We can therefore write the theoretical travel time calculation as a system of linear equations:

$$dt_{th} = A_t ds, \quad (3.15)$$

where dt_{th} and ds are respectively the delay times (times with respect to the average arrival times) and the slowness variations. The matrix A_t contains the length of each ray in the crossed blocks, corresponding to the derivative $\partial t_i / \partial s_j$.

Also in this case, the measurement errors are modeled by a white Gaussian noise ε_t , with zero mean and an unknown variance $\sigma_t^2 I_{n_t}$, where n_t is the size of the vector dt (number of measured arrival time data) so that

$$dt = A_t ds + \varepsilon_t. \quad (3.16)$$

Assuming again that the vector ds is Gaussian with an unknown mean μ_s and an unknown variance $\sigma_s^2 I_{n_s}$, where n_s is the size of the vector ds (number of blocks crossed by a sufficient number of rays). Exactly as above, we can use SAEM for estimating $\theta_t = (\mu_s, \sigma_s^2, \sigma_t^2)$ and for recovering the unknown sequence ds .

3.3.3 Joint inversion of gravimetric and teleseismic data

We have seen that SAEM can be used for estimating independently the hyperparameters of both models, and for recovering independently the two non observed sequences $d\rho$ and ds . If a prior information concerning the joint distribution of these two sequences is also available, we can hope to improve the inversion. Indeed, it is often assumed that there is a linear relationship between velocity and density variations. We assume that such a relationship exists and that it may vary with depth, e.g. due to temperature variations. We define a matrix b and a random vector η so that

$$ds = b d\rho + \eta, \quad (3.17)$$

where η is a white Gaussian noise with mean zero and an unknown covariance matrix Γ_η . We suppose further that the variance of η and the linear correlation coefficient between $d\rho$ and ds are both constant within each layer of the model. Then, both the covariance matrix Γ_η and the matrix b of the linear coefficients are diagonal and constant by blocks, with the same block structure corresponding to the layers. Note also that usually the relationship is formulated between the velocity and the density:

$$dv = B d\rho.$$

There is then a relationship between b and B such that

$$B = -bv^2,$$

where v is the mean velocity of a layer.

Let $\sigma_\eta^2(j)$ be the variance of η and $b(j)$ be the linear coefficient in layer j . Let further J be the number of layers. Then, the parameters of the model we want to estimate are $\theta = (\mu_\rho, \sigma_\rho^2, (b(j), \sigma_\eta^2(j))_{1 \leq j \leq J}, \sigma_g^2, \sigma_t^2)$.

For a given value of θ , the joint posterior distribution of $(d\rho, ds)$ is still Gaussian, the mean $m(\theta)$ and the variance $\Sigma(\theta)$ of this joint distribution are computed in the appendix. At iteration k , the three steps of SAEM become:

- **S-step:** draw $(d\rho^{(k)}, ds^{(k)})$ with a Gaussian distribution with mean $m(\theta_k)$ and variance $\Sigma(\theta_k)$.
- **A-step:** the additional observed data set and the linear relationship (3.17) between the non-observed data leads to additional components in the sufficient statistics. Therefore, one has to update $(s_{k,1}, s_{k,2}, s_{k,3}, s_{k,4}, u_{k,1}, u_{k,2}, u_{k,3})$ as follows:

$$\begin{aligned} s_{k,1} &= s_{k-1,1} + \gamma_k \left(\sum_{i=1}^p d\rho_i^{(k)} - s_{k-1,1} \right) \\ s_{k,2} &= s_{k-1,2} + \gamma_k \left(\sum_{i=1}^p \left(d\rho_i^{(k)} \right)^2 - s_{k-1,2} \right) \\ s_{k,3} &= s_{k-1,3} + \gamma_k \left(\|dg - A_g d\rho^{(k)}\|^2 - s_{k-1,3} \right) \\ s_{k,4} &= s_{k-1,4} + \gamma_k \left(\|dt - A_t ds^{(k)}\|^2 - s_{k-1,4} \right). \end{aligned}$$

For each layer j , let C_j be the set of blocks that belong to the layer j and that are crossed by a sufficient number of rays. Then, $(u_{k,1}(j), u_{k,2}(j), u_{k,3}(j))$ are defined as follows:

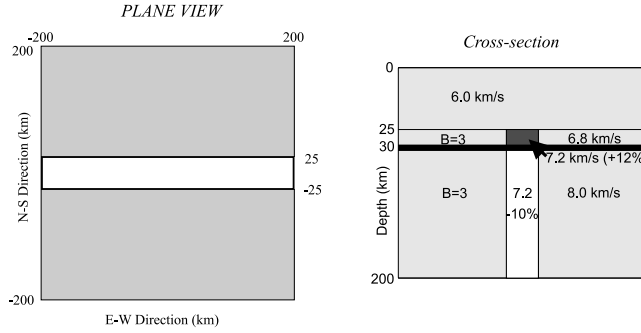
$$\begin{aligned} u_{k,1}(j) &= u_{k-1,1}(j) + \gamma_k \left(\sum_{i \in C_j} d\rho_i^{(k)} ds_i^{(k)} - u_{k-1,1}(j) \right) \\ u_{k,2}(j) &= u_{k-1,2}(j) + \gamma_k \left(\sum_{i \in C_j} \left(d\rho_i^{(k)} \right)^2 - u_{k-1,2}(j) \right) \\ u_{k,3}(j) &= u_{k-1,3}(j) + \gamma_k \left(\sum_{i \in C_j} \left(ds_i^{(k)} \right)^2 - u_{k-1,3}(j) \right). \end{aligned}$$

- **M-step:** compute θ_{k+1} according to

$$\begin{aligned} \mu_{\rho_{k+1}} &= \frac{s_{k,1}}{n_\rho}, \quad \sigma_{\rho_{k+1}}^2 = \frac{s_{k,2}}{n_\rho} - \mu_{\rho_{k+1}}^2 \\ \sigma_{g_{k+1}}^2 &= \frac{s_{k,3}}{n_g}, \quad \sigma_{t_{k+1}}^2 = \frac{s_{k,4}}{n_t} \\ b_{k+1}(j) &= \frac{u_{k,1}(j)}{u_{k,2}(j)}, \quad \sigma_{\eta_{k+1}}^2(j) = \frac{1}{n_j} \left(u_{k,3}(j) - \frac{u_{k,1}^2(j)}{u_{k,2}(j)} \right), \end{aligned}$$

where n_j is the number of blocks in layer j which are crossed by a sufficient number of rays.

FIG. 3.1: Schema of the model.



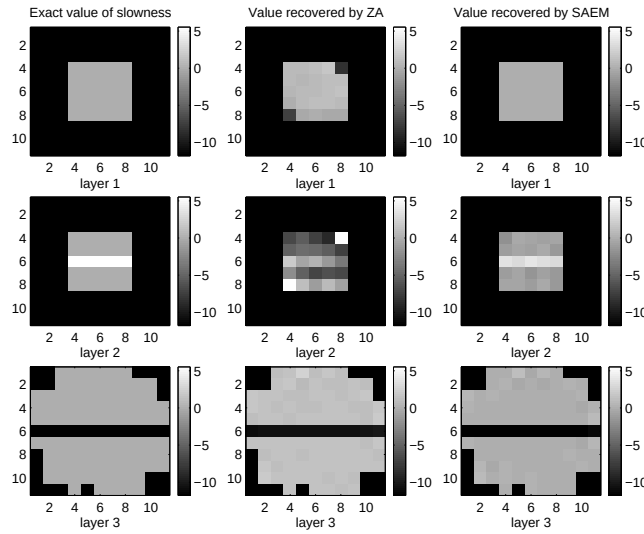
3.4 Synthetic model examples

For the test calculations, we used a model which is inspired from the results of teleseismic tomography in active rift, specifically the East-African Rift in Kenya (Achauer and the Krisp Teleseismic Working Group (1994)). A more or less undisturbed but thinned crust overlies a mantle lithosphere with a strong negative velocity anomaly due to ascending hot material. The anomalous structure is simulated by a 50 km wide vertical stripe in the center of the model running in EW direction (see Figure 3.1). The model has three distinct horizontal layers: a 25 km thick crust without velocity and density anomalies is situated above a 5 km thick layer with a strong positive velocity and density anomaly in the model center, representing the area where thinned crust is replaced by fast and dense mantle material. The velocity anomaly in this layer has an amplitude of +0.4 km/s (6% of the background velocity of 6.8 km/s), the density anomaly 133 kg/m^3 (4.6% of the background density of 2900 kg/m^3), resulting in a correlation factor B of $3 \text{ m}^4/(\text{s} \cdot \text{kg})$. The third layer has a thickness of 170 km and simulates a low velocity - low density channel cutting vertically through the whole lithospheric mantle. This last layer is treated in different examples as one single layer or subdivided into 4 layers. The velocity anomaly has been fixed to -0.8 km/s (10% of the background velocity of 8.0 km/s), whereas the density anomalies vary in the different models between -266 and -133 kg/m^3 so that the correlation factor B varies between 3 and 6.

This model is the same as the one used in Zeyen and Achauer (1997), called "ZA" in what follows, so that we can easily compare the results. It is on the one hand a geologically realistic model in terms of anomaly amplitudes and distributions, on the other hand, it is also a kind of worst case model for the pure velocity inversion, since it is very prone to be affected by vertical smearing. Especially the thin high-velocity layer underlying the rifted crust will not be resolved with existing methods as we will show in the following parts.

The synthetic data (delay times and Bouguer anomaly) were produced with the same software used for the forward calculation during the inversion. The seismic delay time calculation is based on the algorithm published by Evans and Achauer (1993) using a simplified ray tracing in which refraction is controlled only by the background velocity and changes in the angle of incidence due to velocity variations within a layer are neglected. The data simulation was done with a realistic event and station distribution on an area of $150 \times 150 \text{ km}^2$, i.e. not all receivers "worked" for all events and the event distribution was not homogeneous for all azimuths. In this way, we produced 1081 delay times. The Bouguer anomaly calculation is based on Cady (1980) and

FIG. 3.2: Velocity variations in % for single inversion of three layers model.



Zeyen and Pous (1993), simulating data on a regular grid with a 10 km spacing for an area of $300 \times 300 \text{ km}^2$, resulting in 961 data points. We added to the synthetic data sets a Gaussian error with standard deviations of 0.03 s for the delay times (corresponding to 3% of the data standard deviation of 0.80 s) and of 2 mGal for the gravity anomalies (corresponding to 3% of the data standard deviation of 71.2 mGal).

We will now present inversion results for the four following situations, which increase in complexity: a) single inversion of seismic data to compare standard Bayesian inversion (ZA) with the new method; b) joint inversion with three layers (the thick mantle layer is implemented as one single layer) with a favorable relation between the number of data and the number of parameters; c) a 6-layers model, subdividing the mantle into 4 layers with different correlation factors in the different mantle layers; d) a 3-layers model with areas where the assumed linear relationship between velocity and density variations does not apply in order to investigate the behavior of the new algorithm in this very unfavorable situation.

a) Single inversion of seismic delay times.

A first simulation was done using only the above described synthetic delay times. We used three layers with lower limits at 25, 30 and 200 km, each layer consisting of 11×11 blocks with $50 \times 50 \text{ km}$ side length (total area of $550 \times 550 \text{ km}^2$). The area is much wider than the actually available data at the surface for two reasons: on the one hand, we avoid edge effects for the gravity calculation, on the other hand, the seismic rays form a cone which in the lower most layer is about 550 km wide. Taking into account that not all blocks in this model are crossed by seismic rays, the problem consists of inverting 151 velocity (or slowness) variations using 1081 delay time data.

Figure 3.2 shows the results for Bayesian (ZA) and SAEM inversion. The Bayesian inversion algorithm results in smearing. The velocity anomaly in the second layer has disappeared. As shown in Zeyen and Achauer (1997) the anomalies in the upper two layers even changed sign for the ACH algorithm. In contrast, the new method reduces smearing considerably.

b) Joint inversion of a three layers model.

In a second example, we took the same three layers model, using delay times and Bouguer

TAB. 3.1: Estimation provided by the SAEM algorithm after 50 iterations for joint inversion of three layers model.

Parameters	Estimation	Parameters	Estimation
μ_ρ	-17	σ_ρ^2	4.7e^3
σ_g^2	4.5	σ_t^2	7.0e^{-4}
B_1	0.5	$\sigma_{\eta_1}^2$	9.4e^{-9}
B_2	2.1	$\sigma_{\eta_2}^2$	1.2e^{-6}
B_3	3.6	$\sigma_{\eta_3}^2$	3.6e^{-7}

anomaly data. The B parameters in the synthetic model were fixed to $3 \text{ m}^4/(\text{kg}\cdot\text{s})$ in all layers. This results in an inversion scheme in which 514 model parameters (363 density variations and 151 velocity variations) are to be determined using 2042 simulated data (961 gravity anomalies and 1081 delay times). This problem is less well constrained than the first one, but with a ratio of 4:1 between data and model parameters, it is still reasonable.

The ZA algorithm resolves the high velocity region in the second layer (Figure 3.4, second column), although the amplitude is too small, but the high density in this layer are not well resolved (Figure 3.3, second column). The overall variance reduction is better than 95% for both data sets. Also the B parameters are reasonably well resolved (0.7 for layer 1 (i.e. near zero as should be expected) 3.3 for layer 2 and 4.4 for layer 3).

The SAEM algorithm (Figures 3.3 and 3.4, right column) gives much better results. Also the amplitude of the velocity and density anomalies in the second and third layer are well resolved, the B values become 0.5, 2.1 and 3.6. Figure 3.5 shows the convergence of the maximum likelihood estimator of the hyperparameters for the SAEM algorithm. Convergence is reached for all estimators latest after 20 iterations.

c) Joint inversion of a six layers model.

A subdivision of the thick mantle layer into 4 separate layers is geologically more meaningful, but it poses additional problems to the inversion algorithm: in comparison to the number of observed data there is now an elevated number of non observed data to be estimated. This leads for the ZA algorithm to stability problems. As can be seen in Figures 3.6 and 3.7 (second column), there is a strong scattering in the resulting density and velocity variations and especially the velocities are not well resolved in the deeper layers. Consequently, also the B values are not correctly resolved at depth. The SAEM algorithm, in contrast, recovers both sets of non observed data quite well (Figures 3.6 and 3.7, right column). The correlation coefficients, however, do not show the depth variation introduced in the synthetic model (the synthetic B changes linearly from 3 in layer 3 to 6 in layer 6), since the inverted densities stay too high in the deep layers, where gravity resolution decreases due to the effect of distance. The convergence of the estimator of the hyperparameters is of the same order as this one obtained in the three layers model.

d) Joint inversion with areas where the linear relationship is disturbed.

In a fourth model, we changed the density distribution in the first layer so that the southern half has a positive density contrast (increased by 50 kg/m^3) without changing the seismic velocities. Therefore, in this layer, there is no linear relationship between velocity and density contrasts, which might geologically correspond to the juxtaposition of two different terrains. In this case, the ZA algorithm behaves well for the mantle layer, but does not recover very well the upper two layers (see Zeyen and Achauer (1997)). The SAEM algorithm, however, deals much better with

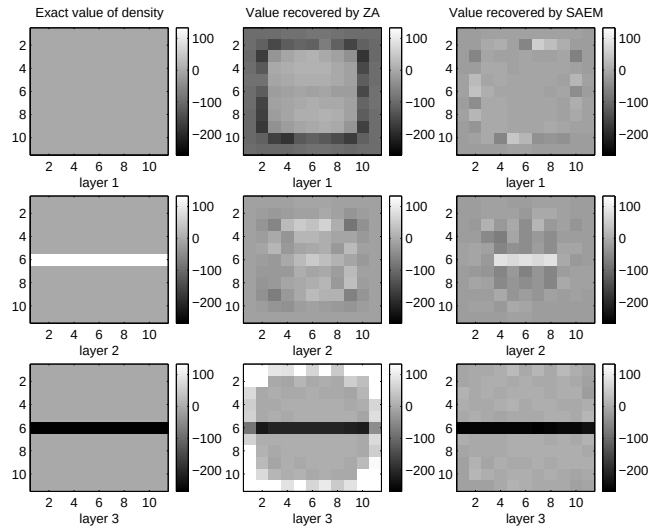
FIG. 3.3: Density for joint inversion of three layers model (in kg/m^3).

FIG. 3.4: Velocity variations in % for joint inversion of three layers model.

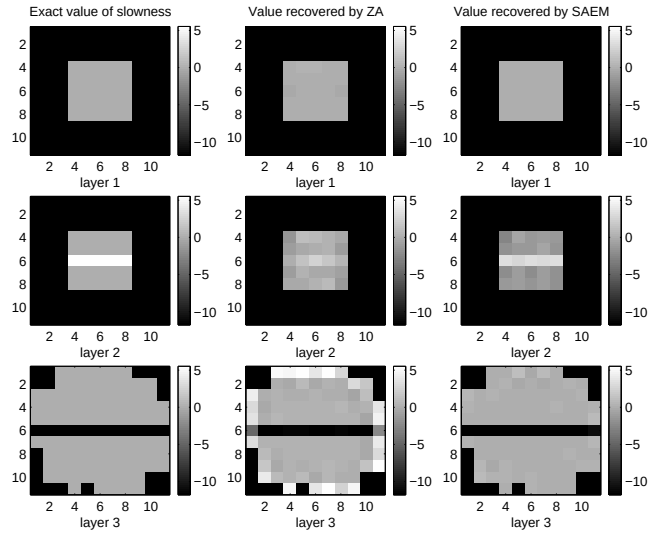
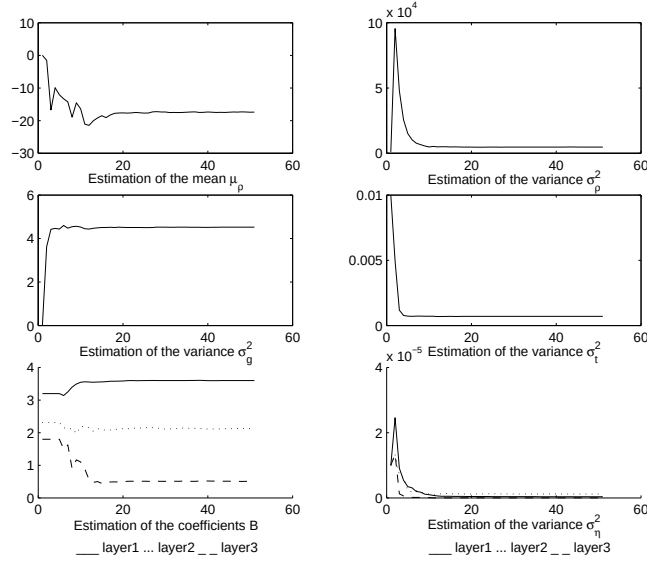


FIG. 3.5: Evolution of the hyperparameters for joint inversion of three layers model during the iteration process. Horizontal axes: iteration number; vertical axes: estimation of the different hyperparameters.



TAB. 3.2: Estimation provided by the SAEM algorithm after 50 iterations for joint inversion of six layers model.

Parameters	Estimation	Parameters	Estimation
μ_ρ	-15	σ_ρ^2	6.1e^3
σ_g^2	3.4	σ_t^2	6.6e^{-4}
B_1	0.1	$\sigma_{\eta_1}^2$	2.3e^{-8}
B_2	2.0	$\sigma_{\eta_2}^2$	6.9e^{-7}
B_3	3.8	$\sigma_{\eta_3}^2$	2.2e^{-7}
B_4	3.6	$\sigma_{\eta_4}^2$	5.7e^{-6}
B_5	2.9	$\sigma_{\eta_5}^2$	8.9e^{-6}
B_6	3.5	$\sigma_{\eta_6}^2$	7.1e^{-6}

FIG. 3.6: Density for joint inversion of six layers model.

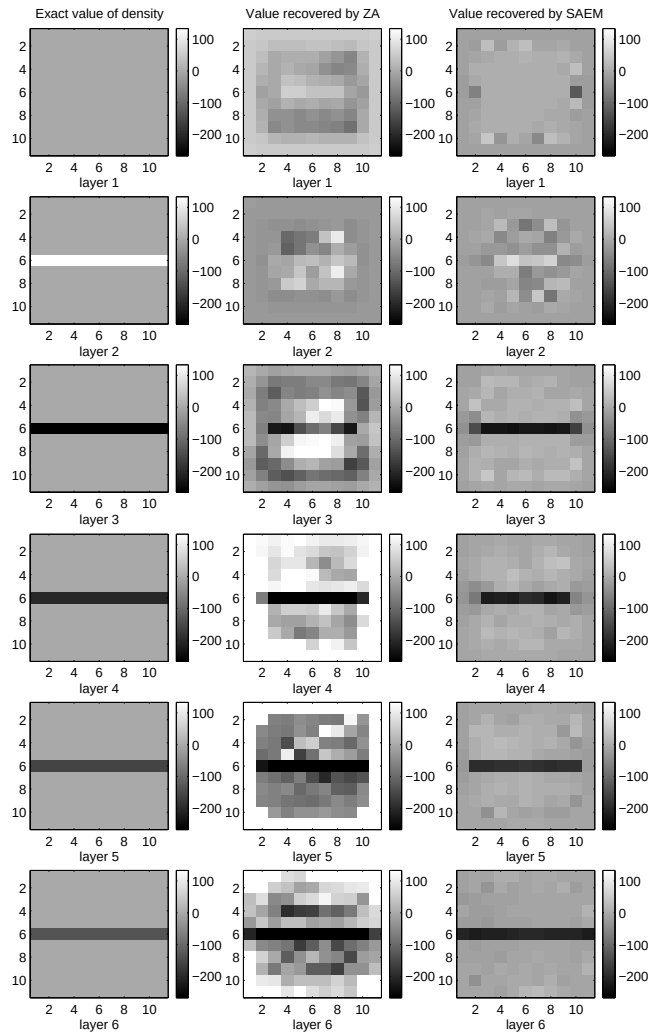


FIG. 3.7: Velocity variations in % for joint inversion of six layers model.

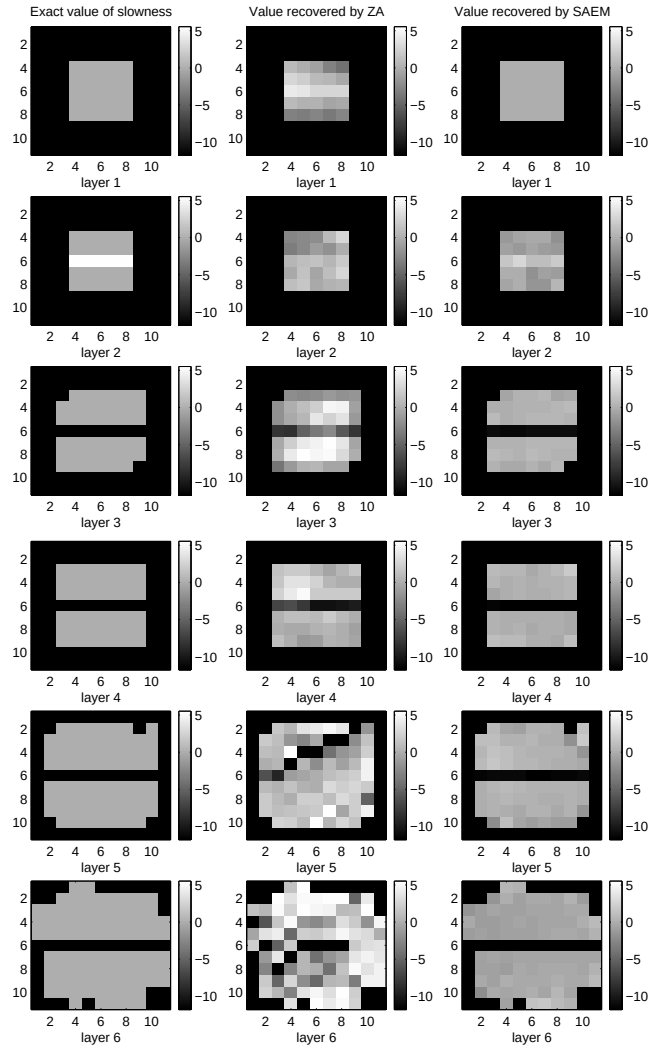
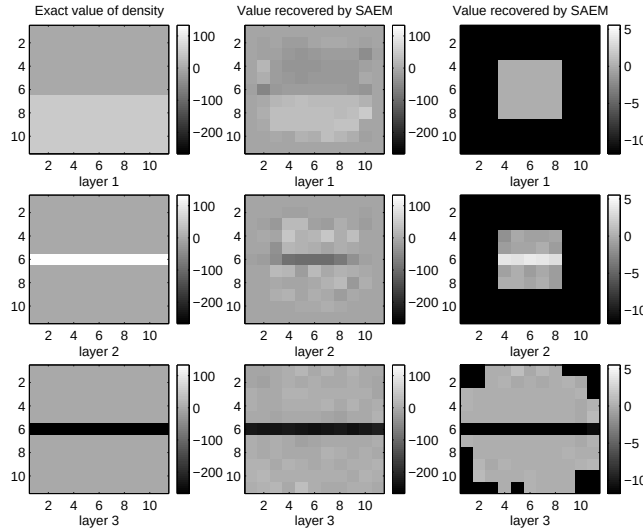


FIG. 3.8: Results for joint inversion of the fourth model.



TAB. 3.3: Estimation provided by the SAEM algorithm after 50 iterations for joint inversion for the fourth model.

Parameters	Estimation	Parameters	Estimation
μ_ρ	-9.8	σ_ρ^2	$3.7e^3$
σ_g^2	6.1	σ_t^2	$7.0e^{-4}$
B_1	-0.12	$\sigma_{\eta_1}^2$	$1.9e^{-8}$
B_2	-1.8	$\sigma_{\eta_2}^2$	$3.3e^{-6}$
B_3	4.0	$\sigma_{\eta_3}^2$	$7.5e^{-7}$

this very unfavorable situation and gives correct results for the velocity variations, although the density variations of the second layer seem to be masked by those of the first layer (Figure 3.8). Note that in Figure 3.8, we did not show the correct velocity model, which is the same as the one of the first two models, however we showed the modified density model. The convergence behavior is not shown for this model, since it is very similar to the one presented Figure 3.5 for the second model.

3.5 Discussion

The presented algorithm is based on a pseudo 3D ray tracer, i.e. the rays do not bend when crossing the limits between blocks with different velocities within the same layer. Instead, the rays are refracted at layer limits with angles corresponding to the average velocities of the layers. If strong lateral velocity variations are present, a full 3D ray tracing is necessary in order to recover correctly the velocity variations at depth. Our algorithm can easily be extended to use such a 3D ray-tracer, however the calculation becomes much more time consuming, since the problem is then strongly non-linear and is based on an iterative procedure in which the Jacobian have to be completely updated after every iteration. A possible way of tackling the full

problem might be to proceed in two steps: first use the presented pseudo-3D algorithm in order to construct a starting model for the full 3D inversion and second use the modeled statistical parameters and velocity-density relationships as constraints for a full 3D inversion.

Another problem which we did not tackle at this stage is the way to define layer thicknesses and block sizes. In the future, some optimization process (see e.g. Kissling, Husen, and Haslinger (2001)) should be integrated into our algorithm in order to increase block sizes where little information (gravimetric and seismic) exists and densify it in areas densely covered by rays or gravity measurements. It could also be considered to use different block sizes for gravity and delay time data inversion in order to account for the decreasing resolution of density distribution with depth.

We have however shown with a series of increasingly complex models that the new method is superior to some similar existent teleseismic delay time inversion methods. Its principal advantage is that the statistical parameters of the model do not have to be chosen by the operator but are treated as unknown. This ensures an objective and optimal result. It also seems to stabilize the inversion in case of unfavorable conditions like a small ratio between the number of observed and non observed data. Especially in the case of joint inversion of gravity and teleseismic delay time data, the automatic determination of the statistical parameters is important, since the different data sets and even more the different parameter sets are characterized by different variances which are generally difficult to estimate a priori.

The method can easily integrate any kind of prior information. For example, if we have some knowledge on the density inside some blocks, we can set the variance of $d\rho$ to $\alpha\sigma_\rho^2$ with α fixed to some arbitrarily small value, and with σ_ρ^2 estimated as above.

3.6 Appendix

Determination of the mean $m(\theta)$ and the covariance matrix $\Sigma(\theta)$ of the joint posterior Gaussian distribution of $(d\rho, ds)$:

For each layer j , let C_j be the set of blocks that belong to the layer j and n_j the number of blocks in layer j . Using Bayes's formula, we get:

$$\begin{aligned} p(d\rho, ds \mid dg, dt; \theta) &\propto f(dg, dt \mid d\rho, ds; \theta) f(d\rho, ds; \theta) \\ &\propto f(dg \mid d\rho; \theta) f(dt \mid ds; \theta) f(ds \mid d\rho; \theta) f(d\rho; \theta). \end{aligned}$$

Thus,

$$\begin{aligned} \log p(d\rho, ds \mid dg, dt; \theta) = C - \frac{1}{2} &\left[n_g \log(\sigma_g^2) + n_t \log(\sigma_t^2) + n_\rho \log(\sigma_\rho^2) + \sum_{j=1}^J n_j \log(\sigma_\eta^2(j)) \right. \\ &\left. + \frac{\|dg - A_g d\rho\|^2}{\sigma_g^2} + \frac{\|dt - A_t ds\|^2}{\sigma_t^2} + \sum_{j=1}^J \sum_{i \in C_j} \frac{(ds_i - b(j)d\rho_i)^2}{\sigma_\eta^2(j)} + \frac{\|d\rho - \mu_\rho\|^2}{2\sigma_\rho^2} \right]. \end{aligned}$$

We see that this posterior distribution is Gaussian. The mean $m(\theta)$ and the covariance $\Sigma(\theta)$ are computed by identification. So we only need to analyze the part of the log-likelihood involving

the vector $(d\rho, ds)$:

$$\begin{aligned} \log p(d\rho, ds \mid dg, dt; \theta) = & \tilde{C}(dg, dt; \theta) - \frac{1}{2} \left[\frac{\|dg - A_g d\rho\|^2}{\sigma_g^2} + \frac{\|dt - A_t ds\|^2}{\sigma_t^2} \right. \\ & \left. + \sum_{j=1}^J \sum_{i \in C_j} \frac{(ds_i - b(j)d\rho_i)^2}{\sigma_\eta^2(j)} + \frac{\|d\rho - \mu_\rho\|^2}{2\sigma_\rho^2} \right]. \end{aligned}$$

Let $Z = (d\rho, ds)$, after some straightforward algebra, we find that

$$\log p(Z \mid dg, dt; \theta) = C(dg, dt; \theta) - \frac{1}{2} [(Z - m(\theta)) \Sigma^{-1}(\theta) (Z - m(\theta))'] ,$$

where

$$\begin{aligned} m(\theta) &= \begin{pmatrix} S \\ T \end{pmatrix} \\ \Sigma(\theta) &= \begin{pmatrix} P & Q \\ Q' & R \end{pmatrix}^{-1}, \end{aligned}$$

with

$$\begin{aligned} P &= \frac{1}{\sigma_\rho^2} I_{n_\rho} + B' \Gamma_\eta^{-1} B + \frac{1}{\sigma_1^2} A_1' A_1 \\ Q &= -B' \Gamma_\eta^{-1} \\ R &= \Gamma_\eta^{-1} + \frac{1}{\sigma_2^2} A_2' A_2 \\ S &= (P - QR^{-1}Q')^{-1} \left[\frac{\mu_\rho}{\sigma_\rho^2} + \frac{1}{\sigma_{A_1' dg}^2} - \frac{QR^{-1}A_2' dt}{\sigma_2^2} \right] \\ T &= R^{-1} \left(\frac{A_2' dt}{\sigma_2^2} - Q'S \right). \end{aligned}$$

□

Chapitre 4

Logsplines density estimation by maximum likelihood in nonlinear inverse problem

Summary

In this chapter, we propose a method for estimating in a non parametric way the density of the missing data in some general inverse problems. We consider a logspline model of a fixed size and the maximum likelihood estimator of this density in this model. So we get a logspline density estimator, which can be approach using the SAEM algorithm coupled with the MCMC method. We also study the convergence of this logspline estimator when the size of the logspline model and the number of observations go to infinity. We illustrate the method on some simulated examples and real data sets.

4.1 Introduction

Let Z be a random variable with density function π having for support a compact interval $\mathcal{I}$ of $\mathbb{R}$. We consider a random sample $z_1, \dots, z_n$ drawn from the same distribution than Z . Instead of observing $z_1, \dots, z_n$, we observe independent random variables $y_1, \dots, y_n$, not necessary identically distributed, taking value in a interval $\mathcal{Y}$ of $\mathbb{R}$, which are related to the unobserved random sample $z_1, \dots, z_n$ through the density functions $(h_i)_{1 \leq i \leq n}$ of $y_1, \dots, y_n$ conditionally to $z_1, \dots, z_n$. Our purpose in this chapter is to estimate by maximum likelihood the unknown density function π from the observations $(y_i)_{1 \leq i \leq n}$ in a non parametric way.

Such indirect problems are known as statistical inverse problems. Since the purpose is to estimate the density function π of Z when the density functions of $y_1, \dots, y_n$ and the density functions of $y_1, \dots, y_n$ conditionally to $z_1, \dots, z_n$ are known, the problem is often difficult to solve. In fact, the solution may be large disturb by a small change of the density function of $(y_1, \dots, y_n)$. This denotes the ill-posed feature of such an inverse problem.

Nevertheless, statistical inverse problem have been often study since they have lots of applications in very different fields, for example gravimetry, teleseismic, pharmacology and signal processing. Some application of a smooth version of the EM algorithm was proposed for density estimation by Silverman, Jones, Nychka, and Wilson (1990), based on the article of Vardi, Shepp, and Kaufman (1985) specific for the positron emission tomography. Schumitzky (1991) presents another EM-like algorithm considering the mean of the posterior densities. More recently, Eggermont and LaRiccia (1995) et Eggermont (1999) have proposed for inverse problem with integral operator a smooth version of EM using a nonlinear smoothing operator involving kernel and obtain an estimator of π which converges in the L_1 -sense when n goes to infinity. Another approach consists in using the standard value decomposition (*SVD*). We can cite here for example the recent works of Cavalier and Tsybakov (2002) and Cavalier, Golubev, Picard, and Tsybakov (2002). Some approaches using logspline were also proposed. Koo and Park (1996) use logspline model coupled with the EM algorithm for estimating the density in deconvolution, whereas Koo and Chung (1998) present an *SVD* type approach. In his article of 1999, Koo proposed in case of deconvolution (i.e. $y = z + \varepsilon$ where ε is independent from z and drawn from a known distribution) a *logspline density estimator* of the density function π of z . He considers the vector space $\mathcal{S}$ of spline functions of a given positive order q on $\mathcal{I}$, namely piecewise polynomials function of degree $q - 1$. Given a subdivision of $\mathcal{I}$, the space $\mathcal{S}$ has a finite dimension J and admits a *B-spline* basis denoted $B_1, \dots, B_J$ (see de Boor (1978)), having the following properties: the $(B_j)_{1 \leq j \leq J}$ are nonnegative and their sum is equal to 1 on $\mathcal{I}$. So for $\theta = (\theta_1, \dots, \theta_J)$ in $\mathbb{R}^J$, he defines the density function π_θ for all z in $\mathcal{I}$ by:

$$\pi_\theta(z) = \exp \left[\sum_{j=1}^J \theta_j B_j(z) - c(\theta) \right] \quad , \quad \text{where} \quad c(\theta) = \log \left(\int_{\mathcal{I}} \exp \left[\sum_{j=1}^J \theta_j B_j(z) \right] dz \right). \quad (4.1)$$

This definition of π_θ ensures that it is a nonnegative function which integrates to one. Koo considers the estimator $\pi_{\hat{\theta}_{ind}}$ of the density function π , where $\hat{\theta}_{ind}$ maximizes in θ the indirect log-likelihood defined as follow:

$$l_{ind}(\theta) = \sum_{j=1}^J \theta_j \hat{B}_j - c(\theta), \quad (4.2)$$

where the sequence $(\hat{B}_j)_{1 \leq j \leq J}$ is obtained from the sequence $(B_j)_{1 \leq j \leq J}$, using the Fourier inversion formula such that the $(\hat{B}_j)$ are appropriated estimators of the (B_j) . He established convergence rates of its estimator when the log-density function $\log \pi$ of Z is assumed to be in a Besov space.

Since the parametric logspline model defined in (4.1) ensures a positive estimator which integrates to one, it is well adapt for the density estimation. So we kept this logspline model to estimate π , but since the method of Koo (1999) is specifically adapt to the deconvolution problem, we will propose another estimator, which is valid in some general nonlinear inverse problems.

The classical parametric maximum likelihood approach would consist in choosing in a fixed parametric model $\{\pi_\theta, \theta \in \Theta\}$ the value of θ in Θ which maximizes the complete log-likelihood of the model defined by:

$$l_n^{com}(\theta) = \frac{1}{n} \sum_{i=1}^n \log[h_i(y_i|z_i)\pi_\theta(z_i)].$$

In cases where the random sample $z_1, \dots, z_n$ is non observed, we consider the value $\hat{\theta}$ which maximizes the observed log-likelihood of the model defined by:

$$l_n(\theta) = \frac{1}{n} \sum_{i=1}^n \log \int_{\mathcal{I}} h_i(y_i|z)\pi_\theta(z)dz.$$

So for a fixed dimension J , the logspline model is a parametric one and we obtain a parametric maximum likelihood estimator of θ depending on n and on J , involving a parametric maximum likelihood estimator of π . From a practical point of view, we approach this estimator $\hat{\theta}_{n,J}$ of θ in the logspline model of size J using the version of the SAEM algorithm presented in the Chapter 1 of this thesis. Moreover, this algorithm is well adapt to logspline model, since it is an exponential families. So under some regularity assumptions of the model, the algorithm converges to a local maximizer of the observed log-likelihood l_n .

We consider the asymptotical properties of the logspline estimator $\pi_{\hat{\theta}_{n,J}}$ of the density π when the number of observations n and the size J of the logspline model tend to infinity. We introduce, assuming its existence, $\theta_{n,J}^*$ which maximizes $E[l_n(\theta)]$ on $\mathbb{R}^J$ and prove, under regularity assumptions, the convergence in probability to zero of $\|\log \pi_{\hat{\theta}_{n,J}} - \log \pi_{\theta_{n,J}^*}\|_2$ when n goes to infinity and also J , but not too fast with n . We consider now the approximation property of $\pi_{\theta_{n,J}^*}$ for π . We prove the convergence in distribution in some particular case.

To illustrate the method, we present some applications to simulated data. The missing data are drawn from a Gaussian distribution or a mixture of two Gaussian distributions. We estimate the density of the missing data from the observations and present also examples where other parameters of the model, such as the variance of the error, are estimated. To end this chapter, we consider again the real data set of the orange trees presented in section 2.3. We estimate the density of the random effect ϕ_i in a non parametric way using the logspline estimate. In the same time, the others population parameters are also estimated.

4.2 Logspline model

We define now precisely the logspline model which will be used. We are interested in piecewise polynomials functions of positive order $q \geq 1$ (i.e. degree $q - 1$) on $\mathcal{I}$. So let $\mathcal{I}$ be equal to $[a, b]$

where $-\infty < a < b < +\infty$ and consider a given knots sequence $\tau = (t_l)_{1 \leq l \leq K+1}$ with $a = t_1$ and $b = t_{K+1}$. Consider now the vector space $\mathcal{S}_{q,\tau}$ of spline functions of positive order q on $\mathcal{I}$, namely piecewise polynomials function of degree $q - 1$ associated to this knots sequence. Then the dimension of $\mathcal{S}_{q,\tau}$ is equal to $J = q + K - 1$ and there exists a B-splines basis denoted $B_1, \dots, B_J$ for $\mathcal{S}_{q,\tau}$ (see de Boor (1978)).

So for a given vector $\theta = (\theta_1, \dots, \theta_J)$ in $\mathbb{R}^J$, we denote for z in $\mathcal{I}$

$$s(z; \theta) = \sum_{j=1}^J \theta_j B_j(z) \quad \text{and} \quad c(\theta) = \log \left[\int_{\mathcal{I}} \exp[s(z; \theta)] dz \right],$$

so that the function π_θ defined as follow

$$\pi_\theta(z) = \exp[s(z; \theta) - c(\theta)] \quad (4.3)$$

is a density function for all θ in $\mathbb{R}^J$. This family is not identifiable since we have for all a real:

$$c(\theta + a) = c(\theta) + a \quad \text{implying that} \quad \pi_{\theta+a} = \pi_\theta.$$

So we set systematically $\theta_J = 0$ in order to get an identifiable family of log-density functions and we denote Θ_J the subspace of $\mathbb{R}^J$ composed of vectors having zero as last coordinate and $\mathcal{M}_{q,\tau}$ the set of densities associated equal to $\{\pi_\theta, \theta \in \Theta_J\}$. From now on, we will consider only subdivisions of $\mathcal{I}$ with subintervals of equal length. We denote the dependence in q and K through J for simplicity. We describe briefly some properties of the B-splines detailed in de Boor (1978):

Proposition 4.1. *For all $1 \leq j \leq J$, the function B_j takes values in the interval $[0, 1]$. Moreover, we have $\sum_{j=1}^J B_j(z) = 1$ for all z in $\mathcal{I}$. By the way, there exists a positive constant D_q independent of the knots sequence, such that for all θ in $\mathbb{R}^J$, the following upper bound holds:*

$$\|\theta\|_\infty \leq D_q \|\log \pi_\theta\|_\infty.$$

We present now the approximation property of the logspline model. For some positive continuous density function f on $\mathcal{I}$, we define $\delta_J = \inf_{\theta \in \Theta_J} \|\log f - \log \pi_\theta\|_\infty$. Then δ_J tends to zero when J goes to infinity (see Stone (1990) for the application to logspline and de Boor (1978) for more details on the links between the convergence rate and the regularity of f).

We define now the observed likelihood corresponding to the logspline model:

$$l_n(\theta) = \frac{1}{n} \sum_{i=1}^n \log \int_{\mathcal{I}} h_i(y_i | z) \pi_\theta(z) dz$$

and the maximum likelihood estimator $\pi_{\hat{\theta}_{n,J}}$ of the density π in this logspline model by:

$$\hat{\theta}_{n,J} = \arg \max_{\theta \in \Theta_J} l_n(\theta).$$

The particular properties of the logspline model let us think that $\pi_{\hat{\theta}_{n,J}}$ will have remarkable properties when n and J tend to infinity. In a first time, we explain how we compute this estimator in practice.

4.3 The SAEM algorithm

For fixed integers n and J , we consider the logspline model $\mathcal{M}_J = \{\pi_\theta, \theta \in \Theta_J\}$ defined above. For all $1 \leq i \leq n$, we denote h_i the density functions of y_i conditionally to z_i . Then the complete log-likelihood corresponding to the logspline model has the following expression:

$$l_n^{com}(\theta) = \frac{1}{n} \sum_{i=1}^n \log h_i(y_i|z_i) + \frac{1}{n} \sum_{i=1}^n \log \pi_\theta(z_i). \quad (4.4)$$

So we will apply the SAEM algorithm to this parametric model in order to approach the estimator $\hat{\theta}_{n,J}$ of θ , that maximizes the observed log-likelihood l_n . To put out the minimal sufficient statistics of the model, we write the developed expression of the complete log-likelihood:

$$l_n^{com}(\theta) = \frac{1}{n} \sum_{i=1}^n \log h_i(y_i|z_i) + \sum_{j=1}^J \theta_j \left[\frac{1}{n} \sum_{i=1}^n B_j(z_i) \right] - c(\theta). \quad (4.5)$$

So we choose as minimal sufficient statistics $\tilde{S}(z) = (\frac{1}{n} \sum_{i=1}^n B_j(z_i), 1 \leq j \leq J)$ and implement the k -th iteration of the SAEM algorithm as follow:

- **S-step:** generate a realization z' using as proposal distribution the prior distribution π_{θ_k} and take z_k equal to z' or to z_{k-1} according to the value of the acceptance probability defined in (1.1).
- **A-step:** update s_{k-1} according to (7).
- **M-step:** update θ_k according to (8).

In this particular case of logspline model, the assumptions needed for the convergence theorem are quite simplified: assumptions **(M1)**-**(M4)** are always satisfied, as well as assumption **(C)** since the minimal sufficient statistic belongs to $[0, 1]^J$. Assumption **(SAEM2)** can be weakened in **(SAEM2')**: the function $\hat{\theta}$ is m -times differentiable, since the function l satisfied always this condition. So we obtain the following result:

Theorem 4.2. *Assume the following assumptions **(M5)**, **(SAEM1')**, **(SAEM2')** and **(SAEM3')** are satisfied. Then, w.p. 1, $\lim_{k \rightarrow +\infty} d(\theta_k, \mathcal{L}) = 0$ where $d(x, A)$ denotes the distance of x to the closed subset A and $\mathcal{L} = \{\theta \in \Theta, \partial_\theta l_n(y; \theta) = 0\}$ is the set of stationary points of l_n .*

Remark. *Concerning the assumption **(M5)**, if the existence of $\hat{\theta}$ is ensured, we systematically have the uniqueness of $\hat{\theta}$, because since the function $\theta \rightarrow c(\theta)$ is convex (see Stone (1990)), the function $\theta \rightarrow \log \pi_\theta$ is concave.*

The theorem ensuring the convergence to a local maximizer of the log-likelihood of the observations can also be applied. Moreover, the assumptions are simplified, since the convexity of $\theta \rightarrow c(\theta)$ implies that the second matrix considered in assumption **(LOC2)** is always positive definite. By the way, the expression of the minimal sufficient statistic involved by the logspline model implies that if assumption **(LOC3)** is satisfied, than the first matrix considered in assumption **(LOC2)** is also always positive definite. We obtain the following particular result:

Theorem 4.3. *Assume the following assumptions **(M5)**, **(SAEM1')**-**(SAEM3')**, **(LOC1)** and **(LOC3)** are satisfied. Then, w.p. 1, the sequence (θ_k) converges to a local maximizer of the observed log-likelihood l_n .*

Remark. *We can also admit some parametric dependence for the conditional densities (h_i) and apply the SAEM algorithm to get estimators of these parameters as well as of θ (see section 4.5.2 for some application).*

4.4 Convergence result

We suppose that given $g_1, \dots, g_n$ the density functions of $y_1, \dots, y_n$, the problem

$$\text{For all } 1 \leq i \leq n, \quad g_i(y) = \int_{\mathcal{Y}} h_i(y|z) \pi(z) dz \quad \text{for all } y \text{ in } \mathcal{Y},$$

has an unique solution π to get an identifiable problem. We define $\Lambda_n(\theta) = E[\ell_n(\theta)]$ as the expected log-likelihood and consider the following assumptions:

- **(H1)** For all J in $\mathbb{N}^*$ and all n in $\mathbb{N}^*$, Λ_n has a local maximum $\theta_{n,J}^*$ in $\mathbb{R}^J$.
- **(H2)** There exist positive constants m and M such that for all $1 \leq i \leq n$

$$m \leq \sup_{(y,z) \in \mathcal{Y} \times \mathcal{I}} h_i(y|z) \leq M.$$

By the way, from now on, we suppose that

$$J = o(n^{\frac{1}{2}-\varepsilon}) \quad \text{for } \varepsilon \text{ in }]0, \frac{1}{2}[. \quad (4.6)$$

Theorem 4.4. *Assume conditions (H1) and (H2) are satisfied. Then on an event whose probability tends to 1 as n tends to infinity, the function ℓ_n has a local maximum $\hat{\theta}_{n,J}$ and we have for all ε in $]0, \frac{1}{2}[$:*

$$\| \log \pi_{\hat{\theta}_{n,J}} - \log \pi_{\theta_{n,J}^*} \|_2 = O_P \left(n^\varepsilon \sqrt{\frac{J}{n}} \right).$$

Remark. *If the density function π belongs to $\mathcal{M}_{J_0}$ for some J_0 , the existence of θ_{n,J_0}^* is ensured since we have $\pi_{\theta_{n,J_0}^*} = \pi$. So let us consider logspline models such that $\mathcal{M}_J \subset \mathcal{M}_{J+1}$ (take for example dyadic subdivisions of the compact interval $\mathcal{I}$). Then π belongs to all $\mathcal{M}_J$ for all $J \geq J_0$ and we obtain the convergence of $\log \pi_{\hat{\theta}_{n,J}}$ to $\log \pi$.*

Remark. *This convergence result shows that we have to choose J not too large for the applications, namely $J = o(n^{\frac{1}{2}-\varepsilon})$ for ε in $]0, \frac{1}{2}[$. We illustrate this remark by some application in section 4.5.1.*

Proof of Theorem 4.4:

The proof is adapted from this proposed by Stone (1990) for the direct estimation problem. We will show that there exists a local maximum for ℓ_n in a neighborhood of a local maximum of Λ_n . So we need some technical lemma similar to these of Stone, which are proved in the appendix.

Given ε in $]0, \frac{1}{2}[$ and J satisfying (4.6), we define $b_n = n^\varepsilon \sqrt{\frac{J}{n}}$ for n in $\mathbb{N}^*$.

Consider then the two following compact subsets of Θ :

$$\Theta_{1,n} = \{ \theta \in \Theta, \| \log \pi_\theta - \log \pi_{\theta^*} \|_2 \leq b_n \}$$

$$\Theta_{2,n} = \{ \theta \in \Theta, \| \log \pi_\theta - \log \pi_{\theta^*} \|_2 = b_n \}$$

We denote no more the dependence of θ^* in n and in J . To begin the proof, let us define for $n \geq 1$, for $b > 0$ and θ in Θ the following event:

$$\Omega_{n,b,\theta} = \{ |(\ell_n(\theta) - \ell_n(\theta^*)) - (\Lambda_n(\theta) - \Lambda_n(\theta^*))| \geq b \| \log \pi_\theta - \log \pi_{\theta^*} \|_2 \}.$$

Concerning the probability of this event, we have the following result proved in the appendix:

Lemma 4.5. *Given $M_1 > 0$, there exists $\alpha > 0$ such that for all $n \geq 1$, for all θ in $\Theta_{1,n}$ and for all $0 < b < M_1/\sqrt{J}$,*

$$P(\Omega_{n,b,\theta}) \leq 2e^{-\alpha nb^2}.$$

By the way, since (4.6), there exists a positive constant M_1 such that for all $n \geq 1$, $b_n \leq M_1/\sqrt{J}$. So for all $0 < b < M_1/\sqrt{J}$, Lemma 4.5 ensures that each of these events has a probability who tends to zero when n goes to infinity. We will now build an event as uniform as possible in θ , without losing the asymptotical property about its probability.

We recall here the Lemma 8 of Stone (1990), which uses topological property of the space $\mathcal{S}$ of splines functions on $\mathcal{I}$. It is a consequence of Lemma 7 of Stone (1986) and of (viii) on page 155 of de Boor (1978).

Lemma 4.6. *Given $\varepsilon > 0$ and $\beta > 0$, there exists a positive constant M_2 such that the set*

$$\left\{ \pi_\theta, \theta \in \Theta \text{ such that } \|\log \pi_\theta - \log \pi_{\theta^*}\|_2 \leq n^\varepsilon \sqrt{\frac{J}{n}} \right\}$$

can be recovered by $O(\exp[M_2 J \log(n)])$ subsets of diameter at most equal to $\beta n^{2\varepsilon} \frac{J}{n}$, where the diameter of such a subset A is defined as $\sup\{\|\log \pi_1 - \log \pi_2\|_\infty, \pi_1, \pi_2 \in A\}$.

So Lemma 4.6 allows us for each $\beta > 0$ to recover the set $\{\pi_\theta, \theta \in \Theta_{1,n}\}$ by a finite number of balls of diameter less than βb_n^2 . So there exist a finite subset I of $\mathbb{N}$, $(\theta_i)_{i \in I} \in (\Theta_{1,n})^I$ such that

$$\text{Card}(I) = O(\exp[M_2 J \log(n)]) \text{ and } \{\pi_\theta, \theta \in \Theta_{1,n}\} \subset \cup_{i \in I} B_\infty(\pi_{\theta_i}, \beta b_n^2) \quad (4.7)$$

where $B_\infty(\pi_{\theta_0}, r) = \{\pi_\theta, \theta \in \Theta_{1,n}, \|\log \pi_{\theta_0} - \log \pi_\theta\|_\infty \leq r\}$.

So if we define the event $\Omega_{n,b} = \cup_{i \in I} \Omega_{n,b,\theta_i}$, we obtain by (4.7)

$$P(\Omega_{n,b}) \leq 2e^{M_2 J \log n - n \alpha b^2}. \quad (4.8)$$

We will choose b later, as a function of n , such that this probability tends to zero when n goes to infinity.

Consider now θ in $\Theta_{2,n}$. There exists i in I such that $\pi_\theta \in B_\infty(\pi_{\theta_i}, \beta b_n^2)$. So let us write the following triangular inequality:

$$\begin{aligned} |(\ell_n(\theta) - \ell_n(\theta^*)) - (\Lambda_n(\theta) - \Lambda_n(\theta^*))| &\leq |(\ell_n(\theta) - \ell_n(\theta_i)) - (\Lambda_n(\theta) - \Lambda_n(\theta_i))| \\ &\quad + |(\ell_n(\theta_i) - \ell_n(\theta^*)) - (\Lambda_n(\theta_i) - \Lambda_n(\theta^*))| \end{aligned} \quad (4.9)$$

Concerning the first term of the right expression in (4.9), we will use the following lemma proved in the appendix:

Lemma 4.7. *For all $n \geq 1$, for all θ_1 and θ_2 in Θ , we have*

$$|(\ell_n(\theta_1) - \ell_n(\theta_2)) - (\Lambda_n(\theta_1) - \Lambda_n(\theta_2))| \leq 2 \|\log \pi_{\theta_1} - \log \pi_{\theta_2}\|_\infty$$

which ensures that:

$$|(\ell_n(\theta) - \ell_n(\theta_i)) - (\Lambda_n(\theta) - \Lambda_n(\theta_i))| \leq 2\beta b_n^2$$

Concerning the second term of the right expression in (4.9), we have on $\Omega_{n,b}^c$ since θ belongs to $\Theta_{2,n}$:

$$|(\ell_n(\theta_i) - \ell_n(\theta^*)) - (\Lambda_n(\theta_i) - \Lambda_n(\theta^*))| < bb_n$$

So we can conclude that

$$|(\ell_n(\theta) - \ell_n(\theta^*)) - (\Lambda_n(\theta) - \Lambda_n(\theta^*))| < 2\beta b_n^2 + bb_n$$

Moreover, we have an upper bound control concerning the comportment of Λ_n around θ^* :

Lemma 4.8. *Given $M_1 > 0$, there exists $\delta > 0$ such that for all $n \geq 1$ and for all θ in Θ such that $\|\log \pi_\theta - \log \pi_{\theta^*}\|_2 \leq \frac{M_1}{\sqrt{J}}$,*

$$\Lambda_n(\theta) - \Lambda_n(\theta^*) \leq -\delta \|\log \pi_\theta - \log \pi_{\theta^*}\|_2^2.$$

which implies that there exists $\delta > 0$ such that

$$\ell_n(\theta) - \ell_n(\theta^*) < 2\beta b_n^2 + bb_n - \delta b_n^2. \quad (4.10)$$

It remains to choose β and b such that the upper bound in (4.10) stays below zero and the upper bound in (4.8) tends to zero.

Let us take for example $\beta = \delta/6$ and $b = \delta b_n/3$, so we obtain

$$\ell_n(\theta) - \ell_n(\theta^*) \leq -\frac{1}{3}\delta b_n^2 < 0.$$

We now check that the probability of $\Omega_{n,\delta b_n/3}$ tends to zero when n goes to infinity. The exponent in (4.8) is equal to $M_2 J \log n - \frac{\alpha \delta^2}{9} J n^{2\varepsilon}$ which tends to $-\infty$ when n goes to infinity. This ensures that the probability of $\Omega_{n,\delta b_n/3}$ tends to zero with n . So we have shown that excepted of an event whose probability tends to zero, the function ℓ_n has a local maximum in $\Theta_{1,n}$, implying the expected result. $\square$

We now study the convergence of π_{θ^*} to π when the size J of the logspline model goes to infinity, motivated by the approximation property of the logspline model. We only present here results in the following particular case: the observed data $y_1, \dots, y_n$ are related to the unobserved data $z_1, \dots, z_n$ through the same conditional density h for all $1 \leq i \leq n$, so that the function Λ no more depends on n . To ensure more accuracy of our notations, we indicate again from now on the dependence of θ^* in J . We define the density function $g_{\theta_J^*}$ by $g_{\theta_J^*}(y) = \int_{\mathcal{Y}} h(y|z) \pi_{\theta_J^*}(z) dz$ for all J in $\mathbb{N}^*$ and the density function g by $g(y) = \int_{\mathcal{Y}} h(y|z) \pi(z) dz$ for all y in $\mathcal{Y}$. We consider now the following assumptions on the densities π and h :

- (H3) The density π is positive on $\mathcal{I}$ and continuous.
- (H4) The conditional density function h is continuous in y and bounded in z .

Proposition 4.9. *Suppose that assumptions (H1)-(H4) hold. Then the density $g_{\theta_J^*}$ converges uniformly to the density g , when J goes to infinity.*

Proof of Proposition 4.9: Under assumption **(H3)**, the approximation property of the log-spline model ensures that for all J there exists θ_J in $\mathbb{R}^J$ such that:

$$\|\log \pi_{\theta_J} - \log \pi\|_{\infty} \leq 2\delta_J,$$

where δ_J tends to zero with J with a rate depending on the regularity of the function $\log \pi$ (see Stone (1990)). For example, if we suppose the function $\log \pi$ q -times differentiable, we obtain $\delta_J = o(J^{-q})$ (see de Boor (1978) for more details). Since the functions $\exp$ and $\log$ are increasing and the integral linear, we obtain for all J :

$$0 \leq \Lambda(\pi) - \Lambda(\pi_{\theta_J}) \leq 2\delta_J.$$

By definition of θ_J^* , we deduce:

$$0 \leq \Lambda(\pi) - \Lambda(\pi_{\theta_J^*}) \leq 2\delta_J,$$

which ensures that $\Lambda(\pi_{\theta_J^*})$ tends to $\Lambda(\pi)$ when J goes to infinity. Moreover, the set of densities $\{g_{\theta_J^*}, J \in \mathbb{N}^*\}$ is equicontinuous: for all J and all (y, y') in $\mathcal{Y}^2$, we have:

$$|g_{\theta_J^*}(y) - g_{\theta_J^*}(y')| \leq \sup_{z \in \mathcal{I}} |h(y|z) - h(y'|z)|.$$

Moreover $\mathcal{I}$ is a compact interval of $\mathbb{R}$ and assumption **(H4)** holds, so this upper bound is finite and ensures that the set of functions $\{g_{\theta_J^*}, J \in \mathbb{N}^*\}$ is equicontinuous. By the way, each density of this set is bounded by M under assumption **(H2)**. So Ascoli's theorem can be apply to this function set and ensures the existence of some uniformly convergent subsequence of the sequence $(g_{\theta_J^*})$. Let us denote $\tilde{g}$ its limit. By the dominated convergence theorem, we obtain that $\tilde{g} = g$. Since the sequence $(g_{\theta_J^*})$ has an unique accumulation point, it implies the convergence of the sequence toward this unique accumulation point and therefore the expected result. $\square$

We define now the following assumption **(H5)**: the missing data problem is of the form:

$$Y = \phi(Z) + \varepsilon, \tag{4.11}$$

where ϕ is a injective function from $\mathbb{R}$ on $\mathbb{R}$, the error term ε follows a given distribution Δ with a non zero characteristic function and the variables Z and ε are mutually independent.

Consider a sequence of random variables (Y_J, Z_J) satisfying **(H5)**, where the variable Z_J has density $\pi_{\theta_J^*}$. Consider also a couple of variables (Y, Z) satisfying **(H5)**, where the variable Z has density π .

Proposition 4.10. *Suppose assumptions **(H1)**-**(H5)** are satisfied. Then the variable Z_J converges in distribution to a random variable Z having density π .*

Proof of Proposition 4.10: Consider the characteristic function of Y_J , for some real t in its definition field:

$$\Phi_{Y_J}(t) = \int_{\mathcal{Y}} \exp(it y) g_{\theta_J^*}(y) dy.$$

By the dominated convergence theorem, $\Phi_{Y_J}(t)$ converges to the characteristic function $\Phi_Y(t)$ of Y . Since Z and ε are supposed mutually independent, as well as Z_J and ε_J for all integer J , we have:

$$\Phi_Y(t) = \Phi_{\phi(Z)}(t) \Phi_{\varepsilon}(t)$$

and also

$$\Phi_{Y_J}(t) = \Phi_{\phi(Z_J)}(t)\Phi_\varepsilon(t).$$

By application of the result of Proposition 4.9 and of the dominated convergence theorem, we deduce that for all real t in its definition field, $\lim_J \Phi_{Y_J}(t) = \Phi_Y(t)$. Since $\Phi_\varepsilon(t) \neq 0$, we obtain the convergence of $\Phi_{\phi(Z_J)}(t)$ to $\Phi_{\phi(Z)}(t)$, implying the convergence in distribution of $\phi(Z_J)$ to $\phi(Z)$. Under assumption **(H5)**, ϕ is injective, so we deduce the convergence in distribution of Z_J to Z . $\square$

4.5 Applications

4.5.1 Application with missing data normally distributed

We consider the following simulated example to illustrate our method. The observations (y_i) are obtained from the non observed data (z_i) at time t by the equality:

$$y_i = \alpha_i \exp(-z_i t) + \varepsilon_i \quad \text{for } 1 \leq i \leq n. \quad (4.12)$$

The coefficients (α_i) are independent and identically drawn from a Gaussian distribution with mean m_α and variance η_α^2 , but are considered as a covariable for the time being. The missing data are also supposed to be independent and identically drawn from a Gaussian distribution π with mean μ and variance γ^2 . The errors (ε_i) are independent and identically distributed from a Gaussian distribution with mean zero and variance σ^2 . Moreover, we suppose that the missing data (z_i) and the errors (ε_i) are mutually independent. We choose as numerical value $n = 100$, $m_\alpha = 5$, $\eta_\alpha^2 = 2.25$, $\mu = 10$, $\gamma^2 = 0.04$ and $\sigma^2 = 0.0028$ to ensure a signal to noise ratio equal to 10.

We apply the algorithm SAEM in a logspline model like presented in section 4.3 to approach the maximum likelihood logspline estimator of the density function π of the missing data z , the parameter σ^2 considered as known.

We choose zero as initial value θ_0 , such that the logspline density estimate is initialized with the uniform distribution on $\mathcal{I}$. We first run the algorithm choosing a large support for π , namely $\mathcal{I} = [-50, 50]$ to locate more precisely where the missing data take their values. Then we adjust the choice of the support according to the results taking into account the part of the interval where the logspline estimate of π is upper 10^{-4} , so we take $\mathcal{I} = [-2.5, 12.5]$ (see Figure 4.1).

For the choice of the size J of the logspline model, we consider the remark after Theorem 4.4. So we test some values for J lower than 10 since we have 100 observations. The results obtained are presented on Figure 4.2. We see clearly that the estimation of the parameters have converged for J equal to 4 and 5, but this is not the case for the values 6 or 7, although the logspline density estimate is not far from the true density π . The best estimation seems to be given for $q = 4$ and $J = 5$, so we will hold these values for the following applications. We draw the attention of the reader to the fact that some curves of the estimate of θ are almost superposed. We also present results for logspline model of different order q , namely we run the SAEM algorithm with histograms, exponential of linear functions and of polynomial functions of degrees 2 and 3. The results are presented in Figure 4.3 and show that the degrees 2 and 3 give similar results.

FIG. 4.1: (a) Estimation of θ , (b) True density π in dotted line and logspline density estimate in plain line for $\mathcal{I} = [-50, 50]$; (c) Estimation of θ , (d) True density π in dotted line and logspline density estimate in plain line for $\mathcal{I} = [-2.5, 12.5]$, ($q = 4, J = 5$).

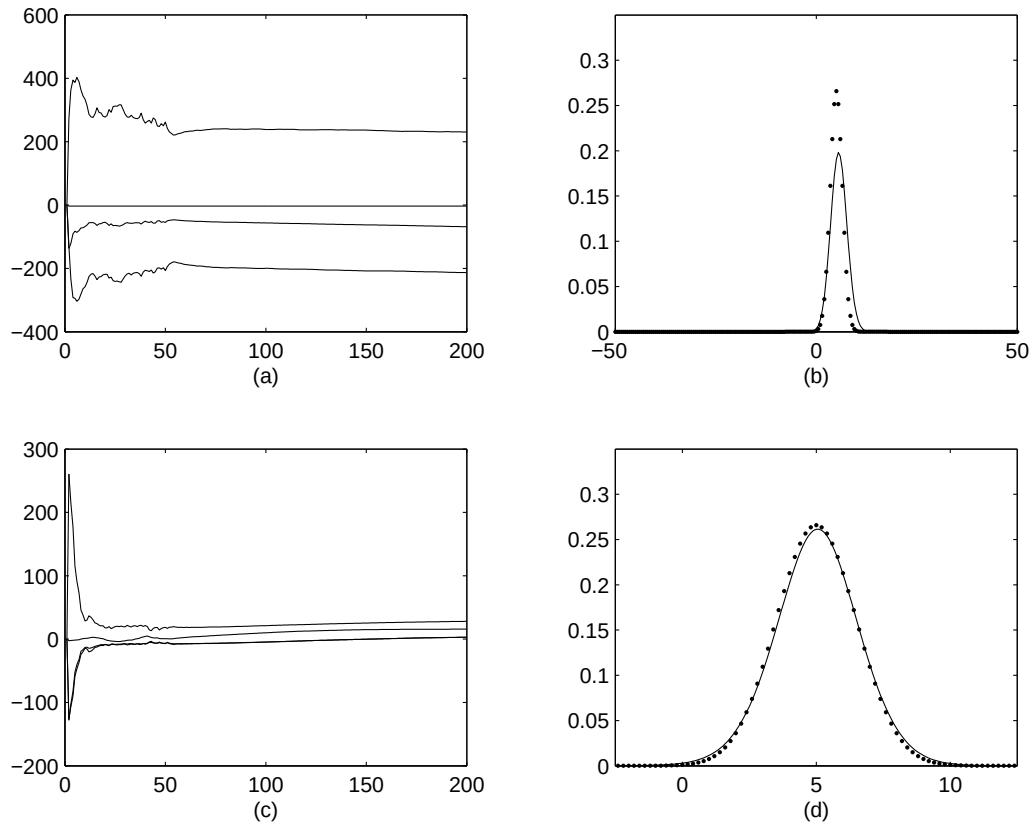


FIG. 4.2: (a)-(d) Estimation of θ for $q = 4$ and J equal respectively to 7, 6, 5 and 4. (e)-(h) True density π in dotted line, corresponding logspline density estimate in plain line.

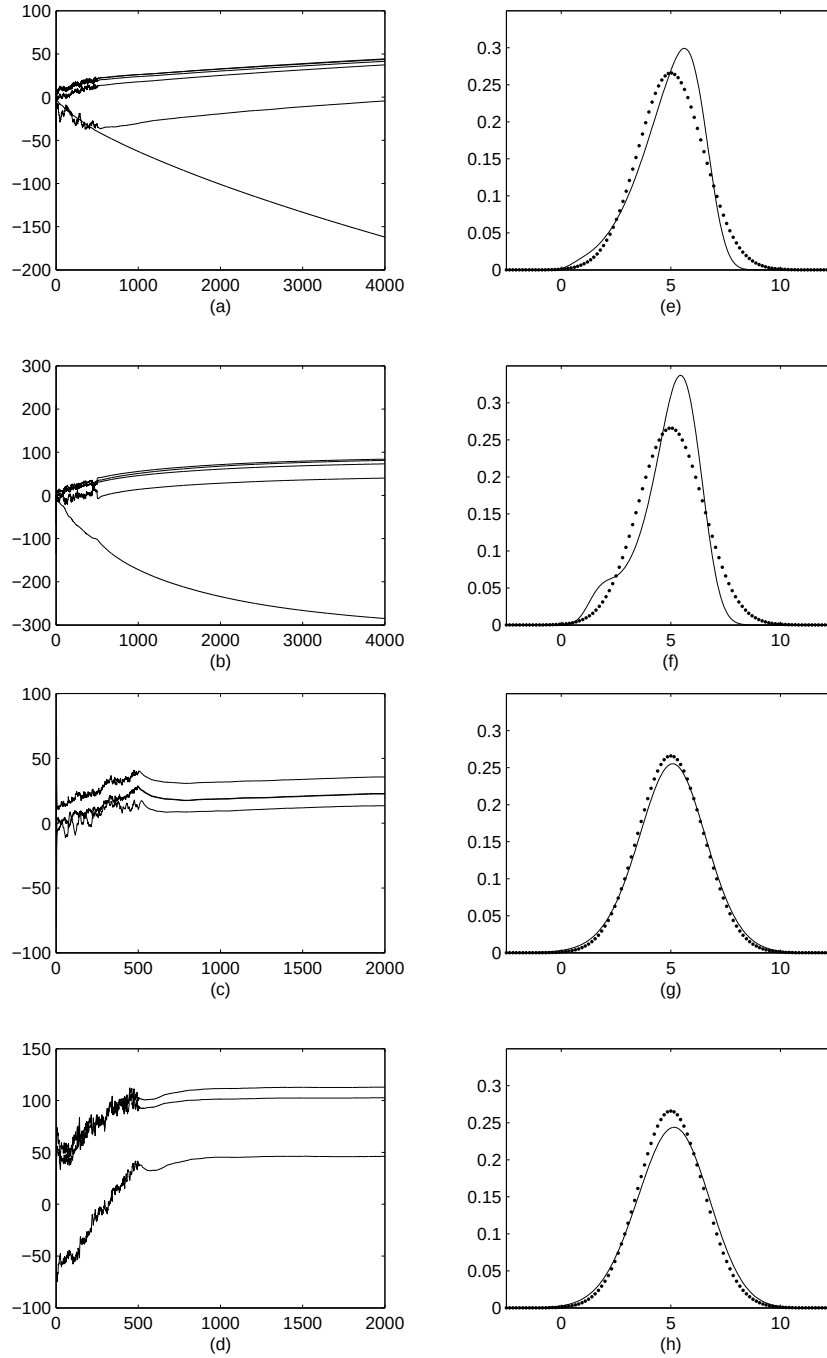


FIG. 4.3: (a)-(d) True density π in dotted line, logspline density estimate in plain line for $J = 5$ and q equal respectively to 1, 2, 3 and 4.

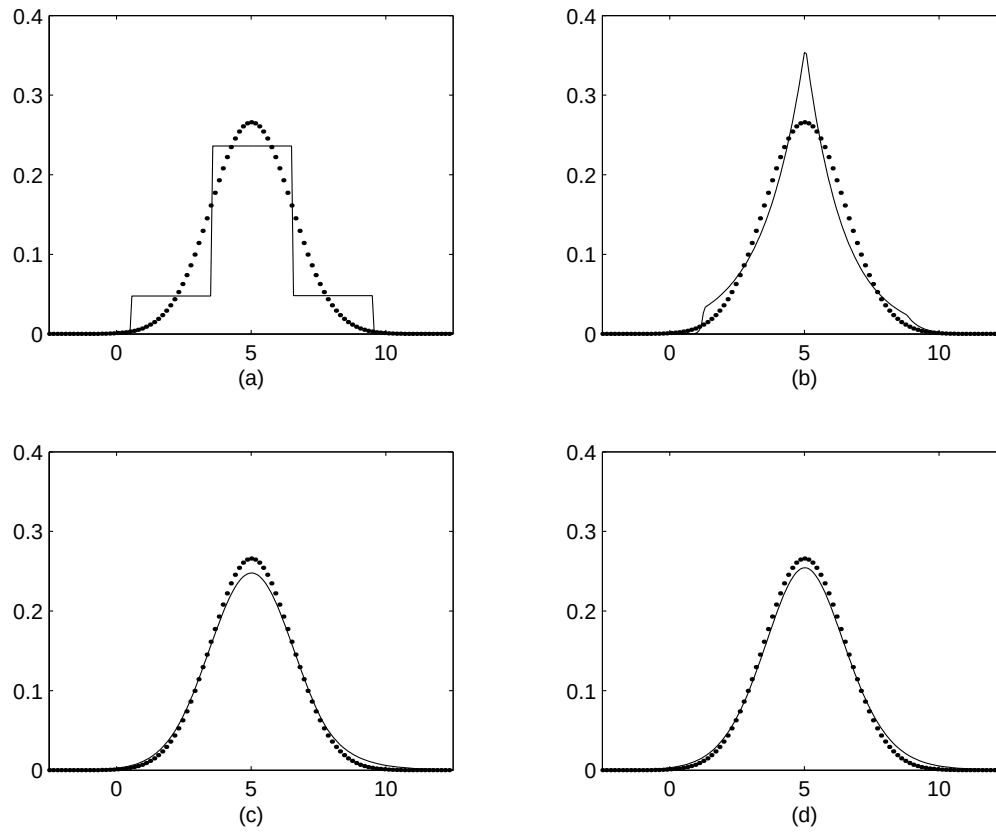
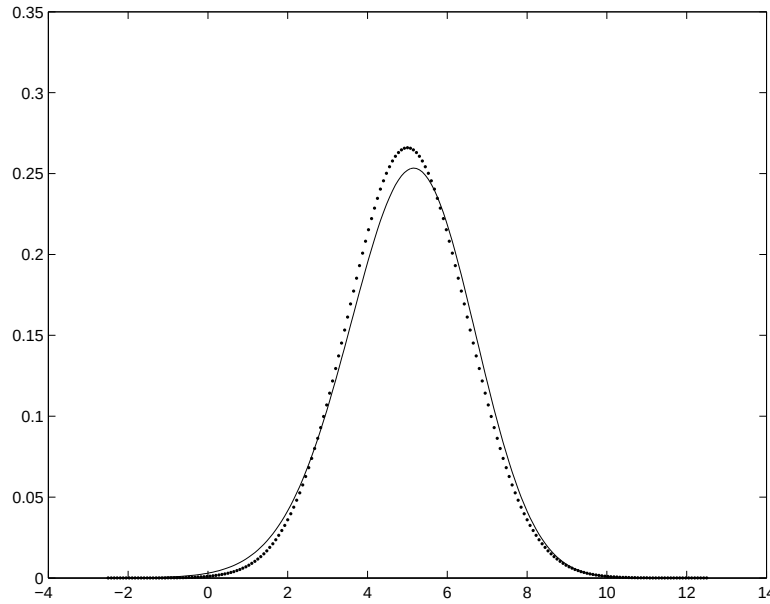


FIG. 4.4: True density π in dotted line, logspline density estimate in plain line ($q = 4, J = 5$).

4.5.2 Application involving estimation of other parameters

We consider again the example presented in (4.12) in the first subsection of this part. We will now suppose that the variance σ^2 of the error is unknown and estimate it with the SAEM algorithm of section (4.3). The logspline density estimate obtained with $J = 5$ and $q = 4$ is presented on Figure 4.4 and is close to the true density π . Like it is shown on Figure 4.5, the estimator of θ is quite stationary after 2000 iterations. We obtain for the estimator of σ^2 the value 0.0025 which seems to be a good estimation, the true value being equal to 0.0028. We consider now that the variance σ^2 of the error and the coefficients (α_i) are unknown. We compute with the SAEM algorithm the logspline density estimate and the estimators of m_α , η_α and σ^2 . The logspline density estimate is also close to the true density π , the variance σ^2 of the error is well estimated since we obtain 0.0027. The estimators obtained for the parameters m_α and η_α^2 are not so relevant, since we obtain 7.63 for the mean and 0.0003 for the variance. The results are presented on Figures 4.6 and 4.7.

4.5.3 Application with missing data drawn from a bimodal distribution

We conserve the model (4.12) and only change the distribution from which the missing data are drawn. We choose a bimodal distribution equal to a mixture of two Gaussian distributions with weights 0.5, means equal to 1 and 2 and variances equal to 0.25^2 . We consider that the variance σ^2 of the error is known. We compute the logspline density estimate for the density of the missing data. The estimator is not far from the true density, the places of the two maxima of the density are quite recovered. The results are presented on Figures 4.8 and 4.9.

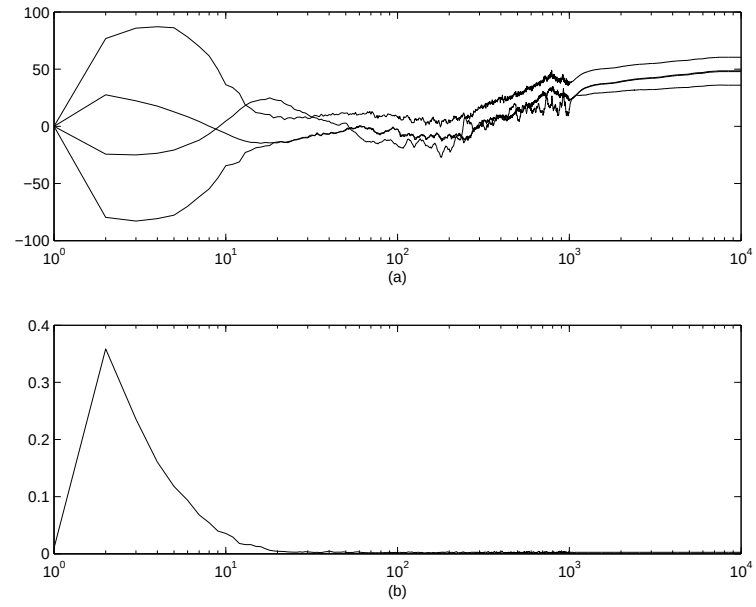
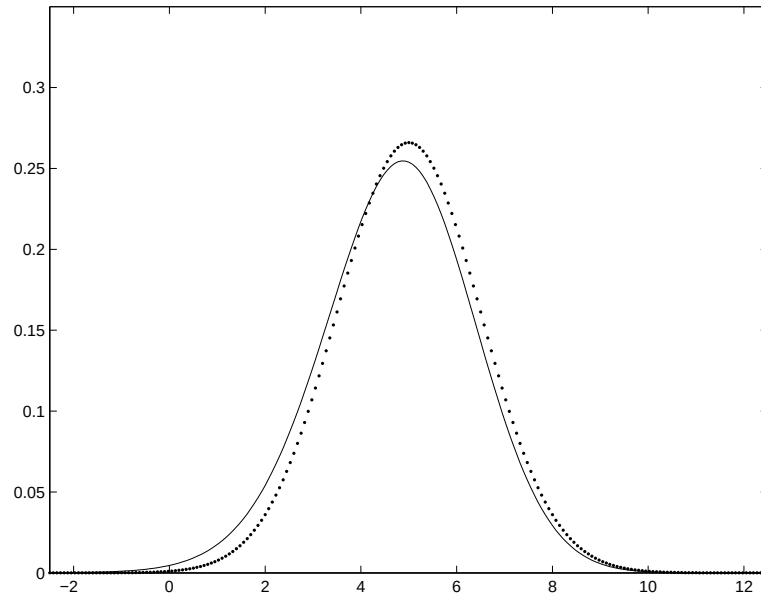
FIG. 4.5: (a) Estimation of the parameter θ , (b) Estimation of the parameter σ^2 ($q = 4, J = 5$).

 FIG. 4.6: True density π in dotted line, logspline density estimate in plain line ($q = 4, J = 5$).


FIG. 4.7: (a) Estimation of the parameter θ , (b) Estimation of the parameters m_α and η_α^2 , (c) Estimation of the parameter σ^2 , ($q = 4, J = 5$).

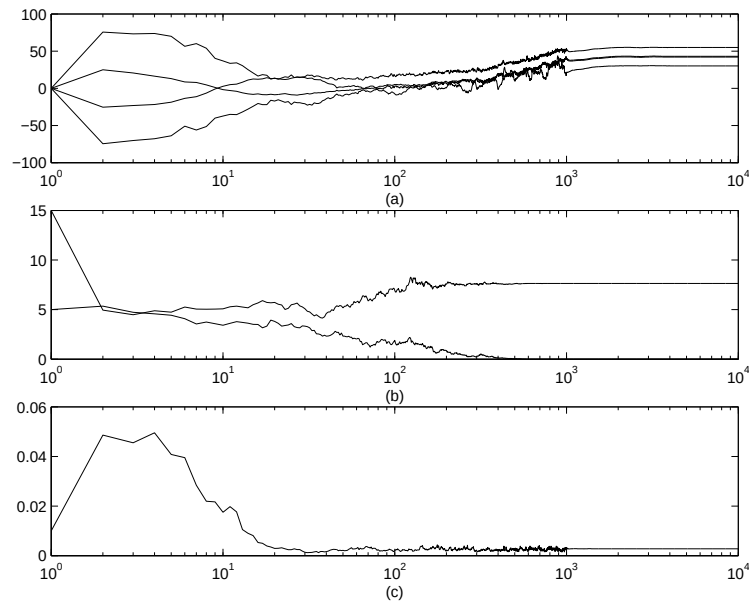


FIG. 4.8: True density π in dotted line, logspline density estimate in plain line ($q = 4, J = 5$).

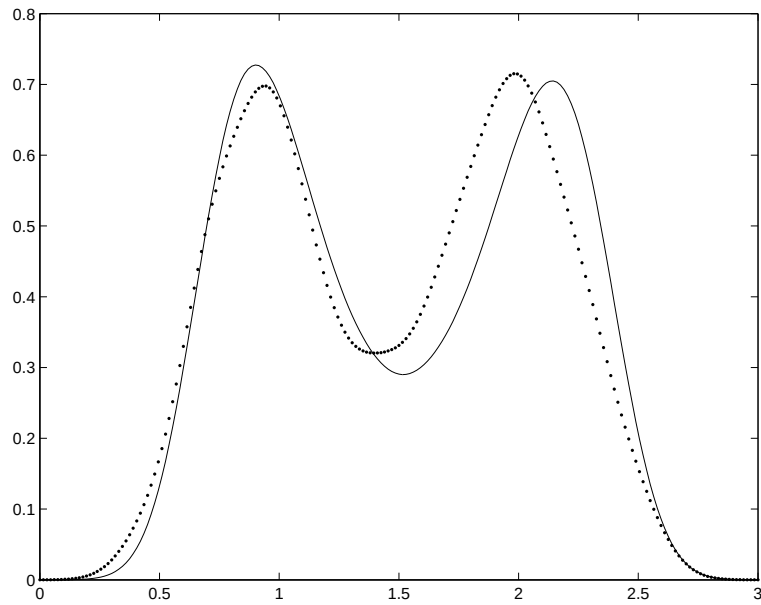
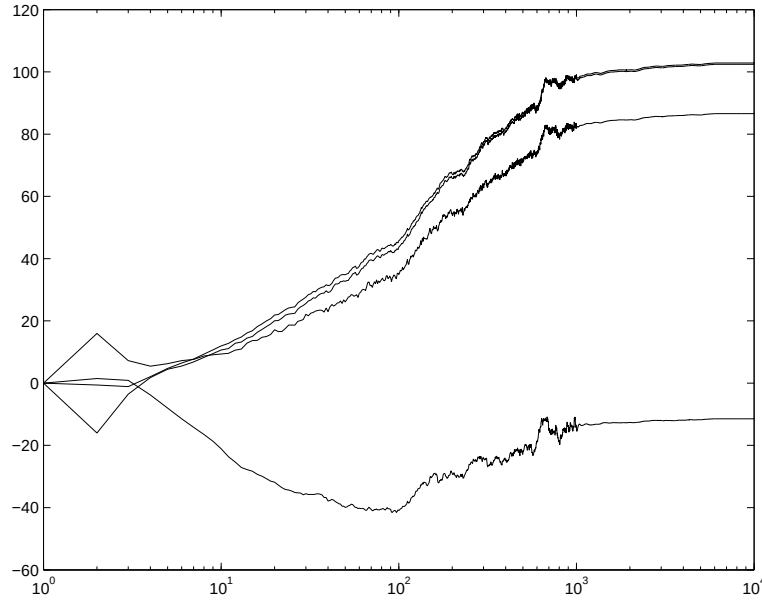


FIG. 4.9: Estimation of the parameter θ , ($q = 4, J = 5$).


4.5.4 Application to the example of the orange trees

We consider again the example of the orange trees presented in section 2.3 of Chapter 2. We recall the model:

$$y_{ij} = C(t_j, \phi_i; \beta_1, \beta_2) + \varepsilon_{ij} \text{ for } 1 \leq i \leq n, 1 \leq j \leq m, \quad (4.13)$$

$$C(t_j, \phi_i; \beta_1, \beta_2) = \frac{\phi_i}{1 + \exp\left(-\frac{t_j - \beta_1}{\beta_2}\right)}. \quad (4.14)$$

We have considered this model as a parametric one, choosing arbitrary a Gaussian prior for the law of the random effect ϕ . We propose now to do a non parametric estimation of the density function π of ϕ , due to the SAEM algorithm applied with the logspline model. We estimate at the same time the others parameters, namely the two fixed effects β_1 and β_2 and the variance σ^2 of the Gaussian errors. We do 10000 iterations of the SAEM algorithm because the parameter θ of the logspline density converges slowly. Figure 4.11 presents the convergence of the parameters using a logarithmic scale for the x-axis. Figure 4.10 represents the density estimate obtained for π and the density function of a Gaussian distribution having mean 192 and variance 1003, these values corresponding to the estimations obtained by the SAEM algorithm in section 2.3 with the Gaussian prior of the distribution of ϕ . The two curves are almost superposed. The estimation of the parameters β_1 , β_2 and σ^2 obtained with this model are detailed in Table 4.1. So we can conclude that, in this example, the Gaussian prior was a well adapt choice.

4.6 Appendix

Proof of Lemma 4.7:

TAB. 4.1: Comparison of estimates obtained with the Gaussian prior of π and with a non parametric estimation of π .

Parameters	β_1	β_2	σ^2
Initial value	650	250	10
Model with Gaussian prior of π	727.36	347.67	61.52
Model with logspline estimation of π	717.15	340.19	61.80

FIG. 4.10: Estimators of the density function π of the random effect ϕ : logspline density estimate in plain line, estimator obtained with Gaussian prior in dotted line ($q = 4, J = 4$).

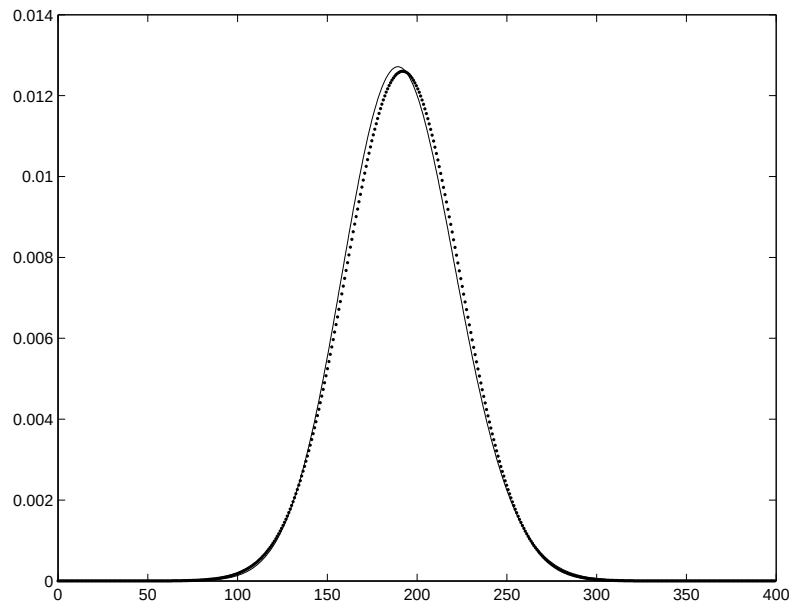
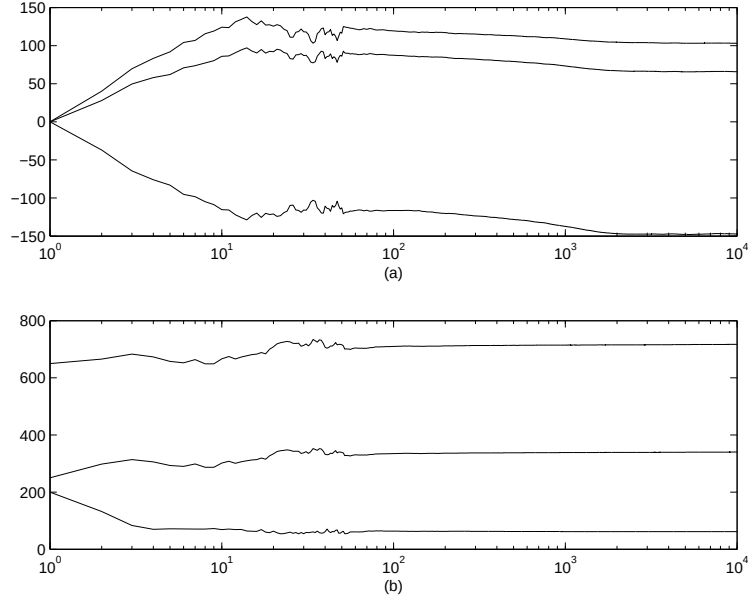


FIG. 4.11: (a) Estimation of the parameters θ , (b) Estimation of the parameters β_1, β_2 and σ^2 ($q = 4, J = 4$). The x-axis uses a logarithmic scale.



We will prove that for all $1 \leq i \leq n$, $\|\log g_{i,\theta_1} - \log g_{i,\theta_2}\|_\infty \leq \|\log \pi_{\theta_1} - \log \pi_{\theta_2}\|_\infty$, which will imply the expected result.

Consider for some $1 \leq i \leq n$, θ_1 and θ_2 in Θ and for y in $\mathcal{Y}$

$$\begin{aligned} \frac{g_{i,\theta_1}(y)}{g_{i,\theta_2}(y)} &= \frac{1}{g_{i,\theta_2}(y)} \int_{\mathcal{I}} h_i(y|z) \pi_{\theta_1}(z) dz \\ &= \int_{\mathcal{I}} \frac{h_i(y|z) \pi_{\theta_2}(z)}{g_{i,\theta_2}(y)} \frac{\pi_{\theta_1}(z)}{\pi_{\theta_2}(z)} dz. \end{aligned}$$

By Jensen's inequality, since $\frac{h_i(y|z) \pi_{\theta_2}(z)}{g_{i,\theta_2}(y)}$ is a density function for all y in $\mathcal{Y}$, we obtain:

$$\log \frac{g_{i,\theta_1}(y)}{g_{i,\theta_2}(y)} \geq \int_{\mathcal{I}} \frac{h_i(y|z) \pi_{\theta_2}(z)}{g_{i,\theta_2}(y)} \log \left(\frac{\pi_{\theta_1}(z)}{\pi_{\theta_2}(z)} \right) dz.$$

So that

$$\begin{aligned} \log \frac{g_{i,\theta_1}(y)}{g_{i,\theta_2}(y)} &\geq - \int_{\mathcal{I}} \frac{h_i(y|z) \pi_{\theta_2}(z)}{g_{i,\theta_2}(y)} \left| \log \left(\frac{\pi_{\theta_1}(z)}{\pi_{\theta_2}(z)} \right) \right| dz \\ &\geq - \|\log \pi_{\theta_1} - \log \pi_{\theta_2}\|_\infty. \end{aligned}$$

By symmetry between θ_1 and θ_2 , we obtain also:

$$\log \frac{g_{i,\theta_2}(y)}{g_{i,\theta_1}(y)} \geq - \|\log \pi_{\theta_1} - \log \pi_{\theta_2}\|_\infty.$$

which implies that

$$\|\log g_{i,\theta_1} - \log g_{i,\theta_2}\|_\infty \leq \|\log \pi_{\theta_1} - \log \pi_{\theta_2}\|_\infty.$$

□

Proof of Lemma 4.5:

This proof repose of Bernstein's inequality (see Hoeffding (1963)) applied to the random variables $X_i = \log g_{i,\theta}(Y_i) - \log g_{i,\theta^*}(Y_i)$, which are independent, bounded uniformly by $\|\log \pi_\theta - \log \pi_{\theta^*}\|_\infty$ since Lemma 4.7. By Lemma 7 of Stone (1986), there exists a positive constant M_3 such that $\|\log \pi_\theta - \log \pi_{\theta^*}\|_\infty \leq M_3 \sqrt{J} \|\log \pi_\theta - \log \pi_{\theta^*}\|_2$ which is the uniform bound we choose for the random variables (X_i) . So by application of Bernstein's inequality, we obtain the following upper bound: there exists a positive constant M_4 such that for all $t > 0$:

$$P(|(\ell_n(\theta) - \ell_n(\theta^*)) - (\Lambda_n(\theta) - \Lambda_n(\theta^*))| \geq t) \leq 2 \exp \left[- \frac{n^2 t^2}{M_4 (\sum_1^n \text{var}(X_i) + nt \sqrt{J} \|\log \pi_\theta - \log \pi_{\theta^*}\|_2)} \right].$$

To conclude to this upper bound, we will show that there exists a positive constant M_5 such that:

$$\text{var}(X_i) \leq M_5 \|\log \pi_\theta - \log \pi_{\theta^*}\|_2^2 \quad (4.15)$$

and conclude , choosing t equal to $b \|\log \pi_\theta - \log \pi_{\theta^*}\|_2$, that there exists a positive constant α such that:

$$P(|(\ell_n(\theta) - \ell_n(\theta^*)) - (\Lambda_n(\theta) - \Lambda_n(\theta^*))| \geq b \|\log \pi_\theta - \log \pi_{\theta^*}\|_2) \leq 2 \exp(-\alpha n b^2),$$

which implies the result. We finally prove the relation (4.15). We have:

$$\text{var}(X_i) \leq E(X_i^2) = \int_{\mathcal{Y}} \left(\log \left(\frac{g_{i,\theta^*}(y)}{g_{i,\theta}(y)} \right)^2 \right) g_i(y) dy,$$

where g_i denotes the density function of Y_i and is equal to $\int_{\mathcal{I}} h_i(y, x) \pi(x) dx$ for all y in $\mathcal{Y}$. So assumption **(H2)** implies that

$$\text{var}(X_i) \leq \frac{M}{m} \int_{\mathcal{Y}} \left(\log \left(\frac{g_{i,\theta^*}(y)}{g_{i,\theta}(y)} \right)^2 \right) g_{i,\theta^*}(y) dy.$$

By the Lemma 1 of Barron and Sheu (1991), we obtain:

$$\text{var}(X_i) \leq \frac{2M}{m} \exp(\|\log g_{i,\theta} - \log g_{i,\theta^*}\|_\infty) \int_{\mathcal{Y}} \log \left(\frac{g_{i,\theta^*}(y)}{g_{i,\theta}(y)} \right) g_{i,\theta^*}(y) dy.$$

We recall here the definition of the Kullback-Leiber divergence defined between two density functions defined as:

$$K(g_1, g_2) = \int_{\mathcal{Y}} g_1(y) \log \left[\frac{g_1(y)}{g_2(y)} \right] dy \quad (4.16)$$

for two density functions g_1 and g_2 on $\mathcal{Y}$ and deduce using Lemma 4.7 that:

$$\text{var}(X_i) \leq \frac{2M}{m} \exp(\|\log \pi_\theta - \log \pi_{\theta^*}\|_\infty) K(g_{i,\theta^*}, g_{i,\theta}).$$

By Lemma 1.1 of Eggermont (1999), we have:

$$\text{var}(X_i) \leq \frac{2M}{m} \exp(\|\log \pi_\theta - \log \pi_{\theta^*}\|_\infty) K(\pi_{\theta^*}, \pi_\theta)$$

and finally applying once again the inequality of Lemma 1 of Barron and Sheu (1991), we conclude by assumption **(H2)**:

$$\text{var}(X_i) \leq \frac{M^2}{m} \exp(2 \|\log \pi_\theta - \log \pi_{\theta^*}\|_\infty) \|\log \pi_\theta - \log \pi_{\theta^*}\|_2^2.$$

Since Lemma 7 of Stone (1986) ensures that there exists M_3 such that $\|\log \pi_\theta - \log \pi_{\theta^*}\|_\infty \leq M_3 \sqrt{J} \|\log \pi_\theta - \log \pi_{\theta^*}\|_2$, since θ belongs to the subset $\Theta_{1,n}$ and since $b_n \sqrt{J} \leq M_1$, we obtain:

$$\text{var}(X_i) \leq \frac{M^2}{m} \exp(2M_1 M_3) \|\log \pi_\theta - \log \pi_{\theta^*}\|_2^2,$$

which is the expected upper bound. $\square$

Proof of Lemma 4.8:

For θ in $\Theta_{1,n}$ and t in $[0, 1]$, let us define $\theta_t = \theta^* + t(\theta - \theta^*)$ and $F_{n,\theta}(t) = \Lambda_n(\theta_t)$. The function $F_{n,\theta}$ is differentiable on $[0, 1]$, so we can write:

$$\begin{aligned} \Lambda_n(\theta^*) - \Lambda_n(\theta) &= F_{n,\theta}(0) - F_{n,\theta}(1) \\ &= - \int_0^1 F'_{n,\theta}(u) du. \end{aligned}$$

If we derive explicitly the function $F_{n,\theta}$, we obtain the following expression for its derivate for all u in $[0, 1]$:

$$F'_{n,\theta}(u) = \frac{1}{n} \sum_{i=1}^n \left[\int_{\mathcal{I}} \pi_{\theta_u}(z) \log \frac{\pi_{\theta^*}(z)}{\pi_\theta(z)} dz - \int_{\mathcal{Y}} g_i(y) \int_{\mathcal{I}} \mu_{i,\theta_u}(y, z) \log \frac{\pi_{\theta^*}(z)}{\pi_\theta(z)} dz dy \right],$$

where $\mu_{i,\theta}(y, z) = \frac{h_i(y|z)\pi_\theta(z)}{g_{i,\theta}(y)}$. Thus,

$$F'_{n,\theta}(u) = \Psi_\theta(u) - \frac{1}{n} \sum_{i=1}^n \Phi_{i,\theta}(u),$$

$$\text{where } \Psi_\theta(u) = \int_{\mathcal{I}} \pi_{\theta_u}(z) \log \frac{\pi_{\theta^*}(z)}{\pi_\theta(z)} dz \quad \text{and} \quad \Phi_{i,\theta}(u) = \int_{\mathcal{Y}} g_i(y) \int_{\mathcal{I}} \mu_{i,\theta_u}(y, z) \log \frac{\pi_{\theta^*}(z)}{\pi_\theta(z)} dz dy.$$

We will first consider $\int_0^1 \Psi_\theta(u) du$ and apply the following lemma proved at the end of the appendix:

Lemma 4.11. *For all θ in Θ , $\int_0^1 \Psi_\theta(u) du = 0$.*

which implies that

$$- \int_0^1 F'_{n,\theta}(u) du = \frac{1}{n} \sum_{i=1}^n \int_0^1 \Phi_{i,\theta}(u) du.$$

Consider now for all $1 \leq i \leq n$ the function $\Phi_{i,\theta}$ and let us derive it to study its variations. For all $1 \leq i \leq n$ and all u in $[0, 1]$, we have:

$$\Phi'_{i,\theta}(u) = - \int_{\mathcal{Y}} g_i(y) \left[\int_{\mathcal{I}} \mu_{i,\theta_u}(y, z) \left(\log \frac{\pi_{\theta^*}(z)}{\pi_{\theta}(z)} \right)^2 dz - \left(\int_{\mathcal{I}} \mu_{i,\theta_u}(y, z) \log \frac{\pi_{\theta^*}(z)}{\pi_{\theta}(z)} dz \right)^2 \right] dy,$$

which is non positive by application of Cauchy-Schwarz's inequality. So for all $1 \leq i \leq n$ the function $\Phi_{i,\theta}$ is non increasing from $\Phi_{i,\theta}(0)$ to $\Phi_{i,\theta}(1)$. So we have for all $1 \leq i \leq n$

$$\Phi_{i,\theta}(1) \leq \int_0^1 \Phi_{i,\theta}(u) du \leq \Phi_{i,\theta}(0),$$

which implies

$$\frac{1}{n} \sum_{i=1}^n \Phi_{i,\theta}(1) \leq \frac{1}{n} \sum_{i=1}^n \int_0^1 \Phi_{i,\theta}(u) du \leq \frac{1}{n} \sum_{i=1}^n \Phi_{i,\theta}(0).$$

Since each function $\Phi_{i,\theta}$ is continuous, by application of the intermediate value theorem, there exists $u(\theta)$ in $[0, 1]$ such that

$$\frac{1}{n} \sum_{i=1}^n \int_0^1 \Phi_{i,\theta}(u) du = \frac{1}{n} \sum_{i=1}^n \Phi_{i,\theta}[u(\theta)]. \quad (4.17)$$

Thus,

$$- \int_0^1 F'_{n,\theta}(u) du = \frac{1}{n} \sum_{i=1}^n \int_{\mathcal{Y}} g_i(y) \int_{\mathcal{I}} \mu_{i,\theta_{u(\theta)}}(y, z) \log \frac{\pi_{\theta^*}(z)}{\pi_{\theta}(z)} dz dy. \quad (4.18)$$

We will now study how the right term is bounded by below by the square of the L_2 norm of the function $\log \frac{\pi_{\theta^*}}{\pi_{\theta}}$. To that purpose, we adapt the proof of Barron and Sheu (1991) for the Kullback-Leiber divergence presented in Lemma 1 of their paper. The result is based of the following inequality derived from the Taylor expansion of e^z :

$$\frac{z^2}{2} e^{-z_-} \leq e^z - 1 - z, \quad (4.19)$$

where z is real and $z_- = \max(-z, 0)$. So we first introduce the following decomposition for all $0 < \lambda \leq 1$:

$$- \int_0^1 F'_{n,\theta}(u) du = \frac{1}{n} \sum_{i=1}^n \frac{1}{\lambda} \int_{\mathcal{Y}} g_i(y) \int_{\mathcal{I}} \mu_{i,\theta_{u(\theta)}}(y, z) \lambda \log \frac{\pi_{\theta^*}(z)}{\pi_{\theta}(z)} dz dy.$$

Let us now write this expression as follows in order to apply the inequality (4.19):

$$\begin{aligned} - \int_0^1 F'_{n,\theta}(u) du &= \frac{1}{\lambda} \frac{1}{n} \sum_{i=1}^n \int_{\mathcal{Y}} g_i(y) \int_{\mathcal{I}} \mu_{i,\theta_{u(\theta)}}(y, z) \left[\lambda \log \frac{\pi_{\theta^*}(z)}{\pi_{\theta}(z)} - 1 + \left(\frac{\pi_{\theta}(z)}{\pi_{\theta^*}(z)} \right)^{\lambda} \right] dz dy \\ &\quad + \frac{1}{\lambda} \frac{1}{n} \sum_{i=1}^n \int_{\mathcal{Y}} g_i(y) \int_{\mathcal{I}} \mu_{i,\theta_{u(\theta)}}(y, z) \left[1 - \left(\frac{\pi_{\theta}(z)}{\pi_{\theta^*}(z)} \right)^{\lambda} \right] dz dy. \end{aligned} \quad (4.20)$$

By application of (4.19) and assumption **(H2)**, we obtain the following lower bound:

$$\begin{aligned}
 & \int_{\mathcal{Y}} g_i(y) \int_{\mathcal{I}} \mu_{i, \theta_{u(\theta)}}(y, z) \left[\lambda \log \frac{\pi_{\theta^*}(z)}{\pi_{\theta}(z)} - 1 + \left(\frac{\pi_{\theta}(z)}{\pi_{\theta^*}(z)} \right)^{\lambda} \right] dz dy \\
 & \geq \int_{\mathcal{Y}} g_i(y) \int_{\mathcal{I}} \mu_{i, \theta_{u(\theta)}}(y, z) \left[\frac{\lambda^2}{2} \left(\log \frac{\pi_{\theta^*}(z)}{\pi_{\theta}(z)} \right)^2 \right] \exp \left[-\lambda \left(\log \frac{\pi_{\theta^*}(z)}{\pi_{\theta}(z)} \right) \right] dz dy \\
 & \geq \frac{m}{M} \int_{\mathcal{I}} \left[\frac{\lambda^2}{2} \left(\log \frac{\pi_{\theta^*}(z)}{\pi_{\theta}(z)} \right)^2 \right] \exp \left[\log \pi_{\theta_{u(\theta)}}(z) - \left(\log \frac{\pi_{\theta^*}(z)}{\pi_{\theta}(z)} \right) \right] dz
 \end{aligned}$$

Since the function $\theta \rightarrow \log \pi_{\theta}$ is concave, we can bound by below this expression as follows:

$$\begin{aligned}
 & \frac{m}{M} \int_{\mathcal{I}} \left[\frac{\lambda^2}{2} \left(\log \frac{\pi_{\theta^*}(z)}{\pi_{\theta}(z)} \right)^2 \right] \exp \left[\log \pi_{\theta_{u(\theta)}}(z) - \left(\log \frac{\pi_{\theta^*}(z)}{\pi_{\theta}(z)} \right) \right] dz \\
 & \geq \frac{m}{M} \frac{\lambda^2}{2} \min \pi_{\theta^*} \exp \left[-\left\| \log \frac{\pi_{\theta^*}}{\pi_{\theta}} \right\|_{\infty} \right] \int_{\mathcal{I}} \left[\left(\log \frac{\pi_{\theta^*}(z)}{\pi_{\theta}(z)} \right)^2 \right] dz
 \end{aligned}$$

So finally, we obtain the following lower bound for the first term of the right expression in (4.20):

$$\begin{aligned}
 & \frac{1}{n} \sum_{i=1}^n \frac{1}{\lambda} \int_{\mathcal{Y}} g_i(y) \int_{\mathcal{I}} \mu_{i, \theta_{u(\theta)}}(y, z) \left[\lambda \log \frac{\pi_{\theta^*}(z)}{\pi_{\theta}(z)} - 1 + \left(\frac{\pi_{\theta}(z)}{\pi_{\theta^*}(z)} \right)^{\lambda} \right] dz dy \\
 & \geq \frac{m}{M} \frac{\lambda}{2} \min \pi_{\theta^*} \exp \left[-\left\| \log \frac{\pi_{\theta^*}}{\pi_{\theta}} \right\|_{\infty} \right] \int_{\mathcal{I}} \left[\left(\log \frac{\pi_{\theta^*}(z)}{\pi_{\theta}(z)} \right)^2 \right] dz
 \end{aligned}$$

We will now prove that λ can be chosen such that the second term of the right expression in (4.20) is non negative. To that purpose, we introduce the function T defined as follow for θ in Θ and for λ in $[0, 1]$:

$$T(\theta, \lambda) = \frac{1}{n} \sum_{i=1}^n \int_{\mathcal{Y}} g_i(y) \int_{\mathcal{I}} \mu_{i, \theta_{u(\theta)}}(y, z) \left(\frac{\pi_{\theta}(z)}{\pi_{\theta^*}(z)} \right)^{\lambda} dz dy.$$

So we obtain the following expression for the second term of the right expression in (4.20):

$$\frac{1}{n} \sum_{i=1}^n \int_{\mathcal{Y}} g_i(y) \int_{\mathcal{I}} \mu_{i, \theta_{u(\theta)}}(y, z) \left[1 - \left(\frac{\pi_{\theta}(z)}{\pi_{\theta^*}(z)} \right)^{\lambda} \right] dz dy = T(\theta, 0) - T(\theta, \lambda).$$

We first consider T as a function of λ and study its variations:

$$\frac{d}{d\lambda} T(\theta, \lambda) = \frac{1}{n} \sum_{i=1}^n \int_{\mathcal{Y}} g_i(y) \int_{\mathcal{I}} \mu_{i, \theta_{u(\theta)}}(y, z) \left(\frac{\pi_{\theta}(z)}{\pi_{\theta^*}(z)} \right)^{\lambda} \left(\log \frac{\pi_{\theta}(z)}{\pi_{\theta^*}(z)} \right) dz dy.$$

$$\frac{d^2}{d\lambda^2} T(\theta, \lambda) = \frac{1}{n} \sum_{i=1}^n \int_{\mathcal{Y}} g_i(y) \int_{\mathcal{I}} \mu_{i, \theta_{u(\theta)}}(y, z) \left(\frac{\pi_{\theta}(z)}{\pi_{\theta^*}(z)} \right)^{\lambda} \left(\log \frac{\pi_{\theta}(z)}{\pi_{\theta^*}(z)} \right)^2 dz dy.$$

Since T is twice differentiable, we do two successive integrations by parts and obtain so:

$$T(\theta, 0) - T(\theta, \lambda) = -\lambda \frac{d}{d\lambda} T(\theta, \lambda) + \int_0^\lambda v \frac{d^2}{d\lambda^2} T(\theta, v) dv.$$

Since the second partial derivate is always non negative, we can bounded by below the previous right term as follows:

$$T(\theta, 0) - T(\theta, \lambda) \geq -\lambda \frac{d}{d\lambda} T(\theta, \lambda)$$

Moreover the first partial derivate is nondecreasing and its value in zero is non positive since we have:

$$\begin{aligned} \frac{d}{d\lambda} T(\theta, 0) &= \frac{1}{n} \sum_{i=1}^n \int_{\mathcal{Y}} g_i(y) \int_{\mathcal{I}} \mu_{i, \theta_{u(\theta)}}(y, z) \log \frac{\pi_\theta(z)}{\pi_{\theta^*}(z)} dz dy \\ &= \Lambda_n(\theta) - \Lambda_n(\theta^*), \end{aligned}$$

by definition of $u(\theta)$ in (4.17) implying (4.18). So we can introduce the following non empty set:

$$\mathcal{L}_\theta = \{\lambda \in [0, 1] \text{ such that } -\frac{d}{d\lambda} T(\theta, \lambda) \geq 0\}.$$

and define $\lambda(\theta) = \sup \mathcal{L}_\theta$. We will now show that $\lambda(\theta)$ defined like this is positive for all θ in $\Theta_{1,n}$. Consider first the case of θ^* . Since $T(\theta^*, \cdot) = 1$, we have $\frac{d}{d\lambda} T(\theta^*, \cdot) = 0$ implying $\mathcal{L}_{\theta^*} = [0, 1]$ and $\lambda(\theta^*) = 1$. Consider now θ in $\Theta_{1,n}$ different from θ^* . Since the function $-\frac{d}{d\lambda} T(\theta, \cdot)$ is continuous and decreasing and since $-\frac{d}{d\lambda} T(\theta, 0) = \Lambda_n(\theta^*) - \Lambda_n(\theta) > 0$, we can deduce that $\lambda(\theta) > 0$ implying

$$-\int_0^1 F'_{n,\theta}(u) du \geq \frac{m}{M} \frac{\lambda(\theta)}{2} \min \pi_{\theta^*} \exp \left[-\left\| \log \frac{\pi_{\theta^*}}{\pi_\theta} \right\|_\infty \right] \int_{\mathcal{I}} \left[\left(\log \frac{\pi_{\theta^*}(z)}{\pi_\theta(z)} \right)^2 \right] dz.$$

Moreover the function $\theta \rightarrow \lambda(\theta)$ defined previously is continuous on $\Theta_{1,n}$, because the function $T(\cdot, \lambda)$ is continuous on $\Theta_{1,n}$ for each λ . This implies that λ_0 defined as $\min_{\theta \in \Theta_{1,n}} \lambda(\theta)$ is positive since $\Theta_{1,n}$ is compact. So we obtain for all θ in $\Theta_{1,n}$

$$-\int_0^1 F'_{n,\theta}(u) du \geq \frac{m}{M} \frac{\lambda_0}{2} \min \pi_{\theta^*} \exp \left[-\left\| \log \frac{\pi_{\theta^*}}{\pi_\theta} \right\|_\infty \right] \int_{\mathcal{I}} \left[\left(\log \frac{\pi_{\theta^*}(z)}{\pi_\theta(z)} \right)^2 \right] dz.$$

By the way, Lemma 7 of Stone (1986) implies that there exists M_3 such that $\left\| \log \frac{\pi_{\theta^*}}{\pi_\theta} \right\|_\infty \leq M_3 \sqrt{J} \left\| \log \frac{\pi_{\theta^*}}{\pi_\theta} \right\|_2$, so we obtain

$$-\int_0^1 F'_{n,\theta}(u) du \geq \frac{m}{M} \frac{\lambda_0}{2} \min \pi_{\theta^*} \exp [-M_1 M_3] \int_{\mathcal{I}} \left[\left(\log \frac{\pi_{\theta^*}(z)}{\pi_\theta(z)} \right)^2 \right] dz,$$

since $b_n \sqrt{J} \leq M_1$, which implies the expected result. □

Proof of Lemma 4.11:

We recall that $\Psi_\theta(u) = \int_{\mathcal{I}} \pi_{\theta_u}(z) \log \frac{\pi_{\theta^*}(z)}{\pi_\theta(z)} dz$ for all u in $[0, 1]$. So we obtain after derivation for all u in $[0, 1]$

$$\Psi'_\theta(u) = \left(\int_{\mathcal{I}} \pi_{\theta_u}(z) \log \frac{\pi_{\theta^*}(z)}{\pi_\theta(z)} dz \right)^2 - \int_{\mathcal{I}} \pi_{\theta_u}(z) \left(\log \frac{\pi_{\theta^*}(z)}{\pi_\theta(z)} \right)^2 dz.$$

We introduce now the function ψ defined as follows for all u in $[0, 1]$:

$$\psi_\theta(u) = K(\pi_{\theta^*}, \pi_{\theta_u})$$

where K is the Kullback-Leiber divergence defined in (4.16). We obtain after twice derivations for all u in $[0, 1]$:

$$\begin{aligned} \psi'_\theta(u) &= K(\pi_{\theta^*}, \pi_\theta) - \int_{\mathcal{I}} \pi_{\theta_u}(z) \log \frac{\pi_{\theta^*}(z)}{\pi_\theta(z)} dz \\ \psi''_\theta(u) &= \int_{\mathcal{I}} \pi_{\theta_u}(z) \left(\log \frac{\pi_{\theta^*}(z)}{\pi_\theta(z)} \right)^2 dz - \left(\int_{\mathcal{I}} \pi_{\theta_u}(z) \log \frac{\pi_{\theta^*}(z)}{\pi_\theta(z)} dz \right)^2. \end{aligned}$$

So we remark that for all u in $[0, 1]$, we have $\psi''_\theta(u) = -\Psi'_\theta(u)$, implying by integration for all u in $[0, 1]$

$$\psi'_\theta(u) - \psi'_\theta(0) = \Psi_\theta(0) - \Psi_\theta(u)$$

Now $\psi'_\theta(0) = 0$ and $\Psi_\theta(0) = K(\pi_{\theta^*}, \pi_\theta)$, so we obtain after a second integration

$$\psi_\theta(1) - \psi_\theta(0) = K(\pi_{\theta^*}, \pi_\theta) - \int_0^1 \Psi_\theta(u) du$$

implying the result since $\psi_\theta(1) = K(\pi_{\theta^*}, \pi_\theta)$ and $\psi_\theta(0) = 0$. □

Références

- Achauer, U. and the Krisp Teleseismic Working Group (1994). New ideas on the kenya rift based on the inversion of the combined dataset of the 1985 and 1989/90 seismic tomography experiments. *Tectonophysics* 236, 305–329.
- Aki, K., Christofferson, A., and Husebye, E. (1977). Determination of the three-dimensional seismic structure of the lithosphere. *J. Geophys. Res.* 82, 277–296.
- Barron, A. R. and Sheu, C.-H. (1991). Approximation of density functions by sequences of exponential families. *Ann. Statist.* 19(3), 1347–1369.
- Benveniste, A., Métivier, M., and Priouret, P. (1990). *Adaptive algorithms and stochastic approximations*. Berlin: Springer-Verlag. Translated from the French by Stephen S. Wilson.
- Birch, F. (1961). The velocity of compressional waves in rocks to 10 kilobars. *J. Geophys. Res.* 66, 2199–2224.
- Brandière, O. and Duflo, M. (1995). Les algorithmes stochastiques contournent-ils les pièges? *C. R. Acad. Sci. Paris Sér. I Math.* 321(3), 335–338.
- Cady, J. W. (1980). Calculation of gravity and magnetic anomalies of finite length right polygonal prisms. *Geophysics* 45, 1507–1512.
- Cappé, O., Douc, R., Moulines, E., and Robert, C. (2002). On the convergence of the Monte Carlo maximum likelihood method for latent variable models. *Scand. J. Statist.* 29(4), 615–635.
- Cavalier, L., Golubev, G. K., Picard, D., and Tsybakov, A. B. (2002). Oracle inequalities for inverse problems. *Ann. Statist.* 30(3), 843–874. Dedicated to the memory of Lucien Le Cam.
- Cavalier, L. and Tsybakov, A. (2002). Sharp adaptation for inverse problems with random noise. *Probab. Theory Related Fields* 123(3), 323–354.
- Celeux, G. and Diebolt, J. (1986). L’algorithme SEM: Un algorithme d’apprentissage probabiliste pour la reconnaissance de mélange de densités. (The SEM algorithm: An algorithm of probabilistic learning for the determination of mixtures of densities). *Rev. Stat. Appl.* 34(2), 35–52.
- Celeux, G. and Diebolt, J. (1990). Une version de type recuit simulé de l’algorithme EM. *C. R. Acad. Sci. Paris Sér. I Math.* 310(3), 119–124.
- Celeux, G. and Diebolt, J. (1992). A stochastic approximation type EM algorithm for the mixture problem. *Stochastics Stochastics Rep.* 41(1-2), 119–134.
- Chen, H. F., Lei, G., and Gao, A. J. (1988). Convergence and robustness of the Robbins-Monro algorithm truncated at randomly varying bounds. *Stochastic Process. Appl.* 27(2), 217–231.

- Concordet, D. and Nunez, O. G. (2002). A simulated pseudo-maximum likelihood estimator for nonlinear mixed models. *Comput. Statist. Data Anal.* 39(2), 187–201.
- de Boor, C. (1978). *A practical guide to splines*, Volume 27 of *Applied Mathematical Sciences*. New York: Springer-Verlag.
- Delyon, B. (1996). General results on the convergence of stochastic algorithms. *IEEE Trans. Automat. Control* 41(9), 1245–1255.
- Delyon, B., Lavielle, M., and Moulines, E. (1999). Convergence of a stochastic approximation version of the EM algorithm. *Ann. Statist.* 27(1), 94–128.
- Dempster, A. P., Laird, N. M., and Rubin, D. B. (1977). Maximum likelihood from incomplete data via the EM algorithm. *J. Roy. Statist. Soc. Ser. B* 39(1), 1–38. With discussion.
- Diebolt, J. and Celeux, G. (1993). Asymptotic properties of a stochastic EM algorithm for estimating mixing proportions. *Comm. Statist. Stochastic Models* 9(4), 599–613.
- Ding, A. and Wu, H. (1999). Relationships between antiviral treatment effects and biphasic viral decay rates in modeling hiv dynamics. *Mathematical Biosciences* 160(1), 63–82.
- Ding, A. and Wu, H. (2001). Assessing antiviral potency of anti-hiv therapies in vivo by comparing viral decay rates in viral dynamics models. *Biostatistics* 2, 1–18.
- Donoho, D. L. (1995). Nonlinear solution of linear inverse problems by wavelet-vaguelette decomposition. *Appl. Comput. Harmon. Anal.* 2(2), 101–126.
- Duflo, M. (1996). *Algorithmes stochastiques*, Volume 23 of *Mathématiques & Applications*. Berlin: Springer-Verlag.
- Eggermont, P. (1999). Nonlinear smoothing and the EM algorithm for positive integral equations of the first kind. *Appl. Math. Optimization* 39(1), 75–91.
- Eggermont, P. and LaRiccia, V. (1995). Maximum smoothed likelihood density estimation for inverse problems. *Ann. Stat.* 23(1), 199–220.
- Evans, J. and Achauer, U. (1993). Teleseismic velocity tomography using the ach method: theory and application to continental-scale studies. In H. Iyer and K. Hirahara (Eds.), *Seismic tomography, theory and practice*, pp. 319–360. London: Chapman and Hall.
- Fort, G. and Moulines, E. (2003). Convergence of the Monte Carlo Expectation Maximization for curved exponential families. *Ann. Stat.* 31(4), 1220–1259.
- Glaznev, V. N., Raevsky, A. B., and Skopenko, G. B. (1996). A three-dimensional integrated density and thermal model of the fennoscandian lithosphere. *Tectonophysics* 258, 15–33.
- Granet, M. and Cara, M. (1988). 3-d velocity structure beneath france in different frequency bands. *Phys. Earth Planet. Int.* 51, 133–152. SEARCH.
- Grizzle, J. and Allen, D. (1979). Analysis of growth and dose-response curves. *Biometrics* 25, 579–628.
- Gu, M. G. and Kong, F. H. (1998). A stochastic approximation algorithm with Markov chain Monte-Carlo method for incomplete data estimation problems. *Proc. Natl. Acad. Sci. USA* 95(13), 7270–7274 (electronic).
- Gu, M. G. and Zhu, H.-T. (2001). Maximum likelihood estimation for spatial models by Markov chain Monte Carlo stochastic approximation. *J. R. Stat. Soc. Ser. B Stat. Methodol.* 63(2), 339–355.

- Hashimoto, Y. and Sheiner, L. (1991). Designs for population pharmacodynamics: Value of pharmacokinetic data and population analysis. *J. Pharmacokin. Biopharm.* 19, 333–353.
- Hastings, W. (1970). Monte Carlo sampling methods using Markov chains and their applications. *Biometrika* 57, 97–109.
- Hoeffding, W. (1963). Probability inequalities for sums of bounded random variables. *J. Amer. Statist. Assoc.* 58, 13–30.
- Humphreys, E. D. and Dueker, K. G. (1994). Physical state of the western u.s. upper mantle. *J. Geophys. Res.* 99, 9635–9650.
- Kissling, E., Husen, S., and Haslinger, F. (2001). Model parametrization in seismic tomography: a choice of consequence for the solution quality. *Phys. Earth Planet. Int.* 123, 89–101.
- Koo, J.-Y. (1999). Logspline deconvolution in Besov space. *Scand. J. Stat.* 26(1), 73–86.
- Koo, J.-Y. and Chung, H.-Y. (1998). Log-density estimation in linear inverse problems. *Ann. Stat.* 26(1), 335–362.
- Koo, J.-Y. and Park, B. U. (1996). B-spline deconvolution based on the EM algorithm. *J. Stat. Comput. Simulation* 54(4), 275–288.
- Laird, N. M. and Ware, J. H. (1982). Random-effects models for longitudinal data. *Biometrics* 38, 963–974.
- Lange, K. (1995). A gradient algorithm locally equivalent to the EM algorithm. *J. Roy. Statist. Soc. Ser. B* 57(2), 425–437.
- Lavielle, M. and Lebarbier, E. (2001). An application of MCMC methods to the multiple change-points problem. *Signal Processing* 81, 39–53.
- Lavielle, M. and Moulines, E. (1997). A simulated annealing version of the EM algorithm for non-Gaussian deconvolution. *Statistics and Computing* 7(4), 229–236.
- Lees, J. and VanDecar, J. (1991). Seismic tomography constrained by bouguer gravity anomalies: Applications in western washington. *Pageoph.* 135, 31–52.
- Lindstrom, M. J. and Bates, D. M. (1988). Newton-Raphson and EM algorithms for linear mixed-effects models for repeated-measures data. *J. Amer. Statist. Assoc.* 83(404), 1014–1022.
- Lindstrom, M. J. and Bates, D. M. (1990). Nonlinear mixed effects models for repeated measures data. *Biometrics* 46(3), 673–687.
- Lines, L., Schultz, A., and Treitel, S. (1988). Cooperative inversion of geophysical data. *Geophysics* 53, 8–20.
- Louis, T. A. (1982). Finding the observed information matrix when using the EM algorithm. *J. Roy. Statist. Soc. Ser. B* 44(2), 226–233.
- Ludwig, J. W., Nafe, J. E., and Drake, C. L. (1970). Seismic refraction. In A. Maxwell (Ed.), *The Sea*, Volume 4, pp. 53–84. New York: Wiley.
- Meng, X.-L. (1994). On the rate of convergence of the ECM algorithm. *Ann. Stat.* 22(1), 326–339.
- Meng, X.-L. and Rubin, D. B. (1992). Recent extensions to the EM algorithm. In *Bayesian statistics, 4* (Peñíscola, 1991), pp. 307–320. New York: Oxford Univ. Press.

- Meng, X.-L. and Rubin, D. B. (1993). Maximum likelihood estimation via the ECM algorithm: a general framework. *Biometrika* 80(2), 267–278.
- Mengersen, K. L. and Tweedie, R. L. (1996). Rates of convergence of the Hastings and Metropolis algorithms. *Ann. Statist.* 24(1), 101–121.
- Mentré, F. and Gomeni, R. (1995). A two-step iterative algorithm for estimation in nonlinear mixed-effect models with an evaluation in population pharmacokinetics. *J. Biopharm. Stat.* 5(2), 141–158.
- Meyn, S. P. and Tweedie, R. L. (1993). *Markov chains and stochastic stability*. Communications and Control Engineering Series. London: Springer-Verlag London Ltd.
- Pinheiro, J. and Bates, D. (1995). Approximations to the log-likelihood function in the nonlinear mixed-effect models. *J. Comput. Graph. Stat.* (4), 12–35.
- Pinheiro, J. C. and Bates, D. M. (2000). *Mixed-Effects Models in S and S-PLUS*. New York: Springer.
- Racine-Poon, A. (1985). A Bayesian approach to nonlinear random effects models. *Biometrics* 41, 1015–1023.
- Ritter, J. R., Jordan, M., Christensen, U. R., and Achauer, U. (2001). A mantle plume below the eifel volcanic fields, germany. *Earth Planet. Sci. Lett.* 186, 7–14.
- Robert, C. (1996). *Méthodes de Monte Carlo par chaînes de Markov*. Statistique Mathématique et Probabilité. [Mathematical Statistics and Probability]. Paris: Éditions Économica.
- Schumitzky, A. (1991). Nonparametric EM algorithms for estimating prior distributions. *Appl. Math. Comput.* 45(2, part II), 143–157.
- Sheiner, L., Hashimoto, Y., and Beal, S. (1991). A simulation study comparing designs for dose ranging. *Statistics in Medicine* 10, 303–321.
- Silverman, B. W., Jones, M. C., Nychka, D. W., and Wilson, J. D. (1990). A smoothed EM approach to indirect estimation problems, with particular reference to stereology and emission tomography. *J. Roy. Statist. Soc. Ser. B* 52(2), 271–324. With discussion and a reply by the authors.
- Sobolev, S., Zeyen, H., Granet, M., Stoll, G., Achauer, U., Bauer, C., Werling, F., Altherr, R., and Fuchs, K. (1997). Upper mantle temperatures and lithosphere-asthenosphere system beneath the french massif central constrained by seismic, gravity, petrologic and thermal observations. *Tectonophysics* 275, 143–164.
- Stone, C. J. (1986). The dimensionality reduction principle for generalized additive models. *Ann. Statist.* 14(2), 590–606.
- Stone, C. J. (1990). Large-sample inference for log-spline models. *Ann. Statist.* 18(2), 717–741.
- Vardi, Y., Shepp, L. A., and Kaufman, L. (1985). A statistical model for positron emission tomography. *J. Amer. Statist. Assoc.* 80(389), 8–37. With discussion.
- Vonesh, E. F. (1996). A note on the use of Laplace’s approximation for nonlinear mixed-effects models. *Biometrika* 83(2), 447–452.
- Wakefield, J. (1996). The Bayesian analysis of population pharmacokinetic models. *J. Am. Stat. Assoc.* 91(433), 62–75.

- Wakefield, J., Smith, A., Racine-Poon, A., and Gelfand, A. (1994). Bayesian analysis of linear and nonlinear population models by using the Gibbs sampler. *J. R. Stat. Soc., Ser. C* 43(1), 201–221.
- Walker, S. (1996). An EM algorithm for nonlinear random effects models. *Biometrics* 52(3), 934–944.
- Wei, G. C. G. and Tanner, M. A. (1990). A Monte Carlo implementation of the EM algorithm and the Poor's Man's data augmentation algorithms. *J. Amer. Statist. Assoc.* 85(411), 699–704.
- Wu, C.-F. J. (1983). On the convergence properties of the EM algorithm. *Ann. Statist.* 11(1), 95–103.
- Yao, J.-F. (2000). On recursive estimation in incomplete data models. *Statistics* 34(1), 27–51.
- Zeyen, H. and Achauer, U. (1997). Joint inversion of teleseismic delay times and gravity anomaly data for regional structures: Theory and synthetic examples. In K. Fuchs (Ed.), *Upper mantle heterogeneities from active and passive seismology*, Volume 17 of *NATO-ASI series 1*, pp. 155–168. Dordrecht: Kluwer Acad. Publ.
- Zeyen, H. and Pous, J. (1993). 3-d joint inversion of magnetic and gravimetric data with a priori information. *Geophys. J. Int.* 112, 244–256.